

When AI Outpaces the Law: A Framework for Ethical AI Governance

Antonio Paganelli
paganelli.a57@gmail.com

Christian Senk
csenk@walsh.edu

Philip Kim
pkim@walsh.edu
Walsh University
North Canton, Ohio 44720, USA

Joseph V. Homan
joehoman3@gmail.com
Robert Morris University
Moon Township, PA 15108, USA

Abstract

Artificial intelligence (AI) has rapidly evolved beyond many existing legal and organizational governance structures, creating ethical, social, and accountability challenges. This paper examines consumer-facing AI systems in high-risk domains and asks how information-systems governance can translate ethical principles into enforceable institutional controls. Using a qualitative comparative legal-policy method, the study analyzes the EU Artificial Intelligence Act (EU AI Act), the U.S. National Institute of Standards and Technology (NIST) AI Risk Management Framework (AI RMF), and two illustrative U.S. disputes involving copyright and privacy. The analysis identifies a recurring gap between voluntary standards, reactive litigation, and the need for operational oversight. It then proposes a three-part governance architecture: NIST-led standards development, a Federal Artificial Intelligence Administration (FAIA) for risk-based review and post-deployment monitoring, and a separate enforcement authority for violations. The contribution is not simply an “FDA for AI”; rather, it combines standards, lifecycle review, and independent enforcement while addressing the speed and distributed supply chain of modern AI systems.

Keywords: Artificial Intelligence Governance, AI Ethics, Public Policy, and Federal Governance Framework

When AI Outpaces the Law: A Framework for Ethical AI Governance

Antonio Paganelli, Christian Senk, Philip Kim and Joseph V. Homan

1. INTRODUCTION

Artificial intelligence (AI) has become a foundational information technology that increasingly mediates decisions, services, and interactions. This paper focuses on large-scale consumer-facing AI systems used in high-risk domains such as healthcare, education, finance, biometric identification, and surveillance. These systems are socio-technical rather than purely technical: outcomes depend on models, data, platforms, deployers, users, organizational controls, and applicable law. That characteristic places AI governance squarely within the information-systems (IS) conversation about decision rights, control, accountability, and the alignment of technology with organizational and societal objectives (Abedin, 2022; Tiwana, 2014).

The paper makes three related contributions. First, it develops an explicit governance lens that connects five principles—legality, transparency, proportionality, responsibility, and predictability—to observable institutional controls. Second, it uses a comparative legal-policy analysis to show how the EU AI Act and NIST AI RMF distribute decision rights, documentation duties, review, and enforcement differently. Third, it derives a three-pronged U.S. framework from the gaps identified in the comparative and case analyses. The framework separates standard setting, lifecycle oversight, and punitive enforcement rather than treating governance as either voluntary self-regulation or a single centralized approval process.

The study addresses three sub-questions: (1) Descriptive: What governance mechanisms do the EU AI Act, NIST AI RMF, and the selected U.S. cases reveal? (2) Analytical: Where do those mechanisms leave gaps when evaluated against the five ethical criteria? and (3) Normative: What institutional design could close those gaps without imposing blanket pre-deployment review on all AI? The paper argues that AI should be governed through proportionate, lifecycle-based safeguards rather than broadly “constrained,” with stronger controls applied where potential harms, affected rights, or systemic dependencies are greater.

2. METHODOLOGY

This study uses a qualitative comparative legal-policy design. The units of analysis are (a) the EU AI Act, (b) the NIST AI RMF, and (c) two U.S. legal disputes: *The New York Times Company v. Microsoft Corporation et al.* and *Lopez et al. v. Apple Inc.* The EU and NIST instruments were selected because they represent contrasting governance logics—binding risk-based regulation and voluntary risk-management guidance—while both are sufficiently developed to support structured comparison. The two cases were purposively selected as illustrative “harm-class” cases: copyright/training-data governance and consumer privacy/data capture. They are not intended to be statistically representative.

Sources were selected from the manuscript’s existing scholarly references, primary legal and policy authorities, and official government materials. Priority was given to peer-reviewed or established scholarly sources for theoretical claims and to primary authorities for legal status. The search and verification process emphasized AI governance, algorithmic accountability, information-systems governance, the EU AI Act, NIST AI RMF, U.S. executive orders, federal legislation, and the selected case dockets. The legal-policy cutoff is July 31, 2026. Secondary sources were used to contextualize issues when primary materials were not readily available, but allegations, party positions, court holdings, and unresolved merits are distinguished.

