



The EU AI Act:

Why it still appears incomplete

Authored By:

SAFE AI Foundation USA

17 August 2026

Abstract

The European Union's Artificial Intelligence Act (EU AI Act) is the first comprehensive legislative framework designed to regulate artificial intelligence across all sectors. It aims to protect citizens, harmonize regulatory standards, and establish a unified governance structure for AI development and deployment. Yet despite its ambition, the Act remains incomplete. Its risk-classification system lacks quantitative cutoff points, its structural design struggles to keep pace with rapidly evolving AI capabilities, and its enforcement mechanisms are fragmented across member states. This article begins with the historical development and timeline of the Act, followed by an examination of its structural architecture. It then analyzes the Act's purpose, capabilities, shortcomings, lack of practicality, and enforcement challenges. The relationship of the EU Act to ISO42001 standard and OECD.AI principles are discussed, followed by an evaluation of the Act's future trajectory and likelihood of success. The article concludes that while the EU AI Act is a landmark achievement, it requires substantial refinement to function as a practical and effective governance instrument for advanced AI systems.

1. Historical Development

The EU AI Act did not emerge suddenly; it is the product of a multi-year evolution shaped by political negotiation, technological change, and growing societal concern. Its origins trace back to 2018 and 2019, when the European Commission began publishing ethical guidelines for trustworthy AI and convening expert groups to explore regulatory options. These early discussions framed AI as a technology requiring not only innovation policy but also rights-based governance, consistent with the EU's broader digital strategy.

Momentum accelerated in April 2021, when the Commission released the first official proposal for an AI regulation. This initial draft introduced the current central and risk-based framework and proposed bans on certain practices such as social scoring and manipulative AI. Over the next two years, the proposal underwent intense scrutiny within the European Parliament and Council. Thousands of amendments were debated, reflecting disagreements over biometric surveillance, the treatment of general-purpose AI models, and the definition of high-risk systems. The legislative process became a negotiation between innovation, safety, and fundamental rights.

By March 2024, the European Parliament formally adopted the Act, and in May 2024 the Council of the European Union gave its final approval. The Regulation was published in the Official Journal of the European Union on 12 July 2024, triggering the countdown for its entry into force. On 1 August 2024, the Act officially entered into force, but its obligations were

designed to apply gradually. The first set of provisions, primarily bans on prohibited AI practices, became applicable on 2 February 2025. These included restrictions on subliminal manipulation, exploitation of vulnerable populations, social scoring by public authorities, certain biometric categorization practices, and emotion recognition in workplaces and schools.

A second major milestone occurred on 2 August 2025, when rules for general-purpose AI models, governance structures, and penalties began to apply. Providers of general-purpose AI models were required to maintain technical documentation, comply with copyright obligations, and, for models with systemic risk, conduct evaluations and incident reporting. Member states were also required to designate national authorities, while the European AI Office and AI Board became operational.

Further obligations will continue to phase in through 2026, 2027, and 2028. Transparency duties, including labelling of AI-generated content and deepfake disclosures, effective from 2 August 2026. High-risk AI systems listed in Annex III, such as those used in hiring, credit scoring, and biometrics, will become fully applicable by 2 December 2027. High-risk AI embedded in regulated products covered by Annex I will apply from 2 August 2028. Some obligations extend even further, with long-stop dates around 2030 for legacy public-sector systems. This timeline reveals both the ambition of the EU and the complexity of translating a broad regulatory vision into operational reality.

2. Structural Overview

The EU AI Act is a layered regulatory architecture composed of interlocking components. At its foundation lie the general provisions, which define the scope of the Regulation, the meaning of an AI system, and the roles of providers, deployers, importers, and distributors. These provisions also introduce AI literacy obligations, requiring organizations to ensure that staff interacting with AI systems possess sufficient understanding of their operation and risks. The core of the Act is its risk-classification framework. AI systems are divided into four categories: unacceptable risk, high risk, limited risk, and minimal risk with each associated with different regulatory consequences. Unacceptable risk systems are banned outright, reflecting the EU's commitment to protecting fundamental rights. High-risk systems, particularly those affecting safety or rights in domains such as employment, education, law enforcement, and critical infrastructure, are subject to stringent obligations. Limited-risk systems must meet transparency requirements, while minimal-risk systems face no specific obligations. See Figure 1.

