

# Balancing risk and social good: proactive law as a strategy in AI governance

**Renátó Vági**

Doctoral School of Law  
Eötvös Loránd University  
Egyetem Square 1-3  
Budapest 1053, Hungary

Wolters Kluwer Hungary Kft.  
Budafoki Street 187–189  
Budapest 1117, Hungary  
[renato.vagi@wolterskluwer.com](mailto:renato.vagi@wolterskluwer.com)

**Abstract:** The rapid evolution of artificial intelligence (AI) technologies has transformed numerous sectors, prompting regulators worldwide to establish normative frameworks to ensure the safe, ethical and beneficial deployment of these systems. Over the past two years, various regulatory approaches to the governance of AI have begun to take shape, sparking an ongoing debate. This reveals the limitations of traditional and reactive regulatory approaches in governing evolving technologies, which will fall short in addressing the unique characteristics of AI. However, a shift towards a more promotive direction has become noticeable in recent decades, largely influenced by principles and methodologies resulting from more reflexive public policies, interdisciplinary research, sociotechnical advances and a pragmatism learned from the proactive legal approach endorsed by Nordic legal scholars. This article explores how the proactive law perspective may further enhance AI regulation. It argues that this approach should actively involve stakeholders in regulatory processes to promote collaboration among developers, regulators and the public and ensure that the development and use of AI and other (future) technologies aligns with societal needs and values.

**Keywords:** *AI governance, AI regulation, European Union AI Act, proactive law*

# 1. Introduction

## 1.1 AI regulation in general

The enthusiasm surrounding AI has been prevalent for over a decade. Yet, it was not until the advent of foundational models like ChatGPT that a widespread discourse on the societal dangers and risks associated with such technologies, as well as concurrent regulatory possibilities, truly began. However, the rapid evolution of those technologies is transforming society, prompting regulators worldwide to establish frameworks that ensure safe, ethical and beneficial deployment. In the past two years, various regulatory cultures have adopted differing approaches to the governance of AI. This work focuses solely on the Western regulatory approaches that have a more direct impact on us.

For instance, federal efforts in the United States of America (USA) focus on maintaining and strengthening its leadership position in the field, particularly emphasizing competitiveness (Mazzucato *et al.*, 2022, p. 1). As a result, apart from a few presidential orders that can be considered ethical guidelines, such as the ‘Executive order on maintaining American leadership in artificial intelligence’ (2019) and the ‘Executive order on the safe, secure, and trustworthy development and use of artificial intelligence’ (2023), and rules concerning high-risk automated decision-making systems such as the *Algorithmic Accountability Act of 2019* (2019), which mainly pertain to the use of AI in public administration, there is no comprehensive federal regulation on AI tools (Wang and Wu, 2024, pp. 13–15). According to Mazzucato *et al.* (2022), this approach leads to a value-extractive economic direction, which could create societal and economic problems, such as harmful uses of AI, excessive market monopolization, loss of social control over technological advancement and adverse impacts on the labor market. Moreover, without a proper regulatory effort to assess the extended societal and economic impacts of AI, unexpected losses, issues and other negative effects could take place. Daron Acemoglu (2021) argues that a proactive approach by the government could contribute to a more precise understanding of societal impacts, thereby stimulating a more human-centered technological development.

While there is no comprehensive federal regulation of AI in the US, individual states have sometimes taken different approaches. For example, the State of California had planned to adopt a holistic regulation focused on safety measures for developers of advanced AI models, the *Safe and Secure Innovation for Frontier Artificial Intelligence Models Act*, known as SB 1047.

Despite passing the state legislature, the bill was vetoed in September 2024, citing concerns about its potential impact on the industry's competitiveness and the need for more flexible solutions (Lesnes, 2024). Nevertheless, several AI-related regulations were enacted in the previous month, including those addressing the prohibition of election-related deepfake content (White, 2024), the protection against AI-generated sexual child abuse material and issues related to AI-generated revenge pornography (Nguyen, 2024). California's legislative power has thus begun to move towards a more content- and societal problem-centric regulatory logic to maintain flexibility in governance.

The European Union (EU) has followed a different path. The *Artificial Intelligence Act* ('Regulation (EU) 2024/1689'; hereafter the *AI Act*, or the Act), adopted by the European Parliament on March 13, 2024, and approved by the Council of the European Union in May 2024, was published in the *Official Journal of the European Union* on June 13, 2024, and entered into force 20 days later. In short, the Act aims to provide a comprehensive framework for AI systems within the EU. It is primarily based on a classic risk-based approach, categorizing AI applications by their societal risk levels and assigning to them different legal treatments. Additionally, the Act introduces the concept of general purpose AI (GPAI) ('Regulation (EU) 2024/1689', Article 3(63)) and establishes specific rules concerning it (Madiega, 2024), which in the words of Chun *et al.* (2024, p. 4) resulted in a "parallel governance track for such AI systems."

This regulation has stirred discussion from its preparatory works with critics warning that its implementation may lead to legal uncertainty and excessive compliance costs (Raposo, 2022; Pagallo, Ciani Sciolla and Durante, 2022). Stakeholders, including human rights organizations and business groups, express concerns about insufficient protections and burdensome regulations (Prifti, Quintavalla and Temperman, 2023). Experts state that unresolved issues could lead other jurisdictions to adopt similar flawed regulatory frameworks (Carnegie Endowment for International Peace, 2024). Some have called for clarity and adaptability in the risk classification system to better reflect the complexities of AI applications. Ebers (2024) even argues that the Act falls short of being truly risk-based in several critical aspects. In his view, one of the main limitations is that it leads to problems because of inadequate risk-benefit analysis, reliance on weak empirical evidence and a rigid classification of risks that does not allow for adjustments based on individual cases (see also Ebers *et al.*, 2021). Further, he states that specific regulations for GPAI models are deemed incompatible with a risk-based

approach, mainly due to challenges in foreseeing all potential risks and applications of such models (Ebers, 2024).

It is also worth mentioning that the Act was initially conceived within a product safety context, but was later integrated with a fundamental rights(-based) agenda at the request of the European Parliament, responding to pressure from the European Commission (in respect to the recommendations of the High-Level Expert Group on AI, namely, the *Ethics Guidelines for Trustworthy Artificial Intelligence* (2019)) and its members' expressed commitment to European principles and values, as commented by Tyrangiel (2024). This was subsequently replaced as the legislative process advanced, which, according to the "rationale", represented a more innovation-friendly perspective (Martínez-Ramil *et al.*, 2023, p. 241). The debate continues regarding which of the two approaches is more effective in addressing the economic and social challenges posed by AI systems (Antonov, 2022). While these discussions are important, as the next subsection will demonstrate, additional factors could enhance the distinction between the holistic sector-based approach and the risk-based rights-based approach, enriching the overall discourse. The focus here will be on preventive and proactive law and legal tools and how they can be applied in the governance of AI (Pohjonen, 2006; Schartum, 2006).

