



Citation: Maggiora, E., & Iacobino, C. (2025) EU Artificial Intelligence Act: risks and opportunities for the European system. *Journal of Emerging Perspectives 2*: 97-101. doi: 10.36253/jep-20886

Received: April 30, 2025

Revised: September 30, 2025

Published: December 20, 2025

© 2025 Author(s). This is an open access, peer-reviewed article published by Firenze University Press (<https://www.fupress.com>) and distributed, except where otherwise noted, under the terms of the CC BY 4.0 License for content and CC0 1.0 Universal for metadata.

Data Availability Statement: All relevant data are within the paper and its Supporting Information files.

Competing Interests: The Author(s) declare(s) no conflict of interest.

EU Artificial Intelligence Act: risks and opportunities for the European system

ENRICO MAGGIORA^{1*}, CLAUDIA IACOBINO²

¹ avvocato, presidente Fondazione dell'Avvocatura Torinese Fulvio Croce, Italy

² avvocatessa, presidente AGAT - Associazione Giovani Avvocati Torino, Italy

e.maggiora@consulenza-impresa.it

*Corresponding author

Abstract. This paper examines the European Union Artificial Intelligence Act (AI Act) as a strategic regulatory response to the rapid and pervasive diffusion of artificial intelligence technologies. Starting from a conceptual framing of AI as an instrument of augmented intelligence rooted in bounded rationality, the contribution highlights how contemporary AI systems adopt satisficing logics through heuristics and large-scale data processing rather than pursuing optimal solutions. The analysis situates the AI Act within the broader EU digital strategy and discusses its risk-based regulatory architecture, which classifies AI systems according to the severity and likelihood of potential harm to fundamental rights, safety, and democratic values. Particular attention is devoted to high-risk AI systems, whose stringent compliance obligations raise significant legal, technical, and economic challenges, especially for SMEs and startups. While acknowledging the risks of regulatory rigidity, innovation slowdown, and market entry barriers, the paper also emphasizes the strategic opportunities generated by the AI Act. These include the consolidation of a trustworthy, human-centric AI ecosystem, the strengthening of the Digital Single Market, and the potential emergence of a global regulatory benchmark through the so-called “Brussels effect.” Ultimately, the AI Act is interpreted not merely as a compliance burden, but as a long-term investment capable of transforming regulation into a competitive advantage for the European system.

Keywords: Artificial Intelligence Act, EU Digital Regulation, risk-based approach, high-risk AI systems, Fundamental Rights Protection, SMEs and Innovation, trustworthy AI, Brussels Effect.

1. THE AI REVOLUTION AND THE NEED FOR REGULATION.

Great revolutions are, for the most part, silent. By the time we realize that an epoch-making change is underway, that change has already occurred, and often, few have seen it coming. In the collective imagination, great revolutions find their seeds of development in a sensational, often violent act that immediately attracts public and media attention, masking and silencing the transformations that take place quietly and which, paradoxically, are the ones destined to make the most noise. The human mind struggles to project and

understand the long-term effects of slow but inexorable transformations and to perceive the real extent of change in the short term.

This is particularly true in the case of technological revolutions, which are never sudden but are the result of a process of gradual development and implementation, the consequences of which, however, are often disruptive.

The advent of Artificial Intelligence (AI) is undoubtedly one of the most discussed and significant IT revolutions in contemporary history. A new language of writing so sophisticated that it can simulate human intellectual processes. It is difficult to give a single, rigorous definition of 'Artificial Intelligence', and it therefore seems appropriate to refer to some authoritative elaborations. Artificial Intelligence is, therefore, '*the study of how to make computers do things that, at present, people are better at*'¹ and, furthermore, '*computational intelligence is the study of the design of intelligent agents*'².

However, at the definitional level, even more significant is the concept of '*bounded rationality*' developed by Nobel Prize winner Herbert Simon, which has had a crucial impact on the emergence and design of AI-based computing systems, forming their theoretical basis. According to Simon, human decision-makers are endowed with limited rationality because they are subject to cognitive constraints, time constraints, and information limitations, as it is impossible to access all relevant information or process it completely. Therefore, decisions that, at least on the surface, appear to lack reasoning and do not guarantee an optimal result are, on the contrary, perfectly suitable (i.e., rational) insofar as human beings, thanks to experience, are aware of their own limitations and the context of reference. Our intellect inhibits the expenditure of resources and does not seek optimal solutions, but satisfactory ones. Human thinking is certainly fallible, but economical in that it aims to avoid wasting time and energy and to achieve a functional result.

Even the latest generation of Artificial Intelligence adopts a *satisficing* rather than maximizing methodology. This approach involves the pervasive use of heuristics and algorithmic cognitive shortcuts, aimed at identifying a pragmatic optimal solution (*good enough solution*) within predefined and stringent time limits, rather than searching for the *perfect solution*, which would be impractical or asymptotic in terms of cost and time.

AI, far from being a mere emulation of human intelligence, thus rises to the rank of a cognitive- amplifying

tool aimed at overcoming the intrinsic cognitive limitations of the decision-maker. AI algorithms exhibit a capacity for massive data processing (Big Data Analysis) that exceeds the threshold of human processing, thus increasing the 'information blade' of Simon's scissors (i.e., the ability to acquire and process relevant information). The primary objective of this integration is not to replace human rationality, but rather to optimize and augment it (*augmented intelligence*).

