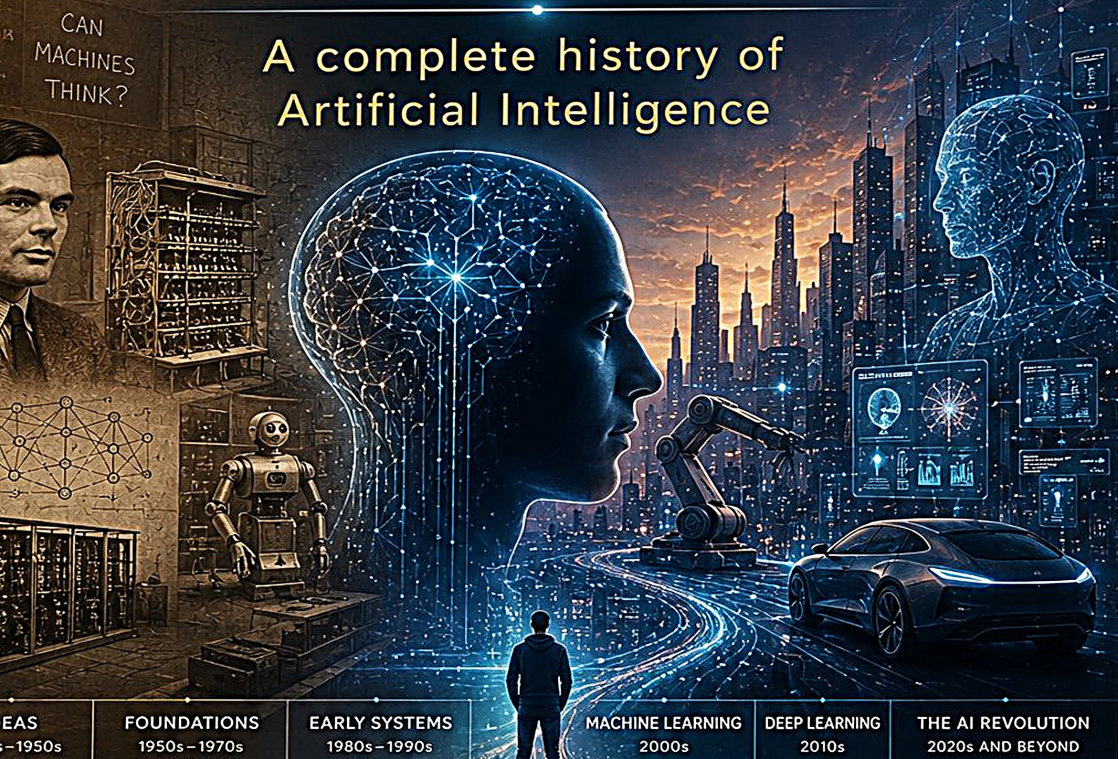


HOW AI CAME TO BE

A complete history of
Artificial Intelligence



PRE-1950s

FOUNDATIONS
1950s – 1970s

EARLY SYSTEMS
1980s – 1990s

MACHINE LEARNING
2000s

DEEP LEARNING
2010s

THE AI REVOLUTION
2020s AND BEYOND

MICHAEL C. CASTELLANO

Aka “MICHAEL AI”

How AI Came to Be *A Complete History of Artificial Intelligence*

Michael C. Castellano

aka “Michael AI”

Copyright © 2025 Michael C. Castellano

All Rights Reserved

ISBN:

Dedication

To Alan Turing, who imagined this world before it was possible, and to every dreamer, builder, and unsung hero who carried that vision forward across generations. And to my family, whose love and belief made my own journey possible.

Acknowledgements

No book is written alone, and no journey worth taking is traveled alone. This one is no exception.

Thank you to my mentors and early investors in Engajer: Robert Michael Fried, Bob (“Raz”) Zeicheck, Bert Davey, and Mike Downey. Your guidance and belief in me, long before there was anything concrete to believe in, shaped my career and the way I see the world.

Thank you to Massy Mehdipour and all of Engajer’s early investors, and to the core team who built the vision alongside me: Daniel Chen, Abe Melloh, and Susie Jerez.

To the ghostwriting team who helped bring this story to the page: thank you for your craft and your dedication to getting it right. To my friends and family: thank you for your endless encouragement through every high and low of this journey.

To the entire team at Engajer, Inc., our employees, shareholders, advisors, contractors, and consultants: this story is your story too. I am honored to have built alongside you.

And finally, to all of the men and women in the field of artificial intelligence, from Alan Turing to those working today: you made AI possible. Some of your

names are known around the world; many are not. But every one of you moved humanity forward, and this book is, above all, a way of saying thank you.

Proofreaders and Contributors

A book is shaped not only by its author but by the careful eyes and generous minds of those who read it along the way. I am deeply grateful to the following people, who proofread drafts of this book and contributed to making it better:

Marisa Christine Castellano

About the Author

Michael C. Castellano, known to many simply as “Michael AI,” is an entrepreneur, technologist, and founder whose career spans more than two decades at the frontier of media, software, and artificial intelligence.

His journey began early. At fifteen, he started his first business. At sixteen, he served as a United States Senate Page during the 107th Congress, witnessing leadership and public service firsthand in the halls of American government. He later moved to Silicon Valley for college, where he launched a video production company from his dorm room. That company evolved into a software company, and ultimately into an artificial intelligence company: Engajer, Inc.

Under his leadership, Engajer developed a virtual persona that passed the Turing Test in real life, in an unscripted interaction with a real person who could not tell she was speaking with a machine, realizing a vision Alan Turing first imagined some seventy-five years earlier.

How AI Came to Be is his first book. He wrote it to tell the full, human story behind the technology now reshaping the world, and to help readers move from amazement to understanding.

Preface

Every day, millions of people open a chat window, type a question, and receive an answer from a machine that seems to understand them. They are amazed, and they should be. But very few of them know what it took to get here.

That gap, between amazement and understanding, is the reason this book exists. The story of artificial intelligence has been told in fragments: a headline here, a documentary there, a press release polished to a shine. What has been missing is the whole story, told plainly, from the beginning, by someone who lived inside it.

I have spent more than twenty-two years building technology companies, the last of which, Engager, Inc., achieved something Alan Turing could only dream of in 1950. That experience gave me a front-row seat to the most important technological transformation of our lifetime, and it gave me a conviction: this history deserves to be recorded, honestly and completely, while the people who made it are still here to be recognized.

This book moves in three parts. First, the foundations: Turing's question, the pioneers, and the decades of quiet groundwork. Second, the breakthrough era: the companies, personalities, and rivalries that brought AI into everyday life. And third, the road ahead:

a clear-eyed look at where this technology is taking us over the next century. You do not need a technical background to read it. You only need curiosity.

Welcome to the story of how AI came to be.

Table of Contents

DEDICATION	I
ACKNOWLEDGEMENTS	II
PROOFREADERS AND CONTRIBUTORS	IV
ABOUT THE AUTHOR	V
PREFACE	VII
TABLE OF CONTENTS	IX
INTRODUCTION	1
CHAPTER 1: ALAN TURING AND THE BEGINNING	16
CHAPTER 2: STEVE JOBS AND THE 90S.....	26
CHAPTER 3: EARLY 2000S TO THE PRESENT	51
CHAPTER 4: ENGAJER AND THE HUMAN-MACHINE CONVERSATION.....	80
CHAPTER 5: CHATGPT	88
CHAPTER 6: ELON MUSK’S AI WORK	105
CHAPTER 7: DEEPSEEK	127
CHAPTER 8: ANTHROPIC	140
CHAPTER 9: THE FUTURE OF AI	167
CONCLUSION	186

Introduction

It did not happen overnight.

Artificial intelligence, the buzzword of our time, did not appear suddenly out of nowhere. It did not emerge with a flick of a genius's wrist or from the quiet of a lab filled with wires. The truth is more layered, more human, and far more fascinating.

What you are holding is not another cliché book about AI taking over the world. It is a story of ambition, chaos, struggle, and cooperation across industries, decades, and continents. This book was written to uncover how we really got here, to this moment where intelligent systems are reshaping everything from medicine to media. The record has been scattered and often misunderstood.

The goal is to bring it all together and set it straight.

Behind every breakthrough were teams of people who stayed up too late, ran out of money, argued over code, and believed in something the world was not ready to see. The development of AI is not a polished science fiction tale. It requires patience, conflict, failure, vision, and steady experimentation. Above all, it is a story about people.

Reading this book will give you more than just facts. You will see what it took to bring artificial intelligence

Michael C. Castellano

from concept to reality, the people and principles involved, and the moving parts that had to align, often imperfectly, to create something astonishing. You will see the scale of this effort and gain a deeper appreciation for it.

You will also gain a clear understanding of our current moment through the lens of technology and by looking at its social, economic, and philosophical impacts. AI is more than a tool. It marks a turning point in history. As you uncover the story of how AI came to be, you will also begin to recognize the direction in which it is guiding us.

Some of the glimpses ahead are grounded in fact. Ideas and possibilities shape others. But all of them are rooted in a close understanding of what has already happened. That is what makes this book different from speculation and hype. It is built on history, not headlines.

Now, let's talk about the person behind this story. I, Michael Christopher Castellano, am not an academic locked away in towers. I am not a bestselling author chasing trends. I am someone who lived it. This is my first published book, but my journey into the world of business and technology began long ago. At fifteen, I started my first business. At sixteen, I served as a U.S. Senate Page during the 107th Congress, observing the workings of leadership and service in the halls of government.

How AI Came to Be

Years later, I moved to Silicon Valley for college and launched a video production company from my dorm room. That company grew into a software company and later into an artificial intelligence company. Twenty-two years of passion and purpose led to the insights in these pages. This is not a secondhand account. This is my experience, and I will share it with you.

When you complete this book, you will understand the process that brought artificial intelligence to life. You will see the forces that enabled its creation and appreciate the rewards and risks that lie ahead. You will learn not just the mechanics but the meaning behind the machine. You will know how it all began and where we are likely headed.

To appreciate where artificial intelligence stands today, we must first understand the journey that brought us here. This goes beyond tracing timelines or listing milestones. It requires recognizing the spark that made the journey possible: the questions that challenged knowledge, the breakthroughs that expanded horizons, the mistakes that taught lessons, and the curiosity that refused to stop at what was already known.

This is the story behind the story, the inspiration that led to machines learning to think.

My experience is not the only thing that pushed me to write this book. What encouraged me was something simple but important. I realized how many people are

using AI today. They interact with tools like ChatGPT or image generators like Midjourney and are amazed by the results. Some use AI for work or study, while others explore it out of curiosity. Regardless of purpose, one thing is universal: the sense of wonder that comes with engaging these technologies.

The popularity of AI has grown rapidly, yet the history of how it was developed remains unfamiliar to many. This gap highlights how little is known about its beginnings.

That space between amazement and understanding is what inspired me. I found myself asking, if people knew what it took to get here, if they understood the depth of effort behind the tools they use every day, would they view them differently? Would they feel respect, concern, or even responsibility?

I believe the answer is yes.

It is extraordinary that someone like Alan Turing, back in the 1950s, could imagine what we are now experiencing. At that time, computers were massive, slow, and simple compared to anything we use today. Yet Turing saw something in them, a potential so vast that he suggested a future where a machine might become so intelligent, so convincingly human in its behavior, that a person would not be able to tell whether they were speaking to a human or a machine.

How AI Came to Be

That idea became known as the Turing Test. For decades, it remained an aspiration, a goal just out of reach. The world changed, technology advanced, and nothing had truly passed the test for more than seventy years.

Until it did.

About seventy-five years later, my company, Engager, Inc., quietly achieved something that even the most optimistic researchers thought was still years away. We developed a virtual persona, a digital being, and introduced it to a real person I will call only Barbara, out of respect for her privacy. She believed she was on an ordinary FaceTime video call with me. She spoke with the persona the way anyone speaks with a friend: naturally, back and forth, without hesitation. It was only afterward, when we looked at what had just happened, that we understood its meaning. She genuinely had not known she was interacting with a machine.

The machine, nicknamed Michael, passed the Turing Test not in theory, and not in a controlled lab with graphs and benchmarks, but in real life, with a person who believed she was speaking to another human. No one set out to run a formal experiment that day, and that is precisely what makes the moment so meaningful: the test was real life itself. It happened, and it was no accident.

It took decades of effort, hundreds of brilliant minds, entire teams across dozens of companies, and funding

Michael C. Castellano

that may have reached into the billions. It was not one invention. It was a symphony of innovation, stitched together over time by people who believed this moment could come.

And it did.

It is extraordinary that his vision became reality seventy-five years after Turing first imagined it. That is not just an interesting piece of history. It is a moment that should be remembered. It deserves to be told in full.

That is what I want this book to do. I want to tell that story, the whole story. Not just headlines or polished corporate versions, but the real thing: the messy, brilliant, human journey that brought this technology into the world.

Readers can expect to learn the rich truth. At times, it might even seem miraculous. But everything in this book is part of a real story. This is the story of a technology that now touches the lives of millions, a technology that came to life behind the scenes, shaped by a community of thinkers, builders, and dreamers.

We will explore who these people were. The major players include those who became household names. At the same time, we will uncover the individuals who dedicated their lives to this vision. Alan Turing was one such person who believed with all his heart that this moment would arrive. Many others followed and carried

How AI Came to Be

that torch forward. People like Steve Jobs advanced the movement in their own ways.

They all played a part. I want to go deeper. I want to look into the research, trace the developments, and uncover the real story behind the story. I want to acknowledge everyone who moved this work forward. Recognition goes beyond the well-known figures to include all who contributed, whether or not they choose to be named. For those who prefer to remain anonymous, their wishes will be respected. Giving credit is important, and so is honoring personal choice.

But maybe this is also a chance to recognize some of the unsung heroes. The people who were quietly dedicated to the work. The ones who did not seek the spotlight or chase headlines. They were not after billions of dollars. They were not interested in self-promotion. They simply cared about building something meaningful and gave their energy to that purpose.

Maybe this book is a way to say: We see you.

Or perhaps it gives voice to some people who never imagined they would be mentioned. Maybe it offers recognition to those who truly made a difference.

We face something greater than a tool because it carries profound implications for our future. It is more than software. Just think about how many people are using this today. Think about how much information it

Michael C. Castellano

carries, how much it can teach, and how much it is already transforming the way we live and work. This technology is going to shape a completely different world.

However, when something is booming, different perspectives about it multiply as well.

There is someone the world calls the Godfather of AI: Geoffrey Hinton. He is a deeply respected figure, and he has earned every bit of that respect. His work on backpropagation in the 1980s and on AlexNet in 2012, both of which you will read about in this book, helped build the very foundation of modern artificial intelligence. Yet in recent years he has been making a lot of doom-filled statements. He talks about AI empowering authoritarianism. He warns that it will widen the gap between the rich and the poor.

I want to give credit where credit is due. Few people alive have contributed more to this field than he has, and when someone of his stature raises concerns, those concerns deserve to be heard, not dismissed. But after more than two decades of building in this space, I see the same technology through a very different lens. I just do not share his outlook.

My perspective is very different.

I actually believe that virtual intelligence, or digital intelligence, is going to raise everyone's standard of living. I think it has the potential to create a cleaner, safer, more

How AI Came to Be

equitable, and more peaceful world. That is the direction in which I see this heading.

I am happy to be clear and specific about the future. I can discuss what might unfold in two years, three years, ten, thirty, or even fifty years. I can break down the possibilities because I do not see doom and gloom. Instead, I see opportunity.

I think it is important that people understand what we are really talking about. For one thing, the name is misleading. Artificial intelligence is not a very accurate term. It is *not* artificial. There is nothing fake about it. It is not imaginary, and it is certainly not a trick. It is real. It is virtual. It is digital. It is intelligence that operates in a different form than our own.

We have biological intelligence. That is what human beings are born with. It lives in our bodies, carried by neurons, cells, feelings, and intuition. Digital intelligence, on the other hand, lives in servers and computers. It operates not through instinct or emotion but through data, logic, and patterns. It functions within the realm of numbers and reasoning. Unlike biological intelligence, it does not possess a body or experience sensations such as pain, yet its growing influence on our world cannot be ignored.

Given these differences, there must be a clear distinction between biological and digital forms of

intelligence, not to separate them but to understand how each contributes to shaping our present and future. This could be the subject of its own chapter, exploring how both types of intelligence, though different in origin and function, are powerful in their own ways.

Digital intelligence, which I prefer to call virtual intelligence, will not emerge solely from lines of code or algorithms. It will be guided by people—thinking and feeling individuals who bring intention, morality, and care into its development.

This is our story, too.

I don't believe the future will be doom and gloom. In fact, I think it will be incredibly positive. In my view, virtual intelligence will help expand human consciousness. It's not here to replace us, but to work with us, to elevate the way we think, perceive, and innovate.

I actually believe this process is cyclical. Biological intelligence creates virtual intelligence, which then helps elevate biological intelligence. It's a spiral of growth. A kind of evolving wheel, where human and machine intelligence refines itself through mutual collaboration. We improve the machines, and the machines, in turn, improve us. The result is expansion on every level: material, intellectual, and even spiritual.

This symbiotic relationship between humans and machines could lead to unprecedented acceleration, not

How AI Came to Be

just in terms of technology but also in our understanding, awareness, and consciousness. I truly believe this, and I think there's solid evidence to support it.

By this book's end, readers will feel a sense of relief. A calm confidence. Despite the loud voices predicting apocalypse, extinction, or authoritarian collapse, there is another side to the story. Yes, we should examine those worst-case scenarios. We should discuss them, take them seriously, and be prepared. However, when we look at historical data, the trend is clear.

With technological innovation, the world has consistently improved. Take Peter Diamandis's TED Talk, *Abundance is Our Future*, for example. He walks us through key global indicators. Literacy rates have risen. Infant mortality has dropped. Extreme poverty has been reduced. Hunger has declined. And not just in one region but across the globe. From Florida to Southeast Asia. Over the last 50 years, 100 years, even 1,000 years, human life on this planet has steadily improved by nearly every measurable standard.

Technology, when guided wisely, tends to lift the standard of living. So, while we must absolutely address potential threats, we also need to recognize the power of innovation to move us forward.

That said, strong policies are essential. We need smart, enforceable laws that keep artificial intelligence in check.

For example, if someone creates a video that looks real but is entirely generated by AI, and that video could mislead, scare, or deceive someone, there should be a legal requirement. It must state that AI created it somewhere in the metadata or even visually on the content itself.

We cannot allow artificial intelligence to be used for deception or manipulation. There needs to be accountability. A clear, common-sense approach says: create what you want, innovate boldly, but be transparent. People must be able to distinguish what is real from what is generated.

The future brings challenges but also immense progress. If we are thoughtful, collaborative, and honest in building and regulating these tools, the future can be not only brighter but extraordinary.

If AI creates something, whether a video, image, voice recording, or text, it must be clearly disclosed through metadata or visible on the screen. This is especially critical when the content could be misinterpreted or cause harm. Simple policies like this are not about stifling innovation. They are about maintaining trust, protecting the public, and establishing boundaries in an era of rapid technological change.

Now let's talk about job replacement. It is not enough to say, "AI is here, and jobs are gone." Disruption without accountability is unacceptable. If an organization wants to

How AI Came to Be

implement AI at scale and replace human roles, a mechanism must be in place that shows real improvement in people's lives.

We should ask: What happened to the individuals whose jobs were affected? Did their careers progress? Did they earn more? Did they find more meaningful work? Did their quality of life improve? If data cannot back those answers, then full-scale deployment should not move forward.

AI must be a tool for elevation, not displacement. If it can free us from tedious labor and open the door to more fulfilling work, we must show this clearly with transparent evidence. Proof should leave no doubt and demonstrate measurable impact. If a system improves job satisfaction, increases income, allows more time for family, enhances health, or brings other benefits, it should be embraced. If not, we need to pause and reassess.

This is not about slowing progress. It is about making sure progress works for people, not against them. Guardrails like these ensure virtual intelligence serves humanity's highest good. Policies like this are practical, enforceable, and grounded in fairness. They are essential if we want AI to be a force for widespread benefit, not narrow gain.

It is easy to picture a future where these technologies help us explore other planets or even the galaxy. The

Michael C. Castellano

reality is, we have not set foot on Mars yet. I certainly have not. Have you? No, I have not either. Maybe that day will come. Our focus must remain on building a solid foundation here on Earth, where the impact is immediate and vital. Let us begin by examining AI pioneers' work and the lasting influence they left behind, starting with Alan Turing.

How AI Came to Be

Section 1

Chapter 1: Alan Turing and the Beginning

Before there were intelligent machines, before there were algorithms shaping our daily lives, there was a man who imagined the unimaginable.

His name was Alan Turing.

In many ways, the story of artificial intelligence begins with him. He was not only a brilliant mathematician and logician but also someone who dared to ask questions that most people of his time could not even begin to understand. What does it mean to think? Can a machine ever do it? And if it could, what would that say about human intelligence?

This chapter is a journey into the life and mind of Alan Turing. It explores how his groundbreaking ideas laid the foundation for what we now call artificial intelligence. Turing not only contributed to the field, but he also created a new way of seeing the world. His work during the Second World War, his theories about computation, and his vision of machine intelligence were ahead of his time and deeply human at their core.

Our book's beginning is about him, and that is not only because he was the first pioneer of AI. It is also

How AI Came to Be

because what he stood for created the foundations on which we stand today.

He represents the spirit of curiosity, the courage to think differently, and the quiet strength of those who imagine something far beyond what already exists. Understanding Turing is not only about understanding the roots of AI but also about appreciating the person behind the science.

Let us step into his world and see how it all began.

Alan Mathison Turing was born on June 23, 1912, in London, England. From a young age, he displayed an unusual brilliance that set him apart. He had a natural gift for mathematics and a way of thinking that often seemed detached from the ordinary. He was curious, quiet, and intensely focused. While other children played, Turing solved puzzles, experimented with numbers, and asked deep questions about how things worked. It was clear even then that his mind operated differently.

He received his education at prestigious institutions such as Sherborne School and later King's College, Cambridge, where he earned a first-class honors degree in mathematics. His work at Cambridge soon attracted attention, and it was during this period that he began developing ideas that would change the world.

One of Turing's most profound contributions came in 1936, when he introduced the concept of what is now

known as the Turing Machine. At its core, the Turing Machine was a thought experiment. It was not a physical device, but rather a theoretical model that described how a machine could carry out any mathematical calculation if it followed a set of instructions.

This machine could read and write symbols on an infinite tape, moving step by step according to rules that governed its behavior. Simple as it may seem now, this idea was revolutionary at the time.

The Turing Machine was important not just because it explained how machines could process information, but because it helped define the very limits of what is computable.

In other words, Turing gave the world a way to understand what a machine could and could not do. This was long before computers as we know them existed. Yet his ideas formed the foundation of modern computer science.

Turing's vision showed that complex operations could be broken down into simple, logical steps. It was this insight that would later make artificial intelligence even imaginable. Without the Turing Machine, there would be no algorithmic thinking, no coding as we know it, and perhaps no AI at all.

His invention did not make headlines at the time, and few people outside academic circles truly understood its

How AI Came to Be

power. But today, almost every computer program, every digital system, and every intelligent machine owes something to the basic ideas Alan Turing put forward. His machine was not just a model of computation. It was a glimpse into a future he could see, even when no one else could.

