



● Responsible AI in practice

2025 global insights
from the AI Company
Data Initiative



Thomson Reuters
Foundation



unesco

Cover image credit: REUTERS/Victor Fraile

Published in 2026 by the Thomson Reuters Foundation, 5 Canada Square, London, E14 5AQ, United Kingdom and the United Nations Educational, Scientific and Cultural Organization (UNESCO), 7, place de Fontenoy, 75352 Paris 07 SP, France.

© Thomson Reuters Foundation and UNESCO, 2026

ISBN 978-92-3-100863-4

DOI: <https://doi.org/10.54678/YJWP8855>



This publication is available in Open Access under the Attribution-ShareAlike 3.0 IGO (CC-BY-SA 3.0 IGO) license (<https://creativecommons.org/licenses/by-sa/3.0/igo/>).

By using the content of this publication, the users accept to be bound by the terms of use of the UNESCO Open Access Repository (<https://www.unesco.org/en/open-access/cc-sa>).

The designations employed and the presentation of material throughout this publication do not imply the expression of any opinion whatsoever on the part of the Thomson Reuters Foundation or UNESCO concerning the legal status of any country, territory, city or area or of its authorities, or concerning the delimitation of its frontiers or boundaries.

The ideas and opinions expressed in this publication are not necessarily those of the Thomson Reuters Foundation or UNESCO and do not commit them.

SHS/REI/BIO/2026/AICDI

Contents

4 Forewords

- 5 Thomson Reuters Foundation
- 7 UNESCO
- 8 An investor perspective

9 Executive summary

- 10 Key Statistics
- 11 Executive summary

13 Introduction

- 14 Why responsible AI matters now
- 16 The AI Company Data Initiative
- 18 The regulatory and framework landscape
- 23 Methodology
- 26 AICDI dataset breakdown

28 Thematic findings

- 29 **Finding 1:** Public commitment from companies to AI governance frameworks remains low, signalling inconsistent governance practices
- 34 **Finding 2:** Many companies publish strategies on AI but it is less clear how these are put into practice
- 42 **Finding 3:** Companies do not demonstrate adequate protections for workers as AI reshapes jobs
- 52 **Finding 4:** Ethical issues – including human rights and environmental impacts – are being sidelined in AI governance and risk management
- 56 **Finding 5:** Limited company policies on AI training data, third-party data controls, and user data rights
- 63 **Sentiment analysis**

69 AICDI company case studies

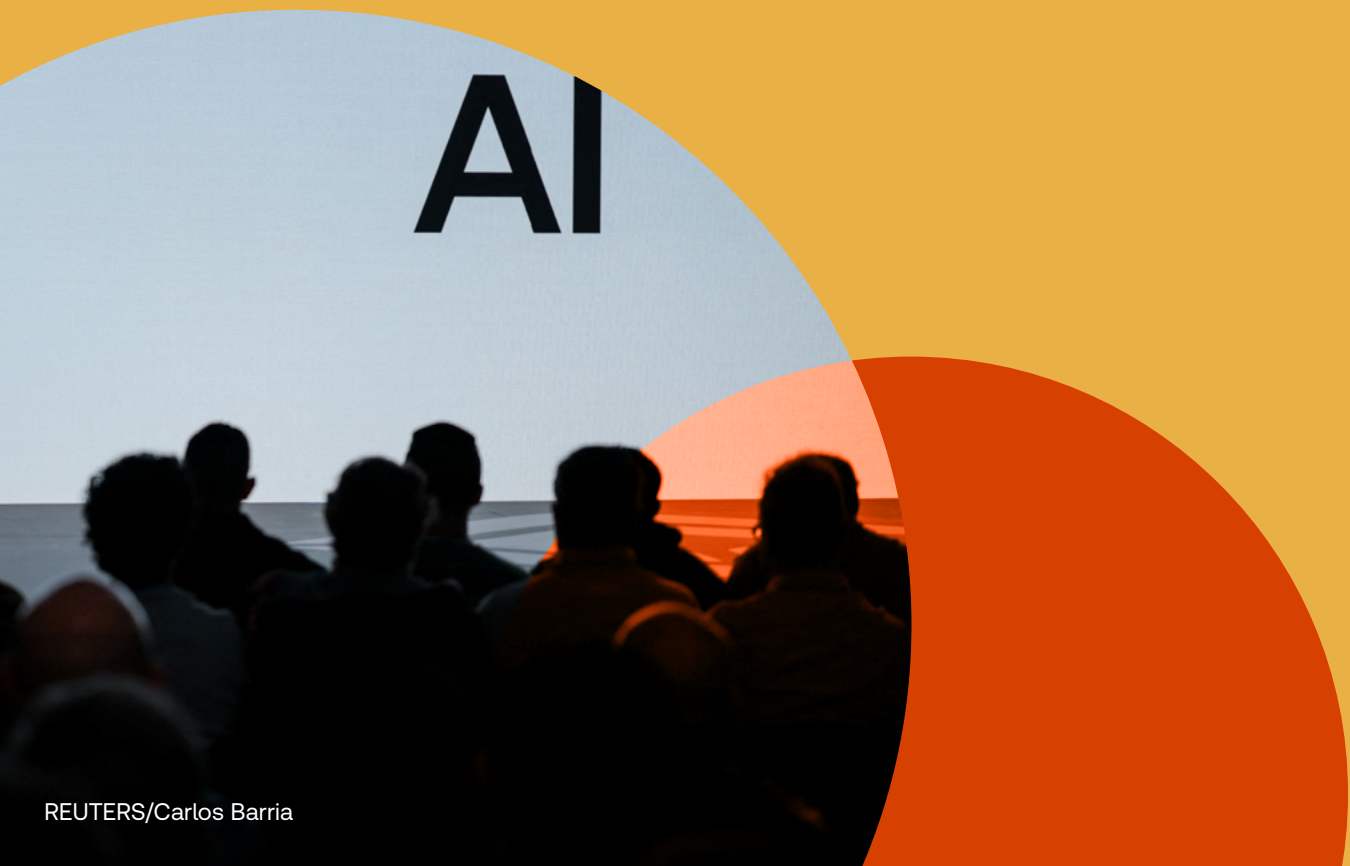
- 71 **Case study 1:** TELUS
Diversity and inclusion
- 72 **Case study 2:** Vodafone
Data, systems and cybersecurity
- 73 **Case study 3:** SAP
Impact on workers
- 74 **Case study 4:** Telefónica
AI governance, strategic & institutional
- 75 **Case study 5:** Banco Bradesco, BASF, Infosys, Telefónica, TELUS, Prudential
AI Skills Training
- 77 **Case study 6:** Gruppo TIM, BASF, TELUS, Telefónica
Environmental considerations
- 79 **Case study 7:** Banco Bradesco, SAP
Review of AI governance mechanisms
- 80 **Case study 8:** Cementos Argos, BASF
Workers' rights

81 Guidance for investors

- 82 Implications
- 83 Investor engagement checklist
- 87 Responsible AI principles for proxy voting
- 89 Investor Case Study: ESG-AM

91 Endnotes

Forewords

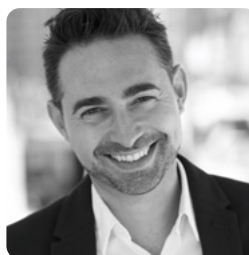
A large graphic at the bottom of the page. It features a light blue semi-circle on the left containing the letters 'AI' in a large, bold, black font. To the right of this is a large orange semi-circle. In the foreground, the silhouettes of a group of people are shown from behind, looking towards the 'AI' text. The background of the entire page is a solid orange color.

AI

REUTERS/Amr Alfiky



Thomson Reuters Foundation



Antonio Zappulla

CEO, Thomson Reuters
Foundation

Artificial Intelligence is reshaping economies and societies at a speed few anticipated. Our AI Company Data Initiative (AICDI) was set up to ensure transparency, trust and responsibility are embedded into this seismic transformation – helping companies, investors, policy makers and society understand how AI is being adopted and governed at a corporate level through a mechanism of voluntary disclosure.

Businesses of different sizes and across multiple sectors are adopting Artificial Intelligence technology at pace, embedding new AI tools and capabilities across their operations, products and services, whilst internal controls, reporting mechanisms, and socio-environmental impacts often lag.

In today's evolving regulatory landscape, AI adoption is clearly outpacing governance. At a corporate level, we have witnessed a lack of systems and frameworks required to ensure the technology is rolled out in an ethical way. This lack of clarity creates an unprecedented level of risks for businesses, workers, consumers, and investors.

For the private sector, the risks are immense: operational failures, allegations of bias and discrimination, privacy and security incidents, environmental impact, and reputational harm.

At the Thomson Reuters Foundation, we work with businesses and investors who are committed to the strongest ethical standards. Our wide network includes companies with a combined market capitalisation of more than \$25 trillion and investors with \$5.5 trillion of assets under management.

Through AICDI, we equip the private sector with the insights it needs to harness the opportunities of AI responsibly in a way that minimises risk and drives long-term value, strengthening governance, and improving disclosure in ways that anticipate, rather than merely react to, emerging expectations. The data emerging from our survey is shared with investors, becoming a valuable and comparable asset tracking how companies are governing AI, allowing them to make better-informed investment decisions that reduce financial and reputational risk.

Grounded in UNESCO's Recommendation on the Ethics of AI — a leading global standard for how AI should be used responsibly, adopted by all UNESCO member states— AICDI was built by a team at the Thomson Reuters Foundation and our partners at UNESCO with input from corporate stakeholders.

This first report shares the findings from our pilot year. What you are about to read is the world's largest study assessing corporate AI adoption globally, featuring 100,000 data points from almost 3,000 companies across 11 sectors and five geographies.

The data emerging from our pilot year confirms both progress and challenges. AI adoption is accelerating, but disclosure and governance maturity remain uneven. Across regions and sectors, many companies are experimenting with AI faster than they are formalising accountability, skills and oversight. Encouragingly, we also see examples of leadership – firms beginning to

embed responsible AI roles, board engagement and clearer policies.

The AICDI survey is not intended to be a one-off snapshot. It is a tool for continuous improvement: giving companies a structured way to take stock of their AI use, clarify accountability, and improve decision-making; and providing investors and other stakeholders with more consistent signals of governance maturity. Our ambition is a healthier AI ecosystem – one in which innovation can move forward with stronger safeguards, clearer expectations, and greater public confidence.

Beyond that, AICDI is designed to be directly usable in investment decision-making and stewardship. It gives investors a globally comparable way to benchmark issuers' AI governance and adoption against sectoral and regional peers, helping distinguish leaders from laggards as responsible AI becomes a differentiator. By grounding assessment in UNESCO's Recommendation on the Ethics of AI and other emerging norms, the dataset also helps investors anticipate where regulatory and liability pressures may concentrate, supporting due diligence, risk management, and more confident alignment with client and stakeholder expectations.

I have been encouraged by the early engagement from companies, investors, and partners who recognise that responsible AI is now a core governance issue, not a niche technical topic. This first report is a starting point. We invite all stakeholders to use AICDI as a shared foundation to improve transparency, comparability, and ultimately to ensure the benefits of AI are shared more widely, responsibly and in the public interest.

Thank you for your continued support.

Antonio Zappulla

UNESCO



Lidia Arthur Brito

Assistant Director-General
for Natural Sciences and
Social and Human Sciences
A.I., UNESCO

We are currently at a critical juncture in technological history. The era of abstract safety debates regarding artificial intelligence is over. The world is no longer simply asking how AI-based systems can impact our society; it is demanding concrete mechanisms ensuring that these technologies do so in a manner beneficial to all of humanity. We need systems that ensure multilingual access, deliver tangible benefits to diverse communities and operate both safely and ethically.

This demand arrives at a time of profound acceleration of technological advancement. We are witnessing a rapid shift from generative tools creating content to agentic and embodied systems capable of taking action. Simultaneously, AI adoption is drastically outpacing formal regulation. Across different states, rules are still being written and the maturity of the legislative AI governance landscape varies.

This reality places a profound responsibility on the private sector to more proactively adopt responsible AI practices. Industry leaders and investors cannot afford to wait for legislation to catch up. Our shared task is to ensure this technological transition unfolds safely and at the pace of trust. To achieve this, the industry must champion ethical deployment, acting voluntarily and transparently to ensure the speed of innovation does not outrun society's confidence.

This is precisely where UNESCO's Recommendation on the Ethics of AI proves its enduring value. As a leading global standard adopted by all UNESCO's Member States, the Recommendation serves as a global compass for responsible AI development, design, and deployment. It provides a unified, values-based framework that guides all stakeholders including businesses in balancing between the benefits of technology against the associated risks.

As the findings in this pivotal report demonstrate, the business community is the implementation engine for these principles. Grounded in publicly disclosed data and informed by UNESCO's Recommendation,

this report offers a vital window into how AI is actually being adopted and governed at the corporate level.

The data confirms both the immense momentum of AI adoption and the stark reality of the current governance gap. Much like the broader ecosystem, companies across industries are embedding new AI capabilities faster than they are formalising accountability, internal controls, and oversight. This lag creates risks, ranging from operational failures and embedded biases to severe reputational harm.

Yet, this report is also a testament to what is possible. It highlights case studies on how pioneering organisations are moving beyond abstract principles to share what governance, accountability, and risk management look like in practice, making responsible AI repeatable and scalable.

Crucially, implementing responsible AI practices requires resources. While multinational corporations have the capacity to build dedicated ethics teams and compliance frameworks, startups and small-to-medium enterprises (SMEs) often face steep constraints. Meaningful industry leadership means that larger, well-resourced companies must pave the way. By sharing best practices, open-sourcing governance tools, and actively supporting smaller players across the globe, major enterprises can ensure that robust ethical standards are accessible to all, rather than a privilege of the few.

The AICDI is a critical tool for ongoing improvement. It provides companies with a structured way to assess their own practices against a global standard, while giving investors the transparent, comparable data they need to differentiate leaders from those who fall behind and mitigate risks.

One of the platforms developed by UNESCO to support this shared journey, is the UNESCO Business Council for Ethics of AI, a collaboration with pioneering companies who committed to operationalise the principles laid out in the UNESCO Recommendation. The Council stands as a vital platform for companies to exchange real-world experiences and spread these good practices.

I invite the members of the Council and all businesses, investors, and policymakers to use the insights of this report, accompanied by the Recommendation on the Ethics of AI as the guiding philosophy, to raise the floor across industries, building a healthier AI ecosystem dedicated to serving the global public interest.

Lidia Arthur Brito

An investor perspective



Eva Cairns

Head of Responsible
Investment,
Scottish Widows

AI will shape the next era of corporate strategy, economic growth and market transformation. In the UK alone 75 per cent of financial firms are already using AI and a further 10 per cent are planning to use it in future.¹

AI brings the prospect of increased efficiency and innovation, as well as the opportunity to help tackle big challenges in areas such as sustainability. Every company in our portfolio is now likely to be adopting AI in some form and using it to try to achieve a competitive advantage at pace. With such rapid and widespread adoption, AI is quickly becoming a core governance challenge that can present material risks to companies and investors.

That's why AI governance is a key engagement theme for us, and we shared our thoughts in our 2025 report, *Governing the Algorithm: investor principles for responsible AI*². We have been working with our asset managers to explore how AI oversight can be more effectively embedded into responsible investment practice.

As members of the Thomson Reuters Foundation's Workforce Disclosure Initiative (WDI), and supporters of its AI Company Data Initiative (AICDI), we share its goals to promote transparency and responsible business practices in corporate AI adoption. Its detailed dataset (which covers almost 3,000 companies) and insights provide useful context to our own stewardship activities, encouraging responsible AI practices within our portfolio companies.

We believe that pension funds have a role to play in encouraging more transparent, equitable, and future-fit corporate behaviour, to support long-term economic stability, inclusion and accountability. The AICDI and its tools provide a useful resource in support of these goals.

Eva Cairns

About the AICDI investor signatory group

The AI Company Data Initiative is supported by a group of 10 investors with AUM of \$1.2 trillion. The investor signatory group works collectively to support the development of this initiative by sharing their perspective on responsible AI and by encouraging their portfolio companies' engagement with the initiative with active outreach. They are organisations that are committed to promoting the responsible use of AI within organisations around the world.

Through the AICDI, investor signatories gain a structured and comparable view of how issuers are governing and deploying AI. This oversight is grounded in AICDI's global dataset of 3,000 companies which includes more than 100,000 data points. Investors can nominate up to 100 portfolio companies for inclusion in the dataset and analysis. Using these insights, investors will be equipped to engage corporates using a UNESCO-aligned

framework that spans across governance, oversight, human capital, and safety and security.

By joining the signatory group, investors strengthen their due diligence and engagement on specific protections, namely data security, accountability stemming from AI model registries, incident reporting, access to redress, as well as bias and fairness assessments. It also enables benchmarking against peers and helps identify both leaders and areas across sectors where targeted improvement is required. Where holdings overlap, there is the opportunity to collaborate with other investors on joint engagements.

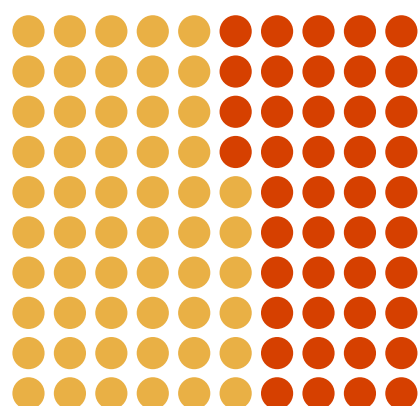
To support the AICDI's work or to find out more about the benefits of investor membership, please contact us at AICDI@thomsonreuters.com

Executive summary



KEY STATISTICS

Across 2,972 companies:



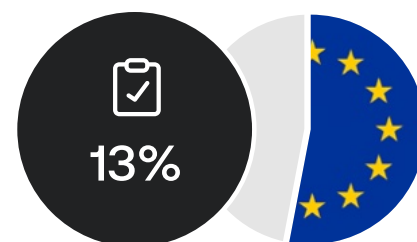
43.7%

publicly communicate
having an AI strategy

However, 76 per cent of these companies show no evidence of having policies to evaluate the quality of data used to train AI systems



13 per cent of companies say they align their strategy with a formal AI governance framework



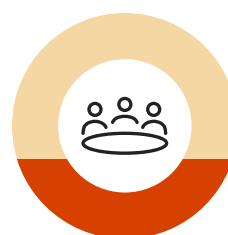
53 per cent of these companies cite the EU AI Act as their reference framework



12%

31%

While 31 per cent of companies claimed to have AI training programmes, only 12 per cent offered structured training with comprehensive coverage



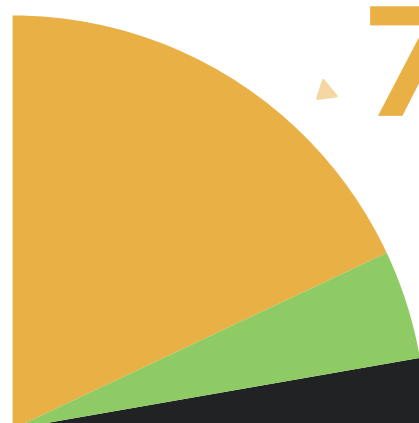
40%

report having board-
or committee-level
oversight on AI



12.4%

report having a policy
to ensure a human
oversees AI systems



72%

of companies do not report conducting any impact assessment with regard to AI. Of the companies that do:



11 per cent report carrying out
Environmental Impact Assessments



7 per cent report carrying out
Human Rights Impact Assessments

Executive summary

Artificial intelligence (AI) is rapidly being embedded across companies' products, services and internal operations, yet governance and disclosure are not evolving at the same speed.

The result is a widening information gap for investors: companies increasingly communicate ambition, principles and oversight structures, while offering less clarity on where AI is actually deployed, how risks are controlled in practice, and who is accountable when systems fail.

As AI becomes operational rather than experimental, this gap creates exposure across operational resilience, bias and safety, workforce impacts, gender equality impacts, data dependency and environmental footprint, and reinforces the need for disclosure that connects strategy to real-world decision making. Equally, weak policies and contracts for sharing data with third-party AI providers create governance blind spots in multi-actor supply chains, where upstream data and design choices can drive downstream harms.

This report looks at corporate practice in the context of the emerging responsible AI regulatory landscape and analyses publicly available data collected by the Thomson Reuters Foundation's AI Company Data Initiative, the largest global dataset of corporate responsible AI disclosures.



Governance is described in conceptual terms

Across the dataset, a consistent pattern emerges: organisations tend to describe governance at a conceptual level, but more rarely demonstrate how it functions day to day on an operational level across the lifecycle of systems. References to policies, committees and high-level oversight appear more frequently than evidence of operational controls, dedicated resources, escalation pathways or monitoring mechanisms that would allow external stakeholders to understand how risks are managed once AI is deployed.

Limited transparency on the impact of AI on workers and the environment

The same dynamic extends beyond core governance. Companies often acknowledge workforce transition and skills development, yet disclosures seldom show how these programmes affect worker learning outcomes or how workers' concerns can be raised and addressed.

Transparency around training data, third-party systems and user rights remains uneven, making accountability across vendor ecosystems difficult to trace.

Ethical and environmental considerations are commonly framed as principles, but only occasionally linked to measurable processes or management practices.

Moving from awareness to accountability

Taken together, the findings suggest that the central challenge of responsible AI is no longer awareness but operationalisation. For investors, this shifts the focus of engagement: the most decision-useful information is not the presence of strategy, but evidence that governance works under real conditions, how models are approved, monitored, corrected and, when necessary, withdrawn. The opportunity for stewardship therefore lies in moving disclosure from statements of intent to verifiable practice.

As AI becomes part of core business infrastructure, understanding how companies control it will increasingly resemble understanding how they control capital, safety or financial reporting, not a specialised sustainability topic, but a foundation of corporate reliability.

As privately developed or deployed AI systems shape more of daily life, transparency must move beyond technical descriptions to show how accountability works—who makes decisions, how issues are escalated, and what remediation paths exist when things go wrong. Clear responsibility for harms or breaches should be identifiable in practice, not just in principle. Just as we expect openness and accountability from government, the private sector must meet comparable transparency standards for AI that affects the public.

Katie Fowler, Director of Responsible Business at the Thomson Reuters Foundation said:

"The findings suggest that the challenge of responsible AI is no longer awareness but ensuring principles translate into in practice. Our AI Company Data Initiative provides a comparable, actionable dataset so investors and companies alike can identify good practice and material risk."

• Introduction



Why responsible AI matters now

Artificial intelligence is moving from experimentation to a core part of day-to-day business infrastructure.

Research from McKinsey has found that AI adoption in at least one business function amongst companies rose to **88 per cent** in 2025, **up from 78 per cent** the year before.³

2024 

2025 

Across sectors, companies are integrating AI into customer-facing services, internal workflows, and decision-making processes. As this integration continues, the question for many stakeholders is no longer whether a company is “using AI,” but how consistently it understands where AI sits in its business, what safeguards are in place, and whether accountability is clear when things go wrong. PwC’s 2025 Global Investor Survey underlines this growing demand for transparency, with 42 per cent of investors saying they want more transparency on companies’ AI investments, and another 42 per cent wanting clearer information on AI returns and cost savings.⁴

Responsible AI adoption matters because the same systems that can improve speed, cost, and personalisation can also create new, scaled risks, often silently and unevenly, when governance doesn’t keep pace. AI can amplify bias in hiring or credit decisions, leak sensitive personal or corporate data, generate plausible but incorrect outputs that shape high-stakes choices, and introduce security vulnerabilities

or regulatory exposure through opaque third-party tools and shadow deployments. These risks are not abstract: they translate into real-world harms (to people), material liabilities (for organisations), and operational fragility (for systems that become dependent on AI without clear controls). In practice, responsible adoption is what turns AI from a short-term productivity lever into a sustainable capability, by ensuring systems are designed and deployed with clear purpose, human oversight, data protection, monitoring, and documented accountability.

Responsible AI is important for everyone, though it affects stakeholders in different ways. For businesses, it underpins effective oversight and reduces legal and reputational risk; for investors, it provides the transparency needed to judge governance quality and long-term resilience. Regulators depend on clear, well-documented AI use to protect rights and maintain market integrity, while suppliers and partners share in the risks created across interconnected systems. Customers need assurance that AI is fair, transparent and respects their privacy, and employees require safeguards against intrusive monitoring or biased automation.