Artificial Intelligence, by virtue of its computational properties and algorithmic nature, exerts a highly pervasive, transformative influence, reshaping the operational structures of every strategic sector of the economic and social fabric: from industrial automation to healthcare, from finance to public administration. Increasingly sophisticated Artificial Intelligence systems, in particular General-Purpose AI Systems (GPAI) and generative AI, have demonstrated extraordinary innovative potential. This rapid technological evolution also raises crucial ethical and legal questions relating to transparency, security, fairness, and the protection of fundamental rights.

The need for supranational regulation of Artificial Intelligence at the continental level has become a legal and strategic imperative to ensure that the penetration and adoption of this technology takes place in accordance with the principles of ethics, safety, and transparency. This need is motivated by the (real) risk of compromising the fundamental values of the European Union, arising from specific systemic dangers such as algorithmic discrimination (linked to *biases* inherent in *datasets*), *mass surveillance*, behavioral manipulation of individuals, and the opacity and poor traceability of synthetic content (e.g. *deepfakes*).

In addition to these critical issues, there is also a need to prevent the fragmentation of the Digital Single Market caused by the emergence of heterogeneous and divergent national regulations. In response to this complex scenario, and in line with the '*A Europe Fit for the Digital Age*' policy strategy, the European Union has developed and adopted the Artificial Intelligence Act³ (abbreviated to 'AI Act'). This regulatory document establishes a harmonized legal framework at EU level, adopting a *risk-based approach*. This methodology provides for the taxonomic classification of AI systems according to the probability and severity of the potential harm they may cause (divided into categories ranging from 'unacceptable risk' to 'minimal risk').

¹ E. Rich, K. Knight, *Artificial Intelligence*, McGraw-Hill, 1991, "The study of how to make computers do things that, at the moment, people are better at".

² D.L. Poole, A.K. Mackworth, R. Goebel, *Computational intelligence: A logical approach*, Oxford University Press, 1998.

³ Full title: "Regulation of the European Parliament and of the Council 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".

The AI Act is not exclusively an *ex-ante* precautionary measure, but rather an ambitious attempt to strike a balance between the demand for technological innovation and the duty to protect fundamental rights. The implicit objective of the Regulation is to project a *de facto* global standard for anthropocentric and trustworthy AI governance, achieving the so-called ‘Brussels effect’⁴ at an international level.

2. THE EU AI ACT

A general overview

The AI Act is one of the pillars on which the European Union’s digital strategy is based, which aims to achieve a virtuous balance between supporting innovation and competitiveness in Member States and protecting the fundamental rights of citizens, workers, and consumers. The process of approving the legislation began on 21/04/2021 with the presentation of a proposed Regulation by the European Commission. Since its inception, the regulatory framework has expressed the intention to establish a harmonized and proportionate regulatory framework for the regulation of Artificial Intelligence within the Union.

This objective has led to the adoption of a risk-based approach, structuring the proposal around a taxonomic classification of AI systems according to the level of potential harm (i.e., risk) they may pose to the safety, health and fundamental rights of individuals. This classification forms the regulatory basis from which the legal definitions, compliance requirements (*ex ante* and *ex post* obligations) and specific burdens imposed on the economic operators involved (suppliers, importers, distributors and users) in relation to the place on the market and use of such systems derive directly.

On 14/06/2023, the European Parliament formalized its negotiating position on the document and, subsequently, on 9/12/2023, the Trilogue composed of the European Parliament, the Council and the European Commission negotiated and signed a provisional agreement on the content of the text of the regulation. At a technical meeting on 21 January 2024, the European Commission made the final version of the text of the regulation available. From that point onwards, the approval process was fairly rapid: in March, the Strasbourg Assembly confirmed the final text by a large majority, and in May, the EU Council definitively adopted the AI ACT in its current version.

The European Regulation on Artificial Intelligence pursues the stated dual purpose of improving the functioning of the internal market and stimulating the spread of human-centric and reliable Artificial Intelligence (AI) systems. Complementarily, it establishes a high standard of protection for health, safety and fundamental rights protected by the Charter of Fundamental Rights of the European Union (including democracy, the rule of law and environmental protection) from the harmful effects of the use of AI systems, while also promoting innovative development. Furthermore, the fundamental objectives pursued by the AI Act include: the establishment of a regulatory framework complementary to the General Data Protection Regulation (GDPR) in relation to specific aspects of personal data protection connected to AI, the regulated management of risks arising from the use of Artificial Intelligence systems, and the promotion of technological development that is both safe and controlled.

The subjective scope of the Regulation essentially concerns two players, in particular: ‘Suppliers’, which includes any entity, whether public or private, that makes the AI system available on the European Union market, and ‘Operators’, which refers to all entities responsible for the use of the AI system in Europe and the resulting output, regardless of their geographical location. In addition to these figures, there are importers, distributors and authorized representatives, who participate in various ways in the circulation or use of AI platforms. Platforms intended for military, defense, national security purposes and those dedicated exclusively to scientific research and development are explicitly excluded.