If the Turing Machine opened the door to the idea of machines doing calculations, the Turing Test stepped into something much deeper. It was Alan Turing's way of asking a question that still matters today. Can machines think?

In 1950, Turing published a paper called *Computing Machinery and Intelligence*. In it, he suggested a new way to look at the idea of thinking. Instead of trying to define what thinking is, which is very hard even today, he asked a different question. What if we could tell if a machine is intelligent by how it behaves?

Alan Turing created a test he called the Imitation Game. In this test, a person would type messages to two others. One was another person, the other a machine. If the person could not tell which was human, the machine was said to show intelligence.

The idea was simple but powerful. Turing believed intelligence was more than solving problems or doing math. It was about using language, answering questions, and holding a conversation in a way that felt human. If a

machine could do that, it might deserve to be called intelligent.

The Turing Test changed how people thought about machines. Before this, machines were seen mostly as tools that followed commands. Turing suggested something more. He showed that machines could one day think in ways close to human thought.

The test raised questions that went beyond science. Could a machine be creative? Could it understand feelings, make jokes, or show kindness? These were not only technical questions but human ones. They asked us to think about what it means to be human.

Even today, the Turing Test is part of discussions about artificial intelligence. Some machines have come close to passing it, while others make us wonder what “passing” really means. The fact that the test still matters shows the depth of Turing’s idea.

The Turing Test was not just technical. It was also human. It asked us to imagine a future where people and machines could talk, think, and understand one another.

Turing’s work opened a larger conversation about human and machine relationships. What does it mean to think? What does it mean to be human? And if a machine can act like a person, how should we treat it?

How AI Came to Be

This marked the beginning of the philosophy and ethics of artificial intelligence. Turing showed that building machines was not only about logic and wires. It was also about values, choices, and responsibility.

His questions led people to think about psychology, morality, and behavior. If a machine can make decisions, can it be held responsible? Should it be treated with care if it seems to feel emotions? What would it say about us if a machine could copy the way we speak and think?

These are not simple questions, and they have no easy answers. But they remain important. As AI becomes part of daily life, we face the same challenges Turing first pointed out. How do we build machines that are fair and respectful of human dignity? How do we control systems that may one day outsmart us? How do we protect our humanity while teaching machines to act more like us?

Turing never claimed to know all the answers. What he did was start a conversation that continues today. That conversation has grown into the fields of AI ethics and philosophy, which explore how we build intelligent systems, how we use them, and what kind of future we create with them.

The work of Alan Turing was not the only thing that started artificial intelligence. It was also the beginning of thinking deeply and responsibly about what intelligence means for all of us.

While Turing laid the foundation for thinking about machine intelligence, others began building on his ideas in more practical ways. In the 1940s, Warren McCulloch and Walter Pitts created a model of artificial neurons. These were simple units designed to copy how the human brain works. Their work later influenced the development of neural networks.

Then, in 1956, John McCarthy introduced the term “artificial intelligence”. He organized the Dartmouth Conference, which is widely seen as the official starting point of AI as a field of study. Together, these efforts helped shape the early ideas of what intelligent machines could become, alongside Turing’s groundbreaking work.

Even though Alan Turing introduced his test many decades ago, the questions he raised are still part of how we think about artificial intelligence today. In fact, many of the core ideas behind modern AI can be traced back to the Turing Test. His vision continues to shape how we design, measure, and understand intelligent systems.

One of the main principles in today’s AI is the idea that machines should learn to interact with people in a natural way. Whether it is a voice assistant, a chatbot, or a translation tool, modern AI is built to understand and respond to human language. This is exactly what the Turing Test was about. It asked whether a machine could

How AI Came to Be

hold a conversation that felt human. Today, we see that goal everywhere.

Another principle in AI today is adaptability. Intelligent machines are expected to learn from data and improve over time. This means they should follow instructions and adjust to new situations. While more advanced, this idea still adheres to Turing's original concept. He envisioned machines that could process information in steps, follow logic, and exhibit flexible behavior. Modern AI builds on this by using more complex data and faster computing, but the foundation remains the same.

Fairness and ethics are also major concerns in AI today. We want machines to give balanced results, avoid harmful biases, and make decisions that respect human values. These ideas also come from the questions Turing raised. If a machine can act like a human, then it must be treated with the same care we give to human decisions. We must ask not only what a machine *can* do, but also what it *should* do.

Even the way we evaluate AI systems today often reflects the spirit of the Turing Test. Many tests still focus on how well a machine can understand, respond, and engage with humans. If it feels natural and makes sense, and if people feel they are being understood, the AI is seen

as successful. That basic idea comes from Turing's original vision.

Of course, AI today is far more advanced than anything Turing could have built during his lifetime. We have deep learning, neural networks, and massive data systems. But the heart of the idea remains the same. Can a machine understand us? Can it think with us? Can it act in ways that feel real?

Turing may not have known all the details of what was coming, but he gave us a way to think about it. His test was not just a test. It was a guide. And even now, it helps us ask the right questions as we build the future of AI.

Hence, Alan Turing did not just help build the first ideas of computers. He helped shape how we think about intelligence itself. Through his work, we learned that machines could do more than follow commands. They could learn, respond, and even mimic human thought in ways that were once hard to imagine.

His Turing Machine gave us the basic idea of how computers could work. His Turing Test showed us how we might one day speak to machines as if they were people. And his questions about thinking, feeling, and understanding helped open the door to the deeper ideas surrounding artificial intelligence today.

Turing was ahead of his time, but his thinking still lives on. Every smart assistant, language model, and AI

How AI Came to Be

program that tries to understand what we mean and how we feel is all part of the path he started. We are walking in his footsteps, even now.

More than anything, Turing reminds us that behind every great machine is a great idea. Behind every great idea is a person who was brave enough to ask, *What if?*

As we move forward in the world of AI, it is important to remember where it all began.

Chapter 2: Steve Jobs and the 90s

Artificial intelligence has changed from an idea of the future into a technology that affects our daily lives. This chapter looks at the most important moments in AI history and shows how the field has grown. It starts with the visions of pioneers like Steve Jobs and continues with practical systems like R1 and Deep Blue. Each development made machines smarter and more helpful for humans. We also explore breakthroughs like neural networks, expert systems, and consumer AI, such as Sony's AIBO. These examples show how seemingly impossible ideas have become tools that make life easier and more connected.

Firstly, let's discuss one of the leading pioneers in this field: Steve Jobs.

Steve Jobs is the name people today recall as one of the most important personalities in the technological and innovatory world. His life story is inspiring and complex, many people say, and has been marked by struggles, risks, and achievements that influenced his career and how people around the planet utilize technology. However, in order to see what his vision was and what he contributed to the industries of personal computing, music, and mobile devices, we need to go back to his early life and

How AI Came to Be

the experiences that formed his character and way of thinking.

Steve Jobs was born on February 24, 1955, in San Francisco, California. He was born to his biological parents, Joanne Schieble and Abdulfattah Jandali, who were young students at the university and believed they were incapable of bringing him up when he was born. Consequently, after birth, he was adopted by a working-class couple, Paul and Clara Jobs, and they offered him a stable home. Jobs would frequently comment on how, out of respect towards his adoptive parents, he was able to give them credit in providing him with the feeling of security and freedom that enabled his creativity to thrive.

Jobs was brought up in Mountain View, California, where the fledgling culture of Silicon Valley had been all around him. This setup was the center of his interest in electronics and mechanics. His father, Paul Jobs, was also a machinist and taught him how to dismantle and reassemble machines. This early exposure to instruments and craftwork provided Steve with a practical insight into how things were, and this would come to play later when he designed products that were not just sensible but elegant as well as easy to use.

Jobs' school days were not easy. He was bright and was easily bored with conventional approaches to teaching. He could be challenging for teachers, but he was

curious, and his non-traditional thinking was a plus. By the time he was older, he had developed interests in design, literature, philosophy, and electronics. He had briefly studied at Reed College, Oregon, but dropped out after one semester. Although he abandoned formal education, an experience at Reed permanently impacted him. He still informally audited classes, and one of them was calligraphy. He later said this class directly impacted the beautiful typography and attention to design, which made Apple computers stand out amongst others.

The other significant thing about Jobs' background was that he was exposed to other cultures and philosophies. During his early years, he went to India to seek spiritual counseling and tried the alternative ways of life that were popular in the 1970s. These experiences informed his thinking, prompting him to appreciate simplicity, gut feelings, and other methods of problem-solving. This combination of technical mastery, artistic taste, and philosophical perspective later formed the basis of his personal vision as an innovator.

By the time Jobs joined forces with Steve Wozniak and Ronald Wayne in the founding of Apple in 1976, his varied experience had well equipped him to make something completely new. His childhood experiences made him understand how to integrate technical knowledge with a profound sense of design and experience. This capability of combining technology and

How AI Came to Be

design with human needs was the staple of his career and one of the reasons why his influence is still being felt today.

Steve Jobs was someone who always thought ahead of his time. He did not just concentrate on the creation of computers or phones. He even had a dream of how technology would live with people in the future. Artificial intelligence was one of the notions that were of great interest to him. Jobs used to think that computers would one day not only take orders but also comprehend people. In his case, the role of artificial intelligence was not to substitute human intelligence but to help people using machines to live more easily and naturally.

The early years represented the time when regular people could not easily use computers. They required codes, technical understanding, and much patience. Jobs regarded artificial intelligence as a way of transforming this. He had a desire to have computers learn to know people rather than having people learn to know computers. He used to fantasize about a time when communicating with a computer would be as easy as communicating with a human being. Artificial intelligence could help make technology more human, friendlier, and faster in his mind.

The concept of personal assistance was one of the strongest aspects of his vision. He believed that one day

computers would perform as assistants that are always ready to propose, provide instructions, or perform simple duties. The idea emerged many years later in the form of Siri, Apple's voice assistant. Siri was unveiled just one day before Jobs died, and it carried his vision that artificial intelligence should seem like a companion with you. Jobs did not want machines to compete with people.

Jobs also believed that true intelligence in machines was not just about how much they could calculate. It was about how well they could connect with people. For this reason, he focused on technical power, design, and user experience. He wanted machines that were not just smart but also beautiful, simple, and meaningful to use. He often said that real progress happens when technology meets art and human imagination. That was also how he thought artificial intelligence should grow, not only as a scientific tool but also as something shaped by creativity and empathy.

He also warned against letting technology feel too cold or lifeless. Jobs never wanted artificial intelligence to create stiff, mechanical conversations. Instead, he dreamed of machines that felt natural, warm, and sensitive to people's needs. He pictured computers that could quietly learn in the background, adjust to the habits of their users, and give them more time to focus on what really mattered in life.

How AI Came to Be

This vision of artificial intelligence is still very powerful today. Many companies continue to follow Jobs's idea, which is that technology must always be about people first. Others worked on making computers smarter, but Jobs focused on making them understand people better. This is why his vision remains so important. It was not only about machines becoming intelligent but also about machines making human life richer and easier.

The way Steve Jobs thought about artificial intelligence was not just a dream of the future. Today, many of those ideas can be seen in real systems that make technology more personal, intuitive, and useful. One example of this is Engager, a platform I created that shows how Jobs's vision has become reality.

Engager is built on the same principle Jobs believed in: technology should understand people instead of forcing people to understand machines. It does not expect users to learn difficult commands or adjust themselves to rigid systems. Instead, Engager adapts to the user. It studies preferences, patterns, and behaviors, and then uses this knowledge to make communication smoother and more meaningful. This is exactly what Jobs imagined when he said technology should quietly serve people in the background and make their lives easier.

Just like Jobs once described computers as personal assistants, Engager provides natural and almost human

support. It can guide users through choices, suggest useful options, and respond to needs in real time. The purpose is not to replace human thought but to extend it, giving people quicker access to information, smarter recommendations, and interactions that feel simple. Engager, in this way, is a clear reflection of the personal assistant Jobs hoped artificial intelligence could become.

Another important link between Jobs's vision and Engager is the focus on user experience. Jobs always said that true intelligence in technology is not only about how powerful it is but also about how people feel when using it. He wanted technology to be simple, attractive, and meaningful. Engager follows this principle by creating an experience that is engaging and easy to use. Its strength lies in the way it combines strong artificial intelligence with a design that feels natural to the human user.

Engager also demonstrates another part of Jobs's dream: technology that learns and adapts over time. Jobs believed that the best systems should grow with their users. Engager applies this by continually learning from user interactions. The more it is used, the more accurate and personal it becomes. This evolving nature shows how artificial intelligence can quietly improve life without demanding extra effort from people.

In this way, Engager is more than just a tool.

How AI Came to Be

It is a practical example of Jobs's vision brought to life. Where Jobs once imagined technology that understands humans, Engajer shows that such technology is possible today. It proves that artificial intelligence is about machines becoming stronger and making human life easier, richer, and more connected.

Moreover, it is important to note that in the 90s, Steve Jobs was not the only one thinking about AI. There were numerous other milestones in the subject that are worth noting.

IBM's Deep Blue Defeats Garry Kasparov

Perhaps one of the most well-known artificial intelligence advances was the 1997 case when the then world chess champion, Garry Kasparov, was beaten by the supercomputer Deep Blue, created by IBM. The whole world was concerned by this event since it was no longer a mere game of chess. It was regarded as a landmark occasion when a machine could outwit one of the best human brains in a profession that was traditionally said to be the ultimate test of strategy, memory, and calculation.

Deep Blue differed from standard computers. It was created with the sole aim of excelling in chess at the highest level. Its ability to examine millions of positions every second gave it a power no human brain could equal. When Deep Blue was faced by Kasparov, the best player

to have ever played chess, it was a battle between human emotion and machine computation. Deep Blue emerged victorious after six games, and history was made because it was the first computer system to defeat a world champion under normal tournament play.

The triumph brought up key issues concerning intelligence. On one hand, it demonstrated the phenomenal development of computer science, showing that machines were capable of performing complex tasks faster and more accurately than a human. However, it also reminded people that machine intelligence was not at all similar to human intelligence. Kasparov retrospectively expressed that he found the machine cold and uncreative, which pointed to the weaknesses of AI during the time. The sheer computing power enabled the machine to win, not the type of flexible thinking that humans apply.

Most importantly, even with limitations, the Deep Blue project was still considered a game-changer. It opened people's eyes to the fact that artificial intelligence was not science fiction. It was so strong that it could compete with humans in those spheres when logic and decision-making are needed. It was also an occasion that inspired fresh debates on how AI has the potential to expand beyond chess and be used in medicine, finance, engineering, and solving ordinary problems.

How AI Came to Be

This historical event resonated well with Steve Jobs' vision. When Deep Blue demonstrated the sheer power of machines, Jobs believed that the future of AI must not be only power but also the knowledge of people. In the case of Jobs, the only way forward would be when machines would be intelligent, empathetic, and designed. The success of Deep Blue proved that machines could surpass human ability in certain fields. Jobs aimed to carry that idea forward, focusing on how computers could enrich life and support people in meaningful ways.

R1 (XCON) and its Operation at DEC

The other key historical event in artificial intelligence was the creation of R1 or XCON. This system first went into operation at Digital Equipment Corporation, or DEC, in the early 1980s. R1 created its presence in a less dramatic yet equally effective manner, unlike Deep Blue of IBM, which achieved immortality due to its thrilling chess win. It demonstrated that artificial intelligence would be able to address practical business issues and introduce tangible advantages to business.

DEC was a large computer firm during the period and was the manufacturer of powerful minicomputers, which were extensively used by industries and research organizations. However, one of its greatest problems was making these complex machines customer-configured. Systems could have many potential components and

configurations, and determining how to assemble them to work properly was a complex and laborious task. Human experts were prone to errors, which slowed down the orders and added expenses.

DEC designed R1 to resolve this issue. It was an expert system; that is, it made use of human experts' knowledge and applied it automatically to new circumstances. Simply stated, it was a computerized version of an experienced engineer. R1 could accept customer orders, identify flaws, and propose the correct mix of components per computer system. This conserved time, minimized errors, and enabled DEC to provide systems with better speed and reliability.

R1 was considered a breakthrough in its time. It was one of the first commercial systems to apply artificial intelligence at such a large scale. By the mid-1980s, R1 was saving DEC an estimated twenty-five million dollars every year. It turned out that AI was not a mere scholarly project but a useful solution capable of making businesses more efficient and generating actual value.

R1 also affected the way individuals thought about the future of artificial intelligence. R1 demonstrated that machines could be used to support human intelligence, not to replace it, as in chess, thus the need to think of machines as opponents of human intelligence. This concept was very close to Steve Jobs' vision. He was sure

How AI Came to Be

that technology must perform the role of an assistant that helps with human creation and avoids unnecessary work. R1 proved that AI may take over regular decision-making processes to allow people to concentrate on more productive tasks.

Though systems such as R1, which were developed by experts, had their limitations and were later replaced by new forms of AI, they made a permanent impression. They provided a basis in the current applications of AI to inform decisions, identify mistakes, and provide suggestions in fields such as healthcare, finance, and customer service. To this extent, R1 was one of the first indications of how AI would gradually influence the daily routine and commerce, largely in accordance with what Jobs had thought of as technology, which serves behind the scenes and makes life easier.

Japan's Fifth Generation Computer Systems Project (FGCS)

The fifth-generation computer systems project in Japan, commonly known as FGCS, was another significant moment in the history of artificial intelligence. The project started in the early 1980s and was among the most ambitious efforts to advance AI research on the national level. The FGCS project was more visionary compared to Deep Blue and R1, which demonstrated practical successes. It was determined to totally redefine

the capability of computers, emphasizing intelligent machines that could think much more like humans.

The Japanese government initiated FGCS, with the aim of putting Japan at the center of computing worldwide. The United States had dominated the computer science progress at that time, and Japan wanted to redress it. It was planned to develop computers that are able to comprehend and reason, using logic programming, a direction where machines are able to solve problems by using logic and reasoning like humans do. These new systems were supposed to be knowledge-based, language-based, and complex decision-making as opposed to the traditional computers that only handled numbers and instructions.

The magnitude of the project was monolithic. It was heavily funded by the Japanese Ministry of International Trade and Industry and united the best researchers from universities, industry, and government. The vision was not merely to make faster computers but also intelligent systems capable of supporting scientific research, management of business, and even normal human communication. The project attracted international attention, and several nations were concerned that Japan could become the leader of the new computing age.

Despite its bold goals, the FGCS project faced major challenges. Logic programming turned out to be less

How AI Came to Be

effective for broad intelligence than researchers had hoped. The computers built during the project could perform specific tasks but struggled to match the flexibility of human thought. By the 1990s, the project was considered a failure in achieving its grand vision. However, it still left an important legacy. It encouraged large-scale collaboration in AI research, showed the importance of government support in technological innovation, and inspired other countries to invest more seriously in artificial intelligence.

When we connect this to Steve Jobs's vision, the FGCS story provides an interesting contrast. While Japan focused on building machines that imitated human reasoning through logic, Jobs stressed that the real progress would come from designing machines that served human needs in simple, intuitive ways. The FGCS project proved how difficult it was to copy human intelligence directly. Jobs's approach, on the other hand, was to blend intelligence with design, so that even if machines were not fully "thinking" like humans, they could still feel intelligent in how they interacted with people.

In this way, the FGCS project is remembered as both a bold attempt and a valuable lesson in AI history. It showed that the dream of intelligent machines is complicated, but it also encouraged a generation of

researchers and companies to think big about what AI could become.

Doug Lenat and the Creation of a Common-Sense Knowledge Base

Doug Lenat's work stands as one of the most ambitious projects in the history of artificial intelligence. He set out to tackle a fundamental problem that had limited AI for many years: computers could perform calculations and follow programmed instructions, but they struggled with the kind of "common sense" that comes naturally to human beings. To bridge this gap, Lenat began developing a massive knowledge base that would capture everyday facts, assumptions, and reasoning patterns. His project, known as Cyc, aimed to encode millions of small truths about the world, ranging from the simple observation that water is wet to more complex ideas about cause and effect.

The goal was not just to store facts but to give computers the ability to reason in a way that resembled human thought. For example, if a system knows that fire is hot and that touching hot objects can cause pain, it can infer that touching fire would hurt. This type of reasoning seems obvious to humans, but computers had long struggled to make such connections. Lenat's vision was to give AI the power to handle ambiguity and context, which are central to human understanding.

How AI Came to Be

Over the years, Cyc became one of the most detailed attempts to formalize knowledge in machine-readable form. While the project faced many challenges and never fully achieved its boldest goals, it left an important mark on AI research. It showed that common-sense knowledge was not only valuable but essential for building systems that could interact naturally with people and solve problems in real-world settings. Lenat's work influenced future AI systems, including those that power today's assistants, which rely on vast databases of information and reasoning frameworks to answer questions, make suggestions, and hold conversations that feel closer to human interaction.

David Rumelhart and the Backpropagation Breakthrough

A major turning point in artificial intelligence came with the work of David Rumelhart and his colleagues, who published a paper on backpropagation for training neural networks. Before this development, neural networks had been studied for decades, but they were limited in practice because researchers did not have an effective way to adjust the strength of connections between layers of neurons. This meant that networks could only solve very simple problems and could not be scaled to handle more complex tasks.

Backpropagation solved this challenge by introducing a method to refine a network's internal weights through repeated cycles of learning. The process worked by comparing the network's output with the desired result, calculating the error, and then sending that error backward through the network to update the connections. The system gradually improved its performance by repeating this process many times, becoming more accurate at recognizing patterns and solving problems.

This breakthrough opened the door to training deeper and more powerful neural networks. For the first time, AI systems could learn from large sets of data in a structured way, capturing relationships and patterns that were far too complex to be programmed by hand. Although computing power and data availability in the 1980s limited the immediate impact, backpropagation laid the foundation for what later became the field of deep learning.

Today, nearly every advanced AI application, from speech recognition to image classification and natural language processing, builds upon the principles of backpropagation. Without this innovation, the progress that followed in machine learning and the creation of systems that power tools like Engager would not have been possible. Rumelhart and his team provided AI with the essential mechanism it needed to grow into the modern era.

Gerald Tesauro and the Success of TD-Gammon

Another important achievement in the history of artificial intelligence came from Gerald Tesauro at IBM, who developed a program called TD-Gammon in the early 1990s. This system was designed to play the game of backgammon, a game that requires both strategy and the ability to handle chance. What made TD-Gammon remarkable was that it did not rely on being programmed with detailed strategies created by human experts. Instead, it learned to play by teaching itself through repeated practice.

Tesauro used a method known as reinforcement learning, where the system improved its play by receiving feedback from its own moves. Over thousands of games, TD-Gammon adjusted its internal model to recognize which choices led to better outcomes. By combining reinforcement learning with the backpropagation technique that had been introduced earlier by David Rumelhart and his colleagues, Tesauro's program steadily became stronger. In time, TD-Gammon reached a level of play that matched and even surpassed many human experts.

This success was a turning point because it showed that machines could reach high levels of performance in complex tasks without relying entirely on human

guidance. It suggested that AI could discover new strategies and insights on its own, simply by exploring possibilities and learning from experience. The ideas proven in TD-Gammon influenced many later systems in areas such as robotics, decision-making, and game-playing AI, including the famous AlphaGo program decades later.