The five ethical criteria are operationalized as observable governance requirements: legality (statutory authority, jurisdiction, due process); transparency (documentation, disclosure, auditability, and contestability); proportionality (risk-tiered controls); responsibility (named actors, reporting, and corrective authority); and predictability (published criteria, target timelines, appeals, and structured

revision). The same criteria are applied to the EU/NIST comparison, both cases, and the proposed framework. This creates a reproducible chain from literature and source selection to case analysis and institutional recommendations.

3. LITERATURE REVIEW

The literature review is organized around the paper's governance problem rather than by application area. It first examines AI governance and IS perspectives, then uses healthcare and education as high-risk examples, and finally synthesizes the unresolved principle-practice gap. This organization better connects the literature to the research questions and to the institutional design developed later in the paper.

3.1 AI Governance and IS Foundations

AI governance literature increasingly treats accountability as an institutional relationship rather than a property that can be achieved by publishing an ethics statement. Floridi et al. (2018), the OECD (2019), and UNESCO (2021) emphasize complementary principles involving human agency, transparency, responsibility, safety, and fairness. Doshi-Velez and Kim (2017) show why technical interpretability is itself an empirical and methodological problem, while Kroll (2018) highlights governance problems arising when data systems are difficult to understand or contest. Abedin (2022), published in the IS field, further shows that explainability can create competing effects: greater explanation may improve trust and accountability but can also expose complexity, uncertainty, or strategic behavior. These studies suggest that principles alone do not produce accountability; organizations need decision rights, evidence requirements, monitoring, and mechanisms through which affected parties can contest outcomes.

3.2 Healthcare

The literature, therefore, points to a principle-practice gap: organizations may endorse transparency, accountability, or fairness while lacking procedures that make those principles observable and enforceable. For AI, the gap is intensified by rapid model updates, opaque supply chains, and platform-mediated deployment. The central IS research issue is consequently not whether ethical principles are desirable, but how governance mechanisms translate them into repeatable organizational and regulatory practices.

Healthcare illustrates why lifecycle governance matters. AI can support diagnosis, treatment planning, and predictive medicine, but patient-facing systems create high consequences when errors, privacy failures, or unanticipated model behavior occur. Research on AI-assisted medical applications demonstrates potential benefits in imaging and clinical decision support (Banerjee et al., 2020; Muntean et al., 2023; Orte Cano et al., 2023). At the same time, patients expect AI applications to reflect their values and preferences rather than displace human judgment (Adus et al., 2023). These findings support proportional review, documentation, human oversight, and post-deployment monitoring rather than a one-time technical certification.

3.3 Education

The literature also cautions against treating application domains as interchangeable. The governance requirements for a medical diagnostic system differ from those for a low-risk recommendation engine. A risk-based IS governance model should therefore specify which decisions are automated, who owns the decision, what evidence is required, and how users can contest or correct outcomes.

Education provides a second high-risk context because AI affects not only efficiency but also learning, assessment, and professional judgment. Research suggests that interactive digital learning technologies can improve comprehension and retention under particular conditions (Hung et al., 2017/2018). However, generative AI creates new questions concerning assessment, academic integrity, and the changing distribution of cognitive work (Chiu, 2023). The relevant governance issue is not whether AI is "good" or "bad" for education, but whether institutions establish clear decision rights, disclosure rules, human review, and assessment practices appropriate to the technology's role.

The education literature also illustrates why governance should remain adaptive. AI capabilities and student practices can change faster than institutional policies. Rather than prescribing one permanent rule, governance should establish review triggers—for example, material model changes, new uses of sensitive data, or evidence of adverse effects—that require reassessment.

3.4 Ethics and AI Governance

Accordingly, the literature review treats healthcare and education as illustrative high-risk contexts, not as exhaustive sector studies. Their role is to demonstrate why AI governance must operate across the technology lifecycle and why the same ethical principles require different operational controls depending on context.

