

# Artificial Intelligence and its Rapid Development Advantages and Disadvantages A Review and Representation of Current Evidence

Firas Rifai

Department of Business Administration  
Al-Zaytoonah University of Jordan, Amman, Jordan  
f.rifai@zu.jo

**Abstract** Artificial intelligence (AI) has moved from a specialist research field into one of the most consequential technologies of the century. This paper asks two questions. How fast has AI actually developed, and by what measure? And what has that speed produced, for whom? Drawing on the Stanford AI Index, McKinsey's global survey of firms, International Energy Agency data, Pew Research Center polling, and peer-reviewed work on algorithmic bias, it finds steep growth in investment, adoption, and model capability, alongside real gains in cancer diagnosis, scientific research, and firm productivity. The same evidence base documents the costs: displacement concentrated in particular occupations, discrimination built into systems already in use, electricity demand rising several times faster than the grid as a whole, a rapid expansion of AI-enabled fraud, and rules that are still being written. The central finding is that the gains and the harms do not land on the same people. Productivity and wage growth cluster in firms already equipped to exploit AI, while displacement, biased decisions, and synthetic-media fraud fall hardest on routine workers, on groups that historical data has always undercounted, and on those least able to verify what they are looking at. Whether the coming decade resembles a broadly shared productivity boom or a narrower one will depend less on the technology than on how fast institutions catch up to it.

**Keywords:** artificial intelligence, machine learning, technology adoption, algorithmic bias, AI governance, labor market.

## 1 Introduction

Artificial intelligence (AI) refers to computer systems that perform tasks once thought to require human cognition: perception, language, planning, decision-making (Russell & Norvig, 2020). The field is nearly seventy years old, but its current phase began around 2012, when deep learning, large training datasets, and cheap parallel computing power converged (LeCun et al., 2015). The last few years accelerated it again. Generative models now produce text, images, audio, and code that most people cannot reliably tell apart from human work.

This has made AI a general-purpose technology in the sense that electricity and the internet are general-purpose. It does not automate one task; it reorganises how work, communication, medicine, education, and government are carried out. The money reflects that. Corporate investment in AI reached USD 252.3 billion in 2024, a 26 percent rise over the previous year, with private investment specifically climbing 44.5 percent (Stanford HAI, 2025). A year later the total had more than doubled again, to USD 581.7 billion (Stanford HAI, 2026). Adoption followed: 88 percent of organisations in McKinsey's global survey reported using AI in at least one business function in 2025, up from 78 percent the year before (McKinsey & Company, 2025).

The evidence this diffusion has produced cuts both ways. AI systems now match or beat specialists on some diagnostic tasks (Wang et al., 2024). They have delivered measurable productivity gains to firms positioned to use them (PwC, 2026), and have sped up work in chemistry, physics, and industrial automation. The same systems encode and amplify social bias (Buolamwini & Gebru, 2018; Obermeyer et al., 2019), displace particular categories of labour (Goldman Sachs Research, 2025), consume electricity at a rate the grid was not built for (International Energy Agency [IEA], 2025), and have made convincing fraud cheap (Deloitte Center for Financial Services, 2024).

This paper reviews the literature and the current statistical evidence to answer two related questions: how rapidly AI has actually developed, and by what measure; and what the principal advantages and

disadvantages of that development have been. Section 2 reviews the relevant literature. Section 3 documents the pace using investment, capability, and adoption data. Section 4 examines the gains across the economy, healthcare, science, and education. Section 5 examines the costs: labour disruption, algorithmic bias, misinformation, environmental burden, and gaps in governance. Section 6 draws out the distributional pattern running through both, and Section 7 concludes.

## 2 Literature Review

The academic study of AI begins with Turing's (1950) proposal of a behavioural test for machine intelligence, and became a formal field at the 1956 Dartmouth workshop. Its founding proposal made a striking claim: that “every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it” (McCarthy et al., 2006, p. 13, originally circulated 1955). The following six decades did not bear that out quickly. The field cycled through bursts of optimism and long “AI winters” in which funding and interest dried up as symbolic and rule-based systems failed to handle real-world complexity (Russell & Norvig, 2020).

The current period traces to the return of neural networks after 2012, when large labelled datasets and graphics-processing hardware let multi-layered networks outperform everything that came before on vision and speech tasks (LeCun et al., 2015). The transformer architecture (Vaswani et al., 2017) provided a second inflection, and underlies the language models driving the generative wave of the 2020s. Together these two developments explain why benchmark performance has improved in jumps rather than gradually (Stanford HAI, 2025, 2026).

A separate, more applied literature measures what AI does to firms and workers. McKinsey's annual State of AI survey (McKinsey & Company, 2025) and PwC's Global AI Jobs Barometer (PwC, 2026) track adoption, productivity, and labour-market effects across firms and occupations, while economic research groups such as Goldman Sachs Research (2025, 2026) model the scale of possible displacement. A consistent finding across this work is that adoption is running ahead of most firms' ability to redesign work around it. McKinsey reports that although 88 percent of organisations use AI somewhere in the business, roughly two-thirds have not begun scaling it across the enterprise at all.

