

Their capital at risk

The rise of AI as a threat to the S&P 500

July
2025



Authors

Sean Greaves

Autonomy is an Institute which creates data-driven tools and research for sustainable economic planning.

Published 2025 by © The Autonomy Institute



The Autonomy Institute
Rowhill Road
London
E5 8EB

Contents

Executive summary	4
Introduction by Will Stronge	7
Background	9
Method	11
Key findings & trends	17
Malicious actors	20
Deepfakes	21
Data security	23
Third-party providers	24
Energy	26
Legislation	28
Export controls	29
Competition	30
Disillusionment	31
Jobs	31
Conclusion	33

Executive summary

Executive summary

This report uses a range of cutting-edge LLM-assisted data techniques to extract **key risk information from S&P 500 company filings**.

Following the recent boom in generative AI, **we examine reported risks from these leading firms related to artificial intelligence**. We clarify the extent to which firms are reporting new AI related risks, what kind of risks are being reported and what these indicate about the broader dynamics of AI in big business.

Our analysis has identified that, **in the past year**:

- ↳ **3 in 4 companies** (380 total) have added or expanded upon **risk concerning AI, indicating a widespread concern with AI related risk**.
- ↳ **1 in 3 companies** (193 total) have added or expanded upon **risk concerning malicious actors using AI**.
- ↳ The number of companies citing **‘deepfake’ as a risk has doubled**, from 16 to 40.
- ↳ **1 in 5 companies** (95 total) have added or expanded upon the **risk of proprietary data or intellectual property being exposed through interacting with AI systems**.
- ↳ **1 in 10 companies** (56 total) have added or expanded upon **risk concerning third-party providers of AI models and software and their vulnerabilities**.
- ↳ **1 in 3 utilities firms** (10 total) have added references to **AI’s increasing energy requirements**.

- ↳ The number of companies citing **EU AI Act – the European Union’s primary legislation on artificial intelligence – related risk has tripled**, from 21 to 67.
- ↳ **1 in 3 companies** (168 total) have added or expanded upon **competitive risk relating to AI**.
- ↳ The number of companies citing **AI bias risk has doubled, from 70 to 146**.
- ↳ **1 in 10 companies** (57 total) have added or expanded upon the **risk of AI failing to deliver intended benefits, success or return on investment**.
- ↳ **Risks to jobs rarely feature among reported risks**, despite being a prominent public concern.

Introduction

**Automation
anxieties
at the top**

Introduction: Automation anxieties at the top

The rapid advancement of artificial intelligence (AI) in recent years, in the form of Large Language Models, has transformed industries, reshaped business models, and introduced new opportunities for innovation. However, alongside these benefits, AI has also brought a growing set of risks, ranging from ethical concerns and regulatory scrutiny to cybersecurity vulnerabilities and operational disruptions. Over the past two years, AI-related risks have gained increasing attention in corporate disclosures, particularly among S&P 500 companies, which serve as a barometer for the broader market and global economic trends.

This report uses computational text analysis combined with large language models to examine the rise in mentions of AI-related risks in S&P 500 company filings and earnings calls, highlighting the key concerns businesses are identifying and the implications for investors, regulators, and policymakers. Understanding this trend is essential for assessing how companies are positioning themselves in response to AI-related challenges and whether current risk management frameworks are evolving at pace with technological advancements. By analyzing these disclosures, this report aims to provide insights into the shifting risk landscape of AI adoption in corporate America and its broader impact on business resilience and governance.

This paper is an expression of Autonomy's ongoing work investigating huge amounts of information drawn from company filings, government contracts and other key corpuses of unstructured data.¹ With the latest available techniques, we can utilise unorthodox datasets to get a temperature check on economies, uncover activities that some would rather stay hidden and better understand the polycrisis that is unfolding around us.

Will Stronge
Chief Executive, The Autonomy Institute



¹ See The Autonomy Institute (2025) 'Democracy and transparency'. Available at: <https://autonomy.work/democracy-transparency/>

Background

Background

The S&P 500 is a stock market index that tracks the performance of 500 leading publicly traded companies listed within the U.S. The index accounts for about 80% of the total market capitalization of U.S. equities, making it a useful reflection of broad market trends amongst the largest companies.

All public companies listed on U.S. stock exchanges are required by law to file an annual report known as a Form 10-K with the U.S. Securities and Exchange

Commission (SEC). One of the most scrutinized sections of this report is *Item 1A: Risk Factors*. This is where companies outline, in extensive detail, the material risks that could negatively impact their business, financial condition, or results of operations.

Risk factors are not merely boilerplate, but a consequential and closely scrutinized text that can expose companies to litigation.

Risk factors reports are dense, running over multiple pages, and serve a dual purpose: 1) to warn investors about genuine threats, and 2) to shield the company from future lawsuits. If a risk is disclosed and later materializes, they can point to prior warnings in the 10-K to show it wasn't misleading shareholders. This creates a particular genre of corporate writing that is in turns part confessional, part legal disclaimer and part foresight.

In recent years companies have faced legal consequences regarding the accuracy of their risk disclosure. In 2023, the United States Court of Appeals for the Ninth Circuit upheld securities fraud complaints against Meta (formerly Facebook), ruling that it misled investors by presenting the misuse of user data by Cambridge Analytica as a hypothetical risk, even though the breach had already occurred. The court found that failing to acknowledge known events gave a false impression of

the company's exposure. Similar rulings against Alphabet and Forescout serve as a reminder that risk factors are not merely boilerplate, but a consequential and closely scrutinized text that can expose companies to litigation.

In the context of AI-driven technological disruption, our aim is to read between the lines, identifying where subtle signals of instability may be starting to take shape.