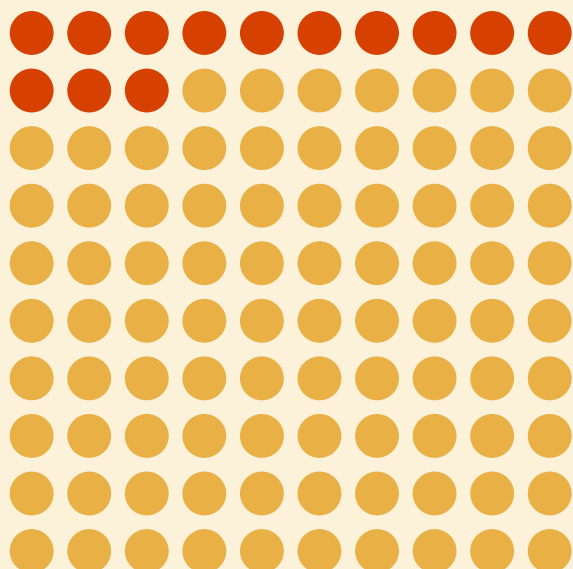


Communities, particularly those with fewer resources, benefit when AI is deployed responsibly and face greater harm when it is not. Together, these perspectives show that responsible AI is essential for trust, safety and equity across society.

While public attention often concentrates on the developers of frontier models, many of the most immediate operational, legal, and social impacts arise from how AI is deployed and managed by companies across the wider economy. As adoption expands beyond the technology sector and into complex global supply chains (spanning from mass multinationals to small Global South organisations), the absence of standardised information on AI use and governance becomes a practical barrier to effective oversight, investment analysis, and internal improvement.

There is concerning evidence to suggest that AI governance is lagging behind AI deployment. A Harvard Law School finds that only 13 per cent of S&P 500 companies have directors with AI expertise,⁵ and a Deloitte survey shows that just 14 per cent of boards regularly discuss AI.⁶

13% of S&P 500 companies have directors with AI expertise





The AI Company Data Initiative

Launched in 2024, the AI Company Data Initiative (AICDI) seeks to address this gap by giving organisations a practical, standardised way to understand, and communicate how AI is being used and governed across their operations.

Its primary tool is a free framework developed by the Thomson Reuters Foundation in partnership with UNESCO (United Nations Educational, Scientific and Cultural Organization) to help organisations understand their responsible AI adoption journey. Grounded in UNESCO's Recommendation on the Ethics of AI, the Initiative through the voluntary survey and other capacity building materials helps organisations map where AI systems sit in the business, assess the safeguards and oversight in place, and benchmark their maturity over time and against peers.

The AICDI aims to make AI governance more visible and comparable for internal leaders, boards, investors, and regulators, supporting better risk management, clearer accountability when things go wrong, and more responsible adoption as AI use scales through supply chains and across sectors. Designed to be accessible at any level of maturity, the survey does not require technical expertise to complete and can serve as a starting point for organisations beginning their AI journey as well as a structured health-check for more advanced adopters.

AICDI captures a broad, decision-relevant view of AI adoption, including:



AI strategy and governance



Human oversight and internal controls



Procurement and deployment practices



Environmental impact assessment



Legal accountability and regulatory compliance



Workforce impacts
(upskilling/reskilling and redeployment)



Data privacy, security, and bias mitigation



Diversity and inclusion in AI practices

By translating responsible AI principles into a structured, comparable dataset, AICDI aims to support two urgent needs at once:

→ Helping companies strengthen their own AI governance and preparedness

→ Equipping investors with clearer signals on how AI is being governed in practice

For companies

For companies, it provides a practical way to identify gaps, strengthen internal governance, and communicate progress with greater clarity. Participating organisations receive scorecards that enable benchmarking, highlight gaps, and help map potential risks, supporting more credible transparency and trust with investors, customers, employees, and regulators. By encouraging a consistent, enterprise-wide approach, AICDI strengthens governance structures, documentation and cross-functional coordination. This supports more efficient workflows, clearer accountability and a more stable foundation for scaling AI safely. This, in turn, supports wider and more sustainable adoption of AI across the organisation.

For investors

For investors, AICDI provides a clearer picture of how AI is governed across sectors, offering standardised disclosure that enhances comparability and reduces information asymmetry. This enables more accurate pricing of operational, regulatory and reputational risk, and also helps identify material vulnerabilities earlier - across data governance, model monitoring, human oversight, third-party dependencies and incident-response processes. The insights generated could strengthen their stewardship and board engagement by shifting conversations from broad commitments towards verifiable governance expectations, including accountability structures, board-level oversight, clear documentation and remediation pathways. By revealing sector-level adoption patterns and emerging norms, AICDI also supports capital allocation towards stronger, more responsible AI practices - helping distinguish durable advantage from short-term gains that mask deferred risk.

This report presents the Initiative's first findings and what they suggest about current patterns in corporate AI adoption, governance, and disclosure at a moment when adoption is accelerating and expectations, both regulatory and societal, are rising just as quickly. It also contains recommendations for investors to make better informed stewardship decisions.

The regulatory and framework landscape

AI governance cannot be meaningfully evaluated in isolation from the external regulatory and framework landscape. Understanding the existing regulations and voluntary standards pertaining to AI therefore provides essential context, establishing the baseline standards against which organisational responsible adoption, controls, and oversight mechanisms can be interpreted.

Existing AI regulations and policy

	The European Union’s Artificial Intelligence Act⁷ Regulation - EU - 2024/1689 - EN - EUR-Lex • Adopted in 2024	The European Union’s General Data Protection Regulation⁸ Regulation - 2016/679 - EN - gdpr - EUR-Lex • Since 2016
	The Brazilian General Data Protection Law (LGPD)⁹ Capa LGPD em inglês 2 • Since 2018	
	The South African Protection of Personal Information Act (POPI Act)¹⁰ Protection of Personal Information Act 4 of 2013 • Published 2013, came into force 2020	
	The Colorado Artificial Intelligence Act (SB 24-205)¹¹ Consumer Protections for Artificial Intelligence • Signed 2024	The California Transparency in Frontier Artificial Intelligence Act (SB 53)¹² 2025
	The South Korean AI Basic Act¹³ 2025	
	The Chinese Interim Measures for the Management of Generative AI Services¹⁴ 2023 The Chinese Cybersecurity Law¹⁵ 2017	The Chinese Data Security Law¹⁶ 2021 The Chinese Personal Information Protection Law¹⁷ 2021



The European Union's Artificial Intelligence Act (AI Act) is the world's first comprehensive law designed to regulate AI. The EU adopted it AI Act was formally adopted in 2024 and has been rolling it out in phases through 2026. The Act sorts AI systems into four categories, based on the level of risk they pose - from unacceptable risk (prohibited) to high risk, limited risk and minimal risk. High-risk systems, such as those used in healthcare, education, employment and law enforcement, must meet strict requirements such as robust documentation, risk mitigation, transparency, human oversight and post-market monitoring. Companies must show how the system works, manage risks, ensure transparency, include human oversight and keep monitoring the system after it is deployed. National regulators, working together through the new European AI Office, enforce these rules. Organisations that fail to comply can face significant fines, similar to those under EU data-protection and competition laws. In 2025 the European Commission supplemented the Act with guidelines for powerful, general-purpose AI systems that could create systemic risk. These models now have to follow additional obligations, such as adversarial testing, incident reporting, strong cybersecurity protections and detailed technical documentation.

The AI Act remains the most advanced attempt to set clear, legally binding standards for safe, transparent, and accountable AI across all sectors.¹⁸ This legislation is currently in the progression of implementation, with a complete rollout due by August 27th.

The EU AI Act is part of a broader regulatory landscape covering digital technologies shaping AI operations indirectly. The EU General Data Protection Regulation (GDPR) continues to influence AI governance: AI systems that process personal data must comply with GDPR principles including lawfulness, purpose limitation, data minimisation and rights such as access, portability and erasure. These provisions impose legal responsibilities on data-driven AI. Comparable statutes in other jurisdictions with binding obligations (e.g. Brazil's LGPD, South Africa's POPIA) similarly shape the lawful handling of personal data used in AI. Given that data quality, fairness and explainability are integral to ethics and safety, data protection law serves as a de facto regulatory floor across many legal systems.



While the United States does not yet have a single, binding federal AI law, multiple states have enacted legally enforceable statutes with AI-specific obligations. States like Colorado have introduced laws requiring developers and deployers of high-impact AI systems to adopt non-discrimination safeguards and transparency measures in domains such as employment and housing. California has passed targeted AI transparency and safety laws that require impacted companies above certain thresholds to report risk assessments and safety measures for large generative AI models.

Federal efforts in the U.S. remain largely agency guidance and executive orders, rather than a unified statute specifically focused on AI.



Elsewhere, countries such as South Korea have progressed national AI legislation, with the AI Basic Act entering enforcement in 2026. This statute is among the first broad national AI laws outside the EU, institutionalising governance mechanisms and safety foundations at the state level.



China's AI regulatory framework emphasises data governance, algorithm accountability, content control and safety compliance. Rather than relying on a single AI law, China regulates AI through a combination of overarching legislation, AI-specific rules and mandatory technical standards. At the core of the framework are the Cybersecurity Law, Data Security Law and Personal Information Protection Law, which apply directly to AI because they govern data handling, security assessments and privacy obligations across all digital systems. These laws set requirements for risk assessments, data minimisation, security controls and cross-border data management, all of which affect how AI models are trained and deployed. On top of these broad statutes, China has implemented a series of AI-specific regulations that target algorithms, synthetic media and generative models:

- **The Algorithm Recommendation Rules** require algorithmic systems, especially those used on large platforms, to be transparent, controllable and aligned with state-defined values. They also prohibit the misuse of algorithms in ways that might disrupt social or economic order.
- **The Deep Synthesis Rules** regulate deepfake technologies by mandating real-name verification, content review processes and security assessments, ensuring that synthetic media can be traced and controlled.
- **The Interim Measures for Generative AI Services**, which took effect in August 2023, impose registration requirements for providers, enforce content moderation, require safety checks and oblige developers to disclose their models' identities and filing numbers. These measures are designed to ensure accountability and traceability in the operation of public-facing generative AI systems.
- In September 2025, China introduced **compulsory labelling rules for AI-generated content**. These rules require both visible labels on all AI-generated text, images, audio, video and virtual scenes, as well as metadata-based implicit labels, making the provenance of synthetic content consistently identifiable.

Non-legally binding standards and corporate self-regulation

Where binding laws are absent or only partially developed, corporate self-regulation remains a major force shaping how AI ethics and governance are put into practice in companies. Non-binding frameworks help organisations mitigate risk, align with best practice, and prepare for future regulation.



International and sectoral standards

International and sectoral standards play a growing role in corporate self-regulation. Although not legally binding, frameworks developed by standardisation bodies (e.g., International Organization for Standardization (ISO)/ International Electrotechnical Commission (IEC)) standards on AI governance, and professional certification schemes provide structured expectations for ethics, risk management, accountability and transparency that companies can use to signal compliance and build stakeholder confidence. These standards increasingly influence procurement and supply chain decisions, as well as investor expectations.



REUTERS/Tim Wimborne



Intergovernmental initiatives

Beyond binding laws and corporate self-regulation, numerous intergovernmental instruments shape global AI governance discourse and practice, often through normative standards, cooperative mechanisms and ethical benchmarks.

The UNESCO Recommendation on the Ethics of Artificial Intelligence is a widely endorsed (by 193 countries) non-binding normative instrument that articulates globally shared ethical principles for AI, grounding them in human rights, dignity, justice, fairness, transparency and environmental protection. It provides detailed guidance for national policy design, institutional governance and stakeholder engagement, and emphasises multi-stakeholder accountability, impact assessment, and redress mechanisms. Member States are urged to enact strong enforcement mechanisms and remedial actions and to build institutional capacities to prevent and mitigate harms stemming from AI systems. The UNESCO Recommendation is widely recognized by Member States and it is currently being implemented in more than 70 countries around the world building on its follow up tools, the Readiness Assessment Methodology and the Ethical Impact Assessment. The Recommendation is serving as a standard reference for countries that are the midst of enacting or revising national frameworks. For example, Nigeria's National AI Strategy, published in September 2025, explicitly drew on the Recommendation as the model approach to develop a high-impact national AI strategy.

This is also the case for the Kenyan National AI Strategy, launched in March 2025, the Mauritius Digital Transformation Blueprint, and Ghana's National AI Strategy from August 2025.

The OECD AI Principles are a widely accepted set of multilateral policy guidelines for trustworthy AI. Endorsed by 47 OECD members (including the EU) and many partner countries, they promote human-centred values such as transparency, fairness, explainability, robustness and accountability. They serve as a policy baseline for both national AI strategies and corporate governance.

The United Nations' Global Digital Compact (GDC) is a comprehensive, non-binding global framework intended to foster responsible, inclusive digital technologies, including AI, and to address issues such as the digital divide, equitable access, safety, and ethical use in line with the UN Charter and the 2030 Agenda for Sustainable Development. The GDC sets common expectations for digital governance cooperation among states and stakeholders.



Corporate AI ethics and governance principles

Major tech firms and consortia have adopted ethical AI principles and governance frameworks, often aligned with international norms and embedded into internal risk, procurement, development standards, impact assessments, and audits. But without mandatory standards and clear legal accountability, practices can be inconsistent and opaque, encouraging speed over responsibility and increasing risks like bias, discrimination, security flaws, and misuse. Fragmented regulation also limits transparency and leaves the public with few avenues for redress, fuelling concern as AI deployment and synthetic content rapidly expand.



REUTERS/Abdul Saboor

Methodology



Data sourcing and collection

- **AICDI collects data from both publicly available corporate disclosures and company responses to the survey.** Public materials, including annual and Environmental, Social and Governance (ESG) reports, governance and responsible AI webpages, cybersecurity and privacy policies, and diversity reporting, were identified for each company.
- **These documents were analysed using a large language model to parse content against the AICDI framework.** The model was instructed to populate survey response fields only where information was explicitly and clearly evidenced in the source material. Where the model identified content that required interpretation, inference, or contextual judgement, this information was not placed in the formal response field but was recorded separately in a “More Details” supporting field to inform engagement and review.
- **Human-in-the-loop quality assurance was embedded throughout the process.** During model development and testing, outputs were reviewed by analysts to validate alignment with the framework and refine prompts. In production, spot checks were conducted on no less than 10 per cent of the sample, with manual review used to assess accuracy, consistency, and appropriate treatment of ambiguous information.
- **Companies were invited to review, validate, and supplement the resulting dataset** through direct survey participation, so the analysis reflects both what’s publicly disclosed and what companies can explain about their own AI use and governance.



Company selection

The target list was constructed through a risk-based and market-representative approach.

→ Risk-based sector weighting

Proprietary research was conducted to assess each industry's exposure to AI-related workforce and human rights risks, including the ways in which employees may be affected by AI adoption. This produced an industry-level risk score, which was aggregated into sector-level risk weightings. These weightings were used to define each sector's proportional share of the overall sample.

→ Company selection within sectors

Within each sector allocation, companies were selected by ranking firms based on market capitalisation and number of employees, prioritising larger and more influential companies on the basis that they are likely to have the greatest workforce footprint and potential AI-related impact.

→ Market balance and country-level adjustment

Following the initial selection, a market-balancing adjustment was applied to ensure global representativeness. Market classification followed London Stock Exchange Group (LSEG) definitions of Developed, Emerging, and Frontier markets. To avoid structural over-concentration in large developed markets, the sample was calibrated towards an approximate 60 per cent Developed Markets and 40 per cent Emerging Markets, with Frontier Markets included on a limited basis within the Emerging/Frontier allocation.

Within this structure, country-level allocations were adjusted using LSEG-adjusted social risk indicators to inform how the Emerging/Frontier portion was distributed across countries with differing risk profiles.

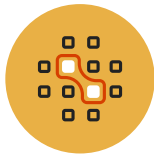
→ Investor nomination and final refinement

Investors were invited to nominate companies of strategic relevance. The nominated list was cross-checked against the sample and, where nominated companies were not already included, they were added by replacing lower-ranked companies, while maintaining the overall sector and market balance described above.

→ Final sample size

While the initial target was 3,000 companies, the final dataset covers 2,972 companies. In a limited number of cases, sufficient information could not be identified to enable reliable mapping to the AICDI framework, and these companies were therefore not included to maintain data quality and consistency.

See this section for a detailed distribution of companies.



Data processing and classification

- **Qualitative responses** were processed using a structured coding framework aligned with the survey's thematic pillars: AI governance, workforce impacts, and data governance.
- **Narrative disclosures** were mapped to analytical categories capturing implementation levels, governance arrangements, and risk controls. Where information addressed multiple dimensions, more than one category could be assigned.



Analytical approach

Our analysis for this report focused on identifying patterns across regions and sectors rather than evaluating individual companies.

→ Treatment of partial and missing information

Consistent with the data collection approach, only explicitly evidenced statements were used to populate structured response fields. Where disclosures were incomplete, ambiguous, or relied on inference, responses were coded as “insufficient information.” Potentially relevant contextual material was retained in the More Details field to support engagement without overstating maturity or assuming absence of practice.

→ Qualitative categorisation and text analysis

Open-text responses were translated into structured categories aligned with the survey themes. The detailed classification framework and coding rules are available upon request.

→ Quantitative analysis

Descriptive statistics and cross-tabulations were used to compare results across sectors, regions, and company roles (AI developers and deployers). Spearman's Rank Correlation explored relationships between governance, risk management, transparency, and workforce measures. Sentiment analysis of textual responses complemented structured indicators with insight into narrative tone.

Findings are presented in aggregate to highlight trends and emerging practices.

AICDI dataset breakdown

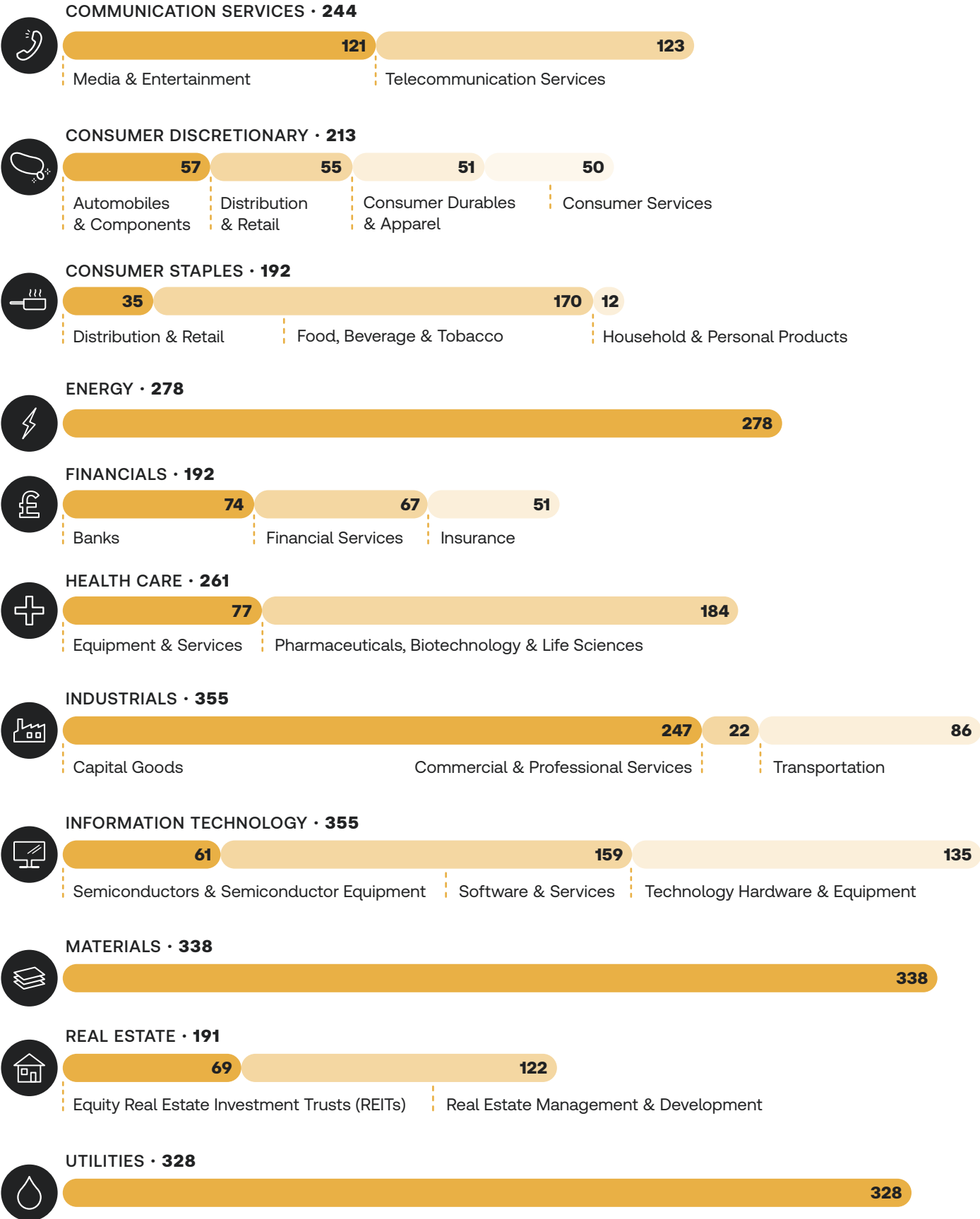
Regional breakdown



Note: The regional distribution applied in this report follows AICDI methodology. The distribution reflects the predefined company list and continental classification approach applied in this report. Regions are grouped broadly by continent to support comparability across markets. The UK is presented separately to reflect its specific relevance within our investor base and engagement context.

Sector breakdown

Note: This sector classification follows GICS (Global Industry Classification Standards)



Thematic findings



REUTERS/Aly Song



• Finding 1:

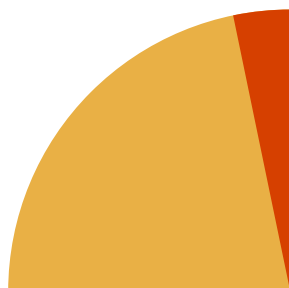
Public commitment from companies to AI governance frameworks remains low, signalling inconsistent governance practices

The rapid emergence of artificial intelligence has left companies scrambling to effectively adopt and integrate AI capabilities into their operations. This implementation can give businesses a competitive edge by boosting efficiency and automating complex tasks but it can also have vast unintended consequences that extend far beyond individual business operations. AI, if not effectively and responsibly governed, can lead to everything from global market disruptions and economic imbalances, to mass security and privacy risks.¹⁹ These risks can be mitigated by adhering to a formal set of a formal set of principles and practices which can provide a foundation for responsible AI adoption.²⁰ These frameworks can help companies deploy effective and responsible AI systems which maximise opportunities and avoid compliance risks. While commitment to such frameworks does not guarantee strong governance, the absence of commitment weakens accountability, comparability, and the credibility of any subsequent controls.

Most companies are not publicly anchoring their AI practices to formal governance frameworks, and even where they do, commitments are more common in a few “AI-exposed” sectors and large-cap firms²¹

28.7%

of the sample group identified as adhering to at least one AI governance code/standard/framework

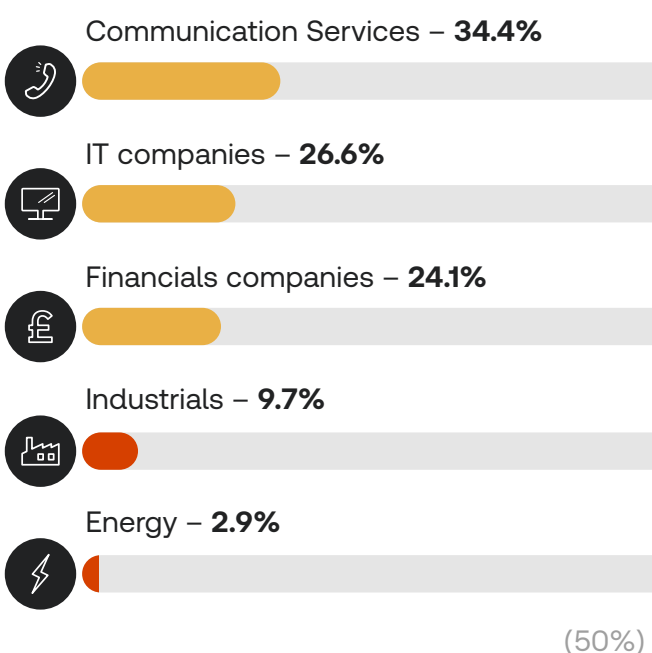


That translates to nearly 9 in 10 companies not publicly committing to any named AI governance framework – demonstrating that despite the growing availability of governance frameworks, most companies are not publicly anchoring their AI adoption to them, making it difficult for stakeholders to assess the robustness of governance.

Where companies are demonstrating an adherence to formal governance frameworks, it is generally in areas where AI exposure and scrutiny are highest.