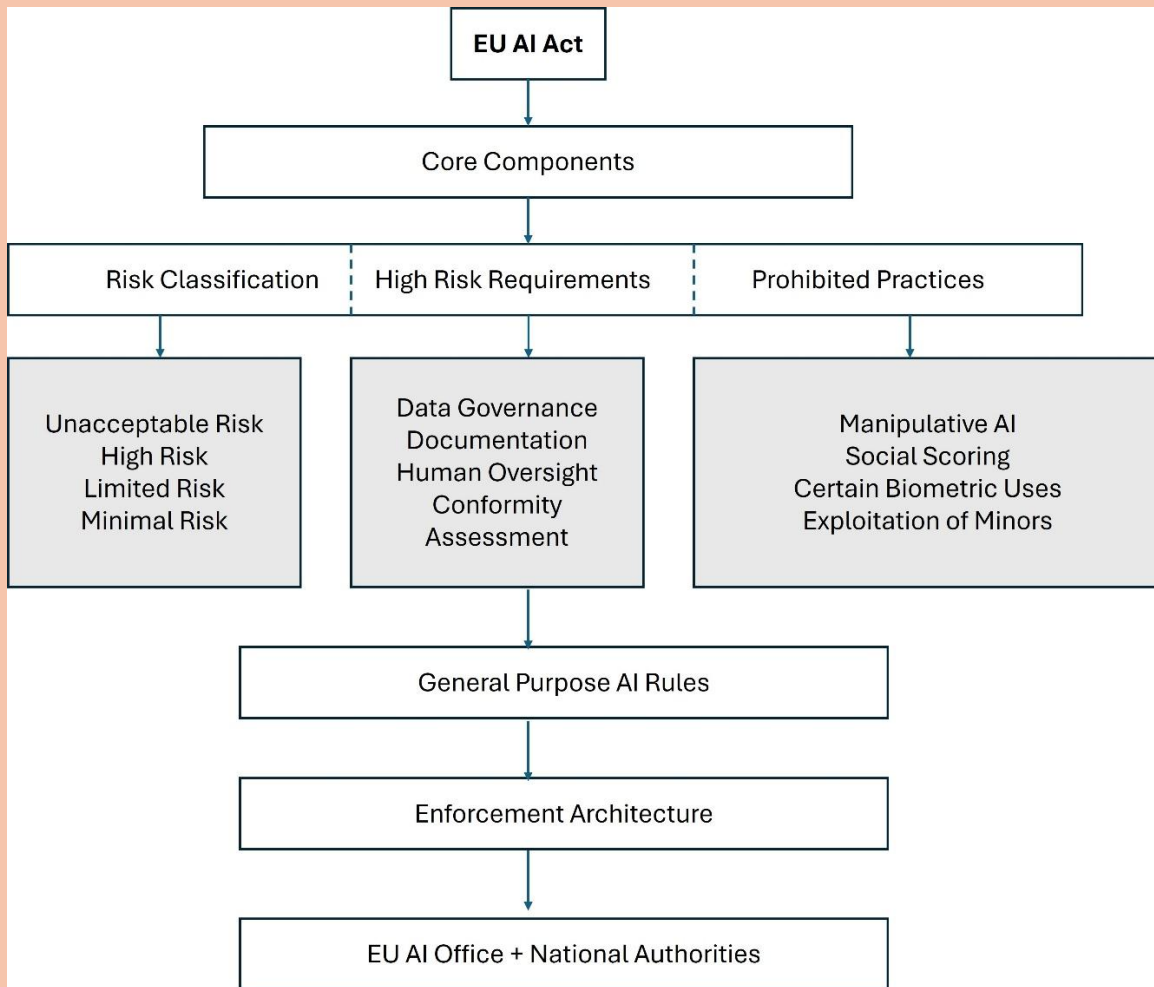


Figure 1 — Structural Overview of the EU AI Act

Surrounding this classification layer are detailed requirements for high-risk systems. These include obligations relating to data governance, documentation, logging, human oversight, robustness, and conformity assessment. Providers must demonstrate compliance before placing high-risk systems on the market, often through notified bodies designated by member states. Parallel to these requirements are specialized rules for general-purpose AI models. Recognizing that large, versatile models can be adapted to many downstream uses, the Act imposes documentation, transparency, and for models with systemic risk, additional evaluation and incident reporting duties [14].

