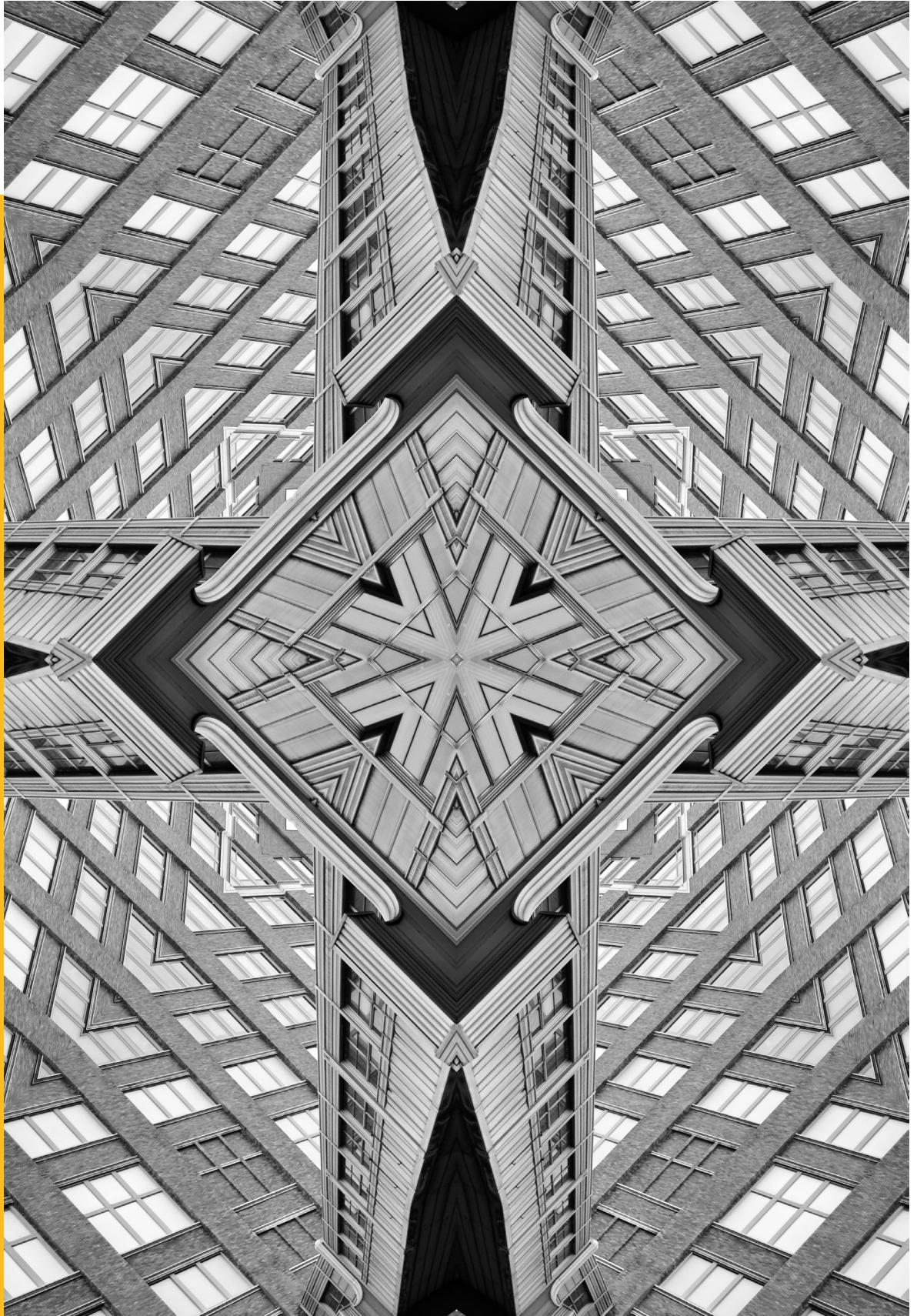


Occasional Paper



ISSUE NO. 486 JULY 2025

© 2025 Observer Research Foundation. All rights reserved. No part of this publication may be reproduced, copied, archived, retained or transmitted through print, speech or electronic media without prior written approval from ORF.

The European Union A.I. Act 2024: Understanding the Context and Exploring its Future

Siddharth Yadav

Abstract

The European Union (EU) Artificial Intelligence Act is a landmark regulation capable of propelling the bloc to the forefront of AI governance. Implementing a novel risk-based tiered approach, the Act has the potential to become a policy template for other states aiming to regulate the development and use of AI. Although the regulation represents a pioneering effort in governing a technology as impactful as AI, it suffers from certain gaps such as rigid categorisations and asymmetric risk-assessment standards. This paper contextualises this regulatory approach within the EU's broader digital policy landscape and analyses the strengths and weaknesses of the proposed risk-based framework.

The rapid development of artificial intelligence (AI) and generative AI platforms in recent years has served as a global wake-up call for the public and governments about the need to manage emerging, transformative technologies. At present, regulators around the world are working to formulate governance frameworks to tackle the highly dynamic field of AI. This dynamism is largely driven by the wide range of projections surrounding AI's future trajectory. Industry leaders such as Sam Altman, CEO, OpenAI, paint an optimistic picture, stating that AI development will lead to the dawn of "The Intelligence Age".¹ Others, meanwhile, highlight the dangers of AI and warn that, without regulation, an "AI Fukushima" may be inevitable.² The Centre for the Study of Existential Risk at Cambridge University, for instance, considers AI a potentially catastrophic threat to human civilisation.³

Bold statements on both technological utopias and dystopias make it difficult to develop a cohesive picture of a plausible future. Yet, regardless of the divergent perspectives, studies have indeed established the need for effective regulation to tackle AI-driven issues such as labour displacement, asymmetric wealth transfer between nations, violations of civil and human rights, and environmental impact.⁴

To address these issues, governments have released guidelines for responsible AI development and participated in multilateral discussions on AI regulation. The results have largely taken the form of prescriptive ethical guidelines and transparency obligations, instead of legally binding rules. In 2021, all 193 member states of the United Nations Educational, Scientific and Cultural Organization (UNESCO) adopted the 'Recommendation on the Ethics of Artificial Intelligence', which provided a roadmap for developing ethics guidelines for AI.⁵ In 2023, the Indian government published a proposal for the Digital India Act (DIA),⁶ informed by the governance principles outlined in NITI Aayog's 2018 National Strategy for Artificial Intelligence.⁷ The DIA proposes a "principles and rules-based" approach, guided by openness, safety and trust, accountability, quality of service, and redress mechanisms.⁸

In the United States (US), the Trump administration has released an Executive Order on ‘Removing Barriers to American Leadership in Artificial Intelligence’, aimed at supporting and deregulating the AI industry.⁹ This pro-innovation deregulatory approach is in contrast to the earlier Executive Order issued by the Biden administration that emphasised safety, accountability, and transparency in AI development.¹⁰ The shift away from AI safety principles is further underscored by the scrapping of the AI Risk Management Framework and Risk Management for Artificial Intelligence and Human Rights developed by the National Institute for Standards and Technology (NIST) and the US Department of State, respectively, under the previous US administration. To date, there has been no federal policy on AI regulation despite the introduction of over a thousand bills on AI policy in the US Congress.¹¹ The White House, however, has opened the federal AI policymaking process to stakeholder consultation and is set to release an AI Action Plan.

In November 2023, the United Kingdom (UK) hosted the Bletchley Park AI Safety Summit. Subsequently, the government released its AI governance consultation outcome in February 2024, recommending five cross-sectoral principles: safety, security, and robustness; appropriate transparency and explainability; fairness; accountability; and governance, contemptibility, and redress.¹² Most recently, in early 2025, France and India co-hosted the Paris AI Action Summit (PAIAS), where the declaration on “open, inclusive, transparent, ethical, safe, secure, and trustworthy” AI development was signed.¹³ PAIAS also marked a point of global divergence on AI regulation with the US and the UK abstaining from signing the declaration while announcing pro-innovation deregulatory approaches that align with their national interests. While European Union (EU) member states unanimously endorsed the PAIAS pledge, the EU simultaneously signalled a tilt towards deregulation. At the Summit, European Commission (EC) President Ursula von der Leyen stated that the EU needed to cut regulatory red-tape and announced an investment package of 200 billion euros to accelerate AI development and adoption.¹⁴ As AI platforms become increasingly capable, their strategic significance for the global economy rises in proportion.