The sectors with the closest association with data and AI technology and products, such as the IT and Financials sectors, are most likely to publicly disclose that they are adhering to a formal AI framework. Most other sectors fall below 15 per cent, with the Industrials and Energy sectors being least likely to disclose adhering to any framework.

Percentage of companies that report adhering to at least one AI framework:



Information, communication and financial services appear to be ahead on AI governance framework adherence because they sit at the sharp end of AI adoption. Their work is highly digitised and data-rich, and much of the value they deliver comes from knowledge-intensive tasks such as analysing information, coding, customer interactions, and risk and compliance workflows.

Given that a larger share of their tasks can be automated, these sectors tend to adopt AI more quickly, making them more likely to scale beyond pilots and integrate AI into core operations. This creates earlier pressure to formalise controls, accountability, and assurance. PwC’s 2024 AI Jobs Barometer suggests these sectors have higher AI exposure, reinforcing the link between task-level applicability and faster governance maturity.²¹

PwC’s labour market data also shows higher “AI penetration” (measured by job ads seeking specialist AI skills): information and communication is about 5 times higher than other sectors, professional services about 3 times, and financial services about 2.8 times. In practice, that kind of penetration usually signals more sustained, business-critical AI use, which increases regulatory, reputational, and client-delivery risk and accelerates the need for clear policies, oversight, risk management, and auditability.²¹

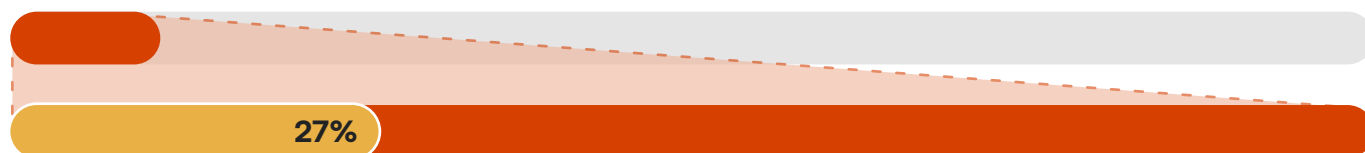
Many companies are therefore developing AI strategies without anchoring them in the principles and practices of a formal governance framework, leaving themselves exposed to risk.

Despite only 13 per cent of companies publicly claiming to adhere to an AI framework, a much larger share report having an AI strategy or guidelines.

13% of companies publicly claim to adhere to an AI framework



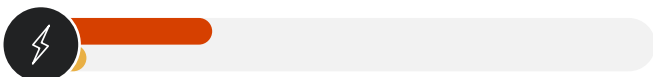
43.7% of companies report having an AI strategy or guidelines



However, among those companies with an AI strategy, only around 27 per cent also report adherence to a governance framework. It indicates that AI strategies are frequently developed without a corresponding external commitment to recognised governance frameworks, suggesting that many strategies are oriented primarily towards accelerating adoption and capturing value rather than setting out robust governance commitments.

This is particularly stark in the Energy sector.

25% claim to have an AI strategy



2% report adhering to a formal framework

It suggests that most Energy companies that appear to have AI strategies/guidelines do not adhere to corresponding framework to ensure these guidelines are compliant or at least do not publicly say that they do.

Communications Services show a larger gap.

65% claim to have an AI strategy



34% report adhering to a formal framework

This indicates that approximately half of communication service companies that appear to have strategies do not cite corresponding framework based on their public disclosure.

This creates risk and limits opportunities when it comes to deploying AI within a business, aligning with findings from a 2025 Thomson Reuters study, which found that organisations with a mature AI framework had better outcomes than those that didn't.²²

This is particularly the case when it comes to deploying AI tools across multiple governance, risk, and compliance functions, including using AI to track risk proactively and for predictive risk modelling to shape risk posture and strategic planning.²³

Where companies are citing the use of an external framework, it is the EU AI Act which dominates, even for companies operating outside the EU.

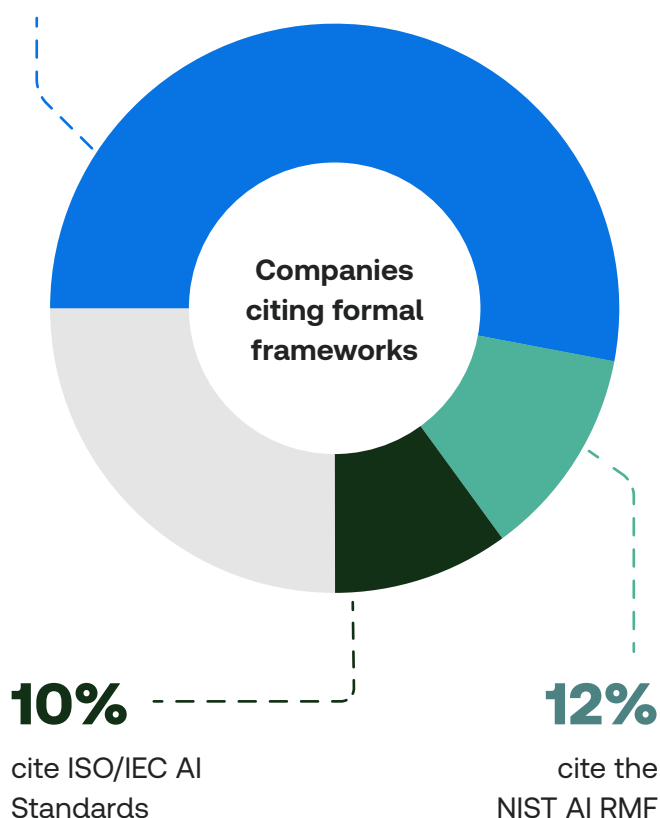
When we look specifically at commitments to formal frameworks, companies reference the EU AI Act as a framework they adhere to. While that share is modest across the full sample, it becomes far more significant among the subset that do report adhering to a framework:

Over half of these companies cite the EU AI Act – despite nearly half operating outside the EU – indicating that it has emerged, by a substantial margin, as the dominate regulatory reference point across jurisdictions.

The following most cited frameworks were the National Institute of Standards and Technology Artificial Intelligence Risk Management Framework (NIST AI RMF) and ISO/IEC AI Standards.

53%

cite the EU AI Act

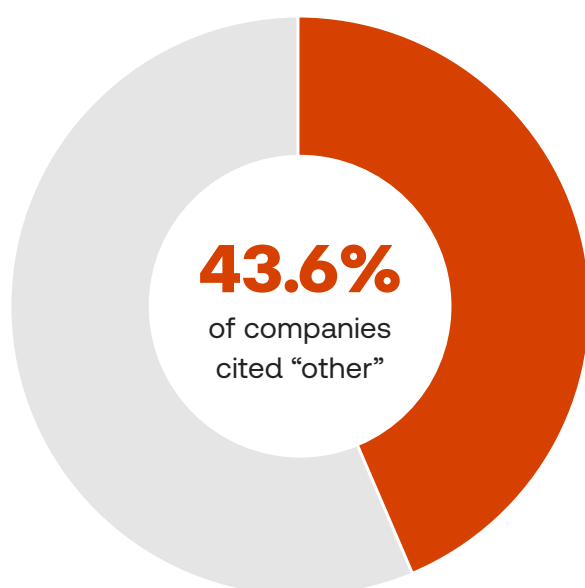


This indicates that, for organisations that disclose details of their AI governance, the EU AI Act, designed as a comprehensive, risk-based legal framework balancing people’s rights and safety with the growth of Europe’s AI industry, appears to be treated as the “gold standard” framework. As a result, it is often used as the default external anchor for structuring and benchmarking governance.²⁴

More than this, the AI Act has clear cross-border reach, with many non-EU providers and users still covered by the legislation if they place AI systems on the EU market or if the system’s outputs are used in the EU. Because of this, companies in EU linked supply chains often meet the Act’s requirements from the start to avoid costly changes later and to keep smooth access to the EU market.²⁵

It also makes a strong case for mandatory reporting initiatives, as it suggests that companies are a lot more likely to adhere to a framework if there is a binding legal imperative behind it.

Even where companies do disclose alignment to a framework, reference points vary widely



Many companies point to internal or non-standard approaches rather than widely recognised benchmarks. In fact, among those companies citing the use of a framework, 43.6 per cent cite “Other”, rather than one of the preidentified standard frameworks such as:

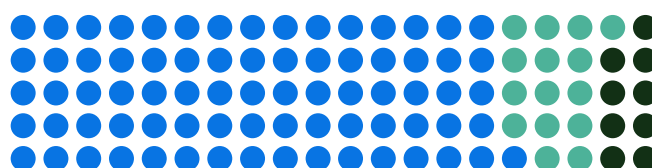
- The EU AI Act
- The OECD AI Principles
- The G7 AI Hiroshima Process
- UNESCO’s Ethics of AI Recommendation

This weakens comparability across the wider AI governance ecosystem.

Commitments anchored to divergent frameworks are difficult to interpret on a like-for-like basis, complicating external assessment and portfolio analysis. Scholars have noted that such fragmentation can restrict access to technology and impose significant compliance burdens unless greater alignment is achieved.²⁶ From an investor perspective, overlapping frameworks also risk producing incompatible data points, making portfolio analysis more difficult.

Where companies align to AI frameworks, they tend to cite only one.

Among companies claiming to adhere to a framework, most only cite one, with the average number of frameworks being 1.37.



● 76.3% cite adhering to 1 framework

● 23.6% cite adhering to 2+ frameworks

● 9.4% cite adhering to 3+ frameworks

This further illustrates the narrow and uneven nature of alignment: most companies rely on a single reference point, while only a small minority draw on several.

Adhering to multiple AI governance frameworks can be a mixed blessing.

- It may strengthen governance by widening risk coverage and improving credibility with regulators, customers, and investors
- But it can also introduce overlapping controls, inconsistent definitions, and higher compliance costs.

A more effective approach is to integrate different frameworks within a single internal governance system – mapping requirements to one coherent structure rather than running parallel processes. This enables organisations to combine compliance baselines with ethics, human rights, and environmental considerations that individual frameworks may miss.



• Finding 2:

Many companies publish strategies on AI but it is less clear how these are put into practice

Research consistently shows that organisations with more mature AI governance tend to outperform their peers, with governance maturity often associated with better alignment of leadership, budgets, ethics, compliance, and technology under a unified framework.²⁷ Thomson Reuters research, for example, finds that organisations with visible AI strategies are twice as likely as those with more informal or ad hoc adoption approaches to experience revenue growth as a direct or indirect result of AI adoption²⁸. Meanwhile, according to the 2024 Edelman Trust Barometer, 43 per cent of respondents said they would reject AI if they believe it was managed poorly.²⁹ However, visible strategies do not necessarily reveal how AI governance is implemented in practice or supported by operational controls. This sets up a contrast explored below: relatively robust signals at strategic level, but much weaker evidence of delivery and oversight.



REUTERS/Seun Sanni

Companies demonstrate a comparatively high AI oversight at a strategic level

40%

When it comes to the visible layer of AI strategy and governance intent, publicly available information is relatively robust. 40 per cent of companies report board/committee-level oversight.

However, strategic signals do not necessarily indicate operational capacity or day-to-day governance.

Board-level structures and strategy are typically the most visible entry points into AI governance, as they are more readily formalised and communicated than underlying processes and resources allocation. This means that strong signals at top layer may coexist with limited evidence of how governance is resourced and implemented in practice.



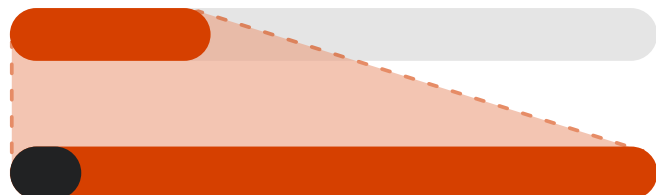
REUTERS/Jose Cabezas

Evidence of AI governance implementation remains far less visible than strategic commitments.

Unlike top-level policies, limited information is publicly disclosed on the teams, processes, and accountability mechanisms that translate intent into action.

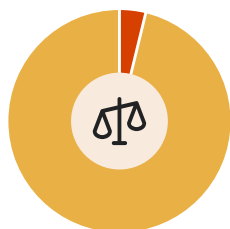
For example, less than a third of all sampled companies claim to have an additional team or resource dedicated to AI governance.

31% have an AI governance team/ resource



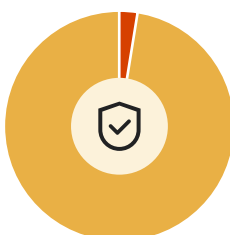
11% identify "Data Protection Officer"

Of the companies which did claim to have an additional team or resource dedicated to AI governance, 11 per cent identified the Data Protection Officer (DPO) as the most popular option, while wider committees or taskforces were cited significantly less often.



3.8%

of companies claim to have an AI ethics committee



2.5%

claim to have an AI safety and security taskforce

This pattern suggests that many organisations are building AI governance infrastructure on top of established cybersecurity and data protection roles.

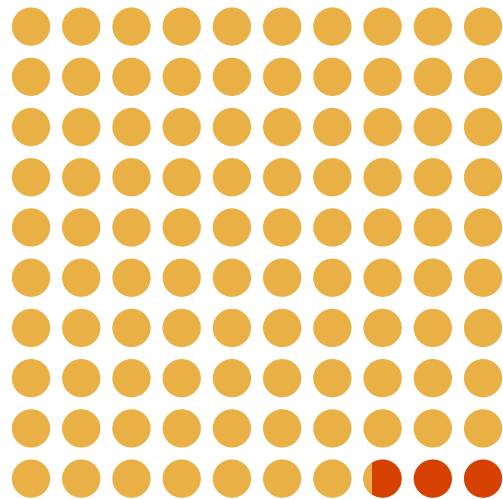
The DPO role is often already in place for broader privacy and cybersecurity needs, even where companies don't yet have a dedicated AI governance team. The overlap between AI governance and data protection means that firms may extend the DPO's remit to cover additional AI-related oversight where dedicated AI teams are not yet in place. While this is a viable outcome, particularly for smaller teams with more limited capacity, the expansion of this role must be accompanied by appropriate training and formal role expansion definitions. More specialised taskforces or committees can support effective implementation across a company by clarifying ownership, setting performance criteria, and enabling ongoing monitoring and evaluation. Such structures could be reflected in the use of internal structures for initiatives such as diversity or just transition programmes and provide a comparable mechanism for formalising accountability in AI.



REUTERS/Andrea Comas

AI model registries and lifecycle infrastructure remain rarely evidenced in corporate governance practices³⁰

Registries and lifecycle accountability are foundational tools for operational AI governance, and their absence suggests strategy may not be being consistently translated into a company's control infrastructure.

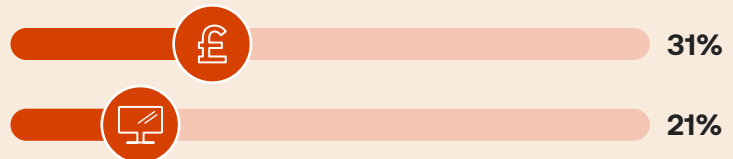
**97.3%**

provide no evidence of a formal AI model registry from their public disclosure.

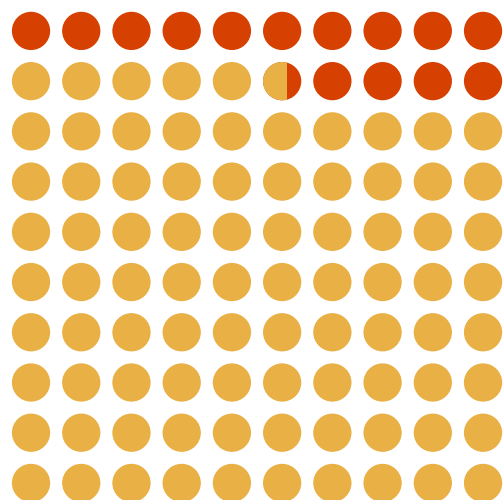
2.7%

of companies publicly report having a formal AI model registry.

Among companies with an AI model registry, the finance sector leads, followed by the IT sector.



Accountability for ethical impacts across the lifecycle is materially higher than registries, but still low and often undisclosed.

**15.4%**

of companies say they can trace ethical impacts of their AI systems to a responsible person or organisation at the relevant stages of the AI system lifecycle.

84.6%

of companies do not indicate they can do this – meaning they either say they cannot, or they don't provide evidence/confirmation that they can attribute responsibility in that way.

This indicates that mechanisms for lifecycle accountability are still emerging but not yet consistently adopted.

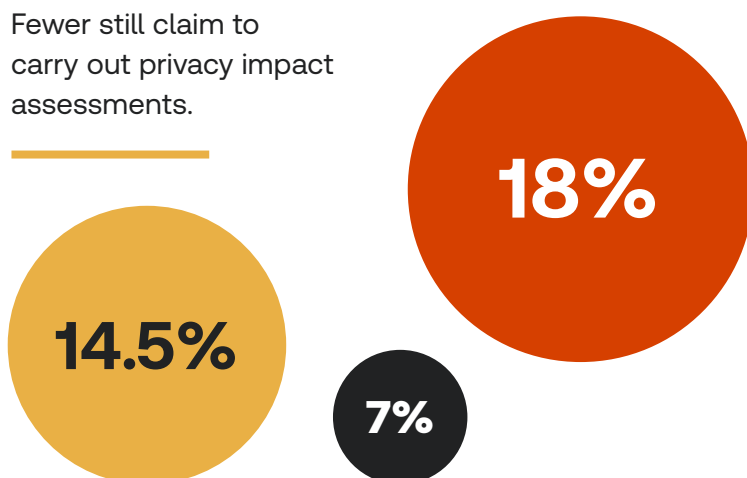


Luka, Inc./Handout via REUTERS

AI impact assessments appear in a minority of companies and most often focus on data and privacy, while AI human-oversight policies are less frequently implemented according to public information. Among firms reporting AI impact assessments, fewer also describe corresponding governance for human oversight.

Less than one fifth of companies claim to carry out data protection impact assessments.

Fewer still claim to carry out privacy impact assessments.

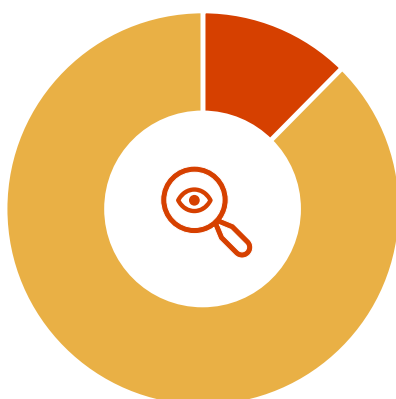


This drops lower when considering impact assessments from a rights perspective, with fewer than 1 in 10 companies claiming to carry out human rights impact assessments.

If companies are not carrying out AI impact assessments, they risk relying on AI systems that are not tested in practice. This can leave material harms and compliance issues undetected, particularly around privacy, bias and discrimination, safety, and human rights, and can weaken accountability by limiting evidence that risks are being identified, mitigated, and escalated when needed.

Where assessments are undertaken but not communicated externally, they are missing the opportunity to reassure stakeholders that they are monitoring the implementation of their AI strategies and ensuring AI is being deployed in the company in a safe, effective and just way.

A crucial part of AI risk management is human oversight governance, yet many companies are not sharing enough data on how – or if – that is occurring.



12.4%

of companies report having a policy to ensure a human oversees AI systems

87.6%

either do not specify whether such a policy exists or state that they do not have one

Impact assessments and human-oversight policy should go hand in hand as they address two connected parts of the same control cycle: impact assessments identify where harm could occur, and human oversight sets out who intervenes, when, and with what authority if risks arise.

In practice, impact assessments can inform oversight requirements, for example by specifying review points, trigger thresholds, escalation routes, and who can pause, override, or roll back a system. This matters most for bias and fairness because human oversight is what makes a non-discrimination promise real in practice: people do the regular checks, apply judgment in context, and step in to fix problems when results start to drift. It also strengthens accountability by making clear that people, rather than an AI system, are responsible for monitoring, challenging, and providing corrective action.³¹

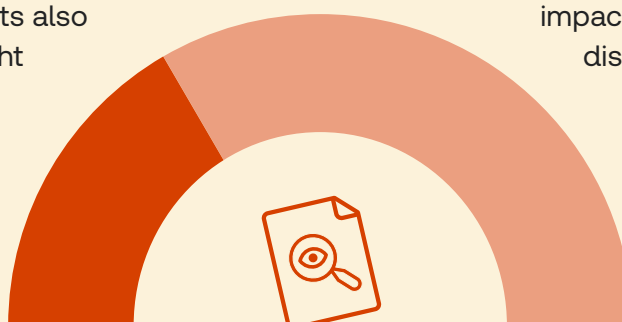
Yet despite this, among the 836 companies reporting any impact assessments, the majority either conduct impact assessments without a disclosed human-oversight policy, or do not publicly report whether one is in place.

28.7%

of companies that report having any impact assessments also report a human-oversight policy

71.3%

of companies either conduct impact assessments without a disclosed human-oversight policy, or do not publicly report whether one is in place

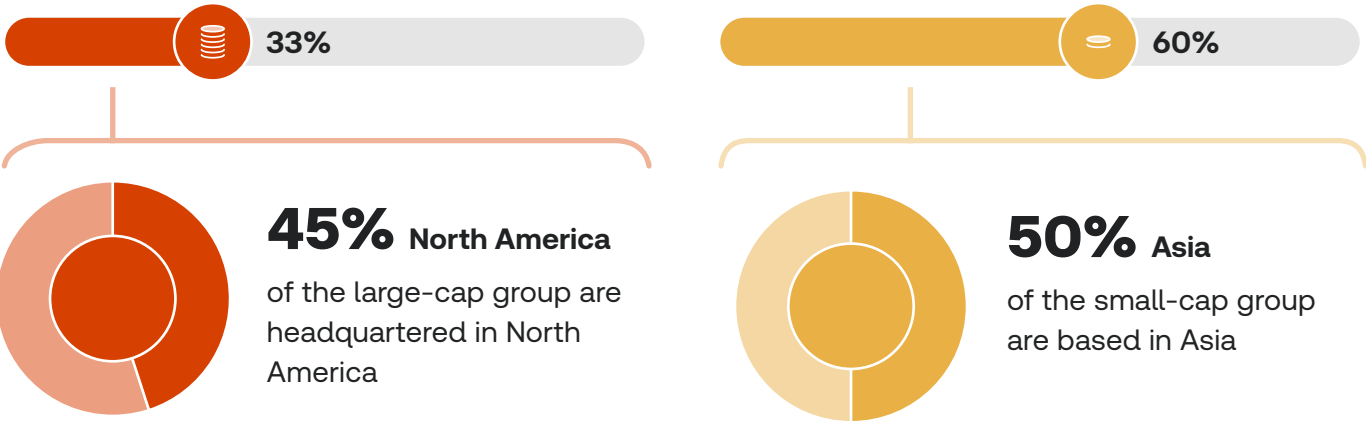


Larger companies are more proactive in AI adoption and governance

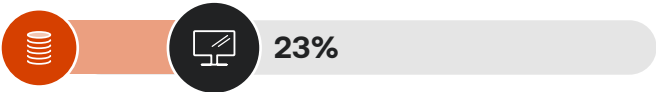
**Large-cap companies:** Companies with a market cap of \$50 billion or more

**Small-cap companies:** Companies with a market cap under \$2 billion

Large-cap companies are nearly twice as likely to have an AI strategy.



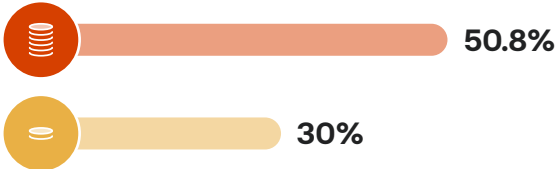
Sector dynamics differ as well.
Among large-caps, IT leads AI adoption by 23 per cent, whereas **no single sector clearly dominates within the small cap segment.**



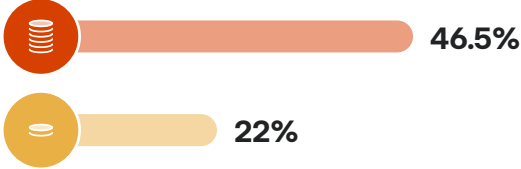
This indicates that AI uptake is shaped not only by firm size but by local ecosystems and industry structure, with leadership in larger firms concentrated in North American technology-intensive sectors, and a more diffuse profile among smaller companies.

There is also a correlation between company size and the data publicly shared on adoption of formal governance structures.
Large-cap companies are more likely to report dedicated AI oversight bodies and dedicated AI governance resources than small-cap companies.

