

In collaboration  
with Accenture



# Advancing Responsible AI Innovation: A Playbook

INSIGHT REPORT  
SEPTEMBER 2025



# Contents

|                                                                                       |           |
|---------------------------------------------------------------------------------------|-----------|
| Foreword                                                                              | 3         |
| Executive summary                                                                     | 4         |
| Introduction                                                                          | 5         |
| <b>Dimension 1: Strategy and value creation</b>                                       | <b>8</b>  |
| Play 1: Lead with a long-term, responsible AI strategy and vision for value creation  | 9         |
| Play 2: Unlock AI innovation with trustworthy data governance                         | 12        |
| Play 3: Design resilient responsible AI processes for business continuity             | 15        |
| <b>Dimension 2: Governance and accountability</b>                                     | <b>18</b> |
| Play 4: Appoint and incentivize AI governance leaders                                 | 19        |
| Play 5: Adopt a systematic, systemic and context-specific approach to risk management | 21        |
| Play 6: Provide transparency into responsible AI practices and incident response      | 24        |
| <b>Dimension 3: Development and use</b>                                               | <b>27</b> |
| Play 7: Drive AI innovations with responsible design as the default                   | 28        |
| Play 8: Scale responsible AI with technology enablement                               | 31        |
| Play 9: Increase responsible AI literacy and workforce transition opportunities       | 33        |
| Conclusion                                                                            | 36        |
| Contributors                                                                          | 37        |
| Endnotes                                                                              | 43        |

## Disclaimer

This document is published by the World Economic Forum as a contribution to a project, insight area or interaction. The findings, interpretations and conclusions expressed herein are a result of a collaborative process facilitated and endorsed by the World Economic Forum but whose results do not necessarily represent the views of the World Economic Forum, nor the entirety of its Members, Partners or other stakeholders.

© 2025 World Economic Forum. All rights reserved. No part of this publication may be reproduced or transmitted in any form or by any means, including photocopying and recording, or by any information storage and retrieval system.

# Foreword



**Arnab Chakraborty**  
Chief Responsible AI Officer,  
Accenture



**Cathy Li**  
Head, Centre for AI  
Excellence; Member of  
the Executive Committee,  
World Economic Forum

Nearly three years into the world's "generative artificial intelligence (AI) moment", AI technologies are rapidly being woven into the fabric of our daily lives, societies and economies. Advancing AI for innovation, human rights and societal benefit demands intentional design, ongoing oversight and active public-private collaboration across stakeholders.

Industry plays a frontline role in AI governance through a series of principles and practices referred to as "responsible AI". However, the maturity in implementing these practices lags behind awareness of its importance, resulting in a responsible AI implementation gap. Business leaders face several roadblocks in addressing this gap that can arise from within their organizations, such as evaluating AI use and risk at scale, as well as those that emerge from navigating the broader jurisdictional environment, such as fragmented regulatory approaches. If unaddressed, this lapse in AI governance is likely to erode confident AI investment, compliance and public trust.

Over the past two years, the Resilient Governance and Regulation working group of the World Economic Forum's AI Governance Alliance, with support from Accenture as its knowledge partner,

has mapped key challenges to the responsible AI implementation gap and co-developed practical mitigations. The alliance itself has become the Forum's fastest-growing community, now comprising over 650 global members from industry leaders, governments, academia and civil society collaborating to drive responsible AI innovation.

This playbook is the product of a sustained and deliberate multistakeholder effort to address a gap in AI governance and enable organizations to advance their responsible AI efforts, building upon our previous report, [Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation](#). By drawing on real-world case studies, comparative policy analyses, and insights from diverse communities and disciplines, this work offers practical actions for business and government leaders to build resilient, scalable and adaptive governance systems for AI innovation.

While this playbook advances industry's practice of responsible AI, building a trustworthy AI ecosystem requires collaboration across all stakeholders. We invite decision-makers from civil society, academia and government to join us in shaping a more resilient global AI governance landscape.

# Executive summary

## A playbook for organization and government leaders to advance responsible AI innovation.

Responsible artificial intelligence (AI) – the practice of developing and managing AI systems that maximize benefits and minimize the risks they pose to people, society and the environment<sup>1</sup> – is fundamental to sustainable innovation. Many organizations report benefits from implementing responsible AI, including improved efficiency and enhanced customer trust.<sup>2</sup> Despite this, research has found that less than 1% of organizations have fully operationalized responsible AI in a comprehensive and anticipatory manner.<sup>3</sup> This gap in responsible AI implementation slows progress, undermines trust in AI technologies and limits their transformative potential.

Successful implementation of responsible AI practices requires actions by business leaders as well as coordination with policy-makers through regulatory clarity, aligned incentives and cross-sector collaboration.

This playbook provides **nine essential plays** across **three key dimensions** of responsible AI. Each play offers complementary actions for the following:



**Organizational leaders** to overcome internal roadblocks



**Government leaders** to address ecosystem challenges

Both sets of actions are critical and maximum impact comes from pursuing them in parallel.

Rather than prescriptive linear steps, the playbook offers a flexible framework that organizations and jurisdictions can adapt to their specific context and maturity level. The result is a practical roadmap for turning responsible AI from aspiration into a competitive advantage – driving innovation while building public trust to enable AI strategies to reach their full potential.

### Dimension 1: Strategy and value creation

**Play 1:** Lead with a long-term, responsible AI strategy and vision for value creation ↗

**Play 2:** Unlock AI innovation with trustworthy data governance ↗

**Play 3:** Design resilient, responsible AI processes for business continuity ↗

### Dimension 2: Governance and accountability

**Play 4:** Appoint and incentivize AI governance leaders ↗

**Play 5:** Adopt a systematic, systemic and context-specific approach to risk management ↗

**Play 6:** Provide transparency into responsible AI practices and incident responses ↗

### Dimension 3: Development and use

**Play 7:** Drive AI innovations with responsible design as the default ↗

**Play 8:** Scale responsible AI with technology enablement ↗

**Play 9:** Increase responsible AI literacy and workforce transition opportunities ↗

# Introduction

Responsible AI practices are critical to unlocking sustainable innovation, yet significant implementation gaps remain.

Artificial intelligence (AI) is transforming the global economy, enabling innovation, growth and social advancement amid heightened technological and geopolitical complexity. Sustainable adoption of AI necessitates an ecosystem of intentionally

designed principles, guidelines and practices – collectively referred to as “responsible AI” – to effectively govern the technology for desirable outcomes, as outlined in Table 1.

TABLE 1 Responsible AI implementation by businesses matters to both organizations and governments

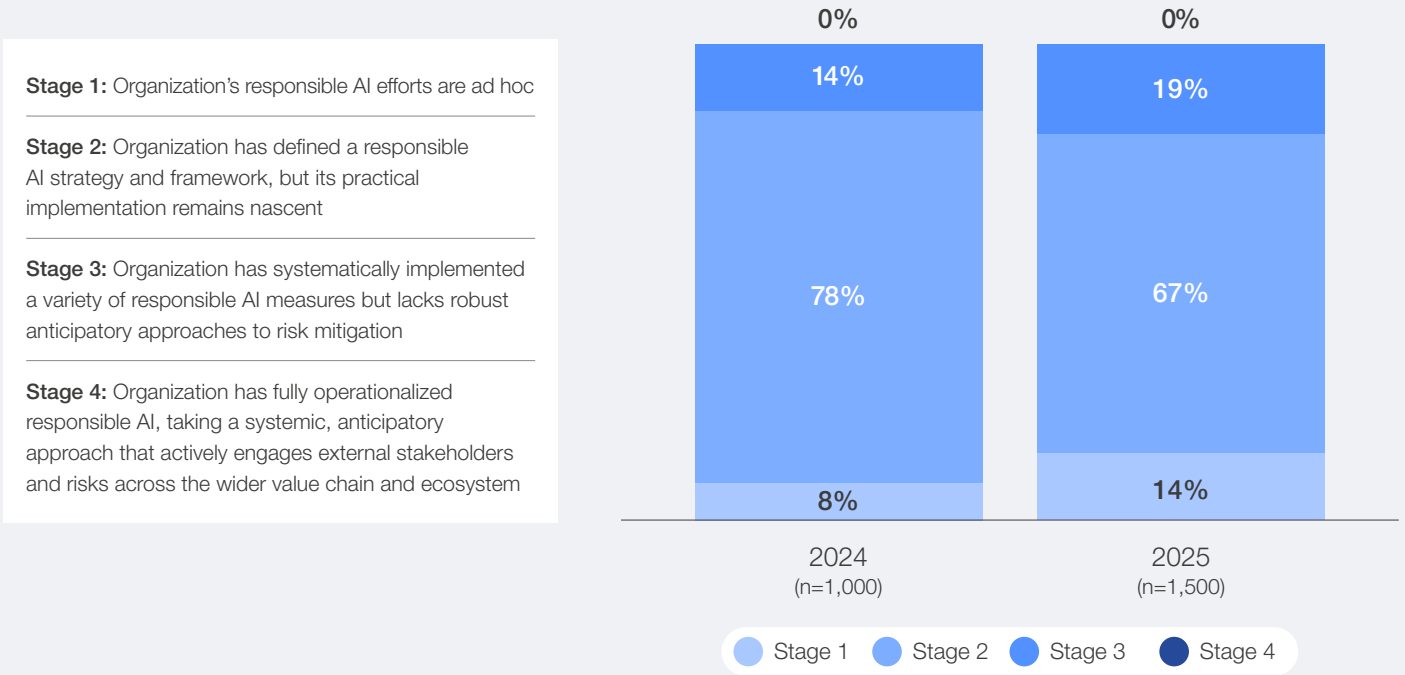
|  For organizations                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |  For governments                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <ul style="list-style-type: none"><li>– <b>Enable business value and competitiveness</b> through reliable and confident AI innovation</li><li>– <b>Increase adoption</b> through enhanced customer/consumer confidence in AI-driven products and services</li><li>– <b>Enhance employee inclusion and retention</b> through trusted AI adoption</li><li>– <b>Ensure robust risk management</b> to mitigate legal, financial and reputational exposure</li><li>– <b>Meet stakeholder expectations</b> of innovation that protects privacy and civil liberties</li><li>– <b>Proactively respond to technological developments</b> such as agentic AI, which depend on comprehensive governance and established trust</li><li>– <b>Reduce regulatory burden</b> with proactive responsible AI implementation</li><li>– <b>Reinforce ecosystem partnerships</b> by aligning standards, facilitating interoperability and enabling long-term innovation</li></ul> | <ul style="list-style-type: none"><li>– <b>Promote economic growth</b> through sustainable AI innovation</li><li>– <b>Safeguard human rights and freedoms</b> that may be challenged with AI adoption</li><li>– <b>Increase public trust in digital transformation</b> by ensuring transparency and accountability of organizations implementing AI</li><li>– <b>Enable inclusive and adaptive policy frameworks</b> built on a foundation of organizational responsible AI maturity</li><li>– <b>Be globally competitive</b> and secure national objectives (e.g. sovereignty, cybersecurity, energy)</li><li>– <b>Promote innovation for all</b> by ensuring equitable access to AI benefits across communities</li><li>– <b>Enhance societal and environmental benefits</b>, especially within critical sectors</li><li>– <b>Improve delivery of government services</b> with streamlined operations and increased decision-making efficiencies</li></ul> |

## The responsible AI implementation gap

Despite increased awareness, responsible AI practices by organizations remain immature. Measured on a four-stage maturity scale, a 2025 survey of 1,500 companies found that 81% remain in the first two early stages of responsible AI. While the number of companies with a stage 3 maturity

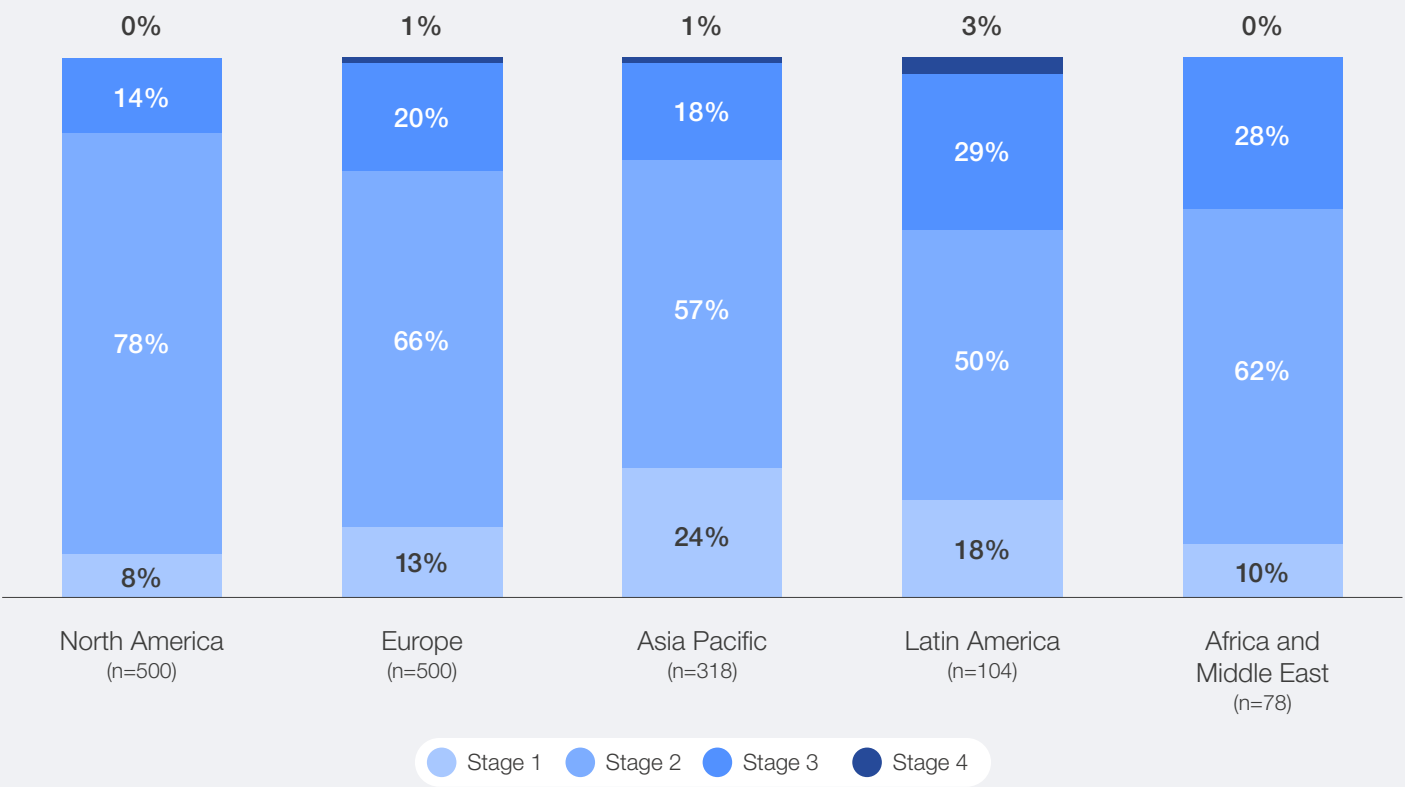
increased from 14% in 2024 to 19% in 2025, less than 1% of companies are at stage 4 (Figure 1).<sup>4</sup> This limited maturity is prevalent across organizations, regions and sectors (Figures 2 and 3). Though the “why” of responsible AI is largely understood, the “how” remains elusive to most organizations.

FIGURE 1 | Global responsible AI implementation in 2024 and 2025



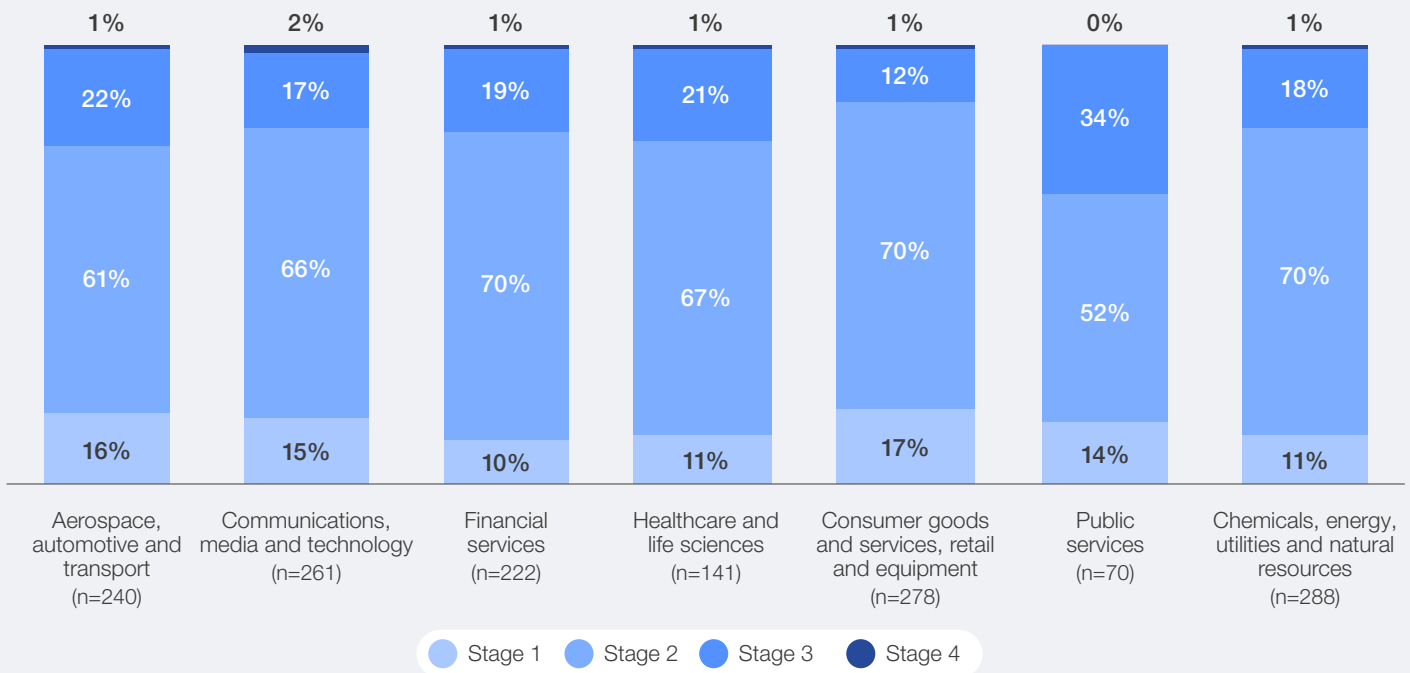
**Source:** Research by Accenture and Stanford Institute for Human-Centered AI; Maslej, N., L. Fattorini, R. Perrault, Y. Gil, et al. (2025). *The AI Index 2025 Annual Report*. Stanford Institute for Human-Centered AI; The responsible AI maturity index has been interpreted according to methods used in: Accenture. (2024). *Responsible AI: From compliance to confidence*. A detailed definition of maturity stages can be found in this report.

