

Wladimiro Gasparri
University of Florence
Department of Legal Sciences – DSG
wladimiro.gasparri@unifi.it
ORCID ID: 0000-0002-9768-5725

Francesca Tesi
University of Florence
Department of Legal Sciences – DSG
francesca.tesi2@unifi.it
ORCID ID: 0000-0001-9279-9646

ARTIFICIAL INTELLIGENCE AND AI ACT: FROM THE INDIVIDUAL TO THE ALGORITHM?*

Abstract: *This paper focuses on the analysis of the AI Act as a risk-based regulatory model for AI aimed at balancing the need for technological development, which can ensure greater efficiency and well-being, with the protection of fundamental rights, for which potential risks must be identified along with appropriate mitigation measures. The goal is to prevent a public decision being left solely to the ‘thinking machine’, an expression of a mechanism of indifferenciation that, through mathematical-computational logic, flattens individual identities (social complexity) onto data, no matter how numerous.*

Keywords: *artificial intelligence, AI Act, public administration, rule of law.*

1. INTRODUCTION. FROM ‘EXPERT SYSTEMS’ TO ‘MACHINE LEARNING METHODS’: THE REASONS BEHIND AI ACT

Artificial intelligence represents a true paradigm revolution that cannot be considered on par with other technological innovations, even recent ones. The

* Report presented at the 6th International Scientific Conference ‘Legal Tradition and New Legal Challenges’ University of Novi Sad, Faculty of Law, Novi Sad 17th-19th October 2024. This paper is the result of the joint work of the two authors, with paragraphs 1-3 and 9-14 attributed to Wladimiro Gasparri and paragraphs 4-8 attributed to Francesca Tesi.

current moment coexists with an ‘anthropological fracture’¹ where numbers become a guarantee of truth, information replaces knowledge and connections replace relationships². The qualitative leap in this change comes from the transition from the logical representation of knowledge, which characterizes so-called ‘expert systems’ and refers to algorithms capable of making inferences and reasoning, to the possibility of applying machine learning methods to machines through access to large datasets (so-called ‘big data’) in an infinitesimal temporal space, enabling the constant development of new criteria for inference between data, prearranged for the assumption of efficient decisions. This is a ‘thinking machine’, a machine capable of learning, which achieves the goals assigned to it without the need for human guidance on how to proceed and sometimes without even the human being having any awareness of what is happening inside the system-machine³.

The availability and processing of these vast collections of data have opened new operational horizons and revived the debate on AI and law, as well as AI and public administration. All of this clearly challenges old certainties and seems to require a rethinking of the legal categories previously known and used, starting with the principle of legality itself: does it make sense to talk about ‘algorithmic legality’ and, if so, in what way?

AI presents a range of challenges to the legal system: its ‘infinite’ potentialities transform both public and private organizations and activities, leading to the potential replacement of humans in favor of the ‘thinking machine’. The enormity of what’s at stake has made it urgent to have a regulatory response: AI is seen as an ‘object’ that must necessarily be regulated to balance the opportunities of technological progress (efficiency and well-being) with risks (manipulation and discrimination), in a perspective of reliability, security, and transparency, with the goal of protecting the fundamental rights of citizens. This approach goes beyond the protection of intellectual property and responsibility for errors induced by

¹ See Antoine Garapon, Jean Lassègue, *La giustizia digitale. Determinismo tecnologico e libertà*, Il Mulino, Bologna 2021.

² Cf. Alberto Andronico, Thomas Casadei, “Introduzione”, *Ars interpretandi* 1/2021, 7-11.

³ Until the end of the last century, research mainly focused on the development of AI models based on the peculiarities of legal reasoning and the knowledge specific to law (the reference in the Italian context is to the research of Vittorio Frosini, *Cibernetica, diritto e società*, Comunità, Roma 1968, and Mario G. Losano, *Giusecibernetica. Macchine e modelli cibernetici nel diritto*, Einaudi, Torino 1969). However, machine learning technologies have upended this landscape: ‘artificial thought’ can no longer be summarized in the “finite sequence of instructions, well-defined and unambiguous, so that they can be mechanically executed and produce a determined result (such as solving a problem or performing a calculation)”; it is no longer limited to applying predefined software rules and parameters, but is characterized by a permeable relationship with the vast amount of data to which it has access in an infinitesimal temporal space, capable of “constantly processing new inference criteria between data and making efficient decisions based on such processing, according to a process of machine learning” (Cons. Stato, sec. III, 25 November 2021, no. 7891).

expert systems of the 1990s, encompassing the dangers that may affect democratic institutions and social relationships as we have known them until today⁴.