Finally, the Act's enforcement architecture consists of the European AI Office, the AI Board, and national competent authorities. These bodies coordinate supervision, issue guidance, and enforce compliance. Penalties for violations can reach significant percentages of global turnover, underscoring the seriousness of the obligations.

3. Purpose of the EU AI Act

The purpose of the EU AI Act is both protective and harmonizing. It seeks to shield individuals and society from harmful AI applications while creating a unified regulatory environment across the EU. The protective dimension is evident in the bans on certain practices and the strict obligations imposed on high-risk systems. The harmonizing dimension is reflected in the effort to replace fragmented national approaches with a single coherent framework.

The Act also aims to provide legal certainty. By defining categories of risk and associated obligations, it gives organizations a clearer understanding of what is required to place AI systems on the EU market. This is particularly important for cross-border deployments, where differing national rules could otherwise create confusion and barriers to trade. Finally, the Act aspires to promote trustworthy AI. By embedding principles such as transparency, accountability, and fairness into legal obligations, it seeks to align technological development with European values and fundamental rights.

4. What the EU AI Act Can Do

In practical terms, the EU AI Act can ban harmful AI practices, regulate high-risk systems, and impose transparency obligations on limited-risk applications. It can require providers of high-risk systems to implement robust risk-management processes, maintain detailed documentation, ensure human oversight, and undergo conformity assessments. It can compel providers of general-purpose AI models to disclose information about training data, comply with copyright rules, and, where systemic risk is present, conduct evaluations and report incidents. The Act also empowers regulators. National authorities and the European AI Office can investigate non-compliance, impose fines, and coordinate enforcement. Over time, the Act can shape industry norms, encouraging organizations to integrate ethical and safety considerations into their design and deployment processes.

5. Shortcomings of the EU AI Act

The most prominent shortcoming of the EU AI Act lies in its risk-classification system. While the four-tier structure appears clear at a conceptual level, it is built on qualitative descriptions rather than quantitative thresholds. Terms such as “significant risk,” “substantial impact,” and “high likelihood” are not defined numerically [10]. As a result, organizations often struggle to determine whether a given system is high-risk or limited-risk, and regulators may interpret the same system differently [9].

The Act's principal risk classification is substantially use-case [12] and context based rather than organized around general frontier-model capability thresholds. A simple model used in a sensitive domain may be classified as high-risk, while a powerful model used in a seemingly benign context may fall into a lower category. This approach can miss the technical realities of Frontier AI, where capabilities such as autonomous decision-making, advanced manipulation, or cyber-offense potential may pose serious risks regardless of the immediate use case.

Issue	Description	Impact
Qualitative definitions	Risk tiers rely on subjective language rather than measurable thresholds.	Creates uncertainty and inconsistent classification.
Use-case-based regulation	Risk depends on context, not capability.	Powerful models may be under-regulated; simple models may be over-regulated.
Ambiguous boundaries	No clear cutoff points between risk categories.	Companies struggle to self-assess compliance obligations.
Regulatory fragmentation	Member states may interpret categories differently.	Leads to uneven enforcement across the EU.
Legal vulnerability	Ambiguity invites disputes and litigation.	Slows deployment and increases compliance costs.

Table 1: Disadvantages of the EU AI Act’s Risk Classification System

6. Practical Challenges

Beyond conceptual issues, the EU AI Act faces serious practical challenges. The absence of measurable thresholds means that compliance often becomes an exercise in interpretation rather than a straightforward application of rules. Organizations must invest in legal and regulatory expertise simply to understand where their systems fit within the framework [13].

The static nature of the regulation also clashes with the dynamic evolution of AI. Frontier models can gain new capabilities rapidly, and emergent behaviors may not fit neatly into pre-defined categories. The Act does not establish a general, continuously updated capability-

monitoring and capability-threshold mechanism comparable to frontier-model capability evaluation frameworks.

