



RÉUNION DES MINISTRES DE
L'INDUSTRIE, DU NUMÉRIQUE
ET DE LA TECHNOLOGIE

MONTREAL

INDUSTRY, DIGITAL AND
TECHNOLOGY MINISTERS' MEETING

TOOLKIT FOR SMALL AND MEDIUM-SIZED ENTERPRISES (SMEs) DEPLOYING ARTIFICIAL INTELLIGENCE (AI)



Government
of Canada

Gouvernement
du Canada

Canada

Introduction

This toolkit builds upon the commitments made by G7 Leaders at the 2025 Kananaskis Summit under Canada's G7 Presidency. It is designed to support small and medium-sized enterprises (SMEs) —including micro-enterprises— in deploying secure, responsible and trustworthy artificial intelligence (AI), in alignment with the Hiroshima AI Process (HAIP) International Guiding Principles for Organizations Developing Advanced AI Systems.

The toolkit incorporates insights and recommendations from the G7 virtual workshop, “Advancing Trustworthy AI Adoption: Leveraging the Outcomes of the G7 Hiroshima AI Process,” held on July 24, 2025, along with feedback from a follow-up survey and contributions from G7 members and key stakeholders, including the International Center of Expertise in Montreal on Artificial Intelligence (CEIMIA) and the Organisation for Economic Co-operation and Development (OECD). The workshop and survey attracted more than 260 participants from 26 countries, representing governments, industries of all sizes, non-profit organizations, and academia.

Across the G7, SMEs are typically defined by number of employees; for the purposes of this workshop and toolkit, SMEs are categorized as enterprises with fewer than 250 employees ([OECD](#)). According to the [OECD](#), SMEs represent 99% of all businesses, employ two out of three workers, and account for 50% to 60% of value added across OECD member states, making their digital transformation essential for economic stability and

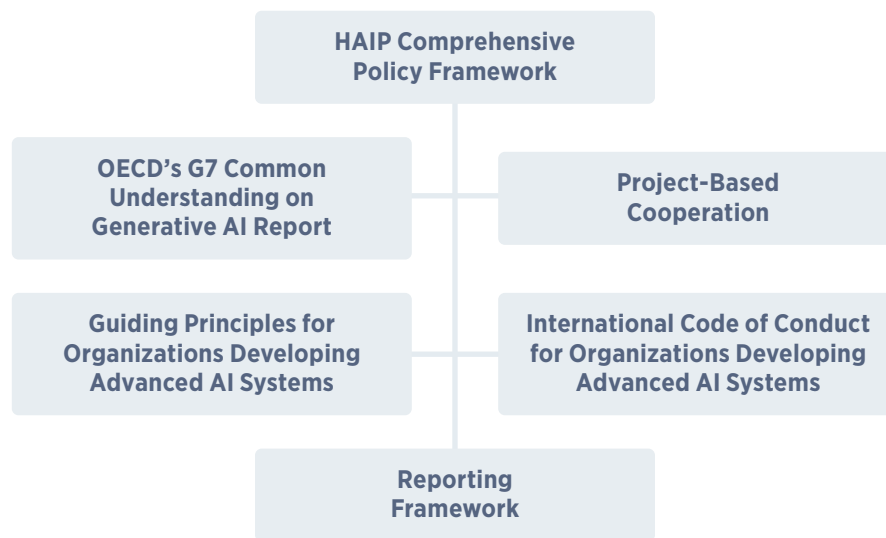
global innovation. However, SMEs may face resource limitations compared to larger firms when implementing trustworthy AI, particularly given the rapidly evolving landscape of AI governance. As G7 countries push to accelerate AI adoption among SMEs, it is crucial they are equipped to deploy it responsibly to build trust that provides assurance to end users and unlocks market opportunities.

For the purposes of the workshop and toolkit, a distinction was made between SME AI developers and SME AI deployers, with the latter as the primary focus. A developer designs, builds, and trains an AI system, while a deployer integrates and uses it in a specific context, and may also be the “user” of an AI system. This distinction thus includes the significant number of SMEs interested in adopting and integrating AI systems into their existing businesses, rather than only those focused on developing AI systems.

The workshop and toolkit aim to address the challenges faced by SME AI deployers seeking to adopt AI responsibly within their businesses. Section 1 provides a brief overview of the HAIP Comprehensive Policy Framework, Section 2 highlights key difficulties SMEs have reported in applying the HAIP Guiding Principles to their specific contexts, while Section 3 focuses on practical challenges within SMEs that can hinder their ability to implement the HAIP Guiding Principles and, ultimately, achieve trustworthy AI deployment.



