

FREE SAMPLE EDITION

This preview contains the first 1 chapter of the complete book.

BOOK DETAILS

- **Title:** Artificial Intelligence Foundations - Concepts, Approaches and Real Use Cases
- **Edition:** 1st Edition
- **Publisher:** Alkademy Books
- **Chapters in full book:** 10

CHAPTERS IN THIS PREVIEW

- Chapter 1: What AI Really Is: Cutting Through the Hype

Get the full book at alkademy.com

Artificial Intelligence Foundations - Concepts, Approaches and Real Use Cases

Alkademy Content Team

1st Edition

Alkademy Books

June 2026

PREFACE

Over the years, I have sat in countless meetings where the word "artificial intelligence" was used to mean almost everything — and, in practice, nothing. I have watched talented product managers propose AI features without a clear sense of what they were actually asking for. I have seen engineers reach for deep learning when a simple rule would have done the job in a fraction of the time. I have listened to executives describe AI strategies that were, at their core, spreadsheet automation dressed in fashionable language. None of these people were unintelligent or careless. They were simply working without a shared foundation — a clear, honest mental model of what AI actually is, how its different approaches work, and when each one genuinely makes sense. That gap, repeated across organizations and industries, is what motivated me to write this book.

This book is for anyone who wants to understand AI clearly and use that understanding in practice. You do not need a mathematics degree or a background in computer science to read it. If you are a product manager trying to evaluate whether a feature should use machine learning or a simpler rule-based system, this book is for you. If you are a developer who has used AI tools but never quite understood what is happening underneath them, this book is for you. If you are a business analyst, a consultant, a team lead, or a student stepping into a world where AI is increasingly unavoidable, this book is for you. The only assumption I make is that you are curious and willing to think carefully. I will handle the rest.

My philosophy in writing this book was simple: explain AI the way it actually works, not the way it is marketed. That means being honest about what AI can and cannot do. It means grounding

every concept in real systems and real use cases rather than abstract theory. It means resisting the temptation to oversimplify on one side — reducing AI to magic — or to overcomplicate on the other, drowning you in mathematical notation before you have any reason to care about it. I wanted to write the book I wished had existed when I was first trying to make sense of this field: one that respects your intelligence, builds your understanding progressively, and always connects ideas back to practical decisions you might actually face.

The book is organized to take you from the broadest picture down to the specific. We begin by establishing what AI really means — its history, its scope, and the common misconceptions that create so much confusion in conversations today. From there, we move into the core approaches that make up the AI landscape. You will learn how rule-based systems work and why they remain relevant despite being considered "old-fashioned" by some. You will develop a genuine understanding of machine learning: what it means for a system to learn from data, how models are trained and evaluated, and where this approach shines and where it struggles. We then go deeper into neural networks and deep learning, demystifying the architectures behind image recognition, language processing, and so much more. The final major section addresses generative AI — large language models, image generation, and the new class of systems that have captured the world's attention — placing them in proper context so you understand both their remarkable capabilities and their real limitations. Throughout, each concept is anchored to concrete use cases drawn from industries including healthcare, finance, retail, and technology.

By the time you finish this book, you will be able to do several things you may not be able to do right now. You will be able to look at an AI application — whether it is a recommendation engine, a fraud detection system, a chatbot, or an image classifier — and identify what kind of approach is likely powering it and why. You will be able to participate in conversations about AI adoption with clarity and confidence, asking the right questions rather than nodding along. Most importantly, you will be able to make better decisions: choosing the right approach for a given problem, setting realistic expectations, and recognizing when AI is genuinely the right tool and when it is not.

I should also be honest about what this book does not cover, and why. This is not a textbook for machine learning engineers who want to implement algorithms from scratch. You will not find detailed mathematical derivations, training code, or deep architectural comparisons between specific model variants. Those topics are important, but they belong in a different book for a different reader. My goal here is foundations — the conceptual clarity and practical judgment

that make everything else easier to learn when you are ready to go further. Think of this book as building the mental scaffolding. What you construct on top of it is up to you.

I am genuinely glad you picked this up. The noise around AI right now is extraordinary, and it can make the subject feel either impossibly complex or suspiciously simple depending on which direction the hype is blowing that week. My hope is that by the last page, you will feel neither overwhelmed nor underwhelmed — just clear. Clear about what AI is, how it works, and how to think about it honestly. That clarity, I believe, is more valuable than any single technical skill, because it shapes every decision that follows. Let's begin.