The burden of compliance is another practical concern. Documentation, risk management, and conformity assessments require significant resources. While large companies may absorb these costs, small and medium enterprises may find them costly and prohibitive, potentially discouraging innovation within the EU.

Practical Issue	Description	Consequence
No measurable thresholds	Risk categories lack quantitative criteria.	Compliance becomes guesswork.
Static regulatory model	Cannot adapt to rapidly evolving AI capabilities.	Fails to govern frontier systems effectively.
Heavy compliance burden	Documentation and conformity assessments are costly.	Disadvantages SMEs and slows innovation.
Misalignment with technical reality	Regulates contexts rather than capabilities.	Blind to emergent dangerous behaviors.
Political compromise	Vague language used to satisfy multiple stakeholders.	Reduces clarity and operational usefulness.

Table 2: Lack of Practicality in the EU AI Act

7. Difficulty of Enforcement

Even if the Act were conceptually and practically flawless, enforcement would remain a formidable challenge. The EU’s regulatory system is decentralized: each member state has its own supervisory authority, with varying levels of resources and expertise. Coordinating enforcement across 27 jurisdictions, while maintaining consistency, is inherently difficult.

The extraterritorial scope of the Act adds another layer of complexity. Because it applies to any AI system affecting EU citizens, organizations outside the EU must also navigate its requirements. This can create friction with other regulatory regimes, particularly in the United States, where AI governance is more sectoral and less centralized.

Enforcement Challenge	Description	Result
Decentralized oversight	27 member states enforce the Act independently.	Inconsistent application and regulatory fragmentation.
Coordination complexity	Requires alignment among national authorities and the European AI Office.	Slow response to emerging risks.
Extraterritorial scope	Applies to non-EU companies affecting EU citizens.	Creates international friction and compliance confusion.
Resource disparities	Some member states lack sufficient regulatory capacity.	Uneven enforcement quality.
Legal disputes	Ambiguity invites challenges in national and EU courts.	Delays implementation and increases costs.

Table 3: Enforcement Difficulties of the EU AI Act

8. Comparisons with ISO/IEC 42001 and OECD AI Principles

The EU AI Act occupies a unique position in the global AI governance landscape, and its relationship with other major frameworks, particularly ISO/IEC 42001 standards and the OECD.AI principles, reveals both alignment and divergence.

ISO/IEC 42001 is a voluntary management-system standard that focuses on how organizations internally govern their AI development processes. It emphasizes documentation, risk management, accountability structures, and continuous improvement cycles. In contrast, the EU AI Act is a legal instrument that regulates AI systems themselves, not the organizational processes behind them. While ISO/IEC 42001 can support compliance with certain aspects of the Act, such as documentation discipline or internal governance, it cannot substitute for the Act’s legally mandated conformity assessments, technical documentation requirements, or high-risk obligations. The EU AI Act does not reference ISO/IEC 42001, nor does it require compliance with it, underscoring the fundamentally different domains they occupy.

The OECD.AI standards, meanwhile, provide a set of high-level principles for trustworthy AI, including fairness, transparency, accountability, and safety. These principles influenced the early ethical discussions that preceded the EU AI Act, and many of the Act's provisions reflect similar values. However, the OECD framework is intentionally broad and non-binding, designed to guide policy rather than impose legal obligations. The EU AI Act transforms these principles into enforceable rules, but in doing so introduces complexity and rigidity that the OECD guidelines avoid. Where OECD.AI emphasizes flexibility and adaptability, the EU AI Act relies on fixed categories and prescriptive obligations. This divergence highlights a tension between global normative frameworks and region-specific regulatory instruments: the former aim to inspire best practices, while the latter attempt to enforce them through law. The EU AI Act aligns with the spirit of OECD.AI but diverges sharply in operational detail.

9. Progression Ahead

As the EU AI Act begins its phased implementation, its current state reflects both promise and uncertainty. The establishment of the European AI Office and the rollout of obligations for general-purpose AI models mark significant steps toward operationalizing the Regulation. Yet many of the Act's most demanding requirements — particularly those governing high-risk systems, remain years away from full enforcement. This creates a transitional period in which organizations must prepare for obligations that are not yet fully defined, while regulators must build capacity to supervise a rapidly evolving technological landscape. The Act's early enforcement will likely focus on prohibited practices and documentation requirements, but the real test will come when high-risk systems and systemic-risk models face scrutiny under the Regulation's more stringent provisions [11].