The absence of a unified worldwide regulatory concept for AI and the criticism that existing regulations are unstable and underdeveloped are attributable to several reasons, despite being relatively new. One primary reason is that AI is an overly broad umbrella term encompassing various technologies and tools, each operating differently, based on different foundations and eliciting different societal impacts, potentially requiring different regulatory logic. Consequently, many experts prefer not to speak of a unified concept, finding it too elusive (Schuett, 2023). Instead, they focus on regulating specific technologies and application areas, tailoring their approaches to the particularities of each.

Another point, highlighted by Mazzucato *et al.* (2022, p. 2), is that to investigate an AI tool beyond its technical and legal compliance and taking into account its full societal and economic effects, one must observe not only the completed model or solution that has been deployed to users but also the entire ecosystem that was affected through its implementation. This includes data preparation, data management, machine-learning operations (MLOps), IT support and often cloud computing. Since these various steps have an impact on AI technologies, the whole ecosystem is important for developing a responsive and responsible regulatory concept.

Moreover, as noted in the EU *AI Act*, GPAI means that a single solution can be utilized across multiple areas with different impacts, thereby complicating its regulation in all sectors. Additionally, one of the essential characteristics of AI systems is their adaptability; they can respond to changes in their environment and, in some cases, alter their operations, which adds to the challenges because there is no fixed standard that they consistently meet at all times and in all situations. With this in view, the traditional reactive approach to legal regulation will prove alarmingly ineffective.

Recent shifts in technology regulation, nonetheless, seem to move towards a more proactive direction, perhaps inadvertently aligning with the principles advanced over the past two decades by the Nordic School of Proactive Law (Magnusson Sjöberg, 2006, p. 2), scholars in the United States of America and others who throughout the years have committed to similar principles. Consequently, the contributions of this school should be emphasized, as they are now more relevant than ever in handling contemporary complexities due to the increasing demand for practical, forward-thinking legal responses extendable to all fields of law.

## 1.2 Integration of proactive law in AI regulation

This study uses the term ‘proactive law’ to describe not only a strategic approach to lawyering but also a ‘universal tool’ for strengthening social value in regulatory activities as basis for legal risk assessments, systems and processes design and the collaborative management of exchange relations (Solarte-Vásquez and Hietanen-Kunwald, 2020). It originated in the late 1990s and aims to supplement the preventive legal practice (Haapio and Varjonen, 1998) that focused on the judicious “detection of early symptoms of legal trouble” to minimize legal risks (Brown, 1956, p. 947). Proactive law looks amply ahead and draws heavily on the tradition of Scandinavian legal realism, particularly in that it emphasizes the importance of examining the social and economic impacts of legal decisions and instruments (Berger-Walliser, 2012, p. 20).

In the context of the social components in the regulation of technological tools, some literature suggests that proactivity is a salient feature of legal instruments, as they do not address legal situations retrospectively and rather show regulatory efforts that emphasize planning by identifying and defining conditions for behaviors, situations and activities before they come into existence. However, this primarily reflects the general *ex-ante* nature of compliance rules (Fried, 2003) and is closely related to the preventive law approach.

To clarify, it is useful to distinguish between the proactive-preventive and promotive, so-called (perhaps not too complimentary) offensive legal approaches (Schartum, 2006, p. 37). Proactivity, by default, tries to minimize risks and prevent harmful consequences. However, proactive offensive law tries to create incentives and tools for a variety of legal relations that help the parties (and/or the stakeholders) achieve their goals and desirable social and economic impact. Also, it is important to restate acknowledgements by Magnusson Sjöberg (2006) about proactivity in law not being a new phenomenon; after all, regulation generally functions as a proactive and reactive steering mechanism, but proactivity is not a material or validity requirement. And yet, “there is much to be gained by further developing means and methods that place an emphasis on the proactive dimensions of the legal domain” (Magnusson Sjöberg, 2006, p. 3). This is especially true because legal experts are not as aware as they should be of the need to associate across fields, refine their performance and adopt new ways of working, not to be left behind in a world where machines have acquired normative power.

Henceforth, proactive ‘offensive’ law is meant to refer to an understanding closer to the classic Nordic school notion for it can aid in creating appropriate regulatory frameworks for technology governance, covering a range of AI systems. This legal perspective may deviate from the traditional regulatory approach and even the *ex-ante* nature of compliance rules in some cases, but it is an option that fits the purpose of contemporary legal systems facing the reality of the digital and AI transformation, as briefly outlined above. AI and other emerging technologies are moving targets for the regulator, the full scope of risks and potential impacts associated with their deployment remains undetermined. Our current strategic capacity lies in enhancing the law-making process by fostering collaboration with stakeholders and adopting a proactive approach that prioritizes positive outcomes.

The article explores how proactive methodologies can be applied to effective AI regulation (Schartum, 2006, pp. 35–51), as officially recommended within the EU (‘Opinion of the European Economic and Social Committee 2009/C 175/05’, 2009). It speaks of the EU *AI Act* to determine the prevalence of preventive and proactive legal tools in the text of the Council regulation and to identify proactive legal approaches that could enhance the normative efficacy of the Act in the region. The paper argues that the proactive law approach would involve stakeholders in the regulatory process, encouraging collaboration between developers, regulators and the public to align AI development with societal needs and values.

Section 2 presents the key preventive tools currently implemented in the *AI Act*. Section 3 discusses the proactive tools already implemented in the European Union's regulation. Section 4 proposes proactive legal tools that the EU could consider adopting to create a more sustainable regulation focusing on positive social impacts. The conclusion summarizes the article's findings.

## 2. Preventive features of the EU *AI Act*

The EU *AI Act* fundamentally employs a risk-based approach, which characterizes much of the regulation as following a preventive logic and an *ex-ante* nature typical of technological compliance rules. Notably, this approach includes the regulation of prohibited AI tools and the classification of risks. This discussion will provide a detailed examination of these two aspects, emphasizing the proactive and innovation-centric critiques of these regulatory elements.

