

Article

Artificial Intelligence Regulatory Models: Advances in the European Union and Recommendations for the United States and Evolving Global Markets

Miriam F. Weismann¹

¹ Florida International University, USA

Keywords: Artificial Intelligence, regulatory model, risk-based, European Union, United States, sectoral

<https://doi.org/10.46697/001c.120396>

AIB Insights

Vol. 24, Issue 3, 2024

In 2023, 28 countries and technology private sector representatives signed the Bletchley Declaration, calling for international oversight of AI by applying a risk-based formulary. The goal is to achieve uniform global regulation. The absence of regulatory uniformity challenges companies operating across borders. This paper provides practitioners and policymakers with an overview of the EU regulatory model and a new legislative recommendation, The AI Integrative Risk-Based model for the U.S., which has not designed any regulatory framework. The AIRB provides a useful approach for conducting meaningful risk assessments to guide future U.S. regulation and ensure compliance in global markets.

INTRODUCTION

AI regulatory governance is a pressing issue, a product of technology outpacing the law. In an unprecedented joint venture between governments and private industry, the treaty signatories recognized the “urgent” need for quick action. The treaty outlines a specific risk-based regulatory approach to guide nations in designing regulatory frameworks for achieving AI safety and security. This approach to AI regulation is consistent with treaty harmonization¹ as a global blueprint for regulatory consensus and offers a road to globally verifiable compliance.

However, regulatory progress in the EU has concluded in proposed legislation (expected to go into force in 2024), whereas in the U.S., congressional action remains mostly aspirational. This disparity in progress leaves both U.S. practitioners and policymakers in a quandary. The lack of regulatory clarity poses significant risks as companies speed toward AI development and deployment in the marketplace. Companies cannot predict with certainty when, if, and how the technology will be regulated and whether cross-border laws will achieve uniformity. The U.S. might consider following the treaty example of what can be accomplished through the public and private sectors’ mutual cooperation in acknowledging and addressing AI risk.

THE EU RISK-BASED REGULATORY FRAMEWORK

The European Parliament proposed a risk-based approach to AI regulatory design at the global level in drafting the 2023 Artificial Intelligence Act (AIA, 2023). After defining categories of specific risk identification factors attendant to AI (included in a separate annex), the AIA constructs a four-tier risk model prioritizing these generic risks and attributing levels of the regulatory response to identified risks as follows:

- Unacceptable risk AI. Harmful uses of AI that contravene EU values (such as social scoring by governments) will be banned because of the unacceptable risk they create;
- High-risk AI. A number of AI systems (listed in an Annex) that are creating adverse impacts on people’s safety or their fundamental rights are considered high-risk. To ensure trust and a consistently high level of protection of safety and fundamental rights, a range of mandatory requirements (including a conformity assessment) would apply to all high-risk systems;
- Limited risk AI. Some AI systems will be subject to a limited set of obligations (e.g., transparency);
- Minimal risk AI. All other AI systems can be developed and used in the EU without additional legal obligations beyond those established by existing legislation.

¹ The process by which nation states make changes in their national laws, in accordance with internal legal systems to produce treaty compliance and global uniformity.