AI ethics scholarship provides a normative foundation, but it also exposes tensions that a governance framework must manage. Floridi et al. (2018) emphasize beneficence, non-maleficence, autonomy, justice, and explicability; the OECD (2019) emphasizes inclusive growth, human-centered values, transparency, robustness, and accountability; and UNESCO (2021) emphasizes human rights, human oversight, fairness, and sustainability. These frameworks overlap substantially with the principles used here, but do not prescribe identical institutional mechanisms. This paper treats its five principles as governance criteria rather than as a claim that they constitute a universal ethical code.

3.5 Literature Synthesis

The literature supports a governance architecture that converts principles into operational requirements. Transparency becomes documentation, testing evidence, change logs, and user-facing disclosures; responsibility becomes named accountable actors and incident-reporting duties; proportionality becomes risk tiers and differentiated review; predictability becomes published criteria and appeal procedures; and legality becomes clear statutory authority and due process. This operationalization is the bridge between AI ethics scholarship and IS governance practice.

The literature establishes three unresolved issues. First, AI governance research recognizes the need for accountability but provides competing views about how much authority should reside with government, firms, or third parties. Second, existing standards such as NIST AI RMF provide flexible practices but are not themselves a comprehensive enforcement regime. Third, legal disputes show that litigation can clarify individual rights without creating a systematic lifecycle governance structure. This paper addresses the gap by linking ethical criteria, comparative governance analysis, and case evidence to an institutional design.

4. LEGISLATIVE POLICY REVIEW

The resulting analytical logic is explicit: the five principles define what responsible governance should accomplish; the EU/NIST comparison identifies how governance authority is currently allocated; and the two cases illustrate where reactive legal mechanisms fail to provide timely, repeatable controls. The proposed three-pronged framework is therefore derived from the preceding analysis rather than introduced independently.

As a legal-policy baseline, this study uses a cutoff date of July 31, 2026. The analysis distinguishes enacted law, agency standards, executive policy, proposed legislation, and judicial proceedings. This distinction is important because the U.S. federal AI landscape changed materially after 2023: Executive Order 14110 was rescinded by EO 14148 on January 20, 2025; EO 14179 directed the removal of barriers to U.S. AI leadership; EO 14365 addressed a national AI policy framework and state-law conflicts; and EO 14409, issued June 2, 2026, emphasized advanced-AI innovation and security. These actions illustrate both the reach and the instability of executive-led governance (The White House, 2025a, 2025b, 2025c; 2026).

The United States also regulates particular AI applications through existing agencies and sector-specific law. FDA oversight of AI-enabled medical devices, financial regulators' authority over regulated institutions, consumer-protection law, employment law, privacy statutes, and other sectoral regimes mean that the U.S. model is not simply "unregulated." The more precise problem is fragmentation: authority is distributed across agencies and legal regimes that may not align around common AI lifecycle

controls.

Congressional action remains limited at the federal level. H.R. 2794, the NO FAKES Act of 2025, was introduced on April 9, 2025, and referred to the House Judiciary Committee. Its stated purpose is to protect intellectual-property rights in individuals' voice and visual likeness, illustrating targeted legislative intervention rather than a comprehensive AI governance regime (H.R. 2794, 2025).

Executive action can move quickly but is inherently less stable than legislation. The rescission of EO 14110 in January 2025 provides direct evidence of this limitation. For IS governance, the implication is that firms cannot treat executive policy as a durable substitute for statutory authority, especially when compliance investments span multiple technology generations.

4.1 EU AI Act and U.S. AI Risk Management Framework Comparative Analysis

NIST AI RMF 1.0 provides a contrasting U.S. governance instrument. Its functions—govern, map, measure, and manage—support continuous organizational risk management, but the framework is voluntary unless incorporated into a contract, regulation, or other binding requirement (NIST, 2023). The U.S. sectoral model can therefore preserve flexibility, but it can also leave gaps when no agency has a clear mandate or when actors have incentives to underinvest in safeguards.

A fixed comparison makes the analytical difference clearer. The EU AI Act is legally binding and risk-based, with defined obligations, supervisory structures, and penalties. NIST AI RMF is voluntary and process-oriented, emphasizing organizational risk management. The EU model supplies stronger external accountability; the NIST model supplies greater adaptability. Neither alone resolves the pace problem for rapidly changing, distributed AI systems.

