



CONVERGENCE OR FRAGMENTATION? A COMPARATIVE ANALYSIS OF AI GOVERNANCE FRAMEWORKS IN THE EU, UNITED STATES, AND CHINA

Mehar Sandhu

PhD Research Scholar, School of Social Sciences, Guru Nanak Dev University, Amritsar, India.

DOI: [https://doi.org/10.56815/IRJAHS.v\(2026\)i2.1-10](https://doi.org/10.56815/IRJAHS.v(2026)i2.1-10)

Abstract:

As artificial intelligence systems increasingly transcend national borders, the regulatory frameworks governing them have developed along strikingly divergent paths. This paper undertakes a comparative analysis of three of the world's most influential AI governance regimes: the European Union's AI Act, the United States' executive order-driven approach, and China's evolving suite of AI regulations. Using a structured comparative framework, the study examines each jurisdiction's regulatory philosophy, risk classification methodology, enforcement mechanisms, and treatment of general-purpose versus narrow AI systems. The EU's AI Act represents a comprehensive, risk-tiered, rights-based model rooted in ex-ante compliance; the US approach reflects a more fragmented, sector-specific, and executive-driven strategy shaped by shifting political priorities and limited congressional action; China's framework combines state-driven algorithmic oversight with strategic industrial policy objectives, emphasising content control and social stability alongside innovation goals. The paper argues that while surface-level convergence is emerging around certain technical concepts — such as risk categorisation, transparency obligations, and safety testing — the underlying regulatory logics, enforcement capacities, and normative priorities remain substantially fragmented. This divergence carries significant implications for multinational AI developers navigating compliance across jurisdictions, for the prospects of international regulatory harmonisation, and for the broader question of whether a global AI governance regime is achievable or desirable. The paper concludes by evaluating pathways toward interoperability, including mutual recognition mechanisms, technical standard-setting bodies, and mini-lateral cooperation, while cautioning against overstating the likelihood of near-term convergence given divergent political economies and strategic interests.

Key words: *AI governance; EU AI Act; algorithmic regulation; comparative policy analysis; regulatory fragmentation; AI executive orders; China AI regulation; global AI governance; regulatory harmonization; risk-based regulation.*

1. Introduction

Artificial intelligence has moved from a niche technical domain to a central axis of geopolitical competition, economic strategy, and social concern. As AI systems are deployed at scale across healthcare, finance, defence, education, and public administration, governments have moved—unevenly and often reactively—to regulate their development and use. Yet unlike prior waves of technology governance, AI regulation is emerging in a fractured global environment, where major powers are pursuing distinct and, in some respects, competing visions of what responsible AI governance should look like.

Three jurisdictions have come to define the contours of this emerging regulatory landscape. The European Union, through its AI Act, has adopted the world's first comprehensive, horizontally applicable legal



framework for AI, built on a risk-tiered classification system and grounded in the EU's broader tradition of rights-based digital regulation. The United States, by contrast, has pursued a markedly different path—one characterized by executive orders, sector-specific agency guidance, and a reluctance (thus far) to enact comprehensive federal legislation, reflecting both the fragmented structure of the American regulatory state and deep political disagreement over the appropriate scope of government intervention in AI markets. China has charted a third course, embedding AI governance within a broader framework of state control, combining targeted regulations on algorithmic recommendation systems, generative AI, and deep synthesis technologies with strategic industrial policy aimed at achieving technological self-sufficiency and global competitiveness.

These divergent approaches raise a pressing question for scholars, policymakers, and industry alike: are we witnessing the early stages of convergence toward a shared global governance paradigm, or is the world settling into a period of durable regulatory fragmentation? The stakes of this question are considerable. For multinational AI developers, divergent compliance regimes create significant operational complexity and legal uncertainty. For policymakers, the absence of interoperable standards raises the risk of regulatory arbitrage, where firms relocate development or deployment to jurisdictions with the least stringent oversight. For the international community, fragmentation complicates efforts to address AI risks that are inherently transnational in character, including disinformation, autonomous weapons, and the governance of frontier AI capabilities with potentially catastrophic risk profiles.

