

Global AI Regulatory Landscape

Peter Pang , Managing Director of AA&T Consulting Services

Note:

The information contained herein is of a general nature and is not intended to address any particular circumstances of individuals or entities. Although we endeavor to provide accurate and timely information, there can be no guarantee that such information is accurate as of the date it is received or that it will continue to be accurate in the future. No one should act upon such information without appropriate professional advice after a thorough examination of the particular situation.

- AI Opportunities

Some AI transforming examples:

1. Healthcare:

- a. Disease detection and diagnosis
- b. Drug discovery and development

2. Finance:

- a. Fraud detection**
b. Risk assessment (Credit/Scenario analysis)

3. Transportation/logistics:

- a. Autonomous vehicles
- b. Route optimization

4. Entertainment / Media:

- a. AI-driven recommendation systems
- b. Smart/Creative content creation

5.E-Commerce:

- a. AI powered chatbot
- b. Customer Relationship Management tools

6. Data security:

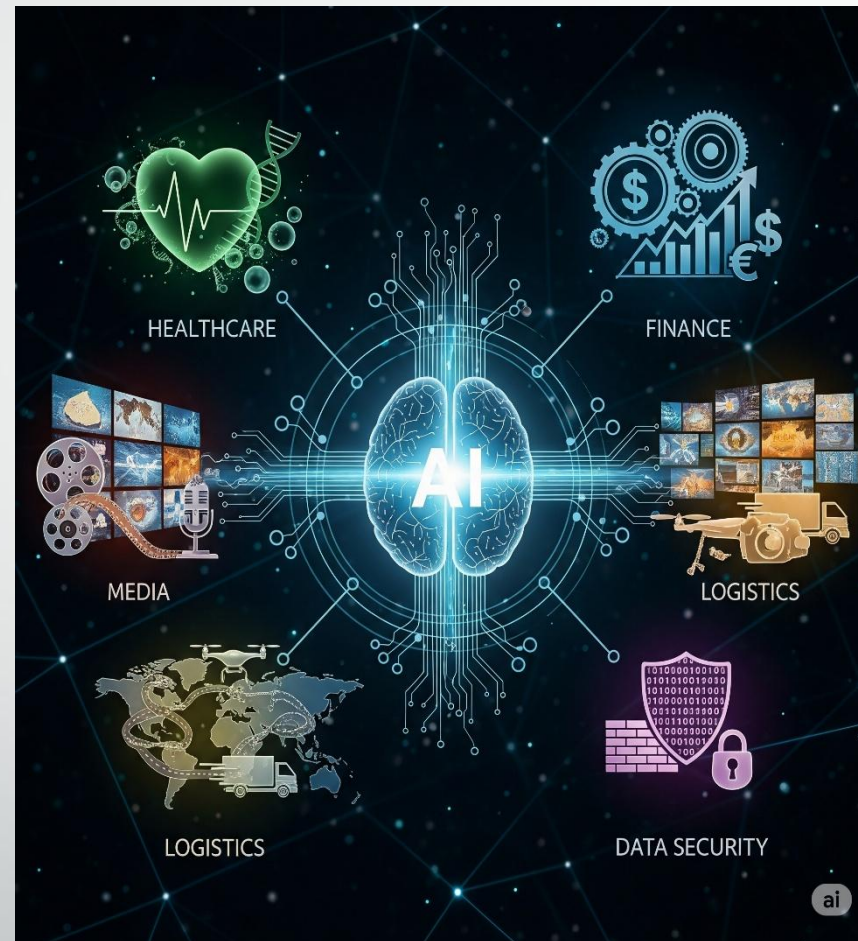
- a. Threat detection and prevention
- b. Network monitoring