A third strand concerns harms. Buolamwini and Gebru (2018) audited commercial facial-analysis systems and found error rates that varied enormously depending on who was being classified, establishing an empirical basis for concern about algorithmic bias. Obermeyer et al. (2019), publishing in *Science*, showed that a widely used US healthcare risk-prediction algorithm systematically understated the needs of Black patients because it used spending, rather than illness, as its target. More recent applied work documents the use of synthetic media in fraud (Deloitte Center for Financial Services, 2024) and the growing electricity demand of the data centres that train and run large models (IEA, 2025).

Finally, a policy and governance literature has grown alongside the technology. The European Union's AI Act (European Union, 2024) is the first comprehensive, risk-tiered framework of its kind and is being phased in through 2027, though the Commission has since proposed slowing parts of that schedule (European Commission, 2025). Public-opinion research from the Pew Research Center (2025, 2026a, 2026b) tracks a public that keeps adopting AI tools while growing steadily more uneasy about them. Read together, this body of work suggests that AI's development has outrun the institutional, regulatory, and organisational structures meant to manage it. That tension is what this paper takes up.

## 3 The Rapid Development of Artificial Intelligence

AI research spans nearly seven decades, but something changed in the early 2010s: capability stopped improving gradually. Table 1 shows three benchmarks documented by the Stanford AI Index, each of which was largely solved within a year of being introduced (Stanford HAI, 2025). SWE-bench is the clearest case: performance on real-world software engineering tasks went from 4.4 percent to 71.7 percent in twelve months.

**Table 1.** Year-over-Year Improvement on Selected AI Capability Benchmarks (2023–2024)

| Benchmark | Domain Tested                         | Improvement (percentage points) |
|-----------|---------------------------------------|---------------------------------|
| MMMU      | Multimodal, college-level reasoning   | +18.8                           |
| GPQA      | Graduate-level science Q&A            | +48.9                           |
| SWE-bench | Real-world software engineering tasks | +67.3                           |

**Source:** Stanford Institute for Human-Centered Artificial Intelligence (2025), The 2025 AI Index Report. Figures reflect the highest reported model score minus the prior year's highest reported score on each benchmark.

Investment tracked capability. Corporate AI investment totalled USD 252.3 billion in 2024, up 26 percent year over year. Private investment within that figure grew 44.5 percent, mergers and acquisitions rose 12.1 percent, and the number of newly funded generative-AI startups nearly tripled (Stanford HAI, 2025). The United States alone accounted for USD 109.1 billion of private AI investment, roughly twelve times China's USD 9.3 billion. The following year corporate investment reached USD 581.7 billion, a further 130 percent rise (Stanford HAI, 2026). Figure 1 sets this against the wider AI market, which Statista Market Insights (2026) projects will grow from roughly USD 255 billion in 2025 to more than USD 1.2 trillion by 2030. That projection deserves a caveat. Market-sizing estimates for AI vary by a factor of three across research firms for the same year, and Statista derives its figures from the funding raised by AI companies rather than from revenue, so the number describes investor appetite more than economic activity.

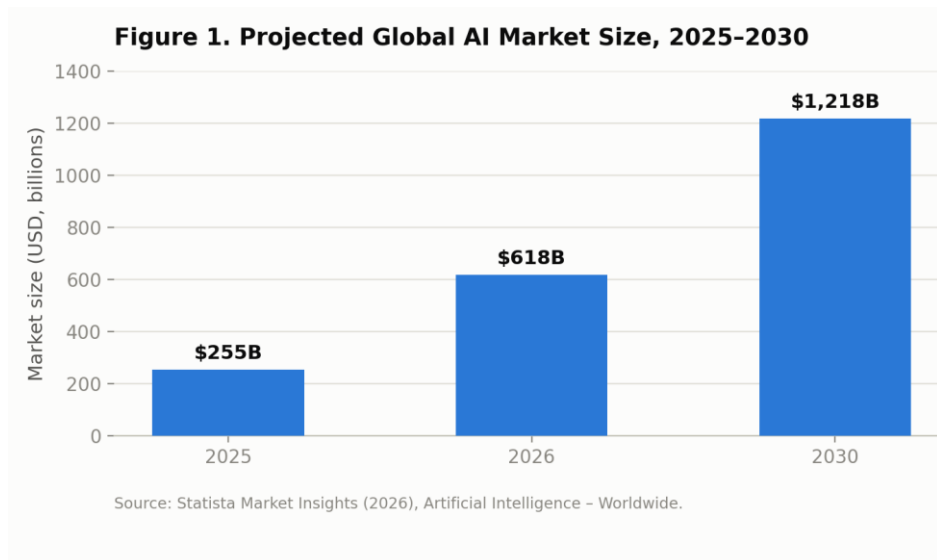


Figure 1. Projected global AI market size, 2025–2030 (USD, billions). Source: Statista Market Insights (2026).