For generative and general-purpose AI, the EU framework is broader than a four-tier classification alone. It includes obligations that attach to general-purpose AI models, governance and documentation requirements, conformity assessment for applicable high-risk systems, market surveillance, and enforcement. The phased application of obligations also illustrates a useful governance principle: implementation can be staged rather than requiring every control to take effect simultaneously.

The NIST approach is valuable because it can be updated without the legislative cycle required for statutory rules. Its weakness is enforceability. The proposed framework, therefore, uses NIST as the standards layer rather than attempting to replace it with a second technical bureaucracy.

The comparison suggests a hybrid design: use flexible technical standards to keep pace with innovation, but attach mandatory review and enforcement to clearly defined high-risk uses. This avoids treating the EU and U.S. approaches as complete opposites and recognizes that U.S. sectoral regulation already supplies binding controls in particular domains.

The policy problem is thus one of governance architecture: how should decision rights, evidence requirements, review authority, and enforcement be distributed across a rapidly changing AI ecosystem? The next sections apply the five ethical criteria and the two case studies to that question.

5. Legality

The five criteria used in this paper are derived from recurring themes in AI ethics and governance literature rather than presented as self-evident principles. They are operationalized below so that the same criteria can be applied to the two cases and to the proposed institutional design.

5.1 Transparency

The criterion is therefore evaluated by asking: Who has legal authority to act? What conduct is covered? What evidence is required? What procedural protections apply?

Transparency concerns access to information needed by users, regulators, and affected parties to understand and contest AI-related decisions. It includes technical documentation, risk assessments, material-change disclosures, and appropriate user disclosures; it does not require public release of

proprietary model weights or trade secrets.

5.2 Proportionality

Proportionality requires that the intensity of governance correspond to the potential severity and scale of harm. Low-risk applications should face lighter controls; high-risk systems should face stronger testing, documentation, human oversight, and monitoring. This criterion also guards against regulatory designs that impose fixed costs that disproportionately burden smaller firms.

5.3 Responsibility

Responsibility assigns accountability to identifiable actors across the AI lifecycle. Foundation-model providers, downstream developers, deployers, and operators may each control different sources of risk. Responsibility cannot be assigned solely to the model developer.

5.4 Predictability

Predictability concerns both regulatory expectations and system governance. Firms need stable rules and published procedures, while users need confidence that materially similar cases will be treated consistently. Predictability does not mean freezing rules; it means making adaptation structured and transparent.

Operationally, predictability requires published review criteria, target timelines, written determinations, appeal procedures, and scheduled reassessment of standards.

Together, the criteria create a practical evaluation test. Legality asks whether authority is valid; transparency asks whether relevant evidence is accessible; proportionality asks whether controls fit risk; responsibility asks who must act; and predictability asks whether decisions and revisions are consistent and contestable. These criteria are then applied to the two cases and used to derive the institutional response.

6. CASE ANALYSIS AND IMPLICATIONS

6.1 Case I: The New York Times Company v. Microsoft Corporation et al.

This case was selected because it represents an archetypal intellectual-property governance problem in generative AI: the relationship among training data, model development, output reproduction, and copyright. It is illustrative rather than statistically representative. The case is especially useful because it tests whether existing copyright doctrine can supply timely governance for a rapidly evolving technical practice.

The New York Times alleges that OpenAI and Microsoft used copyrighted articles without authorization in developing AI systems and that some outputs could reproduce portions of Times content. The defendants have asserted defenses, including fair use. These are allegations and defenses, not findings of liability.

On April 4, 2025, the U.S. District Court for the Southern District of New York issued an opinion that allowed some claims to proceed while dismissing others. The opinion did not resolve the ultimate fair-use question. As of July 31, 2026, the merits of the broader dispute remained unresolved. The procedural posture demonstrates why the paper treats litigation as evidence about governance gaps rather than as a substitute for a comprehensive regulatory regime.

The case exposes a transparency and responsibility problem as well as a copyright problem. Training practices, data provenance, memorization, model architecture, and output controls are distributed across technical and organizational actors. Litigation can determine liability for particular conduct, but it does not by itself create a repeatable system for documenting training data, evaluating risks before deployment, or monitoring material model changes.