To confirm that the nature of the debate is geopolitical, transcending borders and impacting the global market, these challenges were first addressed by the European legal system, even before national ones. The first regulatory interventions mostly came from jurisprudence. In this area, European legislation is very broad, and most recently, it was enriched by Regulation (EU) 2024/1689 of 13 June 2024, which establishes harmonized rules on artificial intelligence (the so-called 'AI Act'). These rules will take full effect by 2026.

2. THE CHARACTERISTICS OF AI ACT: THE 'HORIZONTAL REGULATORY APPROACH' AND THE 'BRUSSELS EFFECT'

Beyond its content, which will be discussed later, in the face of a technology capable of fueling a kind of 'permanent revolution', where the idea of 'dominating' it seems like an ambitious utopia and with the 'natural' risk that the regulation itself may prove to be delayed, the choice made by the AI Act was neither a 'defensive' one, concerned with regulating AI applications in specific sectors or with reference to certain subjects, nor a 'proactive' one, aimed at establishing rules to facilitate the development of new technologies in sectors deemed strategic. Instead, the choice was to regulate AI as a whole (the so-called 'horizontal regulatory approach') considering the need to establish general rules for a phenomenon regarded as new and with unclear boundaries⁵.

European regulation appears, first and foremost, to be marked by a geopolitical strategy: "to improve the functioning of the internal market". The goal appears to be twofold: on one hand, to position the EU as a leader in regulatory production in the field of AI, so as to establish a global reference and a model for other legal systems (the so-called 'Brussels Effect'), through the creation of a "uniform legal framework in particular for the development, the placing on the market, the putting into service and the use of artificial intelligence systems". On the other hand, "to support innovation" through policies promoting the development, use and dissemination of "human-centric and trustworthy artificial intelligence, while ensuring a high level of protection of health, safety, fundamental rights as

⁴ Such regulation is all the more necessary because the new AI technologies simultaneously give rise to new knowledge (understood as the ability to analyze and predict, also with the possibility of generating errors and inaccuracies) and new powers (understood as the ability to control and direct), primarily in the hands of large corporations and private entities. Not only individuals, but even public organizations, when lacking the necessary knowledge and adequate resources, find themselves at a disadvantage, which risks producing the phenomenon of so-called 'regulatory capture'.

⁵ Giusella Finocchiaro, "La regolazione dell'intelligenza artificiale", *Rivista trimestrale di diritto pubblico* 4/2022, 1088 *et seq.*

enshrined in the Charter of Fundamental Rights of the European Union, including democracy, the rule of law and environmental protection, against the harmful effects of AI systems within the Union” (*Recital* 1 and Art. 1, Reg. (EU) 2024/1689).

The European legislator’s perceived need is to “accelerate the process of development and the placing on the market of AI systems” (*Recital* 141), through a change that is accepted and shared by individuals and the business system: both must “be able to trust the technology they interact with, have a predictable legal environment and rely on effective safeguards protecting fundamental rights and freedoms” (COM(2018) 237 final of 25 April 2018, para. 3.3). Therefore, the objective is to design, develop and distribute ‘safe, reliable and ethical’ AI systems and applications, anchored to the respect for the principles and values of the European legal system and the guarantee of individuals fundamental rights (*Recital* 8 and COM(2018) 237 final, para. 4).

3. THE ‘VARIABLE LEVEL OF AUTONOMY’ AS A TYPICAL CHARACTERISTIC OF AI SYSTEMS AND THE RISKS OF ‘SELF-MOTION’

The first legally significant aspect raised by the use of AI technologies and any attempt to regulate them is the normative definition of the object ‘artificial intelligence’, which is directly linked to the identification of risks (and thus the remedies). There is no commonly accepted definition of what AI is today. The solution adopted at the European level is the result of a long debate and represents a sort of convergence point of trends that have emerged internationally, with the starting point being the formula proposed by the OECD since 2019.