Adoption at the level of individual organisations moved just as fast. Eighty-eight percent of organisations told McKinsey they used AI in at least one business function in 2025, against 78 percent a year earlier (McKinsey & Company, 2025; see Figure 4 in Section 4.1). Sixty-two percent were experimenting with or scaling autonomous “AI agents” able to complete multi-step tasks with limited supervision, but only 23 percent were scaling such systems in production, and in no single business function did more than about 10 percent report agents fully scaled. Deployment still trails experimentation by a wide margin. Physical automation grew alongside the software: roughly 541,000 industrial robots were installed worldwide in 2023, about three times the rate of a decade earlier (Stanford HAI, 2025). The investment, capability, and adoption data all point the same way. AI is not improving steadily; it is diffusing through the economy at a rate with few precedents, and, as the next two sections show, producing benefits and risks at the same time rather than in sequence.

## 4 Advantages of Artificial Intelligence

The spread of AI has produced benefits that can be measured. This section reviews five: economic productivity and enterprise value; healthcare and medicine; scientific discovery and industrial automation; education; and everyday convenience and accessibility.

### 4.1 Economic Growth, Productivity, and Enterprise Value

The clearest quantitative evidence of AI's economic benefit is firm-level productivity data. PwC's (2026) Global AI Jobs Barometer, built on more than a billion job advertisements across 27 countries, finds that companies in the most AI-exposed sectors recorded 34 percent labour productivity growth in 2025 relative to 2018, against 24 percent for the least exposed. The gap at the sector level is real but modest. The more striking pattern sits inside that group. The top fifth of the most AI-exposed companies achieved average productivity growth of 163 percent over the same period, nearly five times the figure for AI-exposed companies as a whole. Exposure to AI, in other words, predicts far less than what a firm actually does with it. Workers with AI skills also command a growing premium, which PwC puts at 62 percent in 2026, up from 57 percent the year before, ranging from 118 percent in consumer markets down to 16 percent in government and public-sector work.

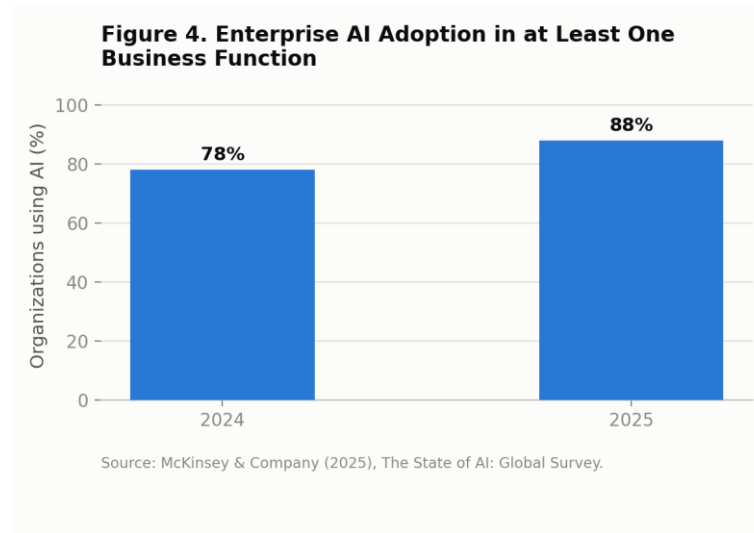


Figure 4. Enterprise AI adoption in at least one business function, 2024–2025. Source: McKinsey & Company (2025).

McKinsey & Company's (2025) global survey shows the same pattern at the organisational level. Eighty-eight percent of surveyed organisations use AI in at least one business function, but only 39 percent can attribute any measurable earnings impact to it, and only about 6 percent qualify as high performers deriving more than 5 percent of EBIT from AI. Those that do report gains in innovation capacity (64 percent), customer satisfaction, and competitive differentiation. Cost savings cluster in software engineering, manufacturing, and information technology; revenue gains cluster in marketing, sales, and product development. Employment in occupations requiring AI skills grew 69 percent against 9 percent for the labour market as a whole, roughly eight times faster (PwC, 2026), even as some routine entry-level work contracted over the same period.

### 4.2 Healthcare and Medicine

Healthcare is where AI's benefits are most directly visible in patient outcomes. For example, an Intelligent Health Care System for Detecting Drug Abuse in Social Media Platforms (Kanan, 2023).

US Food and Drug Administration approvals of AI- and machine-learning-enabled medical devices rose from 6 in 2015 to 223 in 2023, which reflects both a maturing technology and growing regulatory confidence in its clinical reliability (Stanford HAI, 2025).