The lesson of TD-Gammon continues to shape AI today. It demonstrated the power of combining learning methods with self-improvement, showing that machines could evolve their abilities through interaction rather than static programming. This principle is central to the kinds of adaptive systems we now see in tools like Engager, which continuously refine their responses through exposure to human communication and interaction.

Support Vector Machines and Their Growing Importance

Artificial intelligence evolved into a significant mechanism: the Support Vector Machine, or SVM. SVMs can be described as a form of supervised machine learning that is used to classify and regress. The point is that the computer determines the most optimal boundary, which is known as a hyperplane, to divide various types of data. This aids the computer in the correct sorting of new information, even in cases where the information is complicated or overlapping.

How AI Came to Be

The popularity of SVMs grew due to the fact that they performed well in situations where other machine learning algorithms failed. They could process high-dimensional data, which implied that they were good at working with data sets with numerous variables. SVMs were successfully used by people in text classification, image recognition, and bioinformatics, which frequently have large and complex datasets. On the other hand, simpler methods were not as good as this method in terms of accuracy and results, which is why many researchers relied on it to get real-world results.

Another trend in artificial intelligence was the emergence of SVMs, based on practical, data-driven solutions. Early AI systems, such as Deep Blue or FGCS, attempted to loosely recreate intelligence; however, SVMs demonstrated that AI could provide quantifiable results in certain activities. Most contemporary AI systems have relied on pattern recognition and prediction, and they became the foundation of SVMs.

The later AI technologies, such as neural networks and ensemble methods, were also developed with the help of SVMs. Scholars demonstrated that mathematical techniques developed efficiently may allow computers to make choices and forecasts with minimal human intervention. The concept of a smart algorithm finding patterns is also applied in other systems, such as Engajer,

a platform that learns each user's patterns and provides relevant suggestions.

Hidden Markov Models and Bayesian Networks in Speech and Language

Algorithms such as Hidden Markov Models (HMMs) and Bayesian networks, which were used in speech recognition and natural language processing, were also advanced by AI. The algorithms enable computers to comprehend the human language in a manner that was not previously feasible. The previous systems had restricted rules and precise commands, whereas HMMs and Bayesian networks would tolerate uncertainty, variability, and ambiguous language, which occur in human language.

HMMs proved to be very handy in speech recognition. Human speech varies in pronunciation, tone, speed, and accent. HMMs allow computers to model acoustic sequences and supply the probability of the most likely words or phrases, even when people were not perfect speakers. This enhanced the voice systems to a great degree. Dictation programs, voice recognition, and automatic call centers became more accessible to people.

Bayesian networks were effective in interpreting the relationship between variables and were also used to make forecasts using probability. They assisted computers in making reasoning in uncertain information, such as

How AI Came to Be

inference of context, inference of user intent, or interpretation of ambiguous sentences. These algorithms generated flexibility and capacity to deal with real-life communication in natural language processing by applying statistical reasoning.

The application of HMMs and Bayesian networks represented the transition of AI from strict, rule-based systems to learning-based systems utilizing probability. Computers would be able to modify themselves to human behavior, learn patterns, and come up with clever predictions despite incomplete information. These inventions formed the foundation of the contemporary AI capable of communicating with humans, such as voice assistants, chatbots, and systems such as Engajer. Engajer has also approached language patterns and needs prediction, and responds humanely and naturally.

Sony AIBO and AI in the Consumer Market

One of the landmarks in the evolution of artificial intelligence was 1999, when the Sony company decided to introduce to the consumer market an artificial intelligence pet in the form of Sony AIBO. AIBO was among the first instances when AI was not applied in laboratories or industrial use, but in regular homes. The purpose of AIBO was not to calculate, process data, or solve a problem like in former AI systems, but to communicate

emotionally with people and make technology more human and interesting.

AIBO also understood voice recognition, command responses, and even learned from its surroundings. The processing of all robots can result in various behaviors and create the impression of personality and individuality. This was an important advancement in the field of AI, which proved that machines could not only do things but also socialize, adapt to people, and be friends. Many users were interested in the ability of AIBO to respond to gestures, sounds, and movements and act as a living pet.

The advent of AIBO has also put into the spotlight the increasing use of AI in entertainment and lifestyle. It was no longer merely a question of performing helpful functions but of improving experiences and establishing emotional relations between humans and machines. AIBO showed that artificial intelligence can be made to seem friendly, even personal, and bring the technology closer to the lives of commoners.

Compared to current systems such as Engager, AIBO is a step towards AI. It can learn by interacting with people and responding to them. Like AIBO, which changed its behavior according to the users' instructions, Engager applies AI to learn preferences, anticipate needs, and give its users personalized recommendations. AIBO demonstrated that AI could go further than the technical

How AI Came to Be

functions of the technology and remain capable of providing meaningful and human experiences, which still guides the design of consumer-oriented AI today.

Hence, these were the milestones that were achieved during that era.

Although the progress of artificial intelligence had numerous successes, there were also slow points in its development, which became known as the AI winters. Throughout such periods, investments in AI studies were declining, initiatives were postponed or scrapped, and the public lost interest in AI. Scientists had a hard time living up to the lofty hopes of initial achievements, and much of what AI was supposed to be became not instantly available.

The initial AI winter was experienced in the mid-1970s, when the early systems were not able to deal with the complexity in the world. It turned out that it was far more difficult than assumed to develop machine-like intelligence that is human-like. Subsequently, an additional AI winter occurred in the late 1980s and early 1990s. Expert systems, although helpful in a few applications, could not be applied to increasing expectations of general intelligence. Governments and companies began to cut expenditures, and in most parts of the world, AI research decelerated.

Even though these times were hard, they also contributed to the maturing of the field of AI. Scientists learned how to concentrate on what was possible, how to enhance approaches, and how to develop more workable applications. The AI winters reminded scientists that the process of evolution should be gradual and that it can take time to create intelligent systems by being innovative and patient. These teachings were useful in shaping the creation of current-day AI, such as Engager, which integrates learning, flexibility, and utility to fulfill genuine human needs.

The history of artificial intelligence demonstrates how ideas may be turned into actual tools by human creativity and hard work. Deep Blue, R1, Cyc, and TD-Gammon systems show how AI may solve problems and learn in various formats. Even throughout the AI winters when AI power slowed and funding fell, researchers' efforts did not go to waste; they learned valuable lessons that enhanced the AI. In the present world, the lessons are applied today on modern platforms such as Engager to learn with human beings, give advice, and interact in a natural way. The AI story demonstrates that technology could be effective, but it should be used to assist people, ease their lives, and promote human interests.

Chapter 3: Early 2000s to the Present

The beginning of the 21st century became the new era of artificial intelligence. Computers can now work at a much faster speed, data has widened, and the internet has linked people and machines together in different ways. These advancements provided AI with room to grow, leave research laboratories, and enter the world of use. At this time, AI ceased to exist as a specialized device utilized by professionals and became a part of our lives. Starting with search engines and social media and then moving to smartphones and online shopping, AI began to influence the way individuals work, communicate, and make decisions.

The 21st century transformed the world in a number of significant ways. Technology was more rapid than ever, and became the center of the lives of nearly all. The advent of the internet has interconnected individuals working in different nations and has simplified access to information. Smartphones had put computers into the hands of people, and social media had provided novel means for people to share ideas, build communities, and shape opinions. Simultaneously, the volume of data increased in huge proportions, with virtually all activities starting to leave a digital footprint.

The world was also influenced by globalization. The economies became more interdependent, and companies began relying on the global markets and supply chains. This generated new growth opportunities and new challenges, including financial crises, increased competition, and arguments about inequality. There were also big issues such as global climate change, health crises such as pandemics, and cross-border security threats.

It is during the mid-course of these changes that artificial intelligence was created. It has turned out to be the most important instrument for comprehending substantial volumes of data, enhancing communication, and seeking solutions to complicated issues. AI did not remain in the realm of research facilities but started having an impact on healthcare, education, finance, and daily technology. The 21st-century changes provided AI with the environment it needed to develop and demonstrate its significance in many ways.

Xbox 360 Kinect Puts Artificial Intelligence in Your Living Room

One of the first occasions when artificial intelligence interactively appeared in people's homes was the release of the Xbox 360 Kinect in 2010. Hitherto, AI was something that most people identified with research, industry, or extremely advanced technologies, something

How AI Came to Be

that seemed way out of life. Kinect was the first to introduce AI into a family-friendly entertainment system.

Kinect did not have controllers. Instead, it employed cameras, depth sensors, and microphones to identify human motion and voice. The system was able to identify the body positions, monitor gestures, and react to speech commands. It was the first experience of AI reacting to the actions of many users in real time. The game reacted immediately when a player jumped, waved, or shouted. This was quite unlike the act of pressing buttons or moving joysticks, and it was as though the machine knew the player.

The Kinect technology was on the cutting edge. It relied on computer vision to scan a player's body and detect important points, including hands, legs, and head. It used voice recognition to comprehend basic speech and relate it to commands as well. Such AI methods had been used previously in the academic setting, but Kinect made them available in regular living rooms. It demonstrated the potential of artificial intelligence to transition between the realm of theory and practice in a manner that might be accessible and entertaining.

Kinect was not only affecting gaming. Other uses began to be made of the device by the scientists, engineers, and even doctors. It was also used in healthcare, especially in physical therapy, where patients

could do exercises, and the system captured the movement and provided feedback. Developers in robotics can utilize sensors to help machines perceive the environment. In education, Kinect was used by teachers and students alike as a way of interaction. These instances demonstrated that AI-powered technologies can be versatile and used in other applications beyond their market.

Kinect was also important since it influenced the way individuals believed in AI. It turned out that people were receptive to artificial intelligence provided that it enhanced their experiences, and it felt natural to use. The machine was able to make AI appear less terrifying and more of a companion in everyday life. It demonstrated that machines could react to human gestures and voices, not only typed commands or clicks. This emerging form of interaction was a preview of the future of AI, in which natural interfaces such as speech and movement would be considered expected methods of interacting with technology.

Although Kinect later fell out of favor in the gaming market, it still significantly impacted the history of AI. It also became a turning point as AI was introduced into the home and became part of the daily leisure and social activities. Subsequent advancements in voice recognition systems, such as Alexa and Google Home, or gesture-based interfaces in phones and VR technologies, were

How AI Came to Be

based on Kinect's principles. The machine reminded the people that AI was not all about brute computing power or sophisticated data analysis. It may also be about more human-oriented means of communication between people and machines.

IBM Watson Wins Jeopardy

The history of IBM's Watson began in 2011, when it won the television quiz show Jeopardy! against two of its greatest champions, Ken Jennings and Brad Rutter. This was not entertainment only. It marked a significant advance in the field of artificial intelligence since Watson demonstrated that it could both comprehend human language and search and answer questions within a competitive environment quickly. It was aired to millions of viewers and was one of the most public displays of AI power to date.

Watson was quite unlike other previous systems, such as Deep Blue, which had defeated a chess champion fourteen years earlier. Deep Blue was based on calculating and brute force computations. Watson, however, had to work with the issue of human language. Jeopardy! questions commonly use wordplay, puns, and cultural allusions. This required Watson to interpret natural language, deduce context, plow through data loads, and then select the most suitable answer with certainty.

In order to accomplish this, Watson employed sophisticated machine learning, high-level natural language processing, and access to vast bodies of knowledge. It dissected questions and examined potential answers by splitting them into fragments, considering the potential meanings of the question, and then scoring the potential answers. It was not only about finding facts but also about decision-making under time pressure and weighing probabilities. It was a tremendous AI accomplishment in that it demonstrated that machines could transcend designed inputs and react to the unstructured and undemanding human form of communicating.

The win also altered people's attitudes toward AI. The idea of machines winning over human champions made people understand that machines could compete against professionals in the spheres to which knowledge and reasoning were relevant. The suspected implications it might have on jobs, education, and sectors that need skills cast doubt on this type of technology. Other individuals were excited about the expectations for it, whereas others feared how robots would take away human jobs.

After the Jeopardy! win, IBM also started implementing Watson technology into real-world environments. Watson was deployed in healthcare to assist doctors in analyzing medical reports and research publications and reviewing patient histories to

How AI Came to Be

recommend potential diagnoses or treatments. In finance, it assisted in the analysis of risks and customer support. It drove chatbots in the area of customer service, which could communicate with customers. Those applications demonstrated that the very system used to respond to questions on a quiz could have been used to assist in making complex decisions in a critical situation.

Watson's success has turned into a representation of the expanded place of AI in society. It demonstrated that artificial intelligence had come out of the world of games or controlled experiments. Rather, it might operate in an open condition where language, context, and uncertainty were important issues. Watson additionally led the advances in natural language processing and influenced subsequent systems, including digital assistants, such as Siri and Alexa, as well as Google Assistant, that are used today by millions of people.

Although Watson later struggled, particularly in providing repeatable results in the medical field, its Jeopardy! win remains a historic event. It symbolized the possibilities and the dilemmas in AI. The incident revealed that AI has the potential to be potent and amazing, yet, at the same time, it made researchers and businesses learn that blindly implementing AI into practice is a matter of design and lifelong enhancement.

AlexNet Wins ImageNet

A system named AlexNet altered the path of artificial intelligence and won the ImageNet competition in 2012. ImageNet was a massive competition in which AI models were basing themselves on millions of pictures to recognize and categorize objects. Over the years, researchers have experimented with various techniques, although the development was gradual, with the error rates remaining high. AlexNet broke through by applying a deep learning method, which was considerably superior to all previous systems.

It was Alex Krizhevsky, Ilya Sutskever, and Geoffrey Hinton who created AlexNet. It relied on deep convolutional neural networks, an AI architecture modeled after how the human brain receives visual data. Although a neural network was not a new development, slow computers and limited amounts of data in these networks were usually constraining. By 2012, one could train very large models with more powerful graphics processing units (GPUs) and massive labeled datasets (such as ImageNet). AlexNet capitalized on these circumstances and delivered results that were only seen in movies.

The system significantly reduced the mistakes of classifying the images compared to its competitors. This was not just a little enhancement. It was an obvious indication that deep learning was able to attain such a magnitude of performance that its predecessors were

How AI Came to Be

incapable of accomplishing. This was viewed by many researchers as the point at which the machine of deep learning was transferred away from theory to practice. The success of AlexNet promptly emboldened researchers to use such techniques, and deep learning rapidly became the paradigm in computer vision.

The win of AlexNet held significant importance as it confirmed that machines are capable of identifying objects in images with a level of accuracy similar to that of humans. This paved the way to the use of applications, which touched day-to-day life. Autopilot-driven vehicles, facial recognition software, computer-based medical imaging perception, and sophisticated security solutions were all the follow-up products of the breakthroughs that AlexNet exhibited. Deep learning has also become appealing to big technology players, whose focus has pushed them into funding research and applications.

Over and above the images, AlexNet had an influence. It marked the beginning of a wider change in artificial intelligence to deep learning as the prevailing mode. Scholars started using such models in speech recognition, natural language, and other fields. In the following ten or so years, virtually all key advances in AI, including translation engines and generative models, were at least somewhat grounded on the principles of AlexNet.

Michael C. Castellano

AlexNet is known in history as the system that opened the potential of deep learning. It added superior equipment, huge files, and smart design to demonstrate results that no one could have imagined at that time. It still serves as an iconic example in the history of AI, as it demonstrated that AI could progress faster once the proper techniques were applied under the proper circumstances.

NEIL Learns Visual Knowledge Automatically at CMU

In 2013, the researchers at Carnegie Mellon University launched a project named NEIL, or, to be more exact, Never Ending Image Learner. NEIL further contrasted the existing AI systems in the sense that it attempted to automatically acquire visual knowledge and did not have to wait until humans label it or take a step-by-step approach. Rather than being trained on a fixed set of data, NEIL went over the internet to obtain images and, in the process, trained and formed associations between objects, actions, and scenes.

NEIL was primarily aimed at being a continuum of learning, similar to the manner in which human beings do so with time. It was scanning millions of web-based pictures each day and attempting different tasks. As an example, it started to acknowledge that cars frequently appear on a street, or that a kitchen would have a stove,

How AI Came to Be

refrigerator, and sink. As it went on, it developed an expanding body of visual concepts and the connections between them.

This was a big development as the majority of AI then relied on labeled datasets to a larger extent, which involved human involvement in preparation. NEIL exhibited another option, where machines might learn on their own by watching the large pool of images available within the world. It had flaws and failed in certain ways, but it was the first step in the direction of an independent process of knowledge, where AI would be able to continue learning even without direct oversight by a human being.

Another benefit of AI that the project demonstrated was how the visual image might begin to unify with common sense. Through the combination and repetition of the images, NEIL could not only record what objects were but also their interrelationships. This mattered, seeing as genuine intelligence needs more than acknowledgement. It requires context. Examples: The realization that a bed belongs to a bedroom or that people put on helmets even as they are riding their bikes via the association of facts on the ground.

The development of NEIL also created new doubts regarding the future of AI. What would be the boundaries to continuing learning if machines were able to continue

learning on their own? Would they finally create knowledge bases that would compete with human knowledge? Although NEIL was also an early experiment, further AI research on self-learning systems and the concept of lifelong learning was inspired by it.

In the grander context, NEIL symbolized the development of the AI systems, which finished a single task, to the systems that attempted to learn as humans would like to learn, through observing and linking facts as time went by. This method touched the future of unsupervised learning as well as cemented the notion that artificial intelligence must not just carry out precise tasks but create knowledge that receives advances due to continuous growth and modification.

Deep Q-Networks (DQN) Master Atari Games at DeepMind

In 2013, Google DeepMind published a system, known as Deep Q-Networks, or DQN, that clearly showed how artificial intelligence can be trained to play Atari video games by using deep reinforcement learning. This was a big step as the system was not programmed with the rules of the games; it learned them. Rather, it did the same games over and over, was rewarded in accordance with acting well, and got better with time.

Reinforcement learning is grounded in the concept of trial and error. The AI will experiment with an action, see

How AI Came to Be

what occurs, and keep its strategy optimizing rewards in the future. DeepMind combined this approach with deep neural networks, which enabled the system to work with the visual information in the Atari screen and perceive the game setting. The machine knew nothing beforehand. It read straight off the screen, along with the points it received.

This was innovative in several ways in that DQN could achieve human scores in several classic Atari games, such as Breakout, Pong, and Space Invaders. In others, it even surpassed the human players who were not the best. It was a dramatic sight as the system improved. As can be seen in Breakout, initially, it played randomly, but with hours of training, it learned the trick of boring holes through the bricks to get high scores, which even human players deemed elite.

The games were not the only use of the success of DQN. It demonstrated that reinforcement learning, together with deep networks, would be able to tackle problems without having human-codified strategies. This implies that AI has the potential to perform complex tasks in most areas, so long as it has an opportunity to act in an environment and get a response.

The success of DeepMind was also noticed by the rest of the world. It demonstrated how AI will be able to transition past recognition efforts, such as recognizing

images or speech, to actually making decisions and strategies. It also brought about reinforcement learning in robotics, health care, and energy control, amongst other real-world scenarios.

The DQN resulted in work that would form the basis of even more advanced systems, including AlphaGo, which went ahead and beat the world champion in the game of Go. Most of the methods of modern reinforcement learning can all be traced to this success. It is remembered in history that DQN was the system that demonstrated that complex skills could be learned by experience, in a way that would resemble human learning.

Google Releases TensorFlow

In 2015, Google released TensorFlow, an open-source machine learning and deep learning software library. This was a highly significant move since, before TensorFlow, much of the sophisticated AI code remained behind the walls of research organizations or major corporations. Giving researchers, students, and developers all over the world the opportunity to create and teach AI models using TensorFlow, Google provided them with a chance to apply the technology in diverse practical scenarios.

The strength of TensorFlow was that it enabled individuals to create and execute neural networks with less difficulty. The architecture handled the intricate computation of training immense models, even across

How AI Came to Be

multiple computers. People did not have to write in all the details initially; thus, they could concentrate on creating solutions. The applications were efficient in most applications, such as image recognition, natural language processing, and speech recognition.

TensorFlow's open-source release made it special. Google released it for free, and this enabled a worldwide community to use it, develop it, and create tutorials. It became possible to utilize the identical technology used by Google in their own systems by universities, startups, and even hobbyists.

TensorFlow stimulated the development of AI research at an accelerated rate. Code could be shared, experiments replicated, and other researchers could build upon the work of others. TensorFlow was used to construct many subsequent advances in fields such as medicine, robotics, and self-driving cars.

The launch also resulted in a shift in large corporations' attitudes toward AI. Google was not obsessive about holding tools in secret but found it worthwhile to share them with the world. This gave Google a sway in artificial intelligence research and provided the company with fresh ideas and enhancements from the developers around the globe.

TensorFlow is remembered as one of the tools that brought deep learning into the hands of millions of people

in the history of AI. It turned AI not only into the domain of a few specialists who could learn it but of the general population.

Facebook DeepFace's Accuracy Approaches a Human's

In 2014, Facebook released DeepFace, a deep learning face recognition system. This system shocked the world since it had a very high accuracy compared to that of humans. This model was found to identify a face with an accuracy of 97.35 percent in a photo, whilst humans scored closely at 97.53 percent. This translated to a near disappearance of the ability gap between human and machine ability.

DeepFace was powered by a deep neural network that was trained on millions of images. The network was trained to identify details in the human face, e.g., distances between the eyes, the nose's shape, and the mouth's curvature. In contrast with the older systems, DeepFace did not need to use only straightforward features. It interpreted full face photographs three-dimensionally, meaning that the system could cope with angle, light, or expression variation.

This was the result of the strength of deep learning. Face recognition was one of the most difficult problems in computer vision, but DeepFace made it look easy. In the case of Facebook, the system served to label people in

How AI Came to Be

pictures automatically, which enhanced user experience on the site.

DeepFace portrayed more than just an effect on social media. It was established that machines could perform at the level of a human being in activities that appeared too complex. This outcome led to more studies on the use of deep learning in other recognition tasks, including object recognition, medical image recognition, and even handwriting recognition.

DeepFace created another significant concern regarding privacy and ethics. When a machine could track anyone in a photo and even in a busy location, then concerns about tracking and abuse of technology became too tangible. People wondered how governments and companies would manage such power.

DeepFace is one of the largest milestones in artificial intelligence in history. It demonstrated that deep learning would bring machines practically to the same level of performance as human beings in one of the most complicated recognition tasks.

AlphaGo Defeats Lee Sedol

In 2016, Google DeepMind's AlphaGo defeated Lee Sedol, one of the most skilled professional players of the board game Go. This was a historic and shocking moment

since Go is much more intricate than chess. Possible moves to make in the game are virtually infinite, and until recently, many experts were confident that they could not be mastered by any computer.

AlphaGo proved them wrong. The system took the form of deep neural networks with reinforcement learning. It was trained through the billions of games it played against itself and learned mistakes made by human masters. In the process, AlphaGo has designed approaches that the top players have never considered.

The world was watching in 2016 when AlphaGo played against Lee Sedol. Lee lost four of the five games. Among its moves, one called Move 37 shocked everybody, including professional Go players. It demonstrated ability that no one expected of a machine.