‘Artificial intelligence’ is defined as “a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments” (Art. 3, no. 1, Reg. (EU) 2024/1689).

The defining and peculiar characteristic of AI systems is, therefore, their variable level of autonomy, which is entirely new compared to those previously known, enabling the machine to perform tasks to achieve specific assigned goals without the need for any ‘guidance’ (management, accompaniment, prompting) by humans, because the machine is ‘*sapiens*’, meaning capable of achieving the goal *ex se*⁶. This feature of AI is related to its ‘inferential capacity’, that is, its ability to “transcend basic data processing by enabling learning, reasoning or modeling” and thus the

⁶ Cf. Germana Lo Sapio, “L’Artificial Intelligence Act e la prova di resistenza per la legalità algoritmica”, *Federalismi.it* 16/2024, 275 et seq.

possibility for the system “to change while in use” (*Recital* 12). The distant statement remains relevant that we are no longer dealing with a system that “enables the execution of a service without the direct intervention of the service provider’s work”, a mechanism through which “a service is performed through an act to be executed by the person requesting the service”. Rather, we are dealing with a “mechanism that contains within itself the principle of its own motion”⁷.

The variable level of autonomy is both a strength and a weakness: autonomy, along with computational power, opens up new possibilities, but it also presents the risk of surpassing the ‘point of no return’: the point beyond which AI acts autonomously. The variable level of autonomy broadens and makes the boundaries of risk indefinite, although there is substantial convergence on some of these boundaries in (and between) the different scientific communities.

First, there is the risk related to the possible loss of human control over AI activity, which directly concerns the so-called ‘algorithmic legality’ (see below). Then, as a consequence of the previous risk, there is the risk of opacity (or incomprehensibility) of decisions, due to the ‘opacity’ of certain mechanisms used by AI, such as techniques based on artificial neural networks organized in multiple layers (so-called ‘deep learning’). Additionally, there is the risk of discriminatory distortions (biases) that could lead to the violation of fundamental rights⁸. Finally, there is the risk of the deviant use of AI systems, both for surveillance purposes and for undermining values and principles foundational to democracy.

It is clear that risk is a variable in the many applications of AI: any regulatory intervention that does not intend to prevent the development of research and hinder its outcomes cannot avoid incorporating the precautionary principle in the field of AI through the identification of proportional measures and regulatory models according to the level of protection sought⁹.

The challenge, therefore, lies in identifying the risks and managing the complexity of AI systems, the starting point of which consists of three key understandings. First, algorithms reflect human knowledge and the data used. Second, algorithms can be deterministic or non-deterministic, leading to unpredictability. Third, AI systems have transformed human activities in numerous sectors, increasingly

⁷ Antonio Cicu, “*Gli automi nel diritto privato*”, *Il Filangieri. Rivista giuridica, dottrinale e pratica* 1901, 561.

⁸ Discriminatory biases can be present in the vast amount of data or be linked to the multiple applications of the same AI system in different contexts or they can “stem from flaws in the overall design of AI systems (including with regard to human oversight) or from the use of data without correcting any distortions (for example, if a system is trained using only or predominantly data related to men, which results in suboptimal outcomes for women)” (COM(2020) 65 final – *White Paper on Artificial Intelligence. A European Approach to Excellence and Trust*).

⁹ In this sense, AI is understood as “a family of rapidly evolving technologies that requires regulatory oversight and a safe, controlled space for experimentation, while ensuring responsible innovation and the integration of appropriate safeguards and risk mitigation measures” (*Recital* 138).

becoming one of the main sources of knowledge and data on which both individual and collective subjects base their decisions¹⁰.

4. RISK MANAGEMENT IN THE AI ACT: A MODEL OF ‘PRECAUTIONARY REGULATION’

The question of ‘what governance for AI’ found an initial answer in the AI Act, which adopted a risk-based approach. This means that it provides different requirements based on the potential risks to individuals arising from the use of AI systems. The greater the intensity and scope of the risk to the user or citizen, both as an individual and as part of an organized collective, the higher the obligations for the user and/or the provider of the system. As a result, the involved parties must be informed and adopt technical and organizational measures that take into account, from the outset, on the model of privacy by design and by default, considering the existence of risks to the individual.