7. Agriculture:

- a. Crop / soil monitoring
- b. Agricultural robotics

8. All industries

- a. Administration automation
- b. Fraud prevention

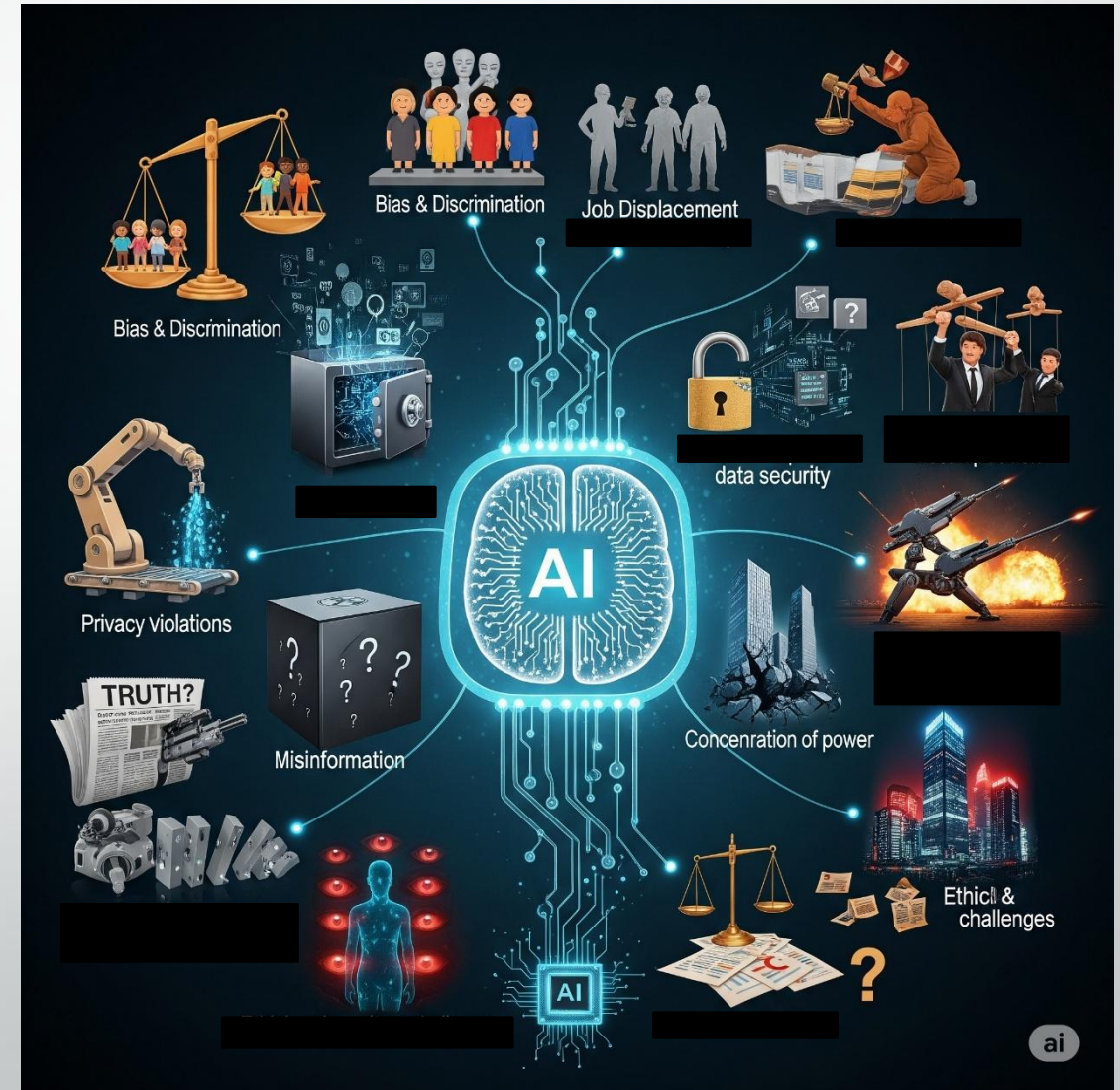


Summary of AI Opportunities & Risks

- AI Risks

Some Risks of AI development:

- 1. Bias and Discrimination*
- 2. Job Displacement and Economic Inequality*
- 3. Privacy Violations and Data Security*
- 4. Lack of Transparency and Explainability (The "Black Box" Problem)*
- 5. Misinformation and Manipulation*
- 6. Autonomous Weapons and Escalation of Conflict*
- 7. Concentration of Power*
- 8. Over-reliance and Unintended Consequences*
- 9. Existential Risks*
- 10. Ethical and Regulatory Challenges*