**Table 2.** FDA-Approved AI/Machine-Learning-Enabled Medical Devices, Selected Years

| Year | Devices Approved | Approx. Growth Since 2015 |
|------|------------------|---------------------------|
| 2015 | 6                | —                         |
| 2023 | 223              | ≈ 37×                     |

**Source:** U.S. Food and Drug Administration data, as reported in Stanford Institute for Human-Centered Artificial Intelligence (2025), The 2025 AI Index Report.

Peer-reviewed clinical research has begun to document systems that match specialist-level diagnostic performance. Wang et al. (2024), writing in *Nature*, describe CHIEF, a pathology foundation model developed at Harvard Medical School. It was pretrained on 15 million unlabelled image tiles, then trained further on 60,530 whole-slide images spanning 19 anatomical sites, amounting to 44 terabytes of data. Tested on 19,491 independent whole-slide images drawn from 24 hospitals and cohorts worldwide, it reached roughly 94 percent accuracy in cancer detection across 15 datasets covering 11 cancer types, and outperformed existing deep learning methods by up to 36 percent. It could also predict patient prognosis and likely treatment response from tissue images alone. Systems like this do not replace pathologists. What they offer is the prospect of extending specialist-level diagnosis to places that have no subspecialist pathologist, and of flagging findings for expert review at a volume no human workforce could match unaided.

#### 4.3 Scientific Discovery and Industrial Automation

AI has also become a research instrument in its own right. Machine-learning models now search chemical and biological space for candidate drugs and materials, analyse large observational datasets in physics and climate science, and guide experimental design in ways that were not feasible at scale before. In manufacturing and logistics this has meant a fast expansion of physical automation: an estimated 541,000 industrial robots were installed worldwide in 2023 alone, roughly three times the rate of a decade earlier (Stanford HAI, 2025). AI-based perception and control have moved robotics beyond the fixed, pre-programmed tasks that previously defined it.

#### 4.4 Education and Personalized Learning

In education, AI-based tutoring and assessment systems are increasingly used to adapt content and pace to individual learners, automate routine grading and feedback, and extend instructional support beyond the classroom. Rigorous large-scale evidence on long-term learning outcomes is thinner here than in healthcare or economic productivity research. What is clear is that the capability improvements documented in Section 3, particularly in language understanding and reasoning benchmarks, have enabled a fast expansion of generative-AI tools for tutoring, writing support, and curriculum design (Stanford HAI, 2025). Uptake among students has outrun institutional response: four out of five US high school and college students now report using AI for schoolwork, while only half of middle and high schools have any AI policy at all and just 6 percent of teachers describe the policies they do have as clear (Stanford HAI, 2026). That gap raises open questions about academic integrity and appropriate use, discussed further in Section 5.

#### 4.5 Everyday Convenience and Accessibility

Outside enterprise and institutional settings, AI is embedded in tools people use without thinking about it: translation, voice assistants, navigation, content recommendation, and accessibility software for users with visual, auditory, or motor impairments. These rarely appear in economic productivity statistics, yet they are among the most widely experienced benefits of the technology.

The ability of ChatGPT, as an AI Tool, to produce outputs that mimic domain-specific roles, such as a lawyer, physician or engineer, allows for the realistic mimicry of expert-level dialogue (Hendawi, 2026).

Roughly half of US adults now use AI chatbots, up from a third in 2024, and the share rises to 66 percent among adults under 30 (Pew Research Center, 2026b). Familiarity has grown considerably faster than trust, a tension Section 5 takes up.

## 5 Disadvantages and Risks of Artificial Intelligence

The same pace of development that produced the benefits in Section 4 has produced a set of well-documented risks. This section reviews five: labour-market disruption; algorithmic bias and discrimination; misinformation and cybersecurity threats; environmental and energy costs; and gaps in ethical and regulatory governance.

### 5.1 Labor Market Disruption

The clearest economic risk is displacement. Goldman Sachs Research (2025) estimates that if current generative-AI use cases were adopted across the economy, about 2.5 percent of US employment would face displacement risk in the near term. Their baseline for the full transition, modelled over roughly a decade, is 6 to 7 percent of the workforce, or around 11 million jobs. They also project a temporary rise in unemployment of roughly half a percentage point during that transition, alongside a 15 percent increase in labour productivity in developed markets once generative AI is fully adopted; a March 2026 update revised the transitional unemployment estimate to 0.6 percentage points (Goldman Sachs Research, 2026). The occupations identified as most exposed are computer programmers, accountants and auditors, legal and administrative assistants, and customer-service representatives. One caveat belongs here. The same bank's economists reported in early 2026 that they still find no meaningful relationship between AI adoption and productivity at the economy-wide level, which sits awkwardly beside the firm-level results in Section 4.1 and is worth holding in mind before treating any of these projections as settled.