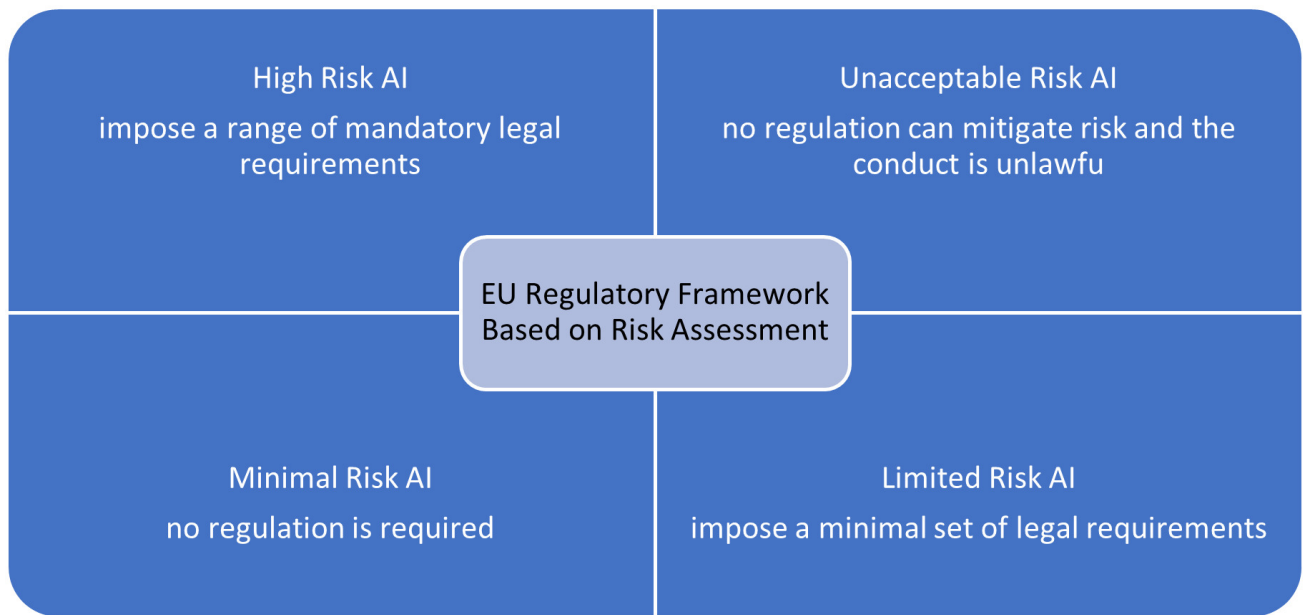


Figure 1. 2023 AIA: EU System of Prioritization of Risk in Regulatory Design

Source: Author

The framework is summarized in [Figure 1](#).

THE U.S. SECTORAL FRAMEWORK

Adopting a risk-based model in the AI domain is still in the early developmental phase in the U.S. However, strides have been made in specific risk identification (such as bias, lack of transparency, privacy concerns, etc.) that are attendant to the use of AI. The U.S. began its foray into AI strategic risk management by Executive Order. Described as a risk-based and sectorally specific model, The AI Leadership Order, Biden [Executive Order 13859](#) (2019), called for a “coordinated Federal Government strategy,” recognizing that AI “will affect the missions of nearly all executive departments and agencies.” The AI Leadership Order further instructed federal agencies to pursue several strategic objectives for “promoting and protecting American advancements in AI.” The salient features of the Order are summarized in [Figure 2](#). Notably, the Order did not require any risk assessment per se but directed agencies to identify risks of AI misapplication in their respective agency. Likewise, the Order did not address risks in the private sector.

Subsequently, the 2023 Biden Executive Order issued. Still a sectoral effort, but more expansive in design and strategy, as indicated in [Figure 3](#), it focuses on four risk areas of concern: AI Safety and Security Risk, Ensuring Responsible Government Use Risk, Equity and Civil Rights Risk, and Risks to Consumers, Patients, Students and Workers. On March 4, 2024, the Office of Management and Budget (OMB) issued a memorandum to advance the goals of the 2023 Executive Order. The memorandum addresses a “subset of AI risks and governance and innovation issues” directly tied to agencies’ use of AI. These identified risks result from “any reliance on AI outputs to inform, influence, decide, or execute agency decisions or actions, which

could undermine the efficacy, safety, equitableness, fairness, transparency, accountability, appropriateness, or lawfulness of such decisions or actions.” The memo specifically prohibits the application of any new agency regulations to non-agency entities and organizations. Thus, as of this date, AI remains unregulated in the U.S.

WHY THE EU HAS OUTPACED THE U.S.