### 2.1 Prohibited AI practices

One of the most 'preventive' regulatory mechanisms within the EU *AI Act* reflects the governance decision of establishing prohibited AI systems and practices. The regulation contains strict prohibitions to ensure the protection of fundamental rights and societal values. Article 5 of the 'Council regulation (EU) 2024/1689' explicitly enumerates the AI systems and practices that are banned from marketing, deployment or use within the Union as follows: AI systems that use subliminal methods or manipulative and deceptive tactics (Article 5, 1(a)); AI systems designed to exploit the vulnerabilities of specific groups or individuals, particularly due to age, disability or socioeconomic status, in ways that could materially distort their behavior and result in significant harm (Article 5, 1(b)); AI practices involving evaluating or classifying individuals or groups based on their social behaviors or personal characteristics, which can lead to unjustified or disproportionate treatment (Article 5, 1(c)). The use of AI to predict an individual's likelihood of committing a crime based solely on profiling or personality traits is restricted, except when such systems support human assessments grounded in objective and verifiable facts (Article 5, 1(d)). It is unlawful to deploy AI to create or expand facial recognition databases by indiscriminately scraping images from the internet or CCTV footage

(Article 5, 1(e)). The application of AI systems to infer emotions in workplaces and educational settings is prohibited unless explicitly justified for medical or safety purposes (Article 5, 1(f)). AI systems that categorize individuals based on biometric data to infer sensitive personal attributes such as race, political beliefs or sexual orientation are strictly prohibited (Article 5, 1(g)). The deployment of real-time remote biometric identification systems in public areas for law enforcement purposes is highly restricted and only permissible under stringent conditions, including prior authorization by a judicial or independent administrative authority, alongside rigorous fundamental rights impact assessments and strict adherence to data protection laws (Article 5, 1(h)).

Although this list reflects the perceptions of the regulators in the EU regarding socially dangerous applications and their awareness of permissible exceptions, under certain circumstances, the express ban is inflexible, because the *AI Act* does not outline the criteria used for this compilation. This suggests it may have been formed arbitrarily or based on a constrained understanding of threats, which are inherently unstable due to the particularities of human interpretation and external factors such as the changing technological environment. This lack of a sharp evaluative framework complicates the applicability of the provisions to systems that might pose serious societal risks but do not correspond to the enumeration of Article 5, making this a partially preventive rather than a proactive rule. Lack of determination hinders the developers' ability to comply and promotes either over-regulatory activities (having to explain, extend, amend and complement legislation) or suboptimal regulatory results (when the regulatory function is neglected).

## 2.2 The AI risk classification

Another significant preventive measure is the classification model embedded within the risk-based approach of the *AI Act*, which categorizes AI systems according to risk levels. The systems that classify as high-risk due to their potential to adversely affect public health, safety or fundamental rights are subjected to stringent requirements before deployment in the EU market. This category includes AI systems used either as safety components of products or as standalone products subject to EU harmonization legislation, especially those requiring a third-party conformity assessment (Article 6(1)). The obligations include the establishment of a continuous risk management and monitoring strategy, ensuring operations with high-quality data

and detailed documentation of the system's performance. Additionally, providers must pass a conformity assessment and obtain CE certification before entering the market, along with demonstrating transparency and implementing mechanisms for human oversight.

The *AI Act* dictates that any AI system may be considered high-risk if it is intended for use as outlined in Annex I of the Council regulation or if it falls under specific use cases mentioned in Annex III unless it can be demonstrated that the system poses no significant risk to individuals' rights or safety. Exceptions also exist for AI systems designed for narrowly defined procedural tasks, enhancements of prior human activities or analytical tasks that do not independently alter decision-making processes without human oversight (Article 6(3)).

The fact that the developers of systems belonging to this category face more responsibilities and must meet additional requirements could guide manufacturers towards developing AI tools that pose lower risks, but also offer lesser societal benefits if introduced to the market. Furthermore, the classification criteria are not entirely clear either. Although the *AI Act* includes provisions mandating the European Commission, in consultation with the European AI Board, to provide guidelines by February 2026 for the practical implementation of such classification, including examples of high-risk and non-high-risk applications (Article 6(5)), the current uncertainty may inhibit innovation. Models and tools being developed now will later be evaluated according to this yet-to-be-clared criteria, thus discouraging investment in ongoing research and development efforts.

It should be said that the high-risk categorization of AI systems will also have a spill-over effect on others. Burri (2022) comments that a low-risk system could qualify as a high-risk one after minor changes are introduced in its development. Moreover, since Annex III of the AI Regulation does not attach technical conditions to the classification but defines the risks based on areas of use, the same technology can be both low-risk and high-risk depending on its application/use. For example, AI for image recognition is a low-risk system, except when used for facial recognition by law enforcement agencies, as then it becomes high-risk (Burri, 2022, pp. 114–115). This ambiguity certainly has a preventive aim but will also have a “deterrence” effect, as it remains unclear on what exact conditions can compliance be ensured when the distance as a matter of legal implications (and obligations) between the two classifications is so substantial.

Furthermore, the EU *AI Act* takes into consideration GPAI, addressing the unique challenges posed by these versatile and powerful systems. Article 51 sets criteria for classifying the GPAI models that entail systemic risks. A model falls into this category if it has high-impact capabilities evaluated on the basis of “appropriate technical tools and methodologies” (Article 51(1)). The Regulation also provides a specific metric for assuming high impact: any model whose training computation exceeds  $10^{25}$  floating-point operations is classified under this category (Article 51(2)).

Manufacturers of such AI systems acquire additional obligations during the development phase. Article 53 consigns the said requirement that GPAI providers maintain comprehensive documentation and ensure transparency, and that they must provide detailed technical documentation on the model’s training, testing and performance evaluation. Furthermore, Article 55 outlines additional responsibilities for providers of GPAI models, identified as bearing systemic risks.

It is in this regard that Ebers (2024) argues that the separate regulation of GPAI tools does not fit into the basic risk categories established in the *AI Act*; it basically creates an overlap. Additionally, it is not known which specific societal challenges the regulator intends to address through this categorization and what social issues it intends to resolve. It can be said that tying the categorization to training computation power does not necessarily convey the social risk of the AI system, as a tool with lower computational training might still pose serious societal and economic risks, whereas computationally intensive models might not present significant ones.

Although not specified as a separate category, open-source AI models are also given special regulatory treatment. The Act articulates specific rules concerning open-source licensed AI systems and states that the regulation does not apply to those released under free and open-source licenses, except when marketed or deployed as high-risk or falling under the scope of Articles 5 or 50 (Article 2(12)). Furthermore, it is established that the obligations defined for GPAI do not apply to providers of AI models that are released under a license allowing free access to the use, modification and distribution of the model and the model’s parameters, including weights, architecture information and usage details that are publicly accessible (apart from those posing systemic risks) (Article 53(2)). This strategy aims to promote the development of open-source models over closed ones by reducing regulatory obligations and creating the perception that open-source models are more transparent and reliable. The rationale is that accessibility and modifiability

automatically increase transparency and accountability, potentially leading to more trustworthy AI applications, which is not necessarily the case.