FIGURE 2 | Responsible AI maturity by region



**Source:** Research by Accenture and Stanford Institute for Human-Centered AI; Maslej, N., L. Fattorini, R. Perrault, Y. Gil, et al. (2025). *The AI Index 2025 Annual Report*. Stanford Institute for Human-Centered AI; The responsible AI maturity index has been interpreted according to methods used in: Accenture. (2024). *Responsible AI: From compliance to confidence*. A detailed definition of maturity stages can be found in this report.

FIGURE 3 | Responsible AI maturity by sector



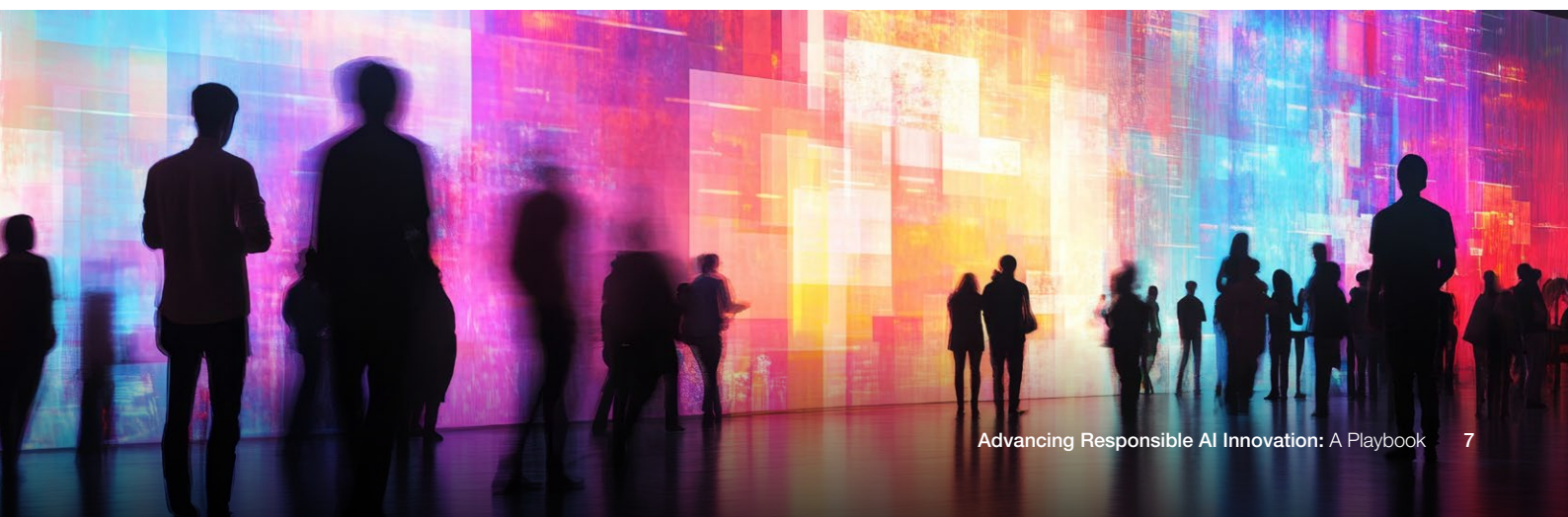
**Source:** Research by Accenture and Stanford Institute for Human-Centered AI; Maslej, N., L. Fattorini, R. Perrault, Y. Gil, et al. (2025). *The AI Index 2025 Annual Report*. Stanford Institute for Human-Centered AI; The responsible AI maturity index has been interpreted according to methods used in: Accenture. (2024). *Responsible AI: From compliance to confidence*. A detailed definition of maturity stages can be found in this report.

## A playbook to advance responsible AI innovation, informed by AI experts

This playbook is designed to help organizations bridge the gap between responsible AI principles and real-world implementation. It emphasizes the importance of a strong, collaborative ecosystem, one that fosters public-private partnerships (PPPs) and international cooperation. Structured around nine plays, across three dimensions, the playbook offers practical guidance for both organizational and government leaders to tackle the internal and related external challenges of responsible AI implementation. While each play outlines coordinated steps and complementary actions for organizations and governments, the recommendations should be considered in the broader context of critical efforts led by other stakeholders, such as academia and civil society.

The playbook is informed by insights gathered from extensive research, working group feedback, expert interviews and input from two in-depth workshops on targeted responsible AI topics:

- Designing Effective Codes of Conduct (Global AI Summit on Africa, Kigali, May 2025), co-hosted with the Government of Japan, with the objective of identifying effective strategies for public-private collaboration in designing and implementing context-specific AI codes of conduct.
- Unlocking AI Innovation through Responsible Data Sharing (AI Governance Alliance Community Meeting, San Francisco, June 2025), with the objective of developing practical approaches to address data sharing barriers.



# Dimension 1:

## Strategy and value creation

Align corporate strategy with responsible AI innovation to create long-term stakeholder value.

This chapter outlines how organizations can embed responsible AI into their strategic core, and the policy tools governments can use to incentivize such practices within organizations.

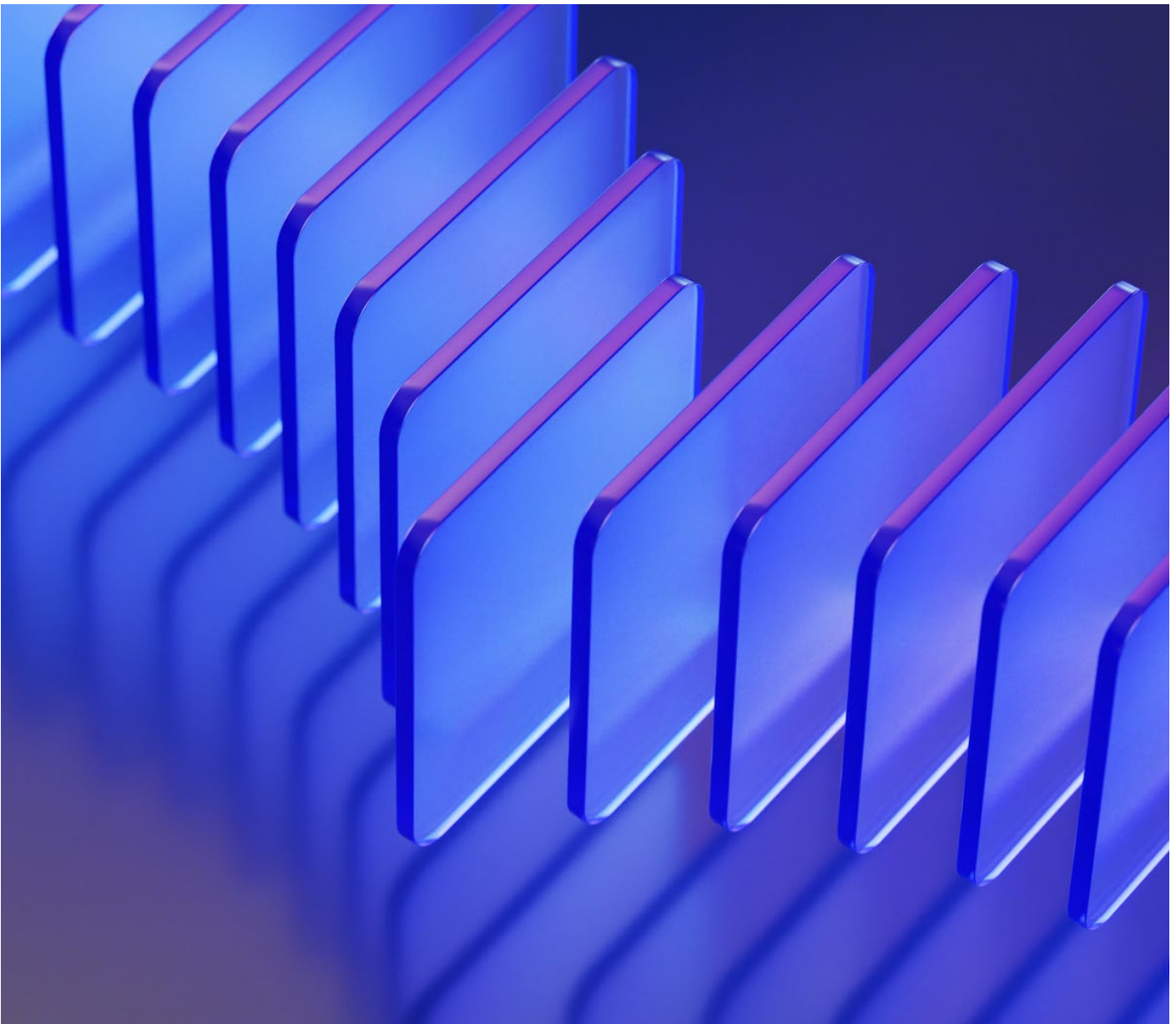
**Play 1:** Lead with a long-term, responsible AI strategy and vision for value creation



**Play 2:** Unlock AI innovation with trustworthy data governance



**Play 3:** Design resilient, responsible AI processes for business continuity

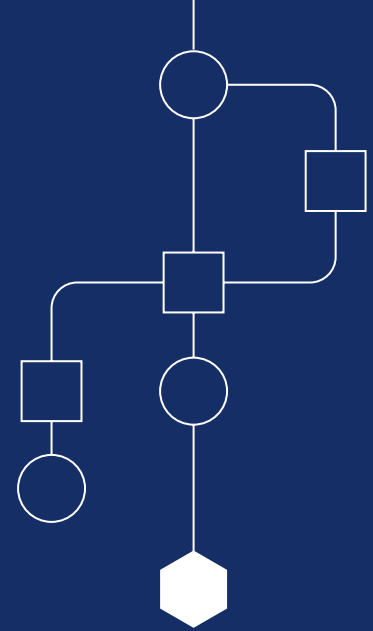




# Play 1

## Lead with a long-term, responsible AI strategy and vision for value creation

To secure both immediate AI opportunities and address evolving risk environments, companies must integrate a responsible AI strategy into their business strategy and AI innovation roadmap. For governments, organizational responsible AI maturity is more than ensuring trust and confidence; it can serve as a foundation for an adaptive AI policy life cycle necessary for new, dynamic AI capabilities like multimodal, robotic, agentic and beyond.



## Organization leaders

### Key roadblocks that arise within the organization

Slow AI adoption, undermining responsible implementation priorities

Return on investment (ROI) pressure, sacrificing ethical safeguards for immediate returns

Insufficient investment in responsible AI talent and tools, preventing organizations from operationalizing principles into scalable practices

Legacy security and IT frameworks and standards that are not adapted to AI risk management

### Actions for organization leaders

- **Embrace the strategic imperative underpinning responsible AI commitments:** Such practices drive significant value (see Table 1) and can yield strong improvements in product quality and contract win rates.<sup>5</sup> To maximize benefits, C-suite and board sponsorship is fundamental to aligning AI governance with the organization's broader strategy, requiring:
  - **Executive education** on the capabilities of AI and the value of responsible AI
  - **One-on-one engagement** with each C-suite member to discuss responsible AI value to their function, emerging compliance requirements and cross-functional alignment
  - **Dedicated AI leadership to own strategy,** buy-in and adoption (see Play 4)
- **Set and socialize responsible AI vision and principles:** These must align with the organizational mission and values and be reinforced by policies, standards, and guidelines supporting adherence and accountability (see Case study 1). Maximizing responsible AI's benefits requires shifting from an abstract bolted-on approach to a methodologically integrated, tested and refined science that ensures systematic and context-specific risk management (see Play 5). For example, Mastercard embeds accountability tools and technical controls into its AI governance programme to systematically evaluate, guide and verify all AI system use across the enterprise. Additionally, leaders must promote a culture of mutual trust where employees view responsible AI as a foundation rather than as an obstacle.
- **Establish dialogue for continuous employee input.** For example, Microsoft and the American Federation of Labor and Congress of Industrial Organizations (AFL-CIO), the largest US labour federation, created the first-of-its-kind AI partnership. The partnership's priorities include direct feedback mechanisms for labour leaders and workers.<sup>6</sup>
- **Be transparent** in the purpose and limits of AI in the organization and how work will be impacted.<sup>7</sup>
- **Tailor training** to enhance the use of AI responsibly (see Play 9); build trust through upskilling and redeploying employees with AI-displaced roles.
- **Align rewards** to responsible performance.

## CASE STUDY 1

### Telefónica's multi-pronged responsible AI strategy

Large organizations often struggle to implement enterprise-wide responsible AI programmes. Rather than starting from scratch, Telefónica, a global telecommunications company, began by building upon its existing risk-oriented privacy governance model and participating in EU-wide efforts to identify AI requirements. After defining a set of AI principles, Telefónica piloted an AI Office, Ethics Expert Group and Responsible AI Champions to steward adoption across teams. The company also tested product and service evaluation methodologies and initiated tailored training and awareness-raising efforts.<sup>8</sup> For cohesive and accountable

adoption, the company established a unified AI governance framework that integrated compliance requirements alongside ethics to evaluate systems for societal impact, human agency and inclusivity.



#### Key insight

Scaling responsible AI requires formalizing principles into a governance model that empowers teams, proactively manages risk, engages cross-functional expertise and supports compliance across diverse regulatory environments.



## Government leaders

### Key roadblocks organizations encounter from the broader ecosystem

Sudden regulatory changes and limited guidance, forcing companies to frequently modify their visions, which erodes the consistency of their strategies

Increased geopolitical competition and insufficient international AI governance cooperation, making companies choose between geopolitical demands for rapid deployment and responsible AI practices

### Actions for government leaders

- **Communicate jurisdictional responsible AI goals:** National approaches and expectations regarding governance and regulation must be clearly articulated to encourage organizations to follow suit in communicating responsible AI practices to their employees. Jurisdictions use one or more approaches to communicate goals. Examples include:

- **National AI strategy:** Brazil's AI Plan (PBIA),<sup>9</sup> the US AI Action Plan,<sup>10</sup> China's AI Action Plan<sup>11</sup> and Costa Rica's National AI Strategy<sup>12</sup>
- **Codes of conduct:** Canada's Voluntary Code of Conduct on the Responsible Development and Management of Advanced Generative AI Systems<sup>13</sup> and the Hiroshima AI Process Code of Conduct (see Case study 7)
- **Guidelines:** Egypt's Charter for Responsible AI<sup>14</sup> and Australia's Voluntary AI Safety Standard<sup>15</sup>
- **Regulation:** The European Union's (EU) AI Act, South Korea's AI Framework Act, and Japan's AI Promotion Act showcase three divergent approaches to comprehensive AI regulation<sup>16</sup>

Enabling confident adoption by industry requires balancing the frequency of goal revisions by governments with predictability. Governments' own implementation of robust responsible AI practices can help set adoption expectations for industry.

- **Incentivize industry implementation of responsible AI:** Jurisdictional incentives can help ensure market goals maintain alignment with public interest<sup>17</sup> and that an organization's responsible AI goals also address macro-level challenges like workforce, environmental and information ecosystem impacts.

Jurisdictions are exploring varied incentive approaches, such as:

- **Preferred procurement:** such as for AI developers ensuring appropriate guardrails in their models<sup>18</sup>
- **Financial penalties or rewards:** including tax incentives, grants or subsidies<sup>19</sup>
- **Standardized frameworks:** incorporating expert-informed risk management approaches and reporting templates for responsible AI practices (see Play 6)
- **Publicity:** recognizing companies that meet the jurisdiction's communicated goals (see Case study 2)

## CASE STUDY 2

### Dubai AI Seal to prevent deceptive practices

When organizations exaggerate the AI capabilities of their products and services (i.e. “AI washing”), it can negatively impact the adoption of, and trust in, the offerings of reputable organizations. The Dubai Centre for Artificial Intelligence (DCAI) launched the “Dubai AI Seal”, which aims to verify AI offering claims for businesses operating in Dubai. Companies are assessed on their activities and services, and their number of AI-specialized employees, projects and partnerships.<sup>20</sup>



#### Key insight

Governments are exploring novel methods (such as seals) as alternatives or complements to traditional incentives (such as penalties) to encourage companies to implement jurisdictional AI governance goals. For jurisdictions that assess both voluntary guidelines and regulatory penalties to be currently unenforceable, seals may offer a middle path: abiding by such guidance is still voluntary but can carry tangible consequences from consumer buying power.





# Play 2

## Unlock AI innovation with trustworthy data governance

Successful AI innovation depends on secure, high-quality and compliant data access with controls on processing, consent, cross-border transfers and AI deployment. Therefore, a modern data foundation must embed security into data workflows and AI systems as well as upgrade traditional security models.



## Organization leaders

### Key roadblocks that arise within the organization

Low data quality and legacy processes, undermining the reliability of company systems

Isolated data ecosystems and fragmented governance, limiting cohesive data strategies and enterprise-wide insight generation

Complex approval processes, hindering internal and external data sharing

Data scarcity, especially in categories that are prone to underrepresentation, impacting model training and risk mitigation strategies

### Actions for organization leaders

- **Implement an enterprise-wide data governance strategy:** Establish guardrails for integrity and compliance that ensure data quality, interoperability and traceability across business units. Deploy data stewards to bridge centralized and decentralized governance approaches (see Case study 1).
- **Reduce silos and streamline approvals:** Enable greater internal access to data insights and the external sharing of data by reducing legacy silos through data mapping exercises and simplifying policy approval processes.
- **Explore approaches to address data scarcity:** Ensure company access to a sufficient volume of high-quality and representative data. Consider these approaches:
  - **Share data between organizations:** Explore the variety of sharing models and assess trade-offs. For example, data trusts are managed by a third party with a fiduciary duty to protect contributors' interests, whereas data cooperatives are member-

owned and governed. Organizations can reduce data training concerns from potential partners by contributing to efforts that standardize AI-related contractual provisions e.g. the Bonterms AI Standard clauses.<sup>21</sup> Proactive communication with the public on the goals and limitations of a data-sharing initiative is needed to secure trust and buy-in.

- **Collaborate on data analysis without sharing raw data:** One such method is federated learning, where a shared AI model is trained locally using data from decentralized edge devices or servers and only the model updates are shared with a central server for aggregation. Another technique employs data clean rooms – controlled environments where organizations upload their data for the retrieval of aggregated and anonymized insights.<sup>22</sup>
- **Synthetic data:** Examine how synthetically produced data may provide an alternative to data scarcity. Care is needed to proactively address governance challenges that synthetic data introduces, such as realism validation with non-replication of private data, provenance documentation and bias replication.