This paper addresses this question through a structured comparative analysis of the EU, US, and Chinese AI governance frameworks. Rather than treating these regimes as static legal texts, the analysis situates each within its broader political, economic, and institutional context, examining not only formal regulatory instruments but also the underlying philosophies, enforcement capacities, and strategic objectives that shape how each jurisdiction approaches AI oversight. The paper proceeds in five parts. Following this introduction, Section 2 reviews the existing literature on comparative technology governance and situates the present analysis within debates over regulatory convergence theory. Section 3 presents the comparative framework and methodology used to analyse the three regimes. Section 4 offers a detailed comparative examination of the EU, US, and Chinese frameworks across key dimensions—regulatory philosophy, risk classification, enforcement mechanisms, and treatment of general-purpose AI. Section 5 synthesises these findings to assess the balance of convergent and divergent forces at play, before the paper concludes by considering pathways toward interoperability and the broader implications for global AI governance.

In examining these three cases together, this paper contributes to a growing body of scholarship seeking to understand not merely what AI regulations say, but what they reveal about the deeper political economies and governance philosophies from which they emerge—and what this means for the prospects of coherent global AI governance in the years ahead.

2: Literature Review

Scholarship on AI governance has expanded rapidly over the past decade, drawing on adjacent literatures in comparative regulation, international relations, and science and technology studies. This section situates the present analysis within three interrelated bodies of work: theories of regulatory convergence and divergence, comparative studies of digital and technology governance, and the emerging AI-specific policy literature.





2.1 Regulatory Convergence Theory

The question of whether regulatory regimes converge or diverge over time has long occupied comparative political economy and international relations scholarship. Classical convergence theory suggests that globalisation, market integration, and shared technical challenges push regulatory systems toward similar outcomes, driven by mechanisms such as competitive pressure, policy learning, and harmonisation through international institutions. Divergence theorists, by contrast, emphasise the durability of domestic institutional structures, path dependency, and distinct political economies, arguing that formally similar regulatory language can mask substantively different enforcement logics and normative commitments. A related literature on the “Brussels Effect” examines how the EU’s regulatory choices can exert outsized influence on global markets even absent formal international coordination, as multinational firms adopt EU standards globally to reduce compliance complexity. Whether this dynamic extends to AI governance—where the EU faces more assertive regulatory competitors in the US and China than it did in prior digital regulation battles—remains contested and is a central question this paper engages.

2.2 Comparative Digital and Technology Governance

A related body of scholarship examines comparative approaches to digital governance more broadly, including data protection, platform regulation, and cybersecurity. This literature has documented how the EU, US, and China have historically pursued distinct regulatory models: a rights-based approach in the EU (exemplified by the GDPR), a market-driven and sectorally fragmented approach in the US, and a state-centric model in China oriented toward social stability and strategic industrial objectives. These prior regulatory battles offer instructive, if imperfect, precedents for AI governance, and scholars have increasingly asked whether the same tripartite divergence observed in data protection will replicate in AI regulation, or whether the distinct risk profile of AI—particularly frontier and general-purpose systems—will produce different alignments.

2.3 AI-Specific Governance Scholarship

Since 2020, an AI-specific governance literature has emerged, addressing questions including risk-based regulatory design, the governance of general-purpose and frontier AI models, and the institutional architecture needed for effective oversight. Much of this literature has focused on individual jurisdictions in isolation—detailed doctrinal analyses of the EU AI Act, critiques of the US’s reliance on executive action absent congressional legislation, and studies of China’s algorithm-specific regulatory instruments. Comparative treatments remain comparatively scarce, and those that exist often focus narrowly on formal legal text rather than the underlying enforcement capacity, political economy, and strategic objectives that shape how regulations function in practice. This paper aims to address that gap by integrating formal regulatory analysis with attention to institutional context and strategic motivation.