Efforts to promote and facilitate transparency in these models require substantial commitment, and this responsibility extends beyond public sector stakeholders, particularly regulators who may lack expertise in the field; thus, there is an urgent need for interdisciplinary collaboration. Section 4 touches upon this point and discusses the potential extensions of these rules more comprehensively along with adjustments to close the gap between technological advancements and societal expectations on the ethical development of AI.

### 3. Proactive aspects of the EU AI Act

So far, this work has emphasized the strong preventive and risk-centered orientation of the *AI Act*, suggesting that it is conceived to prohibit and restrict AI systems and activities that are deemed socially dangerous and/or harmful. However, the Act also incorporates proactive tools and features aimed at fostering socially acceptable practices based on or supported by these technologies. These include the establishment of regulatory sandboxes, initiatives to enhance the transparency of AI, rules about data governance and ways to develop practical guidelines, such as codes of practice.

#### 3.1 AI regulatory sandboxes

One of the European Union's most proactive and innovation-promoting regulatory measures in the *AI Act* is the endorsement of regulatory sandboxes for AI development (Allen, 2019). It mandates that Member States set up at least one national AI regulatory sandbox by August 2, 2026, which can be formed individually or in collaboration with other Member States. The European Commission technically and operationally supports these sandboxes, ensuring a consistent level of capability and expertise throughout the Union (Article 57(1)–(2)).

Sandboxes provide a controlled environment for developing, testing and validating AI systems prior to their market entry (Article 57(5)). Competent authorities offer participants guidance, supervision and support to help identify and mitigate risks, particularly those concerning fundamental rights, health and safety (Article 57(6)). The sandboxes aim to smooth the

conformity assessment process, with exit reports and written confirmations from competent authorities being positively considered by market surveillance authorities and notified bodies (Article 58(1)). Their design emphasizes inclusion, providing broad and equal access to participants, encouraging SMEs and especially startups. Access is granted to these smaller entities free of charge, except for exceptional costs that may be recovered proportionally. Furthermore, the sandboxes promote collaboration among stakeholders in the AI ecosystem, including standardization organizations, testing facilities and research laboratories (Article 58(2)).

The *AI Act* allows for testing high-risk AI systems in real-world conditions outside the regulatory sandboxes as well, before officially deploying an AI system in the market, thus combining caution, flexibility and innovation stimulus. Seeking real-world validation of AI systems under this regulatory oversight may indeed effectively increase the chances to handle risks (Article 60).

### 3.2 Regulations on AI transparency

Another proactive aspect of the *AI Act* concerns the regulation of transparency matters related to individual models and AI tools. Transparency is a key requirement of trustworthy AI in the current institutional framework of the EU (High-Level Expert Group on AI, 2019), and is measured by factors that determine the accessibility of stakeholders to a system's operations and its data governance mechanisms (Dignum, 2019, p. 54). It is crucial because, as explained in the literature, it closely relates to accountability (Felzmann *et al.*, 2020, p. 3338), and is another key requirement of trustworthy AI (High-Level Expert Group on AI, 2019, pp. 19–20). In addition, transparency is demonstrated to correlate with trust (Hoff and Bashir, 2015). Accordingly, the EU mandates comprehensive transparency provisions for high-risk AI systems designed to foster trust, accountability and responsible innovation in developing and deploying these systems.

For instance, Article 13 mandates that high-risk AI systems be designed and developed to ensure “sufficient” transparency, to enable the interpretation of the systems' outputs accurately and use them appropriately; without a clear understanding of the AI system's functionalities and limitations, the responsible and effective use of AI technologies would not be possible. The same Article stipulates that high-risk systems must be accompanied by comprehensive instructions for use, provided in proper digital format. Instructions have to be concise, complete, correct and clear, making them

accessible and understandable to deployers (in this case, the users), which integrates not merely protective but also proactive principles explored and promoted by clarity doctrines (Flückiger, 2016), for instance within the legal design and transaction design communities (Haapio, Barton and Corrales Compagnucci, 2021; also, in general, Corrales Compagnucci, Haapio and Fenwick, 2022) and specifically endorsed by public policy ('Opinion of the European Economic and Social Committee 2009/C 175/26', 2009, pp. 27, 28). By detailing aspects such as the provider's identity, the system capabilities with performance metrics and potential risks, the Act seeks to ensure that deployers are well-informed (Article 13).

Article 50 extends transparency obligations to providers and deployers of certain AI systems, particularly those that interact directly with natural persons or generate synthetic content. They are required to disclose to individuals when they interact with an AI system unless it is evident to a reasonably well-informed person (Article 50(1)). This provision aims to help build public trust in AI technologies and facilitate widespread adoption and innovation.

For AI systems that produce synthetic audio, image, video or text content, the regulation mandates that outputs be marked in a machine-readable format to show that they are artificially generated or manipulated (Article 50(2)). This measure addresses concerns about misinformation and the potential misuse of AI-generated content, such as deepfakes. It also mandates that deployers of emotion recognition or biometric categorization systems inform individuals exposed to these systems, ensuring compliance with data protection regulations (Article 50(3)), corresponding to key requirements derived from ethicality and legality as characteristics of trustworthy AI in the EU (Göksal and Solarte-Vásquez, 2024, pp. 4–7).

These transparency rules for AI models are not fully determined but sufficiently forward-looking by signalling that the quality of information and disclosures affects the trustworthiness of systems and the trust in AI. While these rules begin to address a pressing democratic and societal challenge: the increasing difficulty of distinguishing between truth and falsehood, particularly in relation to AI-generated deepfake content, more meditated measures and an extended normative treatment are necessary. The regulatory framework of AI could also include incentives for supporting open-source models, which is discussed in detail in Section 4 below.

### 3.3 AI data governance

Yet another useful regulatory feature in the Act that aligns with a proactive legal perspective emerges from the key requirement of data governance for AI models (High-Level Expert Group on AI, 2019, pp. 19–20). Article 10 speaks of the conditions for developing high-risk AI systems, emphasizing the importance of the quality of datasets and responsible data management practices in the efficacy (robustness) and ethical integrity of these technologies. Article 10 mandates that AI systems involving model training be developed using the training, validation and testing datasets that meet strict criteria, appropriate for the system’s intended purpose (Article 10(1)).