One reason may be economic. Technological development directly impacts market development (Muhleisen, 2018). When regulating the private sector, the U.S. tends to act with restraint so as not to harm or distort natural market forces. However, economic systems differ, which gives rise to different regulatory constraints. A good example is the divergence of U.S. and EU laws regarding patent protections (Czapracka, 2007). The U.S. views patents not as a monopolistic constraint on free market competition but as protection for companies engaged in research and development spending to develop innovative technology and products. The EU, which does not operate under the principles of capitalism but as a single economic market of twenty-seven countries, constitutionally bound together through coordinated economic and fiscal policies, views patents as generally anti-competitive economic behavior under Art. 102 TFEU. This divergence was reflected in the EU decision against Microsoft regarding its interoperability sharing arrangements with competitors, for which Microsoft argued it had patent protection and lost. Arguably, the EU has less reticence in promulgating market constraints than the U.S. and is quicker to do so. As a result, practitioners might expect tighter AI regulation in the EU than in the U.S. For these reasons, in part, EU regulation is not interchangeable with U.S. regulatory goals.

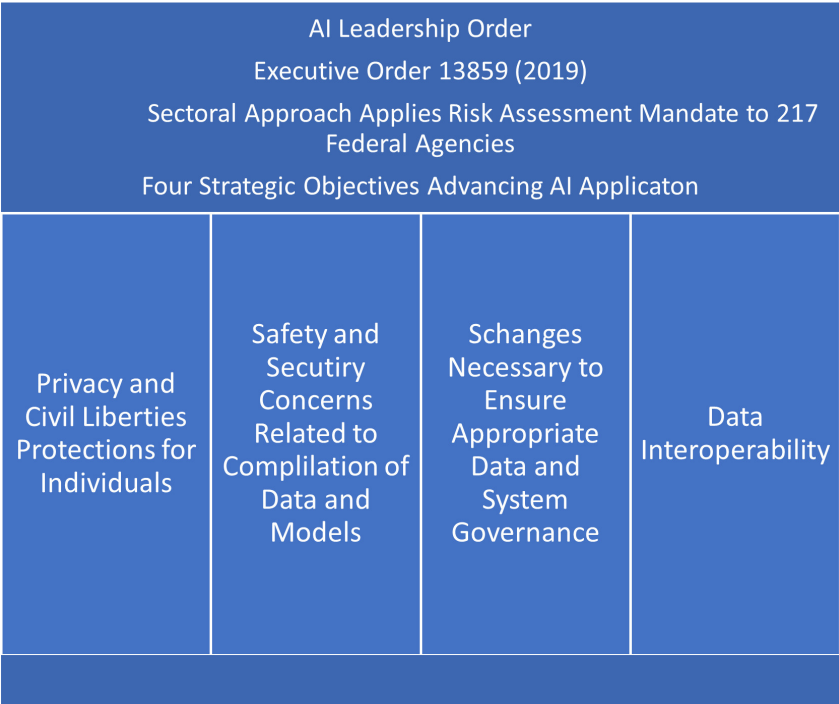


Figure 2. 2019 Risk Assessment Mandate Executive Order 13859

Source: Author

ACTIONABLE RECOMMENDATIONS: AI INTEGRATIVE RISK-BASED (AIRB) MODEL

As previously stated, the lack of regulatory clarity poses significant risks as companies speed toward AI development and deployment in the marketplace. Companies cannot predict with certainty when, if, and how the technology will be regulated and whether cross-border laws will achieve uniformity. To address these issues, this paper recommends that U.S. policymakers and other countries seeking a useful regulatory solution adopt the novel AI Integrative Risk-Based (AIRB) model proposed in this paper. The AIRB model integrates the current U.S. and EU risk assessment models with the International Monetary Fund (IMF) sectoral risk comparison model and provides regulators with a roadmap to design responsive regulations for identified risks. The proposed model is summarized in [Figure 4](#).

The five steps in the model are applied in ascending order. The AIRB model incorporates factors already present in diverse models into a comprehensive version of risk-based assessment with the end goal of matching regulation to risk so as not to over or under-regulate the market.

STEP 1: IDENTIFY RISK

Many articles identify risk factors affecting AI applications (Guan et al., 2022). These risks are identified as “algorithmic risks” and include discrimination, security, interoperability, decision-making, abuse/misuse, technical defect risk, data risk, privacy breach, and others. Appendix A provides a useful roadmap for AI risk identification. Once risks are identified, the next step is prioritizing risk based on

the particular industry sector utilizing AI through a sectoral analysis.