Introduction

In this policy landscape, the EU Artificial Intelligence (EU AI) Act is an important piece of legislation that establishes the world's first legally enforceable regulatory framework for AI systems within the EU. In the global technology sector, the EU has come to occupy a leadership position in technology regulation. The EU General Data Protection Regulation (GDPR), implemented in 2018, quickly became a clarion call for data protection and the gold standard for measuring government policies globally.¹⁵ As the first enforceable legislation for AI development, the EU AI Act has the potential to become a governance template for other nations. To understand its relevance and potential impact, this paper examines the EU AI Act and its risk-based approach to regulating AI systems according to the level of risk they pose.

The following sections explore the AI Act within the broader context of the EU's Digital Strategy. The first section will focus on various directives and policies adopted by the EU throughout the past two decades that established the regulatory priorities outlined in the AI Act. The second section will explore the AI Act itself as well as the policy discussions leading to its ratification. The last section will analyse how industry leaders, and expert and civic groups have responded to the AI Act and its various provisions.

Contextualising the EU's Digital Strategy

In August 2024, the EU released the Artificial Intelligence Act, the world's first formalised legal and regulatory framework for governing AI systems.¹⁶ Initially proposed by the European Commission (EC) in April 2021, the AI Act came into force following a series of dialogues with member states, citizen groups, and industry consortiums.¹⁷ The AI Act is designed as a framework to govern AI systems across the public and private sectors. However, systems operating in domains such as national security and law enforcement remain exempt under specific conditions.¹⁸ The AI Act is the latest iteration of the broader Digital Strategy adopted by the EU over the past 25 years to harmonise the digital legislative frameworks of its member states and formulate a unified governance and enforcement mechanism.¹⁹

e-Commerce Directive

The introduction of the e-Commerce Directive in 2000 was the first notable step taken by the EU to develop its digital strategy for the 21st century.²⁰ The Directive was designed to “create a legal framework to ensure the free movement of information society services between Member States,”²¹ with its scope remaining limited to online information services, online retail, online advertising, online professional services, and online contracting.²² Articles 14 and 15 of the Directive set liability rules for online service providers, reflecting a preference for an outcome-based approach over preventive measures. Article 14 requires them to remove access to illegal content upon acquiring knowledge of it, exempting them from liability if they act promptly;²³ and Article 15 prevents EU member states from imposing general monitoring obligations.²⁴

The e-Commerce Directive acted as a cornerstone of the EU's digital strategy in the 2000s, as highlighted in the EU's i2010 policy framework released in 2005.²⁵ However, the preconditions for the risk-based approach of the EU AI Act subsequently emerged from an increased emphasis on end-user safety and broad-spectrum consumer rights protections in the EU's Digital Agenda 2020 and accompanying policy reforms. This policy framework aimed to facilitate the secure and efficient movement of digital goods and services across national borders, improvement of the quality of

Contextualising the EU's Digital Strategy

networks and services, establishment of a single consolidated EU market, and protection of human and civil rights.²⁶ Several steps concerning consumer rights were taken.

The GDPR and the Digital Single Market

The EU's preventive approach is exemplified in Regulation (EU) 2016/679 or the General Data Protection Regulation (GDPR) released in 2016.²⁷ The GDPR was introduced based on the stipulation in Article 8 of the EU Charter of Fundamental Rights on the protection and responsible processing of EU citizens' personal data.²⁸ Signalling the EU's precautionary approach, the GDPR puts forth regulations on the collection and processing of personal data,²⁹ enshrines rights of data subjects, and establishes the principle of "data protection by design" for online service providers.³⁰

Subsequently, the Digital Single Market (DSM) Directive, released in 2019, addressed the deregulatory approach of the e-Commerce Directive regarding limited liability of online service providers. Article 17 of the DSM Directive specifies that online service providers must take steps to acquire authorisation from copyright holders before making user-generated content available in the public domain to prevent copyright infringement.³¹ As legal scholar Gerald Spindler notes, this provision identifies service providers as active participants as their operation includes "making available to the public" content that may be protected by copyright, thus holding them liable for copyright infringement taking place on their platforms.³² Although this approach imposed on large service providers the monumental task of filtering and categorising millions, if not billions, of user-generated items posted on their platforms every day, it highlighted a policy shift towards a precautionary approach in the EU's Digital Strategy.

Digital Services Act Package

The next set of regulations aimed at protecting the rights of users was the Digital Services Act Package, comprising the Digital Market Act (DMA) and the Digital Services Act (DSA), adopted by the European Parliament in 2022 and released as a single set of rules applied to the EU market.³³

Contextualising the EU's Digital Strategy

The DMA aims to identify ‘gatekeepers’ in the digital market and prevent the creation of monopolies and so-called ‘digital walled gardens’ by enshrining a set of compliance obligations.³⁴ In terms of impact, the European Commission initiated proceedings against US-based companies Alphabet, Apple, and Meta in March 2024, based on concerns of non-compliance with the obligations specified in the DMA.³⁵

While the DMA established the liability of ‘gatekeepers’, the DSA sought to introduce liability standards that the e-Commerce Directive had left unaddressed.³⁶ Media scholar Amélie P. Heldt has identified disinformation and its associated harms as primary causes necessitating the ratification of stricter liability standards, as presented in the DSA.³⁷ Learning from the criticisms of legislations proposed by member states, the European Commission opted to take a more controlled approach with the DSA.³⁸ It did not add general monitoring mandates but expanded the operating standards for service providers regarding the takedown of illegal content.³⁹ The DSA demanded that service providers make online recommender systems more transparent. Additionally, the DSA introduced compliance obligations for risk management and assessment for very large online platforms (VLOPs).⁴⁰ Online traders and sellers are now also subject to wide-ranging due diligence and transparency obligations.⁴¹

The EU’s regulatory approach, in place till it was superseded by the AI Act, provides the rationale behind the outcome-oriented and risk-based framework within which the EU AI Act was crafted. First, the amendments and reforms made to the provisions laid out in the e-Commerce Directive of 2000 targeted specific domains such as copyright (through the DSM strategy), indicating a reluctance to overregulate and a collaborative attitude towards the industry. Through the DSA, the liability system for service providers also focused on improving end-user redress mechanisms and establishing compliance regimes, instead of imposing proactive content monitoring mandates.⁴² Second, compared with the legislations passed by countries like Germany and France that were criticised as attempts at “overpolicing”, the European Commission took an ostensibly balanced approach with the DSA limiting new obligations to online marketplaces.

Contextualising the EU's Digital Strategy

Third, the DSA and the GDPR were enshrined based on the principles stated in the EU Charter of Fundamental Rights. The consumer-focused approach of the DSA, the GDPR, and even the DMA makes a strong reappearance in the risk-based categories outlined in the 2024 AI Act and the preceding discussions between policymakers, expert groups, business consortiums, and civic groups.