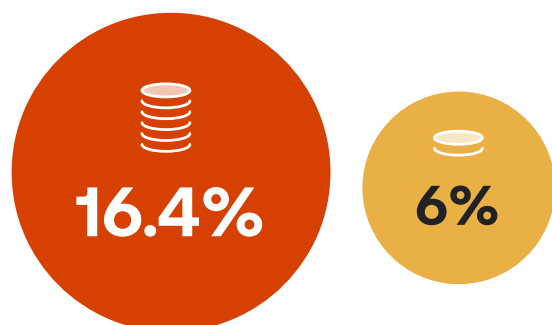
Percentage of companies that report having dedicated AI oversight bodies



Percentage of companies reporting dedicated AI governance resources



The largest companies also report independent AI ethics resources more often.

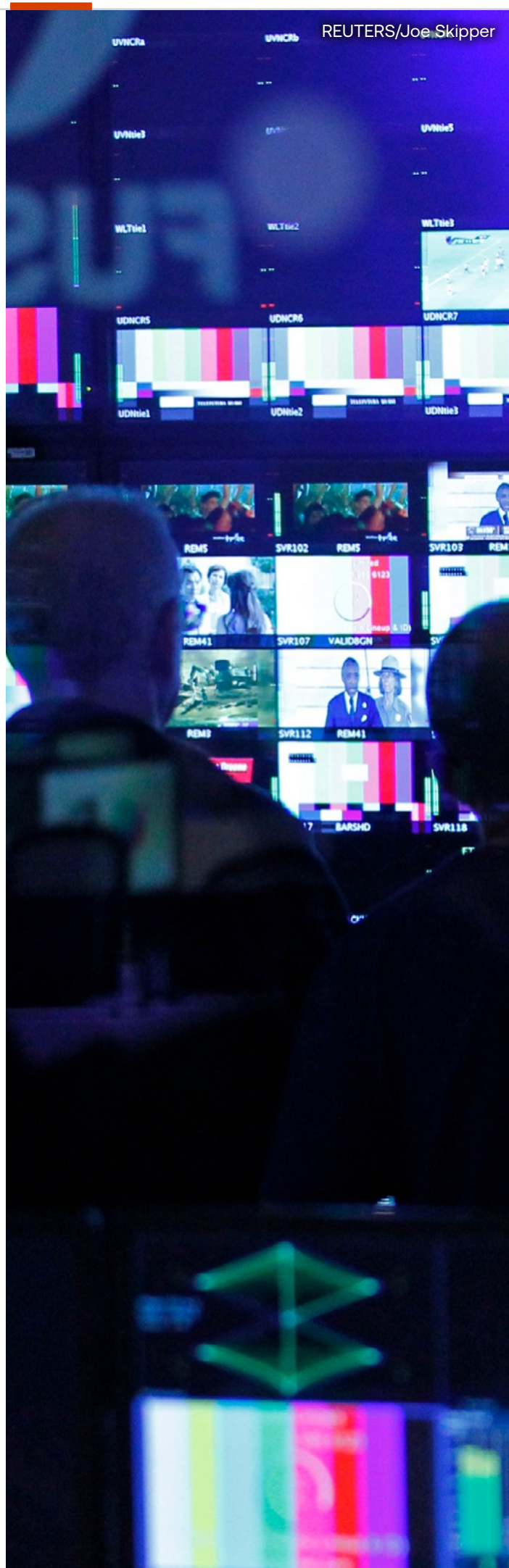


16.4 per cent of large-cap companies report independent AI ethics resources versus 6 per cent among the smallest.

It is unsurprising that organisations with greater capacity and specialisation tend to deploy more developed AI governance. They typically have more budget and staff time to dedicate to governance, risk, and compliance; clearer accountability structures to coordinate across business units; and deeper in-house expertise.

It does showcase, however, the challenge small to medium businesses face in building proportionate AI governance while operating with lean teams, limited compliance bandwidth, and less access to specialist legal and technical support, which can push governance to become reactive rather than embedded.

This strengthens the case for practical, scalable support for smaller firms, such as simplified governance frameworks, off the shelf policies and documentation templates, shared assessment tools, and affordable external advisory or sector led guidance. In practice, this can be complemented by extending existing compliance roles, such as the Data Protection Officer function, to provide interim AI related oversight where dedicated AI teams are not yet in place and resourcing is constrained.





REUTERS/Krishnendu Halder

• Finding 3:

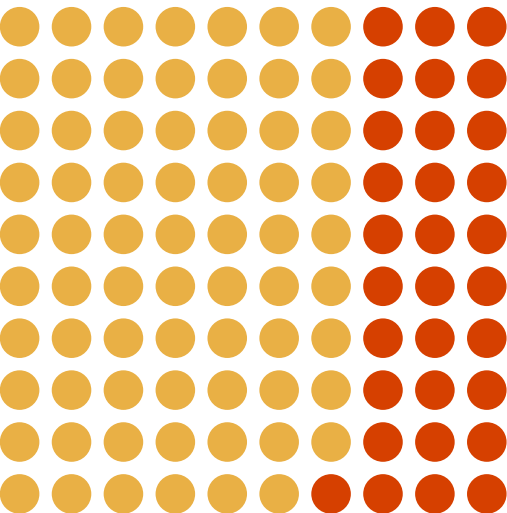
Companies do not demonstrate adequate protections for workers as AI reshapes jobs

AI is rapidly reshaping work both across individual workplaces and within the wider global economy, with some jobs declining while new ones emerge. While there are projections from companies like McKinsey that that AI will create 20-50 million new jobs globally by 2030³², we are also hearing constant news of organisations like Amazon, UPS and Pinterest reducing their workforce because of AI.³³ IMF research also finds almost 40 per cent of global employment is exposed to AI, with half the roles benefitting from AI integration and the other half facing lower labour demand, leading to lower wages and reduced hiring.³⁴ In parallel to this, without sufficient oversight, AI can threaten workers' rights, amplifying bias, increasing surveillance and work intensity, and enabling inhumane decision-making at scale.

To alleviate the risk to workers, companies need to understand where and how their workforce is being affected and respond with protections such as reskilling, retraining, redeployment, and transition support. Where AI has resulted in a worker protection being violated, there also needs to be access to effective remedy. This is only possible with solid workforce data and feedback loops to show where workers are being affected.

The data suggest that employees are ill-prepared for the impact AI is having on the future of work

According to our dataset, most companies are not demonstrating that they are preparing employees for AI-driven change or protecting them from AI-related workplace risks.



The International Labour Organisation (ILO) has identified education and reskilling as the key to ensuring that artificial intelligence is beneficial to workers.³⁵

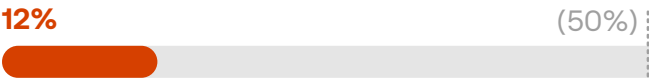
31%



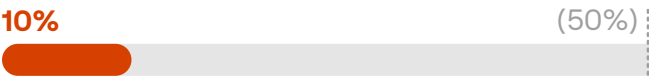
Yet despite this, **less than a third of companies evidenced that they offer training and/or reskilling programmes for employees adapting to an AI-integrated workplace.**

Even with the 31 per cent where training programmes exist, there is a vast variation in the scope and depth of the training offered.

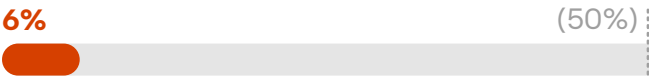
For many the training programmes are not enterprise-wide or structured but ad-hoc, or limited to leadership roles.



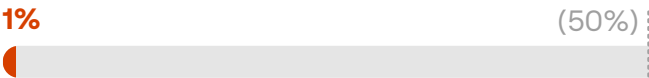
of all sampled companies claim to offer comprehensive and/or tiered training coverage, which means structured training available across the organisation, while others showed a less structured approach:



evidenced ad hoc training coverage, for example, informal or one-off training.



evidenced role-specific training coverage, focused on training targeted at specific teams or functions.



evidenced leadership-focused training coverage, targeted primarily at executives or managers.

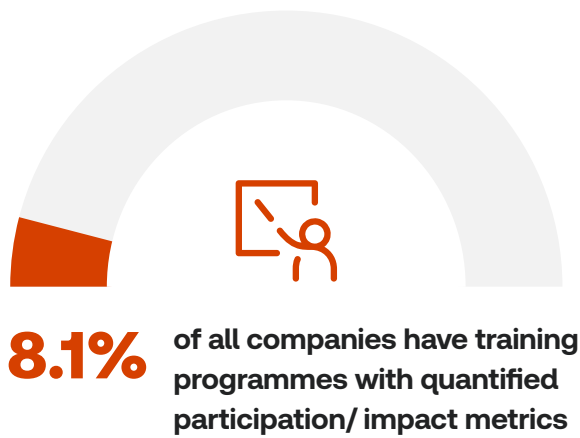
This non-standardised, unstructured approach to AI training can heighten risks for workers by leaving those most exposed to AI-driven change (particularly frontline and non-technical staff) without the baseline awareness, practical guidance, and support they need to use AI tools safely. In addition, they remain unaware of how AI may affect their work and are not instructed on how to raise concerns or seek remedy when problems arise.

Workforce preparedness differs significantly according to company sector.

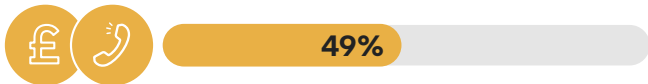


Preparedness is concentrated among large-cap firms and in sectors closest to AI development and deployment (including IT, Communications, and Financials).

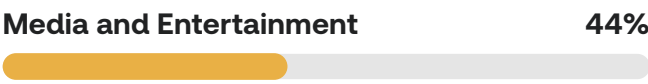
Across the wider market, the prevailing pattern is either no evidence of training or reskilling, or initiatives that lack robust measurement, suggesting workforce readiness is often treated as a capability add-on rather than a tracked, change-management programme.



Financial and Communications companies are the most likely sectors to offer training programmes for employees adapting to an AI-integrated workplace, with almost half claiming to do so.



Within the Communications sector, training activity is uneven – Telecommunications companies lead the way, while Media and Entertainment firms report relatively lower levels.



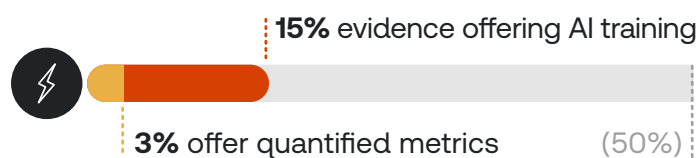
Given that media and entertainment are directly shaped by AI, this gap raises concerns that parts of the sector facing the greatest disruption may not yet be investing proportionately in workforce readiness

Financial and Communications companies are also the most likely sectors to evidence measuring the impact of that training, with 14-15 per cent offering quantified training metrics. IT companies trail them slightly.



While these sectors tend to offer more AI-related training, far fewer evidence available regarding quantified measure of its outcomes. The absence of systematic evaluation makes it hard to tell whether AI-related worker training is actually building usable skills and improving day-to-day worker outcomes, such as better job performance, safer workflows, fewer errors, and greater confidence and autonomy, or whether it is mainly a check-box exercise that raises awareness without changing practice.

Comparatively, Energy and Real Estate companies were the least likely to share data on AI-related training.



While AI-driven decisions directly shape energy use, infrastructure, and community outcomes, limited workforce preparation carries wider implications. Without stronger capability building, adoption may outpace the ability to manage emissions, safety, and transition risks embedded in these business models, which could translate into operational disruption, reputation exposure and mispriced long-term liabilities.

Disclosures raise concerns about labour protections in the context of AI, with limited public information clouding the full picture

Within the dataset, only 14 percent of companies evidenced that they had policies in place to mitigate the negative impacts of AI systems on workers. This means the majority of companies (86 per cent) either have no policies in place, or do not publicly communicate it.

14% of companies evidenced having policies to mitigate negative impacts on workers



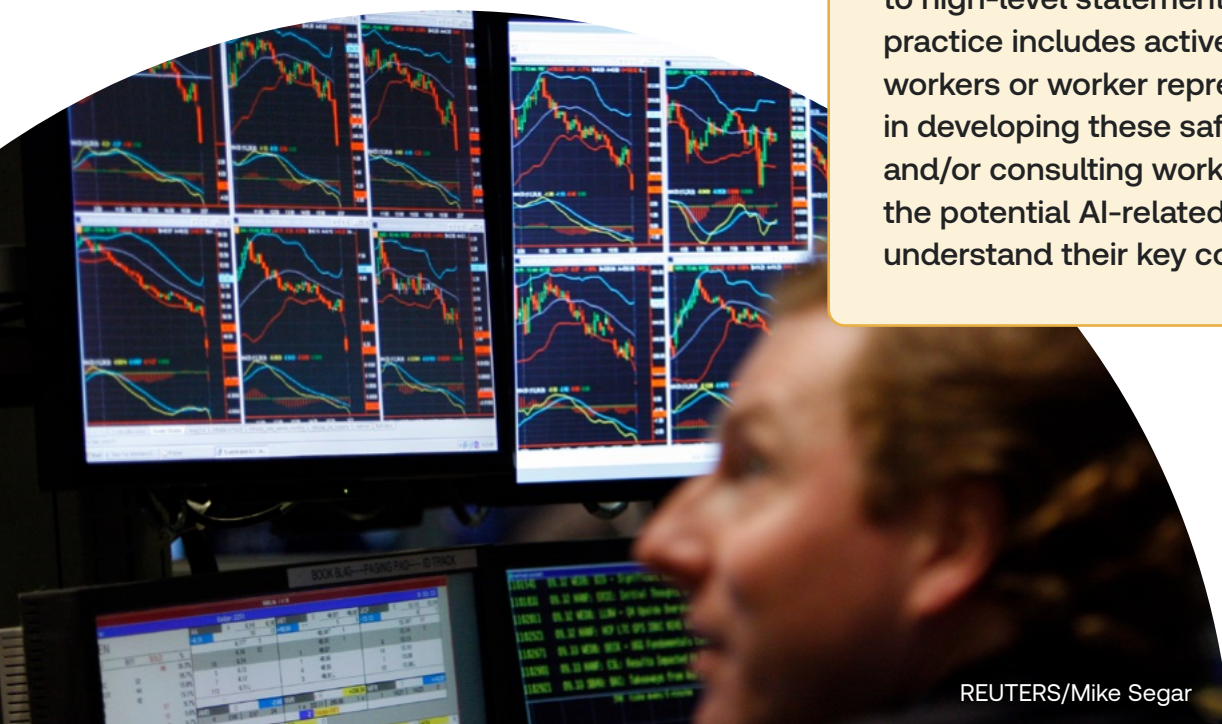
75% of this cohort cited multiple categories

Even within these companies that have some form of safeguarding policies, the information provided is often too high-level to assess their practical effect. Many describe bundles of controls without specifying how these measures operate in practice, what risks they address, or how compliance is monitored. 75 per cent of this cohort, for example, cited multiple categories and 9 per cent of them gave little to no detail on what the safeguards were.

This leads to a lack of specificity and measurability when it comes to the safeguard policies and processes, meaning that companies lack the infrastructure to detect if the safeguards are working or not.

For example, broad statements such as a commitment to protect their employees' privacy, or protect them from AI bias, are too general to enforce or evaluate in practice. They do not specify what actions will be taken, what technologies are covered, or how compliance will be assessed. In contrast, a clearly defined worker protection policy, for example to ban the use of intrusive AI tools in employee monitoring, or a requirement that any algorithmic screening tools undergo an annual bias audit, sets explicit boundaries and creates observable criteria. Such concrete measures can be tracked, tested, and reported on, enabling meaningful oversight and accountability.³⁶

The recommended practice is setting specific, measurable, and enforceable safeguards, as opposed to high-level statements. Best practice includes actively engaging workers or worker representatives in developing these safeguards, and/or consulting workers on the potential AI-related issues to understand their key concerns.



Very few companies share information on an AI-related internal complaints mechanisms, leaving it unclear whether such safeguards are in place.

When assessed on whether the company has an internal mechanism for the submission and review of employees' complaints in relation to AI:

2.3%

evidenced having a complaints mechanism

98%

did not showcase evidence on if they have complaints mechanism



The findings suggest many organisations lack the mechanism for AI related complaints beyond the broad generic complaint and this is compounded by low awareness of where AI systems may infringe employees' rights and protections. As a result, AI-related harms and emerging risks can be misclassified, under-reported, or go unnoticed altogether.

Approached to AI-related internal complaints mechanisms.

The most common approach is multiple channels, followed by a general whistleblowing mechanism. Fewer companies still cite a direct reporting channel, a hotline or a tech committee.

Multiple channels – **56%**



General whistleblowing mechanism – **26%**



Direct reporting channel, a hotline – **7%**



Tech committee – **3.2%**



Despite this, there is a case for building appropriate grievance mechanisms tailored to addressing harms caused by AI specifically.

According to research from Chatham House, these should be tailored to remedy harm suffered due to AI's capacity for operation at scale which risks infringing the rights of large numbers of people at once, for example:

Any AI-related safeguards need to be aligned with the existing regulatory requirements

In the UK context, the right to privacy remains a fundamental consideration in the workplace, even as employers adopt increasingly sophisticated monitoring technologies, including AI-enabled tools. Any monitoring activity must comply with data protection requirements by being proportionate, necessary, and clearly linked to a legitimate organisational purpose. Employers are expected to provide transparent information about the nature, scope and rationale of any monitoring, ensuring that workers are aware of how and why their data is being collected and used.

Covert monitoring is subject to particularly stringent safeguards

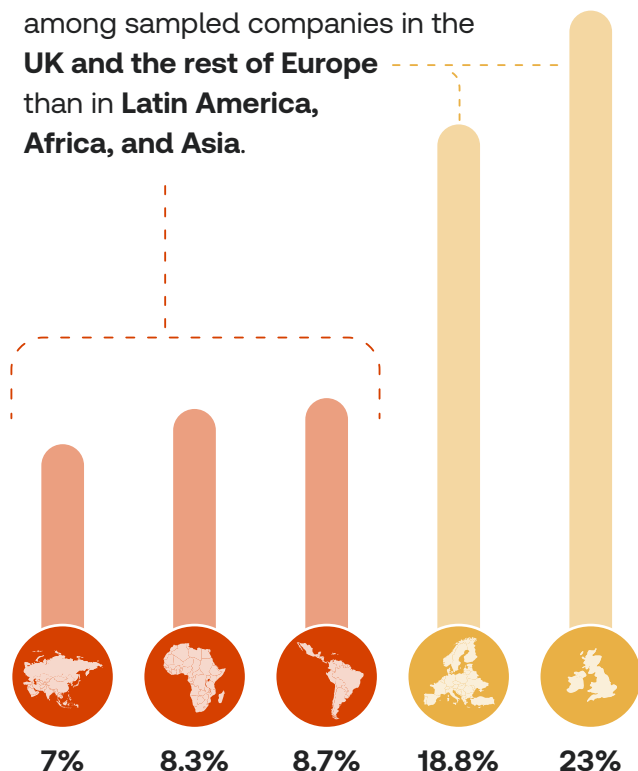
It is permitted only in exceptional circumstances, such as where there is a substantiated suspicion of serious misconduct, and must be narrowly targeted and limited in duration. Prior to undertaking such monitoring, organisations are required to conduct a Data Protection Impact Assessment to evaluate potential risks to individuals' rights and freedoms.

These protections are equally applicable in remote / home-working environments

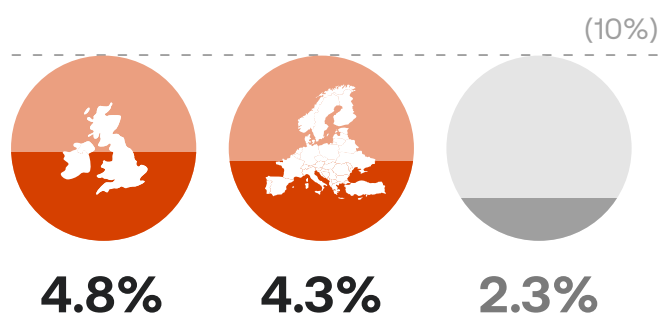
Employees' privacy rights extend to their home, and employers should avoid intrusive practices that blur the boundary between professional and personal spaces. Clear organisational policies and governance measures are essential to ensure that monitoring practices remain compliant with UK data protection law and uphold the core principle of respecting workers' privacy.

Regionally, European companies demonstrate the most transparency on AI safeguarding of workers.

Safeguard disclosure is more common among sampled companies in the **UK and the rest of Europe** than in **Latin America, Africa, and Asia**.



Similarly, the presence of an AI-related complaint mechanism is higher in the UK and the rest of Europe than in average.



This potentially reflects stronger or more mature regulatory and governance expectations around workplace rights, including transparency requirements, and corporate accountability in the UK/Europe, as well as greater institutionalisation of formal reporting channels.

Nevertheless, with under 10 per cent of responses indicating the presence of AI-related complains mechanisms for each region, there is a gap in providing workers with a clear pathway to raise any potential issues specific to the use of AI in the workplace.

The absence of such dedicated mechanisms has several implications: it can limit organisations' ability to detect and address early signs of harm, undermine employee trust in AI deployments, and reduce transparency around how AI-related risks are escalated and managed. It might also increase the likelihood that any potential issues go unreported and unresolved if and when they occur.

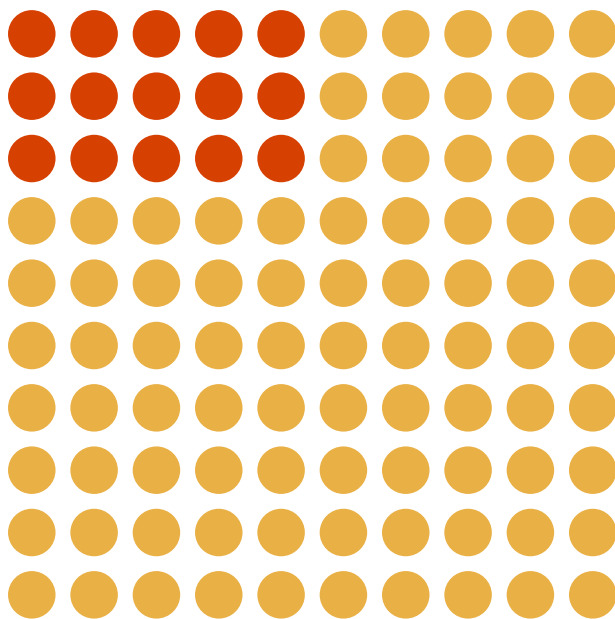
Overall, the data suggests that the global landscape still lacks the safeguards needed to ensure that employees can meaningfully exercise their rights in the context of AI.

For companies that are deploying AI in HR processes, many are lacking data on whether it is inclusive



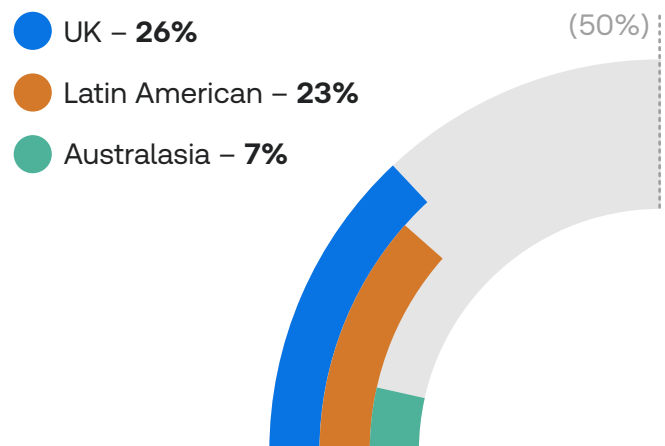
AI adoption in HR remains relatively limited across the sample and is concentrated in recruitment and broad HR toolsets.

15% of companies publicly say they use AI-powered HR tools

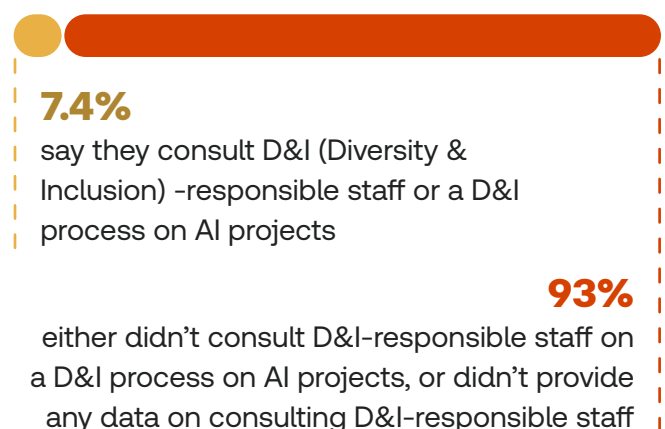


85% of companies did not publicly communicate on this matter

Regionally, UK firms were among the highest of those providing clear evidence of reported HR-AI use, Latin American firms were also more likely to track AI use in HR. Comparatively, companies in the Australasia region were least likely to provide clear data on if AI is being used in HR.