To all the lifelong learners, who never stop exploring, questioning, and growing.

To those who dare to teach themselves, and then teach others.

This book is for you.

Contents

Preface	i
1 What AI Really Is: Cutting Through the Hype	1
1.1 What AI Really Is: Cutting Through the Hype	1

CHAPTER 1

WHAT AI REALLY IS: CUTTING THROUGH THE HYPE

1.1 What AI Really Is: Cutting Through the Hype

Imagine picking up a newspaper on any given morning and counting how many times the phrase “artificial intelligence” appears. You might find it in an article about a hospital using software to detect tumors, in an advertisement for a smartphone camera that “intelligently” adjusts lighting, in a think piece warning that robots will soon take every job, and in a financial report crediting AI for record profits at a technology company. The word is everywhere, yet almost every usage carries a slightly different meaning. For some writers, AI is a sophisticated algorithm quietly optimizing a supply chain. For others, it is the harbinger of a post-human civilization. This gap between casual usage and technical reality is not merely a matter of semantics — it shapes business decisions, government policy, career choices, and public trust. Before we can use AI wisely, we must understand what it actually is.

1.1.1 Defining Artificial Intelligence

The term “artificial intelligence” was coined in 1956 by computer scientist John McCarthy, who used it to describe the scientific goal of making machines behave in ways that would be considered intelligent if observed in humans. That original definition was deliberately broad, and it has remained so. Today, AI refers to a wide field of computer science and engineering focused on building systems capable of performing tasks that, when done by humans, require some form of cognitive ability — things like recognizing speech, understanding language, making decisions, solving puzzles, and learning from experience.

This breadth is both the field’s greatest strength and the primary source of public confusion. AI is not a single product, a single algorithm, or a single company’s technology. It is an umbrella term covering dozens of subfields, methodologies, and tools. A spam filter that classifies your email is an AI system. So is the navigation software routing a delivery truck around a traffic jam. So, in principle, would be a hypothetical machine capable of passing a doctoral examination in any subject. These systems share a label but differ enormously in scope, capability, and underlying mechanism. Treating them as equivalent is like calling a bicycle and a commercial jet the same because both are “transportation.”

Example: Consider how Google Translate works. When you type a sentence in French and receive an English translation, you are interacting with an AI system. That system has been trained on billions of translated documents and uses statistical patterns to predict which sequence of English words most closely captures the meaning of the French input. It performs this task impressively well — often well enough that the result requires no editing. But the system has no understanding of French culture, no awareness that it is translating, and no ability to explain why it chose one word over another. It is performing a task that requires human-like intelligence, but it is doing so through pattern matching at a massive scale, not through comprehension. This distinction will matter enormously as we explore what AI can and cannot do.

A Working Definition for Practical Purposes

For the purposes of this book, we will use the following working definition: *Artificial intelligence is the design and deployment of computational systems that perform tasks typically associated with human cognitive abilities, including perception, reasoning, learning, and decision-making.* This definition is intentionally functional rather than philosophical. It focuses on what AI systems

do rather than on whether they are “truly” intelligent in any deep sense — a question that philosophers and scientists have debated for decades without resolution.

This functional framing is useful because it keeps attention on outcomes and capabilities rather than on metaphysical debates. When a hiring manager asks whether a particular AI tool can screen resumes effectively, the relevant question is not whether the tool “understands” a resume the way a human recruiter does. The relevant question is whether it produces useful, accurate, and fair results. Grounding AI in function rather than in imitation of human consciousness also helps avoid one of the field’s most persistent traps: anthropomorphism, the tendency to project human qualities onto systems that do not possess them.

Key Insight

The word “intelligence” in “artificial intelligence” does not mean the same thing as human intelligence. It refers to the capacity to perform specific cognitive tasks effectively. A chess-playing program is intelligent in the narrow sense that it makes strategic decisions, but it cannot tie its shoes, feel pride after winning, or understand why chess matters to people. Keeping this distinction in mind will protect you from both overestimating and underestimating AI systems.