Given their legal and financial significance, even subtle year-on-year changes to Risk Factors are closely tracked by investors and researchers for signs of shifting threat perceptions. These small edits can serve as early indicators of emerging concerns or evolving strategic priorities. Research by Morgan Stanley and others has found that companies making notable adjustments to their risk language often underperform, while consistency is more often associated with stability. Therefore in the context of

AI-driven technological disruption, our aim is to read between the lines, identifying where subtle signals of instability may be starting to take shape.

Method

Method

We identified AI-related changes to risk factors via a three step process:

- ↳ **Change extraction** in which we identify textual changes between filings.
- ↳ **Filtering** to only those examples concerning AI-related subject matter.
- ↳ **Trend identification** in which we prompted an LLM to provide a classification of these changes.

This method is outlined in more detail below.

1. Change extraction

First we need to identify changes to risk factors comparing the latest filing with the previous.

Existing proprietary tools like [Alphasense](#) and [TipRanks](#) can visualise changes to risk factors at varying levels of granularity, but no adequate open-source solution currently exists. We've therefore built a custom data pipeline to detect changes at the word level.

We downloaded the two most recent 10-K filings as of May 1st 2025 for all 503 securities listed in the S&P 500. While the index tracks 500 companies, three of them (Alphabet Inc., News Corp, and Fox Corporation) issue two classes of stock, bringing the total to 503 filings.

From each filing, we extracted the text of 'Item 1A. Risk Factors'; the section of the filing which describes risks.

As noted earlier, changes to risk factors can be subtle yet significant; sometimes the addition of a single word shifts the context entirely. To capture differences at this level of detail, we computed textual differences between the same section or sentence across fiscal years. This allows us to generate visualisations where changes stand out with clarity.

Consider the following sentence from Intel's latest filing:

For example, threat actors ~~may leverage~~ **are leveraging** emerging AI technologies to develop new hacking tools and attack vectors, exploit vulnerabilities, obscure their activities, and increase the difficulty of threat attribution.

Excerpt from Intel's 10-K (2025 / 2024)

Compared with the previous year, we can see the language is more definitive in its appraisal of threat actors. This could suggest a number of things – perhaps Intel is experiencing an uptick in AI-powered phishing attempts within the past year. Regardless, such changes offer clues for analysts seeking to probe deeper into a company's evolving threat landscape.

To isolate genuine additions to reporting, we first needed to establish sentence-level alignment between reports – distinguishing sentences that correspond across years from those that represent true insertions or deletions. This involved a multi-step matching process:

- ↳ **Exact matches:** we began by identifying sentences that appeared verbatim in both years. These served as anchor points, helping to segment the text and build a structural reference frame.
- ↳ **Fuzzy matches:** between these anchors, we compared unmatched sentences using Levenshtein distance to find pairs that had been reworded or slightly modified, but retained their core meaning.
- ↳ **Residual sentences:** any remaining unmatched sentences with low similarity scores were classified as either newly added or deleted content.

Once sentence alignment was complete, we computed fine-grained diffs between matched pairs, highlighting precise word- or character-level changes. This allowed us to distinguish between the introduction of entirely new risks, the removal of outdated ones, and more subtle rephrasings that reframed existing concerns.

2. Filtering for changes concerning AI risk

The second step of our method was to identify if these changes comprised a meaningful update or expansion to AI-related risk, specifically. Our evaluation method goes beyond simple keyword matching, using a language model to address two key questions:

- ↳ Is the change relevant to AI?
- ↳ Is the change meaningful, rather than cosmetic?

To ensure the model performs optimally in accessing each change, this is presented as clearly and unambiguously as possible, with several preprocessing steps:

Change isolation. A single sentence can contain multiple distinct changes, risks, or complexities. By expressing textual differences in XML format, we can assign unique IDs to each change, providing a clear reference to label AI-relevant changes

Consideration of context. Sentences often rely on surrounding text for their meaning, especially when referencing branded tools or initiatives that imply but don't explicitly mention AI. To address this, we employed a smaller, more cost-effective language model (Claude 3 Haiku) to generate a short context statement for each altered sentence, providing just enough background that the altered sentence could be understood independently of its position within risk factors.²

Once all changes were labelled with unique IDs and context statements generated for their sentence, each sentence was processed by Claude 3.5 Sonnet. The LLM was instructed to identify the IDs of any changes relating to AI that *meaningfully expanded existing risks OR added new risks*.

In the following example, for instance, we can say that changes ["f13a2b7e", "9d80c4af"] are relevant to AI but are purely cosmetic, meaning they do not meaningfully expand or add new risks. On the other hand there are two changes representing distinct expansions in risk: these are 1) risk of Nth-party suppliers causing data breaches through their use of AI ["d3fa8b2e"] and 2) an increase in AI-driven phishing attempts ["3e97cd10"]:

2 Inspired by Anthropic's [contextual retrieval](#) method.

<unchanged>In recent years, we have expanded our use
 of</unchanged> <removed id="f13a2b7e">AI</removed>
 <added id="9d80c4af">artificial intelligence</added>
 <unchanged>within our software. This may lead to an increased
 risk of data breaches related to the use of AI by our</unchanged>
 <removed id="b7e9132c">employees</removed> <added
 id="6a5fdb9">employees, business partners,</added>
 <unchanged>and</unchanged> <removed id="e29cb407">business
 partners.</removed> <added id="d3fa8b2e">Nth-party suppliers.</
 added> <unchanged>We have also noticed an increase in cyber</
 unchanged> <removed id="a64fb29d">incidents.</removed> <added
 id="3e97cd10">incidents, including phishing attempts that appear to
 leverage AI.</added>

3. Trend identification

After extracting all changes from the Risk Factors sections that meaningfully expanded each company's AI-related risk disclosure, we classified these changes into specific categories of AI risk.

