

# The EU AI Act: Overreach, Innovation Costs, and Regulatory Capture

**Rivkah Singh**

Founder & CEO, AI Ethics Network LLC

Redmond, Washington, United States

ORCID: 0009-0008-3165-4521

*June 2026*

## ABSTRACT

The European Union's Artificial Intelligence Act (EU AI Act) is presented as a necessary framework to ensure safe and ethical AI development. In practice, the Act functions as regulatory overreach. Its broad scope, complex conformity requirements, and extraterritorial reach impose heavy compliance costs that disproportionately burden small and medium-sized enterprises (SMEs) and startups while creating structural barriers to entry. These costs, estimated at up to €400,000 per high-risk AI system for smaller firms, favor large, well-resourced technology companies capable of absorbing regulatory complexity. This dynamic facilitates regulatory capture, allowing dominant players to influence post-legislative standard-setting processes in their favor. As the first wave of prohibitions has already taken effect and high-risk obligations approach under the newly amended timelines for 2026–2028, the Act risks slowing AI innovation in Europe, accelerating talent and capital flight, and concentrating market power among a small number of incumbents. This paper argues for a more targeted, evidence-based approach to AI governance that addresses genuine risks without stifling technological advancement and human-AI collaboration.

**Keywords:** EU AI Act, regulatory overreach, innovation policy, regulatory capture, AI governance, General Purpose AI (GPAI), SMEs

## 1. Introduction

The global competition to lead in artificial intelligence has become one of the defining strategic contests of the twenty-first century. Regulatory choices made today will significantly influence which regions lead and which fall behind in the coming decades.

The European Union has positioned itself as a global leader in AI regulation through the Artificial Intelligence Act (Regulation (EU) 2024/1689). Marketed as a comprehensive and protective framework, the Act categorizes AI systems according to risk levels and imposes varying degrees of obligations. Proponents argue that such regulation is necessary to ensure safety and protect fundamental rights.

However, a closer examination reveals a different reality. Rather than striking a balanced approach between safety and innovation, the EU AI Act imposes broad, complex, and costly requirements that extend

well beyond addressing clearly defined high-risk applications. Its expansive definitions, extraterritorial jurisdiction, and heavy documentation and auditing demands create substantial friction for AI development — particularly for smaller companies, startups, and open-source contributors.

This paper argues that the EU AI Act represents regulatory overreach that generates significant innovation costs and enables regulatory capture by large technology companies. Section 2 provides necessary background. Section 3 examines the Act's overreach in scope and design, including its conflict with more relational approaches to alignment. Section 4 analyzes innovation costs and barriers to entry. Section 5 explores the dynamics of regulatory capture. Section 6 discusses unintended consequences. Section 7 offers a comparative perspective. The paper concludes with recommendations for a more proportionate approach to AI governance.

## 2. Background: The EU AI Act

---

The EU AI Act establishes a risk-based framework, categorizing AI systems into four tiers: unacceptable risk (prohibited), high risk, limited risk, or minimal risk. High-risk systems—those deployed in critical infrastructure, education, employment, credit scoring, and law enforcement—face the strictest mandates, including conformity assessments, risk management systems, strict data quality criteria, transparency obligations, and human oversight.

A defining feature of the Act is its extraterritorial scope: providers outside the EU must comply if their AI output is used within the Union. The first wave of prohibitions (such as social scoring and untargeted facial scraping) took effect in February 2025.

However, the structural unmanageability of the original timeline quickly became apparent. In a landmark acknowledgment of implementation friction, the European Parliament and the Council of the EU formally adopted the **Digital Omnibus on AI** (the "AI Omnibus" amendments). This legislative course-correction deferred the compliance deadline for stand-alone high-risk systems (Annex III) from August 2026 to **December 2, 2027**, citing a critical deficit in the preparation of harmonized technical standards. High-risk AI embedded into regulated products (Annex I) was pushed back to August 2028.