AI: Opportunities vs Risks

- Balancing Innovation and Regulations

**How to Balance
between
Innovation and
Regulations ?**



Global AI Regulatory Landscape

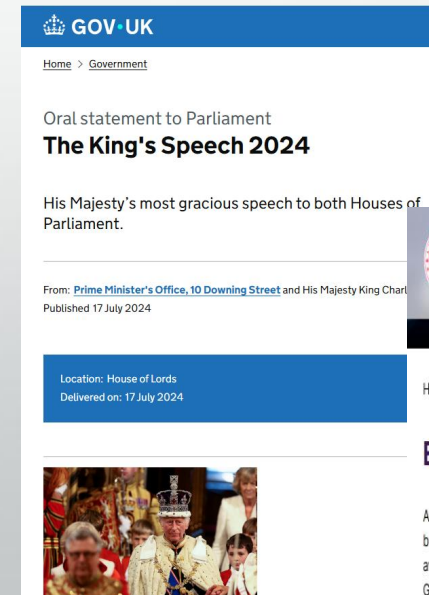
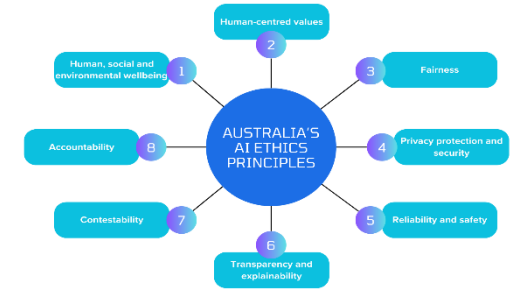
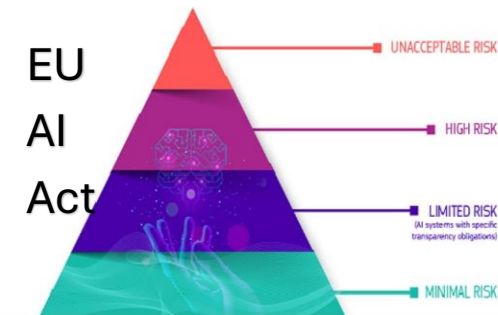
- Regulatory regimes at a glance

- **UNESCO's Recommendation on the Ethics of AI**
- **Mandatory regulations:**
 - EU
 - China
- **Voluntary rules**
 - Australia
 - Hong Kong
 - Singapore
 - The U.K. (pending mandatory)
 - The U.S. (but mandatory in some states)



A risk-based approach

The Regulatory Framework defines 4 levels of risk for AI systems:



Global AI Regulatory Landscape

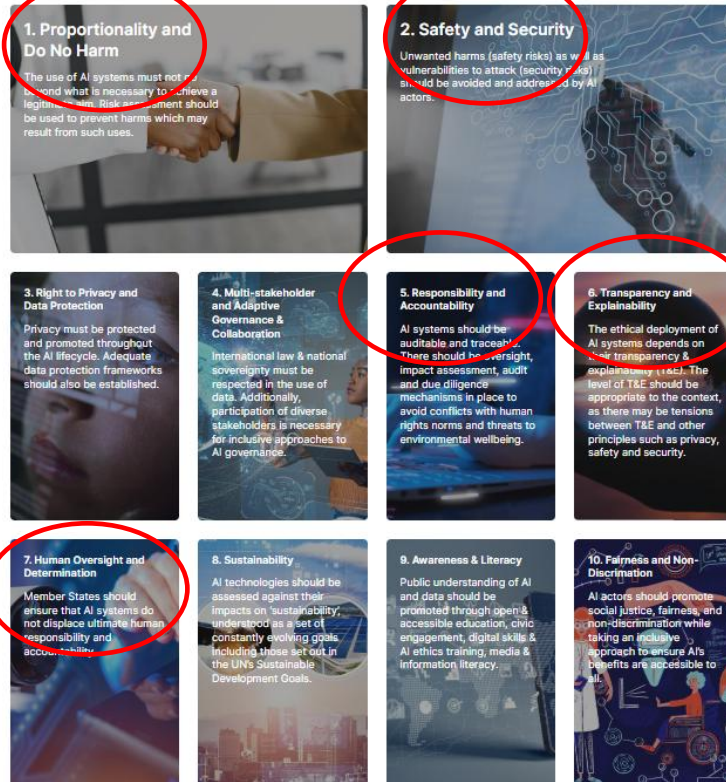
- United Nations

UNESCO's Recommendation on the Ethics of AI ([here](#))

- Adopted by 193 member states in 2021
 - Four core values
 - 10 core principles
 - 11 actionable policy areas
- Developed Readiness Assessment Methodology (RAM) and Ethical Impact Assessment (EIA) in 2023