This approach is already a characteristic of data protection regulation (Reg. (EU) 679/2016 of 27 April 2016), but in contrast to it, Reg. (EU) 2024/1689 does not provide universally applicable rules for the use of AI systems. Instead, it adopts a precautionary and differentiated regulation based on the specificities of the ‘AI product’. Prohibitions and rules are tailored to the level of risk present or presumed to arise (thus, simply potential) in relation to the different “intended purpose”, which refers to the use the provider, as the developer of the AI system, assigns to the system (Art. 3, para. 1, no. 12). No technologies are specifically allowed or excluded, nor are products or services favored or penalized based on the technology used. In this way, the regulation is characterized by heterogeneous obligations and limits on the multiple parties involved in the ‘chain of artificial thought’, “whose intensity varies according to the level of danger”¹¹.

The precautionary principle is applied proportionally, aiming to strike a balanced compromise between the need, on one hand, to establish strict rules to protect fundamental rights and freedoms by prohibiting the use of AI that could endanger them, and, on the other hand, to foster research and innovation in a crucial sector for the economy and global competition, allowing for ‘acceptable’ risks.

Furthermore, this forward-looking regulatory framework seems to present a double awareness. The first concerns the possibility that AI might be developed for general purposes and applied in contexts and for objectives initially unforeseen,

¹⁰ Cf. Andrea Simoncini, Samir Suweis, “*Il cambio di paradigma nell’intelligenza artificiale e il suo impatto sul diritto costituzionale*”, *Rivista di filosofia del diritto* 1/2019, 94.

¹¹ Daniela Messina “*La proposta di regolamento europeo in materia di Intelligenza Artificiale: verso una ‘discutibile’ tutela individuale di tipo consumer-centric nella società dominata dal ‘pensiero artificiale’*”, *MediaLaw* 2/2022, 216.

as it can be integrated into other AI systems. For this reason, the regulation provides specific rules for AI ‘models’ and ‘systems’ for ‘general purposes’, intended for both direct use and integration into other systems (Arts. 51 *et seq.*).

The second refers to the highly evolving nature and rapid development of AI-related technologies. In this sense, to ensure the regulatory framework remains up-to-date, the regulation delegates broad legislative powers to the Commission under Art. 290 TFEU to modify, among other things, the conditions under which an AI system should not be considered high-risk and the list of high-risk AI systems (Arts. 6, para. 6, and 7, para. 1).

This set of choices has an undoubted strength in that it organizes the matter through the formulation of rules that, at the same time, aim to steer its development. However, it also presents weaknesses likely tied to the very nature of the regulated subject, with respect to which the risk-based approach seems ‘weak’ because AI systems are hardly reducible to ‘dangerous products’ since the regulatory framework may quickly become outdated due to their dynamic and evolving character. Moreover, AI models intended for ‘general purposes’ are becoming increasingly widespread.

Within this framework, Reg. (EU) 2024/1689 identifies four different levels of AI system risk, with corresponding obligations for the involved parties: prohibited systems that are unacceptably risky, high-risk AI systems that are acceptable but subject to obligations, minimal-risk systems, and, in a residual manner, zero-risk systems. For each of these categories, the regulation outlines prohibitions, minimum design and development requirements, transparency obligations and measures to promote self-regulation. In addition to these four categories, the regulation, as already mentioned, defines specific rules for ‘general-purpose AI models’, identifying those ‘with systemic risk’. Finally, it should be noted that this regulation, besides not prejudicing the competences of the member States regarding national security, does not apply to AI systems if and to the extent that they are placed on the market, put into service or used with or without modifications exclusively for military, defense or national security purposes, regardless of the entity carrying out such activities, nor to AI systems or models, including their outputs, developed and put into service solely for scientific research and development purposes (Art. 2, para. 3 and 6).

5. AI SYSTEMS BANNED DUE TO UNACCEPTABLY HIGH RISKS: WHEN DETAIL CREATES EXCEPTIONS

The first category concerns AI systems that are banned because they are unacceptably risky: these are AI systems that present a risk deemed unacceptable due to their potential to cause such high danger that it becomes incompatible with the shared values and principles at the European level and thus with the protection

of the fundamental rights of individuals who might be subject to personal data processing by such systems.