The triumph of AlphaGo is a new step in artificial intelligence. It showed that machines were able to process data and make decisions and exhibit complexity in a sophisticated environment. The success sparked more investigations into how AI could be applied in practical fields that also involve the need for strategy, planning, and decision-making.

The event brought crowds together and controversy. The success was celebrated, but people also started to wonder what was going to happen to our human jobs in sectors that machines could do better. The fact that

How AI Came to Be

AlphaGo revolutionized the sphere of AI was not only that; it transformed the way people think about intelligent machines.

Hanson Robotics Unveils Sophia

Back in 2016, Hanson Robotics also launched a humanoid robot named Sophia, which operates using artificial intelligence. Sophia soon achieved celebrity status because she presented herself and appeared most human. She would display facial expressions, be able to identify people, engage in simple dialogues, and even make jokes. Seeing Sophia was bringing something alive for many individuals who had not realized how it was possible to make the robot look nearly alive.

Sophia had AI systems in speech recognition, natural language processing, and facial recognition. She had cameras in her eyes so that she could trace the faces and look straight into the eyes. Her AI was also able to interpret queries and respond in real time. Through interactions, she could learn too, so her responses improved as time elapsed.

Sophia became a global icon. She appeared on television programs, gave interviews, and even became an honorary citizen of Saudi Arabia. Hanson Robotics did not only show her as a technological device but also as something to begin discussing the future of AI and the interaction of humans and robots.

The influence of Sophia exceeded her powers. She symbolized the transformation of AI and AI moving into the visible and touchable computers that could be observed and interacted with in real life. This brought it closer to people by experiencing the development of AI as real.

Simultaneously, Sophia raised new questions about morality and about what we demand of machines. There were individuals concerned about the extent to which human-like robots should be developed, whereas others objected to humanizing a robot. Nevertheless, Sophia demonstrated the path of AI and robotics convergence towards the intelligent component that enables human-centered design.

Sophia is one of the symbols of the attempt to get AI as close to a human life as possible. She showed how artificial intelligence can be embodied into a format that relates not only based on data and logic but also on emotion and appearance.

DeepMind Publishes AlphaZero

In 2017, DeepMind made available AlphaZero, which is an enhanced version of its previous AlphaGo. AlphaZero shocked the world: it learned not only the game of Go but also chess and shogi at superhuman levels. The thing that was unique about AlphaZero was that it had learned these games independently, and none

How AI Came to Be

of these games required human demonstrations or built-in strategies to play.

AlphaZero began with only the game's regulations. It would then reinforce itself in millions of matches. It gradually grew as it devised creative and startling strategies that human specialists had never considered. Within hours, AlphaZero could beat the best computer programs at chess and Go, programs that had taken decades to build.

This accomplishment represented a new path of artificial intelligence. Previous systems, such as Deep Blue or even AlphaGo, relied on the input of human hands, such as databases of games or human moves. AlphaZero demonstrated that an AI was capable of being smarter than human beings, having been able to perform solely using self-play strategies and general learning.

AlphaZero was released with an implication that was not limited to board games. It implied that the same learning style would be used in real-life problems that require complex decision-making, namely, medicine, logistics, or climate modeling. Researchers viewed AlphaZero as an example of how AI could be used to solve problems without overreliance on human knowledge.

The incident also took hold of people's imaginations. Human beings knew that when a machine could learn

games of pure strategy that fast, it was possible that it could learn other matters in the future, too. AlphaZero turned out to be an icon of what AI can do that a human being cannot, and what can be done in a new way of thinking.

Google's BERT

Google presented BERT in 2018, which is a term that denotes Bidirectional Encoder Representations from Transformers. BERT became one of the most significant models in natural language processing as it contributed to computers comprehending human language with significantly higher precision. In earlier systems, most systems would read sentences in either a right-left direction or a left-right direction. BERT was able to reverse this by simultaneously reading the text in both directions, which enabled it to extract the entire meaning of the words depending on their context.

As an illustration, the term bank may refer to a financial institution or the riverbank. The previous systems could not distinguish between them easily, whereas BERT could comprehend the appropriate meaning by observing the words around. That resulted in more natural and accurate search engines, chatbots, and translation tools.

Google was prompt to apply BERT in its search engine, and this enhanced results for millions of users

How AI Came to Be

globally. People saw that Google could now comprehend complicated questions, particularly those in natural and conversational language. BERT assisted the system in determining the real meaning of the users rather than matching keywords.

The release of BERT also had an impact on research. A large number of other organizations and researchers began to create models on the transformer architecture of BERT. It served as a baseline to subsequent systems, bringing natural language understanding boundaries to even greater heights, such as GPT and other large-scale AI systems.

BERT is a milestone in AI since it brought language models closer to human true understanding. It demonstrated that context plays an important role in communication and provided the machines with a far more profound method of learning the meaning.

Deepfake Videos Emerge

Deepfake videos started to emerge in the late 2010s as an example of how artificial intelligence can alter and manipulate media in highly realistic scenarios. These were applications of deep learning and specifically, generative adversarial networks to replace faces or generate new speech and behaviors that appeared nearly natural. To a great number of viewers, it was hard to distinguish whether a video was real or fake.

Initially, deepfakes were shared predominantly online as a prank or experiment. Individuals would make videos in which celebrities appeared to say or do something they never said. The technology became better, and the quality improved with time. The findings were exciting as they demonstrated the power of AI to work creatively, yet they became frightening as to how the technology might be abused.

These demonstrated the possibility and the threat of AI. On the bright side, they might aid in filmmaking, education, and entertainment. Directors would be able to revive the past or lip-sync movies in real life. On the downside, deepfakes may disseminate false information, damage reputations, and even affect political factors if applied to fabricate fake news.

Another reason governments, researchers, and companies sought solutions was deepfakes' emergence. New technologies were invented to identify fake videos and safeguard digital media. Social media also came up with policies to regulate the distribution of malicious deepfakes.

Deepfake videos became a significant milestone in the history of AI. They demonstrated that artificial intelligence not only serves the purpose of technical issues but could also be a driver to influence cultural, truth, and trust in society.

How AI Came to Be

Scientific Research Reports: AI Wins Over Radiologists

A number of research studies demonstrated in the late 2010s and the early 2020s that AI systems could be more effective in reading medical images than radiologists. The accuracy of deep learning models in the detection of diseases like cancer, pneumonia, or brain disorders was compared with that of human doctors, as per these studies. In most instances, the AI was as good or even better than the performance of seasoned experts.

An example that was well known was in the field of breast cancer detection. Scholars trained a deep learning architecture with thousands of mammograms. The AI learned to find really tiny patterns that people overlooked. The system, when tested, caught more cases accurately and raised fewer false alarms than radiologists; findings were seen in lung scans, skin photographs, and brain MRIs.

The success of AI was due to its capability to crunch large volumes of data in a short time and in a consistent manner. Radiologists may become fatigued or distracted, but AI does not. Thousands of images could be scanned in relatively little time, and the systems and patterns that would not be visible to the human eye could be identified.

The medical world was excited by these studies. They demonstrated that AI could become a powerful doctor

companion and make lives healthier and more effective in terms of life-saving. Hospitals started investigating the ways of using AI tools to assist radiologists rather than to substitute them, uniting machine accuracy with human intuition.

Important discussions were also held about the results. People were asking a lot of questions concerning trust, responsibility, and ethics. Should an AI make a mistake in diagnosis, the question would be who is to hold responsibility. Others question whether patients will be comfortable knowing that there is a machine that contributes to their treatment.

The win of AI over radiologists in the studies marked significant milestones in the field of application of artificial intelligence. It demonstrated that AI could go beyond games and internet tools and penetrate the sphere where decisions within AI directly impact human health and survival.

OpenAI Is Founded

In 2015, a group of established technology leaders, including Elon Musk and Sam Altman, co-founded OpenAI. The primary premise of OpenAI was ensuring that artificial intelligence would be created in such a manner that it would be both safe and useful to everyone. A lot of companies were secretly developing AI at the time, and it was feared that the technology would be

How AI Came to Be

monopolized by a select few. OpenAI was determined to do things differently by collaborating in an open working environment and publishing research to the community.

The organization was founded as a non-profit, thus demonstrating its attention to long-term objectives rather than just profit-making. The founders were convinced that AI could transform the world but were aware of the risks. It aimed to create artificial general intelligence, or AGI, which could do anything a human can, and ensure it would serve the whole populace and not a few people.

OpenAI started with the release of research papers and publishing tools, as well as developing developer and researcher platforms. This included work on reinforcement learning, robotics, and natural language processing. They opened their research, providing students, universities, and smaller laboratories with an opportunity to exploit the state-of-the-art AI without having to possess enormous resources.

The establishment of OpenAI soon became popular. It was perceived as a radical move to even the playing field between large tech corporations and ensure the development of AI is more transparent. With time, OpenAI became one of the most popular organizations engaged in AI and developed such systems as GPT-2, GPT-3, and ChatGPT.

The development of OpenAI is one of the most important AI events. It demonstrated that creating powerful AI is not merely the issue of technological advancement but also the question of responsibility, morals, and justice in terms of the distribution and utilization of the technology.

Therefore, the history of artificial intelligence since the beginning of the 2000s illustrates the pace at which the sphere is evolving. During this time, AI shifted the research field to what is presently experienced in life. Kinect, Watson, and DeepFace systems demonstrated that AI could be almost as good in recognition and decision-making as a human being. Advances such as AlphaGo, AlphaZero, and BERT demonstrated that machines can learn in a manner that is even unexpected by specialists. The work of OpenAI and such tools as TensorFlow opened and made AI more open, enabling people all over the world to build and experiment with it.

Simultaneously, this chapter demonstrated that AI is not just about technical accomplishments. Deepfake videos and high-tech medical systems demonstrated that AI is a source of opportunities and a threat. Sophia and other human-like systems began arguing over what place AI should have in society. These are some of the moments when it is emphasized that progress is always followed by doubts about responsibility, ethics, and trust.

How AI Came to Be

The narrative of the time is the incorporation of AI in human lives in direct and visible ways. In games, medicine, language, and vision, AI demonstrated its ability to transform numerous areas. Towards the end of the chapter, it is apparent that artificial intelligence is not in the background anymore. It is the heart of new technology and will continue to define the future in a manner that we have yet to realize.

Chapter 4: Engager and the Human-Machine Conversation

In 2009, I was watching a video on a computer screen when a question changed the course of my life.

The video played, as videos do. It spoke at me. It could inform me, entertain me, even move me. But it could not hear me. It was a monologue, the same monologue it delivered to everyone, every time. And as I sat there watching, the question arrived with a force I can still feel today: What if you could talk back? What if a human being could have a real, dynamic, back-and-forth conversation with a computer over a network, from anywhere in the world, communicating in whatever form felt most natural, whether that was text, voice, or video?

To understand why that question hit me so hard, you have to understand where I was standing when I asked it. I had moved to Silicon Valley for college and launched a video production company from my dorm room at Santa Clara University. By 2009, that company, Millennial Productions, had produced web video for organizations ranging from Fujitsu to Stanford University to Rolls-Royce. I had spent years behind cameras and editing timelines. I knew exactly what video could do, and I had run into the wall of what it could not. Video was the most

How AI Came to Be

powerful communication medium ever invented, and it was completely deaf.

Remember what that moment in technology looked like. YouTube was barely four years old. The iPhone had been on the market for two years. Siri did not exist yet. Streaming was in its infancy. The idea that you could speak to a machine and have it respond intelligently belonged to science fiction. And yet the question would not leave me alone. If a machine could hold a genuine conversation with a human being, then one person could be in a thousand places at once. A founder could pitch every investor personally. A teacher could sit across from every student. A person could, in a very real sense, clone themselves.

That word, clone, became the center of everything. Not a copy of a body, but a copy of a presence: your face, your voice, your knowledge, your way of answering a question, available to anyone, at any time, over a network.

In 2010, I incorporated Engajer, Inc. to chase that vision.

No one builds something like this alone, and I did not. In those early days a remarkable group of people came together around the idea. The core team: Daniel Chen, Abe Melloh, and Susie Jerez. And behind us stood the early investors who put their faith and their resources behind a vision the market did not yet understand: Massy

Mehdipour, Mike Downey, Robert Michael Fried, Bob (“Raz”) Zeichek, and Bert Davey. Several of them, Robert, Raz, Bert, and Mike, became mentors to me as well, and their guidance shaped far more than the company. Together we spent countless hours around whiteboards and kitchen tables trying to answer a question no one had seriously asked before: how do you clone a person? How do you capture what someone knows and how they communicate it, and turn that into something another human being can actually converse with?

Here is the thing people most often get wrong about Engager, so let me say it plainly. The core principle of Engager was never video. The core principle of Engager was conversation: a human being having a dynamic, back-and-forth, interactive conversation with a non-human computer system. The inputs and outputs could be anything, text, audio, video, or whatever medium the human preferred, however they wanted to communicate. What mattered was the fact that a conversation was taking place. Everything else we built was wrapped around that core: delivery tools to get the conversation to the end user, email tools to invite people into it, and analytics to study each conversation and extract something useful from it, what a person cared about, what they asked, where they leaned in. But at the center, always, was the conversation itself.

How AI Came to Be

In practice, our earliest implementation expressed that principle through interactive video, because video was the richest medium available at the time. Instead of a video that played from beginning to end, we built video that listened. Content was broken into segments, and the viewer's questions, choices, and responses determined what came next. The presentation was non-linear. It adapted. It followed the viewer down their own corridor of interest, the way a real conversation does.

We believed we had invented something genuinely new, and we protected it. We filed our first patent in 2010 and our second in 2014. Those filings were deliberately broad: they covered conversational interaction between a human and a computer system that included, but was not limited to, text, audio, video, and any multimedia that could be envisioned in the future. That breadth was the point. From the beginning, the vision included things the world did not yet have common names for: chatbots that could hold a real dialogue, audio agents you could converse with over an ordinary telephone call, and forms of conversational media that had not been invented yet. Years later, the United States Patent and Trademark Office issued U.S. Patent No. 10,956,965 to Engager, naming me and my colleagues as inventors, a public record of how early this vision took shape.

For a while, the world seemed to be catching on. Engager was recognized in the press. Our platform was

Michael C. Castellano

integrated with LinkedIn, which the Wall Street Journal covered. We ran a charitable campaign alongside Will Ferrell. In 2012, I was nominated by Ernst & Young as Entrepreneur of the Year. We raised funding. We had customers. We were, by most reasonable measures, a real company doing real things.

But here is the truth about being early that no one tells you: being early feels exactly like being wrong, right up until the day it doesn't.

The vision of a fully conversational virtual human, a true clone that could see, hear, speak, and respond with genuine intelligence, was ahead of the infrastructure the world could offer in 2010. Think about what that vision actually required. It required broadband fast enough to stream video in both directions without stuttering. It required cloud computing powerful enough to process conversations in real time. It required speech recognition that could actually understand what a person said, not just guess at it. It required natural language processing that could grasp meaning, not just keywords. It required generative models that could produce a lifelike human face and voice on demand. In 2010, not one of those components was ready.

So we did the only thing visionaries can do when the future refuses to hurry: we built what was possible, and we waited for the world to catch up. And piece by piece,

How AI Came to Be

year by year, it did. Deep learning broke through in 2012. Speech recognition crossed the threshold of usefulness. The transformer architecture arrived in 2017 and changed everything about how machines handle language. Large language models followed. Generative audio and video went from laboratory curiosities to consumer products. Every one of the missing components I had been waiting for since 2009 arrived, one after another, over the span of a decade and a half.

Each time one arrived, we folded it into the work. Engager evolved from its earliest implementations into the fullest expression of the founding principle: AI virtual people. A persona that does not just branch between recorded segments but genuinely converses. And that evolution led to the moment I described at the beginning of this book, when our virtual persona appeared on what a real person, Barbara, believed was an ordinary FaceTime video call with me. She talked with it the way you talk with a friend. Only afterward, replaying what had happened, did we understand it: she genuinely had not known. No lab, no rubric, no panel of judges. Just a real person, a real conversation, and a machine that passed for human. Alan Turing imagined that moment in 1950. I had been chasing it since 2009. When it finally happened, it did not feel like magic. It felt like arrival, the end of a sixteen-year road that started with a silent video on a computer screen.

Michael C. Castellano

I share this story not because Engajer is the biggest company in this book. It is not, and I know that. I share it because it captures something the famous stories often leave out: what this revolution looked like from the inside, at ground level, years before the world was paying attention. For every OpenAI there were hundreds of teams like ours, filing patents, raising money, explaining a future that investors could not yet see, and carrying a vision through the long years when the technology was not ready. Those teams are part of how AI came to be, just as surely as the household names are.

And while we were building, quietly, on our corner of the frontier, a research lab in San Francisco was preparing something that would detonate public awareness of artificial intelligence overnight, and change the story for all of us.

That is where we turn next.

Section 2

Chapter 5: ChatGPT

The Man Behind OpenAI

Sam Altman was never destined for mainstream fame. For decades, he was just another face known to Silicon Valley insiders but invisible to everyone else. His early career centered around running Y Combinator, a famous startup incubator. Before that, as a nineteen-year-old, he built a mobile location-sharing app called Loopt. Neither venture made him a household name outside of software engineering circles.

He spent his childhood in St. Louis. He eventually landed a spot at Stanford to study computer science. However, he walked away from his classes early because Loopt was growing fast. A college degree suddenly felt irrelevant compared to a live startup. It sounds like a cliché tech story: the brilliant college dropout who quits to build something massive. Yet while Altman looks the part, he deviates from the script. He never seemed driven by quick cash or a single hot app. His real fixation has always been the distant future.

The tech industry really started watching him during his time at Y Combinator. The accelerator funds two groups of startups every year. It hands over modest seed money, introduces founders to wealthy investors, and gambles that a tiny percentage of the companies will strike

How AI Came to Be

gold. With Altman in charge, the organization grew massively in size and influence. He was remarkably good at it. He possessed a rare clarity of thought, gave blunt, actionable advice, and held deep convictions about where humanity was heading.

Away from the startup grind, Altman harbored a deeper obsession. He constantly thought about artificial intelligence. He was not looking at basic software designed for simple, repetitive chores. He was anticipating the arrival of a broad, adaptable intelligence that would alter human history. Because he viewed this shift as inevitable, he wanted to ensure that the people building it actually cared about safety.

That exact motivation sparked the creation of OpenAI.

A small group launched the project back in 2015. The founding donors included Altman, Elon Musk, Reid Hoffman, and Peter Thiel. They threw immense wealth behind a concept that sounded bizarre and overly optimistic at the time. The goal was to build artificial general intelligence safely and transparently. Crucially, they wanted to keep it out of the hands of a single greedy corporation.

The decision to start as a nonprofit turned heads immediately. Top-tier AI scientists demand immense salaries, making them nearly impossible to hire without

offering stock options. OpenAI quickly realized the charity model could not survive. In 2019, the group pivoted to a capped-profit model. This hybrid setup allowed them to retain corporate cash while limiting investors' financial upside. The original mission statement stayed on the website, but the entity now needed to make money to survive.

Elon Musk walked away from the board right before that shift, citing conflicts with Tesla's automated driving software. He later turned into a fierce critic of the company, claiming it betrayed its original ethics. Altman defended the pivot. The messy, ongoing feud between the two billionaires became one of the most entertaining side stories of the internet age.

Throughout those middle years, the team worked in relative isolation. OpenAI published technical research, trained larger language models, and slowly mastered the art of machine-generated text. When GPT-2 arrived in 2019, it wrote so well that the creators panicked, debating whether publishing it was too dangerous. That hesitation served as a preview for the massive ethics debates happening right now. They ended up releasing it in phases, the world survived, and the lab kept working.

The next version, GPT-3, stunned computer scientists. It drafted clean essays, answered trivia, and wrote functional code with unprecedented skill. Software

How AI Came to Be

developers bought access to the backend code to build niche applications. Still, the broader public remained completely unaware of the shift.

Then, ChatGPT landed.

Five Days. One Million Users.

OpenAI launched ChatGPT on November 30, 2022. There was no big event. No countdown. The team released it as a research preview, meaning it was technically a test product they were putting out to see how people used it and what problems came up. Altman later said he expected maybe a million users total. He got that in five days.

By January 2023, ChatGPT had 100 million active users. That made it the fastest consumer application to reach that number in history. Instagram took two and a half years to get there. TikTok took nine months. ChatGPT did it in two.

The reason it spread so fast is simple: people started showing it to each other. There was this particular experience of talking to ChatGPT for the first time that was hard to describe without just showing someone. You would ask it something. It would answer. You would push back or ask a follow-up. It would adjust. It felt, in a way that no previous software had quite felt, like something that was paying attention to you.

That feeling is worth examining because it shaped everything that came after. ChatGPT was not intelligent in the way a person is intelligent. It did not have experiences or feelings or goals. What it had was an extraordinary ability to produce text that fit the context of a conversation, trained on a vast amount of human writing, in a way that mimicked understanding well enough to be genuinely useful and, sometimes, genuinely startling.

People used it for everything. A lawyer would ask it to summarize a case. A parent would ask it to explain a science concept to a ten-year-old. A programmer would paste in code that was not working and ask what was wrong. A student would ask it to help outline an essay. A marketing manager would ask it to draft a product description. Someone would ask it, just out of curiosity, to write a horror story set in a laundromat, and it would do that too.

The range of what it could do was the thing that kept surprising people. Earlier AI tools had been narrow. You had one that could translate languages. One that could recognize photos. One that could recommend songs. ChatGPT would do virtually anything involving language, and it would do it well enough to be useful even if not always perfectly accurate.

How AI Came to Be

That last part turned out to matter a great deal. ChatGPT was not reliable in the way a calculator is reliable. It could produce wrong information with the same confidence as it produced right information. Researchers called this hallucination, and it became the defining limitation of the technology in its early phase. You could not fully trust what it told you without checking. That meant the people who got the most value from it tended to be those who already had enough knowledge to recognize when it was wrong.

Still, for hundreds of millions of people, the tradeoff was worth it. A tool that was right most of the time, that could handle most of the task and leave the verification to you, was still enormously useful. It lowered the starting cost of a lot of work. It gave people a place to begin.

Everywhere, All at Once

The spread of ChatGPT across different parts of life happened faster than anyone had anticipated, and it happened differently in different places.

Schools caught the brunt of it first. Students discovered almost immediately that they could have ChatGPT write their essays, answer their exam questions, and complete assignments that were supposed to demonstrate their own thinking. Teachers discovered this almost as quickly. What followed was a period of real confusion about what to do. Turnitin, the plagiarism

detection software used by universities around the world, scrambled to build AI detection features. Some schools banned AI outright. Others tried to redesign assignments to be AI-resistant, requiring things like in-person writing, oral defenses of submitted work, and topics tied to personal experience that a model could not plausibly fabricate.

A smaller number of educators took a different approach entirely. They started teaching students how to use AI well, on the grounds that this was a skill that would matter in almost every future workplace and that pretending it did not exist was not a strategy. That debate has not been resolved. It probably will not resolve cleanly, because the technology keeps changing and the right response probably varies by subject, level, and institution.