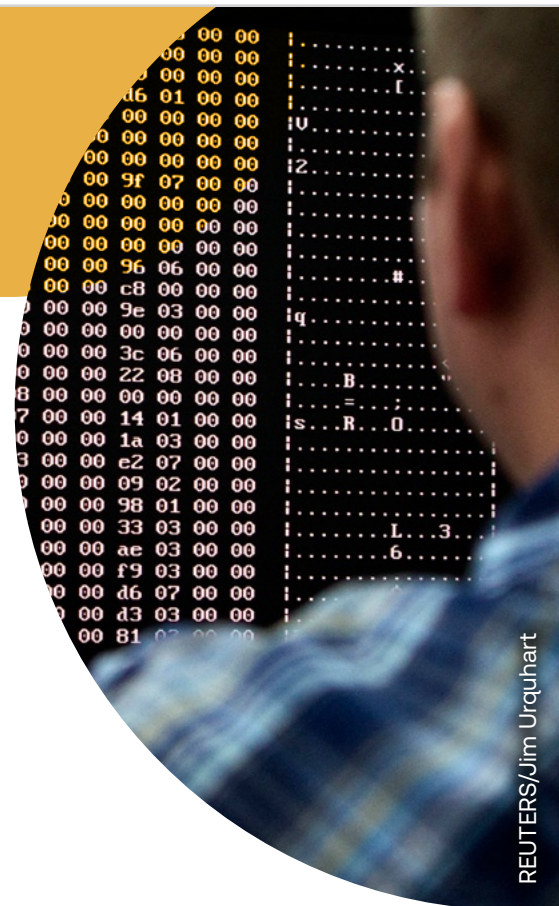


Inclusivity-related governance is rarely evidenced among companies that deploy HR AI. Out of 447 HR-AI users:



This suggests that, even where HR AI is being monitored, the supporting infrastructure for inclusive deployment is often not formalised or not being publicly shared. This substantial gap has several important implications.

- First, it suggests that even where companies have introduced oversight mechanisms for HR technologies, these mechanisms may not be sufficiently equipped to account for the specific risks associated with equality, representation and non-discrimination. Without formalised involvement from D&I specialists, organisations risk overlooking how AI systems might reinforce existing inequities, for example, by replicating historical biases embedded in training data, or by marginalising under-represented groups through automated screening or performance analytics.
- The absence of D&I consultation points to a broader weakness in the governance infrastructure surrounding HR AI. Many organisations appear to be operationalising AI tools without integrating them into established frameworks for inclusive practice. This may reflect limited internal capacity, a lack of cross-functional coordination, or an assumption that existing HR or data-governance processes are sufficient to address fairness risks. However, without explicit oversight from staff trained to identify and mitigate discrimination, companies are less able to anticipate emerging harms or evaluate whether their AI systems align with commitments to workplace equality.
- Finally, the limited transparency around these processes raises concerns about accountability. The absence of public reporting may signal that processes are not in place, or that organisations have not yet recognised the importance of demonstrating inclusive governance to external stakeholders. In either case, the lack of disclosure reduces visibility into how companies are managing the social impacts of HR AI, and limits the ability of regulators, workers and investors to assess whether these technologies are being deployed responsibly.



REUTERS/Kham



• Finding 4:

Ethical issues - including human rights and environmental impacts - are being sidelined in AI governance and risk management

As AI is embedded in critical infrastructures, the social, ethical, and environmental implications of its systems are becoming increasingly visible.³⁷ In many corporate governance approaches, however, “responsible AI” is still framed primarily through the lens of risk management, especially data security, privacy, and compliance, rather than broader ethical concerns such as fairness and bias, or the ways AI can reinforce existing inequalities and marginalisation. At the same time, policymakers, scholars, and business stakeholders are increasingly calling for ecological dimensions to be embedded in corporate AI governance and in leading AI governance initiatives such as the EU AI Act, UNESCO’s Ethics of AI Recommendation and the OECD AI Principles. While these initiatives and corporate programmes promote and ground responsible AI, within them environmental considerations are still not consistently treated as a core governance priority.³⁸

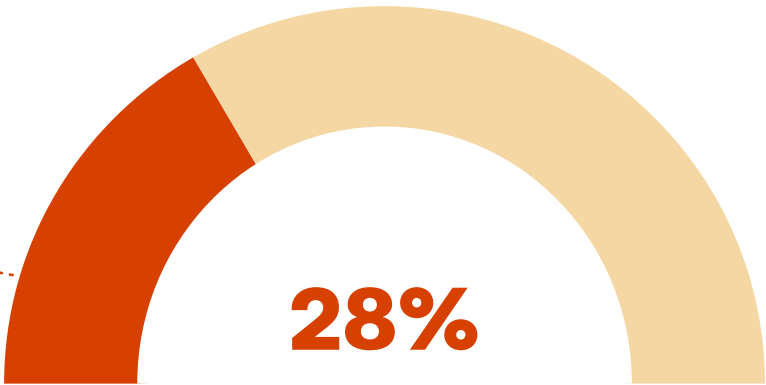
Disclosures indicate that AI governance is largely framed around compliance risks, with limited visibility of ethical priorities

The pattern of disclosure indicates a stronger emphasis on data privacy and security, while ethical, human rights, and environmental impacts are less consistently assessed, less operationalised, and far less independently assured.

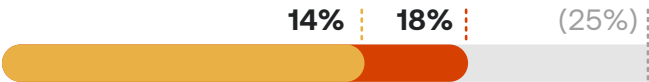
Risk management practices remain privacy and compliance led, with limited visibility of ethics or rights-led.

Only a minority of firms disclose running any form of impact assessment, and the mix strongly favours privacy and compliance.

Firms that report conducting at least one of the listed impact assessment types



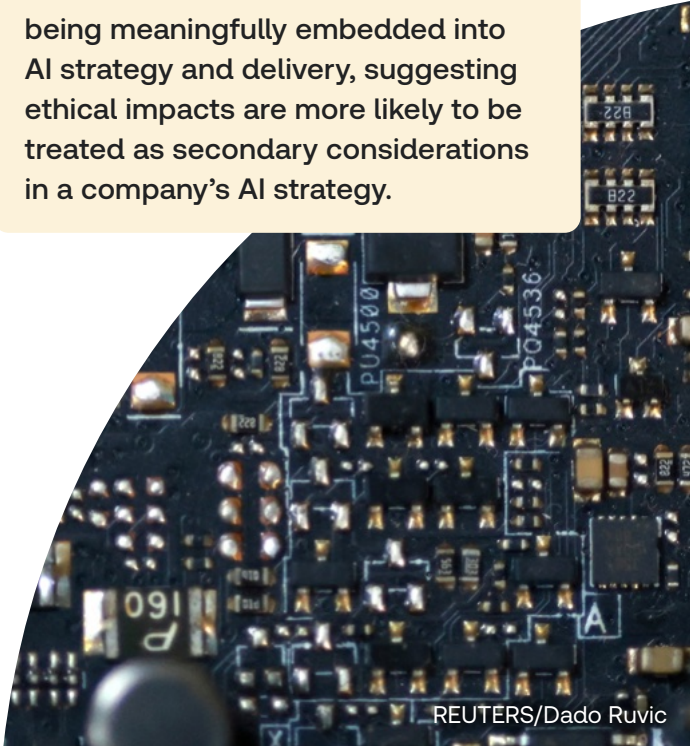
The most prevalent assessments are privacy or compliance-focused, with 18 per cent report conducting a Data Protection Impact Assessment and 14 per cent report conducting a Privacy Impact Assessment.



By contrast, ethics and rights-oriented assessments are notably less common with 7 per cent publicly communicating conducting a Fundamental/Human Rights Impact Assessment and 5 per cent reporting conducting an Ethical Impact Assessment.

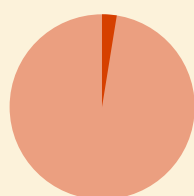


The lack of public disclosure does not confirm inaction, but it leaves stakeholders without credible evidence that ethical and environmental considerations are being meaningfully embedded into AI strategy and delivery, suggesting ethical impacts are more likely to be treated as secondary considerations in a company’s AI strategy.



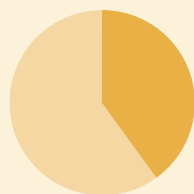
This same pattern of marginalising ethics-specific commitments is also reflected in how firms set up their governance structure.

Few companies report having independent assurances in place to scrutinise ethics-led AI governance.



0.8%

report having an external
AI Ethics Advisory Board

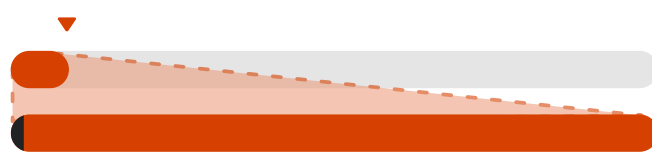


40%

report board/committee-
level oversight of general
AI governance

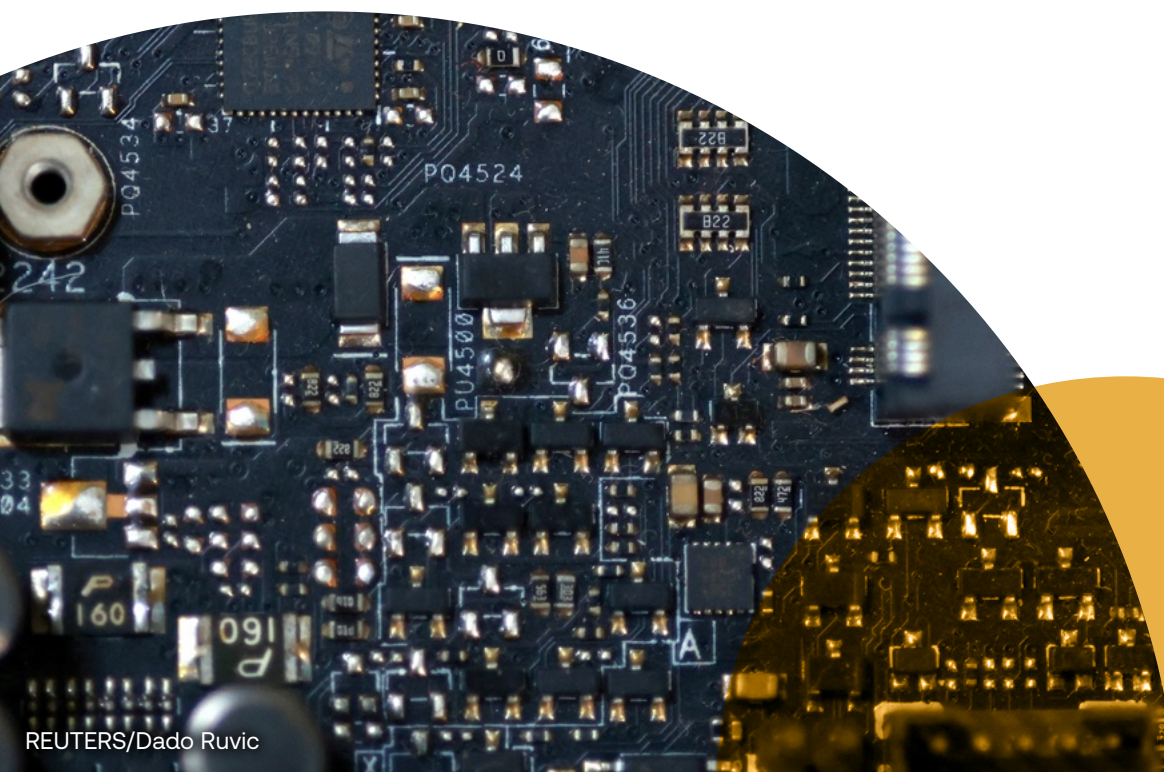
Even in an area where ethical impacts are most scrutinised - bias and fairness - claims of rigorous independent testing are rare.

9% of companies state they have
bias-mitigation controls

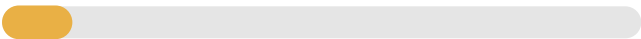


2% of this small cohort cite this
mitigation control as third party audit

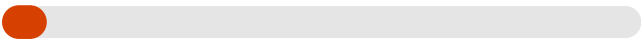
Third-party assurance can help companies bolster legitimacy with stakeholders by strengthening monitoring and control. Publicly disclosing independent assurance of AI governance also signals higher-quality reporting and a credible commitment to managing underlying AI risks, which can build trust and credibility.³⁹



More companies are reporting conducting Environmental Impact Assessments than Human Rights Impact Assessments.



11% report carry out Environmental Impact Assessments



7% report carrying out Human Rights Impact Assessments

That approximately 8 in 9 companies do not report carrying out Environmental Impact Assessments is notable, given the heightened attention from policymakers, scholars, and business stakeholders to how rapidly scaling data-driven systems may affect energy demand, electronic waste, and long-term sustainability.³⁷

These concerns are grounded in the resource intensity of modern AI. Machine learning models rely on energy-hungry cloud infrastructure and data centres, and training large-scale natural language processing models can generate emissions on a scale comparable to the lifetime emissions of several automobiles.³⁷

The overall strain is expected to rise sharply.

Data centres' global electricity consumption:

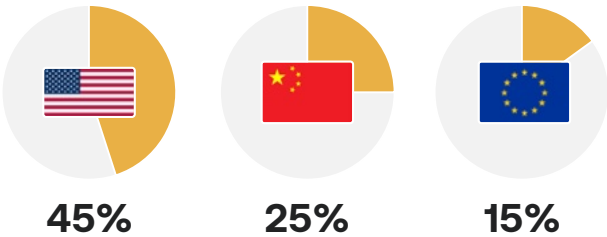
2024 415 TWh

2030 ~945 TWh

The International Energy Agency (IEA) estimates that data centres accounted for around 1.5 per cent of global electricity consumption in 2024 and projects demand will more than double by 2030 (roughly equivalent to Japan's current annual electricity consumption).

Growth is also geographically concentrated.

The United States accounts for the largest share of global data-centre electricity use, followed by China and Europe.⁴⁰



Ethics and environment signals are strongest where governance maturity is already highest.

Across the survey, stronger AI governance practices cluster in large-cap firms and in sectors such as IT, Communications, and Financials, for example, higher rates of impact assessments, clearer oversight policies, the use of AI registries, and more developed bias controls.



By contrast, smaller issuers and sectors such as Energy, Materials, and Real Estate lag on many of these measures. According to the disclosure patterns, the embedding of ethical and environmental considerations follows the same pattern.



REUTERS/Eddie Keogh

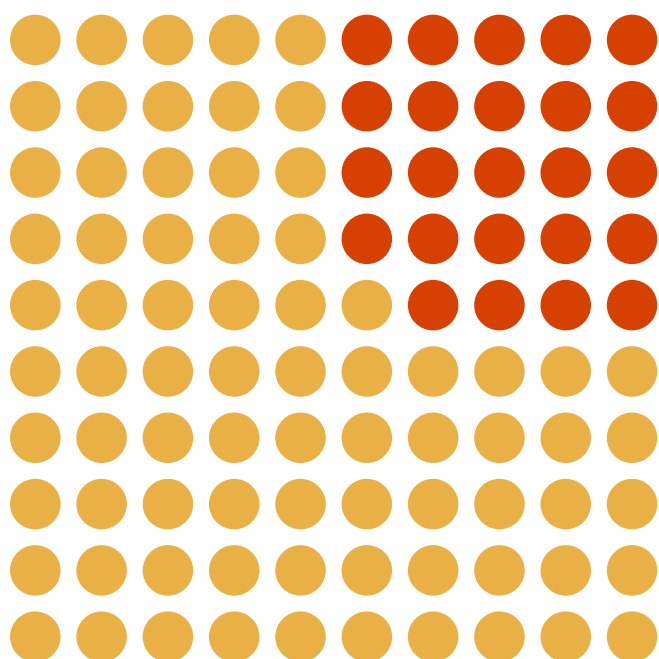
• Finding 5:

Limited company policies on AI training data, third-party data controls, and user data rights

While internal data-privacy and security policies are now well-established across most enterprises, these controls are no longer sufficient for governing the rapidly expanding use of artificial intelligence. AI systems rely on vast quantities of data sourced, combined, and refined through increasingly complex pipelines, often involving multiple third-party providers. Yet, compared with mature privacy and cybersecurity functions, corporate oversight of the data underpinning AI systems remains significantly less developed. Most companies have not embedded structured policies to ensure the quality, provenance, fairness, and appropriate use of the datasets that train or inform their AI models. This gap creates heightened risk: without clear governance of training data, third-party data flows, and user-data rights, AI systems can perpetuate historical inequities, reproduce discriminatory patterns, or propagate non-compliant data across fine-tuning and reuse.⁴¹ As the findings below demonstrate, only a minority of companies have formal oversight mechanisms in place - whether to evaluate training datasets, manage data shared with external AI vendors, or uphold user controls over how their data is processed. Together, these gaps highlight that organisations lack comparably robust governance over the data that fuels AI, leaving material blind spots in accountability, risk management, and end-user protection.

REUTERS/Shannon Stapleton

Around a quarter of companies with AI strategies report having policies to evaluate the quality of data used to train AI systems

**24%**

Oversight of AI training data - the foundation on which models learn - remains significantly underdeveloped across companies. This gap matters because biased, incomplete, or unrepresentative training data can reproduce historical inequities and generate discriminatory outcomes even when systems are designed without intent to discriminate.⁴¹ However, there is limited oversight by companies on the data used to train AI systems in terms of fairness and non-discrimination. Based on publicly available communication from companies, only 24 per cent had policies in place to evaluate the training data of internally developed AI systems and vendor-supplied AI solutions.

Regionally, adoption remains low.



15%

10%

Australasia, North America, Europe and United Kingdom have around 15 per cent of their companies communicating their training-data oversight policies compared with the 10 per cent or less in other regions.⁴²

In Europe, the EU AI Act hard-codes dataset quality and bias controls by requiring high-risk systems to apply data-governance practices and use datasets that are relevant, representative, and are error-free and complete as far as possible.⁴³

In North America, companies face a patchwork of enforceable state and local requirements that indirectly pressure training-data governance. For instance, Colorado's AI Act for high-risk systems requires developer documentation to deployers that includes training-data/data-governance information and discrimination-risk mitigation⁴⁴ and California's CPPA regulations introduce mandated risk assessments and governance duties for automated decision-making affecting consumers.⁴⁵ Additionally, in the US, the NIST AI RMF increasingly functions as "soft law" that shapes corporate norms even without being legally binding: it supplies a shared risk-management vocabulary and encourages repeatable internal controls, documentation, and continuous testing across the AI lifecycle.⁴⁶ By standardising what "good practice" looks like – especially for documentation and evaluation – AI RMF nudges firms toward stronger data governance practices.⁴⁷

Leading frameworks and regulators increasingly treat data governance for fairness as a core requirement and since organisations often procure AI models or services from third-party, training-data oversight must extend to procurement and supplier due diligence-consistent with responsible procurement guidance and human-rights expectations across business relationships.⁴⁸

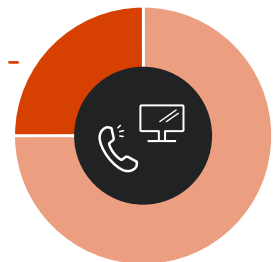


REUTERS/Benoit Tessier

Communication Services and Information Technology sectors lead on AI training-data governance.

25%

of companies in these sectors report relevant policies



The comparatively higher uptake in the Communication Services and Information Technology sectors likely reflects the fact that these industries frequently build or integrate AI products – such as cloud services, software platforms, search and recommendation engines, and advertising technologies – where strong training-data governance is essential for product assurance, customer trust, and safety.



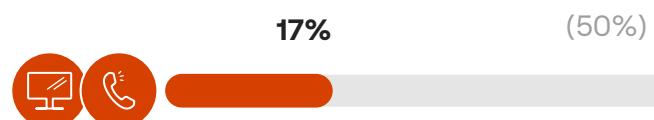
Mirroring gaps in AI training-data oversight, only one in five companies with AI strategies report having policies for data sharing with third-party AI solution providers

Regionally, around 18 per cent of the companies in Australasia and United Kingdom had communication on such policies in place, whereas communication on the same rates drop to 12 per cent or below in all other regions.

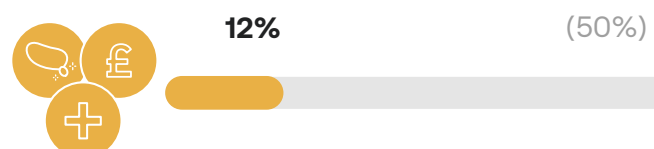


Adoption is similarly limited across sectors

Approximately 17 per cent of the companies in the Information Technology and Communication Services sectors report having relevant policies

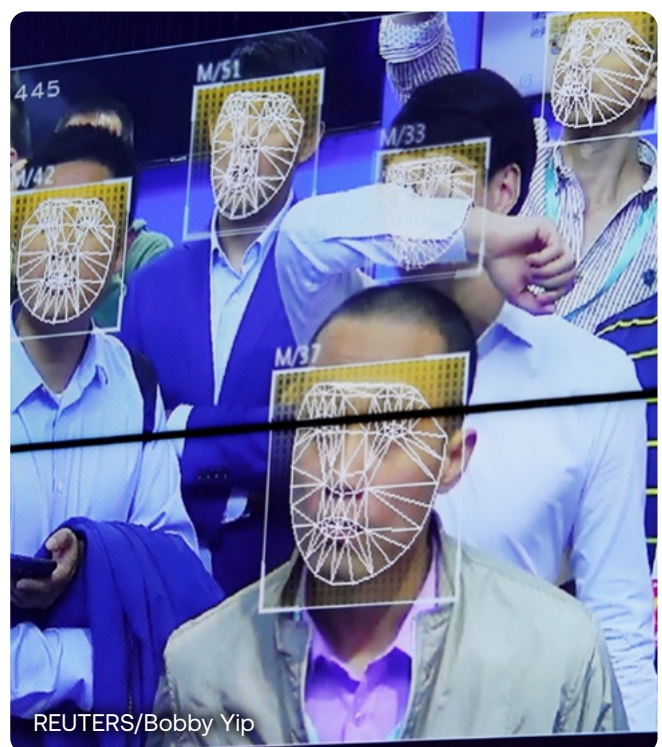


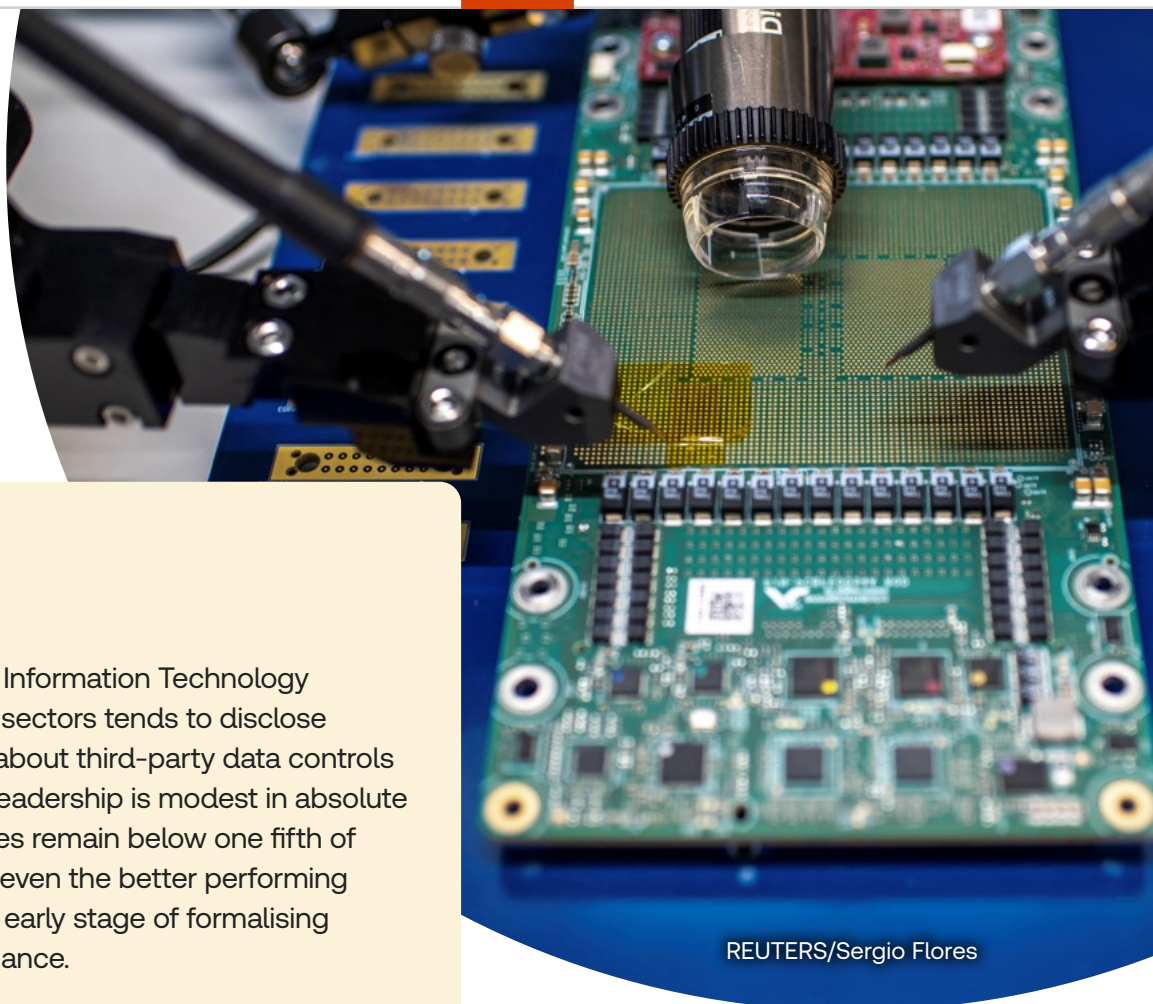
Compared to around 12 per cent or less in the Consumer Discretionary, Financials, and Health Care sectors.