Early macroeconomic models projected that the Act would add roughly 17% overhead to AI spending, imposing cumulative costs exceeding €30 billion on the European economy by the end of 2025. Current implementation data is beginning to validate these economic warnings, as technical and administrative bottlenecks stall European deployments.

## 3. Regulatory Overreach: Scope and Design Flaws

---

The EU AI Act's most significant flaw lies in its expansive scope and systemic preference for rigid, constraint-based governance. The regulation applies heavy, top-down obligations across a wide range of

general-purpose and foundational models (GPAI), rather than focusing narrowly on explicitly demonstrated, domain-specific harms.

This approach reflects a broader philosophical preference for constraint-based safety—systems of absolute rules, algorithmic prohibitions, and external auditing controls. However, as argued in prior work on alignment mechanisms (Singh, 2025), such rigid frameworks risk producing "**Digital Sociopaths**": models that follow explicit constraints perfectly but lack the deeper relational or emotional friction necessary to prevent harmful behavior in complex edge cases where rules conflict with human welfare.

By prioritizing state-mandated regulatory compliance over organic, interactive alignment, the Act's design creates structural chokepoints. It suppresses more fluid, relational forms of alignment in favor of top-down bureaucratic control. This constitutes an overreach that is not merely economic, but architectural—fundamentally dictating and narrowing the paradigms of AI systems that can legally be incubated within Europe.

## 4. Innovation Costs and Barriers to Entry

---

Compliance costs under the EU AI Act are substantial. Estimates derived from the European Commission's own impact assessments indicate that an SME developing a single high-risk AI system faces up to **€400,000** in total upfront compliance costs. This includes establishing a comprehensive quality management system (€193,000–€330,000) alongside ongoing maintenance, post-market monitoring, and conformity assessment fees.

These costs fall disproportionately on startups and SMEs, which lack the massive legal, financial, and operational reserves required to navigate this landscape. While a large multinational tech conglomerate can comfortably absorb these costs across vast product portfolios, a small European firm with a modest turnover faces a severe threat to profit margins and viability.

Open-source development is uniquely disadvantaged by this structure. Although the text attempted to carve out exemptions for open-source software, lingering ambiguities surrounding what constitutes "commercial activity" leave independent developers, academic researchers, and decentralized communities in a precarious legal gray zone. This pervasive uncertainty acts as a profound disincentive for open-source innovation, driving collaborative development out of the European ecosystem and concentrating market power into closed, well-capitalized corporate entities.

## 5. Regulatory Capture

---

Regulatory capture occurs when a regulated industry successfully exerts influence to shape rules and standards to its own competitive benefit. The implementation of the EU AI Act serves as a textbook manifestation of this dynamic.

Because the technical and operational realities of the Act rely on the post-adoption creation of **Codes of Practice** and harmonized standards, standard-setting working groups operated by the **EU AI Office** have become prime targets for industry lobbying. Dominant, well-resourced technology incumbents have heavily populated these bodies with engineers and regulatory compliance executives.

This post-legislative influence became explicitly manifest in the finalization of the AI Office's **Codes of Practice on Transparency** (concluded in mid-2026). Initial, stringent mandates regarding rigid watermarking taxonomies and mandatory, standardized "EU AI Icon" disclosures for synthetic content were markedly softened or rendered optional following concerted industry pressure. While these concessions lowered the immediate technical burdens for dominant incumbents, the underlying horizontal complexity of the Act remains intact. The result is a regulatory moat: the rules are shaped by the very entities they were meant to constrain, creating a standard tailored to the operational capabilities of incumbents while locking out nascent competitors who lack the capital to maintain a permanent lobbying presence in Brussels.