Objectives of the AI Act

In 2020, the White Paper on Artificial Intelligence, published by the European Council, highlighted the risks posed by AI systems to the fundamental rights of EU citizens outlined in the EU Charter.⁴³ Additionally, the paper emphasised the need for increased investment in the AI sector to support EU's competitiveness against the tech sector in North America and Asia.⁴⁴ Since 2018, the need for increased investment and broader socioeconomic adoption of AI has been a central theme in AI policy discussions, as outlined in the European Commission's Coordinated Plan on Artificial Intelligence.⁴⁵ To address these issues, the EU AI Act aims to establish a horizontal regulatory framework across the EU to prevent the fragmentation of the EU single market due to the distributed nature of AI development and deployment.⁴⁶ The AI Act is based on provisions from the Treaty on the Functioning of the European Union (TEFU), specifically Article 16 that enshrines the protection of EU citizens' personal data and Article 114 that authorises the EU legislator to harmonise national laws and regulations for establishing and ensuring the functioning of the EU internal market.⁴⁷

Prior to the AI Act, the European Commission used a 'soft-law' approach based on the 2019 Ethics Guidelines on Trustworthy AI, prepared by the High-Level Expert Group on Artificial Intelligence (AI HLEG), set up by the Commission in 2019. The AI HLEG proposed seven non-binding guidelines: human agency and oversight; technical robustness and safety; privacy and data governance; transparency; diversity, non-discrimination, and fairness; societal and environmental well-being; and accountability.⁴⁸ Based on the AI HLEG guidelines, public consultations, and an impact assessment study published in 2020,⁴⁹ the European Commission adopted a tiered risk-based approach for regulating AI systems.

Regulating AI Through a Risk-Based Approach

The AI Act presents itself as adopting "a clearly defined risk-based approach" to enshrine "proportionate and effective" set of rules for AI systems.⁵⁰ Novelli et al. have noted that risk-based approaches involve

three phases: risk assessment and categorisation; impact assessment; and risk management.⁵¹ By adopting a risk-based approach, the regulation first considers risk as a regulatory concern.^a The risk categories of the AI Act are designed by calculating the foreseeable risk of AI systems causing degrees of harm to the health, safety, and fundamental rights of individuals. Based on the calculation, AI systems seen as posing unacceptable risk are prohibited. Systems seen as posing tacitly acceptable risks are further categorised into ‘high-risk’, ‘moderate risk’, and ‘limited risk’ systems.

In terms of risk management, law and economics scholar Cary Coglianese states that risk-based approaches can broadly take four forms: eliminating all risk; reducing risk to an acceptable level; reducing risk till costs are feasible; and balancing risk reduction with cost of regulation.⁵² As the following sections will show, the AI Act relegates the most ubiquitous types of AI systems (e.g., LLMs and chatbots) to ‘moderate risk’ and ‘limited risk’ categories with lighter compliance obligations, and reserves the higher categories with stricter rules for systems posing the risk of severe harm. The tiered design of the AI Act helps reduce unnecessary regulatory burden and cost of regulation. By enshrining a ‘proportionate and effective’ framework, the AI Act seeks to avoid overregulation and create a conducive environment for innovation.

Risk Category: Unacceptable Risk

The definition of AI systems in the Act is based on their ability to operate with varying levels of autonomy from human intervention due to their adaptiveness and self-learning capabilities.⁵³ The Act further characterises AI systems as having the ability to derive models or algorithms and draw inferences through the process of “obtaining the outputs, such as predictions, content, recommendations, or decisions, which can influence physical and virtual environments.”⁵⁴ AI systems that fall within the purview of the Act may either be products in themselves or components of a product.⁵⁵

As the first priority, the regulation explicitly bans the use of AI systems that pose unacceptable risk to the safety, livelihoods, and rights of EU people.⁵⁶

a The AI Act defines ‘risk’ as the likelihood of foreseeable or possible harm transforming into actual harm, whether it is physical, economic, or psychological.

The scope of this category extends to systems that employ subliminal techniques and are used for social scoring, and those designed to conduct real-time biometric identification in public spaces.⁵⁷ The regulation defines subliminal techniques as “manipulative or deceptive techniques that subvert or impair person’s autonomy, decision-making or free choice in ways that people are not consciously aware of.”⁵⁸ For systems employing subliminal techniques that pose unacceptable risk, the regulation cites examples such as ‘machine-brain interfaces’ and virtual reality, revealing a future-oriented and anticipatory approach towards dangers that are in the speculative realm so far.⁵⁹ While efforts to minimise such dangers should be appreciated, this categorisation does raise the question of whether the AI Act is instituting unnecessary roadblocks for platforms (‘machine-brain interfaces’ in this case) that are still in their infancy.

Another type of AI systems that are prohibited under the AI Act involves those designed for ‘social scoring of natural persons’,⁶⁰ likely an implicit reference to the social scoring system in China that has received widespread scrutiny.⁶¹ The regulation bans AI systems that use data points on “social behaviour in multiple contexts or known, inferred or predicted personal or personality characteristics” to classify individuals or groups in ways that can lead to detrimental outcomes.⁶² In a move reminiscent of the 2002 sci-fi film, *Minority Report*, the AI Act also prohibits predictive assessments by AI regarding a person’s potential future criminal behaviour or their use in legal proceedings against persons who have not actually committed any crime.⁶³ The prohibition of social scoring is related to the unacceptable risks of mass surveillance, prompting a ban on the indiscriminate scraping of facial images from the internet or CCTV footage to expand facial recognition databases.

High-Risk Systems

Following the prohibition of AI systems posing unacceptable risk, the AI Act enshrines the classification of high-risk AI systems. Such systems are distributed along eight categories.⁶⁴ These include biometric identification technologies like remote identification and emotion recognition, permitted for law enforcement under strictly defined conditions. High-risk classification also covers AI used in critical infrastructure management,

educational and vocational training contexts, and employment contexts such as candidate recruitment, evaluation, and monitoring. Systems determining access to public services, benefits, emergency services, credit scoring, and insurance are also classified as high-risk. Additionally, law enforcement tools for victim or crime assessment, polygraph-like tools, evidence evaluation, and risk assessments also fall in this category. AI systems used in migration and border control for risk assessment and application processing as well as in judicial systems for legal research or influencing elections and voter behaviour are included.⁶⁵ General-purpose AI models trained with over 10^{25} floating point operations (FLOPs) may also be deemed high-risk under the AI Act.⁶⁶

Limited-Risk and Minimal-Risk Systems

The AI Act enshrines categories of ‘limited risk’ and ‘minimal risk’ for AI systems that do not pose any significant foreseeable threat to the rights and freedoms of people. The ‘limited risk’ category includes systems like chatbots, synthetic content such as deepfakes, AI-edited or AI-altered content, and biometric and emotional identification systems that require user consent.⁶⁷ As the name suggests, ‘minimal risk’ systems are those that pose no meaningful threat to users. This category includes tools such as spam filters and ad blockers.⁶⁸

Exemptions and Compliance Obligations

The aforementioned use cases of high-risk AI systems (stated in Annex III of the AI Act) have some built-in exemptions for systems that only perform procedural tasks and inconsequential actions.⁶⁹ Providers of AI systems that do not meet the exemption criteria have to meet extensive compliance requirements. General-Purpose Artificial Intelligence (GPAI) models with high compute are subject to similar requirements and must additionally provide appropriately detailed summary of the data used for training the models. Compliance requirements for AI systems in the limited-risk and minimal-risk categories are confined to transparency obligations. For generative AI models and systems used to generate synthetic content, providers must watermark the AI-generated or AI-edited content, inform users that they are interacting with an AI system, and acquire user consent when deploying systems that require biometric identification.⁷⁰

Clear Objectives and Regulatory Priorities

As policy scholar Daniel Mügge notes, the dynamism of the AI sector forces policymakers to rely on speculative projections about the risks and benefits of AI systems.⁷¹ For instance, despite being the global leader in terms of research and innovation, the US still lacks a federal regulation on AI. The AI regulatory landscape in the US consists of a patchwork of various state-level regulations. In contrast, the EU AI Act, with its risk-based approach, provides a clearer set of regulatory objectives and priorities for the EU to harmonise and “rationalise” government interventions.⁷² Additionally, as previous sections have shown, the AI Act uses a balanced and proportionate approach for risk identification and classification, and for distributing regulatory burden, allowing for an efficient use of government resources.⁷³ The Act has also been recognised for promoting the concept of collaborative governance for AI systems.⁷⁴ Efforts on holding multistakeholder public discussions and formation of diverse expert groups ahead of the release of the regulation increase confidence and trust in the regulatory process.