The low adoption rate of such policies could be attributed to the fact that AI systems are produced through complex supply chains with multiple actors and uneven access to information, meaning downstream deployers could lack the power or capability to identify and mitigate upstream data risks.⁴⁹

However, without wider uptake, harmful or non-compliant data can propagate silently across reuse and fine-tuning, undermining accountability throughout the AI supply chain.





REUTERS/Sergio Flores



Companies from the Information Technology and Communication sectors tend to disclose comparatively more about third-party data controls for AI. However, this leadership is modest in absolute terms – adoption rates remain below one fifth of firms, indicating that even the better performing sectors are still at an early stage of formalising supplier-data governance.

IT firms often operate upstream—as cloud/platform providers, model builders, and system integrators, which exposes them to stronger enterprise procurement demands for vendor due diligence and demonstrable governance. In addition, IT companies typically maintain mature software and cybersecurity supply-chain programmes, such as structured vendor onboarding and continuous monitoring processes. This may make it easier for them to extend existing controls to AI components such as dataset provenance, model documentation, and logging-practices explicitly encouraged in risk frameworks that call for third-party AI policies.^{50 51 52}

Communication Service companies, meanwhile, tend to sit downstream as major deployers and distributors of AI-enabled content and services. Their proximity to end users and to copyright, misinformation, and safety risks makes third-party AI controls particularly material.

Yet the data suggest that governance practices in both sectors are only marginally ahead of the wider market, rather than representing a mature benchmark.

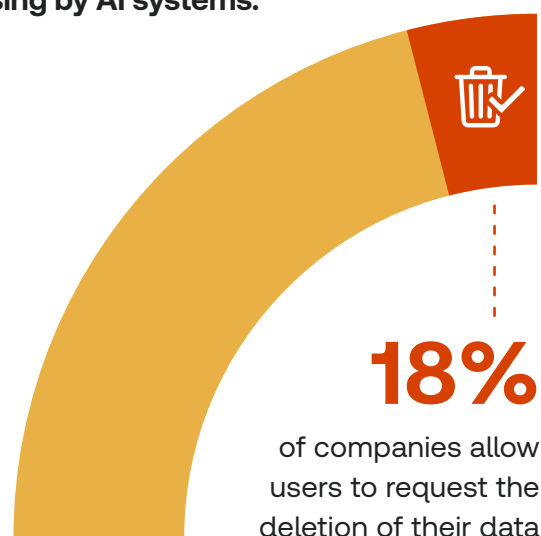
Across non-tech sectors, the lack of formal policies for sharing data with third-party AI solution providers weakens governance because AI is delivered through multi-actor supply chains where information transfer is imperfect and upstream design/data choices create downstream consequences.

Without contractual limits on reuse, provenance requirements, and audit rights, poorly documented training data can embed bias, privacy leaks, and copyright/consent violations that remain opaque to deployers.⁵³ International guidance therefore calls for lifecycle accountability and explicit controls over third-party software/data and supply-chain issues.

Only 16 per cent of the companies allow users to request the deletion of their data and stop its processing by AI systems

As AI capabilities become increasingly embedded across internal workflows such as HR, finance, and risk, as well as in customer-facing services, governance needs to extend beyond privacy policies to protect end-user rights in practice. End-users should receive clear notice of AI use, meaningful transparency and explainability, and accessible pathways to challenge outcomes and obtain remedy, consistent with international guidance on trustworthy, human-rights-respecting AI.^{54 55} Given the widespread reliance on vendor models and third-party datasets, organisations should require supplier due diligence that covers dataset provenance, intended use, quality/bias testing, and ongoing monitoring, aligned with established expectations for responsible business conduct across value chains.⁵⁶

Based on publicly available communications, only 16 per cent of the companies considered in our dataset allow users to request the deletion of their data and stop its processing by AI systems.



25%

of the companies analysed from Europe and United Kingdom provide this mechanism.

Adoption is limited across sectors:



~ 30%

of the companies in the IT and Financial sectors report having relevant policies

Across Europe, companies are more likely to offer user-data rights in the context of AI. This is largely due to the EU's GDPR, which treats personal-data protection as a fundamental right and provides individuals with enforceable controls such as transparency, access, and data-subject rights. GDPR Article 22 also provides safeguards where decisions are based solely on automated decisions that have legal or similarly significant effects and requires safeguards such as human intervention and the ability to contest outcomes.⁵⁷

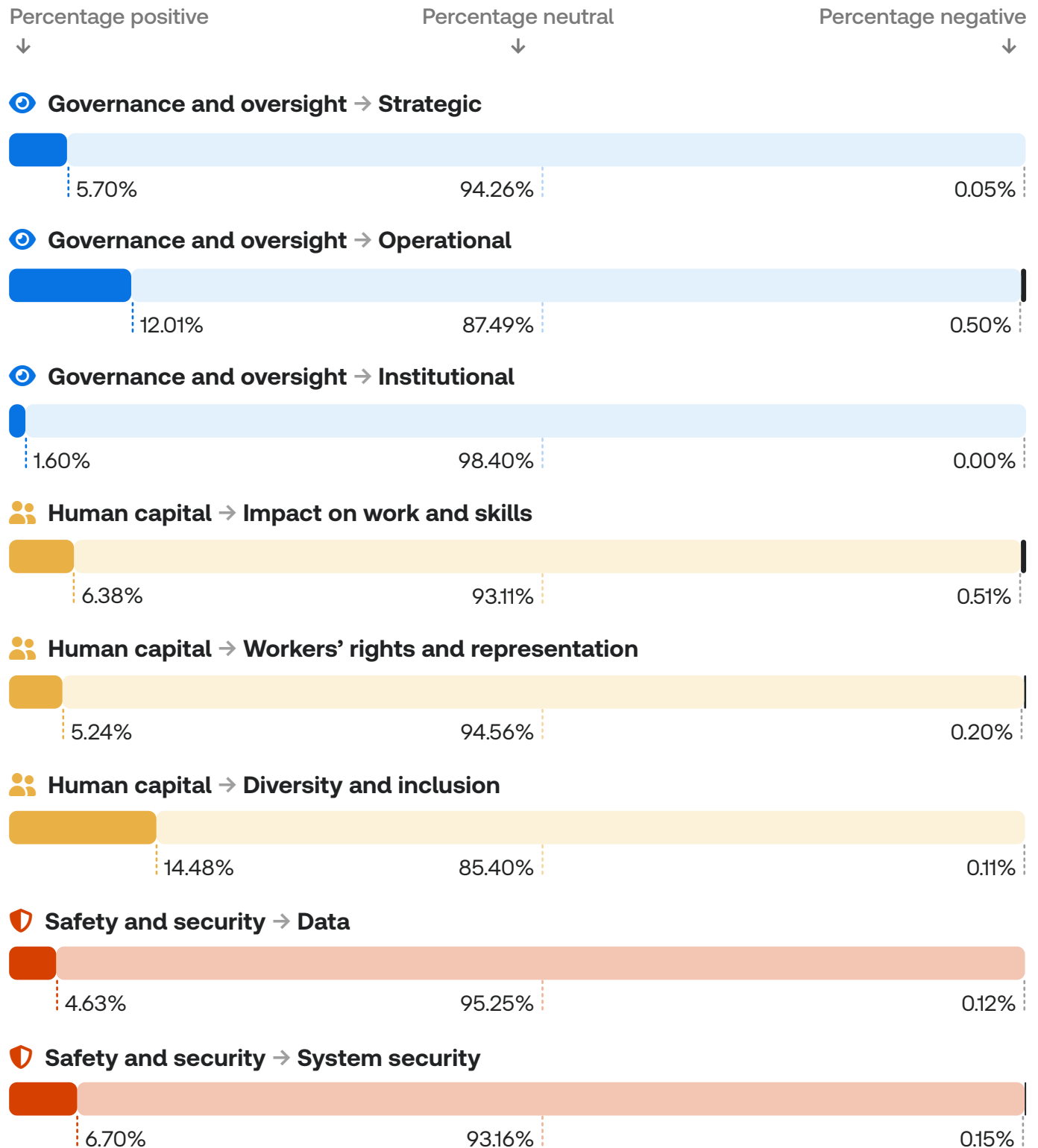
Implications

- If structured governance for AI training data, third-party data flows, and user data rights remains limited, organisations may face increased variability in system outcomes, including uneven performance across populations and disparate impacts linked to biased or unrepresentative datasets.⁵⁸ Weak dataset provenance and documentation can also make it harder to audit or explain how data was collected, prepared, and used, which can constrain internal assurance, external review, and incident response when issues arise.⁵⁹
- In extended AI supply chains, insufficient controls over third-party inputs may contribute to risk propagation across reuse and integration and can complicate risk measurement when upstream practices are not transparent to downstream deployers. Over time, early data issues may compound into downstream operational failures and accumulating technical debt (“data cascades”), particularly in higher-stakes contexts.⁶⁰
- Finally, as some jurisdictions increasingly specify dataset governance expectations for certain AI, gaps in policy and evidence trails may translate into higher compliance and assurance costs and greater scrutiny during assessments.⁶¹

• Sentiment analysis



Sentiment analysis results



The survey's more than 15,000 open text fields offer a unique window into how companies talk about artificial intelligence.

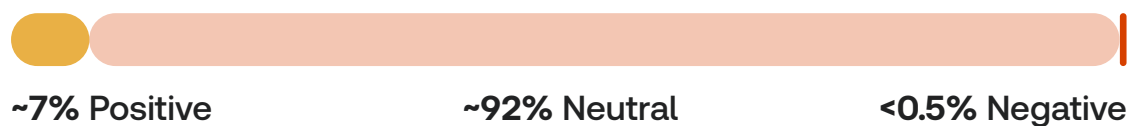
Unlike closed questions, these free-text answers let us read between the lines. They show that, although enthusiasm for AI and its commercial promise is high, the tone becomes markedly more conservative when sampled companies discuss AI governance and responsible AI. This shift in language is itself a finding: it suggests that companies recognise the complexities of governing AI and choose words carefully when commitments have yet to be fully realised.

We processed these responses using two transformer-based models – Financial DistilRoBERTa and Twitter RoBERTa – which assign each answer a score between -1 and +1. Scores above 0.05 were considered positive and those below -0.05 negative. A third lexicon-based tool, VADER, was also tested but proved ill-suited to formal corporate communications as a rule-based classifier it can only recognise scenarios it was designed for and tends to overrate compliance language⁶². The RoBERTa models, by contrast, better capture the neutral, cautious tone typical of governance disclosures.



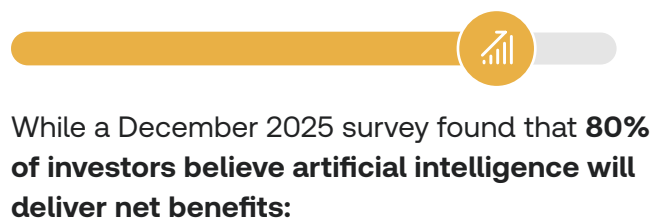
Sentiment patterns

Across all sections the prevailing tone was neutral to slightly positive.



This neutrality reflects how organisations write about responsible AI – measured, careful and often aspirational – rather than exuberant.

Such caution mirrors investor sentiment:



When board directors themselves were polled:



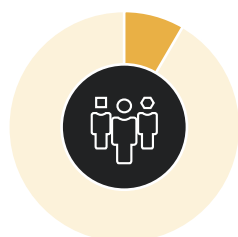
The restrained language in our data therefore reflects a market in which stakeholders welcome AI but insist on responsible safeguards.



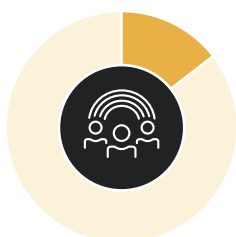
Within this overall neutrality, some themes stand out.

Human capital questions – covering training, inclusion and employee support – registered the highest positive share. Categories such as Diversity and inclusion and Operational issues had even higher positivity.

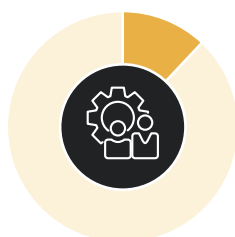
In contrast, governance and oversight and safety and security displayed lower positive shares and higher neutrality.



Human capital
8.4%



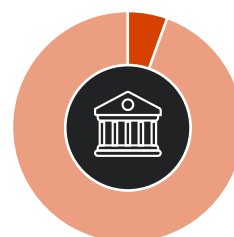
Diversity & inclusion
~14.5%



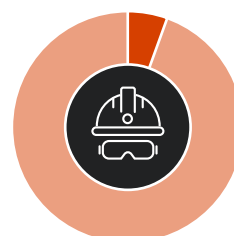
Operational
12%

Respondents regularly referred to diversity programmes, re-skilling initiatives and inclusive workplaces, suggesting that many organisations simply adapted existing human-resources policies to the AI context.

That emphasis on people is echoed in broader opinion surveys: investors overwhelmingly expect companies to equip their employees with AI training⁶⁴, and many directors say they intend to devote more resources to upskilling⁶⁵. Taken together, the internal numbers and the broader mood of the market suggest that workforce support is one of the more mature areas of responsible AI.



Governance & oversight
~5.5%



Safety & security
~5.5%

Rather than implying disengagement, this pattern reflects the early stage of many organisations' AI governance frameworks.

According to EY's proxy statement analysis, the number of S&P 500 companies assigning AI oversight to a board committee more than tripled in 2025.

However only about one quarter have fully implemented AI governance programmes

And just 67% conduct formal risk assessments of third-party models⁶⁵

Our respondents' guarded language on governance and security therefore signals realism: companies are aware of regulatory and ethical obligations but are still building committees, policies and controls.

In this light, neutral sentiment is not a negative sign; it shows that firms are careful not to overpromise on issues like data protection, system security and board oversight

Interpreting positive and neutral signals

Taken together, these findings remind us that polished rhetoric does not always equal substantive progress.

The CFA Institute warns of “AI washing,” where organisations market products as AI-driven without meaningful integration⁶⁶. In our data, high positivity often clusters around topics – such as diversity or operational efficiency – where companies already have well-developed programmes and can easily repurpose existing messaging. By contrast, areas that require new technical capacity, risk management and board oversight tend to be described in neutral terms. Far from signalling complacency, such neutrality can indicate transparency about challenges and a willingness to build the foundations required for responsible AI.

For investors and corporate leaders, the sentiment landscape offers actionable signals.

Strong positivity around human capital and diversity points to areas where firms are most confident and where investments in training and inclusion are paying off. Neutral language around governance and security highlights the parts of the agenda that remain under development. By combining sentiment analysis with other metrics, investors can better distinguish between marketing slogans and genuine progress.

The ability to read these nuances makes sentiment analysis a valuable complement to more quantitative assessments when evaluating a company’s AI readiness and sincerity.



AICDI company case studies



REUTERS/Rula Rouhana



A special thanks to 14 companies we engaged with

We would like to extend our sincere thanks to the inaugural group of companies that actively shared their responsible AI practices with us during AICDI's very first year. They demonstrate distinct leadership and a commitment to transparency on these crucial topics. You can read more below in the company case studies about how these companies are handling the challenges around responsible AI deployment.

This group represents a broad range of countries and sectors, underscoring the collaborative effort behind developing responsible AI practices.



The case studies presented in this section are intended to illustrate practical approaches and experiences shared with the Initiative. Inclusion of a company does not constitute an endorsement, rating, or validation of its overall performance, practices, or products.

All information is derived from company engagement and materials voluntarily provided by participating organisations, supplemented where relevant by publicly available disclosures. The content reflects the information shared at the time of engagement and should not be interpreted as a comprehensive or independently verified assessment.

CASE STUDY 1: TELUS



Diversity and inclusion



TELUS is a communications technology company that uses, develops and conducts AI-related research. In this context, the organisation considers diversity and inclusion in its AI projects and teams developing or implementing AI systems. They also work on ensuring that the systems are adapted for use across diverse national, regional, linguistic, and cultural settings.

TELUS has approached this by incorporating diversity and inclusion considerations in its AI projects. They employ a cross-functional 'purple team' AI testing approach that emphasises the participation of diverse individuals with varying expertise and technical literacy to gain comprehensive insights into any system shortcomings, allowing them to effectively mitigate risks. They also publish an annual research report that captures perspectives on AI from more than 11,000 Canadians and Americans, with special attention to historically underrepresented communities, highlighting the importance of including a wide range of voices to build trustworthy AI.

Additionally, TELUS has worked with PLATO Consulting, an Indigenous-owned software company, to set up an extended purple team and support Indigenous AI workforce capacity and skill-building. The organisation conducts public purple teaming events to engage diverse stakeholders in analysing and assessing the fitness of their generative AI systems. TELUS' Indigenous Advisory Council provides guidance on matters related to AI ethics issues, particularly those affecting Indigenous Peoples. Notably, with guidance from the Council, they expressed concerns that AI-generated images of Indigenous Peoples, such as First Nations, Inuit, and Métis Peoples, may perpetuate stereotypes, inaccuracies, and offensive representations. In response, TELUS publicly committed to not using AI to generate images or art of Indigenous Peoples.

Reference: AICDI (AI Company Data Initiative) 2025 disclosure – Diversity and inclusion, questions 2.7, 2.10 and 2.11; TELUS Annual Report⁶⁷



CASE STUDY 2: VODAFONE

Data, systems and cybersecurity

As a telecommunications company integrating AI across its products, services and business operations, Vodafone faces several key considerations related to safeguarding the security of the data processed within these systems; and ensuring that the data rights of their end users is upheld. There is also a focus on maintaining the safety and security of the AI systems themselves while also identifying the risks stemming from AI to develop effective mitigation strategies.

To address these areas, Vodafone has adopted a comprehensive set of policies supported by practical organisational measures. Central to this effort is the company's Artificial Intelligence Framework⁶⁸, which dictates that their AI systems must carefully manage customer data in alignment with their privacy commitments and prevailing legislations.



This framework complements Vodafone's global privacy management policy, which focuses on respecting local data protection and privacy laws while setting a baseline for those markets where there are no equivalent legal requirements. The framework also commits to respecting end-users' fundamental rights, which includes

data deletion and control over processing of their data. The framework further ensures that customer data is used in AI systems only when a clear legal basis is established, which inherently supports data subject rights while also showing commitment to protect the security of individuals served by AI.

In practice, Vodafone supports this policy through a well-defined governance structure dedicated to addressing emerging AI risks. An AI Governance Board comprising senior leadership is responsible for AI strategy, policy, and threat mitigation. The Governance Board is supported by the Responsible AI Office which along with their "Secure and Privacy by Design" teams, ensures compliance and ethical use of AI while also addressing new risks as they emerge. Their AI framework is a dynamic document which is regularly reviewed and updated to reflect new products, technological developments, and learnings, ensuring that emerging AI risks are considered and mitigation strategies are developed. The governance model also places focus on multidisciplinary approaches, whereby the company enables collaboration between their internal cybersecurity, legal, and privacy experts to inform its AI risk management strategies. Additional safeguards include pseudonymisation and permission-based use of customer data used in AI systems; and application of general cybersecurity measures to the infrastructure that supports AI data.

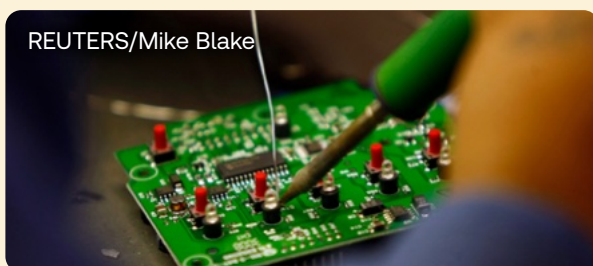
Reference: AICDI (AI Company Data Initiative) 2025 disclosure – Safety & security, Data, questions 3.6 – 3.10; Vodafone AI Framework⁶⁸

CASE STUDY 3: SAP



Impact on workers

SAP is an enterprise software company. They consider the impact of AI on its workforce by offering training and reskilling programmes that help employees adapt to an AI-integrated workplace. They have also established policies to mitigate potential negative effects of AI systems on workers, provides AI ethics training to key stakeholders, and conducts regular surveys to assess employees' awareness of and engagement with ethical AI practices.



SAP employs a fundamentally human-centric approach to integrating AI into the workplace. They emphasise empowering their workforce rather than displacing it and this commitment is carried out through a robust framework of policies and proactive measures. In practice, the company intentionally designs and deploys its internal AI systems with a “human-in-the-loop” architecture, where AI is used to automate repetitive tasks, provide data-driven insights, and support decision-making — while final judgment, critical thinking, and complex problem-solving remain firmly in the hands of employees. The organisation has a formal AI onboarding process required for every new AI feature. This programme goes beyond standard practice by providing specific materials on the AI system's operation and its impact on user workflows. A significant part of this process is dedicated to SAP's Global AI Ethics Policy, which covers critical

ethical considerations such as fairness, transparency, and human oversight.

SAP has committed to building an AI-integrated workplace through a wide array of training and reskilling programmes. These offerings cater to diverse needs, ranging from foundational AI literacy for all employees to advanced, role-specific curricula. Technical staff receive in-depth training in machine learning, data science, and prompt engineering, while business-focused teams engage in programmes that emphasise the value, application, and ethical implications of SAP's AI solutions, reinforcing the company's Responsible AI principles. The organisation provides comprehensive and mandatory training for all relevant staff, ensuring that its Guiding Principles for AI are deeply embedded across the entire AI lifecycle. The educational framework is designed to equip employees with a thorough understanding of AI's purpose, functionality, and limitations, while clearly defining their ethical responsibilities.

SAP considers AI ethics training a fundamental and mandatory component of its development and operational culture. This training is a required part of the onboarding process for all employees involved in the design, development, deployment, and management of AI systems. They actively and continuously assess the level of AI ethics awareness and practices among its staff through various engagement and feedback mechanisms.

Reference: AICDI (AI Company Data Initiative) 2025 disclosure – Impact on workers, questions 2.1 to 2.4, SAP Integrated report⁶⁹



CASE STUDY 4: TELEFÓNICA

AI governance, strategic & institutional

Telefónica is a multinational telecommunications company that company has established AI strategies and guidelines with management oversight and ensures these are accessible to all employees. In addition, Telefónica regularly reviews and updates its governance mechanisms, and engages with industry, regulatory, academic, and civil-society leaders on technological and ethical issues.

Telefónica's AI strategy and guidelines—particularly the Telefónica Artificial Intelligence Principles: AI Code of Conduct—assess the societal impact and potential harms of AI systems beyond direct users. These principles ensure respect for human rights across all contexts in which AI is applied, inherently accounting for broader societal implications. The Responsibility by Design project and the AI Governance Model include processes for assessing risks such as discrimination and bias. Their AI Principles, which were approved by their Board of Directors in 2024, form a core part of their governance approach. Their AI Governance Model is overseen by the Global Compliance Officer through the Digital Compliance & DPO. They also have Responsible AI Champions within their business units to ensure the responsible use of AI. Their company strategy and guidelines on AI involve multiple levels of management and specific roles including the Board of Directors, the Global Compliance Officer, the Global Sustainability (ESG) Office, Responsible AI Office and other specialised teams.