1.1.2 The Landscape of AI: A Field, Not a Technology

One of the most important conceptual shifts a newcomer to AI must make is recognizing that AI is a *field of study and practice*, not a single technology. Just as “medicine” encompasses surgery, pharmacology, psychiatry, and dozens of other specialties, AI encompasses machine learning, natural language processing, computer vision, robotics, expert systems, and much more. Each of these subfields has its own history, methods, open problems, and communities of researchers.

Understanding this landscape prevents a common error in business and policy discussions: assuming that a breakthrough in one area of AI automatically translates to progress in another. When a deep learning model achieves superhuman accuracy in diagnosing diabetic retinopathy from eye scans, that is a genuine and significant achievement. But it tells us very little about whether AI systems can soon hold a nuanced conversation about a patient’s emotional response to a diagnosis. Vision and language are different problems, addressed by different techniques, and progress in one does not imply equivalent progress in the other.

The Nested Relationship: AI, Machine Learning, and Deep Learning

Three terms appear together so frequently that many people treat them as synonyms: artificial intelligence, machine learning, and deep learning. They are not synonyms — they are nested concepts, each one a subset of the one above it.

Artificial intelligence is the broadest category. It includes any technique that enables a machine to mimic cognitive behavior. This includes rule-based systems, search algorithms, probabilistic reasoning, and learning from data. **Machine learning** (ML) is a subset of AI that focuses specifically on systems that improve their performance on a task through exposure to data, without being explicitly programmed with rules for every situation. Rather than a programmer writing “if the email contains the word *lottery*, mark it as spam,” a machine learning system is shown thousands of examples of spam and non-spam emails and learns to distinguish them by identifying statistical patterns. **Deep learning** is a subset of machine learning that uses artificial neural networks with many layers — hence “deep” — to learn representations of data. Deep learning has driven most of the dramatic AI advances of the past decade, from image recognition to language generation, because it can automatically discover useful features in raw data without requiring humans to specify what to look for.

Table 1.1: Comparing AI, Machine Learning, and Deep Learning

Concept	Scope	Real-World Example
Artificial Intelligence	Broadest field; any technique enabling machine cognition	IBM Deep Blue playing chess using search trees and evaluation rules
Machine Learning	Subset of AI; systems that learn from data	Email spam filters trained on labeled message datasets
Deep Learning	Subset of ML; multi-layer neural networks	OpenAI’s GPT models generating human-like text

Example: When Apple’s Face ID unlocks your phone by recognizing your face, it is using deep learning — specifically a type of neural network trained on facial geometry data. This is also machine learning, because the system improved through training on examples. And it is also AI, because face recognition is a cognitive task. All three labels apply simultaneously. The appropriate label depends on the level of specificity your conversation requires.

Other Important Subfields

Beyond machine learning and deep learning, the AI landscape includes several other subfields worth knowing. **Natural language processing** (NLP) deals with the ability of machines to understand, interpret, and generate human language. It powers chatbots, translation services, sentiment analysis tools, and voice assistants. **Computer vision** focuses on enabling machines to interpret and make decisions based on visual input — images and video. It underlies autonomous vehicle navigation, medical imaging analysis, and quality control in manufacturing. **Robotics** combines AI with physical systems, enabling machines to perceive their environment and take physical actions in response. **Expert systems** encode human domain knowledge in the form of rules and use logical inference to make decisions — an older approach that remains useful in fields like medical diagnosis and legal reasoning.

Each of these subfields has produced practical tools that are already embedded in everyday life. Recognizing their distinct natures helps you ask better questions when evaluating AI products and proposals: Which subfield does this tool belong to? What kind of data does it need? What are the known limitations of this approach?

1.1.3 A Brief History: How We Got Here

AI did not emerge fully formed from a Silicon Valley laboratory. It has a rich and sometimes turbulent history spanning more than seven decades, marked by periods of extraordinary optimism, painful disappointment, and eventual resurgence.

The field's formal origins lie in the mid-twentieth century, when mathematicians and early computer scientists began asking whether machines could simulate reasoning. Alan Turing's 1950 paper "Computing Machinery and Intelligence" posed the now-famous question "Can machines think?" and proposed what he called the Imitation Game — later known as the Turing Test — as a practical criterion for machine intelligence. If a machine could converse with a human judge well enough that the judge could not reliably tell it apart from another human, Turing argued, we would have grounds to call it intelligent. The Turing Test has been criticized and refined many times since, but it introduced a conceptual framework that shaped AI research for generations.