Another strength of the Act is its realist stance on the global AI race. Although the regulation repeatedly and consistently emphasises risks of AI and prevention of harm to individuals, the policy discourse surrounding the AI Act has been focused on ensuring the “broadest possible uptake of AI in the [EU] economy.”⁷⁵

The AI Act correctly leaves room for further amendments and iterations to its risk categories based on annual reviews, avoiding misidentification of AI systems and unnecessary regulatory hurdles for AI developers. However, as the following sections will show, the realist stance of the AI Act also leaves room for drawbacks, such as rigid risk categories and lacklustre risk-benefit analyses, that question the Act’s purportedly proportional approach. Additionally, implementation challenges can hinder the effectiveness of the regulation.

Lack of a Cohesive Regulatory Vision

As mentioned earlier, the AI Act is an attempt to implement a “proportionate and effective” regulatory approach designed to responsibly promote competitiveness in the EU tech sector and reduce regulatory burden on AI developers. However, stimulating continent-wide innovation requires a cohesive economic approach complementing regulatory measures. This includes establishing adequate implementation and compliance mechanisms, securing risk-tolerant funding for startups and small and medium-sized enterprises (SMEs), investing in infrastructure to support the compute and energy requirements of AI developers, and public consultations that equitably represent stakeholders’ interests. Although the EU AI Act is an attempt at formalising a streamlined horizontal approach to prevent regulatory overlap and friction between EU member states, critics have noted that the Act in conjunction with guidelines like the 2025 AI Code of Practice (the Code) and the 2022 AI Liability Directive falls short of establishing an innovation-friendly framework.⁷⁶

The Code has received pushback from AI sector leaders and rights groups for proposing untenable compliance requirements and underrepresenting civil society stakeholders and SMEs during consultation rounds.⁷⁷ Moreover, the Code (in addition to the AI Act) will be applicable to AI systems if their output is used within the EU, irrespective of the location of their developers or deployers.⁷⁸ Consequently, the US Mission to the European Union has characterised the Code as a proxy tariff on US AI firms that are leading global AI development due to its cross-border scope.⁷⁹ Unlike the GDPR that popularised the Brussels Effect⁸⁰ and positioned the EU as a leader in global tech regulation landscape, attempting to control the trajectory of AI development—wherein industry leaders are almost exclusively American or Chinese firms⁸¹—may have little impact beyond delaying the entry of foreign AI models into the EU market and pushing native enterprises to more permissive markets.

The implementation of the AI Act has also been experiencing roadblocks across member states as the European Committee for Standardization (CEN) and European Committee for Electrotechnical Standardization (CENELEC) reported that the development of technical standards to be

used by AI developers for compliance—expected to be released in August 2025—will likely be delayed by a year.⁸² The AI Liability Directive, proposed by the European Commission in 2022 to harmonise civil liability rules for AI-related harms, has also been rescinded in 2025 due to disagreements among EU member states. Further, the appointment of fundamental rights bodies by EU member states mandated in the AI Act has also been running behind schedule, signalling a lack of shared priorities.⁸³

Generality and Ambiguity of Risk Categories

The primary drawback of the AI Act stems from its effort to harmonise regulations across the EU. The optimisation towards harmonisation limits the regulation in many cases to the use of blanket categories and ambiguous language. For instance, the definitions of ‘harm’ and ‘high-risk’ remain unclear, given that Annex III identifies entire fields of application as high-risk while simultaneously exempting AI technologies solely used for military purposes.⁸⁴ Veal and Borgesius have noted the problematic approach to assessing harm caused by AI systems;⁸⁵ the Act sees harm as an outcome occurring in a small and defined timescale caused by a single or a small set of definitive events. Instead, they argue, “harm can accumulate without a single event tripping a threshold of seriousness, leaving it difficult to prove.”⁸⁶ Such “cumulative harms” often depend on several factors such as personal inclinations of users and algorithmic optimisations towards user engagement.⁸⁷

A narrowly defined concept of harm thus dilutes enforcement by ignoring long-term behavioural distortions. An example of this issue can be found in Annex III of the Act, which describes the scope of the ‘high risk’ category. The Act states that systems intended to be used for affecting “voting behaviour” will be categorised as ‘high risk’.⁸⁸ However, the Act does not provide any examples of use cases or operational scenarios depicting how AI systems may affect people’s voting behaviour. This omission raises questions about such systems as algorithmic recommender systems.

The AI Act classifies AI systems used for creating deepfakes and synthetic information and recommender systems as limited-risk and minimal-risk systems, respectively. However, algorithmic systems deployed by social

media platforms have been playing an increasingly important role in sociopolitical discourse; their role in exacerbating political polarisation has been acknowledged by researchers.⁸⁹ For instance, Elon Musk's social media platform X has been identified as playing a non-trivial role in affecting the 2024 US presidential election as well as exacerbating political tensions in foreign nations.⁹⁰

Algorithmic recommender systems are optimised to maximise engagement and outrage is an excellent indicator of engagement.⁹¹ If a malicious actor does intend to deploy AI systems in this manner, as Veal and Borgesius have argued, the cumulative effects of such systems may manifest in timescales beyond the scope of the AI Act.⁹² This lack of clarity in the Act raises questions about the regulatory stance on AI systems used on social media platforms.

Arguably, the most striking example of blanket categorisation in the absence of empirical evidence is the prohibition of subliminal manipulation techniques in the AI Act. The regulation specifically mentions platforms like brain-computer interfaces (BCIs) and virtual reality (VR) as posing unacceptable risk. The risk factor is the ability of BCIs and VR to manipulate end users, channelling stimuli outside of conscious perception.⁹³ Although attempting to mitigate such entrenched risks is useful, it should be noted that BCI and VR platforms that could pose such risks do not exist yet.

The efficacy of advanced BCI platforms currently being developed is limited and highly specific to neuromuscular rehabilitation and restoration of motor functions.⁹⁴ A promising modality for BCI development gaining traction today focuses on vision restoration.⁹⁵ Any BCI technology for this purpose will require a two-way connection between a machine and the nervous system, specifically the visual cortex. The presence of this connection also opens a possible avenue for subliminal manipulation. This possibility has led to the emergence of policy considerations regarding the protection of 'neuro rights'.⁹⁶ Simultaneously, the Act relegates AI tools like recreational chatbots to limited- and minimal-risk categories. However, recent events, such as a lawsuit against the company Character AI over acute mental health risks posed by its chatbot, indicate that AI manipulation may not be limited to subliminal stimuli.⁹⁷ Such cases, while

specific, highlight the presence of unprecedented risk factors and unknown unknowns, necessitating a more flexible categorisation to achieve the AI Act's objective of proportionality.

A responsible approach to handling such specialised domains will require ongoing dialogue and evaluation of platforms on a case-by-case basis. Instead, the AI Act paints with a broad brush, without clarifying any parameters or criteria for what causes a risk factor to become unacceptable.⁹⁸ A premature prohibition of modalities that are central to technology platforms will do little beyond stifling innovation. This limitation in the AI Act also extends to VR, identified as a platform that can pose unacceptable risk, despite its market penetration not yet reaching a level that warrants such concerns.

The two instances mentioned above raise the broader question of why this risk category was outlined in the first place, if it specifies use cases that have not developed to cause any actual harm. One possible explanation is that the objective of the AI Act is to present itself as a safeguard against unethical practices, fostering a more receptive environment for public and private investments in AI.⁹⁹ Two interrelated elements are relevant in this respect: the risk acceptability and the broader ethical principles of the AI Act.