### Key roadblocks organizations encounter from the broader ecosystem

Undefined data governance standards for generative AI, challenging data minimization principles

Fragmented data governance due to factors such as AI nationalism, the nascent adoption of AI standards and conflicting regulations, affecting data sharing and interoperability

Limited incentives and lack of standards in emerging data exchange markets, preventing proper oversight of data quality, provenance tracking and fair compensation

Legal ambiguities, generating resistance from companies to share data for fear that recipients will use it to train AI models

Data market concentration generating information asymmetries and stifling competition, particularly for small- and medium-sized enterprises (SMEs) and start-ups

### Actions for government leaders

- **Clarify data governance to account for generative AI:** Assess the impact of generative AI on how businesses are incentivized to collect, retain, use and monetize data. Identify and address gaps in current data governance and content management policies. Consider affordances needed for vulnerable or marginalized populations, such as protecting indigenous data sovereignty rights<sup>23</sup> or children's rights<sup>24</sup> (see Play 7).
- **Promote open and inclusive data ecosystems:** Develop and harmonize policy and regulatory frameworks to enable responsible data sharing, including legal definitions and guidelines for emerging models like data trusts and cooperatives. For example, the EU's Data Governance Act supports trusted data intermediaries and promotes data altruism, laying the groundwork for new stewardship and sharing models.<sup>25</sup> Communicate legal clarity regarding when data can be used to train models to incentivize sharing without fear of exploitation. In the US, the AI Action Plan directs the National Science Foundation and Department of Energy to create secure compute environments for controlled AI access to restricted federal data, alongside the creation of an online portal for a demonstration project.<sup>26</sup>
- **Enable secure sharing through:**
  - **Experimentation support:** Establish regulatory sandboxes to test sharing models without facing compliance risks and provide compliance-by-design tooling.
  - **Mutually beneficial data-sharing markets:** Promote conditions of clear rules on use, value measurement, contributor rights and compensation, including for aggregators and individuals. Financial markets offer a proven blueprint: just as analysts evaluate stocks and shareholders earn dividends, data markets could employ analysts to assess data quality while compensating data owners for their contributions.
  - **Shared data infrastructure:** Facilitate secure, ethical and sovereignty-respecting access to high-quality domestic and cross-border datasets (see Case study 3).
- **Address synthetic data macro-challenges:** While synthetic data offers an alternative to data scarcity, it requires addressing challenges: incentivizing foundation models to revise usage policies that prohibit synthetic data production, providing transparency into biases and limitations, and preventing negative externalities following mass adoption (e.g. model collapse).<sup>27</sup>

## CASE STUDY 3

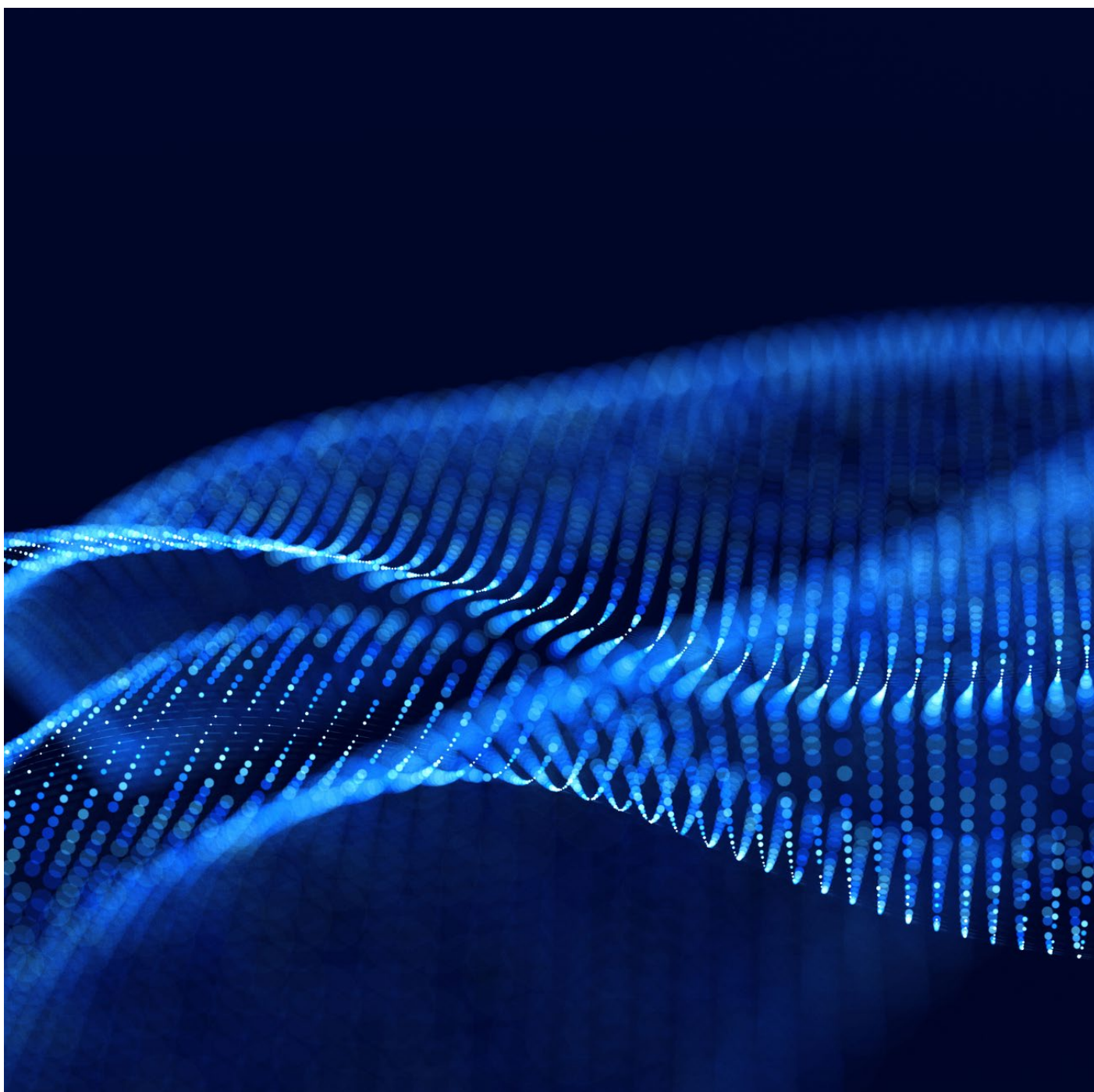
### Legal framework for hosting foreign data in Saudi Arabia

Fragmented legal frameworks introduce compliance barriers, particularly regarding data sovereignty, limiting a jurisdiction's ability to attract foreign investment to support AI economic development. In 2025, Saudi Arabia announced a Global AI Hub law that introduces three models to host foreign data within its borders while balancing data sovereignty rights with international cooperation.<sup>28</sup> The law enables Saudi Arabia to terminate any arrangement for reasons related to security, sovereignty or diplomacy, with a built-in transition period to facilitate data migration.



#### Key insight

Such an approach signals a shift towards territory-sensitive, sovereignty-respecting data infrastructures. New business opportunities are opened, but risks may still need to be assessed by organizations – such as ensuring legal adaptability, technical compliance and risk preparedness in environments where sovereignty-related constraints can rapidly shift regulatory control.

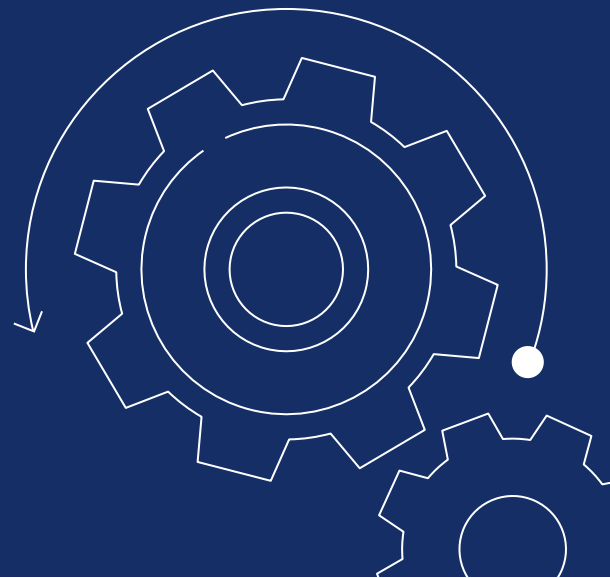




# Play 3

## Design resilient responsible AI processes for business continuity

Resilience is needed to future-proof organizations' AI strategies and their governance, ensuring their adaptability to AI convergence with other technologies,<sup>29</sup> emerging AI architectures, models and capabilities, and regulatory shifts. Organizations must also embed resilience, ensuring novel risks and opportunities are tackled as they arise.



### Organization leaders

#### Key roadblocks that arise within the organization

Problems in interpreting how current regulations affect AI

Ambiguity in the interpretation of AI regulations, creating uncertainty among companies about how to invest in responsible AI

#### Actions for organization leaders

- **Invest in strategic foresight:** Reduce uncertainties with methodological approaches, such as:
  - **Horizon scanning:** A structured evidence-gathering process that explores the strategic environment for early signals of change.<sup>30</sup>
  - **External engagement:** With regulators, standards and norms-making bodies to anticipate evolving policy directions, and with industry, academia and civil society to stay abreast of technological and societal developments.
  - **Scenario planning:** Methods that prompt organizations and individuals to look beyond their assumptions of the future.
- **Adopt a resiliency framework:** Prepare for unknown unknowns, including through adopting a resiliency framework that wraps around risk management:
  - **Contingency planning:** Identify each system that is critical to safety, mission, business, and security<sup>31</sup> and which, if impacted, could pose significant damage to the organization, its offerings, or the public. Prepare business continuity and contingency plans.
- **Knowledge sharing:** Prioritize multi-directional knowledge sharing (e.g. incident reporting or policy changes) that capture evolving AI uses, risks, and opportunities. However, balance policy adaptation frequency with predictability to enable adherence.
- **Balance global consistency and local responsiveness:** Multinational organizations must anchor AI governance in non-negotiable global principles while balancing local adaptation that reflects cultural and regulatory realities.
- **Prepare for fragmentation but invest in interoperability:** Organizations operating across jurisdictions should harmonize multiple AI risk frameworks by mapping common and competing elements to a unified master control set and crosswalk customized to the organization.<sup>32</sup> Organizations can enhance interoperability by engaging in international forums like the International Standards Organization (ISO), adopting broadly recognized standards and sharing best practices (see Play 5).

## CASE STUDY 4

### Infosys “comply up” general standard

Infosys adopts a “comply up” strategy, applying the highest global AI compliance standards – like those in the EU AI Act – across all operations worldwide. This unified approach eliminates complexity from fragmented regulations while exceeding client expectations, as partners increasingly demand robust, responsible AI practices regardless of local requirements. Infosys proved this model’s effectiveness with data privacy, where compliance with the California

Consumer Privacy Act (CCPA) created a strong baseline for global operations.



#### Key insight

Adopting the highest responsible AI standards across all jurisdictions streamlines operations and ensures consistent compliance regardless of the regulatory regime.



## Government leaders

### Key roadblocks organizations encounter from the broader ecosystem

Tensions arising from conflicting laws and overlapping authorities, creating difficulties in law enforcement and compliance<sup>33</sup>

AI regulatory fragmentation and policy instability, generating prohibitive compliance costs<sup>34</sup> and threatening confidence in investments in responsible AI practices

### Actions for government leaders

- **Resolve regulatory tensions and ambiguities between sectoral and cross-cutting AI regulations:** Provide organizations with clear guidance on compliance requirements.<sup>35</sup>
- **Prototype AI governance frameworks:** Enhance policy efficacy and feasibility, and mitigate externalities (e.g. economic or rights/freedoms infringements), through policy prototyping, which borrows design and research practices from products and services.<sup>36</sup>

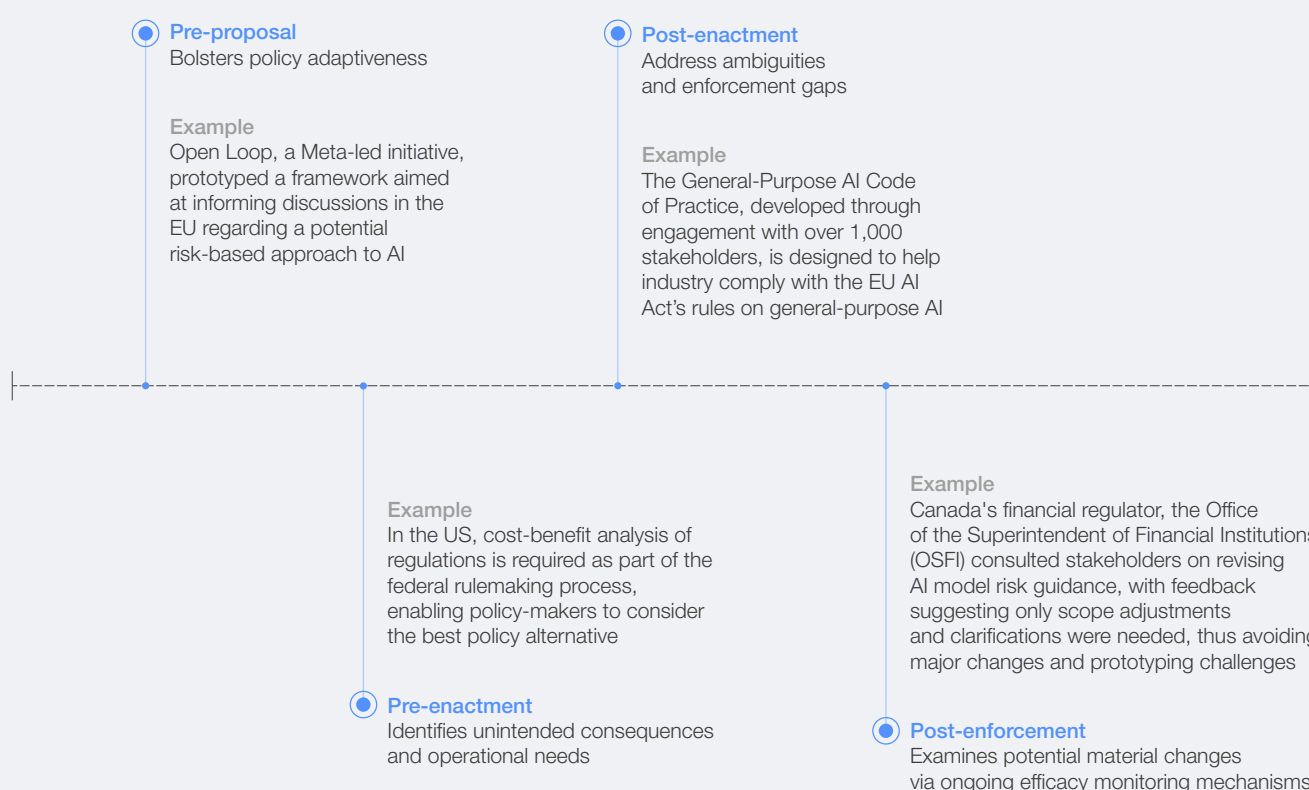
Best practices include:

- **Incentivized participation:** Across organization size, sector, expertise and the public to ensure prototyping considers all impacted parties and their concerns (e.g. intellectual property loss).
- **Clear criteria:** Set goals, metrics and benchmarks for success upfront.
- **Robust methods:** Avoid testing in isolation from existing policies, policy-making processes and enforcement practicalities. Layer prototyping approaches and prototype at multiple stages (see Figure 4). Refine through agile iteration cycles and feedback loops.

- **Transparent process:** Document and communicate decisions, changes and rationale throughout the process. Provide sufficient time for submission and review of feedback.
- **Independence:** Prototyping should be adopted in a manner that bolsters rather than impedes a policy-making process, representing and benefiting the entire population. For example, organization participation for ulterior motives (e.g. regulatory capture or dilution of policy accountability) should be deterred.
- **Promote jurisdictional interoperability through multilateral AI governance frameworks:** Set shared principles, standards and certification protocols to drive innovation and safety while respecting national interests. Help businesses make sense of multiple frameworks through developing crosswalks e.g. the National Institute of Standards and Technology (NIST) *AI Risk Management Framework* (RMF) and Japan AI Guidelines for Business.<sup>37</sup> Consider participation in multilateral forums that enable international cooperation (e.g. the World Economic Forum’s AI Governance Alliance and the Commonwealth Artificial Intelligence Consortium), as well as collaboratively working towards reducing fragmentation (e.g. China-Pakistan AI Cooperation efforts on innovation and governance).<sup>38</sup>

- **Provide foundational support to enable organizations in responsible AI:** Some approaches could include:
  - **Resources and tooling:** Singapore provides governance frameworks, tools, training and certifications, such as the Model AI Governance Framework and AI Verify Foundation.
  - **Ecosystem champions:** Canada has established three non-profit AI Institutes (Amii, Mila and the Vector Institute) that serve as third-party facilitators of cross-sector interaction and public-private engagement to align priorities and share best practices.
- **Sandboxes:** The UK's Financial Conduct Authority (FCA) offers services for safe experimentation to firms in various stages of AI, from discovery to use.<sup>39</sup> Firms gain regulatory support and validation for confident adoption, while the FCA gains practical insights to shape future oversight and policy. Another example is the Government of India's effort, in collaboration with the Centre for the Fourth Industrial Revolution India, to develop a roadmap for establishing an AI sandbox ecosystem tailored to India's unique needs and sectoral priorities.<sup>40</sup>

FIGURE 4 When to prototype policies and regulatory frameworks



**Sources:** Meta Open Loop. (2021). *AI Impact Assessment: A Policy Prototyping Experiment*; US Library of Congress. (2024). *Cost-Benefit Analysis in Federal Agency Rulemaking*. <https://www.congress.gov/crs-product/IF12058#>; European Commission. (2025). *General-Purpose AI Code of Practice now available*. [https://ec.europa.eu/commission/presscorner/detail/en/ip\\_25\\_1787](https://ec.europa.eu/commission/presscorner/detail/en/ip_25_1787).

# Dimension 2:

## Governance and accountability

Increase and incentivize organizational capacity for responsible AI.

This chapter addresses the internal governance structures and external incentives needed to empower leadership and ensure that responsible AI becomes a durable organizational capability.

**Play 4:** Appoint and incentivize AI governance leaders



**Play 5:** Adopt a systematic, systemic and context-specific approach to risk management



**Play 6:** Provide transparency into responsible AI practices and incident responses

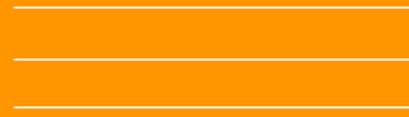
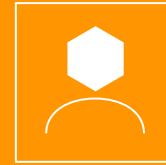




# Play 4

## Appoint and incentivize AI governance leaders

Responsible AI senior leaders enable robust governance frameworks that provide boards of directors with assurance of regulatory compliance across the enterprise, consistent risk thresholds and strategic business alignment.