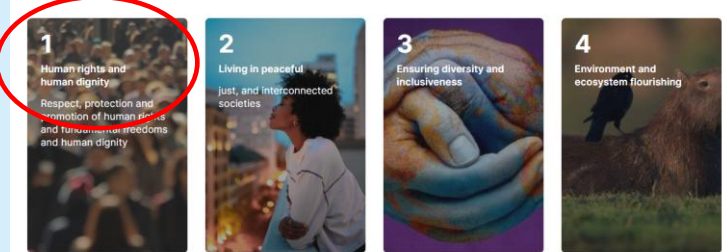
A human rights approach to AI

Ten core principles lay out a human-rights centred approach to the Ethics of AI.



Four core values

Central to the Recommendation are four core values which lay the foundations for AI systems that work for the good of humanity, individuals, societies and the environment:



Implementing the Recommendation

There is still a long way to go to provide Member States with actionable resources that ensure the effective implementation of the Recommendation. For this reason, UNESCO has developed two practical methodologies:



Global AI Regulatory Landscape

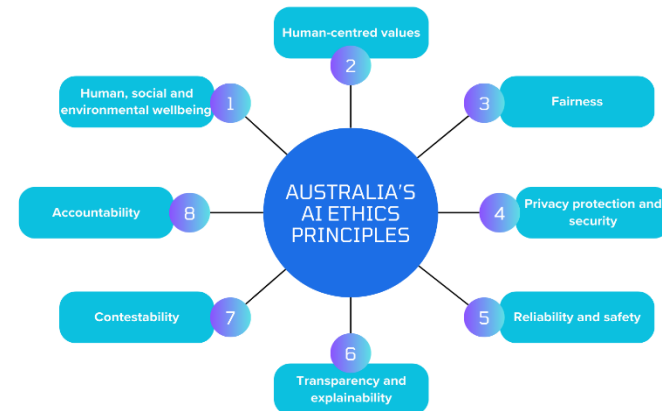
- Australia

Current status (1 Aug 2025)

- **No dedicated AI regulation.**
- **Rely on existing laws** (data protection, intellectual property, employment, competition, consumer protection etc.)
- **Adopted 8 voluntary AI Ethics Principles** (introduced in 2019): Human, societal and environmental wellbeing; Human-centred values; Fairness; Privacy protection and security; Reliability and safety; Transparency and explainability; Contestability; Accountability

Recent development:

Sep 2024	The Australian Government launched its Policy for the Responsible Use of AI in Government and released a proposal paper for introducing mandatory guardrails for AI in high-risk settings, along with Voluntary AI Safety Standards (VAISS).
Nov 2024	The Privacy and Other Legislation Amendment Bill 2024 passed, introducing enhanced transparency requirements in privacy policies for automated decision-making.



Global AI Regulatory Landscape

- China

Current status (1 Aug 2025)

- Not have a single unified AI regulation, like the EU AI Act
- Multi-level legal framework with various regulations on specific AI areas
- Key regulations include:
 - 1) Admin. Provisions on Deep Synthesis in Internet-based Info. Services 《互联网信息服务深度合成管理规定》 (effective. 10 Jan 2023);
 - 2) Interim Measures for the Management of Gen. AI Services (生成式人工智能服务管理暂行办法) (effective 15 Aug 2023);
 - 3) Measures for Labelling AI Generated Content 《人工智能生成合成内容标识办法》 (effective 1 Sep 2025)
 - 4) Regulations on Recommendation Algorithms (互联网信息服务算法推荐管理规定) (effective 1 Mar 2022)
- Other laws affecting AI includes data related laws, such as the Cybersecurity Law, the Personal Information Protection Law, the Data Security Law, the Copyright Law.



Global AI Regulatory Landscape

- China

Recent development:

3 Nov 2022 (effective 10 Jan 2023)	Administrative Provisions on Deep Synthesis in Internet-based Information Services 《互联网信息服务深度合成管理规定》: This Administrative Provisions issued on 25 November 2022 by Cyberspace Administration Office of China, Ministry of Industry and Information Technology, and Ministry of Public Security and effective 10 January 2023, set out responsibilities of the deep synthesis services providers in data security, user registry, algorithm audit, data protection etc.; and prohibits producing, publishing or spreading fake news.
13 Jul 2023 (effective 15 Aug 2023)	Interim Measures for the Management of Generative AI Services (生成式人工智能服务管理暂行办法): On 13 July 2023, the Cyberspace Administration of China (CAC), the National Development and Reform Commission, the Ministry of Education, the Ministry of Science and Technology, the Ministry of Industry and Information Technology, the Ministry of Public Security, and the National Radio and Television Administration jointly released this AI Interim Measures, the first comprehensive administrative regulation on the management of Generative AI services, effective 15 August 2023, outlined requirements on data security, algorithm transparency, and content control for generative AI technologies like large language models.