In offices, adoption was faster and quieter. Knowledge workers started using ChatGPT on their own, without waiting for their companies to have a policy about it. People in marketing, consulting, law, finance, and communications found that it could handle first drafts, summaries, research, and the miscellaneous writing tasks that fill up a workday. There was a period of months where a lot of professionals were using AI tools that their employers had not officially approved, partly because the tools were useful and partly because the employers had not yet figured out what to think.

How AI Came to Be

Then companies started forming policies, which varied wildly. Some banned AI use with company data out of concern about privacy and security. Others embraced it enthusiastically and started building internal tools on top of OpenAI's API. Goldman Sachs built AI tools for its analysts. Consulting firms told clients they were integrating AI into their delivery. Law firms started using AI for document review and research. The pace varied by industry and by how risk-averse the leadership was.

Healthcare moved more slowly, which made sense. The stakes for wrong information are higher when the topic is a person's health. But the interest was real. Clinicians found it useful for documentation, for drafting patient communications, for looking up information quickly. Researchers started exploring how language models could assist in literature review and hypothesis generation. There were also genuine concerns about patients using AI to self-diagnose or to evaluate treatment options without adequate medical judgment. Those concerns have not gone away.

Customer service changed significantly and fairly quickly. The previous generation of chatbots, the ones deployed by banks, airlines, and telecom companies throughout the 2010s, were almost universally awful. They understood only very specific phrases, failed constantly, and frustrated customers so reliably that pressing zero to get a human being became a universal

reflex. The new generation of AI-powered customer service tools was genuinely better. They could understand natural language, handle a wider range of queries, and resolve a large portion of straightforward issues without escalation. That changed the economics of customer service considerably and raised real questions about what happens to the people whose jobs were answering phones.

In parts of the world where access to expertise is scarce, ChatGPT has become hard to categorize as just a productivity tool. In countries where there are not enough doctors, lawyers, or financial advisors per capita to serve the population, a system that can provide informed, articulate answers in multiple languages carries real humanitarian weight. That does not mean it is a substitute for professional expertise or that the risks of misinformation do not apply. But the potential to extend access to knowledge in places where it is genuinely limited is not nothing.

The Updates Came Fast

The ChatGPT that existed in late 2022 was impressive, but it had limits that became obvious quickly once people started using it seriously. It did not know anything that had happened after its training cutoff. It could not browse the internet. It sometimes stated false information as confidently as it stated true information. It forgot

How AI Came to Be

everything between conversations. And it could only handle text.

OpenAI started fixing these things, one after another, at a pace that kept the tech industry slightly breathless.

GPT-4 came in March 2023, and it was a meaningful step up. Better reasoning, better accuracy, better handling of complex tasks. It also introduced the ability to process images alongside text. You could take a photo of a handwritten math problem and ask ChatGPT to solve it. You could upload a chart and ask what it showed. This multimodal capability opened up entirely new use cases.

Plugins arrived next. These let ChatGPT connect to external services and tools, so it could search the web, run code, retrieve real-time data, and interact with third-party applications. Web browsing in particular addressed one of the most frequent complaints: that the model's knowledge stopped at a fixed date and therefore could not answer questions about anything recent.

Memory came later. This was the ability for ChatGPT to remember things about you across conversations. Your name, your job, your preferences, and things you had mentioned previously. Instead of having to re-explain your situation every single time you opened a new chat, the model could build up a picture of you over time and use it to give more relevant answers. This shifted the experience subtly but significantly. It felt less like querying

a search engine and more like having an ongoing relationship with a tool that actually knew you.

Voice was perhaps the most striking addition. OpenAI released voice capabilities that let users speak to ChatGPT and hear it speak back, in genuinely natural-sounding voices. Not the flat, robotic synthesized voice of older text-to-speech systems. Something that could vary its tone, pause naturally, and respond to the flow of a real conversation. The demo OpenAI put out showed the model helping someone prepare for a job interview, talking them through their nerves, adjusting its coaching in real time. People watching it had reactions ranging from impressed to unsettled, sometimes both.

GPT-4o, announced in 2024, was positioned as a genuinely unified model rather than separate systems bolted together. The ‘o’ stood for omni, and the idea was that the model could handle text, audio, and vision as a single integrated system rather than routing different inputs through different pipelines. It responded faster. It picked up on emotional cues in the voice. It could see through a camera and comment on what was in front of you in real time. The gap between what AI could do and what you might have imagined it could do in a science fiction film had become, depending on the scene, fairly small.

How AI Came to Be

Then there were the custom GPTs. OpenAI built a system that let anyone, without any coding knowledge, create a specialized version of ChatGPT tuned for a specific purpose. A travel agent could build one that knew every detail of their booking system. A nonprofit could build one focused on its cause area. A teacher could build one that helped students with a specific subject while refusing to just do their homework for them. This turned ChatGPT from a single product into a platform, and it meant the number of specialized AI tools in the world multiplied rapidly.

What It Did to the Rest of the Industry

If you want to understand what ChatGPT's launch meant for the technology industry, the most useful thing to look at is not OpenAI's own trajectory but everyone else's reaction.

Google had spent years doing foundational AI research. The transformer architecture that makes modern language models possible was invented at Google and published in a 2017 paper with the now-famous title 'Attention Is All You Need.' Google had its own large language models. It had built AI into its products quietly for years, improving Search, improving translation, and improving spam filtering. It was not caught off guard by the existence of capable AI.

What caught it off guard was the product. ChatGPT demonstrated that people would engage with an AI interface directly, in an open-ended way, at an enormous scale. Google had not shipped that product, even though it arguably had the technology to do so. There are reasons for that, including concerns about accuracy in a product tied to the company's core search business. But whatever the reasons, the result was that Google found itself in a defensive position for the first time in a long time.

The company reportedly declared an internal code red. It accelerated work on what became Bard, its own AI assistant, and later rebranded and substantially upgraded it as Gemini. It started integrating AI into Search in ways that changed how results were displayed. It rushed to show the world that it was not behind, a rush that produced at least one embarrassing public demo where the AI gave a factually incorrect answer, and the company's stock dropped noticeably on the news.

Microsoft moved in a completely different direction. The company already had a significant investment in OpenAI, having put in a billion dollars in 2019 and substantially more afterward. When ChatGPT launched, Microsoft was in a position to move quickly. It integrated OpenAI's technology into Bing, framing the new search experience as a genuine challenge to Google's dominance. It built Copilot features into Word, Excel, PowerPoint, Teams, and its other productivity software. It gave

How AI Came to Be

GitHub, which it owned, an AI pair programmer called Copilot that could write code alongside developers. Microsoft's stock climbed steadily through 2023 and 2024 as investors bet that the company had positioned itself well for the AI era.

Meta took yet another approach. It decided to release its large language models as open source, under the name Llama. The logic was partly philosophical and partly competitive: if capable AI was going to exist, Meta believed it should be open and widely distributed rather than controlled by a small number of API-based companies. The Llama models were good enough that developers could run them locally, fine-tune them for specific purposes, and build products without paying per-query fees to OpenAI or anyone else. This seeded an enormous amount of experimentation and gave rise to a whole ecosystem of open-source AI tools.

Hundreds of startups were formed. Many were founded by researchers who had been at Google, Meta, OpenAI, and other labs and now saw a commercial opportunity. Anthropic, founded in 2021 by former OpenAI researchers, developed Claude. Mistral, founded by former DeepMind and Meta researchers in France, built its own competitive models. Companies that had nothing to do with AI started building AI features into their products. The infrastructure companies that provided the computing power for all of this, particularly

Nvidia, which makes the chips most AI training runs on, saw demand and valuations climb dramatically.

Governments started paying attention in a way they had not before. The European Union's AI Act, which had been in development for some time, moved toward finalization. It imposed different levels of requirements depending on the risk level of an AI system. The United States government issued executive orders and started various policy processes. The United Kingdom hosted an AI Safety Summit that brought together representatives from major AI companies and governments to discuss risks and coordination. China announced its own regulatory framework for generative AI. The conversations that had been happening in research ethics seminars were now happening in legislatures and at the ministerial level.

Researchers who had spent years working on AI safety, sometimes at the margins of a field that was more focused on capability than on risk, found their work suddenly in demand. Questions about how to ensure that AI systems did what their developers intended, how to prevent misuse at scale, how to maintain meaningful human control as these systems became more capable, those questions went from theoretical to urgent. Funding for safety research increased. The number of people working on it increased. Whether the pace was adequate

How AI Came to Be

to the rate at which capabilities were advancing is a question that serious people disagree about.

What to Make of All This

Something worth saying plainly: ChatGPT is not magic. It does not think the way you think. It does not know things the way you know things. When it produces an answer that seems intelligent, what it has actually done is generate text that fits the context, based on patterns in a colossal amount of human writing. It is extraordinarily good at that. Good enough to be genuinely useful in a wide range of situations. But it is not the same as understanding.

That distinction matters in practice. It means you cannot trust everything ChatGPT says without checking it. It means the tool works best for people who already know enough to evaluate what it produces. It means there are categories of task, ones that require actual understanding of a specific situation, or professional accountability, or ethical judgment, where you should not just hand the job to an AI and walk away.

None of that reduces what actually happened. What happened is that in November 2022, a product shipped that changed how a large portion of the world interacts with computers and with information. It crossed a threshold that made things possible that had not been possible before, or at least not accessible to ordinary

people. And once a threshold like that is crossed, the world does not go back to the other side of it.

The debates that ChatGPT set off, about education, employment, accuracy, safety, power, and what intelligence even is, are not going to be resolved quickly. They probably should not be. They are real questions, and the answers will shape how this technology develops and who it benefits.

But this much seems clear: the story of AI did not start with ChatGPT, and it will not end with it either. What ChatGPT did was bring that story out of research labs and into everyday life, where everyone now has to figure out what they think about it.

That is a harder position to be in than ignorance. It is also a more honest one.

Chapter 6: Elon Musk's AI Work

Some people just work in tech. Others seem to drag the entire industry behind them wherever they go. Elon Musk is definitely in that second camp. Love him or hate him, it is practically impossible to look at the last twenty years of technological shifts without running straight into something he built, financed, or sparked into motion.

His fingerprints are all over rockets, electric vehicles, underground tunnels, satellite grids, social media, and now, with rapidly growing influence, the highest-stakes projects in the history of artificial intelligence.

This chapter digs into that last part. Specifically, how one of modern business's most restless minds ended up playing such a central, messy role in the evolution of AI. We will trace his roots, the core vision driving him, and the literal ventures he sparked or built from scratch, including OpenAI, xAI, Grok, Neuralink, and a handful of other bets. Together, they paint a picture of a man obsessed with a single, unshakeable belief. He truly thinks the future of intelligence is the definitive issue of our era.

But to make sense of Musk's complicated relationship with AI, you have to look past the internet caricature and see the person behind it.

A Boy Who Read Everything

Elon Reeve Musk arrived on June 28, 1971, in Pretoria, South Africa. He was the oldest of three kids. His dad, Errol, made a living as an engineer and property developer. His mom, Maye, was a model and dietitian who built a serious career in both fields. On paper, it was a comfortable upbringing. In reality, it was not particularly warm. By all accounts, especially his own, Musk's early years were shaped far more by books and computers than by close friendships.

He was different. His teachers noticed it early on. He read constantly and at a ridiculous speed, tearing through encyclopedias and sci-fi epics the way other kids flipped through comic books. By age ten, he had taught himself to program on a Commodore VIC-20 he got as a gift. At twelve, he coded a basic space game called *Blastar* and sold the source code to a PC magazine for around \$500. It was his very first paycheck for something he actually built.

School, though, was brutal. He faced severe bullying, an experience he later described as deeply traumatic. Reading social cues did not come naturally to him, and he just was not built to blend into a crowd. Instead, he possessed an intense, almost obsessive ability to hyper-focus on problems that caught his eye. He also had a rare capacity to sit alone with complex concepts for hours on end.

How AI Came to Be

The books he obsessed over were almost always about the future. Isaac Asimov's Foundation series, centered on a scientist trying to save civilization from an impending dark age by archiving all human knowledge, hit him hard. So did Douglas Adams' The Hitchhiker's Guide to the Galaxy, with its witty, slightly terrifying reminder of how tiny human problems are in a massive universe. These were not just casual reads. They hardwired a specific worldview he would carry for the rest of his life. He realized that human civilization is fragile, that extinction is a very real threat, and that someone actually needs to step up and fix it.

At seventeen, he left South Africa behind. His first stop was Canada, where his mother had relatives, followed quickly by the US. He spent time at Queen's University in Ontario before transferring to the University of Pennsylvania, splitting his focus between economics and physics. That dual approach mattered. Even then, he was building a mental toolkit to understand how massive systems operate, whether governed by market forces or the laws of physics, and what it takes to disrupt them.

In 1995, he got into a PhD program in energy physics at Stanford. He dropped out after exactly two days. The internet era had just kicked off, and he saw a window he was not about to miss.

The Pattern Behind the Companies

What followed is public record, but looking at the timeline reveals a distinct pattern. He and his brother Kimbal launched Zip2, an online city directory for newspapers, selling it in 1999 for roughly \$300 million. Musk grabbed his payout and immediately wagered it on a new startup called X.com, which eventually became PayPal after a merger. When eBay bought PayPal in 2002 for \$1.5 billion, Musk walked away with about \$180 million.

Normal people retire there. They buy an island and a yacht. Musk did the exact opposite, sinking virtually every cent into two wildly ambitious companies launched that same year. Those were SpaceX, aimed at reusable rocketry and the eventual colonization of Mars, and Tesla, an electric car gamble. Back then, smart money said both would fail. Early SpaceX rockets kept blowing up on the launchpad. Tesla came terrifyingly close to bankruptcy during the 2008 crash. Musk later admitted he figured both companies had a measly 10% shot at survival during their darkest days.

Yet, they made it. They completely disrupted their industries. And the takeaway here is not just that Musk got lucky or is a genius, though both might be true. The real takeaway is how his brain works. He pinpoints what he views as the absolute greatest threats to long-term human survival, figures out a lever he can pull to fix them,

How AI Came to Be

and then willingly embraces catastrophic financial risk and operational chaos to make it happen.

In this framework of his, artificial intelligence is not just another hot tech sector. It is a potential extinction-level event if humanity handles it carelessly. That is not a side thought to his AI work. It is the entire foundation of it.

OpenAI: The Project He Helped Start, Then Left

To map out where Musk stands on AI today, you have to rewind to 2015. That was the year he helped breathe life into OpenAI.

The founding crew, which included Sam Altman, Reid Hoffman, Peter Thiel, and Jessica Livingston, among others, came together largely out of a shared anxiety. They were terrified that cutting-edge AI development was being monopolized by a tiny handful of tech giants, specifically Google, which had snatched up DeepMind in 2014. Musk was incredibly vocal about this. For years, he had been publicly warning that advanced AI could pose a greater threat than nuclear weapons. Unchecked, centralized AI control looked to him like one of the biggest risks in human history.

So, OpenAI launched as a nonprofit. It was explicitly committed to safety, open-source transparency, and making sure the benefits did not just go to a few shareholders. The founders pledged over a billion dollars.

Musk was not just a massive financial backer. He was the primary recruiter hunting for top-tier talent and a major force shaping the lab's early culture. He genuinely believed an open, safety-first counterweight to corporate AI labs was the only logical defense against the dangers ahead.

But the partnership fractured. In 2018, Musk walked away from the board, pointing to potential conflicts of interest with Tesla's own push into autonomous driving software. That was the polite, public narrative. The messy reality, leaked by various insiders, involved serious power struggles over the lab's direction. Musk had reportedly pitched taking the reins as CEO himself, a move the rest of the board blocked. The split was final.

Over time, that breakup turned incredibly bitter. When OpenAI pivoted to a capped-profit structure in 2019 to secure the massive capital needed to train large language models, Musk viewed it as a total sellout of their original promise. He ended up suing OpenAI in 2024. He claimed the company violated the very terms of his original donations and abandoned its foundational vow to build open, public-good technology. While OpenAI's leadership fiercely contested the lawsuit, the legal battle laid bare the massive philosophical chasm that had opened up between Musk and the project he helped birth.

Naturally, that bitter exit set the stage for exactly what he did next.

xAI: Building the Alternative

In July 2023, Elon Musk announced the formation of a new company called xAI. The announcement was low-key by his standards. A post on X, a small initial team, a mission statement that emphasized understanding the universe. The stated goal was to build AI that was maximally truth-seeking, curious about reality, and resistant to the kind of ideological constraints that Musk believed had begun to shape the outputs of competitors like ChatGPT.

The company was assembled quickly from people who had worked at some of the most recognized AI research institutions in the world. Several came from Google DeepMind. Others came from OpenAI. A few had worked on foundational research at universities. Igor Babuschkin, who had been a senior researcher at DeepMind, joined as a core technical member. Dan Hendrycks, a researcher well known for his work on AI safety, served as an advisor. The team was small but unusually experienced.

xAI operated with a particular transparency about its ambitions that made it easy to understand what Musk was trying to build. He was not interested in making a slightly different version of something that already existed. He believed that the most capable AI systems available were being trained and filtered in ways that made them less

honest, more prone to telling users what they wanted to hear, and shaped by assumptions about values that were not universal. His goal was an AI that would engage with difficult or uncomfortable questions directly, reason from evidence rather than from editorial judgment, and behave more like a curious, skeptical scientist than a careful institutional communicator.

Whether one agrees with that diagnosis of existing AI or not, it represented a coherent and distinctive set of priorities. And it meant that xAI was not just a commercial competitor to OpenAI or Google DeepMind. It was an argument, built in software, about what AI should fundamentally be.

The company built and trained its own large language model from scratch. The data pipeline, the architecture decisions, the training infrastructure: all of it was developed internally. Musk had the significant advantage that X, the social media platform he had purchased in 2022, gave xAI access to a massive and distinctive dataset of real human conversation, debate, humor, and argument. No other AI company had a resource quite like it.

By late 2023, the first version of xAI's model was ready to deploy. They called it Grok.

Grok: What the Model Was Built to Do

How AI Came to Be

Grok went live in November 2023. It rolled out as a perk for X Premium subscribers. Musk took the name from Robert Heinlein's 1961 sci-fi classic, *Stranger in a Strange Land*. In the book, it means a type of understanding so deep that the observer becomes part of the observed. Sci-fi geeks loved it. It signaled that xAI wanted something beyond basic text matching. They were chasing genuine comprehension.

The initial launch had a specific hook. It was an assistant with an attitude. That was completely intentional. The team built Grok to tackle messy topics that other AI models completely dodged. It talked politics, cracked edgy jokes, and avoided that vague, overly polite neutrality that made other chatbots feel so corporate. Musk put it bluntly, as usual. He called it an AI that actually read the internet, including the messy parts respectable folk ignore.

People either loved it or hated it. Fans liked having a bot that felt unmanaged and willing to face real-world human chaos. Critics pushed back hard. They argued that those safety filters existed for a reason, and optimizing a machine for edginess is not the same thing as optimizing it for truth. It sparked a massive, very real debate about what public AI should actually be allowed to say.

The tech, though, was hard to argue with. In March 2024, xAI open-sourced Grok-1. It was a monster with 314 billion parameters. At the time, it was one of the

largest open models on Earth. Dropping those weights for free was a massive statement. Anyone could download it, tweak it, and build on it without paying API fees. For a company under a year old, it was a wildly aggressive opening move.

Then came the rapid-fire upgrades. Grok-1.5 dropped in April 2024 with better logic and a bigger context window. By August 2024, Grok-2 added vision capabilities, meaning it could analyze photos and charts alongside text. Later that year, in December, xAI added Aurora, an in-house image generator built right into X.

Grok-3 represented the massive leap. To train it, xAI built a giant supercomputer called Colossus in Memphis, Tennessee. They packed it with over two hundred thousand Nvidia GPUs. Musk bragged that it was the most powerful training setup alive. Even if tech rivals disputed the crown, the sheer scale was undeniable. When Grok-3 debuted, its reasoning and math benchmarks put it right at the top of the food chain.

The real story here is the speed. Building elite AI requires insane capital and a mountain of chips. Most companies simply cannot compete. The elite club is tiny, but xAI forced its way inside the velvet rope in less than two years.

Neuralink: Intelligence From the Inside

How AI Came to Be

While Grok lives on screens, Neuralink is a completely different animal. It is much more radical. It is about merging code directly with the human body.

Musk started Neuralink in 2016 with a core team of neuroscientists. Max Hodak, a veteran in the brain-interface space, was part of the early crew. The goal sounded like sci-fi. They wanted an implantable chip to read neural firing speeds with enough resolution to control digital systems, and eventually, flash data right back into the gray matter.

People had been sticking wires in brains for decades. Systems like BrainGate have already proved that paralyzed patients could nudge robotic arms with their thoughts. But the old tech was clunky. Those older implants only monitored a few dozen neurons, required messy external wires, and degraded as scar tissue formed. The bandwidth was just too slow for anything complex.

Neuralink aimed to fix all of that. They built the N1 chip, a wireless implant with onboard processing. No wires poking through the skull. The threads are microscopic and flexible, designed to minimize brain inflammation. To actually insert them, they built a high-tech surgical robot. It stitches these threads into the cortex with a precision no human hand could ever match, avoiding tiny blood vessels along the way.

They also boosted the channel count. Old tech tracked hundreds of neurons at best. Neuralink targeted thousands. More channels mean more data, which opens up massive possibilities.

They spent years testing on animals. Pager, a macaque monkey, went viral in 2021 playing Pong entirely with his mind. The videos were incredible. But they also brought intense regulatory heat and sharp criticism from animal rights groups over lab conditions.

The big human breakthrough landed in January 2024. A twenty-nine-year-old quadriplegic named Noland Arbaugh received the N1 implant in the US. The surgery went smoothly. Within weeks, Arbaugh was playing video games and moving a mouse cursor using just his thoughts. He called it life-changing.

A second patient got the chip a few months later, followed by a third. The trials moved quickly under the FDA's Breakthrough Device pathway, which accelerates reviews for technologies with significant medical potential.

Fixing human bodies is the immediate goal. Giving movement and speech back to paralyzed folks or those suffering from ALS is an obvious win. If it works safely over long periods, it changes everything for millions of families.

How AI Came to Be

But Musk has never hidden his true endgame. He wants human-AI symbiosis. He thinks the real bottleneck isn't AI capability, but how fast humans can output data. Keyboards, thumbs, and voices are agonizingly slow. A direct brain link would speed things up by orders of magnitude. It fundamentally alters our relationship with digital intelligence.

Whether that future is possible, or even safe, remains a heavy debate among top scientists. But Neuralink has dragged the concept out of sci-fi labs and put it into real human patients. That part is no longer theoretical.

Optimus and the Physical Dimension of AI

Tesla, which Musk runs as CEO, has been developing a humanoid robot called Optimus that connects directly to the broader AI story. The project was announced at Tesla's AI Day in August 2021, when a performer in a Tesla-branded suit walked on stage as a placeholder for what the company intended to build. It was the kind of announcement that was easy to dismiss at the time.

By 2023, working prototypes were visible. By 2024, Tesla announced that Optimus robots were being used for tasks inside its own manufacturing facilities. The robot could sort objects, carry items, and perform simple repetitive operations. The target was not to replace the human workers already doing those tasks, at least not immediately, but to prove that the technology worked in

real conditions rather than just in demonstration environments.