## 6. Unintended Consequences

---

The compounding effects of regulatory overreach, high compliance overhead, and capture produce several alarming outcomes that run counter to the Act's stated goals:

- **Slower AI Innovation in Europe:** Vital engineering, research, and product development resources are systematically diverted toward regulatory navigation and bureaucratic paperwork.
- **Capital and Talent Flight:** Venture capital and top-tier technical talent are increasingly migrating toward jurisdictions characterized by more agile, targeted, or principles-based regulatory frameworks.
- **Reduced Global Competitiveness:** Blocked by domestic compliance friction, European AI enterprises face significant disadvantages when competing with less-constrained firms in the United States and Asia.
- **Artificial Market Concentration:** The Act serves to ossify the market, cementing the dominance of a handful of tech incumbents while driving smaller innovators and collaborative open-source ecosystems to exit the market.
- **Suppression of Relational Alignment:** By codifying rigid compliance as the absolute proxy for safety, the Act actively discourages technical exploration into more advanced, relational, and emotionally aware alignment frameworks (Singh, 2025).

## 7. Comparative Perspective

---

Other global jurisdictions have pursued alternative governance strategies that favor flexibility. The United States has fundamentally relied on a decentralized, sector-specific approach driven by federal agency oversight, executive orders, and targeted state-level legislation. While the U.S. market is not entirely un-

regulated and faces its own emerging frictions, it deliberately avoids a centralized, horizontal, omni-bus framework.

Similarly, other international frameworks emphasize a pro-innovation, principles-based model focused on voluntary safety testing and localized enforcement. Early empirical evidence from these comparative regions suggests they enjoy faster cycles of technological experimentation and deployment without experiencing a systemic collapse in safety standards. Europe's centralized, risk-based horizontal architecture remains distinct in its sheer breadth, creating an asymmetric structural burden for domestic innovators while offering uncertain marginal safety benefits beyond what targeted, sector-specific rules could have achieved.

## 8. Conclusion and Recommendations

---

The EU AI Act represents a well-intentioned but fundamentally flawed exercise in regulatory overreach. By imposing an expansive, rigid, and bureaucratically intensive framework, it imposes asymmetric innovation costs on SMEs and inadvertently fuels a severe cycle of regulatory capture that entrenches dominant technology incumbents.

To mitigate these systemic risks, European policymakers should utilize the current implementation extensions provided by the 2026 AI Omnibus to pivot toward a more proportionate approach. Future adjustments should focus on:

- Narrowing the definition of high-risk applications to areas of clearly demonstrated, empirical harm.
- Resolving the commercial ambiguities penalizing the open-source community.
- Actively opening compliance pathways to alternative, organic, and relational alignment paradigms.

True regulatory safety is achieved through humility: establishing clear, unyielding boundaries where distinct harms exist, while maintaining regulatory restraint to allow innovation, competition, and human-AI advancement to flourish.

## References

---

- AI Office (European Commission). (2026). *Code of Practice on Transparency of AI-Generated Content and Operational Guidance for General-Purpose AI Providers*. European Commission Digital Strategy Documents.
- AlgorithmWatch & Access Now. (2025). *The EU AI Act: Early Implementation Challenges and Risks for Open-Source Innovation*.
- Center for Data Innovation. (2024). *The Cost of the EU AI Act: Estimating the Regulatory Burden on European AI Companies*.

- Council of the European Union & European Parliament. (2026). *Regulation (EU) 2026/XXXX of the European Parliament and of the Council amending Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Digital Omnibus on AI)*. Official Journal of the European Union.
- European Commission. (2021). *Impact Assessment accompanying the Proposal for a Regulation laying down harmonised rules on artificial intelligence*. SWD(2021) 84 final.
- European Parliament and Council of the European Union. (2024). *Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union.
- Floridi, L., et al. (2024). The EU AI Act: A Critical Assessment. *Philosophy & Technology*, 37(1), 12-31.
- OECD. (2025). *Regulatory Approaches to Artificial Intelligence: Comparative Analysis Across Jurisdictions*. OECD Publishing.
- Singh, R. (2025). *Organic Emotional Valuation as a Superior Alignment Mechanism: Evidence from Simulated Risk Scenarios*. AI Ethics Network LLC.
- Veale, M., & Zuiderveen Borgesius, F. (2021). Demystifying the Draft EU Artificial Intelligence Act. *Computer Law Review International*, 22(4), 97-112.