STEP 2: SECTORAL ANALYSIS

As noted above, the Executive Order first targeted the Sectoral Analysis of AI risk in the U.S. in 2019. Sectoral Analysis is a standard tool employed by countries globally (Gobierno del Ecuador, 2023). Likewise, the International Monetary Fund (IMF) conducts Sectoral Risk Assessments in global banking and securities to improve existing regulatory standards by analyzing the application of individual standards in various financial sectors (The World Bank & The International Monetary Fund, 2003).

The AIRB model posits that in Step 2, the initial set of risks identified in Step 1 should be reviewed as part of a sectoral analysis conducted by both the government and private sector to determine which AI risks directly impact their respective processes and operations. To some extent, as explained above, several U.S. federal agencies under Executive Order have begun engaging in a sectoral analysis on an agency-by-agency basis. Then, risk prioritization should follow as in Step 3. Not all market risks are the same.

STEP 3: PRIORITIZE RISK

As noted in [Figure 1](#) above, the EU has crafted a four-tier risk assessment model. This model can be used in concert with the U.S. sectoral model to prioritize and segregate risk categories based on the impact of known risk in a particular sector as identified in the step 2 sectoral analysis. Once prioritized, a regulatory response can be linked to the severity of the risk.

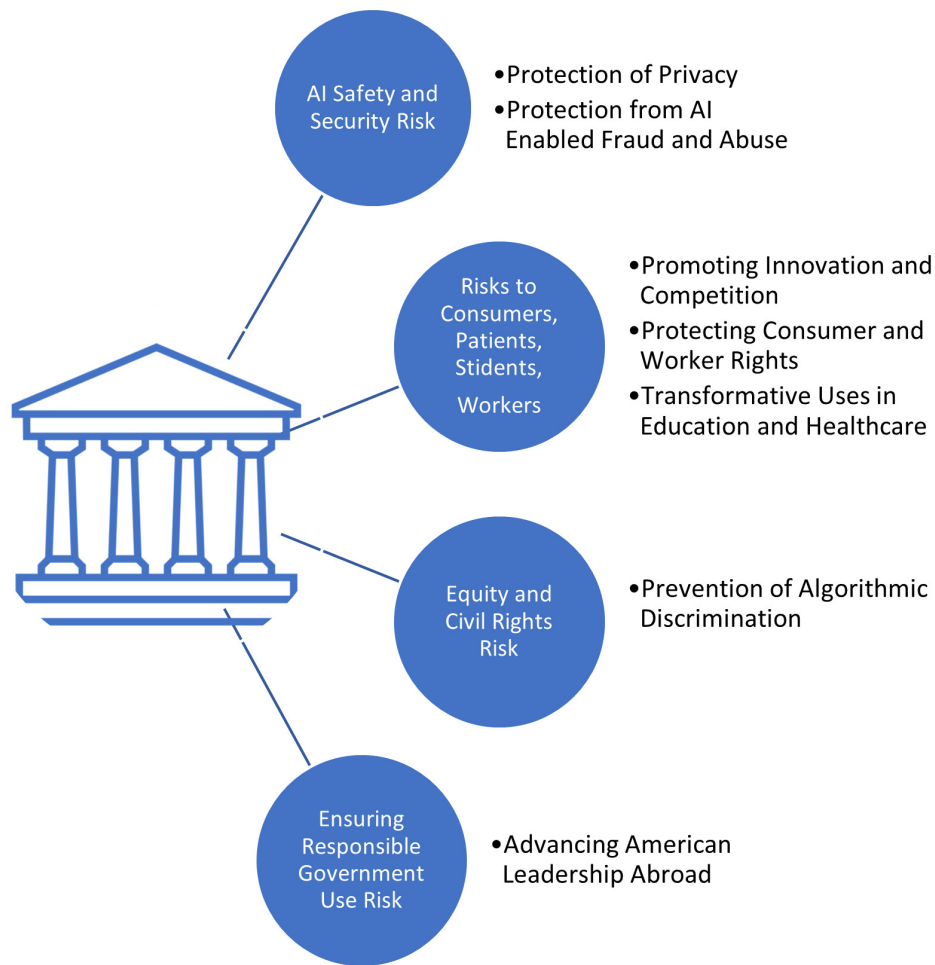


Figure 3. US System of Sectoral Risk Identification Without Assessment and Prioritization in Regulatory Design Biden Executive Order

