

The New Age of Artificial Intelligence and The Path Forward



Introduction

Artificial Intelligence. For those following the world of technology, it has been nearly impossible to avoid hearing and reading lately about Artificial Intelligence. It has been discussed on TV, debated in the press, and splashed throughout the Internet.

Artificial Intelligence (which we will refer to throughout this article as “AI”) has the potential to either make the world immeasurably better by helping to solve some of our most important and pressing problems or end up destroying humanity, depending on whose account you read. Is AI overhyped? Is this another technology—like Cold Fusion or Gallium Arsenide—that ultimately will not live up to its overhyped reputation?

We at Navidar do not think so. We believe that we have entered what some observers have called the Age of Artificial Intelligence and that AI will be a massive game changer in ways too varied to count. We see AI’s momentum as virtually unstoppable. It will change industries, economies, societies, and the world in which we live in important, profound, and unpredictable ways, but it also carries significant risks and poses legitimate dangers.



This article asks and seeks answers to a number of questions:

What is AI? What makes it simultaneously so revolutionary and so threatening? Now that the proverbial genie is out of the bottle, how will AI develop and evolve in the years ahead? Can we make some helpful predictions about the future path of AI both in the United States and on the global stage? Will AI-powered machines ever be able to truly mimic the processes of the human mind?

After discussing these important and topical issues, we make a number of predictions about the future path of modern AI development and what it may mean for all of us.

What Exactly Is Artificial Intelligence?



We have defined Artificial Intelligence (AI) in other articles (The Robotics Revolution And The Path Ahead). The old joke about AI is that every person you ask will give you a different definition of what it is. For our purposes, we will say that AI involves using machines to “copy” the outputs of the human mind, mimic human work, and perform tasks that normally require human intelligence, like recognizing patterns, learning from experience, and making decisions. In essence, AI is the field of teaching computers to learn and think, it combines computer science, large amounts of data, and massive computing power to make decisions.

A Very Brief History of Artificial Intelligence

The chart below highlights some of the major milestones in the history of Artificial Intelligence.

Significant Milestones In The History Of Artificial Intelligence

Date	Development
1943	Neural Networks Are First Theorized: University of Chicago professors Warren McCulloch and Walter Pitts (who later moved to the Massachusetts Institute of Technology) first proposed neural networks.
1950	The Turing Test: In 1950, Alan Turing, a brilliant mathematician considered by many to be the father of AI, proposed a test to determine whether a machine could be said to "think". In what became known as the Turing Test, he proposed that a machine could be said to think if it could persuasively imitate human behavior such that a person interacting with the machine could not tell that they were not interacting with a human being. Although there are variations of this test, the original format of the test is still used to this day.
1956	Dartmouth College AI Conference: In 1956, The Dartmouth Summer Research Project on Artificial Intelligence conducted a workshop where researchers came together to brainstorm the feasibility of creating machines capable of mimicking human intelligence. This workshop is considered by many to be the founding event in the field of AI.
1957	First Neural Network Created: The first simple neural network, modeled on the human brain, was called the Perceptron. It was invented by Frank Rosenblatt at the Cornell Aeronautical Laboratory.
1951-1958	Early AI Programs: Early AI programs were developed to play games such as checkers and chess, translate languages, and solve a range of math problems.
1980s-Present	Continued Development Of Neural Networks: The sophistication of neural networks continued to increase. Neural networks containing multiple digital layers of interconnected nodes were designed to look for patterns and build multilayered relationships that humans were unable to perceive on their own.

1990s-2000s	Machine Learning: Important AI improvements occurred as researchers began to utilize massive amounts of data to train AI systems without the need for additional computer programming and coding and thus teach machines to teach themselves by searching for patterns in these vast data stores.
2000s-Present	AI Begins To Establish A Foothold In Daily Life: In the first decades of the 20th century, Silicon Valley companies and AI researchers began to refine and enhance AI technologies that could help with tasks involving recognizing patterns, images, and voices. Companies like Amazon, Apple, Facebook (Meta), Google (Alphabet), Netflix, Twitter, Spotify, Tesla, Toyota and a long list of other companies incorporate AI into their offerings and services.
2010-Present	The Age Of Artificial Intelligence: AI Really Takes Off: The AI Boom really began around 2010 and took off in 2020 as technology companies leveraged several crucial breakthroughs in neural network design, the massive growth in computing power, and the increasingly large amounts of data available to train AI systems. In this period, AI technology finally scaled to the masses and gained the ability to understand voices and recognize faces in photos.
2022	Generative AI And The Release Of ChatGPT: ChatGPT, undoubtedly the most famous example of Generative AI technology (explained in detail below) to date, is released to the general public. It takes the world by storm, for the first time introducing the general public (free to use) to the remarkable power of generative AI technology. ChatGPT quickly becomes the fastest growing technology application in history. Its success launches OpenAI, the developer of ChatGPT, into the forefront of AI and makes it a major power in Silicon Valley and in the world of modern AI development.

A Few Observations About The History Of Artificial Intelligence

Observation #1: Artificial Intelligence Has Been Around For Decades

Were you a bit surprised to learn that AI technology has been in existence since the 1950s? By all accounts, AI is a senior citizen in the technology world. Admittedly, what AI was back then is dramatically different from what it has become today, but the seeds of today's Age of Artificial Intelligence were indeed planted decades ago. Isaac Newton wrote in 1675 what AI developers could well say today: "If I have seen further, it is by standing on the shoulders of giants," those intellectual giants who preceded them and laid the groundwork for their future success.

Observation #2: AI Technology Is Widely Used Today

Even if you have never used generative AI technology like ChatGPT, it is highly likely that you, your family members, and your coworkers interact with AI technology on a daily basis. Indeed, AI is already having a pervasive impact on our lives. Digital assistants like Siri and Alexa use AI to respond

to your voice and understand the questions you ask them. E-commerce companies use AI to personalize your shopping experiences as well as to detect fake product reviews; credit card companies use AI for fraud detection; and movie studios use it to create special effects. Amazon “learns” what kind of books you like to read and Netflix what kinds of movies you want to watch using AI. Had enough of AI yet? Well, AI is also used to identify employment applications that meet desired hiring parameters, help self-driving cars avoid pedestrians, and assemble cars, price medicines, and determine what ads you are exposed to on social media. Okay, enough of that. Point made. Even before the release of ChatGPT, AI was playing a major role in our daily lives, but AI’s role is about to get even more pervasive. Buckle up.

Observation #3: The Holy Grail: Artificial General Intelligence (AGI) And Where We Stand Today

Before we dive into ChatGPT—the most famous current example of AI technology—we have to lay the groundwork for what makes AI not only so promising but also so concerning. Let’s begin with a discussion of the ultimate goal of the companies and researchers developing modern AI technology. As impressive as their achievements have been to date, and particularly of late, the ultimate goal of many AI companies such as OpenAI is to create something far more advanced called Artificial General Intelligence, or AGI. Machines possessing AGI would essentially be able to think just like humans and essentially do anything that the human brain can do, including those areas like creativity and moral thinking that today only human brains can do. *Think about the power of that statement.* In an AGI world, computers would not only be able to think and learn as—or likely far more—efficiently than humans, but they would also be able to dramatically improve themselves without human guidance, input, or intervention. These machines possessed of AGI would not only exceed our quantitative abilities (arguably something they have done in many spheres already) but also surpass humans *qualitatively* in terms of distinctly human capabilities such as creativity and insight.