2.4 Gaps Addressed by This Study

Taken together, this literature leaves three gaps that the present paper addresses. First, existing comparative work rarely applies a consistent analytical framework across all three major jurisdictions simultaneously, making systematic comparison difficult. Second, much of the literature treats regulatory convergence as a binary outcome rather than examining which specific dimensions of governance (e.g., technical standards, risk taxonomies, enforcement mechanisms) are converging while others diverge. Third, few studies explicitly connect AI governance outcomes to the deeper political economies—democratic versus authoritarian governance, market liberalism versus state capitalism—that shape regulatory choices in each





jurisdiction. This paper's comparative framework, presented in the following section, is designed to address each of these gaps.

3: Comparative Framework and Methodology

This paper employs a qualitative comparative policy analysis to examine the AI governance frameworks of the European Union, the United States, and China. Rather than treating each jurisdiction as a self-contained case study, the analysis applies a consistent set of comparative dimensions across all three regimes, enabling systematic identification of points of convergence and divergence. This approach follows established methods in comparative regulatory scholarship, which favour structured, focused comparison over purely descriptive case narration, allowing for generalisable insights while retaining sensitivity to jurisdiction-specific institutional context.

The EU, US, and China were selected as cases on the basis of three criteria: regulatory influence, market size, and philosophical distinctiveness. Collectively, these three jurisdictions host the majority of the world's leading AI developers and represent the largest AI-relevant markets by both consumer base and capital investment. More importantly for this study's purposes, each jurisdiction represents a distinct governance paradigm—rights-based regulation in the EU, market-driven fragmented oversight in the US, and state-directed techno-industrial policy in China—making the three cases a theoretically productive basis for comparison. Other jurisdictions with emerging AI frameworks, such as the UK, Canada, Singapore, and South Korea, are acknowledged as relevant but excluded from primary analysis in order to preserve analytical depth over breadth; where relevant, these frameworks are referenced comparatively in the discussion section.

Each jurisdiction's governance framework is examined across five comparative dimensions, selected to capture both the formal-legal and political-economic dimensions of AI regulation. The first is regulatory philosophy and legal basis, meaning the foundational normative orientation of the regime—whether centered on fundamental rights protection, market competitiveness, national security, or social stability—and its primary legal instruments, whether statute, executive order, or administrative regulation. The second is risk classification and scope, which considers how each regime defines and categories AI systems by risk level or application domain, and whether the framework applies horizontally across all AI applications or vertically to specific sectors. The third dimension concerns the treatment of general-purpose and frontier AI, examining how each jurisdiction addresses foundation models and general-purpose AI systems as distinct from narrow, application-specific systems, including any compute thresholds, capability-based triggers, or systemic risk designations. The fourth dimension addresses enforcement mechanisms and institutional capacity, including the administrative bodies responsible for oversight, the enforcement tools available to them—such as fines, licensing, audits, or outright bans—and the practical capacity of these institutions to monitor and enforce compliance. The fifth and final dimension concerns transparency, accountability, and rights provisions, encompassing requirements related to disclosure, explainability, human oversight, and the avenues for redress available to individuals affected by AI systems.

The analysis draws on primary legal and regulatory texts, including the finalised EU AI Act and accompanying guidance documents, US federal executive orders and associated agency guidance such as NIST frameworks and OMB memoranda, and China's Cyberspace Administration regulations governing algorithmic recommendation, deep synthesis, and generative AI services. These primary sources are supplemented by secondary sources, including regulatory impact assessments, academic commentary, and reporting on early enforcement actions where available. Given the fast-moving nature of AI regulation, the





analysis reflects the state of each framework as of the time of writing and explicitly notes areas of ongoing regulatory development.

