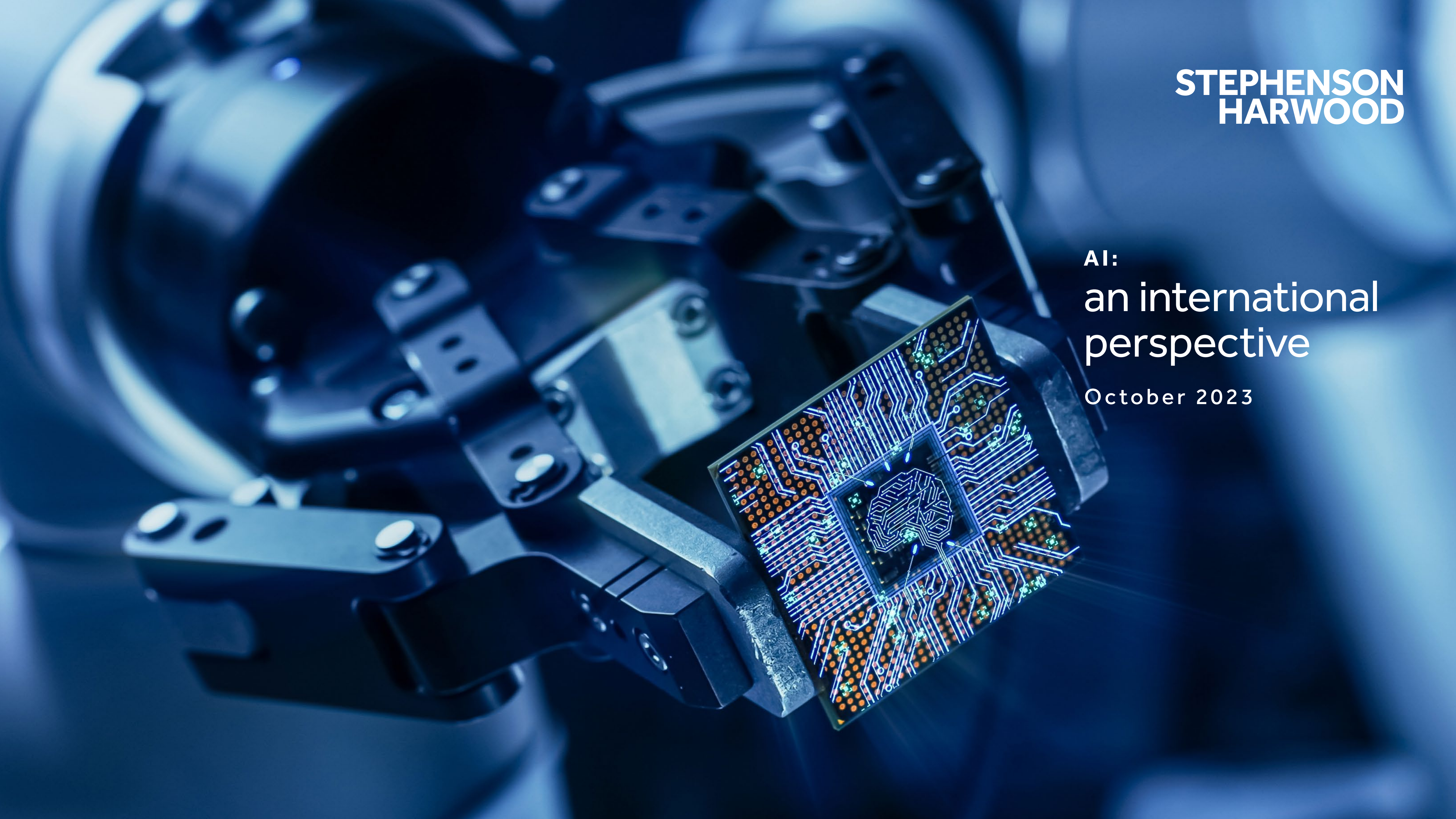


A blue-tinted photograph of a robotic arm holding a microchip. The microchip is square and features a complex circuit pattern with a central brain-like shape. The background is blurred, showing mechanical components of the robotic arm.

**STEPHENSON
HARWOOD**

AI:
an international
perspective

October 2023

INTRODUCTION

Artificial intelligence technology (“AI”) is transforming various aspects of business and society around the world, with applications ranging from generative AI and machine learning, to low-code and no-code AI.