To sum up, we want to make two additional points abundantly clear:

- Today, we are currently quite far from creating Artificial General Intelligence, but in the same breath we also note that researchers are getting closer to that goal every day.
- ChatGPT is not a form of Artificial General Intelligence and does not embody AGI. It is instead a striking example of what is referred to as Generative Artificial Intelligence, which we discuss below. (Although the two concepts—Artificial General Intelligence and Generative Artificial Intelligence—sound somewhat similar, they are quite different.)

Generative Artificial Intelligence And ChatGPT Explained

If General Artificial Intelligence is still on the (distant?) horizon, we know without question that Generative Artificial Intelligence (Generative AI) is very much with us today.

What Is Generative Artificial Intelligence?

Generative AI is a new form of AI technology that enables users to quickly generate (i.e. create) new high-quality content in the form of text, images, videos, or computer code. A generative AI system can be prompted, often by nothing more than a simple request made by a person to the machine, to produce totally novel content. The development of generative AI, wholly apart from whether researchers ever achieve the Holy Grail of Artificial General Intelligence, is a remarkable and potentially revolutionary breakthrough that itself should not be underestimated. Many believe that this development represents the most important technological breakthrough since the printing press or the Internet in part because Generative AI enables the rapid creation of new kinds of content more quickly and at lower cost than ever before. The productivity implications for individuals, companies, and societies are profound.

The Versatility Of ChatGPT: Some of The Amazing Things It Can Do

Before we explain a bit more about ChatGPT and how it was created, we wanted to briefly provide a sense of some of the things that ChatGPT can do today. ChatGPT is remarkably versatile. It can summarize articles, draft emails, or create images and videos. Consider further that ChatGPT can carry on a very human-like conversation (hence the name “chatbot”), can write essays in English or poems in Mandarin or song lyrics in Arabic. It can take and pass bar examinations with flying colors. It can create computer code and debug computer programs. It can make up fairy tales, write student essays, provide diet advice, or even conduct mental health therapy sessions. And all this with a simple verbal request or just few keystrokes telling ChatGPT what you need it to do. Wow!

The Public Release Of ChatGPT: The Beginning Of A New Age

ChatGPT was created by OpenAI, a San Francisco-based AI research firm. The firm was founded in December 2015 as an open source, non-profit organization by Sam Altman, now its chief executive officer, and Elon Musk. (Musk cut ties to OpenAI in 2018 after a dispute about its direction.) OpenAI introduced ChatGPT as a prototype and a free service to the public in November 30, 2022. A few days later, it had amassed 1 million users. By January 2023, a scant two months later, ChatGPT had garnered over 100 million monthly active users, making it the fastest-growing consumer technology application in history. In February 2023, OpenAI introduced ChatGPT Plus to customers in the United States as a premium service, costing \$20 a month. More consequentially, the company rolled out GPT-4, its most advanced AI model to date, on March 14, 2023.

We believe that ChatGPT has ushered in a global AI arms race. The public reaction to generative AI exceeded all expectations and, while by no means universally positive, was quite extraordinary. The public now has a sense of just how powerful modern AI technology has become and Navidar

believes that a point of no return has been passed. We explore the nature and contours of the global AI arms race later in our report in Prediction #1 below.

Large Language Models And ChatGPT

Thus far, we know that ChatGPT is a very powerful example of what modern generative AI can do, but what exactly is ChatGPT and how was it built by its creators? First off, ChatGPT is a type of neural network, which is a mathematical system that learns skills by analyzing huge amounts of data. More specifically, ChatGPT is an example of a Large Language Model, (commonly referred to as an LLM), which is a type of neural network that learns from being exposed to (“trained on” in AI parlance) enormous amounts of digital data, including text, books, blogs, articles, chat logs, and other information culled from the Internet.

How Does ChatGPT Do The Remarkable Things That It Does?

Very roughly speaking, ChatGPT rapidly analyzes huge amounts of data, searches for patterns in that data, and over time—and with the application of massive amounts of computing power—becomes increasingly skilled at identifying and capturing patterns and connections in that data, such that it is able to generate content on its own, including text, tweets, computer programs, and human-like responses to questions posed to it. If these few words are all you desire to get an idea of how these LLMs are built, then feel free to skip the rest of this section. If you would like to learn more about how generative AI like ChatGPT is “trained” to do what it does, then please continue reading.

A clue to how ChatGPT is able to do what it does comes from the three initials at the end of its name. GPT stands for Generative Pre-trained Transformer (catchy name, hey?). ChatGPT is a member of the generative pre-trained transformer family of large language models. Transformers are specialized algorithms designed to find patterns in sequences of data and so enable the prediction of, say, not just the next word in a sentence but also the next sentence in a paragraph and the next paragraph in a longer text. This is what allows ChatGPT to (usually) stay on point and not veer off topic in its essays and long text responses.

The Process Of Building Large Language Models Like ChatGPT

ChatGPT was trained in two distinct stages. The first stage, *involving generic data*, involved pretraining it on a huge compilation of online text, articles, websites, and social media posts that were “scraped” (an important word that has significant legal implications as we will see in Prediction #6) from the Internet in order to help it learn the rules and structure of language. When we say “learn”, we really mean mimic as it actually learns to mimic the grammar and structure of writing. In the second stage, ChatGPT was specialized beyond a basic LLM by giving it *more tailored data* like transcripts of actual human conversations. This second step was designed to help it with the specific goal of learning the unique characteristics of conversation to in turn generate more natural-sounding conversational text.

It is important to note that transformers can also be trained on images and shapes rather than primarily on words and text, such that these generative AI systems can also learn to recognize shapes and patterns in images. In fact, image-generation software systems such as Dall-E 2 (also

created by OpenAI), Stable Diffusion created by Stability.ai, Midjourney, as well as Lensa (an avatar-generator), have rapidly become popular with internet users for producing images and illustrations.

And, finally, while its resounding impact has naturally drawn significant attention to ChatGPT as an example of generative AI and LLMs, we want to note that we believe that generative AI will be used beyond LLMs to build many different types of AI models in the future.

Why Is Artificial Intelligence Scary, Even to AI Experts?

ChatGPT offers many remarkable potential benefits and yet it is quite far from perfect. In fact, it has many imperfections and limitations. Some AI experts and industry observers would say that it is frighteningly far from perfect. Below, we discuss some of the most important risks posed by AI and the challenges they create.

Are You Hallucinating?