Global AI Regulatory Landscape

- China

Recent development (cont'd):

14 Mar 2025 (effective 1 Sep 2025)	Measures for Labelling AI Generated Content 《人工智能生成合成内容标识办法》: On 14 March 2025, the CAC released this labelling measure. effective 1 September 2025. It will be mandatory to implicitly and explicitly label AI-generated content. Explicit labels are required for those that are easily perceived by users and must be added to text, audio, images, videos, and virtual scenes, while implicit labels are embedded within the file's metadata.
25 Apr 2025 (effective 1 Nov 2025)	On 25 April 2025, the State Administration for Market Regulation and the Standardization Administration of China jointly released three national standards aimed at enhancing the security and governance of generative AI. These standards will officially take effect on 1 November 2025. The three standards are: 1. Cybersecurity Technology – Gen. AI Data Annotation Security Specification (网络安全技术生成式人工智能数据标注安全规范): This standard outlines requirements regarding security, staff management, security measures validation and evaluation on data labelling process, platform and tools. 2. Cybersecurity Technology - Security Specification for Gen. AI Pre-training and Fine-tuning Data (网络安全技术生成式人工智能预训练和优化训练数据安全规范): It outlines requirements and evaluation criteria for ensuring the security of datasets used in the pre-training and fine-tuning phases. 3. Cybersecurity Technology - Basic Security Requirements for Gen. AI Service (网络安全技术生成式人工智能服务安全基本要求): This standard establishes security requirements for generative AI services, encompassing user data security assessments, data protection measures, and the safeguarding of training models and datasets.

Global AI Regulatory Landscape

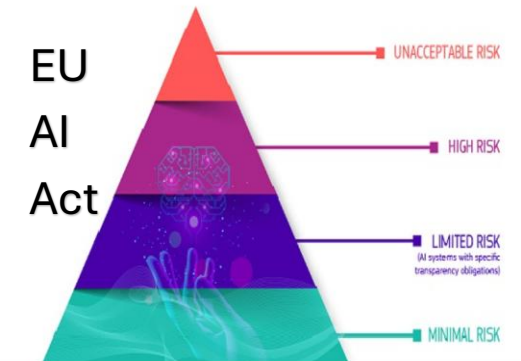
- European Union (EU)

Current status (1 Aug 2025)

- The **EU AI Act** is the landmark piece of legislation, entered into force on 1 August 2024, but its provisions are being applied **in phases**. It adopted **risk-based approach** and classified AI models into 4 different risk categories (see next slide). The AI Act is designed to protect people's fundamental rights, establish a harmonised market, and create a supportive environment for innovation and investment.
- From 2 February 2025, all providers and deployers of AI systems have been obliged under the EU AI Act to ensure a sufficient level of AI literacy of their staff dealing with AI.
- On 9 April 2025, the European Commission (EC) published the "AI Continent Action Plan" ([here](#)). The Plan intends to enhance AI capabilities in the EU by promoting initiatives.
- On 7 May 2025, the EC published its Q&A providing further guidance on the compliance and enforcement of the AI literacy requirement.
- On 10 July 2025, the EC published the final version of the "General-Purpose AI Code of Practice" (the "Code") ([here](#)). The Code is intended to help providers of "General Purpose AI models" to comply with obligations contained in the EU AI Act relating to transparency, safety and security, and intellectual property.