PwC's (2026) jobs barometer shows the disruption is uneven. Jobs requiring AI skills have grown roughly eight times faster than the labour market overall since 2019. But an analysis of 2.4 million US entry-level roles found that early-career positions with high AI exposure are now seven times more likely to demand skills traditionally associated with senior staff, such as judgement and leadership. These "seniorised" entry-level roles grew 35 percent since 2019, while other entry-level roles fell by about 10 percent. The bottom rung of the ladder is being removed for some workers and raised for others. McKinsey & Company (2025) adds that 51 percent of surveyed organisations have experienced at least one negative consequence from AI deployment, most often inaccurate output, and that the average organisation now actively manages four distinct AI-related risks, up from two in 2022. Risk awareness is rising. It is not rising as fast as deployment.

### 5.2 Algorithmic Bias and Discrimination

A substantial body of peer-reviewed research shows that AI systems inherit social bias rather than removing it. In an influential audit of three commercial facial-analysis systems, Buolamwini and Gebru (2018) found error rates reaching 34.7 percent for darker-skinned women against a maximum of 0.8 percent for lighter-skinned men. The worst-performing system showed a 34.4-percentage-point gap between those two groups. Aggregate accuracy looked respectable throughout, and that is precisely the point: a system can report a strong overall number while failing almost completely for a particular group. Biased training data is sufficient to produce that outcome on its own, without anyone intending it.

In healthcare, Obermeyer et al. (2019), publishing in *Science*, examined a risk-prediction algorithm used to manage care for millions of US patients. Black patients were assigned systematically lower risk scores than White patients with comparable levels of underlying illness. The cause was a single design choice: the algorithm used historical healthcare spending as a proxy for healthcare need, and less money had always been spent on Black patients with the same conditions. The effect was large. Correcting for it would have raised the share of Black patients automatically flagged for extra help from 17.7 percent to 46.5 percent. What happened next is the more useful part of the finding. The manufacturer confirmed the result on a national dataset of 3.5 million patients, then worked with the researchers on a replacement label combining cost predictions with predicted chronic conditions. That change reduced racial bias in the algorithm's outputs by 84 percent. The bias was not an inherent property of the model. It was the consequence of one identifiable decision, and it proved fixable once identified.

### 5.3 Misinformation, Fraud, and Cybersecurity

Generative AI has made convincing synthetic media cheap, which has opened new routes for fraud and misinformation. Roughly 500,000 video and voice deepfakes circulated online in 2023; by 2025 the figure was estimated at around 8 million. That estimate originates with the detection firm DeepMedia and has been widely reported since (as compiled in DeepStrike, 2025). Businesses lost an average of nearly USD 500,000 per deepfake-related incident in 2024. Looking forward, Deloitte's Center for Financial Services projects that generative-AI-enabled fraud losses in the United States will climb from USD 12.3 billion in 2023 to USD 40 billion by 2027, a compound annual growth rate of 32 percent (Deloitte Center for Financial Services, 2024). It is worth being precise about what that last figure covers: AI-enabled fraud in general, of which synthetic media is one component, not deepfakes alone.

**Table 3.** Selected Indicators of Synthetic-Media and Generative-AI Fraud



| Indicator                                          | 2023             | Projected 2025–2027   |
|----------------------------------------------------|------------------|-----------------------|
| Deepfake videos and voice clones shared online     | ≈ 500,000        | ≈ 8,000,000 (2025)    |
| Deepfake-enabled fraud attempts                    | +3,000% (YoY)    | Continued growth      |
| Projected U.S. losses, generative-AI-enabled fraud | USD 12.3 billion | USD 40 billion (2027) |

**Sources:** deepfake volume estimates originate with DeepMedia and are compiled in DeepStrike (2025), Deepfake Statistics 2025: The Data Behind the AI Fraud Wave; US fraud loss projections from Deloitte Center for Financial Services (2024). The Deloitte projection covers generative-AI-enabled fraud in general rather than synthetic media alone.

These are not hypothetical figures. In February 2024 a finance employee at the engineering firm Arup joined what appeared to be a routine video call with the company's chief financial officer and several colleagues. Every other participant on the call was an AI-generated deepfake. The employee approved 15 separate transfers totalling roughly USD 25 million before the fraud came to light. Human perception offers little defence here. Korshunov and Marcel (2020) found that high-quality manipulations fool most viewers, and, more troublingly, that people's confidence in their own judgement bears almost no relation to whether they are right. Automated detection is not a complete answer either: on the Deepfake-Eval-2024 benchmark, which uses real-world synthetic content rather than curated laboratory data, purpose-built detection tools lost 45 to 50 percent of the accuracy they showed during validation (as compiled in DeepStrike, 2025).

#### 5.4 Environmental and Energy Costs

The computational scale of modern AI has turned data-centre electricity into a material environmental concern. The International Energy Agency (2025) estimates that data centres consumed roughly 415 terawatt-hours (TWh) worldwide in 2024, about 1.5 percent of global electricity consumption, with the United States, China, and Europe accounting for 45, 25, and 15 percent of that total respectively. The growth rate matters more than the level. Data-centre electricity consumption has risen around 12 percent a year since 2017, more than four times faster than total global electricity demand.