## Organization leaders

### Key roadblocks that arise within the organization

Highly unstructured governance, unclear accountability, and insufficient top-down guidance affecting the implementation of “responsibility”<sup>41</sup>

Problems in identifying and hiring leaders with interdisciplinary knowledge

Misalignment between business functions, inducing misunderstandings about AI definitions, risks, responsibilities and success metrics

### Actions for organization leaders

- **Appoint a senior AI governance leader and cross-functional AI governance body:** This provides confidence in AI risk oversight and compliance and supports alignment with broader business objectives.

#### Key factors include:

- **Set AI governance as the leader’s primary responsibility:** If resources only afford designation as an additional responsibility, ensure the individual can be sufficiently dedicated to the role.
- **Evaluate where to house AI governance:** Examine the trade-offs of housing governance responsibilities within an existing or new function. Organizations report variability in focus areas for individuals assigned to responsible AI, such as privacy, ethics and risk, or analytics.<sup>42</sup> Ensure leaders are resourced to act cross-functionally to advance end-to-end governance.

- **Promote cross-functional alignment:** Intentionally align on terminology and expectations by defining key AI terms accompanied by examples, including edge cases.
- **Segregate duties:** Separate responsibilities between delivery and assurance teams and distribute discrete critical functions of authorization, custody, record-keeping and reconciliation across independent teams. Prevent any single entity from holding unchecked control over AI processes or assets to reduce risks of errors, fraud and regulatory breaches.
- **Use a phased approach to maturing AI governance:** Start with a centralized governance model (e.g. a cross-functional committee) to ensure consistency and accountability. As practices mature, evolve towards a more federated or hybrid model that empowers business units with context-aware oversight (see Case studies 1 and 5).

## CASE STUDY 5

### e&'s structured approach to cross-functional AI governance

As AI adoption accelerated across functions, e& (formerly Etisalat Group) identified a need for structured governance that could guide decentralized use case owners in navigating complex risks. e& established an AI Governance Steering Committee – with representatives from data privacy, cybersecurity, enterprise risk and technology – to provide advisory support, risk reviews and escalation paths for use cases. Regular cross-functional refreshers help the committee remain aligned with evolving standards.



#### Key insight

Embedding governance through functional steering ownership increases early-stage risk flagging. Further, ongoing regulatory refreshers ensure AI risk awareness remains actionable and enabled across decentralized teams.



## Government leaders

### Key roadblocks organizations encounter from the broader ecosystem

Unclear responsibility allocation across the AI value chain, creating systemic risks and deterring organizations from establishing responsible governance frameworks

Lack of shared accountability mechanisms among AI stakeholders, diluting oversight and weakening governance across the AI supply chain

### Actions for government leaders

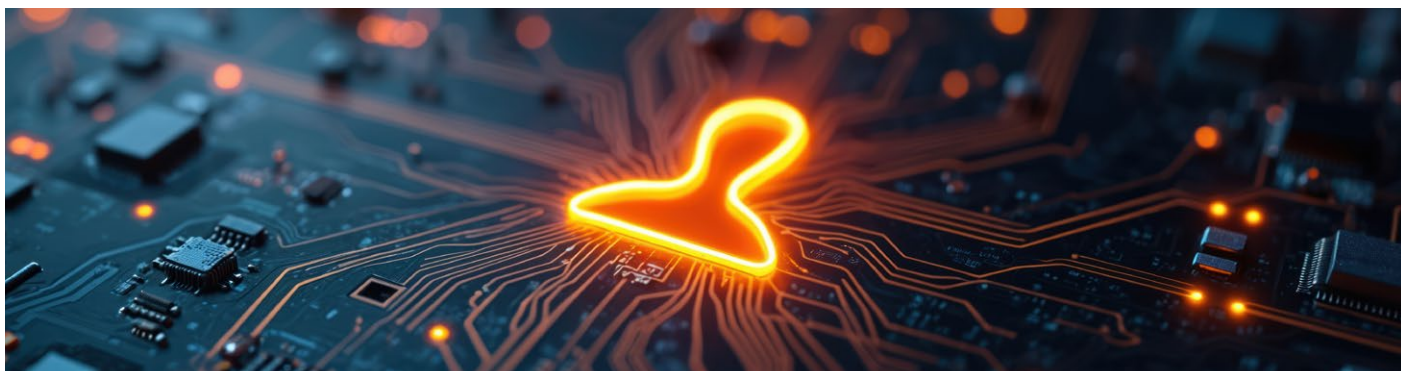
- **Increase responsibility clarity on AI supply chain:** Motivate companies to allocate responsibility internally by clarifying responsibility at the supply chain level.

#### Key actions include:

- **Examine the problem:** Understand the varied underlying challenges to clarifying allocation, especially in the generative AI era, to enhance the efficacy of solutions.<sup>43</sup>
- **Promote ecosystem actions:** Addressing supply chain challenges requires coordinated efforts. Governments should incentivize industry evaluations

and benchmarks, define criteria for responsibility transfers and advance international alignment on responsibility allocation norms.

- **Exemplify responsible AI leadership:** In addition to setting up AI leaders across government functions – such as a national chief AI officer, AI leaders in government agencies or a cross-cutting chief AI officers council<sup>44</sup> – publicly appoint or designate senior governance leaders with well-defined responsibilities. A supporting body, such as Canada's proposed AI ethics review board, can enhance framework adoption by providing responsible AI guidance to higher-risk or impact projects.<sup>45</sup>

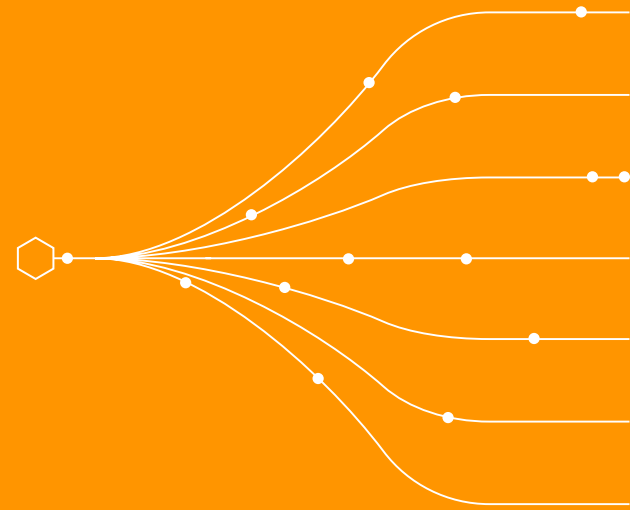




# Play 5

## Adopt a systematic, systemic and context-specific approach to risk management

The business implications of unmanaged AI risk exposure are far-reaching. A systematic, systemic and context-specific approach is needed to align responsible AI decision-making with risk-exposure and tolerances specific to the organization's business size, sector, jurisdiction, operational structure and other contextual attributes.



### Organization leaders

#### Key roadblocks that arise within the organization

Misperception of responsible AI maturity, creating an overestimation of progress in responsible AI implementation<sup>46</sup>

Underestimation of risk management, viewing it more as a niche technical challenge than as an enterprise responsibility

Low prioritization of risks, affecting AI risk mitigation and management measures, especially for organizations with limited resources despite the various AI risk management frameworks available

Outdated procurement reviews, preventing the assessment of risks from AI vendors and third-party software with AI features

#### Actions for organization leaders

- **Conduct a maturity assessment:** Companies should assess the current state of their responsible AI implementation. For example, the Global System for Mobile Communications Association's (GSMA) Responsible AI Maturity Roadmap is an industry-led initiative to help telecommunications organizations adopt and measure responsible and ethical approaches to AI.<sup>47</sup>
  - **Repeat:** Assess regularly to identify improvements, as well as responsible AI impacts and gaps that emerge with the evolving landscape.
  - **Communicate:** Provide the public with transparency into the state of the organization's responsible AI practices (see Play 6).
  - **Tailor high-level external frameworks to organizational contexts:** Invest in adapting generalized risk assessment frameworks to internal control structures, define sector-specific risk scenarios, and integrate standardized and repeatable risk management processes into the organizational value chain and AI life cycle checkpoints: design, development, procurement, deployment and decommissioning (see Case study 6).
- Best practices include:**
- **Comprehensive:** Review governance structures, policies, standards, risk management processes, technical safeguards, workforce capabilities, data practices, accountability mechanisms and alignment with responsible AI principles.
  - **Context-specific:** Perform assessments that are tailored to the context.

- **Make use of emerging context-specific guidance:** Organizations should consider participating in community-based working group efforts that are under way to interpret actor-agnostic risk management frameworks for specific business contexts, such as MLCommons and the OWASP Generative AI Security Project. Examples of context-specific frameworks include:
    - **Activity-based:** *Risk Management Framework for the Procurement of AI Systems (RMF PAIS 1.0)*, adapted from traditional risk management frameworks (e.g. ISO 31000).<sup>48</sup>
    - **Size-based:** *Responsible AI Startups (RAIS) Framework*, providing guidance for the venture capital industry for investing in early-stage companies.<sup>49</sup>
    - **Sector-based:** Monetary Authority of Singapore's *Artificial Intelligence Model Risk Management* information paper, providing guidance for financial institutions.<sup>50</sup>
- Contradictory recommendations could emerge from using multiple context-specific frameworks, which can be mitigated by AI literacy and accountability.

## CASE STUDY 6

### Adapting the NIST AI RMF at Workday

Workday aligned cross-functionally to map its existing controls covering policy, risk evaluation and third-party tool assessments to the NIST AI RMF's categories and taxonomy. An AI Advisory Board, including C-suite executives, steered the programme, managed edge cases and enforced reporting lines between developers and governance teams. Workday also implemented an RMF-based responsible AI questionnaire for evaluating third-party AI tools and updated its data sheets for transparency.<sup>51</sup>



#### Key insight

By operationalizing NIST AI RMF into standardized templates and tooling, organizations can continue to refine their risk management approach while embedding controls into existing processes to advance responsible, transparent and risk-aware AI deployment at scale.



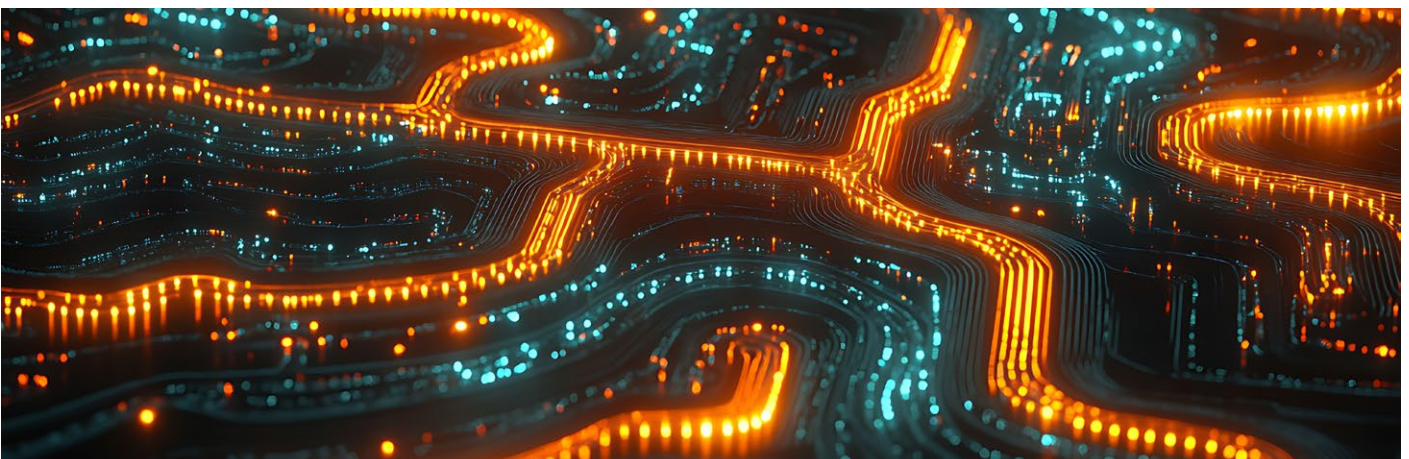
## Government leaders

### Key roadblocks organizations encounter from the broader ecosystem

Competing incentives to self-assess responsible AI maturity, creating fear that limitations in their responsible AI capabilities will expose them to legal liability

Limited awareness of industry best practices, affecting their risk management strategies

Difficulty in adapting industry- or actor-agnostic risk management frameworks, such as NIST or ISO, due to high workload, complexity of guidelines, and limited human and financial resources of organizations<sup>52</sup>



## Actions for government leaders

- **Drive development of context-specific frameworks:** Where possible and practical with local norms, interpret widely accepted frameworks. This can prevent a fragmented landscape, where each industry and sub-industry develops its own AI governance terminology, definitions and practices.

Recommended actions in developing context-specific frameworks include:

- **Engage stakeholders:** Gather diverse expertise across sectors, company sizes, and impacted end-users and communities, e.g. the US AI Action Plan requires NIST to convene a broad range of stakeholders to accelerate the development and adoption of domain-specific national standards for AI systems.<sup>53</sup>
- **Map context:** Identify risks, opportunities, regulatory requirements and context-specific guidance, such as those provided by an industry group or civil society organization.
- **Draft framework:** Detail policies with accountability structures, use cases and best practices. Provide tiered governance models with essential controls and progression paths as organizations grow.
- **Prototype:** Gather public feedback on drafts and pilot-test with organizations (see Play 3). Ensure ongoing evaluation and refinement.
- **Incentivize adoption:** Provide framework training and a Q&A channel for communication with policy-makers to enhance uptake. Test rewards for adoption, such as compliance recognition.
- **Promote open sharing of responsible AI resources:** Sharing best practices, case studies and tools can enable access to context-specific insights and prevent unnecessary trial-and-error lessons already learned by other

organizations. For SMEs, access to responsible AI tools is critical to implementation. Single-entry-point resource repositories should be developed or enhanced when they already exist. Examples include the Organisation for Economic Co-operation and Development's (OECD) Catalogue of Tools and Metrics for Trustworthy AI,<sup>54</sup> Canada's AI and Data Governance Standardization Hub<sup>55</sup> and the World Bank Group's AI-in-Government Case Study Repository.<sup>56</sup>

Factors behind a successful repository include:

- **Scope:** Clear objectives such as target audience, resources and outcomes
- **Accessibility:** Multi-language support and navigable search and retrieval with relevant labels
- **Quality:** Reliable, curated and up to date
- **Comprehensiveness:** Diverse contributions
- **Inclusivity:** Caters to diverse expertise and businesses with varying levels of maturity and resources
- **Transparency:** Governance and curation
- **Sustainability:** Long-term funding and maintenance
- **Incentivized contribution:** Reputation promotion, privileged access to additional content, IP and confidentiality assurances
- **Cooperate internationally to reduce the digital divide:** Insufficient participation from global majority countries in international AI governance discourse can lead to significant knowledge gaps about AI risks and opportunities. Additionally, cooperation is needed to ensure that AI addresses, rather than exacerbates, current structural limitations and power imbalances for the global majority related to infrastructure, data and talent.<sup>57</sup>



# Play 6

## Provide transparency into responsible AI practices and incident responses

For industry and government leaders alike, transparency is foundational to trust, legitimacy and regulatory preparedness. Expectations rely on evidence of oversight, mitigation and continuous improvement. As governments begin mandating AI transparency requirements, companies that proactively develop reporting mechanisms will be better positioned.



### Organization leaders

#### Key roadblocks that arise within the organization

Gaps in the continuous monitoring of AI impacts and their downstream effects, reducing early detection and mitigation

Unassessed third-party AI tools, limiting the ability to accurately track AI risks across the enterprise

A lack of consensus on responsible AI technical standards and of contextually relevant criteria in assessment frameworks, failing to account for risk variation by sector and use case,<sup>58</sup> complicating efforts to effectively benchmark and audit AI systems across jurisdictions

A lack of enterprise-wide protocols, impacting escalation processes for identifying and reporting AI incidents – these remain inconsistent and reactive

#### Actions for organization leaders

- **Champion employee self-reporting:** Support an environment of information sharing and transparency. Develop accessible mechanisms for employees to raise concerns or report incidents related to AI.
- **Establish incident response plans:** Define standardized typologies of AI incidents (e.g. harm to users, environmental overconsumption, fairness violations) and set disclosure thresholds that trigger internal reviews or external reporting. One potential resource is MIT's AI Incident Tracker, a tool that uses AI to process reports from the Responsible AI Collaborative's AI Incident Database before categorizing them with established frameworks, as well as risk and harm severity assessments.<sup>59</sup>
- **Prioritize custom tests and metrics over generic benchmarks:** Increase compliance and reduce risk exposure by encompassing domain- and application-specific risk areas and regulated activities. Prioritize inclusive benchmarks that account for diverse user bases to improve assessment of reasoning, ethics and linguistic depth across global contexts.<sup>60</sup>
- **Provide transparency into responsible AI practices:** Document all AI use cases in AI inventory reporting systems, in terms of use, purpose, data sources and ownership. Maintain an AI risk registry to track potential and realized risk and mitigation guidelines. Use transparency instruments to provide insight into the organization's responsible AI practices (see Table 2). With increasing expectations of reporting on responsible AI practices, companies need to proactively adapt and translate their internal governance policies for a public audience.



### Key roadblocks organizations encounter from the broader ecosystem

Limited incentives to report on responsible AI practices, discouraging companies from sharing information for fear of facing reputational and liability issues

Lack of standardized incident reporting protocols, impeding the collection of reliable and comprehensive data, critical for preventing and mitigating future incidents

Opacity of AI's environmental impact: 84% of generative AI use is done through undisclosed models.<sup>61</sup> As AI adoption grows, data on the environmental impacts is increasingly scarce, fuelling misinformation and public misconceptions.

Ethics washing occurs when companies overstate their capabilities in responsible AI, creating an uneven playing field where genuine efforts are discouraged or overshadowed by exaggerated claims.

Static benchmarks becoming misaligned with emerging risks, especially as many popular benchmarks are reaching saturation points or suffer from a lack of transparency, reproducibility and real-world relevance

### Actions for government leaders

- **Assess the state of responsible AI practices in the industry:** Policy-makers must understand the state of responsible practices by AI providers and industry users within their jurisdiction. Such assessments can:
  - **Incentivize organizations to measure maturity:** Build awareness of the actual state of responsible AI practices within the organization.
  - **Support evidence-based policy:** Educate policy-makers on industry practices and responsible AI implementation challenges.
  - **Prevent unnecessary regulation:** In cases where enough companies demonstrate proactive and sufficient risk management.
  - **Provide insights into forthcoming AI capabilities:** Stay abreast of developing AI to assess potential opportunities and challenges of jurisdictional interest, such as national security.

Jurisdictions should consider the advantages and limitations of various reporting instruments when incentivizing industry to report responsible AI practices (see Table 2). Layering multiple instruments can help offset trade-offs and bolster overall efficacy (see Case study 7). Mandating reporting in select instances may offset participation challenges. Additionally, governments should support academia, civil society, and third-party efforts to assess the state of responsible AI practices.