These developments pose interesting legal challenges, as well as opportunities, for businesses that are developing or implementing AI products and services. However, it is becoming clear that there is a complex patchwork of regulation emerging rather than a common approach, and businesses looking to develop products for international application need to be aware of how the web of AI-regulation is emerging.

Notwithstanding these new, specific, AI regulations that are on the horizon, businesses also need to consider how existing laws and regulations – such as those on data protection and intellectual property - apply to their AI activities. Moreover, businesses should be aware of the ethical and social implications of AI, and follow guidance and best practices issued by relevant regulators and industry bodies. The ‘first mover’ advantage of observing the guidance and best practices is likely to pay dividends in the longer term.

In this insight, we explore how the international landscape for AI regulation is evolving across key tech-focussed jurisdictions. An overarching observation is that many of these proposed regulatory frameworks will have some form of extraterritorial effect. This will create a challenging matrix for international businesses when it comes to managing compliance across their global footprint.

“
As the pace of AI development accelerates, the challenge of creating robust, fair, and transparent regulatory frameworks is of paramount importance, and the stakes are high. It remains to be seen how this plays out globally.

SIMON BOLLANS, HEAD OF TECHNOLOGY

EUROPE

The EU is leading the way in proposing a comprehensive and harmonised regulation of AI. In April 2021, the European Commission (“**Commission**”) introduced the draft [Artificial Intelligence Act](#) (“**AI Act**”), which aims to ensure that AI systems placed on the market or put into service in the EU are safe and respect existing laws on fundamental rights and values.

In its current form, the AI Act adopts a risk-based and proportionate approach to the regulation of AI, whereby AI systems are classified and regulated according to their potential impact on human rights, safety, and democracy. In particular, “low-risk” AI systems will face minimal transparency obligations, while “high-risk” AI systems will have to comply with strict requirements on data quality, human oversight, risk management, and conformity assessments. AI systems that represent an ‘unacceptable risk’ will be outright prohibited.



The AI Act is still subject to negotiations and amendments by the EU institutions and member states, and is expected to be formally adopted by late 2023 or early 2024, and to come into force by mid-2025 at the very earliest. Some of the key issues that need to be resolved include:

- **the definition of AI:** there is debate over whether AI should be defined broadly to cover a wide range of uses, or narrowly to focus on more complex and sophisticated systems;
- **self-assessment of risk level:** there is concern over whether developers should be able to determine the risk level (which corresponds to the level of regulation) of their own AI systems, or whether there should be more external oversight and verification;
- **prohibited AI systems:** there is disagreement over which AI systems should be banned or restricted, such as those that use biometric surveillance in public spaces or manipulate human behaviour; and
- **enforcement:** there is uncertainty over how the proposed coordination between various national and EU authorities, such as the Artificial Intelligence Board and the AI Office, will work in practice and ensure consistent and effective enforcement action.

The AI Act will have extraterritorial effect, meaning that it will apply to any AI system that is offered or used in the EU, regardless of where it was developed or produced. This will have significant implications for businesses operating globally, as they will have to ensure compliance with the EU rules as well as any other applicable national or regional laws. The burden of enforcement will be passed to member states (with each state required to designate or create a regulatory body) with EU-wide issues coordinated by the Commission, advised by a new AI Council.

For more information, please read [our Insight on the EU's AI Act](#).

The UK government has established various bodies and initiatives to advise and guide its AI policy and governance approach, such as the Office for Artificial Intelligence and the Centre for Data Ethics and Innovation. In a continuation of the UK's National AI Strategy, in March 2023, the UK's Department for Science, Innovation and Technology published a white paper (the "[White Paper](#)") outlining a 'pro-innovation approach' to AI regulation with a principles-based framework. It suggested the UK Government will take a light-touch, decentralised approach to AI regulation.

The White Paper outlined a set of cross-sectoral principles intended to encourage responsible AI design, development and use. These principles would be applied by existing sector-specific regulators to deliver a consistent and proportionate approach, whilst affording a degree of flexible interpretation depending on the sector risks and regulatory objectives.

The five key principles focus on:

- 1 safety, security and robustness
- 2 transparency and explainability
- 3 fairness
- 4 accountability and governance; and
- 5 contestability and redress.



As with the EU AI Act, the UK framework would have extra-territorial effect and would apply to those developing, deploying and using AI systems across the UK, irrespective of whether they are based in the UK. If the UK continues with its principle based approach, it remains to be seen whether compliance with the EU AI Act (assuming this forms some form of global benchmark) will in itself meet the UK requirements – we suspect it will be more nuanced, and will require businesses to consider the implications of both.