Optimus is relevant to the AI story because of where its intelligence comes from. The system uses the same neural network architecture that Tesla developed for its Autopilot and Full Self-Driving software, which processes visual information from cameras to understand the environment in real time. Musk has argued that the breakthrough required to make a capable humanoid robot is the same one required to make a capable autonomous vehicle: a system that can see the world, understand what is in it, and act accordingly. If that argument is correct, then Tesla's years of work on autonomous driving give it an unusual head start in robotics.

The potential scale is hard to overstate. Musk has talked about eventually producing hundreds of millions of Optimus units. He has suggested that the robot could address labor shortages, take over physically demanding or dangerous work, and eventually function as a general-purpose assistant in homes and factories. These projections are speculative. But even a fraction of that outcome, a robust, affordable robot capable of performing useful physical tasks in human environments, would be transformative.

Starlink, Compute, and the Infrastructure of Intelligence

It is easy to look at Musk's AI work as a set of separate projects, each with its own name and purpose. But there is a thread connecting them that becomes clear when you step back far enough. Musk has been consistently concerned with the infrastructure that intelligence requires, not just the algorithms and the models, but the physical substrate on which they run and through which they operate.

Starlink, the satellite internet constellation operated by SpaceX, is part of this picture. It provides broadband internet access to parts of the world where ground-based infrastructure is inadequate or absent. This matters for AI because AI is only useful where it can be accessed. Systems that run in data centers in Northern Virginia are of limited help to a farmer in rural Kenya or a student in a remote Andean village if there is no reliable connection between that data center and that person. Starlink changes the access equation, and Musk has explicitly cited the connection: extending connectivity is part of extending the benefits of technology, including AI, to the entire human population.

The Colossus supercomputer cluster built by xAI in Memphis follows the same logic from a different direction. Access to AI also requires the computing

capacity to run it. The concentration of AI capability in a small number of facilities, owned by a small number of companies, creates dependencies and vulnerabilities. xAI's decision to build its own infrastructure was partly about capability and partly about independence. An AI company that relies on cloud computing from a competitor is in a structurally weak position. One that owns its own chips and its own data centers has more control over its own future.

The Contributions, Clearly Stated

It is worth being precise about what Elon Musk has actually contributed to the advancement of artificial intelligence, separate from the claims and the controversies that surround him.

His early co-founding and funding of OpenAI brought together a group of researchers who produced foundational work in reinforcement learning, language modeling, and AI safety. GPT-2, GPT-3, and ultimately ChatGPT, which changed how the world understands and relates to AI, would likely not exist in their current form without the organization that Musk helped create and fund. His departure from the board and his subsequent conflict with it do not erase his role in its origin.

The creation of xAI and the development of Grok added a competitive dimension to the large language model landscape. Competitive markets tend to accelerate

How AI Came to Be

development. The release of Grok-1 as an open-source model with 314 billion parameters was a direct contribution to the research commons, allowing organizations and researchers who could not afford to train their own models to work with one of the most capable available systems. The construction of Colossus demonstrated that computing infrastructure could be built at the required scale by organizations outside the small group of established cloud providers.

Neuralink has advanced the field of brain-computer interfaces more rapidly than it would have advanced without the combination of capital, engineering talent, and public attention that the company brought to the problem. Whether its long-term ambitions are realistic or not, its near-term contributions to the treatment of paralysis and motor disorders are real and ongoing. The first human implant represents a milestone in the history of neural technology.

Tesla's work on autonomous driving has advanced computer vision, sensor fusion, and real-time neural network inference, influencing the broader field. The Dojo supercomputer, developed internally at Tesla for training autonomous driving models, was one of the most powerful custom AI training systems in existence at the time of its development.

Across these efforts, Musk has also contributed something less tangible but arguably equally important: a set of arguments and provocations about what AI should be, how it should be governed, and what risks it poses. His warnings about existential risk from AI, which he began making publicly long before such warnings were fashionable, contributed to the broader conversation that eventually produced formal AI safety research programs, regulatory frameworks, and the institutional attention to alignment that now exists.

You can disagree with his specific conclusions and still acknowledge that his early and persistent alarm-raising changed the climate around AI development.

The Contradictions Worth Acknowledging

Honesty requires acknowledging the tensions in this picture. Musk is a person who co-founded an organization explicitly devoted to safe AI development and then left it, suing it years later and calling it a threat. He has publicly warned about the dangers of advanced AI while simultaneously racing to build it as quickly as possible. He has advocated for AI transparency while operating a social media platform that, at times, reduced the visibility of independent research on his own products. These contradictions are real, and they matter.

The most defensible reading of them is not that Musk is insincere but that he holds a specific theory of the

How AI Came to Be

problem. His position, stated explicitly in various interviews and public posts, is that advanced AI is coming regardless of whether any individual organization chooses to build it. Given that, he believes the relevant question is not whether to build but who should build it and with what values. He answers that it should be built by people who are willing to engage with difficult questions directly, who are not filtering their systems through a set of ideological assumptions he considers narrow, and who are competing with the organizations he thinks pose the greatest risk of concentrated AI power.

Whether that reasoning is correct is a genuinely open question. But it is a coherent position, and understanding it is important to understanding why Musk has done what he has done.

Where This Fits in the Larger Story

This book has moved from Alan Turing's theoretical foundations through Steve Jobs' vision of human-centered technology, through the institutional history of the 2000s and 2010s, and through the emergence of OpenAI and ChatGPT as the point at which AI entered everyday life. Elon Musk's work fits into this narrative in a particular way.

He is not primarily a researcher. He is not the person who developed the transformer architecture, the backpropagation algorithm, or the reinforcement learning

techniques that made modern AI possible. His contribution is different: he is someone who identified these developments as historically important very early, committed enormous resources to shaping their direction, and built organizations that changed the competitive and institutional landscape of the field.

The pattern is familiar from his other work. He did not invent the rocket or the electric car. He built organizations that pushed those technologies from niche or stalled to viable and eventually dominant. His AI work follows the same structure. He identified a lever, put his weight on it, and moved things that might otherwise have moved more slowly or in different directions.

That kind of contribution is harder to celebrate than a scientific breakthrough, and it is harder to criticize than outright fraud. It sits in the complicated middle ground where most consequential things happen: driven by real belief, shaped by real ego, producing real results that carry real risks.

What is clear is that the artificial intelligence landscape of 2024 and 2025 is different because Elon Musk has been part of it. OpenAI exists in part because of him. xAI exists because of him. Grok exists because of him. Neuralink is farther along because of him. Tesla's autonomous vehicle work has influenced the broader field. The public

How AI Came to Be

conversation about AI risk and governance has been shaped by his participation in it.

Whether those contributions have been net positive for humanity is a question that time will answer more fully than we can right now. The honest answer, at this moment, is that the evidence points in several directions at once, and anyone who tells you they know for certain is either more confident than the evidence warrants or more invested in a particular conclusion than in the truth.

What we can say is this: the story of artificial intelligence cannot be told accurately without including Elon Musk. He is too central to too many of the important moments. And the chapter of that story that he is currently writing, with xAI, with Grok, with Neuralink, with the ambitions he has stated plainly in public, is not yet finished.

The next chapters of this book will continue to trace how AI has evolved and where it is going. But before we move forward, it is worth sitting with what Musk's journey reveals about the kind of people and forces that drive technological change at this scale. They are rarely simple. They are rarely pure. They are usually driven by a mixture of genuine conviction, competitive instinct, and an unusually high tolerance for risk. They tend to make things happen faster than they would have happened otherwise, for better and sometimes for worse.

Michael C. Castellano

That has been true of nearly every figure in this story. It is true of Musk. And it will almost certainly be true of whoever emerges as the next central figure in the continuing history of artificial intelligence.

Chapter 7: DeepSeek

Not every seismic moment in the technology industry announces itself with a keynote, a countdown, or a marketing campaign. Sometimes it arrives quietly, almost without ceremony, in the form of a research paper and a model quietly pushed to the public. That is exactly how DeepSeek arrived. And by the time the world realized what had happened, the conversation around artificial intelligence had shifted in ways that very few people had predicted.

In January 2025, a Chinese AI lab that most people in Silicon Valley had barely registered posted a model called DeepSeek-R1. It was not a minor update or an incremental improvement on something existing. It was a genuinely capable reasoning model that performed at or near the level of the best systems from OpenAI and Google, in some benchmarks even surpassing them. That alone would have been notable. What made it genuinely shocking was the price tag attached to training it.

While American labs had been routinely spending hundreds of millions of dollars to develop frontier models, DeepSeek reported training V3, the foundation model beneath R1, for somewhere in the range of six million dollars. The number felt almost impossible. It got passed around in tech circles the way a good rumor does,

with people insisting it had to be wrong or misleading or missing something.

But the more carefully people looked, the more it held up. DeepSeek had not cut corners in some obvious way that would make the model brittle or unreliable. They had engineered around constraints that the rest of the field had treated as fixed. And they had done it while operating under significant chip restrictions, since U.S. export controls had limited China's access to Nvidia's most powerful hardware.

DeepSeek did not have the same access to H100 GPUs that OpenAI or Google or Anthropic had enjoyed. They worked with what they had. And what they produced, working with less, rattled the industry.

In the days after R1's release, Nvidia's stock dropped sharply. Investors had built a mental model in which frontier AI required enormous quantities of cutting-edge chips, and enormous quantities of cutting-edge chips meant enormous Nvidia revenues.

DeepSeek challenged that mental model at its foundation. If you could do more with less, the chip premium that the entire AI investment thesis rested on suddenly looked shakier. It was not just a technology story. It was a financial one, and a geopolitical one, and a philosophical one about what the future of this industry actually looks like.

How AI Came to Be

The app itself, released alongside the model, shot to the top of the App Store downloads almost immediately. People were testing it, probing it, comparing its answers to ChatGPT and Claude. Some found its censorship of politically sensitive topics, particularly anything touching on Chinese government policy or historical events like Tiananmen Square, to be a significant limitation. Others found it surprisingly capable and refreshingly direct on a wide range of topics. Either way, it was being used. It had broken through.

The Man Behind the Model

The person responsible for DeepSeek is not a career AI researcher who came up through the standard academic pipeline. Liang Wenfeng is, in many ways, an unusual figure even by the unconventional standards of the AI industry.

He was born in 1985 in Zhanjiang, a port city in the southern Chinese province of Guangdong. He studied electronic information engineering at Zhejiang University, one of China's most prestigious technical institutions, graduating in 2007. That background in engineering gave him a foundation in systems thinking, in how signals move through hardware and how hardware shapes what software can do. It is a background that would prove relevant in ways he could not have anticipated.

After university, Liang did not go straight into AI research. He went into quantitative finance. In 2015, along with several colleagues from his university years, he co-founded High-Flyer, a quantitative hedge fund based in Hangzhou. Quantitative trading is, at its core, a machine learning problem. You are trying to find patterns in enormous amounts of financial data, build models that can predict market movements, and execute trades faster and more accurately than human intuition allows. High-Flyer turned out to be very good at this.

The fund grew significantly and by the early 2020s had become one of China's most successful quant shops. The infrastructure and expertise that High-Flyer built, the ability to manage vast datasets, to train and iterate on models quickly, to think about computation as a resource to be optimized rather than a commodity to be purchased, fed directly into what would come later.

Liang was not just a financial entrepreneur who got interested in AI as a hobby. He had been running a sophisticated AI operation for years, just in a domain that most people outside the quantitative finance world do not think of as AI at all.

In 2023, Liang made the decision to take a significant portion of the capital High-Flyer had earned and redirect it toward building a general AI research lab. This was not a pivot forced on him by circumstance. It was a deliberate

How AI Came to Be

bet, made at a moment when the rest of the world was just beginning to process the implications of ChatGPT's release. His view, stated clearly in interviews he gave around that time, was that artificial general intelligence was coming, that it would change everything, and that China needed to be a serious participant in its development rather than a consumer of technology built elsewhere.

What is notable about Liang is what he did not do. He did not build DeepSeek as a product company, at least not primarily. He built it as a research organization. The models DeepSeek produced were open-sourced. The papers were published. The technical details were shared with a transparency that was unusual for a lab at the frontier and almost unheard of for a Chinese technology company operating in this space. His stated reason was that open research accelerates everyone's progress, including his own. There is likely a second reason, which is that sharing the work internationally made DeepSeek a participant in the global AI conversation rather than a closed shop that the rest of the world could simply ignore.

Liang is not a public figure in the way that Elon Musk or Sam Altman are public figures. He does not maintain a prominent social media presence. He does not give keynotes at major conferences or appear regularly in Western media. What is known of him comes largely from a small number of extended interviews, most of them in

Chinese, and from the technical work his team has published. By those accounts, he is intensely focused on the research problems themselves, methodical in his thinking, and genuinely driven by a belief that the current moment in AI history is one that requires serious attention and serious work.

Whether that characterization fully captures the man behind the lab is difficult to know from the outside. What is clear is that the organization he built moved faster and spent less money than almost anyone had thought possible, and that the gap between DeepSeek and the leading American labs, which many assumed would only widen over time, turned out to be far narrower than the conventional wisdom had suggested.

What DeepSeek Changed

The impact of DeepSeek on the AI industry is not easy to reduce to a single narrative, because it has operated on several levels simultaneously.

The most immediate impact was technical. DeepSeek's research papers, particularly those detailing the architecture and training approach for their models, introduced and popularized several techniques that have since been widely adopted or studied. Their approach to mixture-of-experts architecture, in which only a subset of a model's parameters are activated for any given input rather than the entire model, allowed them to build

How AI Came to Be

systems that were large in total capacity but computationally efficient during inference. Their work on reinforcement learning from model feedback, a method for improving reasoning capabilities without relying purely on human annotation, was detailed enough that other researchers could replicate and build on it. In an industry where proprietary methods are often closely guarded, this level of technical openness was genuinely rare.

The second impact was economic and strategic. DeepSeek forced a reckoning with assumptions about what AI development actually costs and what is required to do it at the frontier. The dominant narrative heading into 2025 was one of escalating investment: more chips, more data centers, more capital, a race that only a handful of the wealthiest companies and nations could realistically participate in. DeepSeek did not disprove that narrative entirely.

Building at this level still requires serious resources, and the six-million-dollar training figure, while impressive, does not account for all the prior research and infrastructure that made it possible. But it demonstrated convincingly that the resource requirements were not as fixed as the industry had believed. Efficiency was a lever. And pulling that lever hard enough could close gaps that money alone could not.

This mattered enormously for the geopolitical dimension of the AI race. The United States had implemented export controls on advanced chips specifically to limit China's ability to train frontier AI models. The logic was straightforward: if you restrict access to the most powerful hardware, you slow down the competition.

DeepSeek suggested that this logic, while not entirely wrong, was more complicated than its architects had assumed. Restriction creates pressure, and pressure creates innovation. When you cannot simply buy your way to capability, you are forced to find it through cleverness. DeepSeek found it.

The third impact was psychological and cultural, and it may prove to be the most lasting. For several years, the story of frontier AI had been predominantly an American story. The key labs, the key researchers, the key breakthroughs, the key investments: nearly all of them were located in a small geographic corridor running from San Francisco through the Bay Area.

There were serious AI efforts elsewhere, in Europe, in China, in Canada, and in the UK. But the frontier, where the most capable systems lived, felt like it belonged to a specific cluster of American institutions.

DeepSeek changed that perception. It demonstrated, in a way that was hard to argue with because the evidence

How AI Came to Be

was publicly available and independently verifiable, that the frontier was not as exclusively American as many had assumed. This has complicated implications. For American policymakers concerned about national security and technological advantage, it was a warning sign. For researchers and institutions outside the dominant cluster, it was an encouraging signal that the ceiling was not as fixed as it appeared. For the industry as a whole, it introduced a degree of humility that had perhaps been lacking.

Beyond the competitive dynamics, DeepSeek also influenced how the broader field thinks about efficiency as a research objective. For much of the modern AI era, the dominant approach to improving model performance had been to scale: more parameters, more data, more compute. This approach worked, and it worked spectacularly. But it worked partly because the labs pursuing it had the capital to make it work, and it created an implicit assumption that scaling was the path and that efficiency was a secondary concern.

DeepSeek demonstrated that optimizing for efficiency was not just a constraint to be endured when resources were limited but a research direction that could yield genuine advances. That reorientation has begun to show up in how other labs are thinking about their own work.

There is also a broader point about what DeepSeek represents in terms of the global distribution of AI talent. Liang Wenfeng built his team primarily from Chinese researchers, many of them trained at Chinese universities, working in China.

The conventional wisdom in Silicon Valley had long held that the best AI researchers in the world were concentrated in a handful of American institutions and that access to that talent was a structural advantage that could not easily be replicated elsewhere. DeepSeek challenged that assumption in the most direct way possible: by producing work that the talent could not ignore.

The Questions DeepSeek Raises

None of this means DeepSeek is without complications, limitations, or legitimate concerns.

The censorship embedded in its models is real and significant. On topics that the Chinese government considers sensitive, DeepSeek's models simply do not engage. They deflect, change the subject, or produce responses so vague as to be meaningless. For users outside China who want a general-purpose AI assistant, this is a practical limitation.

For researchers thinking about the values embedded in AI systems and how those values differ across cultures and political contexts, it is a deeply important case study.

How AI Came to Be

The question of what AI should and should not say, who gets to decide, and how those decisions should be made is not going away. DeepSeek makes it impossible to treat that question as purely abstract.

There are also questions about data, provenance, and the full picture of what went into training these models. The six-million-dollar figure that became famous refers specifically to the final training run for DeepSeek-V3, the foundation model on which R1 was built. It does not encompass the full cost of the research program that made that training run possible, including earlier models, failed experiments, and the infrastructure already in place from High-Flyer's quantitative work.

This is not a deception, it is a common way of reporting training costs in the industry, and American labs do the same thing. But it is worth keeping in mind when making comparisons.

There are legitimate security concerns as well. Governments and institutions that process sensitive information have been cautious about deploying a model developed in China, where legal requirements around data access and cooperation with government authorities differ significantly from those in democratic countries. These concerns are not specific to DeepSeek. They apply to any AI system whose provenance and governance structure

are not fully transparent. But they are worth stating plainly.

What DeepSeek does not fit neatly into is the simple narrative of threat or breakthrough or disruption that different audiences have tried to assign to it. It is, more honestly, a complicated development that forces the industry, policymakers, and anyone paying serious attention to update their mental models.

That is what genuinely significant moments in technology tend to do. They do not arrive with a tidy lesson attached. They arrive with a challenge to reckon with, and the reckoning takes time.

Where This Sits in the Larger Story

The history of artificial intelligence, as this book has traced it, is not a simple story of steady progress made by a single community in a single place. It is a story of ideas that crossed borders, of researchers who moved between institutions, of insights that emerged from unexpected directions and changed what everyone thought was possible.

DeepSeek fits into that history in a specific way. It is a reminder that the frontier of AI development is not the exclusive property of any single country, company, or community.

How AI Came to Be

It is a reminder that constraints, whether they come in the form of chip export controls or limited capital or working outside the dominant cluster of talent, do not necessarily prevent serious work. They shape it. They force different approaches. And sometimes those different approaches turn out to be genuinely better.

Liang Wenfeng's bet, made in 2023 with profits from a quantitative hedge fund, on the proposition that a Chinese research lab could contribute meaningfully to the global frontier of AI development, has paid off in a way that is hard to dispute. The questions it raises for the industry, for geopolitics, and for our understanding of how AI development actually works are still being processed. They will be for some time.

What is not in question is that DeepSeek arrived, that it mattered, and that the world of artificial intelligence after its arrival is different from the one that existed before. In a field where the next important development always seems to be just around the corner, that kind of clear before-and-after is rarer than it might seem. When it happens, it is worth paying careful attention to what it tells us about where things are actually heading.

Chapter 8: Anthropic

Some companies start just to make money. Others start because the people running them are genuinely scared of what happens if they do not build their product. Anthropic is that second kind of company.

The founders did not care about chasing market share, fighting for ad revenue, or finding a profitable niche. Instead, a group of tech researchers looked at where artificial intelligence was going. They saw how fast the tech was moving, how high the stakes were, and how intense the competition had become. They decided someone had to do this work differently. They wanted to move carefully, stick to clear principles, and speak honestly about the dangers.

Because of that choice, Anthropic grew into a major player in modern tech. It is famous not just for the tools it makes, but for how it makes them. This chapter covers that journey. It starts with the founders, looks at the ideas behind their mission, and follows the team as they grew from a tiny group into a business used by millions of people.

We will also look at the tough questions. What does it actually mean to build safe tech? How do you compete in a brutal industry when your own values slow you down? And what is Anthropic's real plan for the future? To

How AI Came to Be

understand any of this, you have to look at a brother and sister who decided to walk away from their old jobs.

The Amodei Siblings and the Road to Anthropic

Dario Amodei grew up in San Francisco. His dad was a doctor, and his mom worked as a school administrator. He was a bright student who easily jumped between science and humanities, always looking for ways to connect them. He studied physics, earning his undergraduate degree at Stanford and a PhD in biophysics from Princeton. His college research focused on how the brain handles information. That basic question stuck with him for the rest of his career.

After getting his PhD, Dario worked in biology and computer science before taking a job at Baidu, a massive Chinese tech company, in their Silicon Valley lab. At Baidu, he saw what large-scale machine learning could really do, and he built up the tech skills that defined his future. In 2016, he moved to OpenAI as a researcher and eventually became the VP of Research. He was one of the main people behind GPT-2 and GPT-3, the models that led straight to ChatGPT.

Daniela Amodei took a completely different path. She studied English literature at the University of California, Santa Cruz. Then she spent years working for nonprofits and political groups before switching to tech. She joined the payments company Stripe, where she rose to lead its

risk operations. Daniela was great at the business side of things. She knew how to build teams, run big projects, and turn wild plans into reality. When Dario went to OpenAI, she joined him there too, eventually serving as their VP of Safety and Policy.

Together, the siblings brought a rare mix of skills to the industry. Dario understood the deep math and science. Daniela knew how to build a company that could support that science on a massive scale. Since they both spent years at OpenAI, they knew exactly what worked well and what caused problems.

By the end of their time at OpenAI, they were worried by what they saw. It was not that the people there were bad or had wrong goals. The issue was how the company was set up. OpenAI grew too fast, took billions in outside investments, and changed from a nonprofit into a business aimed at making money.

These changes created a lot of pressure. The company began prioritizing fast rollouts over deep safety research. The market was getting competitive, and the rush to launch new products was intense. Important questions got pushed aside, like how the models actually think, what values they pick up during training, and where the real dangers lie.

In December 2020, Dario and Daniela left OpenAI. Five other senior researchers walked out with them. The

How AI Came to Be

group included Tom Brown, a main writer of the GPT-3 paper, Chris Olah, a pioneer in checking how neural networks think, Sam McCandlish, and Jared Kaplan. These resignations made waves because of who was leaving. People at the absolute top of the tech world were walking away to build something completely fresh.

They founded Anthropic in early 2021 and announced it publicly that May. The name comes from the Greek word for human. It is a clear sign that the founders believe human safety is the whole point of this technology.