- **Standardize and incentivize risk and incident reporting:** Promote compliance, data quality and

insights gathering across jurisdictions through harmonized taxonomies, safe harbour provisions, and interoperable disclosure platforms that encourage transparency while safeguarding innovation. The level of disclosure for risks and incidents may vary depending on the audience and availability of expertise and resources required to analyse information.<sup>62</sup> Disclosures must balance data privacy and security, particularly when reporting incidents related to vulnerable populations. Participation incentives for reporting could include access to other organizations' reported incidents or mandated disclosure. For example, the EU AI Act requires general-purpose AI model providers of high-risk systems to "track, document and report relevant information about serious incidents and possible corrective measures to address them."<sup>63</sup>

- **Drive the evolution of benchmarks and standardize validation:** Access to updated code, test sets and validation methods is needed to ensure companies and regulators base decisions on accurate metrics for system performance. Convene industry, academia and civil society to identify benchmarks and standards for AI safety assessments across industries and contexts. For example, Singapore's Global AI Assurance Pilot gathered 17 organizations from 10 industries and nine countries to co-develop norms and practices for generative AI testing.<sup>64</sup>
- **Facilitate environmental transparency disclosures:** Considerations must be given to voluntary or mandatory industry-wide measures that include publishing impact data (e.g. energy, carbon, water) across the AI value chain, integrating AI's environmental costs into corporate reporting and procurement, and developing standardized verification processes (see Case study 2).<sup>65</sup>

TABLE 2 Instruments for reporting on responsible AI practices, by content

| Content type                                                                     | Instruments                                                                                                                                                                                                                                                                                                                                                                          | Considerations (non-exhaustive)                                                                                                                                                                                                                                                                                                                                                                       |
|----------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Commitments</b><br>How an organization says it will implement responsible AI  | <ul style="list-style-type: none"> <li>– <b>Individual:</b> Informal (blogs, speeches) or formal commitments (principles, policies, frameworks) e.g. Perplexity Acceptable Use Policy<sup>66</sup></li> <li>– <b>Joint:</b> Commitments from multiple organizations (see Case study 7)</li> </ul>                                                                                    | <b>Advantages:</b> <ul style="list-style-type: none"> <li>– Agile method for signalling norms</li> <li>– Flexible to organization context</li> </ul> <b>Limitations:</b> <ul style="list-style-type: none"> <li>– Low adoption with varied adherence</li> <li>– Limited public evidence correlating responsible AI commitments with implementation<sup>67</sup></li> </ul>                            |
| <b>Claims</b><br>How an organization self-reports its responsible AI practice    | <ul style="list-style-type: none"> <li>– <b>Reports:</b> Detailing practices e.g. Microsoft 2025 Responsible AI Transparency Report<sup>68</sup></li> <li>– <b>Cards:</b> Insights into the development, governance and safety of an AI model, system of models, or service e.g. Cohere Command R and Command R+ Model Card<sup>69</sup></li> </ul>                                  | <b>Advantages:</b> <ul style="list-style-type: none"> <li>– Provides a benchmark for other companies</li> <li>– Promotes feedback and accountability</li> </ul> <b>Limitations:</b> <ul style="list-style-type: none"> <li>– Variability can hinder standardized comparisons across multiple companies</li> <li>– Self-reporting bias may occur</li> </ul>                                            |
| <b>Evidence</b><br>How an organization substantiates its responsible AI practice | <ul style="list-style-type: none"> <li>– <b>Certifications:</b> A review typically aligned with a set criteria e.g. Anthropic certified by Schellman Compliance, LLC against ISO/IEC 42001:2023<sup>70</sup></li> <li>– <b>Sandboxes:</b> Third-party controlled or monitored environments for AI testing e.g. the United Arab Emirates regulatory sandboxes<sup>71</sup></li> </ul> | <b>Advantages:</b> <ul style="list-style-type: none"> <li>– Provides credibility if certified by a reputable party</li> <li>– Incentivized adoption in pursuit of market differentiation</li> </ul> <b>Limitations:</b> <ul style="list-style-type: none"> <li>– Variability in certification methods risk practice fragmentation</li> <li>– Costly to implement and address renewal needs</li> </ul> |

## CASE STUDY 7

## The Hiroshima AI Process International Code of Conduct and Reporting Framework

In 2023, under Japan's presidency, the G7 launched the Hiroshima AI Process, resulting in the *Hiroshima Process International Code of Conduct for Organizations Developing Advanced AI Systems*.<sup>72</sup> This voluntary code promotes ethical, transparent and secure practices. To reinforce accountability, the G7 and OECD introduced a voluntary reporting framework in 2025 for organizations in member and partner countries. While initial reports were submitted, variations in detail and transparency highlighted limitations in consistency and comparability. The framework's voluntary nature also raised challenges in participation and adherence. The G7, under its Canadian presidency, is exploring additional incentives and clearer guidance. There is also a

forum led by Japan for broader collaboration through the Hiroshima AI Process Friends Group, which now comprises 56 countries and regions.<sup>73</sup> Increasing participation by organizations across diverse jurisdictions will also require reporting requirements to consider language and timing.



### Key insight

Commitments to voluntary frameworks alone are insufficient for ensuring transparent and accountable responsible AI practices by organizations. They likely require the layering of instruments to assess claims (see Table 2), such as standardized reporting.

# Dimension 3:

## Development and use

Oversee the life cycle of responsible AI development, acquisition and use.

This chapter focuses on indispensable technological tools, technical standards and ongoing governance enablement across the AI life cycle, from acquisition and design to deployment and ongoing monitoring.

**Play 7:** Drive AI innovations with responsible design as the default



**Play 8:** Scale responsible AI with technology enablement



**Play 9:** Increase responsible AI literacy and workforce transition opportunities

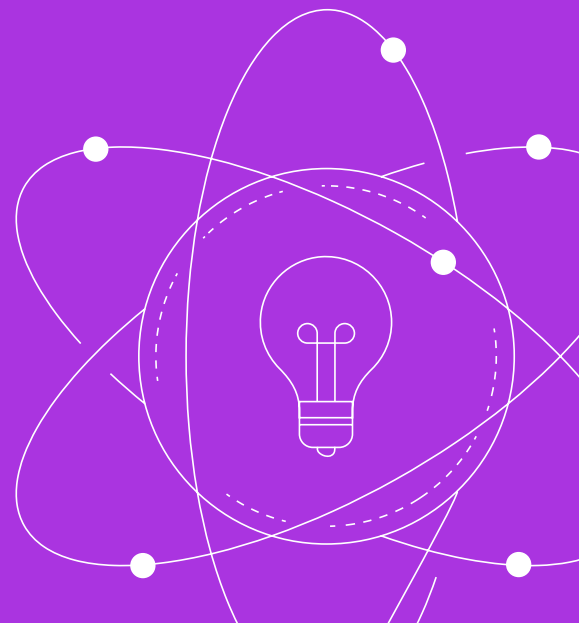




# Play 7

## Drive AI innovations with responsible design as the default

For responsible AI implementation to succeed at scale, organizations must reconfigure the foundational conditions that shape how AI is designed into products and services. Without integrating responsible design principles, even well-intentioned human-AI interaction design methodologies can erode user trust and social well-being.<sup>74</sup>



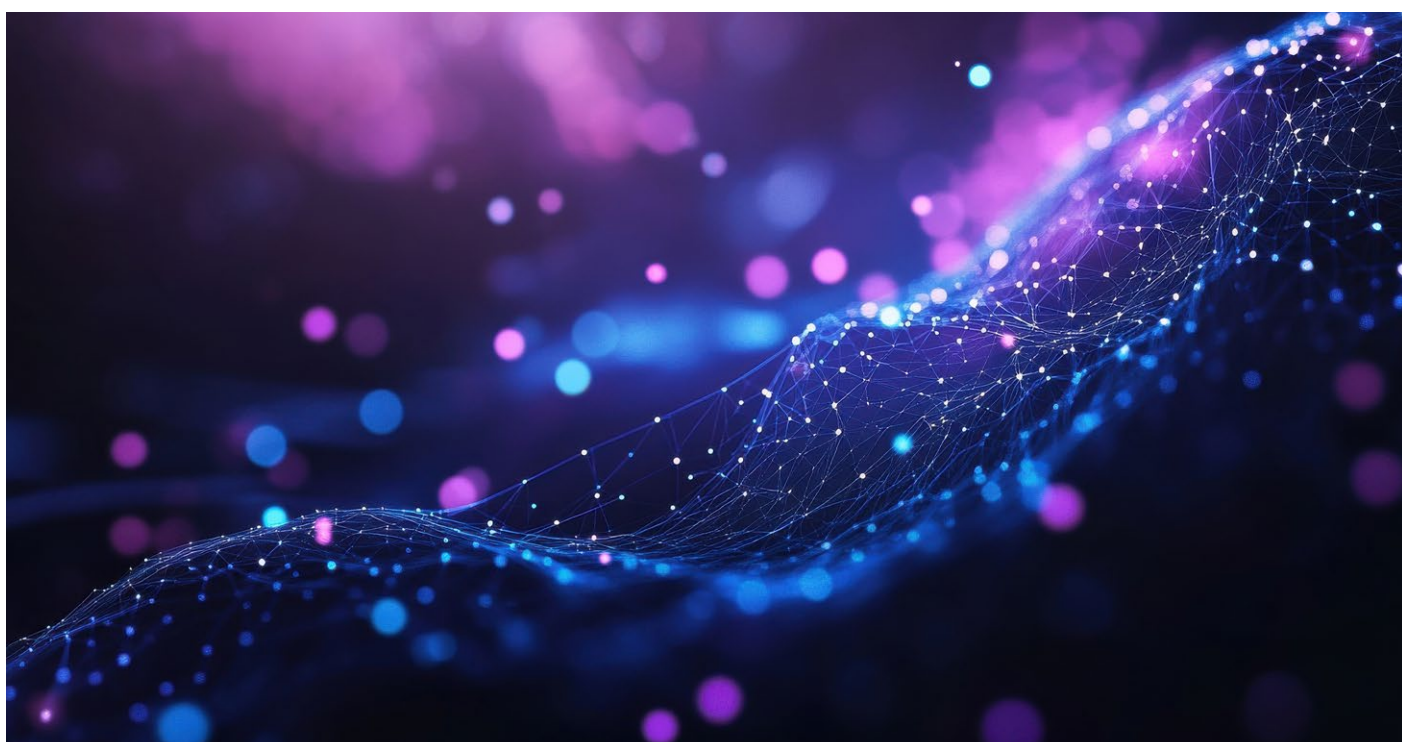
### Organization leaders

#### Key roadblocks that arise within the organization

A narrow focus on user understanding, well-being and agency<sup>75</sup> in human-AI interaction design, resulting in employees lacking awareness of the system's capabilities and limitations, over- or under-trust in AI outputs, flawed decisions/outcomes, uninformed or forced consent, limited transparency and an inability to provide feedback or challenge decisions

Over-optimization based on positive user experience metrics, while neglecting harmful edge cases, can bias modern design approaches such as user-centred, behavioural or frictionless design, enabling manipulation, misuse, addictive use, polarization, and privacy or safety risks

Lack of leadership support and weak incentives for responsible innovation, perceived as antagonistic to fast-paced, agile workflows, forcing teams to navigate short- and long-term trade-offs alone, and avoiding more comprehensive design and development cycles for responsible AI practices



## Actions for organization leaders

- **Prioritize and resource responsible AI design practices:** Efforts to encourage and adequately resource responsible design practices within the organization include:
  - Embed responsible design into performance metrics, resource allocation and recognition programmes.
  - Encourage employees to question existing design approaches and instil ethical and compliant measurements of success.
  - Design for potential negative outcomes by identifying risks and failure scenarios. Build systems with resilience and mechanisms to “fail safely” to ensure continuity and minimize impact when issues arise.<sup>76</sup>
  - Re-evaluate products already deployed<sup>77</sup> to assess gaps in responsible design.
- Integrate responsible AI criteria into procurement and third-party risk management processes to mitigate downstream risks and signal responsibility expectations.<sup>78</sup> Including confidence scores, limitation warnings or a reduced authoritative tone can mitigate the impacts of hallucinations.
- **Build awareness and ownership of established design principles:** Increase understanding of design-specific risks and mitigations. Assign responsible AI stewards across product teams (see Case study 2) and integrate multi-disciplinary design teams into the AI development life cycle.
- **Empower users as partners in responsible AI:** Engage users (e.g. employees, customers and partners) to contribute to responsibility throughout the AI life cycle (see Case study 8). For instance, experts from MIT and Stanford University proposed a new framework that allows third-party users to disclose flaws and monitor AI developers’ responses and resolutions.<sup>79</sup>

## CASE STUDY 8

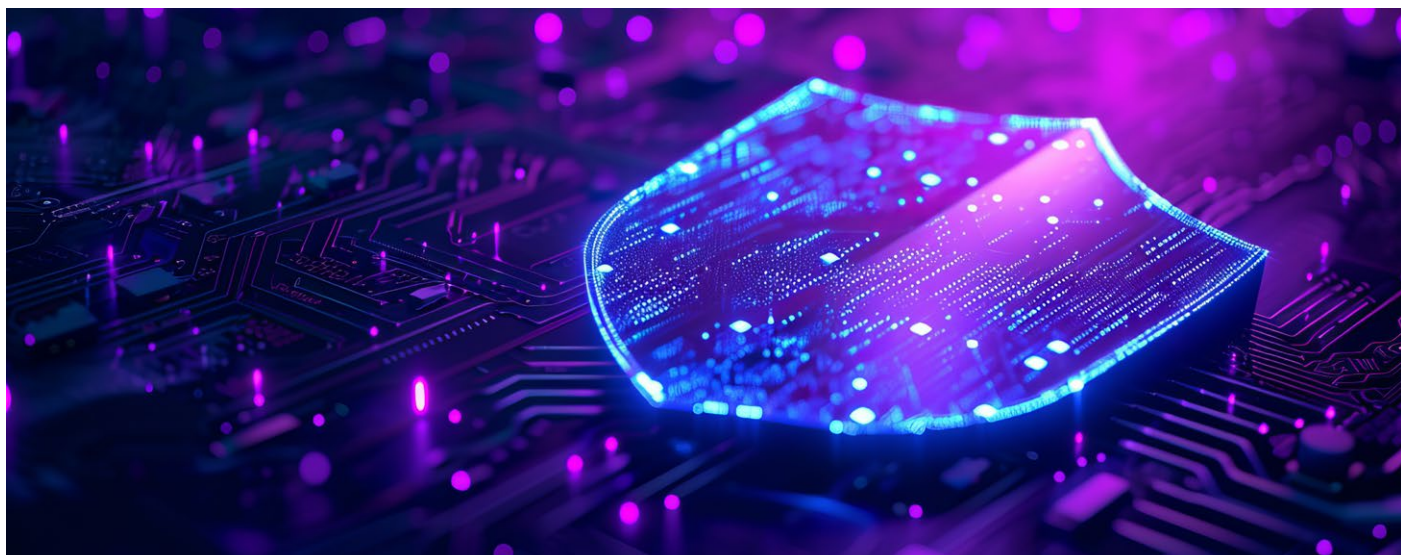
### Co-designing with children for responsible AI innovation

Current AI product development often lacks sufficient consideration of children’s rights and well-being, leading to potential issues with inappropriate content, bias and unequal access. The Alan Turing Institute, in collaboration with the LEGO Foundation, developed a participatory research process to explore how generative AI impacts children.<sup>80</sup> The project – which surveyed over 1,700 children, parents, caregivers and teachers – conducted school-based workshops to capture children’s direct experiences with tools like ChatGPT and DALL·E. The report recommends a child-centred approach to generative AI, including the meaningful involvement of children in the design process.



#### Key insight

The research revealed that children both understand the implications of generative AI and are eager to shape its future, sharing concerns about misinformation, environmental impact and online safety. Children favoured socially beneficial AI uses and opted for creative offline alternatives when available. This study demonstrates that involving users as active partners in product design provides valuable insights to identify or mitigate risks and harms.





### Key roadblocks organizations encounter from the broader ecosystem

Absence of guiding principles, benchmarks, and shared accountability structures, impacting responsible AI design and implementation.

New AI industries without design standards, such as AI therapy and companionship (which are emerging as the number-one generative AI use case<sup>81</sup>), highlight sensitive data collection and privacy concerns and pose unique challenges in terms of trust and security, drawing attention to the need for metrics that assess risk across psychological, ethical and social dimensions.

Misaligned expectations across the AI value chain, from third-party vendors to the organization's specific responsible AI practices, leading to friction or inconsistencies in downstream use.

Reliance on venture capital or corporate backing as AI funding models, prioritizing short-term monetization and market success over long-term governance of products that promote the common good.

### Actions for government leaders

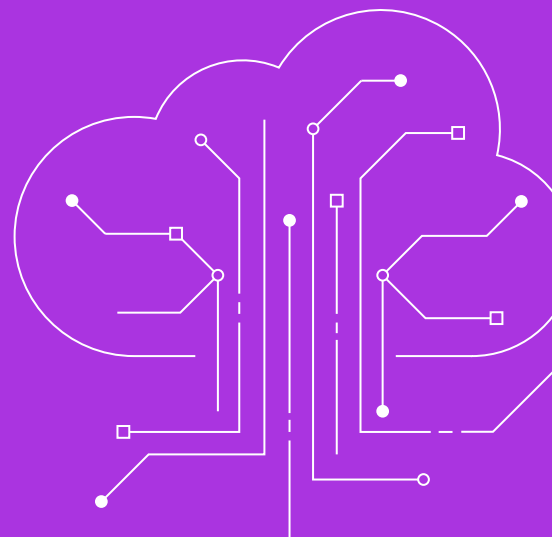
- **Take a socio-technical approach to risk management:** Evolve government AI frameworks, policies and regulations to move beyond narrow technical engineering perspectives and consider the role of broader societal forces in determining AI's outcomes.<sup>82</sup> Example approaches include:
  - Fund interdisciplinary research on AI's economic, social, environmental and political effects.
  - Ensure employees have a voice in the deployment of workplace AI technologies, including protecting organizing rights, strengthening whistleblower protections and prohibiting surveillance practices that deter collective action or expression.
  - Prevent outsized influence from any individual stakeholder group in deciding what constitutes risk or harm, or to what values AI should be aligned.<sup>83</sup>
- **Harmonize responsible design standards for AI:** Collaborate across borders and work with the design community and impacted stakeholder groups to create consensus around design risks and mitigation approaches (see Case study 8). Develop public toolkits to drive awareness and fund sandboxes to experiment with safety-centred user experience (UX) innovation. Encourage adoption of standards and frameworks for impacted stakeholder groups. For example, AI products used by children need design standards for age-appropriate interfaces, explainability and safeguards against manipulation, false or misleading information.<sup>84</sup>
- **Address evolving human-AI interaction impacts:** Adopt a multi-pronged approach, which could include:
  - Informing ethical design standards with multi-disciplinary research that assesses impacts across diverse stakeholder groups, such as child-<sup>85</sup> or older adult-facing<sup>86</sup> products offering AI companionship.
  - Proactively examine emerging areas of human-AI interaction, such as AI use in neurotechnology.
  - Assess impacts on data practices, including collection and monetization of sensitive data, such as for engagement-based design. Address gaps in data governance policies (see Play 2).
  - Creating public campaigns to increase awareness of the benefits and risks, including AI literacy in education systems.
  - Working with multilateral bodies to enforce broad international adherence to human rights.
- **Incentivize the diversity of business models:** Encourage approaches to alternative revenue generation opportunities that can deliver AI products with greater human alignment and evaluate models based on measures of success beyond profit and engagement metrics, such as contributions to scientific advancement and/or societal well-being. Enable academia and civil society to participate in public-interest frontier AI R&D with public compute, data access, and focused research grants, to offset the high costs associated with AI initiatives.