For a deep dive into the content of the White Paper, [read our Insight here](#). The consultation on the White Paper closed on 21 June 2023 and more detailed guidance is expected to follow by the end of the year.

One key development following the release of the White Paper has suggested the UK's approach to AI is still evolving. In July 2023, the House of Lords' Communication and Digital Committee announced its [inquiry](#) into large AI language models. The inquiry will evaluate the work of the government and regulators into AI so far (including an inspection of the adequacy of the

White Paper) and consider how the UK should respond to AI over the next three years. It will also look into how the UK's approach compares to other jurisdictions, including the US and China, and whether international coordination regarding the regulation of AI will be different.

UK regulators are also beginning to respond to and consider the impact of AI on key areas of impact. For example, The Bank of England, together with the Prudential Regulation Authority and to the Financial Conduct Authority, published a [Discussion Paper](#) focusing on the regulation of AI and machine learning in the UK financial services industry. The discussion paper is aimed at creating a broad-based and structured discussion on the key challenges with the use and regulation of AI in the sector. One key takeaway of the discussion paper is that these regulatory authorities are considering how the existing sectoral rules, policies and principles should apply to AI systems and how to identify key gaps in the regulations. This demonstrates the patchwork of legislation affecting AI development and the particular challenges for businesses operating in this space.

MIDDLE EAST

The UAE government has – so far – not produced formal AI regulation or legislation. However, it has been vocal in its support for the commercial prioritisation of utilising AI. In October 2017, the UAE government launched its [Strategy for AI](#) and formed the UAE Council for AI. The strategy focuses on establishing the UAE as a world leader in the use of AI by 2031, rather than looking at the regulation of AI. To support this objective, the government also released, in April 2023, the [Generative AI Guide](#) which provides practical advice for businesses on how to incorporate AI into their practices, and consists of 100 AI applications and implementations of AI in the media industry. That said, as outlined below, there are a number of programs, toolkits and bodies set up in the region, forming a patchwork of initiatives and incentives.

In line with the UAE's AI strategy 2031, the Dubai International Financial Centre (DIFC), in collaboration with the UAE AI Office, unveiled the AI and coding license and incentives to encourage establishing companies with such capabilities. Companies acquiring this license gain access to the DIFC Innovation Hub, the region's foremost cluster of FinTech and innovation entities, and the potential for employees of such companies to secure UAE Golden Visas.



The DIFC is also set to establish the Dubai AI & Web 3.0 Campus, poised to be the MENA region's largest congregation of AI and technology firms. This campus is designed to house Research & Development facilities, accelerator programs, and synergistic workspaces, thereby enticing AI companies to establish, develop, and expand their operations. Furthermore, in an endeavour to bolster its transformation into a digital society, Dubai will grant commercial licenses specifically tailored for companies in the AI sector. Administered by the AI and Web 3.0 Campus, these licenses will come with a 90% subsidy, emphasizing the commitment to fostering digital business growth.

On 7 September 2023, the [DIFC Enacts Amended Data Protection Regulations](#) emphasizing the handling of personal data by AI and similar systems. The DIFC's focus is not on algorithm content but rather on the organizations utilizing these AI systems. This update stands as the pioneering regulation in the region, addressing the processing of personal data via systems like AI and machine learning. It re-emphasises the DIFC as a hub for fostering a flexible environment for AI technology development that upholds data ethics.

Digital Dubai, the body responsible for the development of Dubai's technology and data policies and strategies, has also released an [Ethical AI Toolkit](#). The AI Toolkit is

designed to provide practical support to the state by providing principles and guidelines for developers of AI technology. Digital Dubai explains that the strategy behind the AI Toolkit is to create a trusted platform for discussing policy development on AI and ethical application. The body hints that in the future, advisory and audit services may be introduced (at the moment, the AI Toolkit has a self-assessment tool for AI developers) and in the long term the guidelines may develop into regulations governing AI development and commercial deployment.

In Qatar, an Artificial Intelligence committee was set up in 2021. The committee released a National AI Strategy that, echoing the UAE's approach, chiefly focuses on how to incorporate AI into Qatar's businesses and economy. The strategy also called for the Qatari government to develop an AI Ethics and Governance framework to address ethical and public policy questions which echoes the approach from Dubai. As yet, no such framework appears to be in the works.