The same Article refers to the data relevance and representativeness. Datasets must be pertinent, sufficiently representative, and as error-free and complete as possible, considering the intended use of the AI system (Article 10(3)). This adds to the accuracy and precision of the AI systems, enhancing their effectiveness and reliability (Article 10(4)). The *AI Act* recognizes that different applications call for specific and contextualized data, so the training has to be expert, precise and forward-looking to account for the factors and variables that may affect the operation of AI systems.

The Act allows regulators to actively influence and define frameworks for the use of data to develop AI systems. The selection of datasets for training is critically important for the functioning of AI and carries normative implications, as these models can impact the socioeconomic environment. This presents a tremendous opportunity for stakeholders to implement the proactive law approach. Martínez-Ramil *et al.* (2023) refer to the potentially discriminatory nature of AI, which is one of the most frequently cited challenges associated with these systems. They argue that Article 10 permits a more proactive shaping of technology if incorporating specific obligations “by design” (Martínez-Ramil *et al.*, 2023, p. 242).

### 3.4 Practical guidelines and codes of practice

Similarly, an effective means for proactively shaping technology’s positive impacts is the creation of practical guidelines. These resources enable legislators, preferably in collaboration with other relevant stakeholders, to model interactions and behaviors by establishing concrete instructions or supported modes for the development and use of AI, restricting directions considered socially risky.

The *AI Act* includes several provisions, accordingly, reiterating the importance of risk management systems (Article 9), quality management systems (Article 17) and fundamental rights impact assessments (Article 27). The regulation refers anew to the governance capacities regarding codes of practice. Article 56 mandates the AI Office to promote and assist in creating the codes. The Act supports multistakeholderism, asking for a broad spectrum of input by inviting all providers of GPAI models and relevant national authorities to participate. Additionally, the inclusion of civil society organizations, industry representatives, academia, downstream providers and independent experts is seen to help build a collaborative environment (Article 56(3)). The AI Office and the AI Board are to verify that the codes of practice have well-defined objectives and include commitments and measures to achieve them. Key performance indicators (KPIs) should monitor progress where appropriate. In addition, the participants producing the codes are expected to report regularly on their implementation efforts and outcomes, measured against these KPIs (Article 56(4–6)).

The codes of practice have not yet been developed and thus the exact texts and their actual impact remain to be seen. However, they represent a powerful tool within the EU's institutional framework, particularly in the hands of a “new type of regulator” composed of collaborative cross-sectoral stakeholder groups. Such a tool enables the creation of alternative and more dynamic rules that can adapt to unforeseen technological changes.

#### 4. Recommendations for enhancing the proactive approach in AI regulation

Sections 2 and 3 explained, on the one hand, the ways in which the EU *AI Act* adopts a risk-based approach, strongly supporting preventive, *ex-ante* regulatory measures, among these the identification and formulation of prohibited AI practices and the risk classification framework. On the other hand, they showed how the Act also incorporates proactive regulatory measures, such as regulatory sandboxes, transparency requirements, data governance obligations and the promotion of less traditional forms of regulatory activities, such as guidelines and codes of practice. Against that background, this section will outline the regulatory elements that could be strengthened or introduced to create an environment that encourages innovation, while more systematically emphasizing the societal benefits of AI.

## 4.1 Supporting AI transparency and open-source models

A persistent challenge in the governance of AI systems, and one that will extend to future technologies, is the lack of transparency in their components, operations and training processes. For many publicly available and widely used models, the specifics of their training data and processes are still unknown. This opacity undermines trust in the system's safety, ethicality, legality and capabilities and presents broader societal risks. Without clear knowledge of potential biases or deficiencies in data governance during development, assessing the social impact of AI systems becomes more difficult.

Open-source models directly address transparency concerns by offering accessibility and modifiability. Users can download, implement and adjust these models locally, provided that sufficient computational resources are available. Public scrutiny of their parameters and performance also makes validation and accountability possible. Furthermore, detailed documentation of training data often accompanies these models, enabling a deeper understanding of their biases, limitations and development processes.

The EU *AI Act* already encourages the adoption of more transparent and accessible models and requires technical documentation for high-risk AI systems. However, additional regulatory and incentive-based measures are necessary to encourage transparency in AI development. Promoting these models strengthens governance and could enhance the competitiveness of European AI firms and industries in other sectors.

## 4.2 Complementing the holistic approach with sector-based regulations

A potential direction for the future development of the EU governance framework for AI is the formulation of sector-based rules designed to address specific needs and standards. The *AI Act* adopts a holistic and relatively flexible regulatory approach, which results in abstractions that may require intensive interpretative efforts. Its broad definition of AI does not adequately reflect the evolving concepts, ideas and technologies within the field, much less the prospect of its integration with other emerging and future technologies. Over the decades, AI has become an umbrella term that, in fact, many researchers prefer to avoid (Schuett, 2023).

Characterizations of AI are now made by attending to common features among these systems (Häuselmann, 2022), a task in which the *AI Act* has performed

well. It outlined general rules based on the system's machine-learning capabilities, degrees of autonomy and ability to adapt to environmental changes. However, in order to move toward more useful specifications, it is essential to consider the expanding range of AI applications and the unique conditions of the fields and contexts (such as regions, legal and political systems, jurisdictions, industries, special regimes and cultures) in which these tools will be deployed. To illustrate the relevance of these factors, the *AI Act* explicitly excludes its application to AI systems that are deployed or used exclusively for military, defense or national security purposes (Article 2(3)). While this provision was established in recognition of the unique nature of those activities, if those AI technologies are also utilized for law enforcement or public security, they must still comply with the provisions of the Act (Recital 24).

Implementing rules focused on large language models (LLMs) and associated tools is especially advisable, as these models have rapidly integrated into daily life, influencing various sectors within just a few months of their introduction, and they are expected to have major economic and societal implications. With LLMs becoming more prevalent, concerns grow about their expanding applications in sensitive areas such as healthcare, law and education. The rapid penetration of LLMs evidences the benefits of good regulatory planning for continuous regulatory adaptation. As noted in other discussions around regulatory oversight, frameworks should not only focus on current models but anticipate future iterations and applications. This proactive approach will help regulations remain relevant and effective in managing the complexities associated with the development, deployment, implementation and use of technology. In sum, preventive and proactive normative elements can operate best when reflected within sector-specific regulations.

### 4.3 Strategizing for environmentally conscious AI practices

The complex dilemma posed by sustainability goals and the race for automation must be acknowledged within the broader discourse on responsible AI development. Addressing these interconnected challenges proactively is imperative to ensure that AI technologies contribute to humanity's sustenance and society's well-being, while minimizing their environmental footprint. A concern that the EU AI Regulation has largely neglected is the immense energy consumption required to meet the computational demands of AI systems, along with the subsequent environmental impact of automation.