Several limitations warrant acknowledgment. Regulatory frameworks in this domain are evolving rapidly, and conclusions drawn at a given point in time may be superseded by subsequent legal or administrative developments. Enforcement data, particularly for recently enacted provisions, remains limited, which constrains empirical assessment of practical regulatory impact relative to formal legal design. China's regulatory environment presents particular challenges for external analysis given constraints on transparency regarding enforcement practices and internal administrative guidance. These limitations are addressed through explicit qualification of claims regarding enforcement efficacy and through reliance, where possible, on multiple corroborating sources.

4: Comparative Analysis

4.1 Regulatory Philosophy and Legal Basis

The three jurisdictions examined in this study proceed from fundamentally different normative starting points, and these differences shape nearly every subsequent aspect of their respective frameworks. The European Union's approach is rooted in its long-standing tradition of rights-based digital regulation, treating AI oversight as a natural extension of the values embedded in instruments such as the GDPR. The AI Act is framed explicitly around the protection of fundamental rights, health, safety, and democratic values, and it takes the form of a binding regulation directly applicable across all member states—reflecting the EU's institutional capacity to legislate comprehensively at the supranational level. The United States, by contrast, has pursued AI governance primarily through executive action rather than comprehensive federal legislation, a pattern driven by the difficulty of securing congressional consensus on technology regulation and by a political culture that remains wary of preemptive regulatory intervention in fast-moving markets. This has produced a patchwork of agency-level guidance, voluntary frameworks such as the NIST AI Risk Management Framework, and sector-specific rules issued by bodies like the FTC and FDA, resulting in a regime that is far more fragmented and reversible than its EU counterpart, given that executive orders can be rescinded or substantially altered by subsequent administrations. China's approach differs from both, embedding AI governance within a broader state-directed framework that treats technology regulation as inseparable from industrial policy and social governance objectives. Chinese regulations on algorithmic recommendation, deep synthesis, and generative AI services are issued primarily through the Cyberspace Administration of China and are explicitly oriented toward maintaining social stability, content control, and ideological alignment, while simultaneously supporting the government's strategic ambition to achieve global leadership in AI technology.

4.2 Risk Classification and Scope

The EU AI Act's defining structural feature is its risk-tiered classification system, which sorts AI applications into categories of unacceptable, high, limited, and minimal risk, with obligations scaling according to a system's potential for harm. This horizontal approach applies across all sectors and use cases, creating a single coherent framework intended to capture the full range of AI applications within the EU market, regardless of the industry in which they are deployed. The United States has adopted no equivalent horizontal risk taxonomy at the federal level; instead, risk is addressed unevenly through sector-specific rules and voluntary guidance, meaning that an AI system used in healthcare may be subject to entirely different oversight than a functionally similar system used in employment screening or consumer finance, depending





on which regulatory body claims jurisdiction. China's regulatory scope has developed incrementally and vertically, with distinct rules issued for specific categories of AI application—algorithmic recommendation systems in 2021, deep synthesis technologies in 2022, and generative AI services in 2023—rather than through a single overarching framework. Notably, Chinese generative AI regulations impose distinctive obligations tied to content and political sensitivity, including requirements that generated content align with “core socialist values,” a risk framing entirely absent from both the EU and US approaches and reflective of China's distinct governance priorities.

4.3 Treatment of General-Purpose and Frontier AI

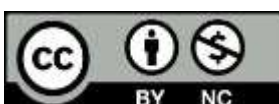
The governance of general-purpose and frontier AI systems represents one of the more technically and politically contested areas of comparison. The EU AI Act addresses general-purpose AI models directly, introducing a tiered system of obligations that intensifies for models designated as posing “systemic risk,” a designation triggered in part by compute thresholds and other capability indicators, with additional requirements around model evaluation, incident reporting, and cybersecurity. The US approach to frontier AI has relied heavily on voluntary commitments secured from leading AI developers, alongside reporting requirements for large-scale compute usage that were introduced via executive action; however, the durability and enforceability of these mechanisms remain considerably weaker than the EU's binding statutory obligations, and their status is subject to change with shifts in executive branch priorities. China's regulatory framework has, to date, addressed general-purpose and foundation models less through dedicated systemic-risk provisions and more through the extension of existing content-oriented and algorithmic filing requirements to large language models, while separately pursuing compute and chip-access controls that function as a de facto, if indirect, mechanism for constraining frontier AI development, driven substantially by the geopolitical context of US export restrictions.