The Saudi Data & AI Authority was set up in Saudi Arabia in 2019 to handle all matters relating to the organisation, development and handling of data and AI. The Authority published [draft guidance](#) on the ethical use of AI, listing seven principles to govern the development and use of AI including fairness, privacy, reliability and accountability. Businesses will also be subject to Saudi Arabia's extensive data protection laws when using personal data in connection with AI products.

INDIA

In 2022, the Indian government proposed the Digital India Act designed to implement regulations across India's whole digital ecosystem (and to replace the Information Technology Act of 2000, which the government has concluded no longer provides the standards of protection required for today's digital landscape).

The proposed Digital India Act will include a chapter focusing on AI technology and it is expected to introduce penal provisions to tackle violations and user harm arising from emerging technologies, including generative AI. The regulation is rumoured to be principles based with a focus on 'high risk' AI systems and encourage accountability and the ethical use of AI based tools. In addition, the legislation is expected to define specific "no-go" areas for companies utilizing AI in consumer applications to prevent severe harm and with penalties for violation.

The draft Digital India Act is expected to be opened for public consultation later in 2023, before progressing to the next stage of implementation. In addition to specific rules regulating the use of AI, it is anticipated that the Act will also introduce new regulations for online safety, e-commerce, over-the-top platforms, ad-tech and other types of digital intermediaries.



CHINESE MAINLAND AND HONG KONG

Chinese Mainland's [interim AI measures](#) (which came into effect on 15 August 2023, the “**Interim Measures**”) represent the country's first specific regulation on AI technology. The measures target those who use generative AI technologies to provide content to the general public in Chinese Mainland. Research, development and internal-facing products are not within the scope of the Interim Measures unless such technologies are used to provide services to the public. They delegate some responsibility for regulation to existing industry regulators to strengthen regulation of AI as it applies within regulated sectors.

The Interim Measures are reasonably stringent in comparison to other jurisdictions, with AI service providers having to comply with the following key obligations (for example):

- conducting security assessments and filing certain information regarding the use of algorithms with the Cyberspace Administration of China where the AI service are capable of influencing public opinion or can mobilize the public;
- ensuring that models use lawful sources of data, do not infringe the intellectual property rights of others, complying with Chinese Mainland's legislations on personal information protection and employing measures to address illegal content;
- tagging the content created by generative AI;



- ensuring content generated by AI is in line with Chinese Mainland's core socialist values;
- not endangering the physical and psychological well-being of others;
- taking effective measures to improve quality of training data; and
- adopting measures increase transparency and to prevent discrimination.

The Interim Measures has a degree of extra-territorial effect, and provide that if AI services providers providing AI services from outside Chinese Mainland do not meet the requirements of the Interim Measures (or other Chinese laws), the Chinese authorities can take technical and other necessary steps to address such non-compliance (such as blocking access in Chinese Mainland to specific services).

The measures are explicitly described as 'interim' and so we expect continuing developments in Chinese Mainland's regulation of AI.

In the Hong Kong Special Administrative Region of the People's Republic of China ("**Hong Kong**"), no AI-specific legislation has been passed to date. However, the Office of the Privacy Commissioner for Personal Data published

the 'Guidance on the Ethical Development and Use of Artificial Intelligence' in August 2021. The [guidance](#) aims to assist organisations in complying with the Hong Kong data protection regime and sets out non-binding ethical principles for the development and use of AI such as accountability, human oversight, transparency and fairness. It includes a self-assessment checklist that organisations can use to determine whether the practices recommended in the guidance have been adopted in the development and use of AI.

The Hong Kong government also [published](#) an Ethical Artificial Intelligence Framework in June 2023, intended to establish a common approach and structure to govern the development and deployment of AI applications (including IT projects with predictive functionality or large language models with training data) by government bureaux and departments as well as for organisations in Hong Kong to take reference from when adopting AI and big data analytics in their projects. It adopts a principle-based approach and includes detailed models and explanations on principles within AI governance. The framework includes a template that organisations can use to reflect on their use of AI and assess the impact of the AI and whether it meets the ethical AI expectations of the framework. At this stage, the framework is not a piece of legislation but it serves as a useful indicator of the direction of travel for AI development in the region.

ASIA

Singapore appeared to take the lead in AI regulation through the launch of [AI Verify](#) in June 2022. This was the first AI governance testing framework and software toolkit in the region for companies to illustrate responsible AI use. AI Verify was expanded to the open-source community in June 2023. Despite this early development, it has since confirmed that it is taking a more cautious approach to regulating the technology, indeed the director of Singapore's Infocomm Media Development Authority confirmed in June 2023 that they currently have no plans to regulate AI in a formal manner.