One of the aspects of this impact is that AI tools demand more resources than other forms of technology, particularly in the case of generative AI, where both model development and usage are energy intensive. As these systems become prevalent, complying with the existing mandate on sustainable data management alone will not suffice.

The EU could develop a strategy that specifies the applications in which the use of large, high-efficiency machine-learning models is justified and in which others would suffice. Such considerations would encourage the adoption of more sustainable and environmentally conscious AI practices (Smith, 2024). Moreover, by delineating appropriate use cases, the EU can promote energy efficiency in AI development and usage, balancing technological advancement with environmental sustainability. The *AI Act* already includes requirements for risk management and compliance monitoring to promote responsible innovation but could not explicitly address the environmental impact of AI or propose sustainable practices in detail. Perhaps transparency requirements and risk management will encourage developers to disclose and mitigate the energy and resource consumption of their AI systems, but there needs to be a more resolute and purposeful normative treatment of the matter, consistent with broader EU sustainability commitments.

Consequently, to guide the environmentally conscious AI development and usage, additional regulatory adjustments are necessary within the realm of AI and technology, and in sustainability-related areas concerning more realistic assumptions, principles and expectations. It is critical for society to regulate proactively and plan systematically for its own sustainable future (Andrade and Solarte-Vásquez, 2024; König, Wurster and Siewert, 2022). This is particularly true for the EU, if proper alignment of its AI objectives, AI regulations and environmental commitments made under the European Green Deal ('Communication COM/2019/640 final', 2019) and the European Climate Law ('Council regulation (EU) 2021/1119', 2021; Andrade and Solarte-Vásquez, 2024, p. 137) is achieved.

Finally, there is an assumption that automation will evolve linearly without challenge. However, some sceptics point out important factors to consider, such as society's capacity to bear the costs of increasing automation and the implications of singularity (Kurzweil, 2024). Overlooking these scenarios weakens the relevance of the *AI Act*. While this paper highlights these issues, it also recognizes its limitations in fully addressing the arguments in depth. Future research will be essential to examine those views and anticipate their implications.

## 5. Conclusion

The widespread discourse on the societal dangers and risks associated with AI truly gained momentum with the emergence and proliferation of ChatGPT, along with concurrent regulatory possibilities. Despite this, a unified global understanding about AI and doctrine on the governance possibilities of this technology is lacking. What is more, the existing regulatory attempts face serious criticism.

One reason is that complexities keep growing on all matters of AI, but regulatory responses should not be delayed. To name a few, AI is an overly broad umbrella term encompassing various new technologies and tools, each operating differently, based on separate foundations and eliciting different societal impacts, which require distinct regulatory logic; some AI systems possess general-purpose potential, leading to varying impacts across multiple applications and complicating their legal treatment; in addition, AI can adapt to changes in its environment and alter its operations, making it challenging to establish fixed standards that apply consistently in all situations. As a result, given these unique characteristics and unprecedented technological capacities, the traditional reactive approach to regulation and the preventive approach prove inadequate.

Early research and policy have endorsed proactiveness in general, and the proactive law and legal practice, pioneered by the Nordic school, who view the role of regulations in society not only as protective, defensive or shielding from harm, but as a strong means for value creation, and more sustainable economic and social relationships. These share purpose with the concept of 'better lawmaking', which emphasizes anticipation, self-initiation and the promotion of positive change. 'Better lawmaking' and the notion of 'good laws' are about dynamic rules that prioritize responsive, accessible and essential objectives, and that consider the real-life needs of the stakeholders and their interactions. Both aim at formulating and advancing regulation that is operationally efficient and meaningful. These perspectives assume that because the successful implementation of rules extends beyond adoption, it requires acceptance and trust, so inclusiveness and participation in regulatory processes can influence positively the effectiveness of the rules. This paper argued that regulation of technology could fundamentally benefit from adopting these principles, explored their applicability to the enhancement of AI regulation and suggested that by embracing proactive regulation, we could better navigate the complexities

of AI technologies while promoting sustainable development and societal well-being.

Whereas the most important normative document in the EU on AI was shown to emphasize preventive, *ex-ante* regulatory tools, it was found to incorporate proactive measures such as regulatory sandboxes, transparency requirements, data governance obligations and practical codes of practice. This work discussed new pro-innovation regulatory elements that should be designed to enhance the societal benefits of AI, including the development of transparent and open-source models, a rational sector-based regulatory framework supplementing the *AI Act*, and the integration and operationalization of environmental sustainability parameters related to the development and use of AI and other emerging technologies. Finally, while the work acknowledged some deficiencies regarding the relationship between automation and sustainability, it recognized limitations in addressing these arguments to remain within its scope, particularly given the insufficient availability of scholarly literature and other reliable data to support claims in this regard.

**Renátó Vági** is a PhD student at Eötvös Loránd University, Hungary. His research interests include the legal applications of AI tools and the impact of AI on the legal field and society. He has managed and led several software development projects aimed at providing AI-based solutions for lawyers, enhancing their work processes. At Wolters Kluwer Hungary, he performs as a legal engineer and legal AI expert, and his main goal is to work in AI-powered LegalTech innovation to make legal research and lawyers' work more efficient.