Asymmetric Ethical Standards

The concept of risk acceptability is crucial for facilitating discussions on AI trustworthiness. This approach was adopted by the AI HLEG while developing ethical principles to guide AI regulation in the EU. Simultaneously, applying the risk acceptability criteria first necessitates identifying the stakeholders for whom the risk should be deemed acceptable. An obvious explanation is that potential risks should primarily be acceptable to the public or end users as long as they align with the law. However, as policy scholars Laux et al. have noted, trustworthiness and risk regulation are multilayered governance factors. They argue that the trustworthiness of an AI regulation depends on “causal relationships”, such as trust in public institutions, confidence in the regulatory process, belief in a regulation's objectives and effectiveness, and society's overall attitude towards technology.¹⁰⁰

The AI Act adopts a paternalist (as opposed to participatory) approach to risk regulation by establishing an “epistemic asymmetry of laypeople versus experts” that could lead to political asymmetry in the future.¹⁰¹ In other words, the risk-based approach of the Act creates room for misalignments between public perception of risk and expert assessments.

Further asymmetries emerge if the AI Act is evaluated in relation to the role of the EU in the global AI race. Despite being home to some of the largest Western economies, none of the world’s 10 biggest AI companies are headquartered in the EU.¹⁰² The 2024 Mario Draghi Report on European Competitiveness, which argued for an aggressive change in the industrial policy of the Union, highlighted the EU’s lagging position in AI development specifically and technology in general.¹⁰³ The report also highlighted the lack of consistent financing for startups and frontier tech SMEs, onerous compliance regimes (including the AI Act), migration of startups to less risk-averse capital markets like the US, and Europe’s dependence on foreign tech infrastructure and supply chains as significant bottlenecks for its innovation ecosystem.¹⁰⁴ Following the release of the report, EU officials have publicly supported and advanced initiatives like the EU-HPC (High-Performance Computing) and EuroStack to accelerate technology development in the region.¹⁰⁵

While the aforementioned initiatives are new, the inertia in the EU’s AI industry was noted in the 2020 White Paper on Artificial Intelligence, published by the European Commission, that acknowledged the EU’s unfavourable position. Nevertheless, the EU has taken the lead in developing governing principles for AI systems, based on ethical principles, while avoiding overregulation. To advance this objective, as mentioned in previous sections, the European Commission established multistakeholder groups like the AI HLEG that formulated ethical guidelines for ‘trustworthy AI.’¹⁰⁶ However, the guidelines were criticised as the composition of the AI HLEG was seen as overrepresenting industry members and underrepresenting academia and civil society.¹⁰⁷ Thomas Metzinger, an ethical philosopher and former member of the AI HLEG, has characterised guiding principles like trustworthiness as attempts at “ethics washing”;^b he also termed AI ethics

^b Ethics washing, similar to the concept of greenwashing, refers to performative actions by companies and regulators that signal a commitment to governing principles like trustworthiness to allay public scepticism. Metzinger’s comments were made in reference to the emphasis placed on governing principles like trustworthiness in the EU AIA while prohibitions against controversial AI platforms like autonomous weapon systems were diluted seemingly due to industry intervention. See: <https://www.cambridge.org/core/journals/asian-journal-of-law-and-society/article/on-the-governance-of-artificial-intelligence-through-ethics-guidelines/992BD33CA7CBBE83E2FBBF6B0179896C>

debates held in the EU as “marketing ploy[s]” to facilitate the creation of future markets for AI developers.¹⁰⁸

The regulatory scope of the AI Act also presents indicators of possible epistemic and political asymmetry. In 2022, the European Parliament Committee on Legal Affairs (CLA) urged the Committee on the Internal Market and Consumer Protection and the Committee on Civil Liberties, Justice, and Home Affairs to remove regulatory exemptions for AI systems used by the national security and military apparatus of EU member states from the AI Act.¹⁰⁹ The Act repeatedly emphasises the protection of consumer rights, civil liberties, fundamental rights, and principles like transparency and trustworthiness as a priority. However, the European Parliament’s generalised approach to maximising harmonisation prevented it from adopting the amendment proposed by the CLA.

When the decision to not adopt the CLA recommendation is seen in conjunction with the previously mentioned prohibition of certain AI systems (concerning subliminal manipulation via BCIs and VR) that have not evolved enough to pose a credible threat to the public, the AI Act falls short of presenting a principled regulatory vision. Nevertheless, shortcomings of the regulation can be addressed through the provisions of the final text of the EU AIA, such as in articles concerning review and evaluation, that mandate the European Commission to annually evaluate the list of prohibited systems, high-risk systems, add new use-cases, modify existing categories and remove categories.¹¹⁰ Given that AI developers and deployers had been given 36 months to set up compliance mechanisms, definitive conclusions on the efficacy of the AI Act will have to wait until impact assessments are conducted in the future.

Recommendations and Conclusion

The EU AI Act is a pioneering regulatory effort to control the unpredictable dynamism of AI development. Being the first legally binding framework for governing AI, the Act is emblematic of the EU's leadership position in global technology governance. Resulting from over two decades of Digital Strategy evolution, the risk-based approach of the AI Act—much like the GDPR—establishes a potential policy template for other governments. By enshrining a tiered classification of AI systems, ranging from unacceptable risk to minimal risk, the Act can increase enforcement efficiency by clarifying and rationalising regulatory priorities.

The AI Act should also be commended for attempting to balance innovation with safety by identifying risks to fundamental rights and consumer rights from AI systems as a regulatory concern. However, the AI Act is not without its limitations: the narrow definition of 'harm' in the AI Act limits harm assessment to events occurring in a short timescale triggered by definitive events;¹¹¹ the risk categories enshrined in the Act at times rely on speculative harms and untested scenarios, potentially imposing unnecessary regulatory obstacles on nascent technologies by anticipating use cases without empirical evidence; and the AI Act's expert-driven tiered approach can create asymmetries between public perception and official assessments of acceptable risk.¹¹²

The following recommendations are designed to address the aforementioned issues with the AI Act.

a. Refine the definition and assessment criteria of 'harm'.

The AI Act needs to address cumulative harm caused by AI systems, instead of relying solely on its current event-based interpretation, which identifies harm through specific incidents. Timescale issues can be addressed through operational scenarios and illustrative use cases for prohibited and high-risk AI systems, as released by the EC, helping clarify regulation enforcement and reduce ambiguity in AI development. Given that the AI Act sees interference in election processes and voting behaviour as a high-risk area, the definition of harm should include cumulative and longer-

Recommendations and Conclusion

term effects on individuals, groups, and democratic processes. Additionally, the EC should establish multidisciplinary forums—in addition to the AI HLEG—comprising economists, human rights scholars, social scientists, and technologists to delineate gradations of harm (such as cumulative or instantiated) and proportionate regulatory responses for future iterations of the Act.

b. Clarify risk category thresholds.

Future iterations of the AI Act need to address certain AI use cases mentioned in the ‘unacceptable risk’ category. Risks posed by AI systems should not be underestimated, particularly when considered in conjunction with invasive technological platforms like BCIs and VR. However, the ambiguous language of the regulation creates unnecessary uncertainties and regulatory walls for nascent technologies that can make contributions to society, particularly through medical applications, for example. If certain AI systems or components are seen as posing unacceptable risk, the EU needs to provide exhaustive explanations and risk benchmarks for such classification through expert and civic consultations. To ensure that the proposed regulatory methods are feasible for mitigating potential future harms, the EU needs to develop evidence-based methodologies to explicate when and how AI systems cross the risk thresholds. Establishing more rigorous and transparent standards for risk assessment and management would allow the EU to avoid overregulation and address tangible risks posed by AI systems.

c. Improve stakeholder participation and diversity.

The explicitly risk-based approach of the AI Act requires standards of risk acceptability to be established in a democratic and transparent manner. A top-down paternalist approach can create misalignments between public expectations and governance mechanisms.¹¹³ The EU needs to move towards a more participatory approach that incorporates public opinion and diverse viewpoints in the regulatory process. This can be achieved by establishing engagement mechanisms for non-industry stakeholders to foster inclusive and democratic dialogues. Beyond expert consultations, the EC should establish deliberation forums for citizens,

Recommendations and Conclusion

consumer advocacy groups, academia, and civil society organisations to raise concerns and interrogate the trade-off between innovation and foreseeable risk. Moreover, online platforms and forums can be created to solicit public feedback, understand concerns, and inform the public about emerging AI technologies and upcoming regulatory changes. By investing in transparent and participatory governance processes, the EU can better align regulations with public expectations to reinforce the legitimacy of its AI governance framework.

d. Create a participatory global regulatory ecosystem.