6.2 Case I: Governance Implications

Evaluated against the five criteria, the case reveals: legality—existing copyright doctrine supplies authority but leaves novel AI-training questions uncertain; transparency—training-data provenance and model behavior can be difficult for outside parties to assess; proportionality—one-size-fits-all rules may burden low-risk uses while missing high-risk reproduction; responsibility—multiple actors may control different stages of the lifecycle; and predictability—uncertain doctrine creates unstable expectations for both creators and developers.

6.3 Case II: Lopez et al. v. Apple Inc.

Lopez was selected as a complementary illustrative case because it represents a different archetypal harm class: consumer privacy and unintended data capture by an AI-enabled voice assistant. The pairing therefore, tests the framework across intellectual-property and privacy risks rather than treating either case as representative of all AI disputes.

The settlement agreement was executed on December 31, 2024. Plaintiffs alleged that Siri-enabled devices could record private conversations following unintended activations and that information was shared with third parties. Apple denied wrongdoing. The court preliminarily approved the settlement on February 10, 2025, and the matter proceeded through the court-approval process during 2025. The settlement should therefore be distinguished from a judicial finding that Apple violated privacy law.

The case demonstrates a limitation of litigation as governance. A settlement can provide compensation and closure for a defined class without establishing a general rule for future voice assistants or other continuously sensing systems. It also illustrates the importance of post-deployment governance because privacy risk may arise from real-world use, unexpected activation, system changes, or organizational data practices.

6.4 Case II: Governance Implications

Against the five criteria, Lopez highlights legality through questions about consent and statutory authority; transparency through uncertainty about when data are captured and how they are used; proportionality through the need to distinguish ordinary device functions from intrusive data collection; responsibility through the relationship among device maker, software provider, contractors, and downstream data users; and predictability through the absence of a general lifecycle rule for unintended data capture.

6.5 Cross-Case Synthesis

The two cases are not offered as a representative sample; they are selected to illuminate two recurring governance problems—intellectual property and privacy—while spanning different stages of the AI lifecycle. The comparison strengthens the argument that litigation is necessary but insufficient: it resolves discrete disputes after harm is alleged, whereas governance must establish repeatable controls before and after deployment.

The analytical bridge can be summarized as follows: five criteria → observable governance gaps → institutional response. Legality and predictability support statutory authority and appeal rights; transparency supports documentation and reporting; proportionality supports tiered review; and responsibility supports named actors, incident reporting, and enforcement. This chain is the principal methodological link between the case analysis and the policy proposal.

6.6 Limits of Litigation and Implications for Legislative Action

Litigation remains important because courts interpret existing law, create precedent, and can provide remedies. Its limitation is temporal and institutional: cases generally arise after a dispute, depend on parties with standing and resources, and produce decisions tied to particular facts and legal doctrines. A regulatory framework can complement litigation by establishing baseline process obligations before deployment and by creating monitoring and corrective mechanisms afterward.

The cases do not demonstrate that courts are ineffective. They demonstrate that courts and regulators perform different governance functions. The proposed framework assigns standards development, lifecycle oversight, and punitive enforcement to distinct but coordinated institutions while leaving courts with judicial review and adjudication.

7. POLICY RECOMMENDATION: A THREE-PRONGED AI GOVERNANCE FRAMEWORK

The proposed framework applies to high-risk consumer-facing AI systems and to foundation-model or platform actors when their systems are used in covered high-risk contexts. It does not create universal pre-market licensing for every AI model. The design responds directly to the pace problem identified in the literature and cases: regulation must be mandatory where risk warrants it, but flexible enough to accommodate frequent technical change.

7.1 Prong One—Standards Development: NIST

NIST should remain the technical standards body because the AI RMF already provides an adaptable vocabulary of govern, map, measure, and manage. Standards should be updated on a scheduled basis and through expedited procedures when material risks emerge. NIST should not decide individual approvals or impose penalties. This separation protects technical standard setting from case-specific enforcement pressures.

The standards layer should define evidence requirements by risk tier: model and system documentation; data provenance and quality evidence where relevant; performance and robustness testing; human-oversight design; cybersecurity and privacy controls; incident-reporting thresholds; and change-management procedures. Standards should distinguish foundation-model providers, downstream developers, and deployers so that obligations follow actual control over risk.