## References

- Acemoglu, D. (2021) 'Remaking the post-COVID world', *Finance & Development*. International Monetary Fund. Available at: <https://www.imf.org/external/pubs/ft/fandd/2021/03/COVID-inequality-and-automation-acemoglu.htm> (Accessed: 13 November 2024).
- Algorithmic Accountability Act of 2019* (2019) [Online]. H.R.2231, 116th Congress (2019–2020). Available at: <https://www.congress.gov/bill/116th-congress/house-bill/2231/all-info> (Accessed: 13 November 2024).
- Allen, H.J. (2019) 'Regulatory sandboxes', *The George Washington Law Review*, 87(3), pp. 579–600. Available at: <https://www.gwlr.org/wp-content/uploads/2019/06/87-Geo.-Wash.-L.-Rev.-579.pdf> (Accessed: 13 November 2024).
- Andrade, A.C.C. and Solarte-Vásquez, M.C. (2024) 'The compatibility between SDGs and the EU regulatory framework of AI', *Journal of Ethics and Legal Technologies*, 6(1), pp. 11–144.
- Antonov, A. (2022) 'Managing complexity: the EU's contribution to artificial intelligence governance', *Revista CIDOB d'Afers Internacionals*, (131), pp. 41–65. Available at: <https://doi.org/10.24241/rcai.2022.131.2.41/en>
- Berger-Walliser, G. (2012) 'The past and future of proactive law: an overview of the development of the proactive law movement', in G. Berger-Walliser and K. Østergaard (eds) *Proactive law in a business environment*. DJØF Publishing, pp. 13–31.
- Brown, L.M. (1956) 'The law office—a preventive law laboratory', *University of Pennsylvania Law Review*, 104(7), pp. 940–953. Available at: <https://doi.org/10.2307/3310432>
- Burri, T. (2022) 'The new regulation of the European Union on artificial intelligence: fuzzy ethics diffuse into domestic law and sideline international law', in S. Voeneky, P. Kellmeyer, O. Mueller and W. Burgard (eds) *The Cambridge handbook of responsible artificial intelligence: interdisciplinary perspectives*. Cambridge: Cambridge University Press, pp. 104–122. Available at: <https://doi.org/10.1017/9781009207898.010>
- Carnegie Endowment for International Peace (2024) 'AI governance research', 7 October. Available at: <https://www.openphilanthropy.org/grants/carnegie-endowment-for-international-peace-ai-governance-research/> (Accessed: 16 November 2024).
- Chun, J., de Witt, C.S., and Elkins, K. (2024) 'Comparative global AI regulation: policy perspectives from the EU, China, and the US', *ArXiv*, article 2410.21279. Available at: <https://doi.org/10.48550/arXiv.2410.21279>
- 'Communication from the Commission to the European Parliament, the European Council, the Council, the European Economic and Social Committee and the

- Committee of the Regions: The European Green Deal', (COM/2019/640 final)' (2019) 11 December. Available at: <https://eur-lex.europa.eu/legal-content/EN/AUTO/?uri=CELEX%3A52019DC0640> (Accessed: 13 November 2024).
- Corrales Compagnucci, M., Haapio, H., and Fenwick, M. (eds) (2022) *Research handbook on contract design*. Edward Elgar Publishing. Available at: <https://doi.org/10.4337/9781839102288>
- Dignum, V. (2019) *Responsible artificial intelligence: how to develop and use AI in a responsible way*. Springer International Publishing. Available at: <https://doi.org/10.1007/978-3-030-30371-6>
- Ebers, M. (2024) 'Truly risk-based regulation of artificial intelligence: how to implement the EU's AI Act', *European Journal of Risk Regulation*, pp. 1–20. Available at: <https://doi.org/10.1017/err.2024.78> (Accessed: 13 November 2024).
- Ebers, M., Hoch, V.R.S., Rosenkranz, F., Ruschemeier, H., and Steinrötter, B. (2021) 'The European Commission's proposal for an Artificial Intelligence Act—a critical assessment by members of the Robotics and AI Law Society (RAILS)', *J*, 4(4), pp. 589–603. Available at: <https://doi.org/10.3390/j4040043>
- 'Executive order on maintaining American leadership in artificial intelligence' (2019) The White House, 11 February. Available at: <https://trumpwhitehouse.archives.gov/presidential-actions/executive-order-maintaining-american-leadership-artificial-intelligence/> (Accessed: 13 November 2024).
- 'Executive order on the safe, secure, and trustworthy development and use of artificial intelligence' (2023) The White House, 30 October. Available at: <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/> (Accessed: 13 November 2024).
- Felzmann, H., Fosch-Villaronga, E., Lutz, C. and Tamò-Larrieux, A. (2020) 'Towards transparency by design for artificial intelligence', *Science and Engineering Ethics*, 26(6), pp. 3333–3361. Available at: <https://doi.org/10.1007/s11948-020-00276-4>
- Flückiger, A. (2016) 'The ambiguous principle of the clarity of law', in A. Wagner and S. Cacciaguidi-Fahy (eds) *Obscurity and clarity in the law*. 2nd edn. Routledge, pp. 31–46.
- Fried, B.H. (2003) 'Ex ante/ex post', *Journal of Law and Contemporary Issues*, 13, pp. 123–160. Available at: <https://doi.org/10.2139/ssrn.390462>
- Göksal, Ş.İ. and Solarte-Vásquez, M.C. (2024) 'The blockchain-based trustworthy artificial intelligence supported by stakeholders-in-the-loop model', *Scientific Papers of the University of Pardubice, Series D: Faculty of Economics and Administration*, 32(2), article 2083. Available at: <https://doi.org/10.46585/sp32022083>
- Haapio, H., Barton, T.D. and Compagnucci, M.C. (2021) 'Legal design for the common good: proactive legal care by design', in M. Corrales Compagnucci, H. Haapio,

- M. Hagan and M. Doherty (eds) *Legal design: integrating business, design and legal thinking with technology*. Cheltenham: Edward Elgar Publishing, pp. 56–81. Available at: <https://doi.org/10.4337/9781839107269.00011>
- Haapio, H. and Varjonen, A. (1998) 'Quality improvement through proactive contracting: contracts are too important to be left to lawyers!', *ASQ World Conference on Quality and Improvement Proceedings*, 52, pp. 243–248.
- Häuselmann, A. (2022) 'Disciplines of AI: an overview of approaches and techniques', in B. Custers and E. Fosch-Villaronga (eds) *Law and artificial intelligence: regulating AI and applying AI in legal practice*. The Hague: T.M.C Asser Press, pp. 43–70. Available at: [https://doi.org/10.1007/978-94-6265-523-2\\_3](https://doi.org/10.1007/978-94-6265-523-2_3)
- High-Level Expert Group on AI (2019) *Ethics guidelines for trustworthy AI*. Available at: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai> (Accessed: 13 November 2024).
- Hoff, K.A. and Bashir, M. (2015) 'Trust in automation: integrating empirical evidence on factors that influence trust', *Human Factors*, 57(3), pp. 407–434. Available at: <https://doi.org/10.1177/0018720814547570>
- König, P., Wurster, S. and Siewert, M.B. (2022) 'Regulating for sustainable artificial intelligence: citizens' support for policy instruments', *SSRN*. Available at: <https://doi.org/10.2139/ssrn.4082262> (Accessed: 13 November 2024).
- Kurzweil, R. (2024) *The singularity is nearer: when we merge with AI*. Random House.
- Lesnes, C. (2024) 'California Governor Gavin Newsom vetoes AI safety bill', *Le Monde*, 30 September. Available at: [https://www.lemonde.fr/en/economy/article/2024/09/30/california-governor-gavin-newsom-vetoes-ai-safety-bill\\_6727751\\_19.html](https://www.lemonde.fr/en/economy/article/2024/09/30/california-governor-gavin-newsom-vetoes-ai-safety-bill_6727751_19.html) (Accessed: 13 November 2024).
- Madiega, T. (2024) 'Artificial Intelligence Act', European Parliament Briefing, European Parliament Research Service, PE 698.792. Available at: [https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS\\_BRI\(2021\)698792\\_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI(2021)698792_EN.pdf) (Accessed: 13 November 2024).
- Magnusson Sjöberg, C. (2006) 'Introduction', in P. Wahlgren (ed) *A proactive approach. Scandinavian studies in law*, 49. Stockholm: Stockholm Institute for Scandinavian Law, pp. 1–7. Available at: <https://www.scandinavianlaw.se/pdf/49-1.pdf> (Accessed: 15 October 2024).
- Martínez-Ramil, P., Bolaños-Frasquet, H., Hamulák, O. and Kerikmäe, T. (2023) 'A multidimensional understanding of EU's digital sovereignty', in D. Ramiro Troitiño, O. Hamulák and T. Kerikmäe (eds) *Digital development of the European Union: an interdisciplinary perspective*. Cham: Springer, pp. 235–247. Available at: [https://doi.org/10.1007/978-3-031-27312-4\\_15](https://doi.org/10.1007/978-3-031-27312-4_15)
- Mazzucato, M., Schaake, M., Krier, S. and Entsminger, J. (2022) 'Governing artificial intelligence in the public interest', *UCL Institute for Innovation and Public Purpose Working Paper Series*, IIPP WP 2022-12. Available at: <https://www.ucl.ac.uk/bartlett/public-purpose/wp2022-12> (Accessed: 13 November 2024).