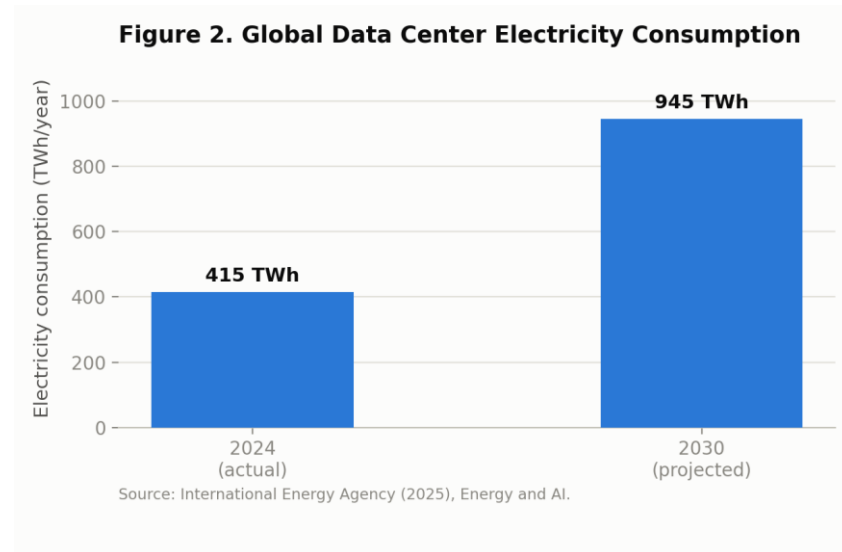


Figure 2. Global data center electricity consumption, actual 2024 and projected 2030 (terawatt-hours). Source: International Energy Agency (2025).

The IEA (2025) projects this will more than double to around 945 TWh by 2030, more than Japan consumes today, and reach roughly 1,200 TWh by 2035. In the United States, data centres are expected to account for nearly half of all electricity demand growth to 2030, and by the end of the decade the country is set to use more electricity for data centres than for producing aluminium, steel, cement, chemicals, and every other

energy-intensive good combined. AI's environmental footprint is not a fixed cost of the technology. It is a cost currently growing faster than the electricity system that has to absorb it, with direct consequences for grid capacity, electricity prices, and emissions in the regions where data centres concentrate.

### 5.5 Ethical, Safety, and Governance Challenges

These risks have developed against regulatory frameworks that are still under construction. The European Union's AI Act entered into force on 1 August 2024, the first comprehensive, risk-tiered legal framework for AI anywhere. It applies in stages. Prohibitions on certain uses and AI-literacy obligations took effect on 2 February 2025. Rules for general-purpose AI models, together with the governance bodies meant to oversee them, followed on 2 August 2025. The decisive date is 2 August 2026: most obligations for high-risk systems listed in Annex III apply from then, transparency requirements begin, and enforcement and financial penalties start. High-risk AI embedded in already-regulated products such as medical devices and machinery has until 2 August 2027, and AI components of certain large-scale EU IT systems have until the end of 2030 (European Union, 2024). Even this schedule is now in question. In November 2025 the Commission proposed, through its Digital Omnibus package, tying the high-risk rules to the availability of harmonised technical standards, which would delay them further (European Commission, 2025). The regulator is asking for more time to regulate. Outside the European Union, approaches remain considerably more fragmented, and enforcement capacity in most jurisdictions lags well behind the deployment described in Section 3.

Public opinion reflects the gap, though not as simply as it is often described. Americans are almost evenly split on whether they trust their own government to regulate AI effectively: 44 percent express a lot or some trust against 47 percent with little or none, and the split runs along party lines, with 54 percent of Republicans and Republican leaners trusting the US to regulate AI compared with 36 percent of Democrats and Democratic leaners (Pew Research Center, 2025). A more recent survey of 5,119 US adults puts it more starkly: 67 percent report little or no confidence in the federal government to regulate AI effectively, up from 62 percent in 2024, and 63 percent say the technology is advancing too fast (Pew Research Center, 2026b). Internationally, however, the picture is not one of uniform distrust. Across 25 countries, a median of 53 percent trust the European Union to regulate AI, against 37 percent for the United States and 27 percent for China, and most people trust their own government more than any of the three (Pew Research Center, 2025). Skepticism about AI governance is real, but it is distributed unevenly, and the EU's regulatory reputation is at present its strongest asset. As Figure 3 shows, the balance of American sentiment tilts clearly toward unease: half of US adults say the growing use of AI in daily life makes them more concerned than excited, against 10 percent who are more excited than concerned (Pew Research Center, 2026a, reporting a June 2025 survey).