This enabled us to identify patterns in how different sectors and firms report on emerging AI-related concerns. To do this, we used Claude Sonnet 3.5 to categorize each change according to one or more categories from the [MIT AI Risk Repository](#)'s Domain Taxonomy of AI Risks. The model was instructed to assign multiple categories where appropriate, allowing us to capture the layered and multifaceted nature of AI risk reporting.

Classification results were then manually reviewed. The most relevant and noteworthy trends are summarized in the findings of this report, while additional data processing steps were applied to derive some final statistics and insights.

Limitations

Our method classifies topics using a large language model – as opposed to the more basic method of identifying keywords – and so some AI-related risks may have been overlooked. In other words, there is likely a greater risk of false negatives rather than false positives. This also means that some of our statistics may be slightly conservative.

Rather than share the full dataset of extracted AI risk changes in this report, we present a curated selection of excerpts that form the basis of our quantitative analysis. Each excerpt included has been manually reviewed and confirmed to meet our relevance criteria.

Key findings & trends

Key findings & trends

3 in 4 companies added or expanded risks concerning AI

380 companies (76% of the S&P 500) added or expanded upon AI-related risks within their most recent 10-K filings.³

Breaking these companies down by Sector (Figure 1), we find that in every sector more than half of companies expanded their AI risk disclosures (Figure 2).⁴ Unsurprisingly, the IT sector had the greatest percentage of AI risk expansion (96% of the total), closely followed by Finance and Communication Services.

S&P 500 companies expanding AI risk disclosure by GICS sector

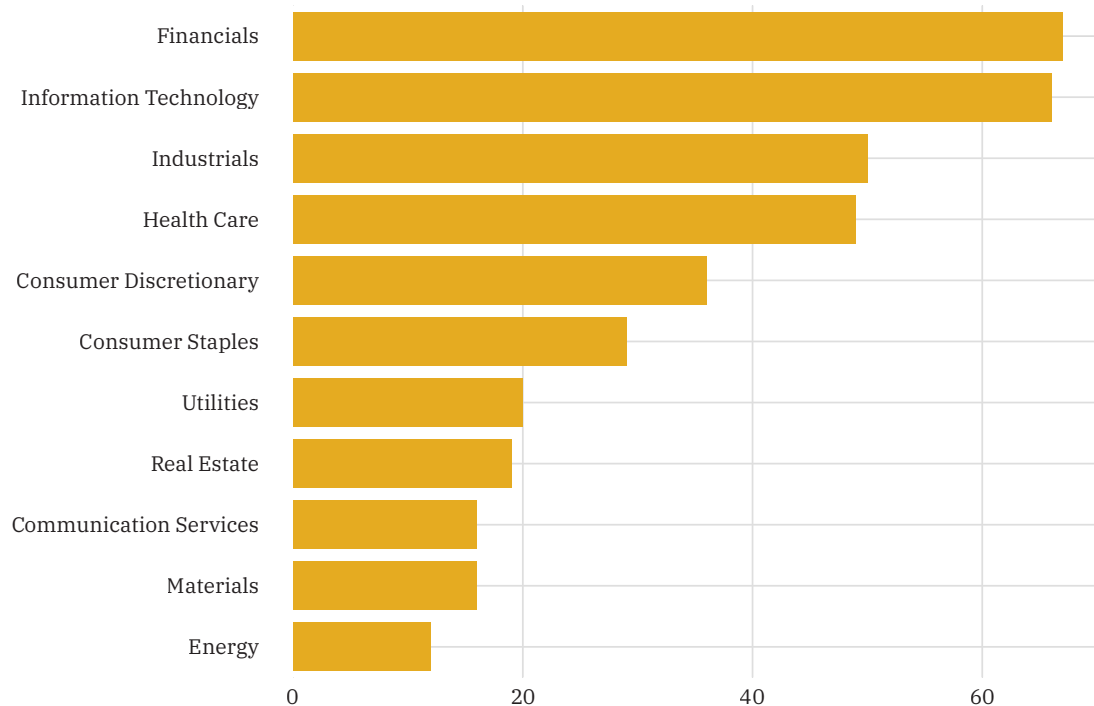
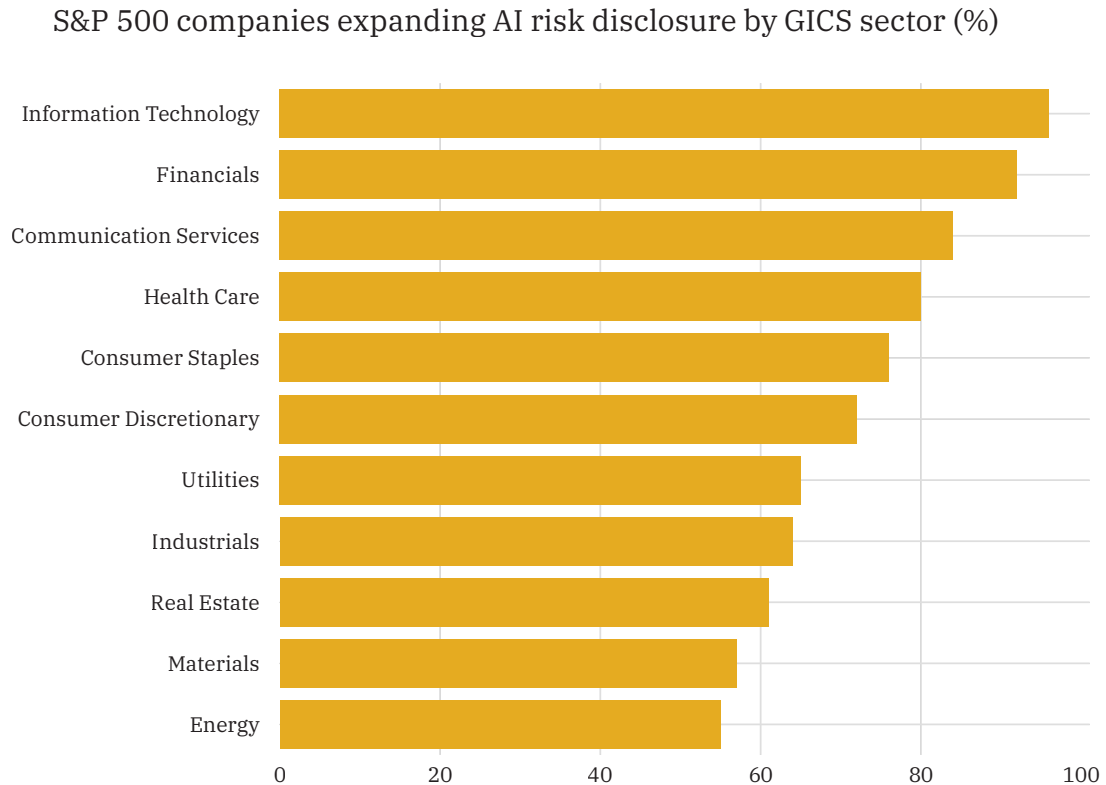