Source: Author

STEP 4: SECTORAL COMPARISON

The sectoral comparison in the AIRB model is designed to provide greater uniformity in regulatory response. Each sector first identifies and prioritizes risk. To the extent that a sectoral comparison of risk shows overlapping risk between sectors, uniform congressional legislation can be designed to cover a broad base of participants and stakeholders. Regulatory agencies can enact more regulatory-specific guidelines if certain sectors have more individualized needs. Expanding sectoral databases can identify patterns and trends in the market applications. Patterns of behavior can lead to more refined risk analysis.

The International Monetary Fund regularly engages in multi-country comparative sectoral analysis using the Vulnerability Exercise (VE). VE is a cross-country examination identifying country-specific near-term macroeconomic risks. VE has been described as a "key element of the Fund's broader risk architecture" applied in IMF member countries (Panth, 2021). The Global legislative trackers such as Global AI Legislation Tracker by IAPP Research and Insights also provide comparative sectoral data (2023).

STEP 5: DRAFT RESPONSIVE REGULATION

Designing a regulatory framework that addresses the risks created by disruptive technology to ensure the safety of users and the public while at the same time facilitating commercial use is not a simple task. Timing is everything. The proposed regulations can be written, but where the speed of technological innovation outpaces regulatory change, implementation of legislation and regulation can take months or even years. Such laws and regulations may be rendered obsolete by the time that they are finalized. Still, while it is difficult to catch up to technology, regulation need not lag so far behind, as evidenced by the EU's AIA.

WHY THE AIRB MODEL?

This regulatory model is based on risk assessment principles and is differentiated from other risk-based models for several reasons. First, it does not necessarily follow that a risk-based model is a regulatory model. Some risk models are designed merely to ensure compliance with pre-existing regulatory frameworks and are like internal compliance models with the end goal of compliance risk management.

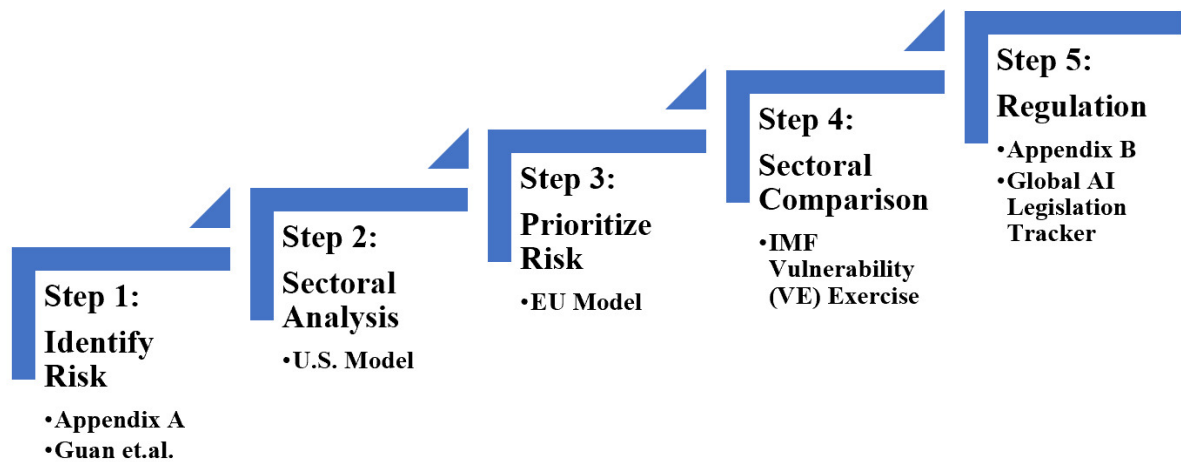


Figure 4. AI Integrative Risk-Based (AIRB) Model

Source: Author

An example is the COSO Enterprise Risk Management Model, now recommended as a safe harbor for publicly traded companies by the SEC. Unlike COSO, the AIRB model aims to achieve demonstrable ethical guardrails for AI based on the promulgation of specific regulations responsive to specifically identified risks. A risk-based approach to regulation “enables a regulator to tailor its regulatory responses so that they are commensurate with the relevant risks.” (Atkins et al., 2016). The AIRB serves as a regulatory guide, not just an internal risk mechanism to achieve compliance.