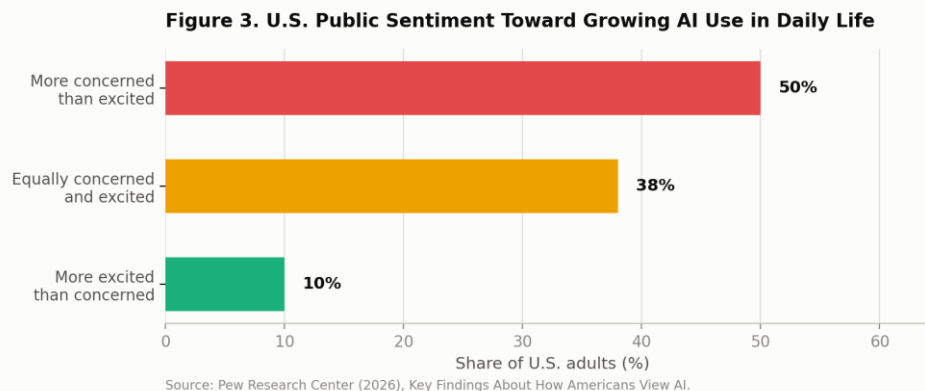


Figure 3. U.S. public sentiment toward growing AI use in daily life. Source: Pew Research Center (2026).

Rapid capability growth, uneven regulatory coverage, and public unease together constitute a distinct category of risk. Even AI applications that are technically reliable and economically beneficial may be adopted, contested, or restricted in ways that are hard to predict, so long as the institutions responsible for oversight keep trailing the technology they oversee.



## 6 Discussion: Balancing Innovation and Responsibility

The evidence reviewed in Sections 3 through 5 supports three broader observations. First, the pace has been genuinely unusual by historical standards. Benchmark performance has improved by tens of percentage points within single years (Stanford HAI, 2025), corporate investment more than doubled year over year (Stanford HAI, 2026), and enterprise adoption climbed from 78 to 88 percent of organisations in twelve months (McKinsey & Company, 2025). Generative AI reached 53 percent population adoption within three years, faster than either the personal computer or the internet (Stanford HAI, 2026). No comparably general technology has spread this quickly at this scale.

Second, the benefits and risks are not evenly distributed. Productivity and wage gains concentrate among firms and workers already positioned to exploit AI. McKinsey (2025) finds only about a third of adopting organisations have progressed from experimentation to enterprise-wide scaling, and only about 6 percent capture significant value. PwC (2026) finds the top-performing fifth of AI-exposed firms achieving productivity growth nearly five times that of AI-exposed firms generally. Meanwhile displacement, algorithmic bias, and exposure to synthetic-media fraud fall disproportionately on workers in routine occupations, on groups that historical data collection has always undercounted, and on individuals and institutions with the fewest resources to verify what is authentic. Concentrated gains alongside diffuse harms is arguably the central policy problem the technology poses, more so than any single capability or failure mode.

Third, the institutional response has lagged the commercial one. The EU AI Act, the most comprehensive framework currently in force, does not begin real enforcement until August 2026, with high-risk product rules following in 2027 and some deadlines running to 2030, and the Commission has already proposed slowing that timetable (European Union, 2024; European Commission, 2025). Public confidence in regulatory capacity is low in the United States and falling (Pew Research Center, 2026b), though notably not uniform across jurisdictions: the EU retains substantially more public trust as a regulator than either the US or China (Pew Research Center, 2025). Closing this gap does not require halting AI development, which the scale of ongoing investment makes unlikely in any case. It requires that organisations deploying AI systems, and the regulators overseeing them, treat bias auditing, energy accounting, and fraud-resistant verification as standard components of deployment rather than optional add-ons. The healthcare literature offers a precedent worth taking seriously. Once Obermeyer et al. (2019) identified the specific design flaw responsible for racial bias in a widely used algorithm, a targeted correction cut that bias by 84 percent. Many of AI's documented harms are addressable in principle, even where the underlying rate of technological change stays fast.

## 7 Conclusion

AI has developed at a pace with few precedents in the history of technology, and the benefits are real and, in several domains, substantial: expanded diagnostic reach in medicine, faster scientific work, and productivity gains for firms and workers able to adopt the technology well. These are documented across multiple independent sources, from peer-reviewed clinical research to industry surveys covering thousands of organisations. The same body of evidence documents concrete, quantifiable harms: displacement risk concentrated in specific occupations, algorithmic bias with measurable disparate impact across demographic groups, a rapidly growing volume of AI-enabled fraud, a data-centre energy footprint expanding several times faster than global electricity demand, and a regulatory apparatus that remains incomplete relative to the scale of deployment already underway, despite recent progress such as the EU AI Act.