In this case, both placing these systems on the market and putting them into service are prohibited, as well as their use within the European Union. More specifically, this category includes AI systems that use “subliminal techniques beyond a person’s consciousness or purposefully manipulative or deceptive techniques, with the objective, or the effect of materially distorting the behaviour of a person or a group of persons by appreciably impairing their ability to make an informed decision, thereby causing them to take a decision that they would not have otherwise taken” (Art. 5, par 1, lett. a); those exploiting “the vulnerabilities of a natural person or a specific group of persons due to their age, disability or a specific social or economic situation, with the objective, or the effect, of materially distorting the behaviour of that person” (lett. b); those intended for “the evaluation or classification of natural persons or groups of persons over a certain period of time based on their social behaviour or known, inferred or predicted personal or personality characteristics” (i.e., social rating practices) (lett. c). Similarly, AI systems that create or expand facial recognition databases through non-targeted scraping of facial images from the internet or closed-circuit television footage (lett. e), as well as, with significant exceptions, predictive policing systems and systems used to analyze a person’s emotions in workplaces and educational institutions (lett. d and f) are also banned. Furthermore, AI systems that biometrically categorize individuals “based on their biometric data to deduce or infer their race, political opinions, trade union membership, religious or philosophical beliefs, sex life or sexual orientation” (lett. g) are prohibited.

Finally, the use of AI systems for “‘real-time’ remote biometric identification” in publicly accessible spaces for law enforcement activities is prohibited (lett. h). In this regard, Reg. (EU) 2024/1689 clarifies the provisions contained in Reg. (EU) 2016/679, distinguishing biometric data related to physical characteristics such as facial images or fingerprint data (Art. 4, no. 14, Reg. (EU) 2016/2016) from ‘real-time’ remote biometric identification data, which includes personal characteristics like voice inflection or gait.

These AI systems are typically used to instantly or without significant delays identify and recognize a large number of individuals or their behavior simultaneously (such as during public demonstrations) without their active involvement, by comparing biometric data of an individual with data in any reference database, regardless of the technology, processes or specific biometric data types used. Although this technology is effective, it is evident that it is particularly invasive of individual rights and freedoms, creating a form of constant surveillance, in addition to the risk of distorted results and discriminatory effects related to age, gender, ethnicity, further exacerbated by the immediacy of the impact and the limited opportunities for control and correction when used in ‘real-time’.

However, this ban is not absolute and includes several specified exceptions in which the use of such systems ensures the protection of a public interest that overrides potential risks, such as the search for crime victims or missing persons or for perpetrators (or suspects) of specific crimes, the prevention of imminent threats to life or personal safety or of terrorist attacks. The use of the AI system in these cases is subject to prior authorization from the judicial authority (or the competent independent administrative authority designated by the Member State), specifying the geographical scope and the time of use, unless there is “a duly justified situation of urgency” (Art. 5, para. 3).

It is clear that these exceptions represent a point of criticism, as they, if interpreted broadly, may end up neutralizing and nullifying the effectiveness of the established bans. Moreover, the absolute nature of the bans in the regulation is somewhat diminished in relation to other cases, given the specific characteristics of the banned AI practices, as the prohibitions are very thoroughly constructed, effectively widening the scope of AI systems that fall outside the ban.

6. ‘HIGH-RISK’ AI SYSTEMS (ACCEPTABLE BUT WITH RESERVATIONS) BETWEEN QUALITY REQUIREMENTS, OPERATOR OBLIGATIONS AND CONFORMITY ASSESSMENT

The second category concerns ‘high-risk’ AI systems. In these cases, the provider, deployer, importer and distributor¹² must comply with a specific set of requirements, first and foremost, the *ex-ante* conformity assessment, aimed at ensuring that these systems meet the health, safety or fundamental rights protection requirements outlined by the Charter of Fundamental Rights of the European Union (CFREU) (*Recital 48*). This also ensures that a proper quality management system is in place for the entire lifecycle of the AI systems and that their use is in line with the instructions for use. In essence, in this case, the goal is to bring these AI systems within ‘sustainable’ levels of risk so that they can be placed on the Union market, or put into service or used, only if they meet certain mandatory requirements that exclude the possible generation of unacceptable risks to those interests considered inviolable by European law.