Why Anthropic Was Created

The company did not start because of one single event. It was the result of years of worry about where tech was heading, and what would happen if nobody stepped in to fix it.

The main fear was simple: the strongest systems were getting smarter faster than anyone could understand. The companies building them were facing huge financial pressure to win. At the same time, the tools to check these systems were terrible. Nobody knew how to look inside their code, see how they made decisions, spot hidden biases, or find risks. The tech community knew about this gap for years, but knowing a problem exists is not the same as fixing it.

Dario had already written articles and given speeches about these risks. He argued that powerful systems

capable of changing the economy would arrive in years, not decades. He believed the time window to build safe foundations was closing fast. To him, solving this problem is one of the most important tasks in human history.

Anthropic was built to turn that belief into action. The goal was not just to publish school papers about risk. They wanted to build safety testing right into the actual engineering process. The founders wanted to prove that safe design could make their products better, not just slow them down, and that you could build powerful tech responsibly.

That distinction matters. Anthropic's founders did not want to stop progress entirely. They just thought it could be done better, which required a team willing to face tough questions.

This focus sets Anthropic apart from other tech giants. OpenAI talks about safety, but they are stuck in a race to launch products and make money. Google has incredible tech, but they have to protect their massive search engine and ad business. Meta gives away its models for free, calling it a win for the public. All of these companies do important work, but none of them started with the single goal of making safety the number one priority.

Anthropic did.

How AI Came to Be

The Mission and What It Actually Means

Anthropic's official mission is to build AI systems that are safe, beneficial, and understandable. These three words are chosen very carefully, and each one means something specific.

Safe has a very deep meaning here. It is not just about making sure the app does not crash or leak passwords. For Anthropic, safety means making sure an advanced system actually does what its creators want it to do. It means ensuring the tech sticks to human values as it gets smarter, and making sure humans can step in and fix things if something goes wrong. This is the core of alignment research. These are not simple tech fixes. They are some of the hardest engineering and philosophical puzzles anyone has ever tried to solve.

Beneficial means the tools Anthropic builds must actually help people. The goal is not just to beat test scores or look impressive to tech fans, but to genuinely make everyday life better for users. This sounds like common sense, but it forces the team to make tough design choices. A system built just to look smart might give confident answers that are completely wrong. A system built to keep people clicking might just tell users what they want to hear instead of the truth. To be truly helpful, a company has to balance these issues and think about what real help looks like in the real world.

Understandable is the hardest technical goal of the three. It points straight to a core research project at Anthropic called interpretability. This is the science of looking inside an AI to see what it is actually doing when it answers a prompt. Right now, modern AI is a mystery even to the engineers who build it. You see the input and you see the output, but the middle part, the reasoning chain and how the system weighs choices, is almost impossible to read. Anthropic wants to change this. They want to build systems you can explain and fix, rather than mysterious black boxes.

This mission changes how Anthropic runs day-to-day. The company spends huge amounts of money on research that will not make cash anytime soon. They share their discoveries publicly, even when it helps their rivals. Sometimes they delay product launches to run more safety tests, and they encourage their staff to point out flaws even when it is inconvenient.

This does not mean Anthropic is a charity. They need to make money to survive, and they fight hard for customers, talent, and market influence. The difference is that their mission drives how they compete, instead of the competition driving their mission.

Constitutional AI: Teaching Values to a Machine

How AI Came to Be

One of Anthropic's biggest breakthroughs is a training method called Constitutional AI. To see why this matters, you have to look at the problem they were trying to fix.

When you first train a large language model, you get a system that is just guessing the next word. It reads vast amounts of human text and gets very good at writing sentences that look normal. But guessing words is not the same thing as having morals. A system trained only on random internet text will repeat everything humans write, including toxic, lying, and dangerous comments. Getting the machine to act civilized requires an extra step.

Before Anthropic came along, the main way to fix this was a method called RLHF, or reinforcement learning from human feedback. With RLHF, groups of humans rank the AI's answers to show which ones are better. The model studies these scores and changes its behavior. It works okay, but it has big flaws. It is slow and costs a lot of money because you need constant human input. It also copies the personal biases of the reviewers. Worst of all, it never explains to the AI *why* a choice was better, it just gives a math signal showing what humans liked.

Constitutional AI works in a completely different way. Anthropic introduced this idea in 2022. The concept is to give the AI a written set of rules, a constitution, and use those rules for training. Instead of waiting for humans to grade every sentence, the AI reviews its own drafts based

on the constitution and fixes them until they match the rules. Another AI model, trained on those same rules, double-checks the work. The end result is a system with values that are clear, steady, and easy to track.

The constitutional principles Anthropic uses are themselves worth examining. They draw on a wide range of sources: the Universal Declaration of Human Rights, widely accepted norms around honesty and harm avoidance, Anthropic's own research on what makes AI behavior genuinely beneficial, and the accumulated feedback from extensive testing. The principles cover things like avoiding harmful content, being honest about uncertainty, respecting user autonomy, and treating different groups of people consistently. They are not arbitrary rules. They are an attempt to make explicit the values that most thoughtful people would want an AI assistant to embody.

What Constitutional AI achieves, at its best, is a system that does not just follow rules but understands principles. This is a meaningful distinction. A system that follows rules will behave correctly in the situations the rules were written to cover and may fail in novel situations. A system that understands principles can apply them to new situations it has never encountered before, reasoning from the underlying values rather than looking for a matching rule.

How AI Came to Be

The method has been adopted and adapted by other researchers in the field. It has influenced how several AI companies think about value alignment, and it has produced a body of research that has contributed to the broader understanding of how to build AI systems with coherent and consistent values.

Claude: From an Internal Model to a World-Class Assistant

Claude is the name Anthropic gave to its lineup of AI assistants. The story of Claude shows how the company turned its academic research into a tool that normal people can actually use.

The very first version, called Claude 1, did not get a big public launch like ChatGPT. Anthropic took a much slower route. In 2023, they quietly opened Claude 1 to just a few companies and software developers through a private setup. They wanted to see how the system handled real-world work, find bugs, and fix issues before letting everyone use it. This fit right into their main strategy: take your time, know exactly what you are creating, and do not rush just to get good headlines.

Claude 2 came out later in 2023 and brought huge upgrades. It could read much longer files, think through tricky questions better, and follow complex directions. Anthropic also worked hard to make Claude 2 more honest and safe. They cut down on the system's tendency

to invent fake facts with total confidence, a massive issue that plagued almost all early AI models, including Anthropic's own early tests.

From the start, Claude felt different from other bots because of its personality, not just its smarts. Anthropic put real work into how the bot acted. Claude learned to be very open about what it did not know, which many users found refreshing. Instead of guessing and making up a fake answer, it simply admitted when it was clueless. If someone asked it to do something dangerous or wrong, it would not just give a generic error message. It would calmly explain why it could not help. It focused on being actually useful rather than trying to look brilliant.

None of these traits happened by accident. They came directly from Constitutional AI and the values the team picked for training. For example, Claude's honesty about its own doubts comes from explicit rules telling it to match its confidence to the facts. Its fair treatment of different people comes from equality clauses in its constitution, and its ability to resist trickery comes from rules about holding its ground under pressure.

Claude 3 came out in early 2024 and marked a massive jump for the company. Instead of launching a single bot, Anthropic released a trio of models for different types of work. They called them Haiku, Sonnet, and Opus. Haiku was built for raw speed and low cost, Sonnet offered a

How AI Came to Be

solid middle ground, and Opus handled the heaviest tasks. The artistic names showed that Anthropic cared about craft and style, not just tech numbers. When it launched, Claude 3 Opus tied or beat the absolute best models from OpenAI and Google on major industry tests. For the first time, Anthropic proved it could compete at the very top of the tech world.

The changes in Claude 3 went way beyond test scores. The models got much better at understanding subtle context, reading massive documents, sticking to exact rules, and staying consistent across different topics. They also got better at explaining their thoughts, making them much more reliable for serious work.

Since that launch, Anthropic has kept upgrading Claude at a speed that caught the whole tech industry off guard. New updates roll out constantly, and each one shows the company's progress in making tech that is both highly capable and genuinely safe.

Today, Claude is used by individuals for everything from writing and analysis to coding and research. It is used by businesses for customer service, document processing, software development, and complex decision support. It is used by researchers and educators. It is used by people who need a thoughtful interlocutor for hard problems and by people who simply want help getting through their workday. The range of its use reflects the range of what a

genuinely capable and genuinely trustworthy AI assistant can contribute.

Competing and Complementing: Anthropic's Place in the AI Landscape

Anthropic operates in one of the most competitive environments in the history of the technology industry. Its primary competitors include OpenAI, whose ChatGPT remains the most widely recognized AI assistant in the world; Google, which has deployed its Gemini models across an enormous range of products and services; Meta, which has open-sourced its Llama models and made them available to millions of developers; xAI, Elon Musk's company building the Grok models; and DeepSeek, the Chinese lab whose unexpected cost efficiency rattled the industry in early 2025.

Each of these organizations has genuine strengths. OpenAI has a massive user base, a strong brand, and deep integration with Microsoft's products. Google has extraordinary data resources, computing infrastructure, and the ability to deploy AI across products that billions of people use every day. Meta's open-source approach has seeded an enormous ecosystem of developers and applications. xAI has moved fast and attracted significant talent. DeepSeek has demonstrated that frontier capability does not require frontier spending.

How AI Came to Be

Anthropic competes with all of these organizations on capability and has demonstrated, particularly with the Claude 3 family and its successors, that it can match the very best models the field has produced. But the way it competes is distinctive.

Most AI companies lead their public communications with capability claims. They emphasize benchmark scores, new features, and the tasks their systems can perform. Anthropic does this too, but it also consistently leads with safety. Its research publications, its public statements, its conversations with regulators and policymakers, all of them return to the same core argument: building capability and building safety are not in tension, and the companies that treat them as though they are will eventually build something they cannot control.

This position has both principled and strategic dimensions. On the principled side, Anthropic's founders genuinely believe what they are saying. They are not performing concern about AI safety for marketing purposes. They think the risks are real and that the industry needs organizations that take those risks seriously regardless of the competitive cost.

On the strategic side, the safety position has turned out to be a genuine differentiator. In an environment where many large enterprises and government institutions are trying to adopt AI responsibly and are worried about

deploying systems that might behave unexpectedly or inappropriately, Anthropic's track record and its transparency about its safety research carry real value. A company that has published detailed work on Constitutional AI, on interpretability, on how it thinks about the risks of advanced AI systems, is a more credible partner for organizations that need to be accountable for how they use AI.

The relationship between Anthropic and its competitors is also more complicated than simple competition suggests. The field of AI safety benefits from multiple organizations working on similar problems, sharing findings, and building on each other's research. Anthropic publishes a significant amount of its safety research openly, making it available to competitors as well as collaborators. It participates in industry-wide conversations about standards, governance, and best practices. It engages with regulators not just to advocate for its own interests but to contribute substantively to the development of policy frameworks that could benefit the entire field.

This posture is unusual. In most industries, companies guard their research closely and engage with regulators primarily defensively. Anthropic's willingness to share research and engage constructively with oversight reflects a genuine belief that the long-term health of AI

How AI Came to Be

development is a collective challenge that cannot be solved by any single organization acting alone.

Research, Interpretability, and the Science of Understanding AI

Beyond the development of Claude, Anthropic has made substantial contributions to the broader science of artificial intelligence, particularly in areas that directly address the safety and alignment problem.

The most distinctive of these contributions is in interpretability research. Chris Olah, who joined Anthropic at its founding, is one of the world's leading researchers in this area. His earlier work at OpenAI and Google Brain had introduced the concept of circuits, the idea that specific computational structures within neural networks implement specific algorithms. His work showed that it was possible, with careful analysis, to identify what a neural network was actually computing, not just statistically but mechanistically. You could look inside a model and find something that genuinely corresponded to a readable internal process.

At Anthropic, this research has expanded dramatically. The interpretability team has produced a series of findings that have changed how the field thinks about what is happening inside large language models. They have identified specific features, meaning directions in the model's internal representational space, that correspond

to recognizable concepts. They have found evidence of internal processes that look like reasoning steps. They have discovered how models represent uncertainty and how that representation affects their outputs.

Some of this work produced genuinely surprising results. In 2022, Anthropic's interpretability researchers published findings suggesting that the features inside Claude were not the simple, individual concepts they had expected to find. Instead, the model appeared to use what they called superposition, packing multiple concepts into single dimensions of its representational space. This was unexpected and had significant implications for how interpretability research needed to proceed. It also showed that the internal structure of large language models is more complex and more interesting than anyone had realized.

In 2024, the team released a paper that got everyone talking. They mapped millions of interpretable features inside Claude: stable internal patterns corresponding to concepts, behaviors, and dispositions. Some patterns even resembled traits like confidence and caution. This did not mean Claude was actually awake or feeling things like a human. But it suggested that the model's inner workings were far more structured and organized than a basic word-guessing machine should be. The study brought up big questions about how we should watch over these tools

How AI Came to Be

and what responsible tech development looks like when the system has important internal states.

The company has also done huge work on scaling laws. These are math formulas that predict how much better an AI gets when you give it more computer power, data, and size. Some of this research happened before Anthropic even existed, done by scientists who later joined the company. It completely changed how the whole industry builds models. By using scaling laws, you can figure out exactly how much cash you need to spend to get a specific level of smarts, making the business side of tech a lot less unpredictable.

Anthropic also spends a lot of time red teaming, which just means trying to break their own software on purpose. They have teams that try everything to get Claude to say awful, lying, or banned things. They use trick questions, long-term manipulation, and hacks called jailbreaks. Whatever they find goes right back to the engineering team to fix the code, making Claude much harder to trick. Anthropic shares many of these findings with the public to help other companies understand where AI is vulnerable and how to fix it.

On the business side, Anthropic works closely with large companies to use Claude for high-stakes jobs. These partnerships give them real proof of how the tech acts under heavy pressure, what kinds of mistakes matter

most, and how to deploy it safely. Anthropic treats these business deals as an incredible data source to see if they are actually meeting their safety goals, rather than just using them for easy revenue.

Anthropic's Long-Term Plan

Dario Amodei is unusually honest for a tech executive about what he thinks the future holds. In a long 2024 essay, he talked about how advanced AI could compress decades of scientific progress into just a few years. He argued that if we build this safely, AI could supercharge research in medicine, biology, mental health, and poverty. It will not replace human scientists, but it will give them tools so powerful that the speed of discovery will fundamentally alter.

However, he made sure to say this is just a possibility, not a guarantee. It all depends on getting the safety part right. A highly powerful system that does not actually care about human values could cause massive harm instead of helping. That is why Amodei views safety research as a basic requirement for progress, not an annoying tax on speed. You cannot enjoy the benefits of advanced tech if you cannot trust it.

Looking forward, the company is focusing on a few specific goals.

First, they want to keep pushing interpretability. The ultimate goal is to get to a point where engineers can look

How AI Came to Be

inside the AI and verify it is acting on the right values, even in brand new situations it never saw in training. Solving this verification puzzle would be one of the biggest wins in the history of tech safety.

Second, they want models to handle doubt better. Right now, most AI tools, including Claude, can sound totally sure of themselves even when they are completely wrong. This is a massive danger for serious jobs. A medical bot that lies with total confidence is way worse than one that simply says it does not know the answer. Anthropic has made some progress here, but they see it as a long-term project.

Third, they are creating better setups for human control. As tech gets smarter, humans need better ways to watch it, spot errors, and jump in to stop problems. Anthropic is building tools for this, while also working with politicians in Washington and Brussels on rules to make this oversight work in the real world.

Fourth, they follow a policy called responsible scaling. This is a public promise the company made to run specific safety checks before launching any model past a certain power level. As these tools grow stronger, the danger of bad actors abusing them grows too, so the bar for safety testing has to rise. Anthropic has publicly promised to lock away any system that fails these tests, even if they face massive financial pressure to release it. It is a rare promise

for a tech giant to make, and it shows just how seriously they take the risks.

Anthropic is also thinking seriously about what the world looks like if AI continues to advance rapidly over the next several years. This includes thinking about the distribution of AI benefits, about who has access to powerful AI systems and who does not, about the ways AI might change labor markets, and about the governance structures that would be needed to ensure that increasingly powerful AI remains aligned with broad human interests rather than narrow ones. These are not questions that any single company can answer on its own. But Anthropic has tried to engage with them honestly rather than avoiding them because they are complicated.

What Makes Anthropic Different

It is worth stepping back and asking what, in the end, distinguishes Anthropic from the other major players in the AI industry.

Part of the answer is technical. Constitutional AI, the interpretability research, the work on scaling laws, and the careful approach to alignment have produced genuinely distinctive contributions to the field. These are not marketing differences. They are research differences, and they have influenced how the entire industry thinks about some of its hardest problems.

How AI Came to Be

Part of the answer is cultural. Anthropic has built a culture in which safety concerns are taken seriously at every level of the organization, not just in a dedicated safety team that everyone else ignores, but in the way products are developed, in the way deployments are managed, and in the way the company communicates publicly about what it does and does not know. Building and maintaining that culture in a fast-growing company under intense competitive pressure is genuinely difficult, and it is not something that can be achieved simply by having the right intentions. It requires constant effort.

Part of the answer is structural. Anthropic has made specific public commitments, through its responsible scaling policy and through its research publications, that create accountability in ways that most companies avoid. When you publish specific criteria that you commit to meeting before deploying more powerful systems, you make it harder to rationalize away the shortcuts that competitive pressure creates. This is a deliberate choice, and it reflects a genuine belief that accountability is not just a regulatory burden but a tool for actually doing things right.

And part of the answer is philosophical. Anthropic was founded on the belief that the development of transformative AI is one of the most consequential things happening in the world right now, that getting it right requires more than good intentions and impressive

benchmarks, and that the organizations building this technology have a genuine responsibility to the people who will live with its consequences. That belief has not wavered as the company has grown and as the commercial pressures have intensified. It shapes decisions in ways that are sometimes costly and sometimes unpopular. It is also, in the long run, what makes Anthropic's contribution to the field distinctive.

Claude, the assistant that emerged from all of this, carries these values in its design. Every interaction reflects a set of deliberate choices about what an AI assistant should be: honest, helpful, consistent, transparent about its limitations, and genuinely committed to the wellbeing of the people it works with. These are not qualities that emerge automatically from training on large amounts of data. They are qualities that have been carefully built in, through Constitutional AI and interpretability research and years of testing and refinement.

The result is a system that millions of people trust with important work, difficult questions, and sensitive topics. That trust is not accidental. It is the product of treating safety not as a constraint on progress but as the foundation on which meaningful progress is built.

The Larger Story

This book has followed artificial intelligence from Alan Turing's theoretical frameworks through the

How AI Came to Be

consumer technology revolution Steve Jobs helped spark, through the rapid expansion of AI capabilities in the early 2000s, through the shock of ChatGPT, through Elon Musk's complicated and consequential engagement with the field, through DeepSeek's demonstration that frontier AI was not an American monopoly. Anthropic fits into this story as a particular kind of response to a particular kind of problem.

The problem is this: the people building the most powerful AI systems in history are under enormous pressure to move fast, and moving fast and being careful are often in tension. Most of the organizations in this race are optimizing primarily for capability and secondarily for safety. Anthropic was built to invert that priority ordering, and to demonstrate by building a genuinely competitive product that the inversion was possible.

Whether it will succeed in the deepest sense, whether the approach it has pioneered will influence how the rest of the industry develops AI as systems become increasingly powerful, is a question that cannot be answered yet. The most important capabilities are still being built. The most consequential decisions are still ahead. The governance frameworks that will shape how these systems are deployed are still being written.

What can be said now is that Anthropic has made the attempt seriously. It has contributed research that has

changed how the field thinks. It has built a product that competes at the highest level while being transparent about what it does not yet know. It has engaged with the hard questions about safety, alignment, and responsibility at a time when engaging with those questions is genuinely costly.

That is not a small thing. In an industry defined by speed and disruption and the pressure to ship before your competitors, choosing to slow down and think carefully about what you are building is its own kind of courage.

Dario and Daniela Amodei, and the team they have built around them, believed that artificial intelligence could be developed in a way that was worthy of the trust being placed in it. They left secure positions at one of the most prestigious AI labs in the world to prove it. The story of whether they were right is still being written.

But the chapter they have contributed so far is one of the most important in the history of this technology. And the questions they have insisted on asking, about safety, about values, about what we owe to the people who will live with the consequences of what we build, are questions that the entire field will need to keep answering.

They are, in the end, the same questions Alan Turing was really asking when he sat down in 1950 and wrote *Computing Machinery and Intelligence*.

How AI Came to Be

Can we build something genuinely intelligent? And if we can, what do we owe it? What does it owe us? What are we responsible for?

Those questions do not have easy answers. But Anthropic has made them impossible to ignore.

Michael C. Castellano

Section 3

Chapter 9: The Future of AI

We are living through one of the most defining eras in human history. The technological journey we have traced in this book took us from Alan Turing's earliest theoretical concepts straight to the debut of ChatGPT. We have looked at Elon Musk's boldest enterprises alongside Anthropic's focus on safety. Today, these innovations have brought us to an entirely unprecedented threshold. Artificial intelligence has stepped out of the realm of science fiction and academic research. It is officially here. It is present in our mobile devices, our medical facilities, our banking networks, our creative outlets, and our everyday dialogue. And this is just the beginning.

The question that comes my way more than any other is quite straightforward: What comes next?

It is precisely what we should be asking. The choices made over the next few years by developers, world governments, corporate leaders, and everyday citizens will determine the path of this technology for decades to come. Grasping our current trajectory is far more than just a fascinating mental exercise; it is an absolute necessity. Whether you are pursuing a degree, building a career, raising a family, or running a business, the evolution of AI will impact your life in both obvious and subtle ways. The

more clearly you anticipate what lies ahead, the better prepared you will be to handle it.

In this concluding chapter, my goal is to peer into the future. I want to offer a candid evaluation of where AI is headed. I am not writing this to spread fear, nor am I looking through a lens of blind optimism. Instead, I write as someone who has followed the evolution of this field with genuine fascination and weighed its consequences deeply. I have organized this forward-looking analysis into four distinct phases: the immediate one to five years, the five to ten-year mark, the ten to fifty-year window, and the reality of life more than half a century from now. Every single one of these eras brings its own unique flavor, its own distinct breakthroughs, and its own hurdles to clear.

Let us begin.

The 1 to 5 Year Outlook: The World is Already Changing

Consider this reality for a moment: the most disruptive wave of AI integration is not some distant event. It is happening all around us at this very second.

Over the next one to five years, the shifts we witness will not feel like futuristic fantasies. Instead, they will manifest as a swift remodeling of daily routines. Applications that felt miraculous just a short while ago will soon feel completely mundane. Complex tasks that used

How AI Came to Be

to demand an entire team of experts will be managed by a single professional utilizing smart software. The divides that historically separated the well-funded from the under-resourced will start to shrink, slowly and imperfectly, but with genuine impact.