Section 1 — Hiroshima AI Process (HAIP) Overview



Launched in 2023 under Japan's G7 Presidency, the [HAIP Comprehensive Policy Framework](#) (the "HAIP Framework") is a G7 initiative that aims to shape the governance of advanced AI systems across the value chain. It comprises four pillars: (1) a report analyzing the priority risks, challenges and opportunities arising from generative AI; (2) [Hiroshima Process International Guiding Principles for Organizations Developing Advanced AI System](#) ("Principles"); (3) the [Hiroshima Process International Code](#)

[of Conduct for Organizations Developing Advanced AI Systems](#) ("Code of Conduct"); and (4) project-based cooperation in AI.

The Principles provide the basic considerations that should be taken into account in responsibly developing and deploying advanced AI systems, including the most advanced foundation models and generative AI systems. The Code of Conduct provides further details on implementation of the Principles, outlining recommended actions for organizations developing advanced AI systems.

Building on the OECD AI Principles, the Principles focus on:

- **Lifecycle:** considering the entire AI lifecycle when mitigating misuse and vulnerabilities with red-teaming and traceability (principles 1 and 2);
- **Transparency:** promoting transparency, information sharing, and reporting for increased accountability (principles 3, 4 and 5);
- **Security:** investing in robust security controls as well as content authentication and provenance mechanisms to identify AI-generated content (principles 6 and 7);
- **Advancing the technology:** investing in research to advance the state of the art in risk mitigation, as well as to solve global and societal challenges and digital literacy initiatives (principles 8 and 9);
- **Standardization:** developing international technical standards (principle 10);
- **Privacy:** implementing protections for personal data and intellectual property (principle 11).

Under Italy's 2024 G7 Presidency and reaffirmed in the [Apulia Communiqué](#) by the G7 Leaders in June 2024, the G7 engaged the OECD to develop a mechanism to monitor voluntary adoption of the Hiroshima AI Process Code of Conduct, providing the impetus for the [HAIP Code of Conduct Reporting Framework](#) ("Reporting Framework"). The Reporting Framework is a voluntary tool with targeted questions to help organizations align with the Code of Conduct, embed safeguards, and promote transparency.

Section 2 — HAIP Guiding Principles & SMEs

A key focus of the workshop was understanding how the Principles can support SMEs in adopting and deploying AI systems. Over the course of the workshop, three Principles were frequently identified as being the most challenging for SMEs to apply: Principles 1, 2, and 10. Overcoming the challenges associated with these Principles would be a key starting point for many SMEs in implementing trustworthy AI deployment and this section highlights relevant considerations that can guide them towards that objective.

PRINCIPLE 1 — Take appropriate measures throughout the development of advanced AI systems, including prior to and throughout their deployment and placement on the market, to identify, evaluate, and mitigate risks across the AI lifecycle.

Considerations: Prior to putting an AI system into use, whether for managing schedules, organizing orders, or any other business function, users should carefully think about the potential risks. Consider what risks may arise in performing the business function and whether the use of an AI tool, and what kind of AI tool, is appropriate for the context. Where an AI system is to be deployed in a function that carries an elevated level of risk, procuring a system from a trusted third party who has taken steps to mitigate the risks that may arise during use (such as testing or providing user manuals that explain how to use the system appropriately) can help.

PRINCIPLE 2 — Identify and mitigate vulnerabilities, and, where appropriate, incidents and patterns of misuse, after deployment, including placement on the market.

Considerations: For deployments of advanced AI systems that carry elevated risk, it is important to continuously monitor for incidents (e.g., unintended effects, errors, or improper use) and act promptly if issues arise. For example, if a business uses an AI-based customer service chatbot, monitoring its responses can ensure that it provides the intended level of service. If the chatbot systematically gives inaccurate information or can be redirected to unintended purposes, this may indicate a malfunction or vulnerability to misuse. In such cases, working with the developer of the chatbot may be needed to mitigate the issue.

PRINCIPLE 10 — Advance the development and, where appropriate, adoption of international technical standards.