4.4 Enforcement Mechanisms and Institutional Capacity

Enforcement architecture varies considerably across the three regimes, with significant implications for practical regulatory impact. The EU AI Act establishes a multi-layered enforcement structure involving national competent authorities, a newly created European AI Office responsible for oversight of general-purpose AI models, and the possibility of substantial fines for noncompliance, calibrated as a percentage of global annual turnover in a manner similar to GDPR enforcement. Whether this institutional architecture will achieve enforcement capacity comparable to its ambitious legal mandate remains an open empirical question, given the resource and technical expertise demands of AI oversight. In the United States, enforcement responsibility is distributed across numerous existing agencies, including the FTC, EEOC, and sector-specific regulators, applying existing legal authorities to AI-related harms rather than operating under a dedicated AI enforcement body, which produces enforcement that is often reactive, jurisdictionally contested, and dependent on the resources and priorities of individual agencies. China's enforcement capacity is generally considered robust relative to its ability to compel compliance from domestic firms, given the state's broader regulatory leverage over the technology sector and its established practices of algorithm registration and content moderation enforcement, though independent verification of enforcement practices is constrained by limited transparency.

4.5 Transparency, Accountability, and Rights Provisions

The EU AI Act includes explicit transparency obligations, including disclosure requirements for AI-generated content, documentation requirements for high-risk systems, and provisions enabling forms of





individual redress, reflecting its grounding in fundamental rights protection. The US framework offers comparatively limited individual rights provisions at the federal level, with meaningful transparency and redress mechanisms emerging primarily through state-level legislation, such as comprehensive privacy laws in states like California and Colorado, resulting in an uneven patchwork of protections that vary significantly by jurisdiction. China's regulations include disclosure requirements oriented primarily toward content labelling and algorithmic filing with regulators rather than individual rights of explanation or redress, reflecting the state-centric rather than rights-centric orientation of the broader regulatory philosophy.

Taken together, this comparative analysis reveals a landscape in which formal regulatory concepts—risk classification, systemic risk designations for frontier models, and transparency obligations—appear with increasing frequency across all three jurisdictions, suggesting a degree of technical convergence. Yet the underlying enforcement logics, institutional capacities, and normative commitments driving these frameworks remain substantially distinct, a tension explored further in the discussion section that follows.

5: Discussion and Synthesis

5.1 Assessing Convergence: Where the Regimes Align

The comparative analysis in Section 4 reveals meaningful, if partial, convergence along several technical and structural dimensions. All three jurisdictions have gravitated toward some form of differentiated treatment based on risk or application sensitivity, even where the underlying taxonomy and criteria differ substantially. Similarly, all three have found it necessary to address general-purpose and foundation models as a distinct regulatory category from narrow AI applications, suggesting that the technical characteristics of large-scale AI systems are exerting a degree of convergent pressure on regulatory design regardless of jurisdiction. Transparency obligations, particularly around AI-generated content disclosure, have also emerged as a near-universal feature across all three regimes, indicating an emerging normative consensus that synthetic content warrants some form of labelling or disclosure. These points of alignment suggest that technical standard-setting bodies, along with the practical realities of how large AI systems are built and deployed, may be functioning as an informal convergence mechanism even in the absence of coordinated international policymaking.