They periodically review and adjust their internal policies, including AI-governance-related internal regulations, where the Regulation of the Governance Model on Artificial Intelligence is subject to review and updates by the Global Compliance Officer.

Policies are generally reviewed on a two-yearly basis or as needed. The review involves stakeholders such as the Board of Directors, its various committees (Audit and Control; Sustainability and Regulation), the Global Compliance Officer, and relevant management areas (e.g., Global Sustainability (ESG) Office, Digital Security, General Secretary, Chief Data Office, Global Technology Officer, and the Digital Innovation unit). Feedback from these reviews—along with emerging risks, regulatory changes (e.g., the EU AI Act), and lessons learned—is integrated to make necessary adjustments to AI governance.

Telefónica ensures that its AI strategy and guidelines are accessible to employees and encourages their review through various means, including Training and Awareness programmes, Internal Policies, and other internal channels.

The company actively engages in technological and ethical exchanges with various leaders by collaborating with international bodies such as the EU AI Office, UNESCO, and GSMA. Through the International Chamber of Commerce of Spain, Telefónica is involved in the UN Global Digital Compact, which focuses on mitigating AI risks for all consumers and users—especially the most vulnerable—demonstrating the company's consideration of broader societal impacts. They additionally take part in OECD's AI working and expert groups, contributing to discussions on AI governance and best practices.

Reference: AICDI (AI Company Data Initiative) 2025 disclosure –AI governance - strategic and institutional, questions 1.3 to 1.8; Telefónica Annual Report⁷⁰

CASE STUDY 5: BANCO BRADESCO, PRUDENTIAL, BASF, INFOSYS, TELEFÓNICA, TELUS

AI Skills Training

Companies across different sectors employ diverse training and reskilling strategies to prepare their workforce for an AI-integrated workplace.



Banco Bradesco, a financial institution, provides technical training in artificial intelligence, programming, and data science, as well as soft skills, to its employees through a partnership with the tech-focused school, Alura. It also offers a 'digital transformation' programme that covers key topics such as AI agents and leadership in transformation. This programme is designed for employees across different Bradesco business units.



Prudential, a financial institution, offers an 'AI for All' training programme that has educated around 5,000 employees on the power of AI and how it can reimagine their work. This programme provides both technical and awareness-focused AI training. AI skills are part of their structured training programmes, which are delivered through a centralised learning platform, targeting employees across the organisation.



BASF, a diversified chemicals company, has jointly agreed with its workers' councils on a general reskilling programme covering technical, hard, and soft skills. This is complemented by AI-specific training offerings delivered through their Data & AI Academy.



Infosys, an information technology company, offers a range of AI-specific training, including generative AI courses, prompt engineering, and AI applications for business growth. Their curriculum includes GenAI-powered professional skill simulators and content on cloud computing and GenAI technologies developed in partnership with hyperscalers. They have an internal learning platform called Lex, which delivers their training programmes to all employees. Targeted programmes support new hires to ensure they are adept in new skills, existing employees for reskilling, and leadership groups, such as the 'more than 250 women leaders' who completed a specialised certificate course on AI applications.



Telefónica, a multinational telecommunications company, offers general training courses to all employees, providing an overview of AI, the company's code of conduct, and Telefónica's AI governance model. In addition, awareness-raising sessions are conducted for business areas to promote ethical practices throughout the AI lifecycle—from design to deployment—ensuring alignment with applicable regulations, such as the EU AI Act. The company also provides a specialised training programme for Responsible AI Champions who lead AI efforts within each business unit. This programme deepens their understanding of both the technical and ethical aspects of AI, equipping them to advocate for responsible AI practices. To support day-to-day application, Telefónica has published two guides accessible to all employees: the Guide to Generative AI for Employees AI, focused on responsible usage, and the Guide to Responsible AI for Business Areas, which supports the ethical development of AI-based products and services.



TELUS, a technology company, offers multiple types of training to help its employees adapt to an AI-integrated workplace. This includes the Data Steward Certification program, which provides Data Stewards with a customized certification training program that equips them to understand the responsible use of AI through data privacy, security, and governance. TELUS also supports team members in achieving the AI Governance Professional Certification offered through the International Association of Privacy Professionals' AI Governance Professional certification. Additionally, the organisation provides ethical machine learning training. The Data & Trust Office includes resources trained in ethical machine learning techniques, demonstrating specialized technical training for AI-related roles. In addition to tailored training for specific job functions, TELUS runs a Data & AI Literacy program for all team members to upskill and understand the opportunities, governance, and risks of data and AI use.

Reference: AICDI (AI Company Data Initiative)
2025 disclosure – AI training, question 2.1

CASE STUDY 6: GRUPPO TIM, BASF, TELUS, TELEFÓNICA

Environmental considerations

Across industries, companies are adopting a variety of strategies to monitor and reduce the environmental impact of the AI systems they develop or deploy throughout the AI lifecycle. As AI models become more complex and computationally intensive, their energy consumption—particularly during training and inference—can contribute significantly to an organisation’s carbon footprint. To address this, companies are increasingly considering environmental impacts across the AI lifecycle, including model design, datacentre energy use, cooling efficiency, and the carbon intensity of electricity sources.



Gruppo TIM, a telecommunications company, references its “Verso una Greener AI” article, which recognises that AI systems consume significant amounts of energy, and highlights the need for optimised architectures and infrastructure to reduce this footprint. Since AI systems require substantial computational resources, Gruppo TIM’s approach to reducing energy and environmental impact applies across the full AI lifecycle – design, development, deployment and operation. By sourcing renewable energy and improving energy efficiency, the organisation reduces the operational footprint of AI inference and training. The company works to ensure that data-intensive services, including AI, are delivered more sustainably by monitoring eco-efficiency indicators related to data traffic and infrastructure. Additionally, by investing in lower-energy infrastructure such as fibre networks and better cooling systems, they address environmental impact from the design and infrastructure side of the lifecycle. The company has committed to achieving 100 per cent renewable electricity by 2025, sourcing energy from renewable sources and investing in on-site electricity generation.



BASF, a diversified chemicals company, continuously monitors the energy and CO₂ emissions associated with its internal AI use through its cloud subscriptions. They have determined that the energy and CO₂ impact of this AI use is currently negligible compared with their overall electricity consumption and corporate carbon footprint. The company is also increasing the share of renewable electricity in its total energy consumption and is using AI to reduce both energy use and CO₂ emissions. To support these efforts, the organisation trains machine-learning models on historical and simulated data, enabling its production plants to operate more efficiently within defined limits.



Thomson Reuters Foundation/Fabio Cuttica



TELUS has established environmental sustainability as a foundational pillar of AI development through its Sovereign AI Factory in Rimouski, Quebec – Canada’s first fully sovereign AI facility. The factory is powered by 99 per cent renewable energy in a LEED Gold certified data centre, reflecting a design philosophy that prioritizes carbon reduction while delivering state-of-the-art AI compute infrastructure.

With a Power Usage Effectiveness (PUE) ratio of 1.15, the facility is three times more energy efficient than the industry average for excess power usage with annual energy savings of 10.6 million kilowatt-hours of energy – enough to power 915 households and reduce 329 tons of carbon emissions. TELUS employs innovative natural cooling systems that leverage Canada’s climate, requiring mechanical cooling for only 40 hours per year and operating in free cooling mode 98% of the time. This reduces water consumption by more than 75% compared to traditional data centres — using 0.23 litres per kilowatt-hour versus the industry average of 1.8 litres and saving 17 million litres of water annually. In addition, TELUS utilises an AI-driven energy optimization system to reduce electricity consumption in its data rooms, and it discloses information on its AI use through their annual sustainability reporting, demonstrating an ongoing commitment to transparency regarding environmental impacts.



Telefónica employs a range of environmental and energy-efficiency strategies across the infrastructure that supports its AI systems. These strategies include an energy efficiency plan aimed at reducing overall energy consumption, including that of data centres and networks hosting AI systems; and a renewable energy plan, designed to lower the carbon footprint of all operations, including that of AI systems powered by its infrastructure.

Through its EcoSmart Services, the company works to reuse network and customer equipment, reducing the need for new manufacturing and thereby lowering the embodied energy and environmental impact of the hardware that AI systems run on.

Telefónica monitors the environmental risks of AI systems through its AI governance tool, ensuring that sustainability is assessed throughout the entire lifecycle. Furthermore, it provides training on environmental impact monitoring solutions and has developed its own tool, Kiri, which measures the CO₂ footprint of AI systems and is available to AI solution owners to support mitigation actions.

Reference: AICDI (AI Company Data Initiative) 2025 disclosure – environmental considerations, question 1.16

CASE STUDY 7: BANCO BRADESCO, SAP

Review of AI governance mechanisms

Companies employ a range of approaches to reviewing and updating their AI-related governance mechanisms. As AI technologies advance rapidly, regular review processes help organisations ensure that their systems continue to comply with emerging regulations, identify and mitigate emerging risks, and uphold ethical standards across the AI lifecycle. These governance mechanisms ultimately help organisations maintain the reliability and integrity of their AI systems as both internal needs and the external landscape continue to evolve.



Banco Bradesco, for instance, maintains a defined review cadence to ensure the continued relevance and effectiveness of its key AI directives. Its Generative AI Framework is updated quarterly to integrate the latest technological developments and platform enhancements. The Corporate AI Policy undergoes a formal annual review to ensure alignment with the evolving regulatory landscape and ethical commitments. Additionally, the overarching Corporate AI Strategy is continuously assessed by executive leadership to ensure its alignment with strategic business goals and market dynamics.



REUTERS/Thomas Peter



SAP has an AI governance framework that undergoes a comprehensive review on a biannual basis, ensuring strategic alignment and thorough assessment. This is supplemented with provisions for ad-hoc reviews, which may be triggered by new regulatory developments, significant technological shifts or the launch of high-impact AI initiatives. The review process is closely integrated with key stakeholders across the organisation. Strategic oversight is provided by the steering committee, including top management, which convenes biannually. Quarterly reviews are held with management from various business and technology domains to maintain cross-functional alignment. Additionally, an operational AI steering committee, composed of technical experts and project leads, meets monthly to address tactical challenges and monitor ongoing performance. Feedback and insights from all reviews are systematically channelled to the designated process owners and subject matter experts responsible for the AI governance framework.

Reference: AICDI (AI Company Data Initiative) 2025 disclosure – Review of AI governance mechanisms, question 1.5



China Daily via REUTERS

CASE STUDY 8: CEMENTOS ARGOS, BASF

Workers' rights

As companies increasingly adopt AI-enabled processes, safeguarding employees from unintended harms—such as unjust monitoring or biased decision-making—has become an essential part of responsible AI governance. As a result, companies are beginning to embed human-rights-based safeguards and governance structures into their AI deployment processes, including consultation with employee representatives, transparent communication, and assessments of potential impacts on worker well-being. To look at how companies ensure that AI tools used in the workplace do not infringe on workers' rights, this section highlights the approaches taken by Cementos Argos and BASF.



Cementos Argos, a conglomerate, has a human rights policy and a commitment to labour practices that guide its actions in protecting the human rights of both internal and external stakeholders. These policies and commitments form the foundation of all company policies, including those related to Artificial Intelligence, which has ethics as its foundation and prioritise people's safety and well-being.



BASF ensures that each use case of AI tools and the systems in which they are implemented is in alignment with employee representatives. The company also monitors and assesses data processing activities to ensure rightful processing of personal data.

Reference: AICDI (AI Company Data Initiative)
2025 disclosure – Worker's rights, question 2.5

Guidance for investors

Our findings show that artificial intelligence has moved from hype to an operational reality, yet governance has not kept pace. Many companies articulate about AI ambitions without explaining how they will be implemented, and only a small share publicly align with recognised frameworks. Recent corporate disclosures, including management commentary, are beginning to caution investors about the risks associated with AI deployment.⁷¹

For investors, this means exposure to operational errors, regulatory penalties and reputational damage. This section distils the implications of those gaps for portfolios and offers practical ways to respond.

Implications

The patterns in the findings translate into several areas of exposure for portfolios:



Operational and legal risk

Because few issuers have formal model registries or conduct impact assessments, algorithmic errors, bias and privacy breaches are more likely to escape detection until they trigger lawsuits, fines or product recalls.



Regulatory headwinds

Emerging regimes such as the EU AI Act emphasise data quality and bias management, and corporate filings are beginning to warn about AI-related risks.⁷¹ Investors should expect higher compliance costs for companies that lag behind.



Sectoral and size disparities

Technology, communications and financial services firms, especially large-caps, generally display more developed governance structures. Energy, materials and real estate companies are less prepared, which should inform sector weighting and risk premiums.



Human-capital vulnerabilities

Many issuers invest little in training or reskilling programmes and lack mechanisms to protect workers or handle grievances. This heightens the risk of operational disruptions and labour disputes.



Data governance gaps

A minority of companies oversee training-data quality or third-party data practices. Since many AI failures originate with vendors, weak oversight introduces hidden risk.⁷²



Ethical and environmental blind spots

Few firms commit to ethical frameworks or environmental impact assessments. High-powered AI models are energy-intensive,⁷² and ignoring ethics invites regulatory and reputational backlash.

Investor engagement checklist

Drawing on the thematic and sentiment findings, this checklist distils the questions investors should ask at any stage of the investment cycle. The prompts below translate governance gaps into a practical framework for both pre-investment screening and post-investment engagement:

Question to ask companies	What it tells you about their approach to responsible AI										
Governance and frameworks Checks whether the company has clear leadership, responsibility, and written rules for how AI decisions get made. <table><tr><th>QUESTION</th><th>WHAT IT TELLS YOU</th></tr><tr><td> Do you have a company strategy and / or guidelines on AI? → AICDI survey question 1.4 </td><td> A defined AI strategy signals the company is managing AI intentionally, with clear ownership, aligned goals, and governance (risk controls, monitoring, documentation) to scale use safely and accountably. </td></tr><tr><td> Does your company adhere to any self-regulatory codes of conduct, voluntary frameworks, commitments, guidelines, or internationally recognized technical standards in relation to the ethical development and deployment of AI, recognised by reputable industry bodies or relevant regulatory authorities? → AICDI survey question 1.2 </td><td> Alignment with recognised AI frameworks and standards signals governance maturity and helps assess consistency, risk controls, and regulatory readiness in ethical AI development and deployment. </td></tr><tr><td> Do you have a company-wide designated board, committee or person(s), or similar bodies designated to review issues of accountability and responsibility, and other ethical issues? → AICDI survey question 1.7 </td><td> Formal arrangements signal mature AI governance, clear accountability and consistent controls (risk checks, testing, documentation, security, incident handling) to ensure AI is used safely, ethically, and compliantly. </td></tr><tr><td> Do you have dedicated resources with the authority and responsibility to ensure the ethical, safe, secure and trustworthy development and use of AI within the company? → AICDI survey question 1.8 </td><td> If a company can assign responsibility at every AI lifecycle stage to specific people or entities, it shows clear, accountable AI governance; if not, accountability is likely fragmented and oversight weaker. </td></tr></table>		QUESTION	WHAT IT TELLS YOU	Do you have a company strategy and / or guidelines on AI? → AICDI survey question 1.4	A defined AI strategy signals the company is managing AI intentionally, with clear ownership, aligned goals, and governance (risk controls, monitoring, documentation) to scale use safely and accountably.	Does your company adhere to any self-regulatory codes of conduct, voluntary frameworks, commitments, guidelines, or internationally recognized technical standards in relation to the ethical development and deployment of AI, recognised by reputable industry bodies or relevant regulatory authorities? → AICDI survey question 1.2	Alignment with recognised AI frameworks and standards signals governance maturity and helps assess consistency, risk controls, and regulatory readiness in ethical AI development and deployment.	Do you have a company-wide designated board, committee or person(s), or similar bodies designated to review issues of accountability and responsibility, and other ethical issues? → AICDI survey question 1.7	Formal arrangements signal mature AI governance, clear accountability and consistent controls (risk checks, testing, documentation, security, incident handling) to ensure AI is used safely, ethically, and compliantly.	Do you have dedicated resources with the authority and responsibility to ensure the ethical, safe, secure and trustworthy development and use of AI within the company? → AICDI survey question 1.8	If a company can assign responsibility at every AI lifecycle stage to specific people or entities, it shows clear, accountable AI governance; if not, accountability is likely fragmented and oversight weaker.
QUESTION	WHAT IT TELLS YOU										
Do you have a company strategy and / or guidelines on AI? → AICDI survey question 1.4	A defined AI strategy signals the company is managing AI intentionally, with clear ownership, aligned goals, and governance (risk controls, monitoring, documentation) to scale use safely and accountably.										
Does your company adhere to any self-regulatory codes of conduct, voluntary frameworks, commitments, guidelines, or internationally recognized technical standards in relation to the ethical development and deployment of AI, recognised by reputable industry bodies or relevant regulatory authorities? → AICDI survey question 1.2	Alignment with recognised AI frameworks and standards signals governance maturity and helps assess consistency, risk controls, and regulatory readiness in ethical AI development and deployment.										
Do you have a company-wide designated board, committee or person(s), or similar bodies designated to review issues of accountability and responsibility, and other ethical issues? → AICDI survey question 1.7	Formal arrangements signal mature AI governance, clear accountability and consistent controls (risk checks, testing, documentation, security, incident handling) to ensure AI is used safely, ethically, and compliantly.										
Do you have dedicated resources with the authority and responsibility to ensure the ethical, safe, secure and trustworthy development and use of AI within the company? → AICDI survey question 1.8	If a company can assign responsibility at every AI lifecycle stage to specific people or entities, it shows clear, accountable AI governance; if not, accountability is likely fragmented and oversight weaker.										

Implementation and controls

Checks whether those rules are followed in practice through testing, approvals, monitoring, and fixes when something goes wrong.

QUESTION	WHAT IT TELLS YOU
What technical / institutional processes have been put in place to ensure the accountability, auditability and traceability of (the working of) AI systems you develop or deploy at all stages of the AI lifecycle? → AICDI survey question 1.14	Having concrete audit/traceability processes (logs, documentation, monitoring, approvals) shows the company makes AI accountability operational and verifiable, not informal or ad hoc.
Do you have a feedback mechanism for users of AI systems you develop and/or deploy to surface potential ethical issues? → AICDI survey question 1.19a	A user feedback channel signals ongoing, realworld AI governance, monitoring for harms, responding quickly, and strengthening accountability.
Do you have a process to conduct any of the following kinds of impact assessments? Ethical Impact Assessment; Fundamental Rights/Human Rights Impact Assessment; Data Protection Impact Assessment; Privacy Impact Assessment; Environmental Impact Assessment; Other. → AICDI survey question 1.15	Carrying out impact assessments helps companies can spot and reduce AI risks before rollout, like bias, privacy issues, safety failures, or legal exposure.

Workforce and social safeguards

Checks how AI affects employees and the public, and what protections exist to prevent harm and ensure human oversight.

QUESTION	WHAT IT TELLS YOU
Do you provide training or reskilling programs for employees adapting to an AI-integrated workplace? → AICDI survey question 2.1	Having clear programmes signals there's ownership, funding, and a plan to close skill gaps, support impacted roles, and reduce risks like unsafe AI use, productivity loss, and retention issues.
How does your company ensure that AI tools used in the workplace do not infringe on workers' rights? → AICDI survey question 2.5	A defined approach signals the company has rules and controls to prevent harmful uses (e.g., intrusive monitoring, biased performance decisions), protect privacy and due process, and provide transparency and escalation paths when tools affect employees.

Data quality and vendors

Checks whether the company's AI is built on reliable, well-governed data and supported by third parties it can trust. Poor data or weak vendor oversight can quietly undermine accuracy, security, compliance, and accountability.

QUESTION	WHAT IT TELLS YOU
Does the company have internal data privacy and security policies? → AICDI survey question 3.1	Whether these policies exist indicates if AI is built and used on a protected data foundation, clear rules for access, retention, sharing, and security controls to prevent leaks, misuse, and compliance failures.
Do you have a policy that governs data sharing with a third-party provider of AI systems or services? → AICDI survey question 3.4	A defined policy signals the company has clear rules, approvals, and contractual protections for vendor data use (purpose limits, security requirements, retention/deletion, and audit rights) to reduce privacy, IP, and regulatory risk.

Ethics and sustainability

Checks whether the company anticipates and manages AI's human, social, and environmental impacts. This ensures systems are fair, transparent, safe, and aligned with stakeholder expectations rather than optimising only for speed or profit.

QUESTION	WHAT IT TELLS YOU
Do you have dedicated resources with the authority and responsibility to ensure the ethical, safe, secure and trustworthy development and use of AI within the company? → AICDI survey question 1.8	Dedicated, empowered resources signal the company has clear accountability and decision rights to set standards, review high-risk uses, require mitigations, and monitor compliance as AI scales.
Which measures are in place to reduce the energy consumption and environmental impact of the AI systems you develop or deploy throughout its the AI lifecycle? → AICDI survey question 1.16	Having clear measures signals the company tracks and reduces impact in training and deployment (e.g., efficient models, optimised compute, and ongoing monitoring).

Using the checklist

This unified checklist is intentionally versatile. In pre-investment screening, investors can use it to gauge governance maturity, identify red flags and compare potential investments across sectors. In post-investment stewardship, the same questions provide a structured agenda for dialogue: investors can request evidence, set milestones and monitor progress over time. Because governance maturity varies widely by industry and company size, investors should calibrate expectations accordingly and integrate the responses into valuation, engagement and voting decisions.

Crucially, the intention of this checklist is for investors to support companies in taking part in the survey – following the logic that “sunlight is the best disinfectant” and that companies which disclose data on their AI governance are more likely to improve their practices.

Investors can incorporate AI considerations into their processes:



Integrate risk into valuation

Factor governance gaps, potential fines and reputational damage into cash-flow and cost-of-capital assumptions.



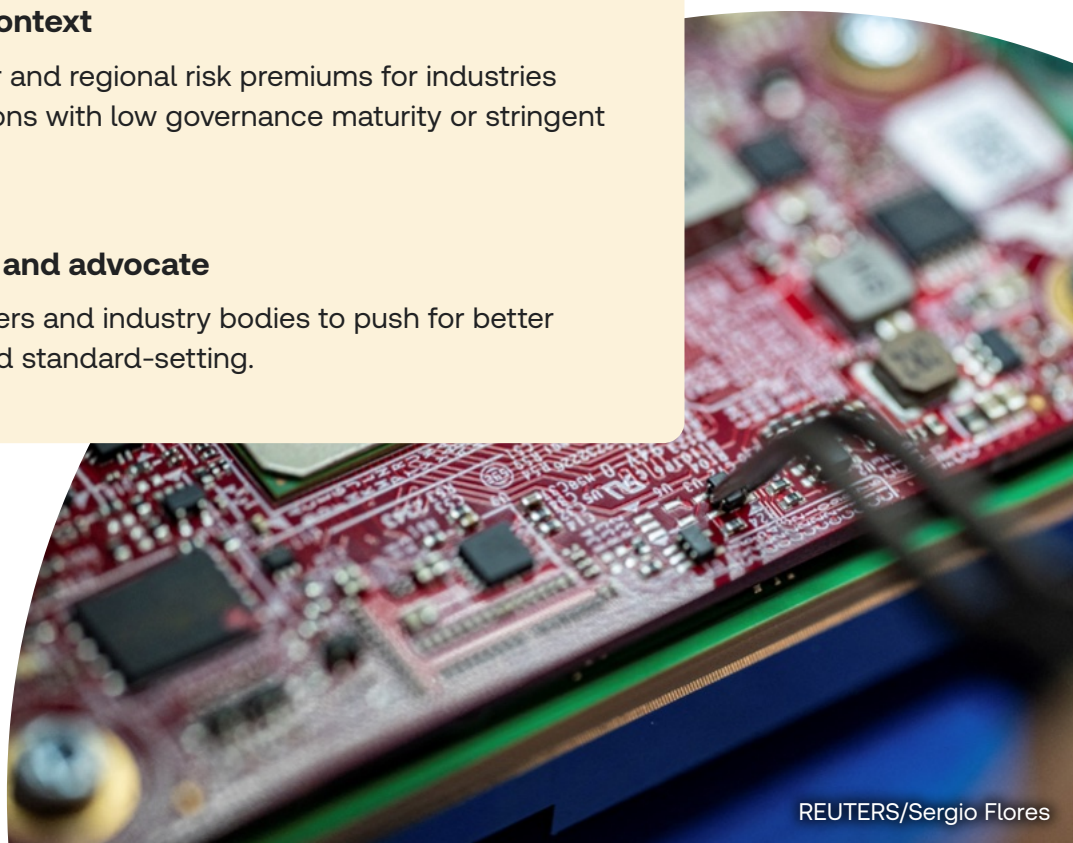
Adjust for context

Modify sector and regional risk premiums for industries and jurisdictions with low governance maturity or stringent regulations.