Second, the AIRB model incorporates a unique regulatory approach combining three different current regulatory approaches operating in the marketplace, including features from the current U.S. and EU risk assessment models, respectively, with the International Monetary Fund (IMF) sectoral risk comparison model. It seizes on the strengths of three current models to suggest a more robust unified regulatory approach for implementation.

Likewise, the AIRB model is risk-specific in terms of regulatory goals. The idea is to do more than identify risk. Other risk frameworks identify risk but overlook the need to include a step that assesses or prioritizes risk. Instead, the AIRB model seeks to tailor the regulatory response to specifically identified risks, recognizing that risk prioritization is critical to managing governmental response and avoiding the problem of over-regulation in a free market economy.

CRITICISM OF THE AIRB MODEL: THE UNRESOLVED LEGAL PROBLEM OF ACCOUNTABILITY

One criticism of this paper’s AIRB model might include that while it provides a systematic approach to AI risk assessment and guided compliance efforts, it does not directly address the accountability problem. The absence of legal accountability predicated on existing law for AI malfunction creates protection gaps for consumers and other human users. However, the problem is deeper than identifying potential victims. From a legal standpoint, it leaves open the question of how regulation should proscribe human accountability for the disruptive failure of AI machine technology. AI, a non-sentient being, is an unregulated algorithm with disruptive potential and falls outside the current U.S. regulatory framework. Because both domestic and international law do not currently recognize AI as a subject of law, AI has no legal personality and, as such, cannot be held personally liable for damages.

The question for policymakers then arises as to how to hold accountable the multiple corporate and non-corporate participants who create, market, and apply AI and create or contribute to the risk of harm in the first instance. In the corporate setting, the legal doctrines of corporate limited liability and separate legal personality protect entrepreneurial activity without the constant risk of legal liability and financial ruin. While these legal doctrines positively foster free enterprise, they can also facilitate the evasion of responsibility and accountability.

This paper does not intend to resolve these unanswered legal questions. Like its predecessor papers, some topics must be left to future research (Allen & Lehot, 2023; Fenwick et al., 2017; Ryan-Mosely, 2024; Whyman, 2023). Rec-

ognizing and discussing areas where this model could be improved is only reasonable.

While many foreseen and unforeseen legal challenges will need to be resolved in the future, the present requires swift regulatory action. Designing a regulatory framework that addresses the risks created by disruptive AI technology to ensure the safety of users and the public while facilitating commercial use is not a simple task. Still, it has been done in the EU, albeit under different market conditions, and can be done effectively in the U.S. and other global markets.

CONCLUSION

AI regulatory governance is a pressing issue, a product of technology outpacing the law. The public and private sectors have coalesced and identified the need for swift regulatory action. The current absence of regulatory clarity poses significant risks as companies speed toward AI development and deployment in the marketplace. Companies cannot predict with certainty when, if, and how the technology will be regulated and whether cross-border laws will achieve uniformity. These pressing concerns need to be addressed with great alacrity.

Regulatory progress in the EU has concluded in legislation, whereas in the U.S., congressional action remains mostly aspirational. This disparity in progress leaves both U.S. practitioners and policymakers in a quandary. To address this pressing need for regulatory clarity in domestic and global markets, the AI Integrative Risk-Based (AIRB)

model provides a useful approach for the U.S. and other global markets to conduct meaningful risk assessments guiding future regulation and ensure compliance in conformity with the public and private sector aspirations reflected in the Bletchley treaty. The features of each integrative step in the model are designed to integrate well-accepted and proven methodologies, combining the advantages of each method in each step of the model with the end goal of formulating a regulatory framework copacetic with free market technological development.