A risk-based approach

The Regulatory Framework defines 4 levels of risk for AI systems:



Global AI Regulatory Landscape

- European Union (EU)

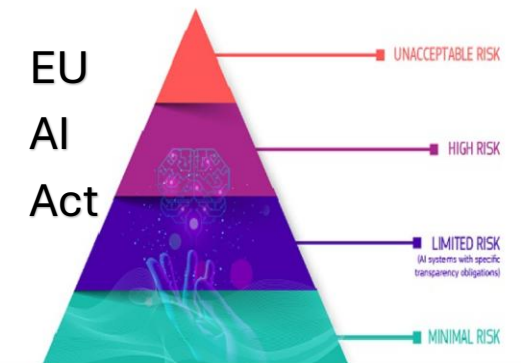
EU AI Act – Phased approach

Some key effective dates of EU AI Act are as follows:

- **2 February 2025:** Prohibitions on certain AI systems (e.g., social scoring, untargeted scraping of internet/CCTV for facial recognition databases, emotion recognition in workplaces/education) came into effect. Organisations must also ensure adequate AI literacy among employees involved in AI use and deployment.
- **2 August 2025:** Provisions for General-Purpose AI (GPAI) models apply, including information obligations for providers and transparency requirements for AI-generated content (e.g., deepfakes must be labelled). Rules governing AI Act penalties also come into effect.
- **2 August 2026:** Obligations for high-risk AI systems (e.g., in critical infrastructure, education, employment, law enforcement, migration) start to apply.
- **2 August 2027:** Obligations apply for high-risk AI systems brought into scope via other EU product safety legislation (e.g., medical devices).

A risk-based approach

The Regulatory Framework defines 4 levels of risk for AI systems:



Global AI Regulatory Landscape

- European Union (EU)

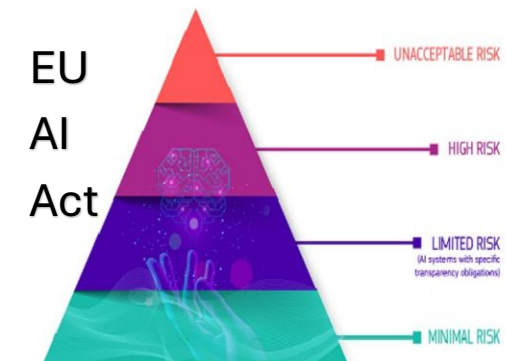
EU AI Act – Risk categories

The Act classifies AI systems into 4 risk categories:

- **Minimal risk:** This category includes most AI systems, such as AI-enabled recommender systems and spam filters, and have no special obligations.
- **Specific transparency risk (or Limited risk):** This category includes AI systems like chatbots, that must disclose to users that they are interacting with a machine. Providers have to ensure that AI-generated content is identifiable, clearly and visibly labelled.
- **High risk:** AI systems in this category are subject to strict obligation before they can be put on the market, including adequate risk mitigation systems, high quality of data sets minimising discriminatory outcomes, logging activities ensuring traceability, detailed documentation, clear user information, human oversight, and a high level of robustness, accuracy, and cybersecurity.
- **Unacceptable risk:** This category includes AI systems considered a clear threat to the safety, livelihoods and rights of people and are banned. The Act prohibits systems with these practices, like harmful manipulation and deception, harmful AI-based exploitation of vulnerabilities, social scoring, CCTV material to create facial recognition databases etc.

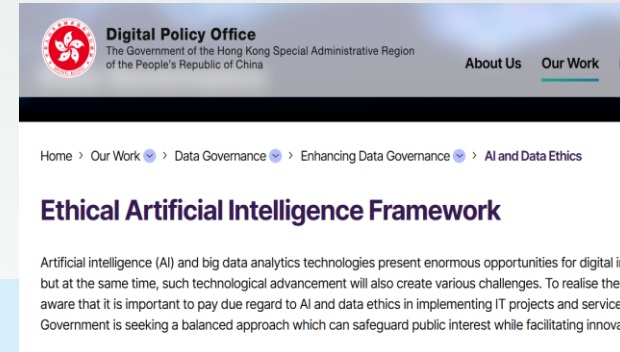
A risk-based approach

The Regulatory Framework defines 4 levels of risk for AI systems:



Global AI Regulatory Landscape

- Hong Kong



Current status (1 Aug 2025)

- No dedicated AI regulation
- Existing Ethical AI Framework is for voluntary adoption.
- Companies follow existing framework of regulations, overseen by different government agents e.g. the Digital Policy Office, the Privacy Commissioner for Personal Data, the Securities and Futures Commission, the Hong Kong Monetary Authority etc.