Looking ahead, the success of the EU AI Act will depend on its ability to adapt to technological change and resolve its internal ambiguities. If the EU introduces measurable capability thresholds, clarifies risk categories, and strengthens coordination among national authorities, the Act could evolve into a robust governance framework. However, if ambiguity persists and enforcement remains fragmented, the Act may struggle to achieve its goals. Its precautionary approach may protect citizens from certain harms, but it may also slow innovation and push frontier AI development to jurisdictions with more flexible regulatory environments. The next few years will determine whether the EU AI Act becomes a global model for AI governance or a cautionary example of regulatory overreach.

10. Conclusion

The EU AI Act is a landmark in the global evolution of AI governance. It represents a serious attempt to align technological development with fundamental rights, safety, and ethical principles. Its risk-based framework, bans on harmful practices, and strict obligations for high-risk systems demonstrate a genuine commitment to protecting citizens and creating a trustworthy AI ecosystem. It has demonstrated clear unity in AI among the EU nations.

Yet the Act is incomplete. Its reliance on qualitative risk categories without clear, quantitative cutoff points undermines legal certainty and invites inconsistent interpretation. Its focus on use cases rather than capabilities leaves it inadequately equipped to address the emergent risks of Frontier AI. Its static design struggles to keep pace with dynamic technological change, and its heavy compliance burdens may dampen innovation, particularly among smaller organizations. Enforcement is complicated by decentralized oversight, resource disparities, and extra-territorial reach.

For the EU AI Act to fulfill its promise, future revisions will need to introduce measurable capability thresholds, dynamic evaluation mechanisms, clearer classification criteria, and more streamlined enforcement structures. Only then can the Act move from being a bold conceptual framework to a truly practical and effective instrument for governing advanced AI systems.