Since AI development and its associated risks are rarely confined to national borders, the EU should proactively coordinate with international partners and standard-setting bodies. This can include establishing a liaison office, in addition to the proposed EU AI Office, dedicated to technology diplomacy and cross-border coordination efforts. Given the EU's pioneering efforts in technology regulation, it could spearhead the creation of international working groups, along with emerging players in the field like India and the UAE, to harmonise transparency standards, align classifications of risk, develop inclusive and diverse ethical standards, and establish cross-border end-user redress mechanisms and interoperability standards for compliance tools.

Collaboration between the EU and emerging frontier technology markets like the UAE and India can be supported through existing projects like the India–Middle East–Europe Economic Corridor (IMEC) that targets connectivity as a core mission.¹¹⁴ The co-hosting of 2025 PAIAS by France and India further indicates Europe's interest in liberalising trade, connectivity, and technology diplomacy with India. As part of its Middle East engagement, the PAIAS also saw the signing of MoUs between France and the UAE for investments in AI and energy infrastructure.¹¹⁵ The three jurisdictions are also more aligned on responsible AI development, compared with the US and the UK, which have shifted focus away from safety and ethical considerations. Using existing alliances and strategic platforms for creating a globally united front through common reference points and interests can help reduce conflicting rules across jurisdictions and encourage responsible innovation.

Recommendations and Conclusion

Adopting the recommendations mentioned above will allow the AI Act to adapt to evolving social-technical realities. The iterative design of the AI Act, with annual review processes and regular impact assessments, will help alleviate concerns regarding definitions and enforcement mechanisms, and enable a fair and equitable distribution of regulatory burdens. The possibility of the AI Act becoming a regulatory template for other jurisdictions will ultimately depend on whether the regulation is able to evolve in conjunction with public sentiment and industry changes. [ORF](#)

Siddharth Yadav is Fellow, *Emerging Technologies*, ORF Middle East.

- 1 Sam Altman, “The Intelligence Age,” September 23, 2024, <https://ia.samaltman.com>.
- 2 Ian Sample, “‘An AI Fukushima is Inevitable’: Scientists Discuss Technology’s Immense Potential and Dangers,” *The Guardian*, November 22, 2024, <https://www.theguardian.com/science/2024/nov/22/an-ai-fukushima-is-inevitable-scientists-discuss-technologys-immense-potential-and-dangers>.
- 3 “Risks from Artificial Intelligence,” Centre for the Study of Existential Risk, University of Cambridge, <https://www.cser.ac.uk/work/research-themes/risks-from-artificial-intelligence/>.
- 4 Kristalina Georgieva, “AI Will Transform the Global Economy. Let’s Make Sure It Benefits Humanity,” International Monetary Fund, January 14, 2024, <https://www.imf.org/en/Blogs/Articles/2024/01/14/ai-will-transform-the-global-economy-lets-make-sure-it-benefits-humanity>.
- 5 “193 Countries Adopt First-Ever Agreement on the Ethics of Artificial Intelligence,” *United Nations News*, November 25, 2021, <https://news.un.org/en/story/2021/11/1106612>.
- 6 Sanhita Chauriha, “How the Digital India Act Will Shape the Future of the Country’s Cyber Landscape,” *The Hindu*, October 9, 2023, <https://www.thehindu.com/sci-tech/technology/how-the-digital-india-act-will-shape-the-future-of-the-countrys-cyber-landscape/article67397155.ece>.
- 7 “National Strategy for Artificial Intelligence,” NITI Aayog, June 2018, <https://www.niti.gov.in/sites/default/files/2023-03/National-Strategy-for-Artificial-Intelligence.pdf>.
- 8 Chauriha, “How the Digital India Act Will Shape the Future of the Country’s Cyber Landscape”
- 9 The White House, “Removing Barriers to American Leadership in Artificial Intelligence,” January 23, 2025, <https://www.whitehouse.gov/presidential-actions/2025/01/removing-barriers-to-american-leadership-in-artificial-intelligence/>
- 10 The White House, *Executive Order on the Safe, Secure and Trustworthy Development and Use of Artificial Intelligence*, November 1, 2023, <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>
- 11 Kevin Frazier and Adam Thierer, “1,000 AI Bills: Time for Congress to Get Serious About Preemption,” *Lawfare*, May 9, 2025, <https://www.lawfaremedia.org/article/1-000-ai-bills--time-for-congress-to-get-serious-about-preemption>
- 12 United Kingdom Office for Artificial Intelligence, *A pro-innovation approach to AI regulation*, Department of Science, Innovation, and Technology, 2023, <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper>.

- 13 Palais de l'Élysée, "Statement on Inclusive and Sustainable Artificial Intelligence for People and the Planet," February 10-11, 2025, <https://www.politico.eu/wp-content/uploads/2025/02/11/02-11-AI-Action-Summit-Declaration.pdf>.
- 14 Justin Hendrix, "Podcast: Paths Diverge at the Paris AI Action Summit," *TechPolicy Press*, February 16, 2025, <https://www.techpolicy.press/podcast-paths-diverge-at-the-paris-ai-action-summit/>
- 15 Giovanni Buttarelli, "The EU GDPR as clarion call for a new global digital gold standard," *International Data Privacy Law* 6, no. 2 (2016), <https://academic.oup.com/idpl/article/6/2/77/2404469?login=false>.
- 16 Directorate-General for Communication, European Commission, https://commission.europa.eu/news/ai-act-enters-force-2024-08-01_en, 2024.
- 17 "EU AIA: first regulation on artificial intelligence," *European Parliament*, June 18, 2024, <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>
- 18 European Parliament, *Artificial Intelligence Act: deal on comprehensive rules for trustworthy AI*, December 9, 2023, <https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>
- 19 Council of the EU. <https://www.consilium.europa.eu/en/press/press-releases/2024/05/21/artificial-intelligence-ai-act-council-gives-final-green-light-to-the-first-worldwide-rules-on-ai/pdf/>, 2024.
- 20 European Parliament, "Directive 2000/31/EC on certain legal aspects of information society services, in particular electronic commerce, in the Internal Market (Directive on electronic commerce)," *Official Journal L* 178, <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32000L0031>
- 21 'Directive 2000/31/EC', p. 2
- 22 'Directive 2000/31/EC', p. 4
- 23 'Directive 2000/31/EC', p. 13
- 24 'Directive 2000/31/EC', p. 13
- 25 'E-Commerce Directive', *European Commission*, June 7, 2022, <https://digital-strategy.ec.europa.eu/en/policies/e-commerce-directive>
- 26 'Digital agenda for Europe', *European Parliament*, April 2024, <https://www.europarl.europa.eu/factsheets/en/sheet/64/digital-agenda-for-europe>