# Play 8

## Scale responsible AI with technology enablement

As AI applications multiply at pace and the risk landscape grows more complex, responsible AI technologies become indispensable – from operationalized platforms to systemic enablement and continuous oversight.



## Organization leaders

### Key roadblocks that arise within the organization

Limited visibility into enterprise-wide AI usage and risks, impacts maintaining a systematic and comprehensive inventory capturing all assets in use

Contending with technical debt from legacy technologies, exposing organizations to heightened AI risks and security vulnerabilities (which also hinder the systematic implementation of technologies designed to integrate trust and regulatory compliance into AI systems)

Human review bottlenecks, preventing automation of risk assessments for AI use cases and resulting in less responsive processes

### Actions for organization leaders

- **Systematize responsible AI:** Identify and use dedicated technology solutions that support the operationalization and scaling of responsible AI tasks, including for and with agentic AI systems (see Case study 9). Examples include:
  - **Real-time monitoring:** Multiple technologies can support continuous AI oversight. A control plane offers centralized governance across distributed systems, while monitoring tools, sensors and agents enable real-time tracking of system performance, security events and adherence to responsible AI and compliance metrics.
  - **AI agents:** These can support in analysing vast threat intelligence and delivering real-time assessments.<sup>87</sup> They may also enhance risk management by scanning and evaluating AI outputs against responsible AI metrics and stress-testing models for alignment.
  - **Red teaming:** Efforts to proactively identify AI system vulnerabilities and ensure resiliency benefit from augmentation with embedded technology solutions to ensure evergreen testing against evolving risks.
- **Hardwire responsible AI controls into enterprise AI infrastructure and solutions:** This incentivizes fluid adoption, accountability and decreases the likelihood of risks being overlooked. Employee upskilling initiatives to make use of responsible AI technologies within workflows may be needed (see Play 9) alongside upgrading legacy systems to a modern digital core. This includes integrating advanced data and AI management tools that support seamless and secure data and AI connectivity across the enterprise.<sup>88</sup>
- **Maintain sufficient human oversight:** To ensure accountability and offset limitations with AI, e.g. hallucinations and reasoning gaps or overreliance on AI outputs. The mandates and cadence of human oversight must adapt to increasingly autonomous and complex agentic AI systems and their potential for unintended consequences. There is an emerging market of platforms that help automate key steps, including AI system registration, risk assessment, requirements assignment and compliance sign-off, while supporting human oversight.

## CASE STUDY 9

### Reinventing AI governance with Accenture's Trusted Agent Huddle

Accenture, a global professional services company, has been reimagining its marketing operations by integrating responsible agentic AI directly into its cloud-based AI Refinery platform.<sup>89</sup> To address increasing demand for faster, smarter campaigns, the organization brought together multiple autonomous agents to streamline traditional marketing processes, cutting planning phase steps by 67% and accelerating time to first draft by 90%. A recent feature, called the Trusted Agent Huddle,<sup>90</sup> has been introduced to facilitate secure and observable agentic collaboration across important ecosystem partners like Writer, Adobe and Salesforce. This is intended to systematize responsible AI

practices directly into daily workflows, governing how agents interact, share data and make decisions.



#### Key insight

Reinventing work through agentic AI shifts the focus from automation to augmentation, unlocking new levels of creativity, speed and strategic impact. These new ways of working will require responsible AI capabilities – like the Trusted Agent Huddle – to be systematically integrated into workflows to ensure accountable collaboration at scale between humans and AI agents.



## Government leaders

### Key roadblocks organizations encounter from the broader ecosystem

Limited incentives for the implementation of responsible AI technologies, focusing on investments in AI innovation rather than the technologies to embed trust and regulatory compliance

Lack of audit mechanisms for third-party AI tools and systems, undermining risk management efforts and hindering responsibility allocation and governance across the AI ecosystem

Investment uncertainties, due to the lack of established interoperability standards between legacy systems and new technologies, and between AI systems and responsible AI technologies, discouraging long-term investments

### Actions for government leaders

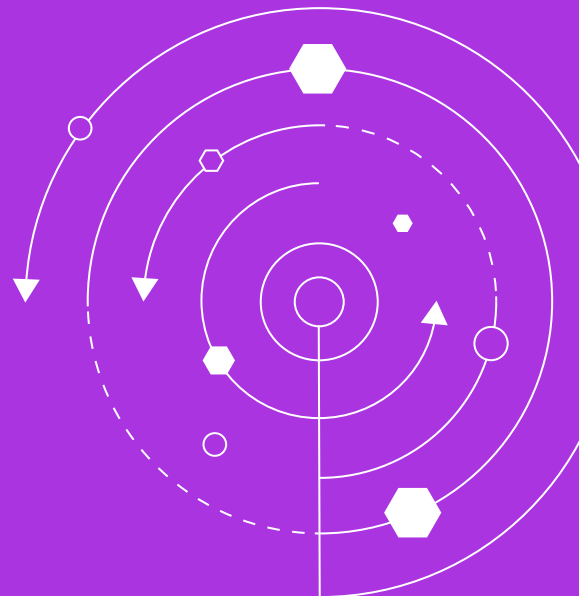
- **Promote R&D of responsible AI technologies:** Motivate a market for responsible AI technologies with signals such as recognition, insurance protection for AI liabilities<sup>91</sup> or minimum design thresholds for AI development (see Play 7).
- **Promote interoperability between responsible AI technologies:** As technology-enabled responsible AI becomes common practice, companies will need common mechanisms to assess each other's approaches. Governments should drive multistakeholder efforts to establish interoperability parameters between partners and upstream and downstream actors. Key components to address include:
  - **Common standards:** Taxonomies, formats and communication protocols for responsible AI metrics and audit data
  - **Interoperable application programming interfaces (APIs):** Shared definitions for bias checks, red-teaming, observability (see Case study 9), explainability, etc.
  - **System-to-system transparency mechanisms:** Traceability, documentation and reporting structures that are comparable across tools
  - **Standardized trust and risk mechanisms:** Dynamic trust assessments between AI agents or systems
  - **Sandboxes:** Environments for safe stress-testing of responsible AI technologies
  - **Multistakeholder governance models:** Collaborations between government, industry, academia and civil society help set norms and resolve cross-border or cross-sector inconsistencies



# Play 9

## Increase responsible AI literacy and workforce transition opportunities

As organizations reinvent themselves around AI, fostering responsible AI literacy and cross-disciplinary skills across the enterprise is critical to prepare for cultural change, capability-building and talent transformation. For governments, investing in AI education is foundational to a public capable of informed decision-making in AI use and to a pipeline that meets increasing business demand for responsible AI experts.



## Organization leaders

### Key roadblocks that arise within the organization

Gap between AI use and risk literacy: 53% of the US population reported using a generative AI tool.<sup>92</sup> However, only 1% were able to correctly answer all questions regarding basic AI literacy,<sup>93</sup> pointing to vulnerabilities for individuals and organizations.

C-suite underestimates worker concerns: 59% of workers express substantial concerns about the impact of generative AI on job security. Yet only 29% of executives assume workers have concerns about job loss.<sup>94</sup> This underestimation can lead to underinvestment in literacy and trust-building approaches.

### Actions for organization leaders

- **Invest in responsible AI literacy across the organization:** Embed literacy regarding AI capabilities, limitations, risks, compliance and ethical considerations into learning and development offerings (see Case study 10). Cross-functional training and change management initiatives are needed to upskill technical and non-technical workers. This can decentralize risk management and enable all employees to be informed users of AI, cognizant of when escalation is needed for support. In the long term, literacy will need to account for evolving AI risks and regulations and for variability in training needs for current AI adopters and an AI-native future workforce.
- **Enhance literacy specificity with defined policies and tooling:** Integrate organizational responsible AI policies and procedures into training programmes. Upskill employees with trainings tailored to approved AI tool use while ensuring transferable skills.
- **Inform leadership decision-making with employee listening:** Define metrics to measure AI adoption, the state of responsible AI practices and trust across the workforce. Then, use those insights to inform workforce transition initiatives. Literacy strategies must reflect and address the employees' various AI concerns that could hinder AI adoption and responsible use (see Table 3). Establish avenues for ongoing employee input to refine trainings (see Play 1).



TABLE 3 | Literacy approaches tailored to different types of employee concerns

| Employee concern                                               | Tailored literacy approaches                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|----------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Don't see relevance of AI tools</b>                         | <ul style="list-style-type: none"> <li>– Ask employees to share examples of how AI is applied in their day-to-day lives, e.g. robotic vacuums or voice assistants</li> <li>– Showcase how AI has been applied in work contexts and the benefits provided</li> </ul>                                                                                                                                                                                                                       |
| <b>Struggle to use AI tools or integrate into work</b>         | <ul style="list-style-type: none"> <li>– Provide hands-on training</li> <li>– Establish peer-based modalities, where employees can share challenges and successes in using AI</li> </ul>                                                                                                                                                                                                                                                                                                  |
| <b>Don't understand risks, limitations and responsible use</b> | <ul style="list-style-type: none"> <li>– Mandate foundational training to all workers</li> <li>– Create role- and tool-based learning so that workers see how AI risk and limitations intersect with their workflows</li> <li>– Make guidebooks of principles and recommended behaviours readily available</li> </ul>                                                                                                                                                                     |
| <b>Fear that AI will lead to job displacement</b>              | <ul style="list-style-type: none"> <li>– Provide transparency regarding where AI is being deployed to replace or augment worker activities</li> <li>– Invest in trust-building with employees, such as by showing how AI has upskilled, rather than replaced, employees</li> <li>– When communicating AI adoption, use careful language to avoid the impression that the company prioritizes AI over people – or that it anthropomorphizes AI agents as equal to human workers</li> </ul> |
| <b>Fear that AI will lead to increased stress</b>              | <ul style="list-style-type: none"> <li>– Balance increased employee output expectations that AI may afford with employee concerns of quality control, talent decisions and stress that can come with scale</li> </ul>                                                                                                                                                                                                                                                                     |

## CASE STUDY 10

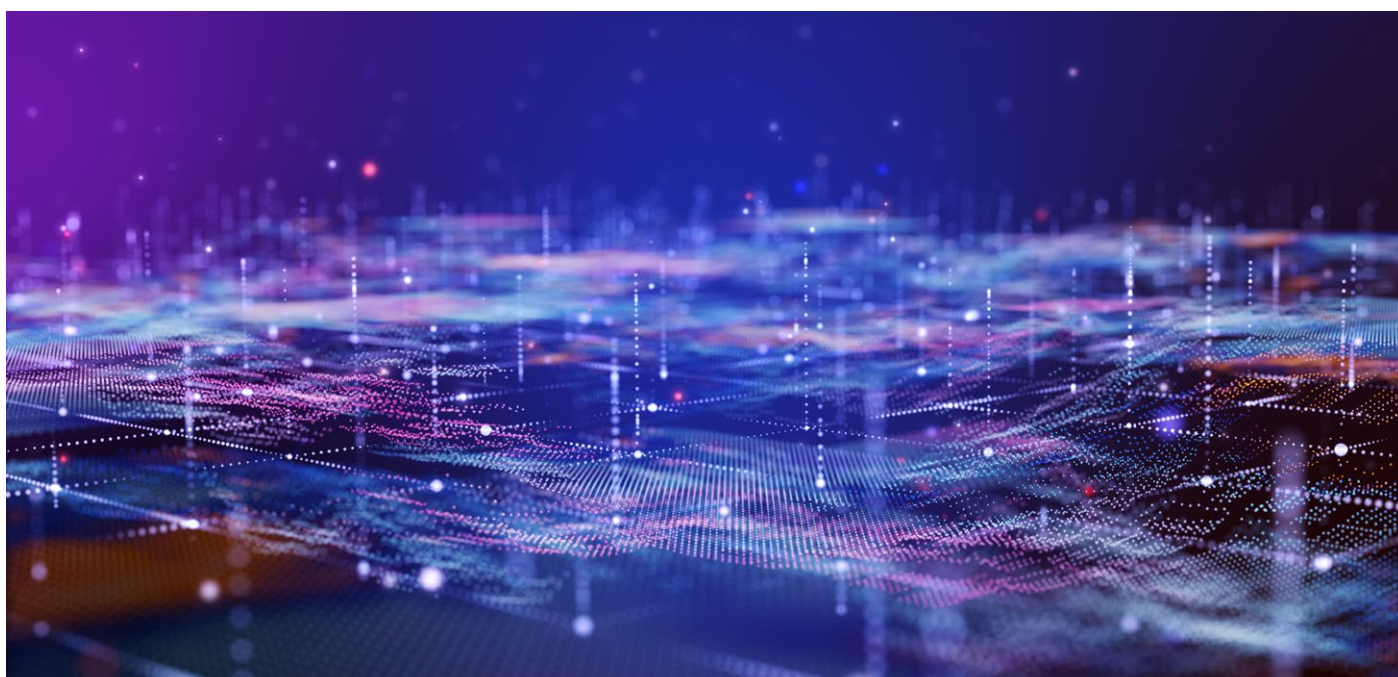
### IKEA's responsible AI literacy programme

Faced with AI's growing impact on retail operations, IKEA recognized an urgent need to equip its global workforce with the skills to interact with AI responsibly. They launched a global AI literacy initiative tailored to employee roles. The programme combines foundational AI knowledge with modules on responsible AI and ethics training. In the programme's first year, over 4,000 employees were trained, with plans to reach 70,000 by 2026 and a company-wide rollout by 2027.<sup>95</sup>



#### Key insight

Organizations should treat responsible AI literacy as a long-term, organization-wide commitment by embedding it into workforce development strategies, encouraging experimentation and integrating change management to support cultural and operational shifts.





### Key roadblocks organizations encounter from the broader ecosystem

The evolving AI literacy gap, caused by the rapid advancement of generative AI, is outpacing standardized training frameworks. It is also creating a critical need for adaptive, resource-efficient literacy programmes that can keep workforce development aligned with cutting-edge AI capabilities while ensuring responsible implementation across diverse organizational contexts.

The skilled workforce gap, caused by a shortage of cross-disciplinary talent necessary for responsible AI, is an issue heightened in regions or sectors with limited access to relevant training or talent.

### Actions for government leaders

- **Promote alignment on responsible AI literacy foundations:** Provide companies with clarity on what constitutes a literacy baseline, such as by defining standards, and support efforts to align and document literacy foundations with experts across sectors.
  - **Create access to literacy and a pipeline of experts:** Enable responsible AI literacy across the general population while supporting specialized technical and socio-technical roles in responsible AI. Support lifelong learning programmes and PPPs and address the unique literacy and access challenges that institutions, researchers and educators face in academia.<sup>96</sup> Jurisdictional approaches reveal various actions to embed responsible AI appreciation into education – from elementary through professional levels – ensuring learners can use AI and understand its societal impact. For example:
    - The European Commission and OECD's AI Literacy Framework (AILit) emphasizes ethical reasoning, creativity and digital responsibility.
    - Rwanda's National AI Policy delineates a multi-year implementation plan for AI literacy.
    - AI Singapore (AISG) embeds responsible AI modules into its AI apprenticeship programme and mid-career training.
    - China's national education guidelines promote integration of ethical AI training and design thinking into a comprehensive digital curriculum.
    - Malaysia's AI Untuk Rakyat (AI for the People) initiative is a self-learning online programme aimed at demystifying AI for individuals across ages and occupations.
    - AI4K12 in the US introduces computational thinking and responsible prompting techniques to develop cognitive and noncognitive skills in young learners.
- Long-term planning is needed to account for curriculum reform, a lag between policy and workforce preparedness, and to ensure that AI literacy is inclusive and accessible, distributed across society.

# Conclusion

Organizations developing or adopting AI have an outsized role in ensuring a trusted ecosystem for sustainable AI innovation and positive impact. Shareholders, policy-makers and consumers alike seek confidence that organizations are implementing robust responsible AI practices. This playbook provides a series of actions to overcome internal and ecosystem roadblocks that organizations often encounter when implementing responsible AI.

Looking ahead, end-to-end responsible AI ensures organizations have the necessary foundations to unlock opportunities in evolving and agentic AI capabilities, where trusted adoption depends on robust governance. Governments, for their part, will need to continue to encourage a context where industry, academia, civil society and the public contribute to a holistic, trustworthy AI ecosystem.

# Contributors

The authors sincerely thank those who contributed their insights via interviews and workshops, as well as contributors not named below. The opinions expressed herein do not necessarily reflect the views of the individuals or organizations involved in the project listed below.