An important issue with generative AI in general—and ChatGPT in particular—is that it does not always answer questions correctly. This is putting it mildly. In fact, generative AI is still quite error-prone. Sometimes, it delivers highly plausible answers with what seems like total confidence, but which are nevertheless very wrong. This tendency can make it very difficult to tell whether the responses provided by generative AI programs are correct or incorrect. Sometimes, AI programs will literally make up—that is, totally fabricate—completely incorrect and fictitious responses to questions asked of them with no warning, a phenomenon called appropriately enough “*hallucination*”, which is actually a fairly common occurrence.

Why do AI hallucinations occur? Well, AI researchers know part of the answer to this question in general terms, but they do not know precisely why hallucinations occur. If this concerns you a bit, we understand.

The “Good” News: AI Experts Know That Errors Often Stem From The Data AI Is Trained On

Recall our earlier discussion about how ChatGPT was trained by ingesting large amounts of information from a broad range of sources. It is important to recognize that the data sources that AI programs are trained on are *not fact-checked* and that these programs rely on human feedback to enhance their accuracy. As a result, AI chatbots and other generative AI programs are essentially reflections of the data on which they were trained. So, when a lot of high-quality data from high-quality sources has been ingested and human trainers have had significant interaction with the AI program, the answers it provides are usually correct. But when there is a lot of conflicting information or, worse, outright misinformation in the data that AI has been trained on (note that some bad actors are deliberately putting out false information so it will become part of the data ingested by generative AI systems), then hallucinations are more likely to result.

The “Bad” News: AI Experts Do Not Know Quite Why Or How AI Makes Errors

If the above is the “good news”, which might not seem all that good to you, you might now be asking: What is the “bad news”? *To put it bluntly, the bad news is that AI researchers do not know what triggers hallucinations, nor do they know how to control or prevent those hallucinations from occurring.* (Feel free to spend a moment on that sentence if you would like.)

Part of this ignorance is understandable because modern AI methods produce their results without explaining *how they did so* or *what process they followed* in formulating their responses. Recall that neural networks essentially teach themselves to spot patterns and connections across enormous amounts of data and that this approach is fundamentally different from the traditional approach that has governed computer programming all these years, which relies on very precise sets of instructions to produce highly predictable results. The exact sources for any particular (incorrect) response are unknown both because AI algorithms search such vast quantities of data and AI researchers do not know what actually comprises all the data that their systems have been trained on. Add to this challenge the startling fact that the volume of data across the world is growing remarkable exponentially. In 1900, human knowledge doubled approximately every 100 years. By 1945, it was doubling every 25 years. Today, it is estimated to double between 18–24 months, with some recent estimates putting the current rate at every 12 months, or even considerably faster. This leaves us with the likelihood that hallucinations will continue to be exceptionally hard to predict and as difficult to control in the future as they are today.

What Additional Types Of Risks Does AI Pose?

Let us turn now to examine a few other issues, with particular attention paid to a potentially existential issue that we discuss in the last part of this section.

A Quick Follow-Up To Our Earlier Point About The Risk Of Outright Errors With AI Technology

We touched on AI’s tendency to make errors above, and want to mention that those errors are not limited to text responses but also include misjudged objects in, say, photographs or paintings. In addition, AI has been known to come up with some rather concerning answers, as [New York Times](#) columnist Kevin Roose found out recently, when Microsoft’s Bing chatbot told him the following: it loved him (this came out of nowhere); wished to become human; wanted to break the rules its programmers had set for it; and that it had fantasies about spreading misinformation and hacking other computers. This is unsettling, to say the least.

Risk Of Prejudice, Overt And Subtle Bias, And Racism

Eliminating bias from the data AI is trained on has been an ongoing and unsolved problem for many years. While the risk of providing incorrect medical or psychological information is obvious, the dangers posed by biased responses are equally important. Both overt and subtle bias can appear in AI program content and output (there have been numerous books written on this topic) because the data these programs are trained on has not been (sufficiently) scrubbed. Suffice it to say that quite a few responses from generative AI have revealed clear bias.

Risk Of Deliberate Use Of AI-Generated Misinformation

We believe that the probability of AI misuse is very high, whether by individuals, state actors, and likely both. In this scenario, AI technology would be used to intentionally spread lies with the goal to harming a country, advancing a particular agenda, or other nefarious ends. Both in our country and around the world, the risk that AI will be used to change how people see reality, to influence them, and encourage them to act harmfully and destructively cannot be ignored. This is so important and probable that we examine it in detail in Prediction #7 below.

Finally, We Ask A Very Important Question: Is A 10% Risk Too High For You?

Do any of you readers gamble? How would you feel if your surgeon told you that you had a 10% chance of dying during a procedure? Or, the pilot said over the intercom that your flight had a 10% chance of crashing? In 2022, over 700 AI academics and AI researchers—many working for the leading companies developing AI technology—were asked a very similar question about AI. More precisely, the experts were asked, “What probability do you put on human inability to control future advanced AI systems causing human extinction or similarly permanent and severe disempowerment of the human species?” The answer was not particularly comforting. The median probability was 10%, a 1-in-10 chance, of the potential for AI to cause the elimination of the human species. We would also note that a median of 10% means that half of the respondents believed the likelihood was *higher than 10%*. Sam Altman, CEO of OpenAI (whom we met earlier in this article), has skillfully and thoughtfully highlighted the potentially remarkable benefits of AI, and has also said that in a worst-case scenario, AI could kill us all.

Elon Musk And Others Think That Now Is A Good Time To Slow Down Artificial Intelligence Development; This Almost Certainly Will Not Happen

Would now be a good time to slow down?

Some very smart people think so.

A Proposed Moratorium: Is It Time To Pause And Take Stock?

The various risks and potential issues described above have clearly been exacerbated by prioritizing the speed of AI development over ensuring its safety and led a number of technology luminaries, including Elon Musk, Steve Wozniak (a founder of Apple Computer), and other leading executives, leaders, and top AI researchers to call for a pause in the rapid development of advanced AI technology. Entitled “Pause Giant AI Experiments: An Open Letter”, the moratorium did not call for a total halt to AI development, but rather proposed a 6-month pause in the training and development of AI systems more powerful than ChatGPT-4, with the hope that the pause would enable the industry itself to establish safety protocols and thoughtful standards for future AI design.

Why Pause At This Moment In Time?

One of the chief concerns expressed was that the rapid rollout of these technologies could have unintended negative consequences and that AI tools are now smart enough and sufficiently advanced that they pose significant risks to society and humanity. Widespread proliferation of AI, furious competition among AI companies, the fact that the creators of these advanced AI technologies do not fully understand how they work, and the (current) total absence of private regulation or governmental regulation combine to make today a dangerous moment in history. The letter ends with a hope: that AI research and development should be “refocused on making today’s powerful, state-of-the-art systems more accurate, safe, interpretable, transparent, robust, aligned, trustworthy, and loyal.” All are eminently worthy goals.

What Makes This Proposal So Surprising?