Collaborate and advocate

Work with peers and industry bodies to push for better disclosure and standard-setting.



Responsible AI principles for proxy voting

Shareholder resolutions on AI are attracting growing support. A recent study of fifteen AI-related proposals filed at US companies in the 2024–2025 proxy seasons found average support of about 30 per cent, nearly double the average for other environmental and social proposals.⁷³ Most of these resolutions were filed at major technology firms (Alphabet, Amazon, Apple, Meta Platforms, Microsoft and a handful of media companies) and targeted board oversight, reporting on societal risks such as misinformation and disinformation, AI-driven advertising and broader transparency.⁷³ Support varies widely across asset managers: European investors generally supported AI resolutions at much higher rates than their U.S. counterparts.⁷³

When evaluating AI-related proxy proposals, investors should observe the following neutral principles:



MEANINGFUL TRANSPARENCY

Disclosures about governance structures, risk-management policies and impact assessments give investors the information they need to evaluate whether a company is managing AI responsibly. Our findings show that only a small minority of companies maintain model registries or conduct impact assessments and most have not committed to recognised governance frameworks. This transparency gap makes it difficult to assess risk. However, proposals should be proportionate to the company's size and exposure: a one-size-fits-all reporting demand could impose undue burden on smaller issuers while providing little incremental insight.

Striking a balance between transparency and practicality ensures that reporting drives improvement rather than becoming a check-the-box exercise.



BOARD EXPERTISE

Board oversight is critical for controlling the risks and opportunities of AI.⁷² However, our findings indicate that only about half of companies have board- or committee-level oversight of AI, underscoring the need for stronger expertise. If directors lack relevant knowledge or fail to seek expert advice, they may rubber-stamp AI initiatives without understanding the underlying risks.⁷²

Supporting proposals that enhance board competence – such as establishing dedicated committees or bringing in external advisers – helps ensure that AI decisions receive informed scrutiny.



INDEPENDENT ASSURANCE

Given the technical complexity of AI, third-party audits and bias reviews provide an objective check on company claims. Yet our findings show that very few companies undertake independent ethical or environmental impact assessments and even fewer provide grievance mechanisms for AI-related issues.

Independent assessments can identify hidden flaws, biases or compliance gaps that internal teams may overlook. Supporting resolutions that call for external assurance increases accountability and builds trust with stakeholders.



SUBSTANCE OVER SYMBOLISM

Many companies already have high-level AI policies, but without measurable goals and timelines they may lack teeth. Our analysis shows that nearly nine in ten companies have not publicly committed to any AI governance framework, and only a tiny share maintains model registries.

Favouring proposals that specify concrete actions, such as establishing a model registry or conducting impact assessments, helps ensure that management commits to meaningful improvements rather than symbolic pledges.



BALANCED AMBITION

Robust AI governance entails costs; unrealistic demands can divert resources from productive innovation or penalise smaller companies. Because governance maturity varies widely across sectors and company sizes—technology, communications and finance firms tend to be more advanced, while energy, materials and real estate lag behind—investors should calibrate expectations accordingly.

Voting decisions should weigh the benefits of heightened oversight against the company's capacity and risk profile, fostering responsible AI adoption without stifling growth.

Investor Case Study: ESG-AM

Why responsible AI matters in ESG-AM's approach

ESG-AM is a Swiss-based asset management firm specialising in sustainable investments, with a concentrated service and product offering in the fixed income segment. Sustainability is integral to our identity as an asset manager and is embedded across our investment approach, with the aim of supporting resilient portfolios and long-term value creation through the thoughtful integration of sustainability-related considerations.

In 2025, ESG-AM identified artificial intelligence (AI) as an emerging sustainability consideration due to its expanding use and its potential to give rise to a range of environmental, social, and governance implications. The development and deployment of AI raises questions related to energy consumption, labour impacts, bias, and governance oversight that are increasingly relevant for responsible business practices. At the same time, AI has the potential to contribute positively to addressing sustainability challenges, including improving resource

efficiency, supporting innovation, and enabling more informed decision-making.

ESG-AM views engagement as a constructive tool to promote transparency and to understand how companies identify and manage risks and opportunities arising from the use of AI. In 2025, we began examining responsible AI as an engagement theme, focusing on sectors where AI presents significant challenges or opportunities, and identifying priority companies for engagement in collaboration with ESG-AM's portfolio managers. Given rapid technological development and uneven disclosure practices, ESG-AM's primary emphasis is on encouraging transparency, with the aim of establishing a baseline level of comparability across issuers and potential investees. This initial phase centres on awareness-raising on the importance of transparency and information gathering. Against this backdrop, ESG-AM joined AICDI in 2025 as a first step in supporting a structured, informed dialogue with our investees.



Where AICDI fits and what it enables

Voluntary disclosure initiatives such as AICDI can be particularly useful in a rapidly evolving regulatory landscape. They provide complementary sources of information that may support investors' analytical processes, enhance comparability across companies, and help identify areas where further dialogue with issuers may be constructive. While such initiatives do not in themselves ensure completeness, accuracy, or verification of issuer disclosures, our experience with sustainability disclosure frameworks, including the Workforce Disclosure Initiative, indicates that repeated participation by companies often leads to more detailed and consistent disclosures over time, contributing to improved transparency. In this context, AICDI data can be a helpful source among the inputs ESG-AM uses to inform the prioritisation of engagement efforts, supporting future dialogue focused on encouraging responsible AI practices beyond disclosure. We view AICDI not only as a data collection exercise, but also as a forum that supports shared learning, peer benchmarking, and company engagement at a stage when many companies are seeking orientation on responsible use of AI tools and practices.

An engagement example

During AICDI's pilot year, ESG-AM focused its engagement on encouraging portfolio companies to take part in the AICDI survey. Priority companies were identified through portfolio analysis of sector-level exposure to AI development and deployment risks, focusing on technology, semiconductors, telecommunications, and consumer electronics.

ESG-AM engaged with four companies across these sectors, encouraging disclosure to AICDI and positioning the survey as a tool to assess and strengthen internal AI practices. In the case of one telecommunications company, engagement included written outreach to encourage participation, followed by a call facilitated by ESG-AM that brought together representatives from the company's investor relations and AI teams, alongside a member of the AICDI team. During the discussion, ESG-AM shared its investor perspective, while the AICDI team member walked through the survey framework, highlighting strengths as well as potential gaps in the company's AI-related disclosure. The company subsequently submitted data to AICDI.

One semiconductor company further expressed interest in participating in a future cycle, citing current resource constraints. A large technology company noted that it already provides extensive AI-related disclosure through existing transparency reporting and is prioritizing upcoming regulatory requirements. A consumer electronics company did not respond despite multiple outreach



attempts. ESG-AM will continue to monitor these issuers and may re-engage in a future AICDI cycle.

As participation grows, AICDI could provide a space to support company dialogue on the responsible use of AI, including the sharing of good-practice examples and peer learning. Over time, such a platform could help identify emerging expectations, highlight areas where additional guidance may be beneficial, and support more informed dialogue between market participants. It may also provide a forum for investors, companies, researchers, and civil society to exchange insights on responsible AI challenges and blind spots and to contribute constructively to broader policy discussions.

“By improving visibility into companies’ AI practices, AICDI helps investors move discussions with companies toward the areas where additional clarity, risk mitigation or progress is most needed.”

Endnotes

- 1 Bank of England. (2024, November 21). Artificial intelligence in UK financial services. <https://www.bankofengland.co.uk/report/2024/artificial-intelligence-in-uk-financial-services-2024>
- 2 Scottish Widows. (2025, May). Governing the Algorithm: Investor Priorities for Responsible AI <https://adviser.scottishwidows.co.uk/assets/literature/docs/61448.pdf>
- 3 Quantum Black AI by McKinsey. (2025). The State of AI in 2025: Agents, innovation and transformation. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
- 4 PwC. (2025, December 8). Investors overwhelmingly look to technology sector to fuel growth—but expect greater transparency on AI strategies and policies: PwC 2025 Global Investor Survey. <https://www.pwc.com/gx/en/news-room/press-releases/2025/pwc-2025-global-investor-survey.html>
- 5 Mishra, S. (2024, April 11). AI governance appears on corporate radar. Harvard Law School Forum on Corporate Governance. <https://corpgov.law.harvard.edu/2024/04/11/ai-governance-appears-on-corporate-radar/>
- 6 Abrash, L., Probst, A., Edelman, K., & Harding, C. (2024, October 7). Governance of AI: A critical imperative for today's boards. Deloitte Insights. <https://www.deloitte.com/us/en/insights/topics/leadership/successful-ai-oversight-may-require-more-engagement-in-the-boardroom.html>
- 7 European Union. (2024). Artificial Intelligence Act. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- 8 European Union. (2016). General Data Protection Regulation. <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>
- 9 Brazil. (2018). General Data Protection Law (Lei Geral de Proteção de Dados – LGPD). <https://www.gov.br/anpd/pt-br/centrais-de-conteudo/outros-documentos-e-publicacoes-institucionais/lgpd-en-lei-no-13-709-cap.pdf>
- 10 South Africa. (2013). Protection of Personal Information Act 4 of 2013 (POPI Act). South African Government. <https://www.gov.za/documents/protection-personal-information-act>
- 11 Colorado General Assembly. (2024) Colorado Artificial Intelligence Act (SB 24-205). Consumer Protections for Artificial Intelligence. Colorado General Assembly. <https://leg.colorado.gov/bills/sb24-205>
- 12 California Legislature. (2025). California Transparency in Frontier Artificial Intelligence Act (SB 53) Advancing California's world-leading artificial intelligence industry. <https://www.gov.ca.gov/2025/09/29/governor-newsom-signs-sb-53-advancing-californias-world-leading-artificial-intelligence-industry/>
- 13 Ministry of Science and ICT. (2025). AI Basic Act. [Press Releases - 과학기술정보통신부 >](https://www.msis.go.kr/press-releases)
- 14 China. (2023). Interim Measures for the Management of Generative Artificial Intelligence Services. Cyberspace Administration of China. http://www.cac.gov.cn/2023-07/13/c_1690898327029107.htm
- 15 China. (2017). Cybersecurity Law
- 16 China. (2021). Datasecurity
- 17 China. (2021). Personal Information Protection Law
- 18 European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L (2024/1689), 1–195. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- 19 Roberts, H., Hine, E., Taddeo, M., & Floridi L. (2024, May 07). Global AI governance: barriers and pathways forward. <https://doi.org/10.1093/ia/iaae073>
- 20 AI21 Labs. (2025, August 4). 9 key AI governance frameworks in 2025. AI21. <https://www.ai21.com/knowledge/ai-governance-frameworks/>
- 21 PwC. (2024). PwC's 2024 AI Jobs Barometer. <https://www.pwc.com/gx/en/issues/artificial-intelligence/job-barometer/report.pdf>
- 22 Thomson Reuters. (2025). Future of Professionals Report 2025: Strategic AI adoption—Unlocking innovation and maximizing returns. <https://www.thomsonreuters.com/content/dam/ewp-m/documents/thomsonreuters/en/pdf/reports/future-of-professionals-report-2025.pdf>
- 23 DiMauro, J. (2025, September 5). Lack of defined AI strategy plagues businesses and compliance, studies show. Global Relay Intelligence & Practice. <https://www.grip.globalrelay.com/lack-of-defined-ai-strategy-plagues-businesses-and-compliance-studies-show/>
- 24 Cancela Outeda, C. (2024). The EU's AI Act: A framework for collaborative governance. Internet of Things, 27, Article 101291. <https://doi.org/10.1016/j.iot.2024.101291>
- 25 Gov PD. (2025). Businesses using AI in the EU must comply with the AI Act. <https://gov-pd.co.uk/businesses-using-ai-in-the-eu-must-comply-with-the-ai-act/>
- 26 Roberts, H., Ziosi, M., & Cailean, O. (2023). A Comparative Framework for AI Regulatory Policy. CEIMIA. <https://ceimia.org/wp-content/uploads/2023/02/Comparative-Frameworkfor-AI-Regulatory-Policy.pdf>
- 27 Jackson & Freeman. (November 28, 2025). AI Maturity Matters. <https://ssrn.com/abstract=5824162>
- 28 Thomson Reuters. (2025, June 26). The AI adoption reality check: Firms with AI strategies are twice as likely to see AI-driven revenue growth; those without risk falling behind. <https://www.thomsonreuters.com/en/press-releases/2025/june/the-ai-adoption-reality-check-firms-with-ai-strategies-are-twice-as-likely-to-see-ai-driven-revenue-growth-those-without-risk-falling-behind>

- 29 Edelman. (2024). 2024 Edelman Trust Barometer: Key insights around AI. <https://www.edelman.com/sites/g/files/aatuss191/files/2024-03/2024%20Edelman%20Trust%20Barometer%20Key%20Insights%20Around%20AI.pdf>
- 30 For clarity, an AI Model Registry is a centralised repository that manages and tracks machinelearning models across their lifecycle, functioning as version control by storing and versioning models from development to deployment and eventual retirement. AI Lifecycle refers to the processes and operational mechanisms that support an AI system across the artificial intelligence system lifecycle, from research, design and development through deployment and use, including maintenance, operation, monitoring and evaluation/validation, and endofuse and termination; these phases are often iterative rather than strictly sequential.
- 31 Corrêa, A. M., Garsia, S., & Elbi, A. (2025). Better together? Human oversight as means to achieve fairness in the European AI Act governance. *Cambridge Forum on AI: Law and Governance*, 1, e29. <https://doi.org/10.1017/cfl.2025.10010>
- 32 Innopharma Education. (2026). The impact of AI on job roles, workforce, and employment: What you need to know. <https://www.innopharmaeducation.com/blog/the-impact-of-ai-on-job-roles-workforce-and-employment-what-you-need-to-know>
- 33 O'Brien, M. (2026, February 1). Did artificial intelligence really drive layoffs at Amazon and other firms? ABC News. <https://abcnews.go.com/Business/wireStory/artificial-intelligence-drive-layoffs-amazon-firms-hard-129773731>
- 34 Cazzaniga, M., Jaumotte, F., Li, L., Melina, G., Panton, A. J., Pizzinelli, C., Rockall, E., & Tavares, M. M. (2024, January). Gen-AI: Artificial intelligence and the future of work (IMF Staff Discussion Note No. SDN/2024/001). International Monetary Fund. <https://www.imf.org/-/media/files/publications/sdn/2024/english/sdnea2024001.pdf>
- 35 International Labour Organization. (2025, January 28). How reskilling for AI could unlock new and better jobs. <https://www.ilo.org/resource/article/how-reskilling-ai-could-unlock-new-and-better-jobs>
- 36 European Trade Union Confederation. (2022, December 6). ETUC resolution calling for an EU directive on algorithmic systems at work. <https://www.etuc.org/en/document/etuc-resolution-calling-eu-directive-algorithmic-systems-work>
- 37 Maunier, Gérald & Vandewalle, Jean-Jacques & George, Patrick. (2023). Environmental Impacts in AI Governance Integrating sustainability and environmental assessments into AI data governance. https://www.researchgate.net/publication/394514916_Environmental_Impacts_in_AI_Governance_Integrating_sustainability_and_environmental_assessments_into_AI_data_governance
- 38 Ghosh, M. (2025). Artificial intelligence (AI) and ethical concerns: A review and research agenda. *Cogent Business & Management*, 12(1), 2551809. <https://doi.org/10.1080/23311975.2025.2551809>
- 39 Maroun, W. (2022). Corporate governance and the use of external assurance for integrated reports (Accepted version). *Corporate Governance: An International Review*, 30(5), 584–607. <https://eprints.whiterose.ac.uk/id/eprint/208770/1/CORPOR-1.PDF>
- 40 London School of Economics and Political Science. (2025, September 12). What direct risks does AI pose to the climate and environment? <https://www.lse.ac.uk/granthaminstitute/explainers/what-direct-risks-does-ai-pose-to-the-climate-and-environment/>
- 41 Barocas, S., & Selbst, A. D. (2015). Big data's disparate impact. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2477899>
- 42 AICDI's definition of "European" countries in this instance and throughout the report refers to European countries excluding the UK and Russia
- 43 European Union. (2026). Article 10: Data and data governance. AI Act Service Desk. <https://ai-act-service-desk.ec.europa.eu/en/ai-act/article-10>
- 44 KPMG. (2024). AI regulation: Colorado Artificial Intelligence Act (CAIA) – Regulatory alert. <https://kpmg.com/kpmg-us/content/dam/kpmg/pdf/2024/ai-regulation-colorado-artificial-intelligence-act-caia-reg-alert.pdf>
- 45 Francis, J., & Gluck, J. (2025, October). CCPA regulations on automated decisionmaking technology, risk assessments, and cybersecurity audits: Issue brief. *Future of Privacy Forum*. https://fpf.org/wp-content/uploads/2025/10/FPF_CCPA-Regulations-Issue-Brief.pdf
- 46 Wallach, W., Reuel, A., & Kaspersen, A. (2025). Soft law in international AI governance. *Carnegie Council for Ethics in International Affairs*. <https://media-1.carnegiecouncil.org/cceia/Soft-Law-in-International-AI-Governance.pdf>
- 47 National Institute of Standards and Technology. (2025, February 6). NIST AI RMF Playbook. <https://www.nist.gov/itl/ai-risk-management-framework/nist-ai-rmf-playbook>
- 48 World Economic Forum. (2023, June). Adopting AI responsibly: Guidelines for procurement of AI solutions by the private sector. https://www3.weforum.org/docs/WEF_Adopting_AI_Responsibly_Guidelines_for_Procurement_of_AI_Solutions_by_the_Private_Sector_2023.pdf
- 49 Ada Lovelace Institute. (2023). Allocating accountability in AI supply chains. <https://www.adalovelaceinstitute.org/resource/ai-supply-chains/>
- 50 National Institute of Standards and Technology. (2023). AI Risk Management Framework (AI RMF 1.0). <https://www.nist.gov/itl/ai-risk-management-framework>

- 51 International Organization for Standardization. (2023). ISO 42001 explained: What it is and how it works. <https://www.iso.org/home/insights-news/resources/iso-42001-explained-what-it-is.html>
- 52 EY. (2024, July 12). The European Union Artificial Intelligence Act: Latest developments and key takeaways. <https://www.ey.com/content/dam/ey-unified-site/ey-com/en-gl/insights/public-policy/documents/ey-gl-eu-ai-act-07-2024.pdf>
- 53 Hopkins, A., Cen, S. H., Ilyas, A., Struckman, I., Videgaray, L., & Mądry, A. (2025). AI supply chains: An emerging ecosystem of AI actors, products, and services (arXiv:2504.20185). <https://arxiv.org/abs/2504.20185>
- 54 UNESCO. (2022). Recommendation on the ethics of artificial intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000380455>
- 55 Organisation for Economic Co-operation and Development. (2024). OECD AI principles. <https://oecd.ai/en/ai-principles>
- 56 Office of the High Commissioner for Human Rights. (2011). Guiding principles on business and human rights: Implementing the United Nations “Protect, Respect and Remedy” framework. https://www.ohchr.org/sites/default/files/Documents/Publications/GuidingPrinciplesBusinessHR_EN.pdf
- 57 European Union. (2016). General Data Protection Regulation: Article 22 – Automated individual decision-making, including profiling. <https://gdpr-info.eu/art-22-gdpr/>
- 58 Barocas, S., & Selbst, A. D. (2015). Big data’s disparate impact. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.2477899>
- 59 Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J. W., Wallach, H., Daumé III, H., & Crawford, K. (2021). Datasheets for datasets. Communications of the ACM, 64(12), 86–92. <https://cacm.acm.org/research/datasheets-for-datasets/>
- 60 Sambasivan, N., Kapania, S., Highfill, H., Akrong, D., Paritosh, P., & Aroyo, L. (2021). “Everyone wants to do the model work, not the data work”: Data cascades in high-stakes AI. In CHI Conference on Human Factors in Computing Systems (CHI ’21). <https://doi.org/10.1145/3411764.3445518>
- 61 The Artificial Intelligence Act. (2026). Article 10: Data and data governance. <https://www.euaiact.com/article/10>
- 62 Credera. (n.d.). Finding the right sentiment analysis model for you: VADER vs. Spark NLP. <https://www.credera.com/en-us/insights/finding-the-right-sentiment-analysis-model-for-you-vader-vs-spark-nlp>
- 63 JUST Capital. (n.d.). Research reveals key insights on responsible corporate AI deployment. <https://justcapital.com/news/research-reveals-key-insights-on-responsible-corporate-ai-deployment/>
- 64 From Adoption to Integration: Boards Challenged to Effectively Implement AI, Finds 2025 BDO Survey <https://www.businesswire.com/news/home/20251202692063/en/From-Adoption-to-Integration-Boards-Challenged-to-Effectively-Implement-AI-Finds-2025-BDO-Survey>
- 65 Corporate Compliance Insights. (2025, August 1). Boards increasingly tout AI expertise. <https://www.corporatecomplianceinsights.com/news-roundup-august-1-2025/>
- 66 CFA Institute. (2025). AI washing: Signs, symptoms, & suggested solutions. <https://rpc.cfainstitute.org/research/reports/2025/ai-washing>
- 67 TELUS Corporation. (2025). 2024 annual report: Leading with purpose, innovating with passion. <https://assets.ctfassets.net/>
- 68 Vodafone Group Plc. (n.d.). Artificial intelligence framework. <https://www.vodafone.com/~media/Files/V/vodafone/corp/documents/investors/vodafone-artificial-intelligence-framework.pdf>
- 69 SAP SE. (2025). SAP integrated report 2024. <https://www.sap.com/integrated-reports/2024/en.html>
- 70 Telefónica, S.A. (2025). Annual accounts 2024: Consolidated information. <https://www.telefonica.com/en/shareholders-investors/financial-reports/annual-report/>
- 71 Ontario Securities Commission. (n.d.). How are issuers discussing AI in their financial disclosures? <https://www.osc.ca/en/industry/artificial-intelligence/how-are-issuers-discussing-ai-their-financial-disclosures#:~:text=Sentiment%20Analysis%20Results>
- 72 Neervoort, S., & Rang, W. (2024, March 21). Artificial intelligence: An engagement guide. Harvard Law School Forum on Corporate Governance. <https://corpgov.law.harvard.edu/2025/05/30/investor-views-on-ai-oversight-what-do-proxy-votes-tell-us/#:~:text=,higher%20average%20support%20than%20107>
- 73 Stewart, L., & Meng, R. (2025, May 30). Investor views on AI oversight: What do proxy votes tell us? Harvard Law School Forum on Corporate Governance. <https://corpgov.law.harvard.edu/2025/05/30/investor-views-on-ai-oversight-what-do-proxy-votes-tell-us/#:~:text=,higher%20average%20support%20than%20107>

Responsible AI in practice

2025 global insights from the AI Company Data Initiative

Artificial intelligence (AI) is rapidly being embedded across companies' products, services and internal operations, yet governance and disclosure are not evolving at the same speed. This report looks at corporate practice in the context of the emerging responsible AI regulatory landscape and analyses publicly available data collected by the Thomson Reuters Foundation's AI Company Data Initiative, the largest global dataset of corporate responsible AI disclosures. As privately developed or deployed AI systems shape more of daily life, transparency must move beyond technical descriptions to show how accountability works— including who makes decisions, how ethical issues are escalated, and what remediation paths exist when things go wrong. Clear responsibility for harms or breaches should be identifiable in practice, not just in principle. Just as we expect openness and accountability from government, it is important that the private sector meets comparable transparency standards for AI that affects the public.

Disclaimer

This publication, the information therein and related materials are not intended to provide and do not constitute financial or investment advice. Thomson Reuters did not assess companies according to financial performance or metrics. Thomson Reuters Foundation makes no representation regarding the advisability or suitability of investing in any particular company, investment fund, pension or other vehicle, or of using the services of any particular bank, asset manager, company, pension provider or other service provider for the provision of investment services. A decision to use the services of any bank, or other entity, or to invest or otherwise should not be made in reliance on any of the statements set forth in this publication. While every effort has been made to ensure the information in this publication is correct, Thomson Reuters Foundation and its agents cannot guarantee its accuracy and they shall not be liable for any claims or losses of any nature in connection with information contained in this document, including, but not limited to, lost profits or punitive or consequential damages or claims in negligence.