In February 2023, South Korea passed a bill titled the 'Act on Promotion of AI Industry and Framework for Establishing Trustworthy AI'. If passed, the act will follow a similar approach to the EU's AI Act by categorising AI products based on risk and requiring 'high risk' AI to meet certain principles, including trustworthiness. The Act will not require government pre-approval for the development of AI technologies. There is no clear timeline yet, but the bill may pass through the legislative process at the end of 2023.



Despite a significant amount of discussion and debate on regulating AI, the US does not have a comprehensive federal law on AI, but rather a patchwork of sector-specific and state-level laws and regulations that may apply to certain aspects of AI. This patchwork of developments includes (but is not limited to):

- [the National AI Initiative Act](#) that came into force in 2021. This provides for a coordinated program across the federal government to accelerate AI research and application for economic prosperity and national security;
- a non-binding [blueprint for an AI bill of rights](#) which was introduced by the White House in October 2022, outlining protections that should be applied with respect to all automated systems that have the potential to meaningfully impact individual or community rights. The blueprint outlines five non-binding principles to guide the design, use and deployment of AI, including safe and effective systems, protecting against algorithmic discrimination, and data privacy;
- the National AI Commission Act which was introduced in June 2023 and, if passed, will establish a commission of 20 experts to draft a proposal for a comprehensive regulatory framework on AI; and



- the latest version of the National Institute of Standards and Technology's (NIST) [AI Risk Management Framework](#), which was released in January 2023. The framework aims to manage the risks of AI and improve the ability to incorporate trustworthiness into the design and use of AI products. The framework outlines how an organisation can implement continuous risk management throughout the lifecycle of an AI product;

Other developments include the US National Science Foundation's announcement in May 2023 of the establishment of [seven new AI research institutes](#) which will further the federal government's aim of advancing a cohesive approach to AI's opportunities and risks. In July 2023, representatives of seven tech companies including Google, Meta and OpenAI gathered at the White House and agreed to a set of [non-binding principles](#) for increasing the safety of AI technology, including third-party security checks and watermarking AI content.

Although limited, there are also some specific US state and City laws relating to AI in and around the workplace. For example:

- in 2019 (as amended in 2022), the Illinois General Assembly passed the [Artificial Intelligence Video Interview Act](#);
- in 2021, the New York [State Senate Bill S2628](#) amended the civil rights law in relation to electronic monitoring (which captures the use of AI); and
- Also in 2021, the New York Local Law [144 of 2021](#) was introduced to regulated automated employment decision tools.

Given this fragmented approach there are calls for federal regulation, and key sector regulators are beginning to respond. In particular, the Federal Trade Commission ("**FTC**") has addressed the importance of regulating AI through the lens of consumer rights, competition and unfair business practices (and the FTC's existing approach to regulation looks set to continue when addressing the challenges of new and disruptive AI technologies). The approach was neatly summed up by FTC chair Lina Khan (as part of a [joint statement](#) with the department of justice, consumer financial protection bureau, and US equal employment opportunity commission) as "there is no AI exemption to the laws on the books, and the FTC will vigorously enforce the law to combat unfair or deceptive practices or unfair methods of competition". On 30 October 2023 Biden's administration issued an Executive

On 30 October 2023 Biden's administration issued an Executive Order on AI encouraging the "safe, secure and trustworthy development and use of AI" (the "Order"). The Order contains a policy roadmap that encourages the responsible and effective use of AI and sets out disclosure requirements and industry-wide obligations for AI systems. The Executive Order focuses on eight key themes: AI safety and security, privacy, equity and civil rights, consumer benefits, workers, innovation and competition, global cooperation on AI, and the US government's use of AI.

Practically, although the Executive Order highlights the support and direction of the White House in this area, the execution of the Executive Order is largely dependent on the implementation of further rules and requirements by various US agencies and bodies (and so does not yet impose any new regulatory requirements perspective). For a summary of the key points raised by the Executive Order, please see our Insight [here](#).

The US is likely to face some challenges and trade-offs in developing and implementing a coherent and effective AI regulation, such as balancing federal and state interests, ensuring public trust and accountability, fostering innovation and competitiveness, and coordinating with international partners and allies. The US may also need to address some gaps and inconsistencies in its existing laws and regulations that may affect the development and deployment of AI, such as a federal data privacy regime.