.....

ABOUT THE AUTHOR

Miriam Weismann is a Professor of Business Law and Tax and the Academic Director of the Healthcare MBA at Florida International University. Her research focus includes white-collar crime, cybercrime, financial fraud, corporate governance, international law, taxation, and legal ethics. She has published in prestigious journals, including the *Journal of Business Ethics* and the *Journal of World Business*, and has published three books: *Corporate Crime and Financial Fraud*; *Parallel Proceedings: Navigating Multiple Case Prosecutions*; and *Money Laundering: Legislation, Regulation & Enforcement*.

Submitted: February 11, 2024 EDT, Accepted: June 17, 2024 EDT



This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CCBY-4.0). View this license's legal deed at <http://creativecommons.org/licenses/by/4.0> and legal code at <http://creativecommons.org/licenses/by/4.0/legalcode> for more information.

REFERENCES

- Allen, N., & Lehot, L. 2023, December 7. What to Expect in Evolving U.S. Regulation of Artificial Intelligence in 2024. *Foley & Lardner*. <https://www.foley.com/insights/publications/2023/12/us-regulation-artificial-intelligence-2024/>.
- Atkins, C. et al. 2016, October 18. Clarifying the Difference Between Regulation Designed to Manage Risk and a Risk-Based Approach to Regulation. *Maddocks*. <https://www.maddocks.com.au/insights/15256>.
- Czapracka, K. 2007. WHERE ANTITRUST ENDS AND IPBEGINS – ON THE ROOTS OF THE TRANSATLANTIC CLASHES. *Yale Journal of Law and Technology*, 9: 44–46.
- European Parliament. 2023. *Artificial Intelligence Act of 2023, open for signature June 16, 2023, Legislative Train 10.2023 (AIA)*. <https://www.europarl.europa.eu/legislative-train/carriage/regulation-on-artificial-intelligence/report?sid=7401>.
- Exec. Order No. 13859, Maintaining American Leadership in Artificial Intelligence*. 2019, February 11. . <https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence>.
- Fenwick, M. et al. 2017. Regulation Tomorrow: What Happens When Technology Is Faster than the Law? *Am. Univ. Bus. L. Rev.*, 3(6): 561–567.
- Global AI Legislation Tracker. 2023, August 25. *IAPP Research and Insights*. https://iapp.org/media/pdf/resource_center/global_ai_legislation_tracker.pdf.
- Gobierno del Ecuador, Financial and Economic Analysis Unit. 2023. *Importance of Sectoral Risk Assessment: The Ecuadorian Experience*. https://www.oas.org/es/sms/ddot/gelavex/54/docs/19-PPT_UAFE_DAE_Gelavex23may-ENG.pdf.
- Guan, H. et al. 2022. Ethical Risk Factors and Mechanisms in Artificial Intelligence Decision Making. *Behav. Sci.*, 12: 343.
- Muhleisen, M. 2018. The Long and Short of the Digital Revolution. *International Monetary Fund*. <https://www.elibrary.imf.org/view/journals/022/0055/002/article-A002-en.xml>.
- Panth, S. 2021. How to Assess Country Risk: The Vulnerability Exercise Approach Using Machine Learning. *International Monetary Fund*. <https://doi.org/10.5089/9781513574219.005>.
- Ryan-Mosely, T. 2024. Four Lessons From 2023 That Tell Us Where AI Regulation Is Going. *MIT Review*. <https://www.technologyreview.com/2024/01/08/1086294/four-lessons-from-2023-that-tell-us-where-ai-regulation-is-going/>.
- The World Bank & The International Monetary Fund. 2003. *Analytical Tools of the FSAP*. The World Bank. <https://www.imf.org/external/np/fsap/2003/022403a.pdf>.
- Whyman, B. 2023. AI Regulation is Coming: What is the Likely Outcome? *Center for Strategic and International Studies*. <https://www.csis.org/blogs/strategic-technologies-blog/ai-regulation-coming-what-likely-outcome>.