Figure 1. S&P companies by GICS sector. Source: Autonomy analysis of S&P 500 filings.

³ Autonomy analysis of S&P company filings. All 10-K excerpts can be found on the [Autonomy Data Unit blog](#).

⁴ Referring to **GICS** sectors.



*Figure 2. Percentage of S&P companies expanding AI risk disclosure by GICS sector.
Source: Autonomy analysis of S&P 500 filings.*

An analysis of companies expanding AI risks by GICS sub-industry provides a more precise picture within AI-focussed industries (Figure 3). In many of these smaller niches, almost every firm wrote about AI. For instance all 14 ‘semiconductor’ companies, 11 ‘asset management & custody banks’ and 9 ‘financial exchanges & data’ companies added further detail on their exposure to AI-related risks.

S&P 500 companies expanding AI risk disclosure by GICS sub-industry

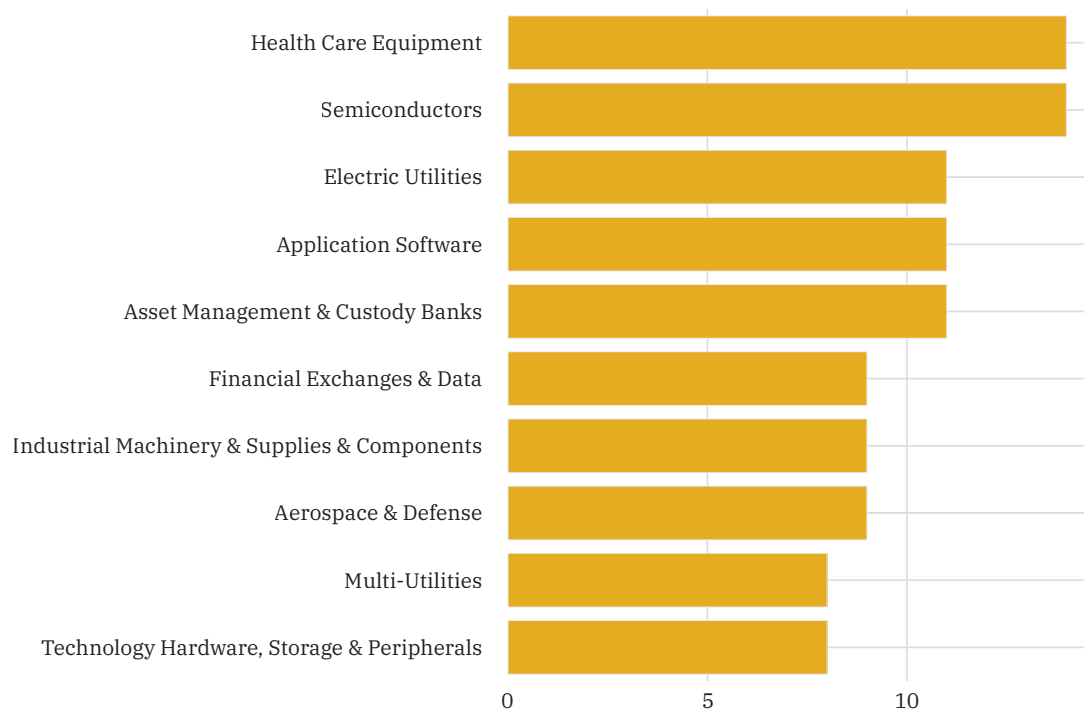


Figure 3. Number of S&P companies by AI related sub-industry. Source: Autonomy analysis of S&P 500 filings.

In the following section of the report we outline a number of trends identified among S&P 10-K risk factors relating to AI.

Malicious actors

1 in 3 companies added or expanded risk concerning malicious actors using AI

193 companies (39% of the S&P 500) expanded their disclosure of risks related to malicious actors leveraging AI. This group included 33 financial firms (45% of all S&P 500 finance companies) and 35 companies in the IT sector (51% of S&P 500 IT companies), highlighting a particular concentration of concern in industries heavily reliant on digital infrastructure.

The most frequently cited genres of threat included digital impersonation, the creation and spread of disinformation, and the use of AI to generate malicious code. Many companies noticed a significant rise in the volume and sophistication of attacks conducted with greater levels of automation, targeting, and coordination. Such attacks may harbour the capabilities to evade detection over long periods of time and erase forensic traces. Companies including Gen Digital, Salesforce, Intel,

and Visa Inc. all subtly updated their risk factors within the past year to establish these threats as no longer hypothetical, but actively encountered.

Additionally, as our market presence grows, we may face increased risks of ~~cyberattack attempts~~ ~~cyberattacks~~ or security threats, and as AI technologies, including generative AI models, develop rapidly, threat actors ~~may use~~ ~~are using~~ these technologies to create new sophisticated attack methods that are increasingly automated, targeted and coordinated and more difficult to defend against.