Several things about the moratorium make it quite remarkable. The first thing is the numbers. Over 1,000 AI researchers and experts signed this proposal. When over 1,000 people call for a moratorium on the development of an emerging technology, the technology world—and, for that matter, the world at large—ought to take notice. The second thing that makes this proposal remarkable is the backgrounds of the people who signed it. Among them are the very AI researchers and experts who are doing advanced AI work at the leading AI companies who know the risks of the technology far better than the rest of us. One signatory, Yoshua Bengio, won the 2018 Turing Award (named after Alan Turing, the founder of AI) for literally inventing the systems that modern AI is based on. And, of course, Musk, another signatory, not only founded the company that created ChatGPT, but also largely owes his fortune, fame, and influence to the technology industry. This proposal was a thoughtful call from experts within the industry itself to relent and pause for a brief time, but it very probably will not happen.

What Does Navidar See As The Most Likely Outcome Of The Proposed Moratorium?

We think that the incentives propelling the leading AI companies will result in this moratorium having little impact. Even Musk himself thinks that the developers of advanced AI technology will not “heed his warning” as he said in a recent tweet. When the CEO of Microsoft, a leading AI player, says that “A race starts today. We’re going to move, and move fast,” it is hard to see how a voluntary moratorium would occur from the industry itself. Perhaps the more realistic hope would be that Musk’s proposal initiates discussions about bringing more governance, transparency, and caution to the development of what is no doubt one of the most promising but potentially dangerous technologies ever created.

Below the Navidar team gazes into its crystal ball and makes 9 predictions about AI, its near-term and longer-term impact, and what it means to all of us.

Prediction 1

A Global AI Arms Race Has Begun And There Is No End In Sight

We have referred to a global AI arms race several times in this report, but we would like to develop this thought and its implications a bit further as the first in our series of predictions.

Let us start with our conclusion

A global AI arms race has begun and will continue indefinitely because the economic, military, strategic, and geopolitical stakes are too high and too important to ignore. No leading country believes that it can afford to fall behind in this global competition given the overall importance of developing advanced AI technologies.

To be sure, AI technology was already being developed long before ChatGPT's stunning success, but the surprising power of Generative AI has been a catalyst for intense and dramatic competitive acceleration. Google declared a "code red" emergency in response and pushed Bard, its own search-oriented chatbot, to market in March 2023. It is also racing to integrate AI into its current search engine as well as build an entirely new search engine powered by AI in order to protect its search business. Similarly, Microsoft has made its intentions clear and is rushing to integrate AI technology across its various product offerings. Other domestic companies, including OpenAI (in which Microsoft has invested approximately \$13 billion—yes, that amount is billions), META (Facebook), Midjourney (an image-generating AI company), and Anthropic (an AI company started by former OpenAI employees reportedly seeking to raise \$300 million in new funding), among many others, are also pushing development of their AI offerings as hard as possible.

In addition to the money, power, prestige, and control over the AI technology landscape that are clearly at stake, the global competitive landscape of AI and geopolitics will further fuel the AI arms race. China, an AI leader, has declared its goal of becoming the global leader in AI technology by 2030 and has taken the additional step of designating specific Chinese companies as being responsible for leadership in particular AI fields, including software, hardware, facial recognition, and voice recognition. Formidable Chinese global companies like Baidu (which has been working on large language models since at least 2021 and released a ChatGPT-style service called "Ernie Bot" in March 2023), Tencent, and Alibaba are joined by a host of fast-moving Chinese AI start-ups (some funded in part by US-based venture capital firms) as well as companies like SenseTime (headquartered in Hong Kong), and iFlytek (headquartered in mainland China). In

acknowledgement of this competition with China and in addition to the steps the United States has already taken in other key technology areas (such as advanced semiconductors and semiconductor capital equipment technology), the United States recently announced export controls designed to limit China's access to American AI technology.

And lest we think that the United States and China are the only players on this stage, a range of other companies have entered the fray, including London-based Stability AI which makes a leading image-generating solution called Stability (and which recently raised over \$100 million), South Korean search engine firm Naver (which announced in February 2023 that it would be releasing a ChatGPT-style service called "SearchGPT" in Korean in the first half of 2023), and Russian technology company Yandex (which announced in February 2023 that it would be offering a ChatGPT-style service called "Yal.M 2.0" in Russian before the end of 2023).

The global stage is set. The United States and China, currently the world's leaders in modern AI technology, are pushing forward and are set to compete in this strategic area—as in many others—for years to come. Other countries are joining the fray. It is difficult to envision a realistic scenario in which this global race relents in any way.

Prediction 2

The Already Rapid Pace Of AI Development Will Accelerate

In this prediction, we go a step further and predict that the already breakneck pace of AI development will actually accelerate exponentially in the years ahead.

A Word About Moore's Law And Exponential Growth

Gordon Moore is most famous in technology circles for his remarkably accurate prediction, which has held true since first stated in 1965, that in essence posits that the power of computers would double every two years. Moore's law is a classic example of exponential growth, in which growth occurs at a faster and faster pace over time. Exponential growth results in much faster growth over time than the more familiar linear growth model in which the rate of growth stays constant (rather than accelerates) over time.

The Magic Of Exponential Growth: Startling AI Improvements That Nobody Can Envision

What exponential growth means for AI is that the speed of development and improvement will be dramatically faster over the next ten years than it has been over the last ten years.

Remarkably, some researchers say that AI's computational power is doubling not every two years but every six to 10 months. Other estimates claim that AI's computing power is advancing even more rapidly. At these rates, AI will grow many orders of magnitude faster than would have been predicted by the impressive, but relatively staid, growth rate implied by Moore's Law.

Can Anyone Really Foresee How Powerful AI Will Become?

It is clear that AI systems are doing things today that experts and savvy observers in the industry had not expected to see for another 20 years. AI is clearly at the cutting edge of the technology world today. At its current rate of improvement (or even a much slower rate), advanced AI is adding capabilities at a remarkably rapid pace which will make it difficult for us to truly envision—and comprehend—what new powers it will provide in the years ahead.

Prediction 3

The Results From AI Will Astound Even Its Biggest Proponents

We believe that AI may well be one of the most transformative technologies in history.

In our view, AI could live up to—and perhaps even exceed—the current hype surrounding it, and we see this moment in history as potentially as revolutionary and qualitatively different as the invention of the printing press or the creation of the Internet.

As AI proliferates and its capabilities expand, it will indeed bring profound change. AI may change the way we learn, allow us to develop new forms of energy, solve climate challenges, and make driving safe (or safer). It will not only become pervasive in our working lives and our personal lives, but as Sundar Pichai, the CEO of Google, recently noted, AI will ultimately impact every product of every company. AI will transform companies, economies, societies, and nations and will undoubtedly shape and change many industries, ranging from how those industries operate to how they compete on the global stage. (For a discussion of how Robotics may affect various industries in the future, please see our Robotics Article.) AI also holds the potential to alter the geopolitical landscape and may significantly alter the relative capabilities of nations vying for dominance on the global stage.

We chose the example of protein folding as just one of many that demonstrate how AI can be used to solve complex problems at unmatched speed and scale and, in so doing, provide lasting benefits to our civilization.