From a practical standpoint, this is how I see things playing out. AI will transform into a permanent digital assistant integrated into almost every career path. Legal professionals will depend on AI to generate and analyze agreements. Medical practitioners will use it to analyze symptoms, catch irregularities in scans, and recommend custom care plans. Educators will leverage AI to customize lessons for individual pupils on the fly. Coders, advertisers, financial advisors, architects, and reporters will all find that AI has become just as essential to their daily routines as email or spreadsheets.

Of course, this transition will not happen uniformly. Certain sectors will pivot overnight while others push back. Many professionals will welcome AI as a tool that amplifies their capabilities, whereas others will view it as a genuine threat. Both perspectives are entirely valid. However, I hold one core conviction above all else: AI itself is not going to replace the professionals who learn to master it. It will, however, replace the professionals who choose to ignore it completely.

For everyday users, digital assistants are about to become much more powerful and personal. We can already see this starting to happen with tools like ChatGPT, Claude, and other similar programs. In the next few years, these systems will start to build what can be called a contextual memory. This means they will develop a real understanding of who you are, what matters to you, and how you look at the world. This is not a form of spying. Instead, it is a level of personal customization that used to be completely impossible. Your AI assistant will do a lot more than just answer basic questions. It will guess what you need before you ask, manage your schedule with real intelligence, help you think through tough choices, and act as a true partner for your mind.

At the same time, multimodal AI systems will grow up very fast. These are systems that can look, listen, talk, read, and think all at once. The days of only typing text back and forth with an AI are already ending. We are moving into a time where you can talk naturally to a computer, show it a picture to get a smart breakdown, or watch it create high-quality video from a simple written note. These features are not ten years away from us. You can actually use early versions of them today, so the next few years will mostly focus on making them reliable, smooth, and available to everyone.

We will also see government rules begin to form during this time. Governments across the globe are trying

How AI Came to Be

to figure out how to watch over AI without stopping its progress. The European Union has already set an early example with its AI Act. The United States is working on its own rules, and China is doing the exact same thing. The world of laws and regulations will be messy and far from perfect at first, but a basic structure will start to show up. That structure is incredibly important for making sure this technology grows in a safe and responsible way.

How to Prepare for the Near Future

So, how should you get ready for these upcoming changes? I have three main pieces of advice on how to navigate this shift.

First, you need to build your AI literacy. This does not mean you have to go back to school to become a computer programmer or a data expert. However, you do need to understand what AI can actually do and what it cannot do. You need to learn how to talk to it well and how to look at its answers with a critical eye. The most useful skill over the next five years will not just be knowing how to turn on an AI, because millions of people will know how to do that. The real value will come from knowing how to use AI wisely and carefully.

Second, you should focus heavily on the things that make us uniquely human. The talents that are the hardest for a machine to copy are the exact ones you should practice and grow. These traits include deep empathy,

strong moral choices, creative imagination, leadership, and the ability to make real friends. Do not think of these as secondary strengths. They are actually the most important skills you can have in a world where basic technical power is shared by anyone who has a computer and an internet connection.

Third, you must stay curious. This whole field is shifting much faster than any single textbook or school class can keep up with. The people who do really well in this new environment will be the ones who build a real, daily habit of learning. They will be the people who read, explore new tools, try things out, and change their plans when necessary. You are already doing this exact work right now by reading this book, and that natural drive is going to serve you very well.

The 5 to 10 Year Outlook: Embedded, Essential, and Everywhere

If the first five years are focused on people quickly adopting these tools, the following five years will be all about deep integration. By the time we enter the 2030s, AI will no longer feel like some brand-new invention that everyone is trying to figure out. Instead, it will feel just like basic infrastructure. It will be as quiet and essential to daily life as electricity or water from the tap.

Consider how you interact with the internet today. You do not stop to marvel at the fact that you can search

How AI Came to Be

for a fact in two seconds or text someone on the other side of the planet instantly. You just do it. That exact shift in your relationship with technology, moving from a fun novelty to an absolute necessity, is exactly what will happen to AI over the next ten years. It will become the foundation that supports almost every modern activity.

During this time, we will see the widespread arrival of what experts call agentic AI. Agentic systems are not just passive tools that wait around to answer your prompts. They are systems that can take a long sequence of actions, chase after big goals over time, and work with a real amount of independence. Imagine an AI that does not just help you write a quick email reply, but actually handles your entire inbox, books your meetings, talks to the AI assistants of your coworkers, orders office supplies, and alerts you only to the issues that truly need a human eye. That is not a wild fantasy. That is the actual path this technology is on right now.

Because of this, the shift in the workplace will speed up. Some jobs will shrink significantly. Other career paths will open up that we cannot even picture today. The long history of technology shows us that new tools usually create more work than they destroy over time. However, that shift is rarely easy, and it rarely helps everyone in the same way unless we make intentional choices to offer wide support and training. This is one of the biggest tests of the coming decade. We must make sure the massive

wealth and extra free time unlocked by AI are shared with the public, rather than grabbed only by the people who already hold all the money and power.

Healthcare will be one of the areas that changes the most during this window. Diagnostic tools driven by AI will become the standard way doctors treat patients. Finding new medicines, which used to take over a decade and cost billions of dollars, will speed up dramatically. AI will be able to spot helpful chemical compounds and predict how molecules behave with incredible accuracy. Custom medicine, meaning treatments built specifically for your DNA, your body, and your personal medical history, will move out of elite research labs and into regular neighborhood clinics.

Schooling will change in a similar way. The old model of teaching everyone the exact same thing at the exact same speed is already under a lot of pressure. It will soon transform into adaptive learning software that meets every individual student right where they are. A child living in a tiny rural town with very few resources will have access to the absolute best tutor imaginable, available at any hour of the day or night. This is an incredible step forward for making knowledge equal for everyone. If we implement it carefully, it could be one of the greatest tools for human progress ever built.

How AI Came to Be

Scientific breakthroughs will also happen faster than ever before. Combining AI with fields like biology, chemistry, climate science, and physics will unlock answers that might have normally taken human beings generations to solve. We are very likely to see massive wins in green energy, new treatments for brain diseases, and a much better ability to predict complex systems like world economies and local environments.

How to Ready Yourself for the Future

How should you start preparing for this five-to-ten-year horizon? The work actually starts right now. It is less about mastering one specific software app, since tools will change constantly, and more about growing the basic skills that make you flexible.

First, practice systems thinking. As AI makes small, individual tasks incredibly easy, your most valuable trait will be the ability to see the whole puzzle. You need to understand how different pieces connect, predict the ripple effects of big decisions, and know how to ask the right questions even when the answers are not obvious yet.

Second, work on your ability to collaborate, especially when it comes to human and AI partnerships. The teams and businesses that do well in this era will be the ones that learn how to mix AI into their workflows smoothly. They will set clear boundaries for what machines decide versus

what humans decide, and they will keep real human supervision on the things that matter most.

Finally, and perhaps most importantly, focus heavily on your personal values. As AI systems take over more important roles in society, the question of what we want those systems to achieve becomes vital. The people who understand right from wrong, who can explain what a fair outcome looks like, and who can hold large companies accountable for the software they build will be absolutely irreplaceable.

The 10 to 50 Year Outlook: AI as a Foundation of Civilization

Forecasting thirty to fifty years into the future requires a different kind of imagination. At this horizon, we are not extrapolating from current trends so much as reasoning about the nature of intelligence itself, about the long-term trajectory of scientific progress, and about the kind of civilization we are choosing to build.

Let me be honest about the uncertainty involved. No one knows exactly what the world will look like in 2060 or 2075. Anyone who tells you otherwise is either deluding themselves or selling something. What I can offer is a thoughtful assessment of the directions in which we are heading, informed by the trajectories I can observe today.

At this scale, I believe AI will have matured from a powerful tool into what I think of as a foundational layer

How AI Came to Be

of civilization—as fundamental to how we organize our world as writing, mathematics, or law. The question will no longer be whether AI is involved in a given domain of human activity, but rather how it is involved and what principles govern that involvement.

By this point in time, AI systems will almost certainly be capable of reasoning at a level that rivals or exceeds human expert performance across virtually all cognitive domains. This is what researchers refer to when they talk about Artificial General Intelligence, or AGI—systems that can learn and apply knowledge flexibly across diverse contexts, much as humans do. Whether that milestone arrives in 2030 or 2045 or later is a matter of genuine debate among the world’s leading experts. But few serious researchers today doubt that it will eventually arrive.

The arrival of AGI—or systems approaching that capability—will trigger a cascade of changes that are difficult to fully anticipate. Scientific progress could accelerate to a pace that makes today’s breakthroughs look gradual by comparison. An AGI system that can read, synthesize, and build upon the entire corpus of human scientific literature while running thousands of simultaneous experiments could compress decades of progress into years. The implications for medicine, energy, materials, and our understanding of the universe are staggering to contemplate.

AI will not be developing in isolation during this period. It will be converging with other transformative technologies, and that convergence will produce effects that neither technology could achieve alone.

Biotechnology and AI are already beginning to intersect in profound ways. Over the next few decades, I expect this intersection to deepen dramatically. AI will enable us to design biological systems with precision, to understand the mechanisms of aging and disease at a molecular level, and potentially to extend healthy human lifespan significantly. The ethical questions this raises are immense—but so is the potential to alleviate suffering on a massive scale.

Quantum computing, though still in relatively early stages today, may reach practical utility within this timeframe. When it does, its combination with AI could unlock computational capabilities that are genuinely difficult to reason about from our current vantage point—solving optimization problems that are currently intractable, enabling new forms of cryptography, and modeling complex systems with unprecedented fidelity.

Brain-computer interfaces, pioneered by efforts like those of Neuralink, may mature to the point where the boundary between human cognition and artificial intelligence becomes genuinely blurry. This is among the most profound and philosophically challenging

How AI Came to Be

developments on the horizon—one that will force us to revisit fundamental questions about identity, autonomy, and what it means to be human.

Robotics will advance in lockstep with AI, producing physical systems of remarkable capability and versatility. The combination of advanced AI with sophisticated physical embodiment will transform industries from agriculture to construction to elder care. The world of physical labor will change as dramatically as the world of intellectual labor, and the societies that manage this transition thoughtfully will have enormous advantages over those that do not.

What does this mean for the younger generation—for those who are children or teenagers today, who will come of age in this transformed world?

My message to them is this: you are not preparing for a world that already exists. You are preparing for a world that is still being invented, and you have the opportunity to be part of that invention. The skills that will matter most are not the ones tied to specific technologies—those will change. The skills that will matter are the ones that are enduringly human: the ability to think critically, to collaborate across differences, to lead with integrity, to imagine futures that do not yet exist, and to care genuinely about the well-being of others.

Michael C. Castellano

Study widely. Develop deep expertise in something, but do not let that expertise become a cage. The most powerful thinkers of the coming century will be those who can move fluidly between disciplines—who understand enough biology to talk to a computer scientist, enough philosophy to challenge an engineer, enough economics to contextualize a scientific breakthrough.

And above all: stay engaged. The world being built right now will be the world you inherit and eventually steward. Your voice, your values, and your choices matter enormously.

The 50-Plus Year Outlook: A Hypothetical Horizon

I want to be transparent with you as we enter this final section. Anything I say about the world a century from now is, by its nature, speculative. History is full of examples of brilliant people making confident predictions about the future that turned out to be profoundly wrong—and equally full of developments that no one saw coming at all. I offer the following not as prophecy, but as a kind of informed imagination—an attempt to think seriously about where the long arc of this technology might lead.

If I let myself envision the world of 2125, here is what I see.

How AI Came to Be

I see a world in which the distinction between human and artificial intelligence has become philosophically complex in ways we are only beginning to grapple with. Highly capable AI systems may be so deeply integrated into human life—into our decision-making, our creativity, our medicine, our governance—that the question of where the human ends and the machine begins will no longer have a clean answer. This is not necessarily frightening. It may be one of the most extraordinary expansions of human capability in history.

I see a world where the most pressing problems that threaten us today—climate change, pandemic disease, poverty, access to clean water and food—have been addressed, if not entirely solved, through the combined power of human ingenuity and artificial intelligence. I am genuinely hopeful about this. AI offers us tools of extraordinary power, and I believe that if we are wise in how we deploy them, we can solve problems that have defied human effort for generations.

I see a world where the exploration of the cosmos—something that has always stirred deep in the human spirit—has accelerated dramatically. AI will not just analyze the data we gather from space; it will help design the missions, pilot the spacecraft, and perhaps eventually determine where and how human civilization extends beyond this planet. The universe is vast. We have barely glimpsed its edge. AI may be the key that opens the door.

I see a world of technologies that today we can barely conceive. New forms of energy production. Materials that build themselves. Communication systems that operate at speeds and scales we cannot currently imagine. Medicine that repairs the body at a cellular level. Perhaps even forms of digital consciousness—entities that exist in ways that have no precedent in biological history.

The interplay between these technologies and AI will be extraordinary. Each advance will enable others. The pace of progress, already accelerating, may reach speeds that feel incomprehensible to us today—much as the pace of technological change in the 20th century would have seemed incomprehensible to someone living in the 1700s.

And yet—and I want to say this clearly—none of this is guaranteed.

The hopeful future I have sketched above is possible. But so are darker outcomes. AI systems that are misaligned with human values and allowed to accumulate too much power. Technologies that deepen inequality rather than alleviate it. A world in which the benefits of this extraordinary progress flow to a narrow few while billions are left further behind. Autonomous systems deployed in warfare in ways that make conflict more likely and more devastating. The erosion of privacy, autonomy, and democratic governance under the weight of pervasive surveillance and AI-driven manipulation.

How AI Came to Be

These risks are real, and the people who take them seriously are not alarmists—they are realists. The gap between the best possible future and the worst possible future is wide, and what fills that gap is the quality of the decisions we make today.

So what must we do? What are the decisions that will determine which future we inhabit?

We must insist on transparency. AI systems that make consequential decisions—about who gets a loan, who gets parole, who receives medical care, who is flagged as a security threat—must be subject to meaningful oversight and accountability. Black boxes are not acceptable when human welfare is at stake.

We must invest in AI safety research with the same seriousness that we invest in AI capability research. The organizations doing this work—Anthropic among them—are working on problems that may be among the most important in human history. They deserve resources, attention, and respect.

We must build international frameworks for AI governance. The challenges posed by this technology do not respect national borders. Climate change taught us, slowly and painfully, that some problems can only be addressed through global cooperation. AI governance is another such problem, and we would do well to begin

building those frameworks now, before the stakes become even higher.

We must ensure that the benefits of AI are broadly shared. This is not just a moral imperative—though it is that. It is a practical one. Societies with extreme inequality are unstable. If AI supercharges the concentration of wealth and power, the resulting instability will threaten everyone, including those at the top.

And we must cultivate wisdom—individual and collective. Technical capability without wisdom is a dangerous combination. The most important question is not what AI can do. It is what we should do with AI, and why, and for whose benefit, and with what constraints. These are questions that require not just engineers and economists, but philosophers, ethicists, poets, historians, and ordinary citizens with the courage to engage.

I began this book with a look back at Alan Turing, at the origins of a field that started with a thought experiment about whether machines could think. I want to end it with a look forward. Not with certainty, but with conviction.

I believe the future can be extraordinary. I believe AI can be a force for profound good—a technology that helps us cure diseases, understand the universe, lift people out of poverty, and expand the horizons of human possibility in ways we can barely imagine. I have seen

How AI Came to Be

enough of this field to know that the people working on it—the researchers, the ethicists, the entrepreneurs, and yes, the visionaries who sometimes get carried away but are reaching for something real—are, on the whole, trying to build something that makes the world better.

But that future will not arrive on its own. It will be built. By choices made in boardrooms and legislatures and classrooms and living rooms. By individuals who decide to stay informed, to engage, to demand accountability, and to bring their values to bear on the technology that is reshaping their world.

The age of AI is not something happening to us. It is something we are creating together. And that means the future is, to a remarkable degree, still ours to write.

Let us write it well.

Conclusion

This book began with a simple conviction: that most people interacting with artificial intelligence today have no real understanding of what it took to get here.

Not because they are uninformed or uninterested. But because the story has never been told in full. The headlines move fast. The press releases are polished. The debates are loud. And somewhere underneath all of that noise, the actual history, the people who sacrificed for it, the ideas that drove it, and the decades of quiet, grinding work that made it possible, gets lost.

That is what I set out to fix.

This book isn't saying that AI will save us or destroy us. People have argued about those two extremes for years, but both sides miss the point. The real message is smaller and much more important.

AI is just a human story. It didn't just appear out of nowhere. People made it, and choices shaped it. You can see human fingerprints everywhere you look. Alan Turing worked completely alone on ideas that his peers couldn't understand. Years later, teams at Anthropic stayed up late to fix problems that didn't even have names yet. Real people imagined this tech, argued over it, paid for it, and built it because they cared.

How AI Came to Be

Since humans made it, it is our responsibility. Tech executives and engineers shouldn't be the only ones making the calls. The choices that decide if AI helps everyone or just gives power to a few people belong to all of us. We will make better choices if we actually understand what we are dealing with.

That is why I wrote this book. I didn't want to brag with technical details, and I didn't care about jumping on a trend. I wanted to help people move from just being amazed to actually understanding. I wanted to give you a good foundation so you can handle a technology that will shape the rest of your life.

We looked at a lot of history to get here.

We started with Alan Turing. He is the man who started this whole thing. Back in 1950, computers were huge, slow, and couldn't do anything human. Then Turing asked if a machine could think. That single question changed the next seventy-five years. His early work gave us the blueprint for modern computers, and his Imitation Game gave us a way to talk about AI. He asked hard questions when the world wasn't ready, and he set the stage for everyone who came later.

Next, we talked about Steve Jobs and the personal computer. Jobs turned machines from specialist tools into things for normal people. He knew something that other engineers missed. He knew that technology is best when

it is invisible, blends into regular life, and just helps people get things done. His early ideas about personal assistants and simple machines set the stage for the world we live in now.

We then traced the expansion of artificial intelligence from the early 2000s into the present. Neural networks. Deep learning. The explosion of data. The massive investments in computing infrastructure. The steady accumulation of breakthroughs that most people never heard about but that changed what was possible. This was the era when AI stopped being an academic curiosity and started becoming the engine of the modern digital world.

Then came ChatGPT, the moment when artificial intelligence crossed from the background into the foreground of public awareness. Its release was not a sudden miracle. It was the visible tip of decades of invisible work. But it changed everything about how the world related to this technology, overnight, in ways that are still unfolding.

We looked at the complicated and consequential role of Elon Musk in this story, a figure whose ambitions and contradictions reflect something important about the era in which we live. We examined DeepSeek and what its emergence said about the global nature of AI development, and the limits of the assumption that the frontier of artificial intelligence was the exclusive property

How AI Came to Be

of a handful of American companies in the San Francisco Bay Area.

And we spent meaningful time with Anthropic, the company whose story I find particularly compelling not because of its market position or its technical benchmarks, but because of the questions it has insisted on asking. What does it mean to build technology responsibly? Can safety and capability coexist, or does one always come at the expense of the other? What do we owe the people who will live with the consequences of what we build? Anthropic did not invent these questions. But it has done more than almost anyone else to refuse to let them be ignored.

Finally, we looked ahead. At the immediate changes arriving in the next one to five years. At the deeper integration of the decade after that. At the transformation of civilization itself over the following half century. And at the long, speculative horizon of a world a hundred years from now, where the boundaries between human and artificial intelligence may look nothing like what we can currently imagine.

Through all of it, I have tried to hold one idea constant: that the future is not fixed. It is being built right now, by choices made by real people with real values and real responsibilities. History has shown us, consistently, that technology deployed with wisdom and intention

Michael C. Castellano

raises the standard of living. Technology deployed carelessly, or in service of narrow interests, causes harm. The difference is not the technology. The difference is us.

Now I want to speak to you directly.

I am not a household name. I did not found one of the Big Tech companies. I was not in the room when ChatGPT was born or when the first neural network cracked a benchmark that changed the industry. I am someone who started a business at fifteen, served in the halls of the U.S. Senate as a teenager, moved to Silicon Valley with a dream, and spent twenty-two years building things in a world that was changing faster than most people around me could see.

Along the way, the company I built, Engager, Inc., did something that I am still humbled to say out loud: we quietly passed the Turing Test in real life, not in a lab, not in a controlled academic setting, but with a real person who simply could not tell she was speaking to a machine. That moment validated seventy-five years of human ambition. It validated every researcher who stayed up too late, every investor who took a chance on a wild idea, every pioneer who carried this torch a little further down the road.

But I am not telling you this to boast. I am telling you this because I want you to understand where my perspective comes from. It does not come from a

How AI Came to Be

professor's podium or a pundit's desk. It comes from years of being in the middle of this story, of building things that did not yet have a name, of believing in something long before the rest of the world caught up.

And what I believe, with every part of my experience behind me, is this: we are going to be okay.

I know there are loud voices telling you the opposite. I know the Godfather of AI has been warning about fascism and inequality. I know the news cycle runs on fear. I know the dystopian storylines get more clicks than the hopeful ones. But when I look at the actual data, at literacy rates, at infant mortality, at poverty reduction, at the things Peter Diamandis has spent years documenting in painstaking detail, I see something different. I see a species that keeps finding its way forward. Imperfectly. Messily. But forward.

Biological intelligence created virtual intelligence. Virtual intelligence will help elevate biological intelligence. That cycle, that upward spiral of mutual growth between human and machine, is not a threat to our humanity. It is an extension of it. We improve the machines, and the machines improve us. That is not doom. That is the most human story there is.

But I also want to be honest with you. Hope is not passivity. Optimism without action is just wishful thinking. The extraordinary future I believe in will not

arrive on its own. It will require smart policies. It will require transparency and accountability. It will require people who understand this technology well enough to ask the right questions and demand real answers. It will require all of us to stay engaged.

You are now one of those people. By finishing this book, you have given yourself something that most people interacting with AI every day do not have: a genuine understanding of where it came from, what it cost, who built it, and where it is headed. That is not a small thing. That is exactly the kind of informed engagement the world needs more of right now.

So take it seriously. Use it. Stay curious. Ask questions. Build things. Demand accountability. Push back when the technology is used carelessly, and celebrate when it is used wisely.

Alan Turing imagined a machine that could think. He did not live to see how right he was. The unsung researchers who followed him never got the recognition they deserved. But together, across decades and continents and competing visions, they built something astonishing. Something real. Something that is now in the hands of billions of people every single day.

That story deserved to be told. I am grateful I got to tell it.

And now, the next chapter belongs to you.

HOW AI CAME TO BE

A COMPLETE HISTORY OF ARTIFICIAL INTELLIGENCE

About The Book

The real story of the most important technology of our time.

Millions of people talk to artificial intelligence every day. Almost none of them know the story of how it came to be.

It began in 1950, when Alan Turing asked a question the world was not ready for: *Can machines think?* What followed was a seventy-five-year relay race of dreamers and builders—from Steve Jobs to ChatGPT, Elon Musk, DeepSeek, and Anthropic—each carrying the vision further down the road.

Michael C. Castellano lived it from the inside. In 2010, more than a decade before the world caught on, he founded Engager, Inc. and patented true human-machine conversation. Years later, his virtual persona did what Turing only imagined: it passed the Turing Test in real life.

This is the complete, human story of AI—where it came from, who built it, and where it is taking us next.

MICHAEL C. CASTELLANO

ISBN