Excerpt from Salesforce's 10-K (~~2025~~/~~2024~~)

The range of actors potentially leveraging AI in cyberattacks is strikingly broad. Huntington Bancshares' growing list includes organized crime groups, terrorist organizations, state-sponsored hackers, hostile foreign governments, disgruntled employees or vendors, activist groups, and entities engaged in corporate espionage as possible perpetrators.

These threats are amplified by rising geopolitical tensions, which have prompted ~~warnings~~ from the U.S. federal government about an increased risk of cyberattacks. Companies involved in national security infrastructure, such as Textron, have reported facing particularly persistent, sophisticated, and well-organized adversaries, often exceeding the threat levels seen in other sectors.

Deepfakes

Companies mentioning deepfakes have doubled

Identity fraud is one of the most frequently cited malicious uses of AI in recent disclosures. 'Deepfakes' – digitally manipulated images, video, or audio that convincingly mimic real individuals – were mentioned by 40 companies during the most recent reporting period, up from 16 the year prior. This continues a steep upward trend in usage of the term across corporate risk disclosures over the past five years.

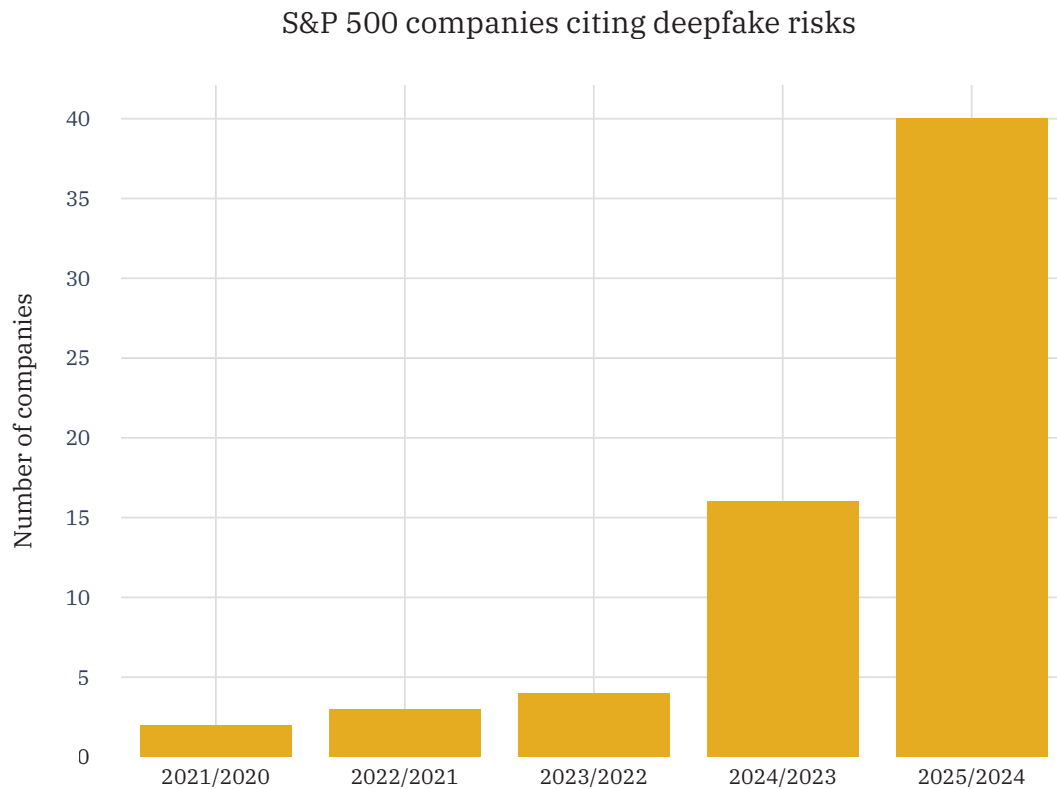


Figure 4. Number of S&P companies mentioning deepfakes since 2020 filings. Source: Autonomy analysis of S&P 500 filings.

The first S&P 500 companies to mention deepfakes were Adobe and Marsh McLennan in 2019, just two years after the term was coined by the eponymous Reddit user known for posting synthetic pornographic images. This pairing of Adobe and Marsh is apt; one company develops the creative tools that can be weaponized to generate convincing images, whilst the other, as a global insurance broker, is attuned to anticipating emerging risks. 6 years on, Marsh warns of just how low the barrier to entry for malicious actors using AI has become. The timeline stretching from the birth of a new technological risk to that same risk becoming an everyday attack vector appears to be narrowing.

Certain individuals are disproportionately subject to impersonation. Alphabet Inc., Microchip Technology and Blackstone all discuss the risks of deepfakes being used to mimic company executives or employees. Ebay Inc. revealed it was targeted by an adversary using AI to impersonate the voice of one of the company's senior leaders in an attack that ultimately proved unsuccessful. In addition, social media and legacy media companies alike warn of information ecosystems further polluted by deepfaked representations of public figures and politicians.

The manipulation of content by bad actors, including the creation of “deep fakes” (videos created with AI to realistically impersonate persons such as journalists or political candidates), could erode audience trust by making it difficult to determine what is real.

Excerpt from Fox Corporation's 10-K (2024/2023)

Data security

1 in 5 companies added or expanded the risk of their proprietary data or intellectual property being exposed through interacting with AI systems

95 companies (19% of the S&P 500) expanded their discussion of data privacy and intellectual property risks associated with the use of AI technologies. As firms increasingly funnel their sensitive information through nascent AI software, from customer records to proprietary business data, they face mounting concerns over potential data leaks.