5.2 Assessing Divergence: Where the Regimes Diverge

Despite these areas of surface alignment, the analysis reveals persistent and arguably widening divergence at the level of regulatory philosophy, enforcement architecture, and institutional durability. The EU's rights-based, horizontally binding approach stands in sharp contrast to the US's fragmented, executive-driven, and comparatively reversible framework, a divergence rooted in fundamentally different constitutional structures and political cultures rather than in disagreement over technical risk assessment alone. China's framework diverges from both in a different direction, subordinating individual rights protections to state and social stability objectives in a manner that has no genuine analogue in either the EU or US approaches. Perhaps most significantly, the enforcement capacity gap between these regimes—driven by differences in institutional resourcing, administrative tradition, and political will—means that even where formal legal language converges, practical regulatory outcomes may continue to diverge substantially. This distinction between textual convergence and functional divergence represents one of this paper's central findings: scholars and policymakers assessing global AI governance trends should be cautious about inferring substantive alignment from surface-level similarities in legal drafting.





5.3 Implications for Multinational AI Developers

The persistence of regulatory fragmentation, even amid partial technical convergence, carries significant practical implications for firms operating across these jurisdictions. Multinational AI developers face the prospect of navigating materially different compliance regimes, including divergent risk classification systems, inconsistent documentation and disclosure requirements, and varying thresholds for systemic risk designation. This complexity creates incentives for firms to engineer compliance strategies around the most stringent applicable regime—a dynamic consistent with the Brussels Effect observed in prior EU digital regulation—though the presence of two other major regulatory powers with distinct and assertive approaches suggests this dynamic may be considerably weaker for AI governance than it was for data protection. Firms may instead face genuine trilateral compliance burdens rather than a single de facto global standard, with associated costs likely to disproportionately burden smaller AI developers relative to well-resourced incumbents.

5.4 Implications for International Cooperation

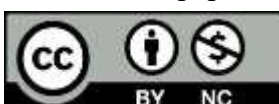
The findings of this analysis carry important implications for the prospects of international AI governance cooperation. The partial convergence observed around technical concepts suggests that targeted, functional cooperation—for instance, through international technical standard-setting bodies, mutual recognition arrangements for conformity assessments, or narrower mini-lateral agreements among like-minded jurisdictions—may be more achievable in the near term than comprehensive treaty-based governance analogous to nuclear nonproliferation or climate frameworks. The deep divergence in underlying political economies and normative commitments identified in this paper suggests that comprehensive international harmonisation remains unlikely absent significant shifts in the domestic political conditions driving each jurisdiction's approach. This is particularly true given that AI governance, unlike many prior domains of international regulatory cooperation, is deeply entangled with great power strategic competition, particularly between the United States and China, a dynamic that actively disincentivizes the kind of regulatory cooperation that might otherwise be expected between major economic powers.

5.5 Pathways Toward Interoperability

Notwithstanding these obstacles, the analysis suggests several plausible pathways toward greater interoperability short of full harmonisation. Technical standard-setting bodies such as ISO and IEC may serve as important convening spaces for de facto alignment, particularly around measurable technical criteria such as compute thresholds and evaluation methodologies. Mutual recognition arrangements, allowing conformity assessments conducted under one jurisdiction's framework to satisfy requirements in another, offer a further pathway to reduce compliance burden without requiring substantive legal harmonisation. Mini-lateral cooperation among democratic jurisdictions with broadly aligned rights-based commitments—potentially including the EU, UK, Canada, and other likeminded partners—may prove more tractable than trilateral cooperation inclusive of China, given the scale of underlying normative divergence identified in this analysis.

6: Conclusion

This paper set out to examine whether the AI governance frameworks emerging in the European Union, the United States, and China reflect a trajectory toward regulatory convergence or one of continued fragmentation. Through a structured comparative analysis across five dimensions—regulatory philosophy, risk classification, treatment of general-purpose and frontier AI, enforcement mechanisms, and transparency provisions—the paper finds that the answer is neither straightforwardly convergent nor straightforwardly





divergent, but rather a more complex pattern in which technical and structural concepts are converging even as underlying political economies, enforcement capacities, and normative commitments remain substantially distinct.