Considerations: Adopting recognized standards can enhance a business's credibility and trustworthiness with stakeholders, clients, and partners. Standards can also support greater alignment across organizations by providing clear expectations and guidelines and can assist in making well-informed strategic decisions. Selecting AI vendors whose products conform to AI standards provides greater assurance that providers are committed to responsible practices, helping reduce potential risks in AI adoption. Finally, SMEs should identify opportunities to be part of discussions setting industry-led AI standards. SME involvement in standard-setting for AI systems can bring unique insights to the development of standards that larger organizations may not have.

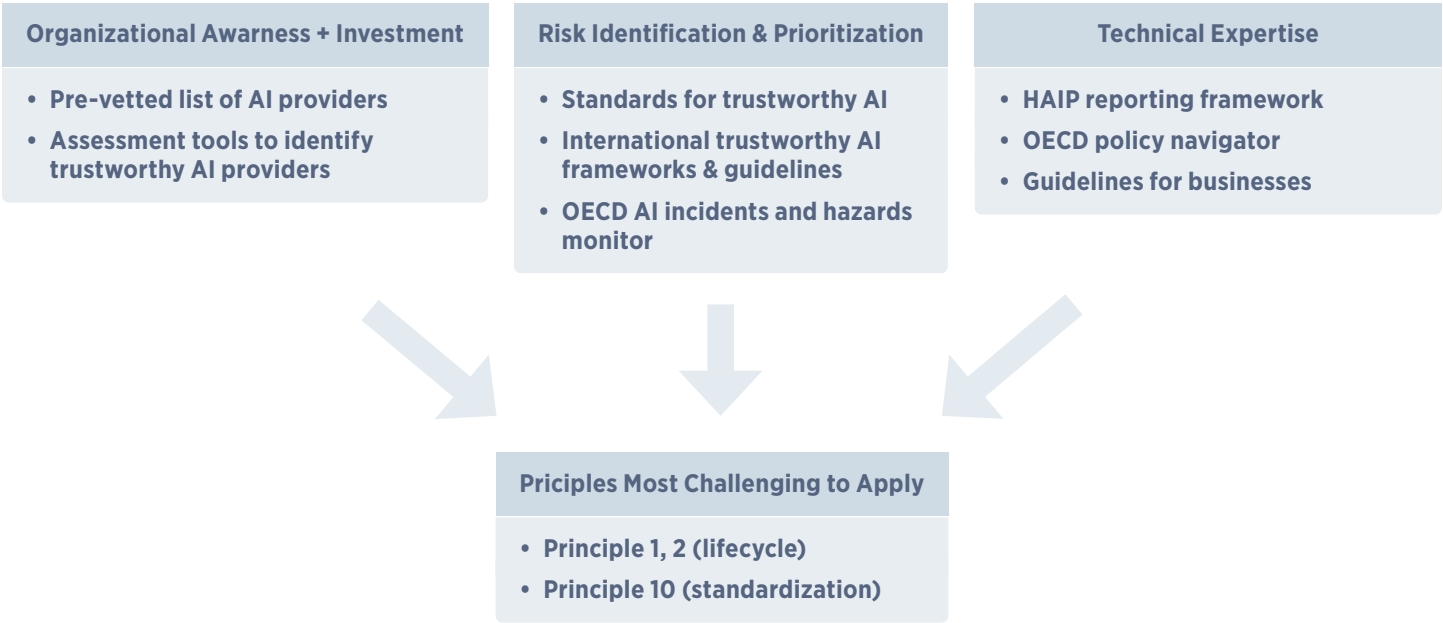
By considering these Principles when beginning their journey toward the deployment of trustworthy AI, SMEs can build a strong foundation for responsible and effective adoption.



Section 3 — Challenges & Tools

The workshop also discussed practical obstacles SMEs may experience when deploying AI in alignment with the HAIP Principles. This provided additional insights into the specific challenges associated with implementing the Principles identified in Section 2 (Principles 1, 2, and 10). These challenges are highlighted in this section with, for each, a list of corresponding resources that can help SMEs overcome them.

Resources for Addressing Challenges



Resources for organizational awareness and investment:

Participants highlighted the potential for misalignment between business functions within an organization, which they indicated tend to limit organizational understanding of and investment in trustworthy AI. Without a clear understanding of how AI could be leveraged to support specific business processes and objectives, decision-makers may miss opportunities that could benefit their organization and customers. Furthermore, they may not have clarity on where there are potential risks when integrating AI into business functions.