## Lead authors

### **Rafi Lazerson**

Associate Manager, Responsible AI,  
Accenture; Project Fellow,  
World Economic Forum

### **Manal Siddiqui**

Senior Manager, Responsible AI – Canada Lead,  
Accenture; Project Fellow, World Economic Forum

### **Karla Yee Amezaga**

Initiatives Lead, AI and Data Governance,  
World Economic Forum

## World Economic Forum

### **Daniel Dobrygowski**

Head, Governance and Trust,  
Centre for AI Excellence

### **Audrey Duet**

Head, Data and AI Innovation,  
Centre for AI Excellence

### **Casey Price**

Specialist, AI/Data Governance and Community,  
Centre for AI Excellence

## Accenture

### **Patrick Connolly**

Responsible AI Research Manager; Project Fellow,  
World Economic Forum

### **Kathryn White Krumpholz**

Managing Director, Innovation Incubation

### **Andrew J.P. Levy**

Chief Corporate and Government Affairs Officer

### **Valerie Morignat**

Senior Manager, Responsible AI

### **Charlie Moskowitz**

Government Relations Senior Manager;  
Project Fellow, World Economic Forum

### **Ali Shah**

Managing Director, Responsible AI

### **Dikshita Venkatesh**

Offering Development and Innovation Specialist;  
Project Fellow, World Economic Forum

## Working group members

### **Lovisa Afzelius**

Chief Executive Officer, Apriori Bio

### **Suhani Akhare**

General Counsel and Corporate Secretary,  
Automation Anywhere

### **Hoda Al Khzaimi**

Associate Vice-Provost for Research Translation and  
Entrepreneurship, New York University Abu Dhabi

### **Abdulaziz Al-Ali**

Head, Centre for the Fourth Industrial  
Revolution Qatar

### **Hassan Aldarbesti**

Executive Director of International Affairs  
Department, Ministry of Information and  
Communication Technology (ICT) of Qatar

### **Uthman Ali**

Associate Fellow, University of Oxford,  
Saïd Business School

### **Norberto Andrade**

Professor and Academic Director, IE University

### **Amir Banifatemi**

Chief Responsible AI Officer,  
Cognizant Technology Solutions US

**Richard Benjamins**

Co-Founder and Chief Executive Officer,  
RAIGHT.ai

**Seth Bergeson**

Senior Manager, AI and Responsible AI, PwC

**Jakub Bielamowicz**

Associate Director, Governmental and Regulatory  
Affairs, Group Compliance, Regulatory and  
Governance, UBS Business Solutions

**Saqr Binghalib**

Executive Director, Artificial Intelligence,  
Digital Economy and Remote Work Applications  
Office of the United Arab Emirates

**Anu Bradford**

Professor, Law, Columbia Law School

**Sandra Braman**

Professor and Senior Scholar, Quello Center,  
Michigan State University

**Melika Carroll**

Head, Global Government Affairs and Public Policy,  
Cohere

**Winter Casey**

Senior Director, Government Affairs, SAP

**Daniel Castano Parra**

Professor, Law, Universidad Externado de Colombia

**Kaan Cetinturk**

Global Chief Information Office and Director of  
Information and Communications Technology  
Division, United Nations Children's Fund (UNICEF)

**Neil Chase**

Chief Executive Officer, CalMatters and The Markup

**Neha Chawla**

Principal Legal Counsel, Infosys

**Simon Chesterman**

Senior Director, AI Governance, AI Singapore,  
National University of Singapore

**Quintin Chou-Lambert**

Office of the UN Tech Envoy, United Nations

**Frincy Clement**

Head, North America Region, Women in AI

**Magda Cocco**

Head, Practice Partner, Information, Communication  
and Technology, Vieira de Almeida and Associates

**Ben Colman**

Co-Founder and Chief Executive Officer,  
Reality Defender

**Amanda Craig**

Senior Director, Responsible AI Public Policy,  
Microsoft

**Renée Cummings**

Data Science Professor and Data Activist  
in Residence, University of Virginia

**Nicholas Dirks**

President and Chief Executive Officer,  
The New York Academy of Sciences

**Eileen Donahoe**

Former US Special Envoy for Digital Freedom  
2023-2024, US Department of State

**Daniel Duke Odongo**

Director, Product; Member of the Executive Team,  
Ushahidi

**Jeff Easley**

General Manager, Responsible Artificial  
Intelligence Institute

**Mark Esposito**

Faculty Affiliate, Harvard Center for International  
Development, Harvard Kennedy School and  
Institute for Quantitative Social Sciences

**Nita Farahany**

Robinson O. Everett Professor of Law and  
Philosophy and Director, Duke Science and Society,  
Duke University

**Max Fenkell**

Vice-President, Government Relations, Scale AI

**Kay Firth-Butterfield**

Chief Executive Officer, Good Tech Advisory

**Katharina Frey**

Deputy Head, Digitalization Division, Federal  
Department of Foreign Affairs (FDFA) of Switzerland

**Huanzhang Fu**

Assistant Director, Strategic Innovation, Innovation  
Centre, International Criminal Police Organization  
(INTERPOL)

**Urs Gasser**

Dean, TUM School of Social Sciences and  
Technology, Technical University of Munich

**Justine Gauthier**

General Counsel; Head, AI Governance,  
MILA – Quebec Artificial Intelligence Institute

**Debjani Ghosh**

President, National Association of Software  
and Services Companies (NASSCOM)

**Danielle Gilliam-Moore**

Director, Global Public Policy, Salesforce

**Nabeel Goheer**

Chief of Asia, Middle East, Europe Region,  
PATH UK

**Mark Gorenberg**

Chair, MIT Corporation, Massachusetts Institute of Technology (MIT)

**Brian Patrick Green**

Director, Technology Ethics, Markkula Center for Applied Ethics, Santa Clara University

**Sam Gregory**

Executive Director, WITNESS

**Koiti Hasida**

Director, Artificial Intelligence in Society Research Group, RIKEN Center for Advanced Intelligence Project, RIKEN

**Dan Hendrycks**

Executive Director, Center for AI Safety

**Yanqing Hong**

Professor, Law School, Beijing Institute of Technology

**Marek Jansen**

Senior Director, Strategic Partnerships and Policy Management, Volkswagen

**Jeff Jianfeng Cao**

Senior Research Fellow, Tencent Research Institute

**Sam Kaplan**

Senior Director and Assistant General Counsel, Palo Alto Networks

**Kathryn King**

General Manager, Technology and Strategy, Office of the eSafety Commissioner Australia

**Edward S. Knight**

Executive Vice-Chairman, Nasdaq

**Agnes Li**

Head of Global AI Legal, Bytedance

**Caroline Louveaux**

Chief Privacy, AI & Data Responsibility Officer, Mastercard

**Shawn Maher**

Global Vice-Chair, Public Policy, EY

**Gevorg Mantashyan**

First Deputy Minister, High-Tech Industry, Ministry of High-Tech Industry of Armenia

**Gary Marcus**

Chief Executive Officer, Center for Advancement of Trustworthy AI

**Brittany Masalosalo**

Global Head of Government Relations, HP

**Mike Mattacola**

Chief Business Officer, CoreWeave

**Gregg Melinson**

Senior Vice-President, Corporate Affairs, Hewlett Packard Enterprise

**Robert Middlehurst**

Senior Vice-President, Regulatory Affairs, e& International

**Satwik Mishra**

Executive Director, Centre for Trustworthy Technology, Centre for the Fourth Industrial Revolution

**Nick Moës**

Executive Director, The Future Society

**Chandler Morse**

Vice-President, Corporate Affairs, Workday

**Henry Murry**

Vice-President, Government Relations, C3 AI

**Miho Naganuma**

Senior Executive Professional, Digital Trust Business Strategy Department, NEC Corporation

**Dan Nechita**

EU Director, Transatlantic Policy Network (TPN)

**Jessica Newman**

Director, AI Security Initiative, Centre for Long-Term Cybersecurity, UC Berkeley

**Michael Nunes**

Vice-President, Payments Policy, Visa

**Adele O'Herlihy**

Deputy General Counsel, Inception

**Florian Ostmann**

Head, AI Governance and Regulatory Innovation, The Alan Turing Institute

**Tiffany Pham**

Founder and Chief Executive Officer, Mogul

**Melisha Pillay**

Risk Executive, Old Mutual

**Oreste Pollicino**

Professor, Constitutional Law, Bocconi University

**Catherine Quinlan**

Vice-President, AI Ethics, IBM

**Roxana Radu**

Associate Professor of Digital Technologies and Public Policy, Blavatnik School of Government; Hugh Price Fellow, Jesus College University of Oxford

**Jayesh Ranjan**

CEO, Industry and Investment Cell; CEO, SPEED; Special Chief Secretary, Youth Advancement, Tourism and Culture, Government of Telangana

**Martin Rauchbauer**

Co-Director and Founder, Tech Diplomacy Network

**Alexandra Reeve Givens**

Chief Executive Officer, Center for Democracy and Technology

**Philip Reiner**

Chief Executive Officer, Institute for Security and Technology

**Andrea Renda**

Senior Research Fellow, Centre for European Policy Studies (CEPS)

**Tom Renwick**

Group Chief Legal Officer, Astra Tech

**Rowan Reynolds**

General Counsel, WRITER

**Andres Rojas Corvalan**

Director, Applied AI Projects, Industry Innovation, Vector Institute

**Nilmini Rubin**

Chief Policy Officer, Hedera Hashgraph

**Arianna Rufini**

ICT Advisor to the Minister, Ministry of Enterprises and Made in Italy

**Crystal Rugege**

Managing Director, Centre for the Fourth Industrial Revolution Rwanda

**Joaquina Salado**

Head of AI Ethics, Telefónica

**Idoia Ana Salazar Garcia**

Professor, Universidad CEU San Pablo; President, OdiselA

**Nayat Sanchez-Pi**

Chief Executive Officer, INRIA Chile

**Roland Scharrer**

Fellow and Visiting Research Scholar, Stanford University

**Thomas Schneider**

Ambassador and Director of International Affairs, Swiss Federal Office of Communications, Federal Department of the Environment, Transport, Energy and Communications (DETEC)

**Robyn Scott**

Co-Founder and Chief Executive Officer, Apolitical

**Var Shankar**

Lecturer, Governance and Responsible AI Lab (GRAIL Lab), Purdue University

**Uyi Stewart**

Vice-President, Inclusive Innovation and Analytics, Mastercard

**Charlotte Stix**

Head, AI Governance, Apollo Research

**Arun Sundararajan**

Harold Price Professor, Entrepreneurship and Technology, Stern School of Business, New York University

**Nabiha Syed**

Executive Director, Mozilla Foundation

**Patricia Thaine**

Co-Founder and Chief Executive Officer, Private AI

**Anne-Lise Thieblemont**

Vice-President, Government Affairs, Qualcomm

**Ulrike Till**

Director, IP and Frontier Technologies Division, World Intellectual Property Organization (WIPO)

**Peter Tucker**

Head of Commercial Legal, Twelve Labs

**Anna Tumadóttir**

Chief Executive Officer, Creative Commons

**Ott Velsberg**

Government Chief Data Officer, Ministry of Economic Affairs and Information Technology of Estonia

**Miriam Vogel**

President and Chief Executive Officer, Equal AI

**Takuya Watanabe**

Director, AI Industry Strategy Office, Ministry of Economy, Trade and Industry Japan

**Andrew Wells**

Chief Data and AI Officer, NTT DATA

**Denise Wong**

Assistant Chief Executive, Data Innovation and Protection Group, Info-communications Media Development Authority of Singapore

**Kai Zenner**

Head, Office and Digital Policy Adviser, MEP Axel Voss, European Parliament

**Arif Zeynalov**

Transformation Chief Information Officer, Ministry of Economy of the Republic of Azerbaijan

**Jason Zhou**

Senior Research Manager, Concordia AI

# Additional acknowledgements

**Jesse Barba**

Head of Federal Government Relations,  
Cengage Group

**Daniela Braga**

Founder and Chief Executive Officer, Defined.ai

**Nic Chavez**

Chief Information Security Officer, IBM

**Saad Farooq**

Director, Digital Relations, e&

**Evi Fuelle**

Global AI Policy Lead, Scale AI

**Tony Gaffney**

Special Advisor to Chair and Former CEO,  
Vector Institute

**Eugenio Garcia**

Director, Department of Science, Technology,  
Innovation and Intellectual Property (DCT),  
Brazilian Ministry of Foreign Affairs

**Katharina Koerner**

Senior Principal Consultant, AI Governance,  
Trace3, Responsible AI and Privacy,  
Tech Diplomacy Network

**Anthony Giuliani**

Global Head of Operations, Twelve Labs

**Alexis Liu**

Head, Legal, Weights and Biases

**Irina Michalowicz**

Senior Director, Policy and Government Affairs,  
Palo Alto Networks

**Casey Mock**

Senior Policy Director, The Anxious Generation

**Bo Viktor Nylund**

Director, UNICEF Innocenti Global Office  
of Research and Foresight, United Nations  
Children's Fund (UNICEF)

**Madan Oberoi**

Chief Security Advisor and Lead for Intelligence  
and Public Safety, APAC, Google Cloud

**Marc-Etienne Ouimette**

Founder, Cardinal Policy

**Timothy Persons**

Principal, Digital Assurance and Transparency of US  
Trust Solutions, PwC

**Halak Shrivastava**

Global AI Public Policy and Regulatory Affairs, Cohere

**Scott Starbird**

Chief Public Affairs Officer, Databricks

**V Valluvan Veloo**

Director, Manufacturing Industry, Science and  
Technology Division, Ministry of Economy, Malaysia

**Roy Wang**

Chief Legal Officer, Eightfold AI

## Accenture

**Chiara Addis**

Research Specialist

**Laetitia Cailleteau**

Managing Director; Europe, Middle East and Africa  
Responsible AI Lead

**Leah Chauhan**

Manager, Responsible AI; UK Government Lead

**Melanie Condron**

Compliance and Ethics Counsel Associate Director

**Shivansh Gupta**

Research Specialist

**Harvin Hilaire**

Management Consulting Analyst

**Ted Kwartler**

Managing Director, Americas Responsible AI Lead

**Ian Lever**

Manager, Management Consulting

**Thomas Niven**

Managing Director, Responsible AI

**Mimi Whitehouse**

Senior Manager, Responsible AI, Health and  
Public Service

## World Economic Forum

**Muneera Alshunaifi**

Seconded Fellow 2024-2025, Centre for the Fourth  
Industrial Revolution Saudi Arabia

**Genta Ando**

Executive Director, Japan External Trade  
Organization; Project Fellow 2023-2025

**Maria Basso**

Head, AI Applications and Impact

**Teysir Bedretdin**

Specialist, AI Governance and  
International Collaboration

**Agustina Callegari**

Initiatives Lead, Technology Governance,  
Safety and International Cooperation

**Samira Gazzane**

Initiatives Lead, AI Competitiveness  
for Regional Collaboration

**Karyn Gorman**

Lead, Communications and Marketing,  
Centre for AI Excellence

**Benjamin Larsen**

Initiatives Lead, Safe Systems and Technologies,  
Centre for AI Excellence

**Hannah Rosenfeld**

Former Specialist, Artificial Intelligence  
and Machine Learning

**Harsh Sharma**

Lead, AI and ML, Centre for the Fourth Industrial  
Revolution India

**Stephanie Smittkamp**

Coordinator, Centre for AI Excellence

**Stephanie Teeuwen**

Former Specialist, Data and AI

## Production

**Laurence Denmark**

Creative Director, Studio Miko

**Will Liley**

Editor, Studio Miko

**Cat Slaymaker**

Designer, Studio Miko

# Endnotes

1. World Economic Forum. (2024). *AI Governance Alliance: Briefing Paper Series*. <https://www.weforum.org/publications/ai-governance-alliance-briefing-paper-series/>.
2. Maslej, N., L. Fattorini, R. Perrault, Y. Gil, et al. (2025). *The AI Index 2025 Annual Report*. Institute for Human-Centered AI, Stanford University. [https://hai.stanford.edu/assets/files/hai\\_ai\\_index\\_report\\_2025.pdf](https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf). See McKinsey & Company Survey, 2024. Figure 3.3.6.
3. Based on research conducted by Accenture and the Institute for Human-Centered AI at Stanford University – academic paper is forthcoming. See Figure 1 and Endnote 4.
4. For 2024 survey data in figure 1 see Accenture. (2024). *Responsible AI: From compliance to confidence*. <https://www.accenture.com/content/dam/accenture/final/accenture-com/document-3/Accenture-Responsible-AI-From-Compliance-To-Confidence-Report.pdf>; 2025 survey data in figures 1, 2 and 3 is based on research conducted by Accenture and the Institute for Human-Centered AI at Stanford University – academic paper is forthcoming. Note: companies with no AI adoption and implementation strategy were excluded from the analysis. 2025 data for responsible AI maturity has been interpreted according to methods used in Accenture's aforementioned 2024 report: Accenture. (2024). *Responsible AI: From compliance to confidence*.
5. Accenture and Amazon Web Services (AWS). (2024). *Thrive with responsible AI: How embedding trust can unlock value*. <https://www.accenture.com/content/dam/accenture/final/accenture-com/document-3/Accenture-Responsible-AI-From-Risk-Mitigation-to-Value-Creation.pdf>.
6. Microsoft Source. (2023). *AFL-CIO and Microsoft announce new tech-labor partnership on AI and the future of the workforce*. <https://news.microsoft.com/source/2023/12/11/afl-cio-and-microsoft-announce-new-tech-labor-partnership-on-ai-and-the-future-of-the-workforce/>.
7. Dobrygowski, D. and G. Moschos. (2025). *Why workers must be at the centre of the digital trust agenda – and 4 ways to get there*. World Economic Forum. <https://www.weforum.org/stories/2025/05/workers-digital-trust-agenda-4-ways/>.
8. Muñoz Marcos, A., J. M. Bolufer de Francia and J. Salado Morelada. (2024). *We have updated our artificial intelligence principles*. Telefónica. <https://www.telefonica.com/en/communication-room/blog/we-have-updated-our-artificial-intelligence-principles/>.
9. Government of Brazil, Ministry of Science, Technology and Innovation. (12 June 2025). *Publicada versão final do Plano Brasileiro de Inteligência Artificial sob coordenação do MCTI*. <https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/noticias/2025/06/publicada-versao-final-do-plano-brasileiro-de-inteligencia-artificial-sob-coordenacao-do-mcti>
10. The White House. (2025). *Winning the Race: America's AI Action Plan*. <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
11. Ministry of Foreign Affairs, People's Republic of China (2025). *Global AI Governance Action Plan*. [https://www.mfa.gov.cn/eng/xw/zyxw/202507/t20250729\\_11679232.html](https://www.mfa.gov.cn/eng/xw/zyxw/202507/t20250729_11679232.html).
12. Ministry of Science, Innovation, Technology and Communications, Government of Costa Rica. (2025). *National Artificial Intelligence Strategy 2024-2027*. <https://www.micitt.go.cr/sites/default/files/2025-02/National%20Artificial%20Intelligence%20Strategy%20of%20Costa%20Rica.pdf>.
13. Innovation, Science and Economic Development Canada. (2023). *Voluntary Code of Conduct on the Responsible Development and Management of Advanced Generative AI Systems*. <https://ised-isde.canada.ca/site/ised/en/voluntary-code-conduct-responsible-development-and-management-advanced-generative-ai-systems>.
14. Ministry of Communication and Information Technology. (20 April 2023). *Egyptian Charter for Responsible AI Launched*. [Press release]. [https://mciit.gov.eg/en/media\\_center/latest\\_news/news/66939](https://mciit.gov.eg/en/media_center/latest_news/news/66939)
15. Australian Government, Department of Industry, Science and Resources. (5 September 2024). *Voluntary AI Safety Standard*. <https://www.industry.gov.au/publications/voluntary-ai-safety-standard>.
16. Paulger, D. (2025) *Understanding Japan's AI Promotion Act: An "Innovation-First" Blueprint for AI Regulation*. <https://fpf.org/blog/understanding-japans-ai-promotion-act-an-innovation-first-blueprint-for-ai-regulation/>.
17. World Economic Forum. (2024). *Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation*. [https://www3.weforum.org/docs/WEF\\_Governance\\_in\\_the\\_Age\\_of\\_Generative\\_AI\\_2024.pdf](https://www3.weforum.org/docs/WEF_Governance_in_the_Age_of_Generative_AI_2024.pdf).
18. Ibid.
19. Ibid.
20. Dubai Future Foundation. (2025). *Dubai introduces AI Seal to certify trusted AI companies*. <https://www.dubaifuture.ae/uncategorized/dubai-introduces-ai-seal-to-certify-trusted-ai-companies/>.
21. Bonterms. (n.d.). *Bonterms AI Standard Clauses (Version 1.0)*. <https://bonterms.com/forms/ai-standard-clauses-version-1-0/>.
22. NTT Data. (1 February 2024). *Accelerate Snowflake's data sharing ecosystem with data clean rooms*. <https://us.nttdata.com/en/blog/2024/january/accelerate-snowflakes-data-sharing-ecosystem-with-data-clean-rooms>
23. United Nations. (7 August 2025). *Ensuring Indigenous Peoples' rights in the age of AI*. <https://www.un.org/en/desa/ensuring-indigenous-peoples%E2%80%99-rights-age-ai>.
24. UNICEF. (November 2021). *Policy guidance on AI for children 2.0*. <https://www.unicef.org/innocenti/reports/policy-guidance-ai-children>.