- 27 European Parliament, “Regulation (EU) 2016/679 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation),” *Official Journal L* 119/1, <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:02016R0679-20160504>
- 28 European Commission, *Legal Framework of EU Data Protection*, https://commission.europa.eu/law/law-topic/data-protection/legal-framework-eu-data-protection_en; European Union, “Charter of Fundamental Rights of the European Union,” *Official Journal of the European Union*, 2012, <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:12012P/TXT>.
- 29 ‘Regulation (EU) 2016/679’, pp. 6-7, 9-11
- 30 ‘Regulation (EU) 2016/679’, pp. 15-17, 19-20
- 31 European Parliament, “Directive (EU) 2019/790 on copyright and related rights in the Digital Single Market and amending Directives 96/9/EC and 2001/29/EC,” *Official Journal L* 130, <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32019L0790>.
- 32 Gerald Spindler, “EU Internet policy in the 2020s,” in *Research Handbook on EU Internet Law* (Cheltenham, UK: Edward Elgar Publishing, 2023), <https://www.elgaronline.com/edcollchap/book/9781803920887/book-part-9781803920887-7.xml>.
- 33 “The Digital Services Act package,” *European Commission*, <https://digital-strategy.ec.europa.eu/en/policies/digital-services-act-package>
- 34 “About the Digital Markets Act,” *European Commission*, https://digital-markets-act.ec.europa.eu/about-dma_en
- 35 “Digital Services Act: Questions and Answers,” *European Commission*, <https://digital-strategy.ec.europa.eu/en/faqs/digital-services-act-questions-and-answers>
- 36 Amélie P. Heldt, “EU Digital Services Act: The White Hope of Intermediary Regulation,” in *Digital Platform Regulation: Global Perspectives on Internet Governance*, ed. Terry Flew and Fiona R. Martin (Switzerland: Palgrave Macmillan, 2022), 69-84, <https://library.oapen.org/bitstream/handle/20.500.12657/56979/1/978-3-030-95220-4.pdf#page=82>.
- 37 Heldt, “EU Digital Services Act,” 70-71
- 38 Two useful examples are the ‘Avia Law’ proposed in France and ‘NetzDG’ proposed in Germany, see: “French Avia Law declared unconstitutional: what does this reach us at EU level?,” *EDRi*, June 24, 2020, <https://edri.org/our-work/french-avia-law-declared-unconstitutional-what-does-this-teach-us-at-eu-level/>; Diane Lee, “Germany’s NetzDG and the Threat to Online Free Speech,” *Yale Law School*, October 10, 2017, <https://law.yale.edu/mfia/case-disclosed/germanys-netzdg-and-threat-online-free-speech>.

- 39 Spindler, “EU Internet policy in the 2020s,” 8.
- 40 European Parliament, “Regulation (EU) 2019/1150 on promoting fairness and transparency for business user of online intermediation services,” *Official Journal L* 186, <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32019R1150>
- 41 “Regulation (EU) 2019/1150”
- 42 Heldt, “EU Digital Services Act,” 72
- 43 “White Paper on Artificial Intelligence – A European approach to excellence and trust,” *European Commission*, February 19, 2020, Brussels, https://commission.europa.eu/document/download/d2ec4039-c5be-423a-81ef-b9e44e79825b_en?filename=commission-white-paper-artificial-intelligence-feb2020_en.pdf
- 44 “White Paper on Artificial Intelligence,” 4
- 45 European Commission, “Coordinated Plan on Artificial Intelligence,” COM(2018) 795 final, December 7, 2018, https://eur-lex.europa.eu/resource.html?uri=cellar:22ee84bb-fa04-11e8-a96d-01aa75ed71a1.0002.02/DOC_1&format=PDF.
- 46 Tambiama Madiega, “Artificial intelligence act,” *European Parliament Research Service*, March 2024, <https://www.iisf.ie/files/UserFiles/cybersecurity-legislation-ireland/EU-AI-Act.pdf>.
- 47 Madiega, “Artificial intelligence act,” 2.
- 48 European Commission, “Ethics Guidelines for Trustworthy AI,” *High-Level Group on Artificial Intelligence*, April 8, 2019, <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- 49 European Commission, “Impact Assessment Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) And Amending Certain Union Legislative Acts,” April 21, 2021, <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021SC0084&qid=1619708088989>.
- 50 European Parliament, “Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act),” *Official Journal*, <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689&qid=1731278257144>.
- 51 “Regulation (EU) 2024/1689,” 7
- 52 Claudio Novelli, Federico Casolari, Antonio Rotolo, Mariarosaria Taddeo and Luciano Floridi, “AI Risk Assessment: A Scenario-Based, Proportional Methodology for the AI Act,” *Digital Society* 3, no. 13 (2024), <https://link.springer.com/content/pdf/10.1007/s44206-024-00095-1.pdf>.

- 53 “Regulation (EU) 2024/1689,” 4
- 54 “Regulation (EU) 2024/1689,” 4
- 55 “Regulation (EU) 2024/1689,” 4
- 56 “Regulation (EU) 2024/1689,” 52
- 57 “Regulation (EU) 2024/1689,” 51-53
- 58 “Regulation (EU) 2024/1689,” 8
- 59 “Regulation (EU) 2024/1689,” 8
- 60 “Regulation (EU) 2024/1689,” 9
- 61 Zeyi Yang, “China just announced a new social credit law. Here’s what it means,” *MIT Technology Review*, November 22, 2022, <https://www.technologyreview.com/2022/11/22/1063605/china-announced-a-new-social-credit-law-what-does-it-mean/>.
- 62 “Regulation (EU) 2024/1689,” 9
- 63 “Regulation (EU) 2024/1689,” 9
- 64 “Regulation (EU) 2024/1689,” 127-129
- 65 “Regulation (EU) 2024/1689,” 127-129
- 66 “High-level summary of the AI Act,” *EU Artificial Intelligence Act*, February 27, 2024, <https://artificialintelligenceact.eu/high-level-summary/>
- 67 “EU AI Act Risk Categories: Each Category Explained,” *Captain Compliance*, May 13, 2024, <https://www.captaincompliance.com/education/eu-ai-act-risk-categories/>
- 68 “EU AI Act Risk Categories: Each Category Explained”
- 69 “High-level summary of the AI Act”
- 70 “High-level summary of the AI Act”
- 71 Daniel Mügge, “EU AI sovereignty: for whom, to what end, and to whose benefit?” *Journal of European Public Policy* 31, no. 8 (2024): 2200-2225, <https://www.tandfonline.com/doi/pdf/10.1080/13501763.2024.2318475>.
- 72 Martin Ebers, “Truly Risk-based Regulation of Artificial Intelligence: How to Implement the EU’s AI Act,” *European Journal of Risk Regulation*, (2024): 1-20, https://www.cambridge.org/core/services/aop-cambridge-core/content/view/E526C1D0D7368F9691082220609D60F4/S1867299X24000783a.pdf/truly_riskbased_regulation_of_artificial_intelligence_how_to_implement_the_eus_ai_act.pdf.
- 73 “Truly Risk-based Regulation of Artificial Intelligence”

- 74 Celso Cancela-Outeda, "The EU's AI act: A framework for collaborative governance," *Internet of Things* 27, (2024): 1-11, https://www.sciencedirect.com/science/article/pii/S2542660524002324?ref=pdf_download&fr=RR-9&rr=8f4aa75c9b53ed08.
- 75 "Coordinated Plan on Artificial Intelligence"
- 76 Oliver Roberts, "EU AI Act's Burdensome Regulation Could Impair AI Innovation," *Bloomberg Law*, February 21, 2025, <https://news.bloomberglaw.com/us-law-week/eu-ai-acts-burdensome-regulations-could-impair-ai-innovation>
- 77 Pieter Haech, "EU Rules for advanced AI are step in wrong direction, Google says," *Politico*, February 10, 2025, <https://www.politico.eu/article/google-eu-rules-advanced-ai-artificial-intelligence-step-in-wrong-direction/>; Cynthia Kroet, "Big Tech watered down AI Code of Practice: report," *Euro News*, April 30, 2025, <https://www.euronews.com/next/2025/04/30/big-tech-watered-down-ai-code-of-practice-report>
- 78 Theophile Maizire, "EU's AI Code of Practice: Third Draft," *techUK*, March 11, 2025, <https://www.techuk.org/resource/eu-s-ai-code-of-practice-third-draft.html>
- 79 Kroet, "Big Tech watered down AI Code of Practice"
- 80 Raluca Csernaton, "The EU's AI Power Play: Between Deregulation and Innovation," *Carnegie Endowment Europe*, May 20, 2025, https://carnegieendowment.org/research/2025/05/the-eus-ai-power-play-between-deregulation-and-innovation?lang=en¢er=europe&mkt_tok=ODEzLVhZVS00MjIAAGai-WxBUpbSvUpmBxHR3_HTYaA49dy6ctmkOW3RFh935KBXLI5y2eroEvmBtAFy2ZjCkfjD0LR071tp3iHkhHpZFdHg80exHkTGEk9.
- 81 Mario Draghi, *The Future of European Competitiveness: Part A | A competitiveness strategy for Europe*, European Commission, 2024, https://commission.europa.eu/document/download/97e481fd-2dc3-412d-be4c-f152a8232961_en?filename=The%20future%20of%20European%20competitiveness%20_%20A%20competitiveness%20strategy%20for%20Europe.pdf.
- 82 Cynthia Kroet, "EU Standards Bodies Flag Delays to Work on AI Act," *Euro News*, April 16, 2025, <https://www.euronews.com/next/2025/04/16/eu-standards-bodies-flag-delays-to-work-on-ai-act>
- 83 Trent Kubasiak, "The Future of the AI Liability Directive in Europe After Withdrawal," *Ethical AI Law Institute*, February 13, 2025, <https://ethicalailawinstitute.org/blog/the-future-of-the-ai-liability-directive-in-europe-after-withdrawal/>
- 84 "Regulation (EU) 2024/1689," 127-129
- 85 Michael Veale and Frederik Zuiderveen Borgesius, "Demystifying the Draft EU Artificial Intelligence Act: Analysing the good, the bad, and the unclear elements of the proposed approach," *Computer Law Review International* 4, (2021): 97-112, <https://arxiv.org/pdf/2107.03721>.