¹² The provider is a natural or legal person, public authority, agency or another body that develops an AI system or an AI model for general purposes or has an AI system or AI model developed for general purposes, with the intention of placing that system or model on the market under its name or trademark, whether for consideration or free of charge. The deployer is a natural or legal person, public authority, agency or another body that uses an AI system under its own authority, excluding cases where it is used in the course of personal, non-professional activities. The importer is a natural or legal person located or established in the Union that places an AI system on the market bearing the name or trademark of a natural or legal person established in a third country. The distributor is a natural or legal person in the supply chain, other than the provider or importer, that makes an AI system available on the Union market.

The category of ‘high-risk’ systems is divided into two groups, depending on whether the AI is incorporated, as a safety component, into one of the products whose manufacture is already regulated by European harmonization legislation (for example, medical devices) (Art. 6, para 1) or is designed to be used as an ‘independent element’ (Art. 6, para. 2).

In the first case, the AI system (a) can be used as a safety component of a product (so-called ‘integrated’ or ‘incorporated’ systems) or (b) can assist the functionality of such a product without being incorporated into it (so-called ‘non-integrated’ systems). In the latter case, it is considered ‘high-risk’ if the product itself is subject to a conformity assessment by third parties for market placement or service introduction according to European harmonization laws, to avoid duplication in the conformity assessment and dilution in the chain of responsibilities.

‘Independent’ AI systems are considered ‘high-risk’ when used in certain predefined sectors, including biometrics, education, employment, access and provision of essential public and private services, migration, justice and democratic processes (Annex III of the Regulation). For each of these, there are ‘subclasses of use’, which include those for biometric identification of individuals, traffic management safety or the provision of water, gas or electricity. The list is ‘open’ as the Commission is empowered to adopt delegated acts to modify the content of Annex III (Art. 7, para. 1).

‘High-risk’ AI systems must comply with a series of specific obligations and requirements that take into account three profiles: the intended purpose of the AI system, the generally acknowledged state of the art in AI and the risk management system (Art. 8, para. 1).

They require the availability of understandable information about their operation throughout their lifecycle to ensure traceability, verify compliance with regulatory requirements and monitor their operation before and after market placement. For this purpose, the automatic recording of events (‘logs’) shall be preserved for the entire lifetime of the system (Art. 12) and technical documentation shall be kept up-to date (Art. 11).

It is also necessary to establish, implement and document a risk management system as an “continuous iterative process planned” that requires “regular systematic review and updating” to ensure its ongoing effectiveness, with the aim of identifying and mitigating risks to health, safety and fundamental rights (Art. 9, para. 1 and 2).

Risk management systems “shall be tested for the purpose of identifying the most appropriate and targeted risk management measures” (Art. 9, para. 6) to ensure the effectiveness of the entire procedure and data sets “shall be subject to data governance and management practices appropriate for the intended purpose” and “shall be relevant, sufficiently representative, and to the best extent possible, free of errors and complete in view of the intended purpose” (Art. 10, para. 1-3).

The commercialization of the ‘high-risk’ AI system shall also be preceded by the preparation of clear and understandable technical documentation demonstrating compliance with the requirements (Art. 11, para. 1 and 2). Finally, it should be noted that the deployer is required to carry out an impact assessment on fundamental rights before use to identify specific risks and the corresponding measures (Art. 27).

This regulation also highlights the need to limit the ‘high-risk’ AI systems by further delineating them according to the individual applications defined in Annex III, thus minimizing “potential restriction to international trade” (*Recital* 46). The reference here is to the provisions stating that while the use of an AI system in the sectors and for the purposes listed in Annex III “shall always be considered to be high-risk where the AI system performs profiling of natural persons” (Art. 6, para. 3, clause 3), outside of this situation, even if used in those ‘sensitive’ sectors, it is not considered ‘high-risk’ when it does not present a significant risk of harming the legal interests protected in these sectors, as it does not materially influence the decision-making process or adversely affect those interests in a substantial manner (Art. 6, para. 3, clauses 1 and 2).