25. European Commission. (n.d.). *European Data Governance Act*. <https://digital-strategy.ec.europa.eu/en/policies/data-governance-act>.
26. The White House. (2025). *Winning the Race: America's AI Action Plan*. <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
27. Shumailov, I., Z. Shumaylov, Y. Zhao, N. Papernot, et al. (2024). AI models collapse when trained on recursively generated data. *Nature*, vol. 631, pp. 755-759. <https://www.nature.com/articles/s41586-024-07566-y>.
28. Saudi Arabia Communication, Space, and Technology Commission (CST). (2025). *CST Publishes a Public Consultation for the Global AI Hub Law*. <https://www.cst.gov.sa/en/media-center/news/N2025041401>.
29. World Economic Forum. (2025). *Technology Convergence Report*. <https://www.weforum.org/publications/technology-convergence-report-2025/>.
30. UN Global Pulse. (2022). *Horizon Scan Manual: A Step by Step Guide*. <https://www.unglobalpulse.org/document/horizon-scan-manual-a-step-by-step-guide/>.
31. Hinchey, M. and L. Coyle. (2010). *Evolving critical systems: a research agenda for computer-based systems*. University of Limerick. <https://hdl.handle.net/10344/2085>.
32. Kavanagh, J. (2025). *AI Governance Controls Mega-map*. <https://www.ethos-ai.org/p/ai-governance-controls-mega-map>.
33. World Economic Forum. (2024). *Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation*. [https://www3.weforum.org/docs/WEF\\_Governance\\_in\\_the\\_Age\\_of\\_Generative\\_AI\\_2024.pdf](https://www3.weforum.org/docs/WEF_Governance_in_the_Age_of_Generative_AI_2024.pdf).
34. World Economic Forum. (2024). *AI Governance Alliance: Briefing Paper Series*. <https://www.weforum.org/publications/ai-governance-alliance-briefing-paper-series/>.
35. World Economic Forum. (2024). *Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation*. [https://www3.weforum.org/docs/WEF\\_Governance\\_in\\_the\\_Age\\_of\\_Generative\\_AI\\_2024.pdf](https://www3.weforum.org/docs/WEF_Governance_in_the_Age_of_Generative_AI_2024.pdf).
36. Meta Open Loop. (2021). *AI Impact Assessment: A Policy Prototyping Experiment*. <https://openloop.org/reports/2021/01/ai-impact-assessment-a-policy-prototyping-experiment.pdf>.
37. NIST Trustworthy & Responsible Artificial Intelligence Resource Center (AIRC). (n.d.). *NIST AI Risk Management Framework (NIST AI RMF) and Japan AI Guidelines for Business (AI GfB)*. [https://airc.nist.gov/docs/FINAL\\_Crosswalk1\\_Terminology\\_RMf\\_GfB.pdf](https://airc.nist.gov/docs/FINAL_Crosswalk1_Terminology_RMf_GfB.pdf) and [https://airc.nist.gov/docs/FINAL\\_Crosswalk2\\_Concepts\\_RMf\\_GfB.pdf](https://airc.nist.gov/docs/FINAL_Crosswalk2_Concepts_RMf_GfB.pdf).
38. Ministry of Foreign Affairs, People's Republic of China (2025). *Joint Statement between the People's Republic of China and the Islamic Republic of Pakistan*. [https://www.mfa.gov.cn/eng/xw/zyxw/202502/t20250206\\_11550159.html](https://www.mfa.gov.cn/eng/xw/zyxw/202502/t20250206_11550159.html).
39. UK Financial Conduct Authority. (2025). *FCA allows firms to experiment with AI alongside NVIDIA*. <https://www.fca.org.uk/news/press-releases/fca-allows-firms-experiment-ai-alongside-nvidia>.
40. World Economic Forum. (2025). *Shaping AI Sandbox Ecosystem for the Intelligent Age*. [https://reports.weforum.org/docs/WEF\\_Shaping\\_the\\_AI\\_Sandbox\\_Ecosystem\\_for\\_the\\_Intelligent\\_Age\\_2025.pdf](https://reports.weforum.org/docs/WEF_Shaping_the_AI_Sandbox_Ecosystem_for_the_Intelligent_Age_2025.pdf).
41. Smith, G., N. Luka, M. Osborne, B. Lattimore, et al. (2025). *Responsible Generative AI Use by Product Managers: Recoupling Ethical Principles and Practices*. <https://doi.org/10.48550/arXiv.2501.16531>.
42. Hirsch, D., J. Ott and A. Westover-Munoz. (2024). *Responsible AI Management: Evolving Practice, Growing Value*. <https://iapp.org/resources/article/ohio-state-report-responsible-ai-management/>.
43. World Economic Forum. (2024). *Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation*. [https://www3.weforum.org/docs/WEF\\_Governance\\_in\\_the\\_Age\\_of\\_Generative\\_AI\\_2024.pdf](https://www3.weforum.org/docs/WEF_Governance_in_the_Age_of_Generative_AI_2024.pdf).
44. The White House. (July 2025). *Winning the Race: America's AI Action Plan*. <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
45. Government of Canada. (2025). *AI Strategy for the Federal Public Service 2025-2027*. [https://publications.gc.ca/collections/collection\\_2025/sct-tbs/BT48-55-2025-eng.pdf](https://publications.gc.ca/collections/collection_2025/sct-tbs/BT48-55-2025-eng.pdf).
46. Accenture. (2024). *Responsible AI: From compliance to confidence*. <https://www.accenture.com/content/dam/accenture/final/accenture-com/document-3/Accenture-Responsible-AI-From-Compliance-To-Confidence-Report.pdf>.
47. Telefónica. (17 September 2024). *GSMA launches responsible AI roadmap for telco industry*. [Press release]. <https://www.telefonica.com/en/communication-room/press-room/gsma-launches-responsible-ai-roadmap-telco-industry/>.
48. Miller, C. L. and G. Waters, G. (n.d.). *Risk Management Framework for the Procurement of AI Systems (RMF PAIS 1.0)*. The Center for Inclusive Change. <https://inclusivechange.hubspotpagebuilder.com/rmf-for-ai-procurement>.
49. Radical Ventures. (2023). *Responsible AI Startups (RAIS) Framework*. <https://github.com/radicalventures/RAIS-Framework>.
50. Monetary Authority of Singapore. (2024). *Artificial Intelligence Model Risk Management – Information Paper*. <https://www.mas.gov.sg/-/media/mas-media-library/publications/monographs-or-information-paper/imd/2024/information-paper-on-ai-risk-management-final.pdf>.
51. NIST Trustworthy & Responsible Artificial Intelligence Resource Center. (n.d.). *Using the AI Risk Management Framework: Workday*. <https://airc.nist.gov/docs/workday-success-story.pdf>.

52. Meta Open Loop. (2024). *Generative AI Risk Management and the NIST Generative AI Profile (NIST AI 660-1)*. <https://openloop.org/reports/2024/10/report-2-nist-generative-ai-profile.pdf>.
53. The White House. (2025). *Winning the Race: America's AI Action Plan*. <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
54. Organisation for Economic Co-operation and Development (OECD). (n.d.). *Catalogue of Tools & Metrics for Trustworthy AI*. <https://oecd.ai/en/catalogue/overview>.
55. Standards Council of Canada. (2025). *Artificial Intelligence and Data Governance Standardization Hub*. <https://ai-standards-normes-ia.ca/en/home>.
56. World Bank Group. (n.d.). *Event Recording: GovTech and Public Sector Innovation Global Forum*. <https://www.worldbank.org/en/events/2024/12/19/govtech-and-public-sector-innovation-global-forum#5>.
57. World Economic Forum. (2024). *Generative AI governance: Shaping a collective global future*. AI Governance Alliance: Briefing Paper Series. <https://www.weforum.org/publications/ai-governance-alliance-briefing-paper-series/>.
58. Reuel, A., P. Connolly, K. F. Meimandi, S. Tewari, et al. (2025). *Responsible AI in the Global Context: Maturity Model and Survey*. Stanford University. <https://arxiv.org/pdf/2410.09985>.
59. MIT AI Risk Repository. (n.d.). *AI Incident Tracker*. <https://airisk.mit.edu/ai-incident-tracker>.
60. Hulagadri, A.V., J. Kreutzer, J. G. Ngui and X. B. Yong. (2025). *Towards fair and comprehensive multilingual LLM benchmarking*. <https://cohere.com/blog/towards-fair-and-comprehensive-multilingual-and-multicultural-llm-benchmarking>.
61. Luccioni, S., B. Hamazaychikov, T. A. de Costa and E. Strubell. (2025). *Misinformation by Omission: The Need for More Environmental Transparency in AI*. <https://arxiv.org/pdf/2506.15572>.
62. 5Rights Foundation. (2025). *Children and AI Design Code*. [https://5rightsfoundation.com/wp-content/uploads/2025/03/5rights\\_AI\\_CODE\\_DIGITAL.pdf](https://5rightsfoundation.com/wp-content/uploads/2025/03/5rights_AI_CODE_DIGITAL.pdf)
63. EU Artificial Intelligence Act. (2025). *Article 55: Obligations for Providers of General-Purpose AI Models with Systemic Risk*. <https://artificialintelligenceact.eu/article/55/>.
64. Singapore Infocomm Media Development Authority & AI Verify Foundation. (2025). *Global AI Assurance Pilot, Annex A*. [annex-a-global-ai-assurance-pilot.pdf](https://www.aivertify.com/annex-a-global-ai-assurance-pilot.pdf)
65. Luccioni, S., B. Hamazaychikov, T. A. de Costa and E. Strubell. (2025). *Misinformation by Omission: The Need for More Environmental Transparency in AI*. <https://arxiv.org/pdf/2506.15572>.
66. Perplexity. (n.d.). *Perplexity Acceptable Use Policy*. <https://www.perplexity.ai/hub/legal/aup>.
67. TechBetter. (2024). *Evaluating AI Governance: Insights from Public Disclosures*. [https://www.ravitdotan.com/files/ugd/f83391\\_b853450bcc274e9ba9454d618ee41a94.pdf](https://www.ravitdotan.com/files/ugd/f83391_b853450bcc274e9ba9454d618ee41a94.pdf).
68. Microsoft. (2025). *2025 Responsible AI Transparency Report*. <https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/msc/documents/presentations/CSR/Responsible-AI-Transparency-Report-2025-vertical.pdf>.
69. Cohere. (n.d.). *Command R and Command R+ Model Card*. <https://docs.cohere.com/docs/responsible-use>.
70. Anthropic. (13 January 2025). *Anthropic achieves ISO 42001 certification for responsible AI*. <https://www.anthropic.com/news/anthropic-achieves-iso-42001-certification-for-responsible-ai>.
71. United Arab Emirates. (n.d.). *Regulatory sandboxes in the UAE*. <https://u.ae/en/about-the-uae/digital-uae/%20regulatory-framework/regulatory-sandboxes-in-the-uae>.
72. European Commission. (2023). *Hiroshima Process International Code of Conduct for Advanced AI Systems*. <https://digital-strategy.ec.europa.eu/en/library/hiroshima-process-international-code-conduct-advanced-ai-systems>.
73. Hiroshima AI Process (n.d.). *Supporters*. <https://www.soumu.go.jp/hiroshimaaiprocess/en/supporters.html>.
74. World Economic Forum. (2022). *Earning Digital Trust: Decision-Making for Trustworthy Technologies*. <https://www.weforum.org/publications/earning-digital-trust-decision-making-for-trustworthy-technologies/>.
75. World Economic Forum. (2024). *Digital Trust: Supporting Individual Agency*. <https://www.weforum.org/publications/digital-trust-supporting-individual-agency/>.
76. Zhou, L., V. Prabhakaran, R. Ramasubramanian, R. Levin, et al. (2007). *Graceful degradation via versions: specifications and implementations*. *Symposium on Principles of Distributed Computing*. <https://www.microsoft.com/en-us/research/publication/graceful-degradation-via-versions-specifications-and-implementations/>.
77. Internet Matters. (2025). *Me, myself and AI: Understanding and safeguarding children's use of AI chatbots*. <https://www.internetmatters.org/wp-content/uploads/2025/07/Me-Myself-AI-Report.pdf>.
78. Opet, P. (n.d.). *An open letter to third-party suppliers*. J.P. Morgan. <https://www.jpmorgan.com/technology/technology-blog/open-letter-to-our-suppliers>.
79. Stanford Institute for Human-Centered Artificial Intelligence (HAI). (2025). *A framework to report AI's flaws*. <https://hai.stanford.edu/news/a-framework-to-report-ais-flaws>.
80. The Alan Turing Institute. (n.d.). *Understanding the Impacts of Generative AI Use on Children*. [https://www.turing.ac.uk/sites/default/files/2025-05/combined\\_briefing\\_-\\_understanding\\_the\\_impacts\\_of\\_generative\\_ai\\_use\\_on\\_children.pdf](https://www.turing.ac.uk/sites/default/files/2025-05/combined_briefing_-_understanding_the_impacts_of_generative_ai_use_on_children.pdf).

81. Eliot, L. (14 May 2025). *HBR's Top 10 Uses Of AI Puts Therapy And Companionship At The No. 1 Spot*. Forbes. <https://www.forbes.com/sites/lanceeliot/2025/05/14/top-ten-uses-of-ai-puts-therapy-and-companionship-at-the-1-spot/>.
82. Chen, B. J. and J. Metcalf. (28 May 2024). *A sociotechnical approach to AI policy*. Data & Society. <https://datasociety.net/library/a-sociotechnical-approach-to-ai-policy/>.
83. Lazar, S. and A. Nelson. (2023). AI safety on whose terms? *Science*, vol. 381, no. 6654. <https://www.science.org/doi/10.1126/science.adi8982>.
84. UNICEF. (November 2021). *Policy guidance on AI for children 2.0*. <https://www.unicef.org/innocenti/reports/policy-guidance-ai-children>.
85. Firmino, S., and S. Vosloo. (9 July 2025). *The risky new world of tech's friendliest bots*. United Nations Children's Fund (UNICEF). <https://www.unicef.org/innocenti/stories/risky-new-world-techs-friendliest-bots>
86. Corzo, C. (28 August 2024). *Seniors welcome help from robot companions but concerns remain, FIU Business research finds*. Florida International University Business. <https://business.fiu.edu/news/2024/seniors-welcome-help-from-robot-companions-but-concerns-remain.html>.
87. Accenture. (2025). *State of Cybersecurity Resilience 2025*. <https://www.accenture.com/content/dam/accenture/final/accenture-com/document-3/State-of-Cybersecurity-report.pdf>.
88. Narain, K., P. Daugherty, K. Schelfaut, D. Wood, et al. (2024). *Reinventing with a Digital Core – Chapter 1: How to accelerate growth through change*. Accenture. <https://www.accenture.com/us-en/insights/technology/reinventing-digital-core>.
89. Kramer, J., K. Close and L. Guan. (2025). *Designing a new agentic collaborative workforce*. Accenture. <https://www.accenture.com/us-en/blogs/data-ai/designing-new-agentic-collaborative-workforce>.
90. Rosenbush, S. (2025). *AI Agents Are Learning How to Collaborate. Companies Need to Work with Them*. Wall Street Journal. <https://www.wsj.com/articles/ai-agents-are-learning-how-to-collaborate-companies-need-to-work-with-them-28c7464d>.
91. Stern, A. D., A. Goldfarb, T. Minssen and W. P. Nicholson. (2022). AI Insurance: How Liability Insurance Can Drive the Responsible Adoption of Artificial Intelligence in Health Care. *New England Journal of Medicine*, vol. 3, no. 4. <https://doi.org/10.1056/CAT.21.0242>.
92. Ognyanova, K. and V. Singh. (2024). *Public awareness and use of AI tools*. National AI Opinion Monitor. Rutgers University. <https://naiom.net/public-reports/NAIOM%20Report%2001%20AI%20Use%202024.pdf>.
93. Ognyanova, K. and V. Singh. (2025). *AI trust and knowledge in America*. National AI Opinion Monitor. Rutgers University. <https://naiom.net/public-reports/NAIOM%20Report%2002%20AI%20Trust%20Knowledge.pdf>.
94. Accenture. (2024). *Work, workforce, workers: Reinvented in the age of generative AI*. <https://www.accenture.com/content/dam/accenture/final/accenture-com/document-2/Accenture-Work-Can-Become-Era-Generative-AI.pdf>.
95. Ingka Group. (2025). *How IKEA is navigating AI literacy in a world without instructions*. <https://www.ingka.com/newsroom/no-manuals-available-how-ikea-is-navigating-ai-literacy-in-a-world-without-instructions/>.
96. World Economic Forum. (2024). *Governance in the Age of Generative AI: A 360° Approach for Resilient Policy and Regulation*. [https://www3.weforum.org/docs/WEF\\_Governance\\_in\\_the\\_Age\\_of\\_Generative\\_AI\\_2024.pdf](https://www3.weforum.org/docs/WEF_Governance_in_the_Age_of_Generative_AI_2024.pdf).



---

COMMITTED TO  
IMPROVING THE STATE  
OF THE WORLD

---

The World Economic Forum, committed to improving the state of the world, is the International Organization for Public-Private Cooperation.

The Forum engages the foremost political, business and other leaders of society to shape global, regional and industry agendas.

---

World Economic Forum  
91–93 route de la Capite  
CH-1223 Cologny/Geneva  
Switzerland

Tel.: +41 (0) 22 869 1212  
Fax: +41 (0) 22 786 2744  
[contact@weforum.org](mailto:contact@weforum.org)  
[www.weforum.org](http://www.weforum.org)