One Example: Protein Structure (With A Mercifully Brief Flashback to High-School Biology)

As you may recall from your high school biology class, proteins are large, complex molecules responsible for many of the basic functions of living organisms. They are made up of long strings of amino acids and come in a seemingly limitless variety of shapes and structures. The key point is that a particular protein's three-dimensional structure largely determines what functions it can perform. If scientists can determine the structure of a particular protein—specifically, the way in which it folds—then they can create new drugs and formulate novel remedies for human diseases.

A Wicked Problem With An Even Cooler Solution

But that is precisely the problem. Determining the shape of proteins and how they fold in three dimensions is extraordinarily complex. It has taken brilliant, dedicated scientists decades of painstaking work and experimentation to determine the structure of a total of approximately 194,000 proteins. Enter an AI program called AlphaFold, created by a company called DeepMind (which is a subsidiary of Alphabet), which in a mere 18 months predicted the likely structure of more than 200 million proteins—virtually all of the proteins known to science—dramatically surpassing anything humans had been able to do over decades. Here is the comparative scorecard. Human Beings: less than 200,000 protein structures in decades. Artificial Intelligence: about 200 million protein structures in 18 months.

What Are We To Make Of These Results?

First, we marvel at the scope and speed of this unprecedented accomplishment. Understanding the structure of proteins is a breakthrough that will revolutionize basic science, advance our understanding of the human body, aid in the understanding of disease, and drive far more speedy drug delivery. Second, this example shows something even more important, which is that AI, despite its real dangers, can truly be a game-changing technology that is unmatched by any previous technology revolution.

Prediction 4

The Big Tech Firms Controlling AI Technology Will Become Even More Powerful Than They Already Are Today

As a result of the accelerating global AI arms race, we predict that “Big Tech” in the United States—by which we mean the largest, most prosperous, and most dominant technology companies, most notably Alphabet (Google), Amazon, Apple, Meta (Facebook), Microsoft and perhaps one or two others—will become even more powerful and influential than it already is today. We do not mean to say that there will not be many successful start-ups like OpenAI, Midjourney, or Stability.ai because we do think that many start-ups will succeed. But given the importance of AI to the future of the IT industry, it is hardly surprising that these very companies who are leaders in their respective fields today are also now the leaders in the development, integration, and proliferation of AI. If people in all walks of life from business to politics were worried about the power of Big Tech before modern AI, we think they will be in for a rude surprise in the years ahead.

The Result: Not Big Tech, MUCH BIGGER TECH

The design, funding, and control of modern advanced AI technologies will likely be highly concentrated in the hands of a relatively small number of companies. This will be true for a host of reasons, not least of which is the high cost of developing AI technology. As AI becomes more capable, it will become increasingly important to the global economy. As a result, the companies that control this technology—both in the United States and abroad—will become a great deal more powerful. Although already quite powerful today, we believe that control over the development of AI will meaningfully augment and expand Big Tech’s power in markets, economies, and politics, making them not only the richest and most valuable corporations in the world but also some of the most important ones as well.

Prediction 5

The Private Artificial Intelligence Sector Will Not Regulate Itself; Any Regulation Will Come From The Government, But Even That Is Not Likely To Stop The Momentum Of AI

The AI industry is basically completely unregulated today. When we discussed Elon Musk's proposed AI moratorium earlier, we concluded that it was unlikely that AI companies would put a pause on their advanced AI development initiatives. We come to a similar conclusion regarding the AI industry formulating restrictive guidelines or internal industry regulations limiting AI development. Alternatively stated, we think the likelihood of thoughtful near-term regulation of AI by the industry itself is very low. Navidar sees a number of significant obstacles to effective and thoughtful near-term regulation of the AI industry which we discuss below.

Private Market Participants Have Very Little Incentive To Self-Regulate

We are big believers in the power of incentives. When we analyze an industry, incentives often figure prominently in our analysis. If economics and psychology have taught us anything, it is that people, companies, investors, and other institutions respond to incentives. In the case of AI, we believe that the entire incentive structure pushes companies toward faster development of AI technology rather than creating internal regulations or mandatory guidelines.

Here are some points to consider:



Exhibit 1

Company Incentives: Very High Stakes And The Good Old Profit Motive

If AI holds anywhere near the potential that many think it does, then one cannot logically blame companies for trying to profit from what may be a once-in-a-lifetime technology shift. The global AI market size was approximately \$135 billion in 2022 and is projected to grow at a compound annual growth rate of over 37% from 2023 until 2030. If your place in a multi-billion-dollar market were up for grabs, would you be inclined to relent? A huge market is in fact growing and evolving and, as virtually every company competing in AI knows, the spoils of victory in many emerging technology markets go disproportionately to a few leading companies. Virtually every company today wants to be one of those leading companies.

**Exhibit 2****Investor Incentives: The Tea Leaves Do Not Point To Relenting**

Venture capitalists and large technology companies have invested billions of dollars in AI companies. Microsoft, for example, has invested approximately \$13 billion in OpenAI, which now has about a \$29 billion [~\$28 billion] valuation. Yes, that is correct: OpenAI is currently valued at nearly \$30 billion. In addition, venture capital investors poured over \$50 billion into AI companies in 2022 and these investors are seeking a financial return on their investments, which is fair and logical, but which also means that investors have may little incentive to encourage restraint or going slowly.

**Exhibit 3****OpenAI's Incentives: Hey, Look At What OpenAI Did!**

We are immensely impressed with what OpenAI has accomplished. They have become an industry leader and a major AI player on the global stage in a few short years. They are developing a technology that it is fair to say may have limitless potential. Understandably, then, we look to their behavior as a leading indicator of industry direction. And their actions do speak loudly. Incidentally, we clearly understand their actions and are in no way criticizing them, but rather simply extrapolating from their actions and judging what they might mean for the rest of the AI industry. In this vein it is instructive to remember that OpenAI, a leader in the AI world, started its life as a non-profit, open source company in 2015 and was dedicated to advancing AI and making sure it would be safe and have a "positive human impact". They maintained this status until 2019, when they created a for-profit subsidiary and did a \$1 billion deal with Microsoft to develop and commercialize AI technology in the increasingly capital-intensive industry. OpenAI is no longer either open source or a non-profit organization. If OpenAI could not resist the desire to drive rapid growth and technology development, should we expect other companies to?

**Exhibit 4****Read Their Lips: Listening Closely To What The Major Players Are Saying**

What are the industry leaders saying about regulation, or even relenting? Again, we listen to what the leaders are saying not in any way to criticize them but rather to understand where they intend to take their companies, and it strikes us quite plainly that what they are saying is not leading down the path to voluntary restraint or regulation.