These risks are particularly acute for those companies that rely heavily on third-party AI vendors, such as OpenAI or Anthropic. There are widespread concerns that these vendors may inadvertently or intentionally use their customers sensitive information to train their models, seeding the possibility that confidential data could resurface in interactions with future users or competitors. The opacity of AI systems and limited contractual control over how data is handled have only amplified these anxieties.

There is also a risk that our confidential information becomes part of a model that is accessible by other third-party AI applications or users as a result of a cybersecurity incident or a third-party AI developer's violation of our vendor engagement terms.

Excerpt from HCA Healthcare's 10-K (2025/2024)

The legal landscape governing data privacy and intellectual property in the context of generative AI is attempting to catch up. Companies are trying to mitigate risk through contractual agreements that impose data security obligations on vendors. However, many acknowledge the limitations of these efforts. Sherwin-Williams, for example, warns that in the event of a data leak or unauthorized use involving AI systems, existing intellectual property law may not provide adequate protection.

Concerns are not only limited to existing IP. Some companies are also grappling with the uncertain legal status of content created using AI tools. GoDaddy notes that “any content created by us using generative AI tools may not be subject to intellectual property protection,” a factor that could hinder the company’s ability to commercialize such content.

Amid this uncertainty, firms are increasingly taking matters into their own hands. Risk mitigation strategies include anonymizing data before feeding it into AI systems, restricting the type of data that can be processed by generative models, or imposing usage limits. Microchip, Amgen, and Brown & Brown all mention restricting employee use of third-party AI tools like ChatGPT. Yet even with internal governance policies in place, these same companies admit that their employees may circumvent such rules, which could result in exposure of sensitive information down the line.

Third-party providers

1 in 10 companies added or expanded discussion of risks concerning third-party providers of AI models and software and their vulnerabilities

56 companies (11% of the S&P 500), added or expanded their risk disclosures to include vulnerabilities associated with third-party providers of AI models and infrastructure. This shift reflects a broader structural reality: the current AI ecosystem is dominated by a handful of providers offering closed-weight, proprietary solutions like OpenAI’s GPT models or Anthropic’s Claude. In contrast, earlier in the decade, many companies built AI capabilities using open-weight, locally hosted frameworks like PyTorch, TensorFlow, or HuggingFace’s Transformers library. The result is a growing asymmetry; while companies increasingly integrate AI into core business operations, they are doing so through opaque systems they neither own nor fully control.

Beyond the now-familiar concerns around data privacy and intellectual property, companies are beginning to articulate a wider array of risks. These include dependence on third parties for model performance, uptime, and pricing; limited transparency into how models are trained; and the potential loss of functionality or data access if providers face legal or operational disruptions. As GE Healthcare warns, “we may have limited rights to access the underlying intellectual property used to create the generative AI model,” which could impair their ability to “independently verify the explainability, transparency, and reliability of the underlying model.” This lack of insight becomes especially critical in regulated sectors like healthcare, where explainability is not just desirable but often required.

Another concern is operational dependency. Companies such as Airbnb flag the risks of outages, data loss, or service disruptions stemming from hosted AI services: “any disruption, outage, or loss of information through such hosted services could disrupt our operations or solutions, damage our reputation, cause a loss of confidence in our solutions, or result in legal claims or proceedings.” In such cases, companies may have limited ability to recover damages from the affected provider.

Moreover, legal entanglements involving AI vendors could create cascading effects. If a third-party provider becomes subject to litigation or regulatory scrutiny, customers may suddenly lose access to essential tools or be forced to find costly replacements. In some cases, third-party licensors of AI technologies may impose restrictive terms or even revoke licenses, forcing companies to seek alternative providers or risk interruptions in product development and service delivery.

Cybersecurity is another rising concern. With the concentration of AI capabilities in a few providers, these entities become increasingly attractive targets for cyberattacks. Companies relying on them may inherit these risks without the ability to mitigate them directly.

Given these vulnerabilities, some companies are beginning to explore strategies to hedge against over-reliance. This includes diversifying their AI toolchain, investing in proprietary capabilities, or retaining the option to shift toward open-source or on-premise alternatives if needed. As dependency deepens, so does the strategic imperative to remain flexible.

AI technology and services require access to high-quality datasets, foundation models, and other AI system components. We currently rely, in part, on third parties to provide these components. In the future, we may face difficulties acquiring the necessary rights from third parties due to market competition and other factors. This challenge could hinder our ability to develop, implement or maintain AI technologies. To overcome this, we may need to invest in alternative strategies, such as forming alliances or developing our own resources.

Excerpt from Cognizant's 10-K (2025/2024)

Energy

1 in 3 utilities companies added reference to AI's increasing energy requirements

10 of the 31 utilities firms in the S&P 500 (32%) added new reference to AI's increasing energy requirements within their risk factors. These risks are focussed around significant projected increases in electricity demand driven by the expansion of AI data centers.

Southern Company, a major gas and electric utility and the second-largest in the U.S. by customer base, notes that “traditional electric operating companies are experiencing projected demand that exceeds recent experience, creating the need for new power generating resources and transmission facilities”. The company attributes the “majority of this demand” to “the power needs and projected power needs of data centers to serve an increasingly digital economy and to support artificial intelligence”.

Similarly, Exelon, the US's largest regulated electric utility, forecasts “substantial increases in load, driven largely by the increasing use of data processing facilities dedicated to artificial intelligence”. Utilities including Dominion Energy and Vistra Corp also identified specific regions, including Loudoun County in Virginia, otherwise known as “Data Center Alley”, the world's largest concentration of data centers, and parts of Texas, where AI data centers are becoming especially concentrated and may shape the focus of future energy infrastructure projects.

To meet this potential surge in demand, utilities recognize the need to rapidly scale energy generation and transmission capacity. Yet this acceleration carries significant risk. CMS Energy, Pinnacle West Capital, Centrapoint Energy, Vistra Corp, and Southern Company explicitly cite concerns about the difficulty of accurately forecasting energy demand for AI, recognizing it may not develop as planned. Vistra Corp offers several factors that could undermine these projections, including changes in technology, more energy efficient AI solutions or slow adoption of AI products and services, economic downturns, or adverse government actions.