7.2 Prong Two—Federal Artificial Intelligence Administration (FAIA)

FAIA would conduct a risk-based review of covered high-risk deployments rather than approve every AI model. Its jurisdiction would attach primarily to high-risk uses and deployment contexts, with authority to require information from foundation-model providers when their models materially affect those uses. This addresses a central weakness of a simple “FDA for AI” analogy: AI risks are distributed across models, applications, users, and contexts (Carpenter & Ezell, 2024).

For high-risk deployments, the applicant would submit the following: intended use and deployment context, system and model documentation, relevant training data provenance, risk assessment and mitigation plan, evaluation and validation evidence, human oversight procedures, privacy and security controls, incident response plan, and a change management plan. Trade secrets and security-sensitive information could be protected from public disclosure while remaining available to regulators under confidentiality rules.

FAIA should use three review pathways: streamlined review for established low-risk changes; standard review for new high-risk deployments; and enhanced review for systems with potentially systemic or rights-sensitive effects. Rather than treating the previously suggested three-to-twelve-month period as a fixed rule, this paper recommends a target-based service standard: routine reviews should have published target times, with expedited review available when evidence is sufficient and the delay itself creates material public risk.

7.3 Prong Three—Enforcement

Applicants should receive written determinations and a defined administrative appeal process. Judicial review would remain available under applicable law. Small firms, academic systems, open-source projects, and research prototypes should receive tailored pathways or exemptions when they do not operate in covered high-risk deployment contexts, while downstream commercial deployment of an open model would remain subject to the relevant risk-based requirements.

Separating technical standards, routine lifecycle review, and punitive enforcement can reduce role conflicts, but it also creates coordination costs. The framework requires formal information-sharing

agreements, common definitions, joint incident protocols, and a lead-agency rule for overlapping jurisdiction. The goal is not institutional fragmentation; it is functional separation with coordinated execution.

7.4 Coordination, Feasibility, and Implementation

FAIA would not replace the FDA, FTC, EEOC, CFPB, SEC, or other sector regulators. Instead, it would establish cross-sector AI governance requirements and defer to existing agencies on sector-specific substantive law. For example, an AI-enabled medical device would remain subject to FDA authority, while FAIA could address cross-cutting documentation, model-change, and incident-reporting requirements. Congress would need to specify preemption rules carefully because state AI laws increasingly address areas such as discrimination, privacy, and consumer protection. EO 14365 demonstrates that federalism and preemption are already central policy questions.

Implementation would require congressional appropriations, specialized technical staff, legal expertise, audit capacity, and secure information systems. Regulatory fees could supplement appropriations but should be structured so that small firms are not priced out of the market. A phased implementation beginning with a limited set of clearly defined high-risk uses would reduce administrative risk and allow the agency to learn before expanding jurisdiction.

7.5 Funding, Governance, and International Coordination

The proposal also faces legitimate objections. Regulatory capture could favor large firms with compliance resources; excessive pre-deployment review could slow innovation; broad definitions of high-risk AI could expand agency jurisdiction; and information requirements could raise First Amendment, trade-secret, or intellectual-property concerns. These risks support narrow statutory definitions, procedural safeguards, confidential regulatory access, periodic review, and explicit sunset or evaluation provisions for new regulatory powers.

Funding should primarily come from congressional appropriations, supplemented by proportionate fees and civil penalties. Congressional oversight, independent audits, public reporting, and conflict-of-interest rules should apply to the institutions. International coordination should focus on interoperability of testing evidence, terminology, and incident reporting, reducing duplicative compliance for firms operating across jurisdictions while preserving each jurisdiction's substantive legal authority.

Existing "FDA for AI" proposals correctly identify the value of pre-deployment testing but also highlight the difficulty of defining the regulated object, anticipating harms, and governing distributed AI supply chains (Carpenter & Ezell, 2024). The present framework introduces three key elements: it regulates high-risk deployment contexts rather than every model; it distinguishes between standards, routine review, and punitive enforcement; and it treats monitoring and controlled model updates as integral to approval rather than as exceptions to it. Its central claim is therefore institutional and lifecycle-oriented, not simply that AI needs a new licensing agency.