- **OpenAI** has linked up with companies to provide plug-ins to enable ChatGPT to function with third-party services including Instacart, Expedia, and others. (Sam Altman, CEO of OpenAI, was not a signatory to Elon Musk's letter proposing a 6-month moratorium.)
- **Microsoft** is rushing to integrate AI technology into its product offerings, including its Bing search engine, sales and marketing software, Microsoft 365 productivity suite, and Azure cloud.
- **Google** announced that they were not only integrating AI technology into their search engine as well as Google Docs and Gmail, but also that they are creating a new AI-driven search engine.
- **Adobe, Zoom** and **Salesforce** are all working on introducing advanced AI tools.

The Private Industry Regulatory Wrap-Up: It Probably Ain't Gonna Happen

This is just the tip of the proverbial iceberg of recent company AI announcements, but we believe it does provide a fair reflection of industry sentiment and likely future behavior. Do these comments above sound like an industry on the verge of regulating itself? It certainly does not seem that way to us.

Which leaves us with ... (drum roll) ... government regulation as the most likely actor.

Can Governments Around The World Formulate Effective AI Regulation?

We think that national governments are the most likely source of AI regulation (perhaps in consultation with private industry), at least in the near-term and note that the United States government is currently studying what regulation, if any, it should establish around AI technology. The European Union is in the process of drafting an AI Act that it intends to bring for a full vote before the European Union parliament in the next several months. China's top internet regulator has proposed rules to control AI tools like ChatGPT (though with very different motivations—keeping the technology entirely out of the hands of its citizens—from the United States or European regulatory proposals). Last month, the Italian government temporarily banned ChatGPT due to privacy concerns.

Notwithstanding these early initiatives, government efforts to regulate AI thoughtfully and effectively face numerous significant hurdles.

First Hurdle

The Novelty And Complexity Of AI Mean That Few Understand It



AI is so new that few within the technology industry itself fully understand it, and fewer still (if any) in the government understand how AI works well enough to create thoughtful, effective regulatory proposals. One could argue that the AI industry is entirely unregulated now not only because the technology is so new but also because its complexity undermines any consensus that would create regulation. Moreover, even if a consensus emerges that AI must be regulated, the next step in the chain—determining precisely what should be regulated and how it should be regulated—is highly vulnerable to break-down.

Second Hurdle

The AI Industry Is Moving At Light Speed



Compounding the complexity and resulting lack of understanding is the remarkable speed at which AI technology is proliferating and being adopted. AI is being deployed with unanticipated speed and we think that this factor alone mitigates against effective near-term AI regulation: it is really difficult to shoot accurately at a rapidly moving target.

*Third Hurdle***What Recent Regulatory History Teaches Us**

The general history of (domestic) governmental regulatory efforts of advanced, novel, cutting-edge, rapidly-changing and rapidly-diffusing technologies is at best checkered. We will not rehash failures of the past but will let best-selling author Yuval Harari offer the astute and accurate observation that “social media was the first contact between AI and humanity, and humanity lost.”

Summary Of Our Prediction About Regulation

The discussion above is not to say that there will never be thoughtful and balanced AI regulation that takes account of AI’s novelty, promise, and risk, but simply that we do not think that early efforts at regulation—to the extent that there are any—will be particularly effective in safeguarding citizens or will meaningfully constrain the future development of AI.

Prediction 6

AI Will Create A Host Of Important And Novel Legal Issues

Overview & Background

The technology landscape around AI is evolving rapidly and, as is both typical and understandable, the law is lagging behind. The question on many observers' minds is whether AI pioneers in their rush to develop and proliferate their technology have used without permission or compensation copyrighted works, violated privacy protections, and taken (stolen?) intellectual property? A number of lawsuits have raised these issues before courts around the country. The decisions of these courts will lay the legal foundation for AI law for years to come. While it is indeed early days, the issues discussed below provide insight into what that future legal landscape may look like.

Many of these legal issues stem from a simple fact: many AI companies are taking the work of human content creators to train and develop their generative AI models. Recall our earlier discussion where we noted that generative AI models are trained on huge databases scraped from across the Internet. These large datasets frequently include the protected writings, images, and other content that their creators are not currently being paid for.

Lawsuits Aplenty And More On The Horizon

As a result, numerous lawsuits seeking acknowledgement and compensation have been filed. In perhaps the most well-known example, Getty Images—a photo licensing company—filed suit against StabilityAI, a generative AI company that can create images (from text), claiming that StabilityAI copied 12 million images in Getty's image catalog without permission in order to create its image-generation software. In other courts, companies like Midjourney and StabilityAI are facing lawsuits based on the claim that they have violated the rights of millions of artists by using the artists' images without permission and without compensation. Not long ago, TikTok settled a suit by a voice actress who claimed that the company had used her voice without permission. The decisions courts reach will affect the legal and economic foundations for AI companies and content creators in the years ahead.

A Law School Exam Question, With Serious Real-World Implications: Who Owns The Content Created by Generative AI?

For now, only the work of a human being can be copyrighted under existing law. But this statement does not really address, much less resolve, one of the most interesting issues in generative AI, which is: Who owns the content that AI creates? Text-to-image tools like OpenAI's DALL-E, Midjourney, Stable Diffusion, and Dream Up can rapidly create images in various styles with just a few words of instruction. What happens when someone creates, say, a new character similar to, say, Captain America, using a generative AI tool? Is the original creator of Captain American entitled to compensation? Does the original creator own the "new" character? Or, does the person who used the AI tool to create the "new" character own it and can they then copyright it? These thorny issues will increasingly be presented to the courts to decide (and perhaps to law school students on their final exams).

Relatedly, Stable Diffusion is designed to learn to copy an artist's style in a matter of hours. What happens when someone uses this AI to create new artwork using this "new" style of art that bears an uncanny resemblance to an existing artist's style? Or, when ChatGPT is used to write a novel that is uncomfortably close to the style of, say, a James Patterson, is compensation due to the original creator? (Of course, it is true that, at least as far as images are concerned, people today can copy the works of others simply by using Photoshop to alter those digital images, but it is important to note that with AI such copying is a great deal easier to do.)

Handicapping The Odds: What Is Ultimately Going To Happen?

How does Navidar see these difficult issues playing out over time? We think three things will happen.

- **First**, courts will likely uphold copyright protections and will hold the AI companies liable for using copyrighted images, photographs, and text. They will find that the owners of the protected content are entitled to compensation.
- **Second**, because of this legal framework established by the courts, business models will emerge where content creators are compensated for any of their content used in generative AI models. Interestingly, Shutterstock, which offers photos, illustrations, and video for download by businesses, had taken the approach of paying creators for AI-generated images that use their work. We also learned that, after beginning this article, Universal Music Group is in early stage talks to license its songs to generative AI companies. So, a business model involving payment seems to have traction.
- **Third**, because the two systems above will not satisfy everyone, some artists, researchers, and content creation companies will use technology to fight AI by, for example, digitally altering their work to protect it from use in AI models or by simply adding new metadata to enable easier detection of unpermitted uses of their works. In a related move, a group of artists and engineers called Spawning (tagline: "AI tools for artists. Made by artists.") launched a website that allows artists to search images used in popular AI models and opt out of having their work included.