The 1956 Dartmouth Conference, organized by McCarthy along with Marvin Minsky, Claude Shannon, and Nathaniel Rochester, is widely regarded as the official birth of AI as a field. The participants were optimistic: they believed that with sufficient effort, machines could be made to

simulate every aspect of human intelligence within a generation. That prediction proved wildly overoptimistic. Early AI systems performed impressively on well-defined, constrained problems — playing checkers, solving algebra word problems, proving logical theorems — but struggled catastrophically when confronted with the messiness and ambiguity of the real world.

Case Study: The history of AI includes two major periods of reduced funding and interest, now called “AI winters.” The first occurred in the mid-1970s when government funding agencies, particularly in the United States and United Kingdom, grew frustrated with the gap between researchers’ promises and actual results. The second AI winter followed the collapse of the commercial expert systems market in the late 1980s. Companies had invested heavily in rule-based AI systems that required enormous effort to maintain and could not adapt when their domain knowledge changed. These winters were not failures of the underlying ideas so much as failures of expectation management — a lesson the current AI boom would do well to remember.

The modern renaissance of AI began in earnest in the early 2010s, driven by three converging forces: the availability of massive datasets (thanks to the internet), dramatic increases in computational power (particularly graphics processing units originally designed for video games), and algorithmic innovations in deep learning. In 2012, a deep neural network called AlexNet won the ImageNet Large Scale Visual Recognition Challenge by a margin so large it shocked the computer vision community and announced that a new era had begun. Since then, progress has been rapid and visible: voice assistants, language models, autonomous vehicles, protein structure prediction, and generative art have all emerged as public-facing demonstrations of what modern AI can achieve.

1.1.4 What AI Is Not: Dismantling Common Myths

If understanding what AI is requires careful definition, understanding what AI is *not* requires equal care. Popular culture, marketing language, and even some journalism have created a set of persistent myths about AI that distort public understanding and lead to poor decisions. Addressing these myths directly is not an exercise in pessimism — it is a prerequisite for realistic and productive engagement with the technology.

Myth One: AI Is Intelligent in the Way Humans Are

The most pervasive myth is that AI systems possess something like human general intelligence — that they understand the world, have intentions, and could, if sufficiently advanced, think and feel as we do. This conflation is understandable. When you have a conversation with a large language model and it responds with apparent empathy and nuanced reasoning, the experience feels uncannily human. But the appearance of understanding is not the same as understanding itself.

Current AI systems, including the most sophisticated language models available today, are what researchers call **narrow AI** or **weak AI** systems. They are optimized to perform specific tasks or classes of tasks, and they do so by finding patterns in data. A language model does not know what the words it generates mean in any experiential sense. It does not know that Paris is a city you can visit, that grief feels a particular way, or that a promise carries moral weight. It has learned statistical relationships between tokens of text and generates outputs that are statistically consistent with its training data. The outputs can be remarkably useful and often indistinguishable from human-written text, but the process producing them is fundamentally different from human cognition.

General AI (also called artificial general intelligence, or AGI) — a system capable of performing any intellectual task that a human can, with comparable flexibility and depth — does not yet exist. It remains a research goal and a topic of vigorous debate about whether it is achievable, how far away it might be, and what its arrival would mean. Treating current AI tools as early versions of AGI is a category error that leads to both inflated expectations and misplaced fears.

Myth Two: AI Is Objective and Unbiased

A second common myth holds that AI systems, because they are mathematical and data-driven, are free from the biases that affect human judgment. This belief is not only incorrect — it is dangerous. AI systems can perpetuate and even amplify human biases, because they learn from data generated by humans in a world shaped by historical inequities.

Example: In 2018, Amazon quietly shut down an AI recruiting tool it had been developing internally. The tool was designed to screen resumes and identify top candidates automatically. The problem was that the system had been trained on resumes submitted to Amazon over a ten-year period — a dataset dominated by male applicants, reflecting the gender imbalance in the

technology industry. The system learned to penalize resumes that included words like “women’s” (as in “women’s chess club”) and to downgrade graduates of all-women’s colleges. The bias was not introduced deliberately; it emerged from the patterns in the training data. Amazon’s engineers tried to correct the problem but ultimately concluded that they could not guarantee the tool would not find other proxies for gender, and the project was abandoned.