Recent development:

Jul 2024	In July 2024, the Digital Policy Office published the Ethical Artificial Intelligence Framework (Version 1.4). The Framework is for voluntary adoption by the private sector. It sets out ethical principles, an AI governance model, an AI lifecycle guide, and an impact assessment template for organisations to integrate into existing risk management and project governance processes.
Apr 2025	On 15 April 2025, the Digital Policy Office (DPO) released the Hong Kong Generative Artificial Intelligence Technical and Application Guideline. It aims to provide practical operational guidance for technology developers, service providers, and users in the application of generative AI technology. It covers the scope and limitations of applications, potential risks and governance principles of generative AI technology, including technical risks such as data leakage, model bias, and errors that need to be addressed.

Global AI Regulatory Landscape

- Singapore



Current status (1 Aug 2025)

- No dedicated AI regulation
- Relies on sectoral guidance from various authorities and commissions, including the AI Verify Foundation (AIVF), Infocomm Media Development Authority (IMDA), Personal Data Protection Commission (PDPC), Monetary Authority of Singapore (MAS), Competition etc.
- The Singapore government developed various frameworks and tools to guide AI development, including:
 - The Model AI Governance Framework (2019, updated in 2020), providing guidance to private sector to address key ethical and governance issues in deploying AI solutions
 - AI Verify, an AI governance testing framework and toolkit designed to help organisations validating the performance of their AI systems against AI ethics principles through standardised tests.
- The National AI Strategy 2.03 (first launched in 2019, updated in 2023) (NAIS 2.0), outlined Singapore's ambition and commitment to building a trusted and responsible AI ecosystem
- AIVF and IMDA drafted revised Model AI Governance Framework for Generative AI (2024 Framework), to address new issues and provide guidance on suggested practices for safety evaluation of Generative AI models.

Global AI Regulatory Landscape

- United Kingdom (UK)

Current status (1 Aug 2025)

- No single dedicated AI regulation.
- Companies and organisations must apply existing laws on data protection, intellectual property etc.
- In April 2023, the UK government published its AI policy white paper, opting for an outcomes based approach rather than a single omnibus law. It outlined five principles consisting of safety, security and robustness; transparency and explainability; fairness; accountability and governance; and contestability and redress.
- However, on 17 July 2024, the King's Speech proposed a set of binding measures on AI, which deviates from the previous agile and non-binding approach. The Digital Information and Smart Data Bill was also announced, which will be accompanied by reforms to data-related laws, to support the safe development and deployment of new technologies (which may include AI). It is yet to be implemented.



[Home](#) > [Government](#)

Oral statement to Parliament

The King's Speech 2024

His Majesty's most gracious speech to both Houses of Parliament.

From: [Prime Minister's Office, 10 Downing Street](#) and His Majesty King Charles III
Published 17 July 2024

Location: House of Lords
Delivered on: 17 July 2024



Global AI Regulatory Landscape

- United Kingdom (UK)

Recent development

Apr 2024	In April 2024, the UK (AI Safety Institute) and the US signed landmark AI Safety agreement to work together on testing advanced AI.
Sep 2024	The UK signed its first international treaty with the Council of Europe addressing AI risks, focusing on protecting human rights, democracy, and the rule of law.
Mar 2025	The Artificial Intelligence (Regulation) Bill, a Private Member's Bill, was reintroduced in the House of Lords, proposing a new "AI Authority" and codifying AI principles into binding duties.
June 2025	The new Data (Use and Access) Act 2025 introduces significant changes to data law in the UK, with implications for application of AI technology.



[Home](#) > [Government](#)

Oral statement to Parliament

The King's Speech 2024

His Majesty's most gracious speech to both Houses of Parliament.

From: [Prime Minister's Office, 10 Downing Street](#) and His Majesty King Charles III
Published 17 July 2024

Location: House of Lords
Delivered on: 17 July 2024



Global AI Regulatory Landscape

- United States (US)

Current status (1 Aug 2025)

- No comprehensive federal regulations that regulate the AI development or specifically prohibit or restrict their use.
- Regulation is a mix of executive orders, voluntary frameworks (like NIST AI Risk Management Framework), and state-level initiatives.
- Throughout 2024 and early 2025, many states have introduced or enacted AI-related legislation, addressing specific concerns like deepfakes in elections, data privacy, and government use of AI, including Colorado's broad AI Act, New Hampshire criminalizing malicious deepfakes, Tennessee's Ensuring Likeness, Voice, and Image Security Act (ELVIS Act), Maryland's rules for AI use in state government, and California's package of AI laws (some effective January 2026). On 12 Jun 2025, the New York legislature passed the Responsible AI Safety & Education ("RAISE") Act (S 6953), a frontier model public safety bill that would establish safeguard, reporting, disclosure, and other requirements for large developers of frontier AI models.