The central conclusion of this review is that AI's development is not one trend with one balance of costs and benefits. It is several trends moving at different speeds. Technical capability and commercial deployment have moved fastest. Regulatory and institutional capacity has moved slowest, and in the European Union's case is now being asked to slow further still. Public trust has moved in the opposite direction from adoption, growing more cautious even as usage grows more common (Pew Research Center, 2026b). Narrowing these gaps, through bias auditing of the kind Obermeyer et al. (2019) demonstrated, energy-efficiency standards for AI infrastructure, fraud-resistant identity verification, and regulatory frameworks that keep pace with deployment rather than following it by several years, will largely determine whether the coming decade of AI resembles the broadly shared productivity gains of earlier general-purpose technologies or a narrower pattern in which the benefits and the harms accrue to different people.

## 8 References

1. Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. *Proceedings of Machine Learning Research*, 81, 77–91.
2. DeepStrike. (2025). Deepfake statistics 2025: The data behind the AI fraud wave. <https://deepstrike.io/blog/deepfake-statistics-2025>
3. Deloitte Center for Financial Services. (2024). Deepfake banking fraud risk on the rise. Deloitte Insights. <https://www.deloitte.com/us/en/insights/industry/financial-services/deepfake-banking-fraud-risk-on-the-rise.html>
4. European Commission. (2025). Digital omnibus package: Proposed amendments to Regulation (EU) 2024/1689. European Commission.
5. European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union.
6. Goldman Sachs Research. (2025). How will AI affect the global workforce? Goldman Sachs. <https://www.goldmansachs.com/insights/articles/how-will-ai-affect-the-global-workforce>
7. Goldman Sachs Research. (2026). How will AI affect the US labor market? Goldman Sachs. <https://www.goldmansachs.com/insights/articles/how-will-ai-affect-the-us-labor-market>
8. Hendawi, Samar,T., Kanan, Elbes, Mohammed, Mughaid, Ala, AlZu'bi, Shadi (2026). Automated prompt engineering pipelines: fine-tuning LLMs for enhanced response accuracy. *Expert Systems with Applications*, Volume 298, Part B, 2026, 129689, ISSN 0957-4174,
9. International Energy Agency. (2025). Energy and AI. IEA. <https://www.iea.org/reports/energy-and-ai/executive-summary>
10. T. Kanan et al., "An Intelligent Health Care System for Detecting Drug Abuse in Social Media Platforms Based on Low Resource Language," in *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 32, pp. 691-703, 2024, doi: 10.1109/TASLP.2023.3294699.
11. Korshunov, P., & Marcel, S. (2020). Deepfake detection: Humans vs. machines (arXiv:2009.03155). arXiv. <https://arxiv.org/abs/2009.03155>
12. LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444. <https://doi.org/10.1038/nature14539>
13. McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A proposal for the Dartmouth summer research project on artificial intelligence, August 31, 1955. *AI Magazine*, 27(4), 12–14. (Original work published 1955)
14. McKinsey & Company. (2025). The state of AI in 2025: Agents, innovation, and transformation. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
15. Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
16. Pew Research Center. (2025). How people around the world view AI. <https://www.pewresearch.org/global/2025/10/15/how-people-around-the-world-view-ai/>
17. Pew Research Center. (2026a). Key findings about how Americans view artificial intelligence. <https://www.pewresearch.org/short-reads/2026/03/12/key-findings-about-how-americans-view-artificial-intelligence/>
18. Pew Research Center. (2026b). Americans and AI 2026: Chatbots, smart devices and views on impact. <https://www.pewresearch.org/internet/2026/06/17/americans-and-ai-2026-chatbots-smart-devices-and-views-on-impact/>
19. PwC. (2026). 2026 global AI jobs barometer. <https://www.pwc.com/gx/en/services/ai/ai-jobs-barometer.html>
20. Russell, S., & Norvig, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
21. Stanford Institute for Human-Centered Artificial Intelligence. (2025). The 2025 AI index report. Stanford University. <https://hai.stanford.edu/ai-index/2025-ai-index-report>
22. Stanford Institute for Human-Centered Artificial Intelligence. (2026). The 2026 AI index report. Stanford University. <https://hai.stanford.edu/ai-index/2026-ai-index-report>
23. Statista. (2026). Artificial intelligence – Worldwide. Statista Market Insights. <https://www.statista.com/outlook/tmo/artificial-intelligence/worldwide/>
24. Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.

<https://doi.org/10.1093/mind/LIX.236.433>

25. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
26. Wang, X., Zhao, J., Marostica, E., Yuan, W., Jin, J., Zhang, J., Li, R., Tang, H., Wang, K., Li, Y., Wang, F., Peng, Y., Zhu, J., Zhang, B., Gil, D. B., Han, X., Fisher, D. E., Lian, C. G., Farahani, N., ... Yu, K.-H. (2024). A pathology foundation model for cancer diagnosis and prognosis prediction. *Nature*, 634, 970–978. <https://doi.org/10.1038/s41586-024-07894-z>