For SMEs beginning their AI journey, choosing services already recognized for trustworthy AI can significantly streamline the process. This approach helps avoid the time and effort of independently evaluating each system’s functionality and risk controls. Furthermore, some governments now offer pre-approved lists or vetting processes for AI providers that demonstrate strong governance, particularly for public procurement. This simplifies adoption, allowing SMEs to focus more on leveraging AI and less on extensive evaluations.

For example, the Government of Canada makes available a public list of AI suppliers who provide responsible and effective AI services, solutions, and products: [List of interested Artificial Intelligence \(AI\) suppliers - Canada.ca](#). Although designed for public procurement, this list can be a valuable resource for SMEs searching for AI suppliers with proven capabilities in relevant contexts and have completed an [Algorithmic Impact Assessment](#) (AIA). Listed suppliers have pre-qualified to deliver trustworthy and effective AI, providing pre-vetting for SMEs, potentially offering more confidence in their expertise and governance practices.

Other countries are pursuing initiatives in this realm. For example, the Government of the United Kingdom (UK) has provided the [Guidelines for AI Procurement](#) to enable effective adoption of AI. The UK also completed public consultations on the [AI Management Essentials](#) (AIME) tool in January 2025, a self-assessment tool intended to help organizations assess and implement responsible AI management systems and processes.



Additionally, the [OECD's Catalogue of Tools & Metrics for Trustworthy AI](#) provides SMEs with a useful directory of resources and assessment tools for identifying AI providers who align with recognized principles of responsible and ethical AI. This catalogue can help SMEs discover suppliers and solutions that demonstrate trustworthy practices, making it easier to evaluate and select AI partners.

Resources for risk identification & prioritization: SMEs often face difficulties in identifying and prioritizing risks associated with AI deployments, particularly in the face of rapidly evolving technology and shifting regulatory requirements. Limited resources can further constrain their ability to manage and mitigate these risks effectively.

The [OECD Framework for the Classification of AI Systems](#) is designed to help organizations identify and assess different types of AI based on their use context, human involvement, functionality, and risk profile. Its primary purpose is to support informed decision-making and responsible AI adoption. Before deployment, the framework can be used to gauge potential risks and the complexity of the solution, enabling the implementation of appropriate safeguards. For example, when considering an AI chatbot for customer service, the framework helps clarify the system's functionality, the level of oversight required, and the risks to be managed. It can be applied by mapping the use case to governance requirements, defining necessary safeguards, and translating these into vendor expectations. The framework should be revisited as the system scales to ensure risk controls remain effective.

[ISO/IEC 42001:2023](#) is an international standard that specifies requirements for establishing, implementing, maintaining, and continually improving an Artificial Intelligence Management System (AIMS) within organizations. It is designed for entities providing or deploying AI-based products or services, ensuring responsible development and use of AI systems. This standard provides clear guidance on establishing policies, risk management practices, accountability structures, and monitoring mechanisms to ensure AI systems are used responsibly and transparently. It is applicable to organizations of all sizes and sectors, providing the tools needed to govern AI responsibly and effectively, and will be updated over time to reflect evolving best practices and technological advancements.

The Canadian Digital Governance Council's [CAN/DGSI 101:2025 standard](#), "Ethical Design and Use of Artificial Intelligence by Small and Medium Organizations," is a new roadmap designed to help SMEs ethically adopt and use AI. This standard provides clear guidelines and requirements for integrating ethics into AI systems, covering key areas like risk management, ethics by design,

deployment strategies and continuous monitoring. It is specifically built for organizations with fewer than 500 employees, providing the tools to confidently use AI while upholding the highest ethical standards, and it will be regularly updated to keep pace with AI advancements.

OECD's [AI Incidents and Hazards Monitor](#) (AIM) helps SMEs spot and prioritize AI risks by surfacing sector-specific incidents, common problems in similar deployments, and providing insights about the common risk areas to watch for. AIM catalogs AI incidents and hazards to show risks and harms, reveal patterns over time, and support a shared, trustworthy understanding of AI. It also offers flexible filtering by country, industry, AI principle, severity, and other factors, so users can quickly spot high-risk areas for specific-use cases.

The Vector Institute's [Responsible Generative AI Governance Guide](#) offers a practical resource for responsible generative AI governance, providing a high-level understanding of challenges and risks through a user-friendly, interactive interface on controls and risk metrics. It supports the identification and mitigation of risks within a project and includes numerous resources for implementation planning. The platform features a built-in AI chatbot that helps navigate all functionalities, including risk assessment, and allows the addition of specific business cases for tailored risk identification.