NRG Energy is perhaps the most comprehensive in its description of this risk, that AI-driven growth may not materialize: “there is no assurance that these forecasts will be accurate or that the anticipated load growth will occur as projected. Factors such as evolving technology, improvements in energy efficiency, changes in economic conditions, shifts in government policy or regulation, and project delays or cancellations by the Company's commercial and industrial consumers (including data center facilities) could reduce or slow demand for electricity relative to current expectations.”

These uncertainties introduce financial risk, particularly the danger of overbuilding and incurring stranded costs. Pinnacle West Capital warns that as “data center and other extra-large customer opportunities evolve and develop, we may also enter into arrangements with customers and potential customers that require us to invest capital and assume credit risk related to such developments and the related generation and transmission investments before we receive any potential return.” This suggests a deepening interdependence between the development of AI and those building the energy infrastructure that enables it to grow.

In addition, volatility in stock prices of perceived significant energy consumers, such as technology companies involved with artificial intelligence or cryptocurrency, or other significant developments with such companies, could cause increased volatility in stock prices of energy utility companies such as Ameren.

Excerpt from Ameren's 10-K (2025/2024)

The impacts of AI and energy interdependence are not just limited to energy utilities. Top US tech firms developing AI and operating their own data centers, such as Meta, Tesla, and Alphabet, added mention of energy access and reliability as critical factors in scaling their AI infrastructure. Tesla writes that “as we continue to develop our artificial intelligence services and products, we may face many additional challenges, including the availability and cost of energy, processing power limitations and the substantial power requirements for our data centers”.

New AI may stretch existing data centres and compute facilities to their limit, requiring levels of power and cooling density they were not designed for. This is also a concern for several real estate investment trusts that operate a range of communications infrastructure and real estate for life sciences. American Tower, Healthpeak Properties, Digital Realty and Alexandria Real Estate Equities, all flag the risk of their facilities being inadequate to handle these new high performance computing workloads. This may require upgrades and alternative use of space through significant capital expenditure.

Legislation

Companies citing EU AI Act-related risks tripled

U.S. companies are grappling with a growing patchwork of AI-related regulations that expanded from just 1 in 2016 to 60 in 2024.⁵ Yet, it is the EU's legislation that has drawn the sharpest rise in corporate attention. The spike in references reflects mounting concern over the compliance burden and potential financial penalties associated with emerging AI laws.

The EU AI Act, which came into force in August 2024, establishes a sweeping, risk-based framework for regulating AI systems. In 2024, 67 companies cited the Act in their filings, up from just 21 the year before. Its provisions will begin rolling out in February 2025, with full enforcement by August 2026. The Act applies to both developers and deployers of AI, including those outside the EU whose systems are used within the bloc. By categorizing AI systems into unacceptable, high, and low risk, the Act imposes tiered obligations, with serious violations subject to fines of up to €35 million or 7% of global annual turnover.

Whilst companies have not yet been hit by fines or litigation relating to EU legislation, there are early instances of companies sharing experience of investigations and litigation from U.S. authorities relating to their use of AI within higher risk domains. Ford mentions its autonomous vehicle and driver assist technologies, including BlueCruise have been subject to government investigations. Elsewhere, healthcare and insurance provider Cigna added mention of the fact it is currently subject to litigation claiming improper use of AI within the medical claims evaluation process.

⁵ Dexcom 10-K filing. (December 2024). Available at: <https://www.sec.gov/Archives/edgar/data/1093557/000109355725000036/dxcm-20241231.htm>

Export controls

Seven IT companies mention US export controls on AI as risks

As geopolitical tensions escalate, AI technologies are becoming a central focus of U.S. export controls. The outgoing Biden Administration's Interim Final Rule on Artificial Intelligence Diffusion (AI Diffusion IFR), initially set to take effect in May 2025, would represent a major expansion of efforts to curb the global spread of advanced computing power. Major chipmakers, including Nvidia, AMD, Micron Technology, and Supermicro, warn the new rules could hurt their financial performance and competitive standing by cutting off access to key markets like China.

The AI Diffusion IFR would expose U.S. providers and the U.S. industry to an enhanced risk of retaliation from other countries, in the form of tariffs, import/export controls, or other regulatory actions.

Excerpt from Nvidia's 10-K (2025/2024)

The AI Diffusion IFR would establish a global licensing regime for high-performance chips and AI systems. The plan divides the world into three export tiers where only trusted U.S. allies and vetted buyers, known as "Verified End Users", would enjoy streamlined access to advanced hardware, while most countries would be subject to stricter licensing controls. U.S. companies would need to obtain case-by-case government authorization before selling to most foreign customers, even if those customers have bought and deployed similar hardware in the past.

The IFR would also introduce a global quota system, placing a cap on the number of advanced chips that could be exported per country per year. Nvidia and Supermicro warn that this would introduce a competitive scramble among U.S. firms for limited export slots, turning what was once an open, demand-driven market into a regulated contest, with potential buyers facing months of delay or outright rejection depending on their location, affiliations, or perceived risk. Companies also suggest the controls could be applied retroactively to hardware sold in prior years, restricting how customers can use, repair, or upgrade AI systems, from companies like Nvidia, that they already own.

Prior export controls appear to have had a chilling effect with some companies relocating warehousing and testing operations out of China. Supermicro and AMD caution that they may soon be unable to export any AI-related products to China whilst Nvidia reports that Chinese regulators have launched retaliatory investigations to establish whether the company is discriminating against local customers due to U.S. export compliance.