- 86 Veale and Borgesius, “Demystifying the Draft EU Artificial Intelligence Act”
- 87 Veale and Borgesius, “Demystifying the Draft EU Artificial Intelligence Act”
- 88 “Regulation (EU) 2024/1689,” 129
- 89 Pablo Barberá, “Social Media, Echo Chambers, and Political Polarization,” in *Social Media and Democracy*, ed. Nathaniel Persily and Joshua A. Tucker, (Cambridge: Cambridge University Press, 2020), 34–55, <https://www.cambridge.org/core/books/social-media-and-democracy/social-media-echo-chambers-and-political-polarization/333A5B4DE1B67EFF7876261118CCFE19>.
- 90 Kanishka Singh and Sheila Dang, “Musk and X are epicentre of US election misinformation, experts say,” *Reuters*, November 5, 2024, <https://www.reuters.com/world/us/wrong-claims-by-musk-us-election-got-2-billion-views-x-2024-report-says-2024-11-04/>; Jill Lawless, “Elon Musk Helped Trump Win. Now he’s looking at Europe, and many politicians are alarmed,” *Associated Press*, 8 January 2025, <https://apnews.com/article/elon-musk-europe-politics-germany-uk-f50d69d0d192a2d81c95f5d64c6d4acd>.
- 91 Matija Franklin, Hal Ashton, Rebecca Gorman, Stuart Armstrong, “The EU’s AI Act needs to address critical manipulation methods,” *OECD.AI*, March 21, 2023, <https://oecd.ai/en/wonk/ai-act-manipulation-methods>.
- 92 “Demystifying the Draft EU Artificial Intelligence Act”
- 93 “Regulation (EU) 2024/1689,” 8
- 94 Elena Losanno, “Brain-Body Interfaces to Assist and Restore Motor Functions in People with Paralysis,” in *Brain-Computer Interface Research*, ed. C. Guger, B. Allison, T.M. Rutkowski, M. Korostenskaja (Springer, Cham: 2024), https://link.springer.com/chapter/10.1007/978-3-031-49457-4_7#citeas.
- 95 Ptitto Maurice, Bleau Maxime, Djerourou Ismaël, Paré Samuel, Schneider Fabien C., Chebat Daniel-Robert, “Brain-Machine Interfaces to Assist the Blind,” *Frontiers in Human Neuroscience* 15, (2021), <https://www.frontiersin.org/journals/human-neuroscience/articles/10.3389/fnhum.2021.638887/full>.
- 96 Jennifer Blumenthal-Barby, “Neuro rights and the right to mental integrity,” *Journal of Medical Ethics* 50, (2024): 655, <https://jme.bmj.com/content/50/10/655>.
- 97 Kate Payne, “An AI chatbot pushed a teen to kill himself, a lawsuit against its creator alleges,” *AP News*, October 26, 2024, <https://apnews.com/article/chatbot-ai-lawsuit-suicide-teen-artificial-intelligence-9d48adc572100822fdbc3c90d1456bd0>
- 98 Barbara Prainsack and Nikolaus Forgó, “New AI Regulation in the EU seeks to reduce risk without assessing public benefit,” *Nature Medicine* 30, (2024): 1235–1237, <https://www.nature.com/articles/s41591-024-02874-2>.

- 99 Veale and Borgesius, “Demystifying the Draft EU Artificial Intelligence Act”
- 100 Johann Laux, Sandra Wachter and Brent Mittelstadt, “Trustworthy artificial intelligence and the European Union AI Act: On the conflation of trustworthiness and acceptability of risk,” *Regulation & Governance* 18, (2024): 3-32, <https://onlinelibrary.wiley.com/doi/pdf/10.1111/rego.12512>.
- 101 Laux et al., “Trustworthy artificial intelligence and the European Union AI Act”
- 102 Marko Marjanovic, “10 Biggest AI Companies in the World [2024],” *FinBold*, May 17, 2024, <https://finbold.com/guide/10-biggest-ai-companies-in-the-world/>
- 103 “The Future of European Competitiveness: Part A | A competitiveness strategy for Europe”
- 104 “The Future of European Competitiveness: Part A | A competitiveness strategy for Europe”
- 105 Francesca Bria, “The Quest for European Technological Sovereignty: Building the EuroStack,” *TechPolicy Press*, October 15, 2024, <https://www.techpolicy.press/the-quest-for-european-technological-sovereignty-building-the-eurostack/>
- 106 European Commission, “Ethics guidelines for trustworthy AI,” April 8, 2019, <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- 107 Mario D. Schultz, Ludovico Giacomo Conti And Peter Seele, “Digital ethicswashing: a systematic review and a process-perception-outcome framework,” *AI Ethics*, 2024, <https://link.springer.com/article/10.1007/s43681-024-00430-9>.
- 108 Thomas Metzinger, “EU guidelines: Ethics washing made in Europe,” *Tagesspiegel*, April 8, 2019, <https://www.tagesspiegel.de/politik/ethics-washing-made-in-europe-5937028.html>
- 109 European Parliament, “Opinion of the Committee on Legal Affairs on the proposal for a regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative Acts,” *Committee on Legal Affairs*, 2021/0106 (COD), September 12, 2022, <https://artificialintelligenceact.eu/wp-content/uploads/2022/09/AIA-JURI-Rule-57-Opinion-Adopted-12-September.pdf>.
- 110 “Regulation (EU) 2024/1689,” 43; “Article 112: Evaluation and Review,” EU Artificial Intelligence Act, <https://artificialintelligenceact.eu/article/112/>.
- 111 Veale and Borgesius, “Demystifying the Draft EU Artificial Intelligence Act”
- 112 Schultz et al., “Digital ethicswashing”
- 113 Manuel Wörsdörfer, “The EU’s Artificial Intelligence Act: An Ordoliberal Assessment,” *AI Ethics* 5, (2025): 263-278, <https://link.springer.com/article/10.1007/s43681-023-00337-x#citeas>.

Endnotes

- 114 Alberto Rizzi, “The infinite connection: How to make the India-Middle East-Europe economic corridor happen,” *European Council on Foreign Relations*, April 23, 2024, <https://ecfr.eu/publication/the-infinite-connection-how-to-make-the-india-middle-east-europe-economic-corridor-happen/>
- 115 “UAE to invest billions in France AI data center,” *France24*, February 7, 2025, <https://www.france24.com/en/europe/20250207-uae-to-invest-up-to-€50-billion-in-massive-ai-data-centre-in-france>



Ideas . Forums . Leadership . Impact

20, Rouse Avenue Institutional Area,
New Delhi - 110 002, INDIA

Ph. : +91-11-35332000. Fax : +91-11-35332005

E-mail: contactus@orfonline.org

Website: www.orfonline.org