- Nguyen, T. (2024) 'California governor signs bills to protect children from AI deepfake nudes', *AP News*, 30 September. Available at: <https://apnews.com/article/ai-deepfakes-children-abuse-7dcf5c566e2a297567f1e148ac2074a4> (Accessed: 13 November 2024).
- 'Opinion of the European Economic and Social Committee on 'The proactive law approach: a further step towards better regulation at EU level' (2009), *Official Journal C* 175/26, 28 July. Available at: <https://eur-lex.europa.eu/LexUriServ/LexUriServ.do?uri=OJ:C:2009:175:0026:0033:EN:PDF> (Accessed: 13 November 2024). Available at: <https://eur-lex.europa.eu/LexUriServ/LexUriServ.do?uri=OJ:C:2009:175:0026:0033:EN:PDF> (Accessed: 13 November 2024).
- Pagallo, U., Ciani Sciolla, J. and Durante, M. (2022) 'The environmental challenges of AI in EU law: lessons learned from the Artificial Intelligence Act (AIA) with its drawbacks', *Transforming Government: People, Process and Policy*, 16(3), pp. 359–376. Available at: <https://doi.org/10.1108/TG-07-2021-0121> (Accessed: 13 November 2024).
- Pohjonen, S. (2006) 'Proactive law in the field of law: a proactive approach', in P. Wahlgren (ed) *A proactive approach. Scandinavian studies in law*, 49. Stockholm: Stockholm Institute for Scandinavian Law, pp. 53–70. Available at: <https://www.scandinavianlaw.se/pdf/49-4.pdf> (Accessed: 15 December 2024).
- Prifti, K., Quintavalla, A. and Temperman, J. (2023) 'Artificial intelligence and human rights: understanding and governing common risks and benefits', in A. Quintavalla and J. Temperman (eds) *Artificial intelligence and human rights*. Oxford University Press.
- Raposo, V.L. (2022) 'Ex machina: preliminary critical assessment of the European Draft Act on artificial intelligence', *International Journal of Law and Information Technology*, 30(1), pp. 88–109. Available at: <https://doi.org/10.1093/ijlit/eaac007>
- 'Regulation (EU) 2021/1119 of the European Parliament and of the Council of 7 June 2021 on establishing the Climate Law (CELEX No. 32021R1119) (2021) *EUR-Lex*. Available at: <https://eur-lex.europa.eu/legal-content/EN/AUTO/?uri=celex:32021R1119> (Accessed: 13 November 2024).
- 'Regulation (EU) 2024/1689 on laying down harmonised rules on artificial intelligence and amending Regulations (EC) no. 300/2008, (EU) no. 167/2013, (EU) no. 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828' (2024) *Official Journal L* 2024/1689, 12 July. Available at: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689> (Accessed: 13 November 2024).
- Safe and Secure Innovation for Frontier Artificial Intelligence Models Act* (2023) Senate Bill 1047, 7 February. Available at: [https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill\\_id=202320240SB1047#:~:text=This%20bill%20would%20enact%20the,to%20promptly%20enact%20a%20full](https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=202320240SB1047#:~:text=This%20bill%20would%20enact%20the,to%20promptly%20enact%20a%20full) (Accessed: 13 November 2024).

- Schartum, D.W. (2006) 'Introduction to a government-based perspective on proactive law: a proactive approach', in P. Wahlgren (ed) *A proactive approach. Scandinavian studies in law*, 49. Stockholm: Stockholm Institute for Scandinavian Law, pp. 35–51. Available at: <https://www.scandinavianlaw.se/pdf/49-3.pdf> (Accessed: 15 December 2024).
- Schuett, J. (2023) 'Defining the scope of AI regulations', *Law, Innovation and Technology*, 15(1), pp. 60–82. Available at: <https://doi.org/10.1080/17579961.2023.2184135>
- Smith, N. (2024) 'Plentiful, high-paying jobs in the age of AI', *Noahpinion*, 17 March. Available at: <https://www.noahpinion.blog/p/plentiful-high-paying-jobs-in-the> (Accessed: 13 November 2024).
- Solarte-Vásquez, M.C. and Hietanen-Kunwald, P. (2020) 'Transaction design standards for the operationalisation of fairness and empowerment in proactive contracting', *International and Comparative Law Review*, 20(1), pp. 180–200. Available at: <https://doi.org/10.2478/iclr-2020-0008>
- Tyrangiel, J. (2024) 'I found the smartest politician on AI. It's no one you'd expect', *The Washington Post*, 20 March. Available at: <https://www.washingtonpost.com/opinions/2024/03/20/ai-europe-regulation-leading/> (Accessed: 26 November 2024).
- Wang, X. and Wu, Y.C. (2024) 'Balancing innovation and regulation in the age of generative artificial intelligence', *Journal of Information Policy*, 14. Available at: <https://doi.org/10.5325/jinfopoli.14.2024.0012>
- White, J.B. (2024) 'Gavin Newsom signs election “deepfake” ban in rebuke to Elon Musk', *Politico*, 17 September. Available at: <https://www.politico.com/news/2024/09/17/newsom-signs-election-deepfake-ban-00179557> (Accessed: 13 November 2024).