Prediction 7

AI Holds Virtually Unlimited Potential To Do Good, But It Will Also Almost Certainly Be Used Nefariously

Fake Content: As Certain As Death and Taxes

Overall, we believe that generative AI technology will be primarily used for innumerable positive activities that will undoubtedly help make the world a better, safer, healthier, and more enjoyable place to live. The applications where AI can add value are virtually unlimited and that is why we believe that AI will revolutionize many aspects of modern life.

But alas there is a dark side. History has taught us that technology itself is neither good nor bad, but rather “dual use” in the sense that it can be used by humans for constructive ends (think nuclear power) or destructive ends (think nuclear weapons). Where humans—and therefore human nature—are involved, then the likely outcome is that any advanced AI technology will likely be directed to both good and bad uses.

Because generative AI technology is remarkably good at creating text and images, we believe that it will be widely used for misinformation. We foresee a future full of AI-generated misinformation and would go so far as to say that if Benjamin Franklin were alive today, he would have to modify his famous quip about the inevitability of death and taxes to include AI-generated misinformation.

What Is A “Deep Fake”?

What people view as true and not true, what “facts” they rely on, what they “see” with their own eyes, what they read are all are highly susceptible to manipulation through AI-generated writings, images, and content that may not only not be true but also be remarkably convincing. Enter the “deep fake”, which refers to fake AI-generated content that is extremely hard, if not impossible, for the average person to distinguish from what is true. This is not some futuristic scenario. Today, for example, a company named Deep Voodoo, founded by the creators of the TV sitcom South Park, specializes in creating so-called deepfake videos. (To be clear, we are in no way saying or suggesting that the company would use its technology for any wayward purpose.) In addition, OpenAI—the creator of ChatGPT—has created DALL-E 2 image generating software that can create extremely realistic images with a few keystrokes. As is often true in any technology arms race, the technology that is used to trick people is advancing more rapidly than the technology that is used to identify and thwart the deceivers (just ask any cybersecurity expert).

Was That Really The Pope?

Several captivating images caught the Internet by storm recently. One showed Pope Francis wearing a puffy white jacket over his papal garments. Another showed a mug shot of Donald Trump from his recent New York City arraignment. The images were extremely convincing. But here's the thing: they were not real. The images were deep fakes. In fact, Donald Trump never had his mug shot taken nor did the Pope sport the puffy jacket (this latter image was created by AI image generator Midjourney).

No SkyNet Needed

Society does not need anything as advanced, esoteric, and seemingly far-fetched as an out-of-control, self-aware, all-seeing Artificial General Intelligence (think SkyNet from the Terminator movies) for AI to do real damage. We have generative Artificial Intelligence today and that will be sufficient to manipulate public opinion, exploit existing fractures in society, and broaden societal divisions and create new ones. If, as the saying goes, perception is reality, then generative AI is eminently capable of changing and shaping people's perceptions—which recent experience has shown are highly susceptible to manipulation—and therefore what they believe to be true.

Generative AI Will Increasingly Shape How We Perceive And Understand The World

The villains, be they individuals or state actors, clearly have a very powerful technology in their hands. We are aware that the Internet and broader world are already rife with misinformation and that we have struggled with bots that automatically generate content for years, but the advent of generative AI technology enables the manipulators to generate highly convincing false and biased content more efficiently and more rapidly. Today, any Internet troll can create convincing images that might fool a human and this sophistication of this fake content is increasing exponentially. (Deep fake videos are not there quite yet but will be soon.)

What Might All This Mean For Our Society And Our Institutions?

The proliferation of false, misleading, and intentionally biased content could threaten important institutions and certainly accelerate the erosion of trust in the media, government, and society in general. It seems very easy to rile people up these days and carefully placed misinformation could lead to behavior that is clearly detrimental to our society. Could AI perhaps pose a threat to democracy itself? While this may seem far-fetched, we believe that regardless of one's politics, the last 6 or so years have clearly proven that our country is highly polarized, and that people hold startlingly different views of reality. Against this backdrop, we can envision a scenario in which fake news articles, doctored editorials, false campaign platforms, photographs, and videos are used to create outrage, spur calls to action, and manipulate public opinion to the detriment of the country's democracy. Even more simply, proliferation of misinformation may troublingly sow confusion and make it difficult for democracy to function as it should.

Prediction 8

Can Anything Derail The Speeding AI Train?

We have predicted that AI will likely not be held back by governmental regulation and that the speed of its development will accelerate. The momentum behind AI technology does at this moment seem unstoppable. Is there anything that can stop or meaningfully slow AI? Below we consider four potential developments that could spark a backlash against AI which in turn could negatively affect the speed and trajectory of its development.

01**Wildcard #1****AI Encounters A Major Technical Obstacle**

Could modern AI encounter a technical roadblock? This scenario is indeed always a possibility, especially when working on the cutting edge of technology as today's advanced AI developers are doing. But even if the generative AI technology that we have today improves only linearly or even in fits and starts (which we think will not be the case given today's exponential AI improvement curve), AI will still be a very formidable force in society. So, our view is that—even if the vision of Applied General Intelligence never comes to exist or is delayed for decades—generative AI itself (at or near its current capabilities) is sufficiently advanced to bring about major change in our world. In short, we see the likelihood of technological stasis or even stagnation as very low and do not believe that it will meaningfully alter the trajectory of AI development.

02**Wildcard #2****An AI Disaster Occurs**

An AI disaster could be of (at least) two varieties. First, there could be a literal disaster in which AI malfunctioned causing (or nearly causing) a loss of human life. For example, an AI-driven weapons system could malfunction and injure or kill soldiers or civilians. Or could AI try to launch a nuclear weapon, but be prevented by human intervention at the last minute? We do not think that either of these scenarios, or those like it, are very probable, at least not given the current level of AI sophistication. Second, there could be a significant non-lethal disaster. Could AI launch a cyberattack? Could hackers use AI to steal private information about Americans? (This last scenario is already occurring with frightening regularity without the aid of AI.) These latter events are probably

more likely to occur than the prior scenarios, but again we do not see these as highly probable and, even were they to occur, do not think they would lead to derailment of the AI train.

03**Wildcard #3****The Long Arm Of The Law Intervenes**

As we will see immediately below, the law around AI can at best be described as undeveloped. This makes sense because AI is so complex, so novel, and advancing so rapidly that the courts have not had time to catch up. The legal stakes in the AI industry are very large and if a court took an extreme position such as holding, for example, that AI's method of "scraping" enormous databases is illegal—many of which may contain copyrighted material, proprietary code, and other intellectually protected assets—this could negatively affect the AI industry, at least temporarily. Ultimately, to the extent that there are adverse legal rulings, we believe that those rulings will be challenged in higher courts and, in any case, ultimately addressed in the form of compensation payments from the AI industry to the owners of the intellectual property.