This example illustrates a principle that runs throughout AI development: a system can only be as fair as the data it learns from and the choices made in designing, training, and deploying it. Recognizing AI’s potential for bias is not an argument against using it — it is an argument for using it thoughtfully, with appropriate oversight and ongoing evaluation.

Myth Three: AI Will Soon Replace All Human Workers

Headlines about AI and employment tend toward extremes. Either AI will create a golden age of leisure as machines handle all drudgery, or it will eliminate jobs on a catastrophic scale, leaving most humans economically redundant. The reality is considerably more nuanced and historically familiar.

Technological change has always disrupted labor markets. The mechanization of agriculture, the invention of the printing press, the rise of factory production, and the computerization of office work all eliminated certain jobs while creating others. AI is following a similar pattern. Tasks that are routine, well-defined, and data-rich are most susceptible to automation. Tasks that require physical dexterity in unpredictable environments, deep contextual judgment, interpersonal trust, creativity rooted in lived experience, and ethical reasoning are far more resistant. The practical question for any worker or organization is not “will AI replace humans?” but “which specific tasks within which roles are candidates for automation, and what new capabilities will be needed?”

Common Mistake

Avoid the trap of treating AI as either a universal threat or a universal solution. The impact of AI on any given role depends on the specific tasks involved, the quality of available data, the regulatory environment, and organizational choices about how to deploy the technology. Blanket statements about AI “replacing” entire professions almost always oversimplify a complex, context-dependent reality.

Myth Four: More Data Always Means Better AI

A widely repeated belief in technology circles is that data is the new oil — the more you have, the more powerful your AI systems will be. While data is genuinely important to most modern AI approaches, the relationship between data quantity and system quality is far from linear. Data quality, relevance, and diversity matter at least as much as volume. A system trained on enormous amounts of low-quality, biased, or irrelevant data will perform poorly regardless of the dataset's size. Moreover, some AI techniques — particularly those involving careful reasoning, structured knowledge, or small-sample learning — do not require massive datasets at all.

The “more data is always better” myth also carries practical risks. It encourages organizations to collect and retain data indiscriminately, raising privacy concerns and creating legal liabilities. It can lead to the mistaken belief that an AI project will succeed if only enough data is gathered, when the real bottleneck may be problem formulation, model selection, or deployment infrastructure. Sophisticated AI practitioners spend as much time thinking about which data to use as they do about how much of it to gather.

1.1.5 Why Getting This Right Matters

The stakes of understanding AI accurately are not merely academic. Misconceptions about AI have real consequences at every level of society. Organizations that overestimate AI capabilities make costly investments in tools that cannot deliver what was promised. Those that underestimate them miss opportunities to improve efficiency, safety, and service quality. Policymakers who conflate narrow AI with AGI may craft regulations that are either too restrictive for current tools or too permissive for future ones. Individuals who believe AI is infallible may defer to its outputs in situations that demand human judgment.

Scenario: A regional hospital system decides to implement an AI tool that predicts which patients are at high risk of being readmitted within thirty days of discharge. The tool performs well in the research setting where it was developed. But when deployed at this particular hospital, it consistently underestimates risk for elderly patients from rural areas — a demographic underrepresented in the training data. If the clinical staff treats the tool's outputs as authoritative rather than as one input among many, some high-risk patients will be discharged without adequate follow-up. The harm is not caused by the AI being malicious or even technically broken; it is caused by a mismatch between the tool's actual capabilities and the assumptions made about

those capabilities during deployment.

Clear thinking about AI also enables better communication across disciplines. When a data scientist, a business strategist, a legal counsel, and a clinical ethicist sit down to evaluate an AI proposal, they bring different vocabularies and different concerns. A shared, accurate understanding of what AI is and is not — grounded in the kind of definitions and distinctions introduced in this chapter — gives that conversation a common foundation. It makes it possible to ask the right questions, surface the right risks, and make decisions that are both technically sound and organizationally responsible.