The [Methodology for the Risk and Impact Assessment of Artificial Intelligence Systems from the Point of View of Human Rights, Democracy and the Rule of Law](#) (HUDERIA Methodology) provides a structured framework and guidance for voluntary evaluation of the potential risks and societal impacts associated with AI systems throughout their lifecycle. Adopted by the Committee on Artificial Intelligence of the Council of Europe, it supports the implementation of human rights, democracy, and rule of law principles in the design, development, deployment, and use of AI. The HUDERIA Methodology comprises four key components: the Context-Based Risk Analysis, which provides a structured approach to collecting and mapping information needed to identify and understand the risks an AI system could pose to human rights, democracy, and the rule of law, while also helping determine whether AI is an appropriate solution for the problem at hand; the Stakeholder Engagement Process, which proposes an approach for engaging relevant stakeholders to gather insights from potentially affected persons and contextualize potential harms and mitigation measures; the Risk and Impact Assessment, which outlines steps for assessing risks and impacts related to human rights, democracy, and the rule of law; and the Mitigation Plan, which describes processes for defining mitigation and remedial measures, including access to remedies



and iterative review. Applicable across sectors and adaptable to different contexts, the HUDERIA Methodology aims to guide both public and private entities in ensuring that AI technologies are trustworthy, human-centered, and aligned with fundamental rights.

Resources for technical expertise: Participants noted that they often have insufficient in-house technical expertise to translate HAIP guidance into practice. This often makes it difficult for SMEs to put trustworthy AI frameworks into practice.

The [HAIP Reporting Framework](#) was developed through a collaborative international effort, combining the strategic leadership of the G7 nations with the technical and policy expertise of the OECD. This partnership aimed to create a practical tool for organizations to comply with the Code of Conduct. To foster valuable peer-to-peer learning and knowledge sharing, the Framework transparently publishes submitted reports online. This feature allows users to not only see the types of questions asked in a practical, real-world context, but also to learn directly from the experiences and risk mitigation strategies employed by other organizations. Users will discover a growing collection of diverse reports from various industries and respondents. It is also particularly relevant for SMEs to note that the OECD is currently developing further initiatives and resources specifically designed to support SMEs in applying the HAIP Code of Conduct effectively.

The [OECD AI Policy Navigator](#) serves as a central repository for tracking public AI policies and regulations globally. Given the rapidly evolving regulatory landscape, the Navigator provides crucial insight into how governments are addressing AI risks.

Regularly updated by experts and official contributors, it offers the functionality to filter by country and includes international AI policy updates, allowing users to tailor information relevant to their jurisdiction. This ensures a reliable reference point for anticipating policy shifts and aligning organizational practices with emerging global standards.

The [EU AI Act Article 62](#) requires EU Member States to provide specific support measures for SMEs, including start-ups, involved in providing and deploying AI. These measures encompass training activities, expert advice, and the facilitation of SME participation in the standardization development process. With the date of entry into force set for August 2026, this obligation will trigger forthcoming support to aid trustworthy AI deployment.

The [AI Guidelines for Business](#), developed by Japan's Ministry of Internal Affairs and Communications and the Ministry of Economy, Trade and Industry, is a helpful resource for SMEs deploying AI, offering practical, non-binding guidance to navigate the complexities of trustworthy AI governance. This integrated document helps readers understand AI risks, aligning with international discussions such as HAIP. It emphasizes a risk-based approach, encouraging proactive identification and management of risks across the AI lifecycle, from initial data considerations to continuous operation. The Guidelines break down key principles, including human-centric design, safety, fairness, privacy, security, transparency, and accountability, and offer specific advice tailored to AI business users, supporting confident deployment, maximization of benefits, and minimization of potential harms while adapting to the rapidly evolving AI landscape.

Conclusion

The collective efforts of G7 members, in close collaboration with the OECD, underscore a strong commitment to supporting SMEs to implement trustworthy AI within their organizations. Through various initiatives, there is a concerted push to equip smaller enterprises with the resources to navigate AI's complexities responsibly. We anticipate that these ongoing efforts will encourage greater engagement from SMEs, including in the development processes of future AI frameworks, ensuring that their unique perspectives and needs are thoroughly integrated for a more inclusive and effective global AI ecosystem.



ANNEX — Hiroshima Process International Guiding Principles for Organizations Developing Advanced AI Systems

1. Take appropriate measures throughout the development of advanced AI systems, including prior to and throughout their deployment and placement on the market, to identify, evaluate, and mitigate risks across the AI lifecycle.