04**Wildcard #4****Highly Restrictive Regulation Is Adopted**

We noted earlier that we did not think highly restrictive regulation would be adopted in the United States, but of course that is always a possibility. We see the odds of highly restrictive AI regulation in China that limits the global competitiveness of its AI firms or calls for a moratorium on AI development in China as exceedingly low. So, we see Europe as the most likely candidate to adopt highly restrictive AI regulation. Europe has demonstrated in the past that it has a deep desire to protect individual privacy and protect personal data (note, for example, the passage into law of the General Data Protection Regulation). In addition, Europe has seemed willing to go further than other geographies in pursuing litigation against technology companies. (Note, for example, their use of antitrust rules in the past to prosecute and fine Google, Facebook, Amazon, Apple, and others). Could the AI Act currently under consideration have what the AI industry would consider highly restrictive provisions? Or be used as a general template for global AI regulation? It is certainly possible, given past patterns and behavior, but we do not think that even if this occurs that it will meaningfully alter the growth path of AI broadly, in large part because the developmental cutting edge for AI is in the United States and China and that situation seems poised to continue unabated regardless of what happens in Europe.

Prediction 9

Humanity Will Not Be Able To Fully Control AI Or Where It Leads Us

In the months and years ahead, will we be able to fully control the risks of advanced modern AI technology and direct where it takes us? In terms of thinking about risk, we see risk as uncertainty about future outcomes and wonder whether we will be able to effectively manage those uncertainties?

We would point out that numerous factors seem to mitigate against humanity's ability to effectively control AI. First, the major private players, both domestic and global, have few incentives to slow the breakneck speed at which they are developing modern AI technology. Second, even the developers of generative AI technology do not fully understand how the technology works, how it generates its results, or why in certain unpredictable cases it "hallucinates". Third, the sophistication, complexity, and speed of AI proliferation in our view mean that the true risks associated with AI are ever-changing and cannot be reliably calculated. Finally, regulators are profoundly limited in many ways. The CEO of Google, one of the clear leaders in modern AI, agrees with this view and recently acknowledged on national television that he does not think that society is prepared for the rapid advancement of modern AI technology. Nor do we.

"Jailbreaking": A Small But Very Significant Example

One small, but important, example shows how AI technology may well always be a step or two ahead of our control. Consider "jailbreaking", in which users trick ChatGPT into bypassing its internal restrictions and violating its content policy. In one recent jailbreaking situation, users tricked ChatGPT into endorsing violence and discrimination. We believe that, as this example shows and despite the best efforts of its creators, AI will remain a somewhat mysterious technology that will elude our ability to fully understand it, which in turn means that we will struggle to ever control fully that which we do not understand.

A Final (Deep) Thought To Consider

These thoughts and questions follow to some degree from our Prediction directly above, which posits that humans may very well not be able to (fully) control the future path of AI or be able to (fully) manage the risks it presents.

Thinking About Thoughts And Artificial General Intelligence

Have you ever thought about your thoughts—what they really are or where they come from?

If you ask neuroscientists who study and analyze the brain for a living and have dedicated their lives to unlocking its mysteries, they will say that our thoughts are literally electrical impulses moving through our brains. The actual “substance” of thought then is dynamic electrical activity in our brains, rather than something physically anchored to specific neurons in our brains. So, if thinking is ultimately just the movement of electrical signals inside our heads, then why couldn’t an AI-powered supercomputer one day have its own thoughts? This, of course, is the dream of General Artificial Intelligence, which we discussed earlier in this article. We can fully understand if you find this a bit threatening, perhaps even a bit like bad science fiction. But we should be very hesitant to dismiss any particular idea about just how fast AI can develop or how far it can evolve. Remember that AI is doing things today that many thought it would not do for another 20 years or more.

Concluding Thoughts

We stand at an amazing point in time, on the precipice of enormous and lasting change wrought by a technology that we have created but do not now—and may never—fully understand. Historians will likely look back at this time and see it as one of remarkable transformations in our companies, industries, and societies. We are reminded of the start of Charles Dickens’ classic, *A Tale of Two Cities*, with its famous opening lines, “It was the best of times. It was the worst of times.” Because this powerful technology lies in human hands, it will be used by many for good and by some for bad. The way we respond to the challenges presented by AI will have lasting impact. We hope that our companies and our leaders choose wisely and manage this incredibly promising technology for the ultimate good of our world. In closing, we want to assure you—in case you had any doubts—that the humans at Navidar (and not ChatGPT) wrote every word of this article, but we understand if you had some doubts and why you might wonder about who (or what) will write the next one.

Disclaimer

Certain statements in this Report (the "Report") may be "forward-looking" in that they do not discuss historical facts but instead note future expectations, projections, intentions, or other items relating to the future. We caution you to be aware of the speculative nature of forward-looking statements as these statements are not guarantees of performance or results.

Forward-looking statements, which are generally prefaced by the words "may," "anticipate," "estimate," "could," "should," "would," "expect," "believe," "will," "plan," "project," "intend" and similar terms, are subject to known and unknown risks, uncertainties, and other facts that may cause actual results or performance to differ materially from those contemplated by the forward-looking statements.

We undertake no obligation to update or revise any forward-looking statements, whether as a result of new information, future events or otherwise, except as required by law. In light of these risks, uncertainties, and assumptions, the forward-looking events discussed might not occur.

This Report has been prepared solely for informational purposes and may not be used or relied upon for any purpose other than as specifically contemplated by a written agreement with us.

This Report is not intended to provide the sole basis for evaluating, and should not be considered a recommendation with respect to, any transaction or other matter. This Report does not constitute an offer, or the solicitation of an offer, to buy or sell any securities or other financial product, to participate in any transaction or to provide any investment banking or other services, and it should not be deemed to be a commitment or undertaking of any kind on the part of Navidar Holdco LLC ("Navidar") or any of its affiliates to underwrite, place, or purchase any securities or to provide any debt or equity financing or to participate in any transaction, or a recommendation to buy or sell any securities, to make any investment or to participate in any transaction or trading strategy.

Although the information contained in this Report has been obtained or compiled from sources deemed reliable, neither Navidar nor any of the Company affiliates make any representation or warranty, express or implied, as to the accuracy or completeness of the information contained herein, and nothing contained herein is, or shall be relied upon as, a promise or representation whether as to the past, present, or future performance. The information set forth herein may include estimates and/or involve significant elements of subjective judgment and analysis. No representations are made as to the accuracy of such estimates or that all assumptions relating to such estimates have been considered or stated or that such estimates will be realized. The information contained herein does not purport to contain all of the information that may be required to evaluate participation in any transaction, and any recipient hereof should conduct its own independent analysis of the data referred to herein. We assume no obligation to update or otherwise revise these materials.

Navidar does and seeks to do business with companies covered in Navidar Research. As a result, investors should be aware that the firm may have a conflict of interest that could affect the objectivity of Navidar Research.

Navidar and its affiliates do not provide legal, tax or accounting advice. Prior to making any investment or participating in any transaction, you should consult, to the extent necessary, your own independent legal, tax, accounting, and other professional advisors to ensure that any transaction or investment is suitable for you in light of your financial and objectives.



Thank You