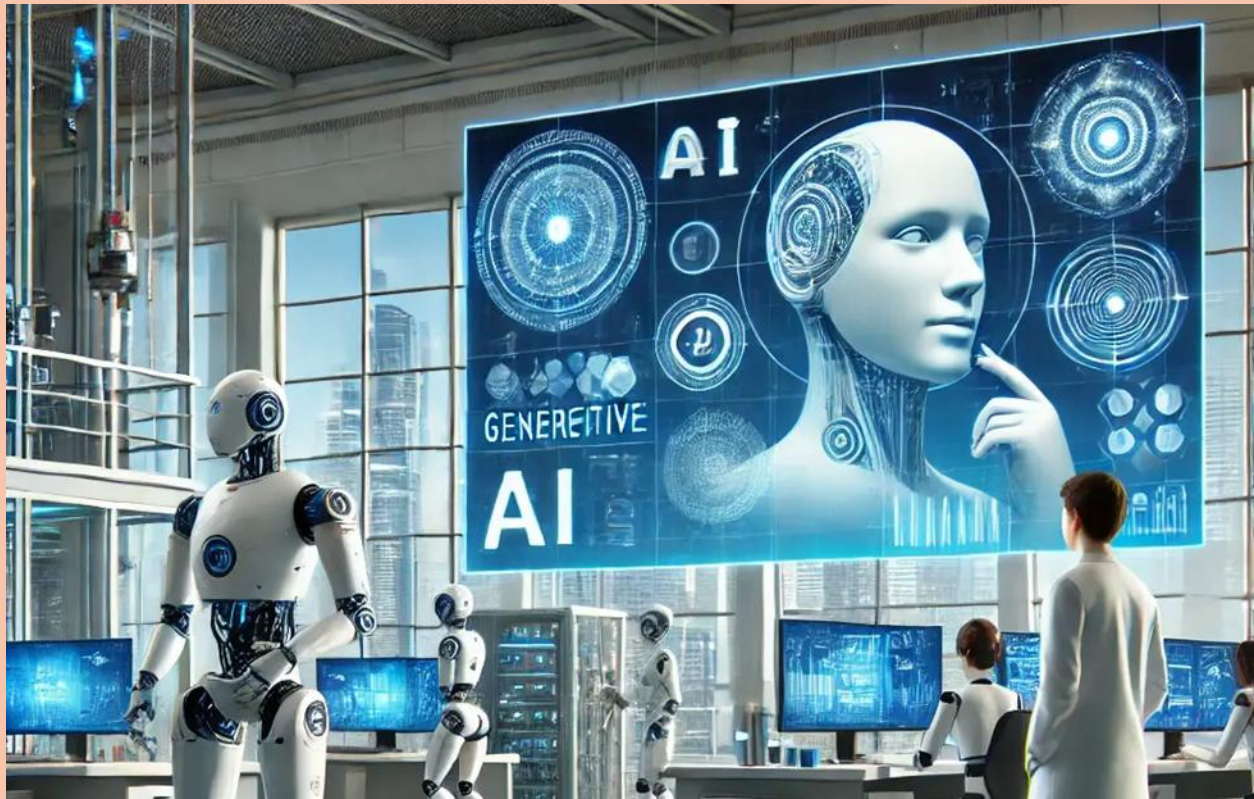
References

1. European Commission. *Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)*.
2. Regulation (EU) 2024/1689, Official Journal of the European Union, 12 July 2024.
3. AI Act Service Desk. *Timeline for the Implementation of the EU AI Act*.
4. Lawbster. *EU AI Act Compliance Timeline: Every Key Date*.
5. OECD AI Policy Observatory.
6. NIST AI Risk Management Framework.
7. Google DeepMind. *Frontier Safety Framework*.
8. SAFE AI Foundation. *SAIOP: Safe AI Operating Point*.
9. Ebers, M. (2025). Truly Risk-based Regulation of Artificial Intelligence: How to Implement the EU's AI Act. *European Journal of Risk Regulation*
10. Novelli, C., Casolari, F., Rotolo, A., Taddeo, M., Floridi, L., et al. (2024). AI Risk Assessment: A Scenario-Based, Proportional Methodology for the AI Act. *Digital Society*, 3, 13.

11. Carey, S. (2025). Regulating Uncertainty: Governing General-Purpose AI Models and Systemic Risk. *European Journal of Risk Regulation*.
12. Gstrein, O. J., Haleem, N., & Zwitter, A. (2024). *General-purpose AI regulation and the European Union AI Act*. *Internet Policy Review*.
13. Teichmann, F. (2026). Risk, Reasonableness and Residual Harm under the EU AI Act: A Conceptual Framework for Proportional Ex-Ante Controls. *European Journal of Risk Regulation*.
14. Bommasani, R., Hau, A., Klyman, K., & Liang, P. (2024). *Foundation Models and the EU AI Act*. RegML 2024.
15. *Rebuilding the pyramid: The AI Act's risk-based approach using a binary decision diagram*. *Computer Law & Security Review*, 2025, 106189.
16. Policy briefs from the Centre for AI Safety and the Future of Humanity Institute.

Disclaimer: This technical report is provided free of charge by the SAFE AI Foundation USA. The information in this report is provided “as it is”, by the SAFE AI FOUNDATION, USA. The use of the information provided here is subject to the user’s own risk, accountability, and responsibility. The SAFE AI FOUNDATION is not responsible for the use of the information by the user or reader. Please reference this article if you wish to cite it elsewhere.

Note: We do not work with consultancy firms or any third parties to produce this document. All copyrights are reserved by the SAFE AI Foundation USA. Donations and partnerships are welcome. Questions and enquires about this work can be directed to the executive director Dr. Chai K Toh at: contact@safeaifoundation.com



SAFE AI FOUNDATION

www.safeaifoundation.com

The SAFE AI Foundation was formed in 2025, in the heart of the Silicon Valley, California, USA. Its mission is to embrace the AI revolution and promote the use of good AI to improve on the quality of life and the economy. It will address the issues of bad AI and its impacts on society. The foundation will inform, educate, and mitigate the dangers of bad AI and help the society to embrace good AI. The foundation will interact with industries, universities and public sector to jointly advance the progress of Good AI. The foundation is a non-profit 501c(3) organization, registered in the State of California, USA.