1.1.6 Building Your Mental Model of AI

A mental model is a simplified internal representation of how something works, accurate enough to guide useful predictions and decisions. You do not need to understand the physics of combustion to drive a car safely, but you do need a mental model accurate enough to know that pressing the accelerator increases speed and that running out of fuel stops the engine. Similarly, you do not need to understand the mathematics of neural networks to make good decisions about AI, but you do need a mental model that captures the key distinctions explored in this chapter.

A useful mental model of AI for a non-specialist might look like this: AI is a broad field of tools and techniques for building systems that perform cognitive tasks. Most current AI tools are narrow — they do one thing well and nothing else. They learn from data, which means they can be biased and can fail when new situations differ from their training. They are powerful amplifiers of human capability in the right contexts but are not substitutes for human judgment in complex, high-stakes, or novel situations. Progress in AI is real but uneven, and claims about AI capabilities should be evaluated critically, with attention to what specific task is being performed, on what data, and under what conditions.

This mental model is not the final word. As you progress through this book, you will refine and extend it with more technical detail. But even at this level of abstraction, it equips you to read an AI news story critically, evaluate a vendor's claims, participate meaningfully in a strategic conversation about AI adoption, and recognize when a colleague's understanding of AI is based on myth rather than reality. That is a significant practical advantage, and it is the foundation on which everything else in this book is built.

Best Practice

When evaluating any AI claim — whether in a news article, a vendor pitch, or an internal proposal — ask four questions: (1) What specific task is this AI system performing? (2) What data was it trained on, and how representative is that data? (3) What are the known failure modes or limitations? (4) Who is accountable when the system makes a mistake? These four questions will cut through most hype and surface the information you actually need.

Key Takeaways

- AI is a broad field focused on building systems that perform tasks requiring human-like intelligence — it is not a single technology, a magic solution, or a monolithic force. It encompasses dozens of subfields, methods, and tools, each with distinct capabilities and limitations.
- Popular media portrayals of AI often conflate narrow AI tools with general intelligence, leading to widespread misconceptions about what current systems can actually do. No existing AI system possesses general intelligence comparable to a human being.
- Machine learning and deep learning are subsets of AI, not synonyms for it. Understanding this nested relationship helps you interpret claims about AI capabilities with greater precision and ask more targeted questions.
- AI systems can reflect and amplify human biases because they learn from human-generated data. Objectivity is not an automatic property of algorithmic systems — it must be actively designed for, tested, and monitored.
- Understanding what AI cannot do is just as important as understanding what it can do. Recognizing limitations around generalization, bias, interpretability, and contextual judgment prevents costly overreliance.
- A clear mental model of AI enables better communication, more realistic expectations, and sounder decision-making — whether you are a practitioner, a manager, a policymaker, or an informed citizen navigating a world increasingly shaped by these technologies.

Exercises

1. **Evaluating AI Claims in the Wild.** Collect five AI-related claims from news articles, product advertisements, or social media posts. For each claim, write a short paragraph that applies the definitions introduced in this chapter: Is the claim referring to narrow AI or implying general intelligence? Is it describing machine learning, deep learning, or a rule-based system? Is the claim specific about what task is being performed and on what data? Does the claim acknowledge any limitations? This exercise will sharpen your ability to read AI coverage critically.
2. **Drawing the Nested Diagram.** On paper or in a digital drawing tool, create three concentric circles. Label the outermost circle “Artificial Intelligence,” the middle circle “Machine Learning,” and the innermost circle “Deep Learning.” In each ring, write one real-world example of a system that belongs to that layer but not to the inner ones (for AI but not ML, for ML but not deep learning, and for deep learning specifically). Annotate each example with two or three sentences explaining why it belongs in that layer. This visual exercise reinforces the hierarchical relationship between these concepts.
3. **The Misconception Interview.** Choose a colleague, friend, or family member who is not a technology specialist and ask them three open-ended questions: What do you think AI is? What do you think AI can do that humans cannot? What worries you most about AI? Take notes on their responses, then review this chapter and identify which of the myths discussed — AI as human-level intelligence, AI as objective, AI as universal job-replacer — appear in their answers. Write a brief reflection on how you might explain the accurate picture to them using the language and examples from this chapter.

You've reached the end of the pre-view.

The complete edition contains all chapters, including:

- Full chapter content with detailed examples and exercises
- Glossary of key terms
- Appendices and reference material
- A comprehensive index

Get the full book at alkademy.com