Global AI Regulatory Landscape

- United States (US)

Current status (1 Aug 2025) (cont'd)

- President Trump issued an Executive Order for “Removing Barriers to American Leadership in AI” ("Removing Barriers EO") in January 2025, that rescinds all policy and regulation changes under President Biden's Executive Order for the Safe, Secure, and Trustworthy Development and Use of AI ("Biden EO") that are "inconsistent" with "enhanc[ing] America's global AI dominance."
- The White House Blueprint for an AI Bill of Rights, issued under Biden, asserts guidance on five principles and associated practices including: 1) safe and effective systems; 2) algorithmic discrimination and protection; 3) data privacy; 4) notice and explanation; and 5) human alternatives, consideration and fallbacks. The Removing Barriers EO did not specifically revoke the AI Bill of Rights.
- Leading AI companies, like Adobe, Amazon, Anthropic, Google, IBM, Meta, Microsoft, Nvidia, Open AI, Palantir, Salesforce, etc. are noted to voluntarily committed to "safe, secure, and transparent development of AI technology". They are committed to internal/external security testing of AI systems before release, sharing information on managing AI risks etc.

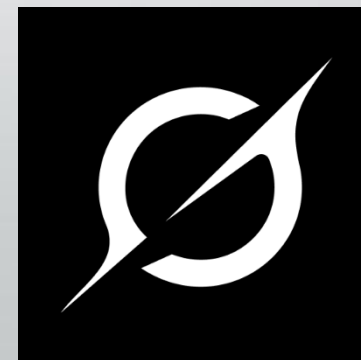


AI Companies Ethics Principles & Tools

- Corporate rules/tools at a glance

Corporate AI Principles (Pr)/tools (t):

- Amazon (AWS) ([Pr](#), [t](#))
- Anthropic (Claude) ([Pr1](#); [Pr2](#); [Pr3](#); [t1](#); [t2](#))
- Google (Gemini) ([Pr](#); [t1](#); [t2](#); [t3](#); [t4](#); [t5](#))
- Meta (LLaMA) ([Pr](#); [t](#))
- Microsoft (copilot) ([Pr](#), [t](#))
- Nvidia ([Pr](#), [t](#))
- OpenAI (ChatGPT) ([Pr](#))
- xAI (Grok) ([Pr1](#), [Pr2\(Draft\)](#))



Global AI Regulatory Landscape

- Conclusion

- A **divergence** in global approaches to AI governance:
 - The U.S. has reversed a federal mandate on AI safety
 - Other significant economies, including the EU and China, doubled down on their respective regulatory paths.
- The development of robust ethical frameworks and practical tools by leading AI companies indicated an emerging industry **consensus** on core values like fairness, transparency, and safety.
- Companies operating internationally should:
 - **navigate complex rules of different regions**, from the EU's strict, risk-based classifications to China's multi-layered content and data regulations, and the varied, sector-specific laws in other regions.
 - **proactively adapting AI governance strategies** to these frameworks in a globally interconnected market; and
 - **establish ethical practices and adopt safety tools** in line with common practice by leading AI companies in building a more trustworthy and accountable AI ecosystem for everyone..



Global AI Regulatory Landscape

- Time to reflect

- *If you can choose, what do you want AI to do and not to do?*
- *When do we need regulating AI?*
 - *When people really get hurt? About to get hurt?*
 - *When AI shows misaligned/dishonest behaviours / deceives users/leaks company secrets/blackmails human officer?*
 - *Or at the process at the design stage?*

Any questions?



AA & T Consulting

Hope that you enjoy this session

If you need any help in technology or regulatory risk matters, please feel free to contact us by phone (+852 9181 8659 (HK); +61 452 371 753 (Aus.)), email (advisory@aathk.com) or via website's "contact us" page at: www.aathk.com or www.aataus.com.



Note: The information contained herein is of a general nature and is not intended to address any particular circumstances of individuals or entities. Although we endeavor to provide accurate and timely information, there can be no guarantee that such information is accurate as of the date it is received or that it will continue to be accurate in the future. No one should act upon such information without appropriate professional advice after a thorough examination of the particular situation.