This finding carries important implications. It suggests that observers should resist the temptation to read formal legal similarity—such as the shared emergence of risk-tiered classifications or systemic risk designations for frontier models—as evidence of genuine regulatory alignment. The EU’s rights-based and institutionally durable framework, the US’s fragmented and executively reversible approach, and China’s state-directed model oriented around social stability and strategic industrial policy each reflect distinct constitutional structures, political cultures, and governance philosophies that are unlikely to converge fully in the foreseeable future, even as the underlying AI technologies they seek to regulate become increasingly similar in technical form. The gap between textual convergence and functional divergence identified in this paper therefore represents an important corrective to overly optimistic accounts of emerging global AI governance consensus.

At the same time, this paper’s findings should not be read as entirely pessimistic regarding the prospects for international cooperation. Targeted, functional convergence—achieved through technical standard-setting bodies, mutual recognition arrangements, and minilateral cooperation among jurisdictions with aligned normative commitments—represents a plausible and potentially productive path forward, even in the absence of comprehensive treaty-based harmonization. Such pathways may prove more durable and more achievable than ambitious calls for a single global AI governance regime, particularly given the extent to which AI governance has become entangled with broader great power strategic competition.

This study is not without limitations, and these point toward productive avenues for future research. As noted in Section 3, the rapidly evolving nature of AI regulation means that the specific provisions analyzed in this paper will likely be superseded or amended over time, and future research should track the evolution of enforcement practice as these frameworks mature, moving beyond the analysis of formal legal text toward empirical assessment of regulatory outcomes. Further comparative research incorporating additional jurisdictions—including the United Kingdom, India, and other emerging AI regulatory powers—would help clarify whether the tripartite divergence identified in this paper represents a durable global pattern or a transitional phase preceding broader realignment. Finally, future work might productively examine the perspectives and compliance experiences of multinational AI developers navigating these divergent regimes directly, offering an empirical complement to the primarily doctrinal and institutional analysis undertaken here.

Ultimately, the governance of artificial intelligence is likely to remain a contested and evolving domain for the foreseeable future, shaped as much by geopolitical competition and domestic political economy as by the technical characteristics of the technology itself. Understanding the specific dimensions along which convergence is occurring—and, just as importantly, those along which it is not—will be essential for policymakers, firms, and scholars seeking to navigate this landscape in the years ahead.

References

- Executive Order 14110, Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, 88 Fed. Reg. 75191 (Oct. 30, 2023) (rescinded Jan. 20, 2025).
- Executive Order 14179, Removing Barriers to American Leadership in Artificial Intelligence, 90 Fed. Reg. 8741 (Jan. 23, 2025).





- Cyberspace Administration of China, National Development and Reform Commission, Ministry of Education, Ministry of Science and Technology, Ministry of Industry and Information Technology, Ministry of Public Security, & National Radio and Television Administration. (2023). *Interim Measures for the Management of Generative Artificial Intelligence Services* [生成式人工智能服务管理暂行办法]. Effective August 15, 2023.
- Cyberspace Administration of China, Ministry of Public Security, & Ministry of Industry and Information Technology. (2023). *Provisions on the Administration of Deep Synthesis Internet Information Services*. Effective January 10, 2023.
- Cyberspace Administration of China. (2022). *Administrative Provisions on Algorithm Recommendation for Internet Information Services*. Effective March 1, 2022.
- Bradford, A. (2020). *The Brussels Effect: How the European Union Rules the World*. Oxford University Press.
- European Parliament and Council of the European Union. (2024). Regulation (EU) 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*, L 2024/1689. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)* (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- Zhang, L. (2023, July 19). China: Generative AI measures finalized. *Library of Congress Global Legal Monitor*. <https://www.loc.gov/item/global-legal-monitor/2023-07-18/china-generative-ai-measures-finalized/>