This includes employing diverse internal and independent external testing measures, through a combination of methods such as red-teaming, and implementing appropriate mitigation to address identified risks and vulnerabilities. Testing and mitigation measures should for example, seek to ensure the trustworthiness, safety and security of systems throughout their entire lifecycle so that they do not pose unreasonable risks. In support of such testing, developers should seek to enable traceability, in relation to datasets, processes, and decisions made during system development.

2. Identify and mitigate vulnerabilities, and, where appropriate, incidents and patterns of misuse, after deployment including placement on the market.

Organizations should use, as and when appropriate commensurate to the level of risk, AI systems as intended and monitor for vulnerabilities, incidents, emerging risks and misuse after deployment, and take appropriate action to address these. Organizations are encouraged to consider, for example, facilitating third-party and user discovery and reporting of issues and vulnerabilities after deployment. Organizations are further encouraged to maintain appropriate documentation of reported incidents and to mitigate the identified risks and vulnerabilities, in collaboration with other stakeholders. Mechanisms to report vulnerabilities, where appropriate, should be accessible to a diverse set of stakeholders.

3. Publicly report advanced AI systems' capabilities, limitations and domains of appropriate and inappropriate use, to support ensuring sufficient transparency, thereby contributing to increase accountability.

This should include publishing transparency reports containing meaningful information for all new significant releases of advanced AI systems.

Organizations should make the information in the transparency reports sufficiently clear and understandable to enable deployers and users as appropriate and relevant to interpret the model/system's output and to enable users to use it appropriately, and that transparency reporting should be supported and informed by robust documentation processes.

4. Work towards responsible information sharing and reporting of incidents among organizations developing advanced AI systems including with industry, governments, civil society, and academia.

This includes responsibly sharing information, as appropriate, including, but not limited to evaluation reports, information on security and safety risks, dangerous intended or unintended capabilities, and attempts by AI actors to circumvent safeguards across the AI lifecycle.

5. Develop, implement and disclose AI governance and risk management policies, grounded in a risk-based approach – including privacy policies, and mitigation measures, in particular for organizations developing advanced AI systems.

This includes disclosing where appropriate privacy policies, including for personal data, user prompts and advanced AI system outputs. Organizations are expected to establish and disclose their AI governance policies and organizational mechanisms to implement these policies in accordance with a risk-based approach. This should include accountability and governance processes to evaluate and mitigate risks, where feasible throughout the AI lifecycle.

6. Invest in and implement robust security controls, including physical security, cybersecurity and insider threat safeguards across the AI lifecycle.

These may include securing model weights and algorithms, servers, and datasets, such as through operational security measures for information security and appropriate cyber/physical access controls.

7. Develop and deploy reliable content authentication and provenance mechanisms, where technically feasible, such as watermarking or other techniques to enable users to identify AI-generated content.

This includes, where appropriate and technically feasible, content authentication such provenance mechanisms for content created with an organization's advanced AI system. The provenance data should include an identifier of the service or model that created the content, but need not include user information. Organizations should also endeavour to develop tools or APIs to allow users to



determine if particular content was created with their advanced AI system such as via watermarks.

Organizations are further encouraged to implement other mechanisms such as labeling or disclaimers to enable users, where possible and appropriate, to know when they are interacting with an AI system.

8. Prioritize research to mitigate societal, safety and security risks and prioritize investment in effective mitigation measures.

This includes conducting, collaborating on and investing in research that supports the advancement of AI safety, security and trust, and addressing key risks, as well as investing in developing appropriate mitigation tools.

9. Prioritize the development of advanced AI systems to address the world's greatest challenges, notably but not limited to the climate crisis, global health and education.

These efforts are undertaken in support of progress on the United Nations Sustainable Development Goals, and to encourage AI development for global benefit. Organizations should prioritize responsible stewardship of trustworthy and human-centric AI and also support digital literacy initiatives.

10. Advance the development of and, where appropriate, adoption of international technical standards.

This includes contributing to the development and, where appropriate, use of international technical standards and best practices, including for watermarking, and working with Standards Development Organizations (SDOs).

11. Implement appropriate data input measures and protections for personal data and intellectual property.

Organizations are encouraged to take appropriate measures to manage data quality, including training data and data collection, to mitigate against harmful biases. Appropriate transparency of training datasets should also be supported and organizations should comply with applicable legal frameworks.