There is the growing sense among U.S. firms that they face a structural disadvantage. While they navigate expanding restrictions and opaque licensing processes, competitors based in Europe, Israel, or China often remain unaffected, prompting global customers to “design out” U.S. tech in favor of less regulated alternatives.

Competition

1 in 3 companies added or expanded their competitive risks specific to AI

168 companies (34% of the S&P 500) added or expanded upon competitive risks directly related to AI. A common concern among these companies is that competitors or new market entrants may outpace them in successfully adopting AI.

Yet, as with any nascent and disruptive technology, a less frequently acknowledged risk lies in moving too quickly. Rushing to release an AI-driven product that is flawed or underdeveloped can itself create significant competitive disadvantages.

Competitive pressures may also drive rapid AI development or deployment, increasing the risk of releasing inadequately tested or unreliable features.

Excerpt from Axon Enterprise's 10-K (2025/2024)

Disillusionment

1 in 10 companies added or expanded upon the risk of AI failing to deliver intended benefits, success or return on investment

Despite rapid advances in general AI capabilities, many companies investing heavily in AI struggle to achieve a clear return on investment, with 57 (11% of the S&P 500) explicitly warning that they may never recoup their spending or realise the expected benefits.

Excitement around AI creates pressure to appear innovative, yet quantifying tangible gains remains difficult for many at this stage. As uncertainty persists, continued investment at current levels may be unsustainable, raising concerns about the stability of the AI investment cycle. Even companies [publicly bullish](#) on AI often strike a more sober tone in their disclosures.

There are significant risks involved in deploying AI and there can be no assurance that using AI in our platforms and products will enhance or be beneficial to our business, including our profitability.

Excerpt from Palantir Technologies's 10-K (2025/2024)

Jobs

Just six firms mention work displacement as an AI-related risk

Although one of the most prominent public anxieties about AI is its potential to automate jobs, either by altering existing roles or displacing workers, this concern appears relatively infrequently in corporate disclosures. In two of the most detailed examples, Accenture notes that replacing some of its services with AI could impact the utilization rates of its professionals, while JPMorgan acknowledges that job displacement due to AI may affect employee morale and retention, requiring active management.

Perhaps from this point onwards, it may be possible to begin observing risks associated with second-order effects of automation. A real estate investment trust based in Palo Alto, California, managing rental properties in the tech hubs of California and Seattle, noted the risk of losing tenants whose jobs could be automated out of existence. This ties AI-driven job losses to a broader concern of potential instability in the property market.

•changes in the general or local economic climate that could affect demand for housing, including ~~layoffs, due to~~ an increase in the use of new technologies ~~and artificial intelligence~~ to replace workers, ~~slowing job growth,~~ and other events negatively impacting local employment rates, ~~tenant dispersion,~~ wages and the local economy

Excerpt from Essex Property Trust's 10-K (2025/2024)

Alongside concerns about workforce displacement, an even larger number of companies cited a different risk: the challenge of attracting or retaining employees with AI expertise. This shift underscores how AI is not only threatening existing jobs but also transforming the broader culture and structure of work.

In the technology industry, there is substantial and continuous competition for highly skilled business, product ~~development,~~ ~~development and~~ technical and ~~other personnel.~~ ~~personnel,~~ particularly in the AI field.

Excerpt from Oracle Corporation's 10-K (2024/2023)

Conclusion

Conclusion

Though seemingly mundane, risk filings offer a window into the internal perceptions of large private companies – such as the S&P 500 firms covered in this report. Indeed, such filings are frequently more candid and insightful than outward facing communications and media.

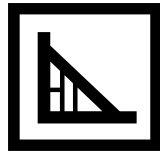
In the case of AI, the number of risks raised, and the frequency with which they are updated, suggests a broad base of concern among these firms. Notably, however, their concerns differ from those which might be held by the general public – or at least from those that appear in media discourse about AI. For instance, we found very little discussion or concern about job displacement from AI. And, where such concerns did appear, they were frequently secondary in nature, such as for Adobe which lists AI work displacement as a reputational risk for their products, or JPMorgan, which discusses the possibility of poor employee morale following displacements.

Instead, firms' concerns are focused, on the one hand, on the potential for AI to harm the business interests of firms via the actions of malicious actors, such as in the case of AI assisted cyberattacks, or through the exposure of sensitive or proprietary data via large language models. On the other, there is also concern from firms developing AI about emerging regulatory and trade regimes which are perceived as potentially impeding their ability to develop or sell their products. This is the case, for instance, among chip manufacturers anxious that export controls will limit their operations in China.

Particularly interesting, perhaps, are filed risks related to reliance on third party providers of AI models and infrastructure. These point to larger, structural, changes to the ways AI is being used within leading companies. The larger models on which cutting edge AI techniques rely cannot always be operated in house (even by S&P 500 companies), and therefore create dependencies on a smaller number of increasingly central firms that offer these services.

But most of all, these risk filings expose anxieties about the potential for AI to transform economies – or worse, that it might fail to do so. Firms note, for instance, that they could fail to keep up with a rapidly developing technical landscape. But risk factors from the cutting edge of the 'AI revolution' can also have a somewhat sobering effect – such as in the case of Palantir stating that “there can be no assurance that using AI in our platforms and products will enhance or be beneficial to our business.” Such statements suggest more moderated expectations than might be found in keynote addresses from the same firms.

At the same time, emerging anxieties about the risk of AI being used against the interest of large firms by nefarious actors or hackers could suggest a migration of momentum away from big tech towards unexpected and – from the perspective of these firms – unhelpful deployments of the technology. That three quarters of S&P 500 companies have updated their risk assessments with regards to AI does not indicate that AI is a technology that big business feels it has mastery over – on the contrary it indicates a growing uncertainty around what its outcomes will be.



**Autonomy
Institute**

**Published 2025 by:
© The Autonomy
Institute**

The Autonomy Institute
Rowhill Road
London
E5 8EB