Given the current fragmented approach in the US (currently evidenced by a patchwork of state data protection regulation), if individual states progress more quickly on a unilateral basis it may prove difficult for the US government to develop a federal framework.

On 1 and 2 November 2023, the UK Government held the first global AI Safety Summit (the "Summit") which brought together international governments, practitioners, civil society groups and experts in the research of AI. The aim of the Summit was to consider the key risks of AI systems, goals around mitigation of these risks and improvements in AI safety. In a significant step towards this objective, twenty-eight governments signed up to a world-first agreement on the safety of AI (the "Bletchley Declaration"). Signatories included countries that are at the forefront of developing AI technologies such as the UK, US, China, the EU, Japan, India, Saudi Arabia, Singapore, and UAE among other governments.

The Bletchley Declaration established risks anticipated with AI systems such as misuse, cybersecurity, disinformation, biotechnology, privacy concerns, and ability for bias. It also articulated a substantial risk on 'frontier' AI systems such as general-purpose AI models. The Bletchley Declaration highlights the potentially 'catastrophic' risk of AI and how classification and categorisation of risk is fundamental to mitigation. The Bletchley Declaration also addressed actors who are developing 'unusually powerful and potentially dangerous' AI systems and said they had a responsibility to ensure the safety of such developments, including secure testing measures.

In summary, the declaration reflects a commitment to harness the potential of AI while addressing its inherent risks through responsible practices. It underscores the importance of building a shared understanding of AI risks, formulating risk-based policies, and fostering scientific research. However, it remains to be seen how the Bletchley Declaration will feed into the development of AI regulation around the world. This will be useful to evaluate at the time of the next Summit, currently intended to be held in the Republic of Korea in May 2024.

CONCLUSION

17

We are beginning to see a patchwork of regulation, consultation, and guidance emerging across key technologically focussed jurisdictions around the world in respect of AI and emerging technologies. These approaches range from mechanisms to incubate the development of AI and establish its potential for commercial, social, and civil life, to prescriptive legislation designed to protect individuals from varying degrees of harm. Usurpingly, they are in no way harmonised.

Given the potential perceived threat of AI on the way we live and work, governments are heeding the call to come together to consider how they can work together to implement cross-jurisdictional frameworks and global principles, such as those highlighted in the [OECD report](#) for the G7 Digital and Tech Working Group. The UK is also looking to pave the way for international cooperation and more coordinated action, through the [AI Safety Summit](#) at Bletchley Park in November 2023. In relation to the Summit, UK technology Secretary Mitchell Donelan said “International collaboration is the cornerstone of our approach to AI regulation, and we want the summit to result in leading nations and experts agreeing on a shared approach to its safe use.”

Whilst these new developments offer exciting changes to the AI development space, existing legislation should not be an afterthought. Existing law (for example, around data protection, employment, financial services, and other areas) should be carefully considered by businesses considering developing or deploying AI systems and processes, in particular those with a global footprint.

For companies looking to exploit the technological advancements of AI and the associated benefits, they will need to monitor legislative developments and how guidance and regulation around the world continues to evolve. We expect this will be challenging, given the fast pace of innovation and the global target audience for many AI products in development. As a starting point to any compliance programme, companies first need to get a clear understanding of what AI the business is using and developing, and start mapping that against key considerations and requirements in an “AI asset register”. This will form the basis for the analysis of what new laws may apply, and how it may impact on the AI in current use or development.

CONTACT US



Simon Bollans
Partner, head of technology
London

T: +44 20 7809 2668
E: simon.bollans@shlegal.com



Melissa Forbes-Miranda
Partner
Dubai

T: +971 4 407 3904
M: +971 54 306 6543
E: melissa.forbes@shlegal.com



Zoe Zhou
Managing partner
Guangzhou

T: +86 20 8388 0590
M: +86 137 6085 6526
E: zoe.zhou@shlegalworld.com

To read more of our latest insights,
[visit our Technology Hub.](#)



Boriana Guimberteau
Partner
Paris

T: +33 01 44 15 82 78
M: +33 6 86 97 40 61
E: boriana.guimberteau@shlegal.com



Katherine Liu
Partner
Hong Kong

T: +852 2533 2717
M: +852 6193 5227
E: katherine.liu@shlegal.com

www.shlegal.com

© Stephenson Harwood LLP 2023. Any reference to Stephenson Harwood in this document means Stephenson Harwood LLP and/or its affiliated undertakings. Any reference to a partner is used to refer to a member of Stephenson Harwood LLP.