The framework is also deliberately modest about novelty. It synthesizes established governance mechanisms, but its contribution is the explicit derivation of the institutional architecture from the five ethical criteria, the EU/NIST comparison, and the two purposively selected cases. The framework should therefore be understood as a testable IS governance model rather than a claim that the paper has invented an entirely new regulatory instrument.

7.6 Limitations and Alternative Approaches

The framework has limitations. The legal-policy analysis is qualitative and does not estimate compliance costs, staffing requirements, or the probability of regulatory capture. The two cases are illustrative rather than representative, and the framework has not been empirically tested with regulators, firms, or affected users. Alternative approaches—including self-regulation, co-regulation, tort litigation, sector-specific regulation, and market-based certification—may outperform a federal agency in particular contexts. Future empirical work should compare these alternatives and test whether the proposed separation of functions improves accountability without creating excessive delay.

8. CONCLUSION

This study used a qualitative comparative legal-policy design to examine how AI governance mechanisms operate across the EU AI Act, NIST AI RMF, and two illustrative U.S. disputes. The analysis found a consistent institutional gap: ethical principles are widely articulated, but authority, evidence, review, and enforcement are unevenly distributed across the AI lifecycle. The cases show why litigation remains necessary, but cannot, by itself, provide proactive, repeatable governance.

The paper's principal contribution is a three-pronged, lifecycle-oriented governance architecture. NIST supplies adaptable technical standards; FAIA supplies risk-based review, monitoring, and corrective authority for covered high-risk deployments; and a separate enforcement function addresses serious violations. The framework differs from a simple FDA analogy by focusing on deployment context, distributed responsibility, controlled updates, and post-deployment oversight.

Future research should test three questions: whether risk-based lifecycle review can maintain acceptable review times as model-update frequency increases; how governance responsibilities should be allocated across foundation-model providers, developers, deployers, and platforms; and whether the proposed separation of standards, review, and enforcement improves measurable accountability relative to self-regulation, sectoral regulation, or litigation alone. The framework is therefore a bounded policy proposal requiring empirical and comparative testing rather than a claim of finality.

9. REFERENCES

- Abedin, B. (2022). Managing the tension between opposing effects of explainability of artificial intelligence: A contingency theory perspective. *Internet Research*, 32 (2), 425–453. <https://doi.org/10.1108/INTR-05-2020-0300>
- Adus, S., Macklin, J., & Pinto, A. (2023). Exploring patient perspectives on how they can and should be engaged in the development of artificial intelligence (AI) applications in health care. *BMC Health Services Research*, 23 (1), 1–14. <https://doi.org/10.1186/s12913-023-10098-2>
- Banerjee, I., de Sisternes, L., Hallak, J. A., Leng, T., Osborne, A., Rosenfeld, P. J., Gregori, G., Durbin, M., & Rubin, D. (2020). Prediction of age-related macular degeneration disease using a sequential deep learning approach on longitudinal SD-OCT imaging biomarkers. *Scientific Reports*, 10(1), 15434. <https://doi.org/10.1038/s41598-020-72359-y>
- Carpenter, D., & Ezell, C. (2024). An FDA for AI? Pitfalls and plausibility of approval regulation for frontier artificial intelligence. *arXiv*. <https://arxiv.org/abs/2408.00821>
- Chiu, T. K. F. (2023). The impact of generative AI (GenAI) on practices, policies, and research direction in education: A case of ChatGPT and Midjourney. *Interactive Learning Environments*. <https://doi.org/10.1080/10494820.2023.2253861>
- Davidson, T. (2023). The danger of runaway AI. *Journal of Democracy*, 34(4), 132–140. <https://doi.org/10.1353/jod.2023.a907694>
- Dixon Jr., J. H. B. (2023). Artificial intelligence—"What hath God wrought." *Judges' Journal*, 62(3), 37–39.
- Donham, B. P. (2023). It's not just about the algorithm: Development of a joint medical artificial intelligence capability. *Joint Force Quarterly*, 111, 50–57.
- Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv*. <https://arxiv.org/abs/1702.08608>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). *AI4People—An ethical*

- framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Hung, I. C., Kinshuk, & Chen, N. S. (2017). Embodied interactive video lectures for improving learning comprehension and retention. *Computers & Education*, 117, 116–131.
- Kassymova, G. K., Klimova, E. K., Garbuzova, G. V., Kivlenok, T. V., Lavrinenko, S. V., & Arpentieva, M. R. (2023). Socio-psychological problems of digital education in subject development. *Journal of Pharmaceutical Negative Results*, 14, 2162–2172. <https://doi.org/10.47750/pnr.2023.14.S02.255>
- Kroll, J. A. (2018). Data science data governance [AI ethics]. *IEEE Security & Privacy*, 16(6), 61–70. <https://doi.org/10.1109/MSEC.2018.2875329>
- Lee, C.-E., Baek, J., Son, J., & Ha, Y.-G. (2023). Deep AI military staff: Cooperative battlefield situation awareness for commander's decision-making. *Journal of Supercomputing*, 79(6), 6040–6069. <https://doi.org/10.1007/s11227-022-04882-w>
- Lopez, F., Lopez, F., Pappas, J. T., & Yacubian, D. (2024). Settlement agreement and release: Lopez et al. v. Apple Inc. United States District Court for the Northern District of California.
- Muntean, G. A., Marginean, A., Groza, A., Damian, I., Roman, S. A., Hapca, M. C., Muntean, M. V., & Nicoară, S. D. (2023). The predictive capabilities of artificial intelligence-based OCT analysis for age-related macular degeneration progression—A systematic review. *Diagnostics*, 13(14), 2464. <https://doi.org/10.3390/diagnostics13142464>
- National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). U.S. Department of Commerce. <https://www.nist.gov/itl/ai-risk-management-framework>
- OpenAI. (2025). Response to The New York Times lawsuit. <https://openai.com/new-york-times/>
- Organization for Economic Co-operation and Development. (2019). OECD principles on artificial intelligence. <https://oecd.ai/en/ai-principles>
- Orte Cano, C., Suppa, M., & del Marmol, V. (2023). Where artificial intelligence can take us in the management and understanding of cancerization fields. *Cancers*, 15(21), 5264. <https://doi.org/10.3390/cancers15215264>
- Pope, A. (2024, April 10). NYT v. OpenAI: The Times's about-face. *Harvard Law Review*. <https://harvardlawreview.org/blog/2024/04/nyt-v-openai-the-timess-about-face/>
- Price II, W. N. (2021). Problematic interactions between AI and health privacy. *Utah Law Review*, 2021(4), 925–936.
- Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).
- Ryan, M. (2020). In AI we trust: Ethics, artificial intelligence, and reliability. *Science and Engineering Ethics*, 26(5), 2749–2767. <https://doi.org/10.1007/s11948-020-00228-y>
- The White House. (2025a, January 20). Initial rescissions of harmful executive orders and actions (Executive Order 14148).
- The White House. (2025b, January 23). Removing barriers to American leadership in artificial intelligence (Executive Order 14179).

- The White House. (2025c, December 11). Ensuring a national policy framework for artificial intelligence (Executive Order 14365).
- The White House. (2026, June 2). Promoting advanced artificial intelligence innovation and security (Executive Order 14409).
- Tiwana, A. (2014). Platform governance. In Platform ecosystems. Elsevier.
- Tsang, E. (2023). AI for finance (1st ed.). CRC Press.
- U.S. Congress. (2025). H.R. 2794, NO FAKES Act of 2025. 119th Congress. <https://www.congress.gov/bill/119th-congress/house-bill/2794>
- UNESCO. (2021). Recommendation on the ethics of artificial intelligence. United Nations Educational, Scientific and Cultural Organization. <https://unesdoc.unesco.org/>
- Xia, Q., Chiu, T. K. F., Chai, C. S., & Xie, K. (2023). The mediating effects of needs satisfaction on the relationships between prior knowledge and self-regulated learning through artificial intelligence chatbot. *British Journal of Educational Technology*, 54(4), 967–986. <https://doi.org/10.1111/bjet.13305>
- Zelevnikow, J. (2023). The benefits and dangers of using machine learning to support making legal predictions. *WIREs Data Mining and Knowledge Discovery*, 13(4). <https://doi.org/10.1002/widm.1505>