Compliance with the obligations under the AI Act for each of the above-mentioned entities is modulated differently according to the specific use of the technology and, crucially, according to the level of intrinsic risk associated with the AI systems adopted by the organization.

AI practices posing an unacceptable risk constitute the fine line between whether an AI system can be considered lawful or not. AI systems that contravene the fundamental values of the Union and that may pose a significant threat to the security and freedom of individuals, through techniques designed to manipulate the consent or decision-making of citizens, manipulating or exploiting people’s vulnerabilities or creating unfair discrimination, are considered to be of ‘unacceptable risk’ and are therefore prohibited.

Chapter III, Section 1 is entirely devoted to AI systems classified as ‘high risk’, for which there is no ban, but very stringent technical and procedural obligations are imposed. An AI practice can be considered ‘high

⁴ Anu Bradford, *The Brussels Effect, How European Union rules the World*, Oxford Press University

risk' when its use can significantly affect people's health, safety or fundamental freedoms. Suppliers of high-risk AI systems must comply with stringent and ongoing obligations and

ensure an adequate level of *compliance*. The placement on the market of such systems is always subject to a mandatory conformity assessment.

Finally, we come to AI practices considered 'low risk', which are subject to much milder obligations, aimed primarily at ensuring the transparency of AI use processes and self-regulation. The risk associated with these systems is not structural but is limited to cases where the user is not adequately informed about their interaction with the machine. As mentioned, the AI Act does not impose any particular legal obligations on such systems but encourages adherence to voluntary codes of conduct for the promotion of ethical and safe practices.

3. RISKS AND OPPORTUNITIES FOR THE EUROPEAN SYSTEM.

The introduction of such a detailed regulatory framework for artificial intelligence is an important strategic choice for the European Union, although it exposes the continental system to a number of intrinsic technical, legal and economic risks that EU institutions and operators in the sector are inevitably called upon to address.

In the author's opinion, one of the most challenging aspects of the Regulation lies in the correct legal classification of AI practices, with particular reference to 'high-risk' systems. This classification presents the greatest degree of uncertainty in terms of definition, even though it is the prerequisite for the activation of the highest *set of compliance* obligations. It is also essential to consider the exponential speed at which AI is evolving. The rigidity of the list contained in Annex III, which defines 'high-risk' systems, risks becoming obsolete with the emergence of new applications (e.g., the interaction between basic models and specific *downstream* systems), making it necessary for the European Commission to constantly and promptly review and update it.

Furthermore, the compliance burdens imposed on suppliers of 'high-risk' systems are extremely detailed and technically complex, involving the entire production chain.

The obligation to guarantee data quality and *governance* requires burdensome procedures to ensure the representativeness, absence of *bias* and traceability (*data provenance*) of training, validation and *testing* data. In addition to the already complex framework described above,

there is also the obligation to maintain detailed automatic records as referred to in Article 12 (*logging capabilities*), which are necessary to enable the traceability of operations throughout the entire life cycle of the system, with obvious implications in terms of *data security* and *privacy* management. Although the entire compliance assessment process is designed to ensure security, it requires considerable resources and specialist expertise, which clearly places a particular financial burden on SMEs.

An excessive compliance burden, particularly for start-ups and SMEs, runs the real risk of slowing down technological innovation in Europe. The high cost of technical and documentary requirements (Art. 16 et seq.), combined with the risk of penalties, could discourage small businesses from developing or using 'high-risk' systems, limiting their competitiveness compared to global *competitors* operating in less regulated jurisdictions.

In summary, the AI Act sets an ambitious ethical and safety standard, but its effectiveness will depend on the ability of European institutions and economic operators to overcome these technical, governance and economic sustainability challenges.

In any case, despite the potential risks described above, the AI Act is not only a burden for European operators and, under certain conditions, also for the legislator, but is also a strategic lever capable of generating significant economic and geopolitical opportunities. Regulation of this area can, in fact, constitute a competitive global advantage, transforming compliance into a real market *asset*. The AI

Act, as was the case with the adoption of the GDPR, can easily generate the infamous 'Brussels effect'. By imposing high *standards* on anyone wishing to operate in the Single Market, the EU is effectively pushing global economic players to adopt its *standards* worldwide, optimizing production processes and avoiding duplication of systems. As the first bloc in the world to establish a comprehensive and binding regulatory framework for AI, Europe can shape the international rules of the game, influencing future legislation in third countries.

The Regulation also defines a global benchmark for trustworthiness and safety. The 'Made in the EU, AI Compliant' label can become an internationally recognized mark of quality, facilitating the export of European technologies. The harmonization of standards at European level eliminates regulatory fragmentation between Member States, reducing *compliance* costs for European companies and creating a digital Single Market for AI. Stringent requirements for 'high-risk' systems impose rigorous development processes that not only increase safety but also enhance the technical quality and reliability of European products, making them more attractive.

The AI Act should therefore be interpreted as a long-term strategic investment: an initial compliance cost that generates a return in the form of consumer confidence, international competitiveness and technological autonomy.