This complex regulation concerning ‘high-risk’ AI applications provides a sort of ‘safety exit’ from what appears, even to the European legislator, to be the actual reality of algorithmic development, characterized by rapid and continuous evolution, driven by geopolitical competition even before economic competition, and increasingly difficult to govern through specific rules. If these rules risk blocking scientific development, the solution is to shift from ‘rules’ to ‘principles’ which are the foundational elements of so-called ‘algorithmic legality’: transparency, human oversight and accountability.

7. ‘CODES OF CONDUCT’ AS A TOOL FOR PROMOTING ETHICAL AND RELIABLE AI IN ‘LOW-RISK’ AI SYSTEMS AND THE RESIDUAL CATEGORY OF ‘ZERO-RISK’ AI SYSTEMS

The third category identifies AI systems that are ‘different from high-risk AI systems’. In this case, no specific legislation is prescribed, but rather, voluntary adoption of ‘codes of conduct’ by providers and deployers is encouraged, with the hope that the requirements set out in Chapter III, Section 2 of the Regulation (including technical documentation, record-keeping, transparency, human oversight and robustness) will also be applied to ‘low-risk’ AI systems, “taking into account the available technical solutions and industry best practices allowing for the application of such requirements” (Art. 95, para. 1, and *Recital* 165).

Among the objectives of the ‘codes of conduct’ are the evaluation and minimization of the environmental impact of AI systems, including energy-efficient

programming and techniques for the efficient design, training and use of AI, as well as promoting AI literacy, particularly for those involved in the development, operation and use of AI systems (Art. 95, para. 2, lett. b and c).

Regarding this category of systems, the market surveillance authority is also expected to intervene whenever there are sufficient grounds to believe that an AI system classified by the provider as not high-risk under Art. 6, para. 3, is actually ‘high-risk’. In this case, it is the responsibility of the surveillance authority to assess the AI system in question concerning its classification as a high-risk AI system based on the conditions set out in the aforementioned article. If the authority determines that the AI system in question is indeed ‘high-risk’, it “shall without undue delay require the relevant provider to take all necessary actions to bring the AI system into compliance with the requirements and obligations” set forth by the Regulation and take appropriate corrective measures within a prescribed period as determined by the authority (Art. 80, para. 1 and 2).

The fourth category is residual and includes all AI applications that present minimal or zero risk (for example, video games or spam filters in email messages), which are excluded from the application of the rules set by the Regulation. They must, however, comply with the general rules applicable to AI, such as those on the protection of personal data, competition, civil liability or consumer rights.

8. ‘AI FOR GENERAL PURPOSES’: THE ATTEMPT TO REGULATE GENERATIVE AI MODELS BETWEEN THE NEED FOR TRANSPARENCY AND PREVENTION OF SYSTEMIC RISKS

Beyond the various AI systems, Reg. (UE) 2024/1689 also addresses ‘AI for general purposes’, that is, those AI models characterized by significant generality and the ability to competently perform a wide range of distinct tasks, regardless of how the model is placed on the market, and which can be integrated into a variety of downstream systems or applications (Art. 3, para. 1, no. 63). It also includes generative AI models, which “allow for flexible generation of content, such as in the form of text, audio, images or video, that can readily accommodate a wide range of distinctive tasks” (*Recital* 99).

The regulation of these AI models is due, primarily, to the dynamic nature of artificial intelligence, its dependence on available data and individual usage contexts, all factors that can strongly influence the value chain. The state of the art sees the growing spread of both general-purpose AI models and machine learning technologies, capable of learning from data after the training phase, creating continuously new risks. It is evident that a timely regulatory response could not limit itself to regulating only the different AI systems but also had to consider this additional and pervasive technology, primarily to subject it to trans-

parency obligations, which, although ‘weak’, appear unavoidable. In other words, above the ‘pyramid taxonomy of risk’ consisting of different ‘AI systems’, there is an ellipsis represented by ‘AI models for general purposes’ (or basic AI models) for the development of special AI systems: essentially, ‘generative AI models’.

The distinction between ‘AI systems’ and ‘AI models for general purposes’ is tied to the functional characteristics of the latter, specifically their “generality and the capability to competently perform a wide range of distinct tasks”. Functionally, these models are “typically trained on large amounts of data, through various methods, such as self-supervised, unsupervised or reinforcement learning” to be used in a variety of environments and for a broad and diverse range of tasks (*Recital* 97). Although they can be used individually, they are essential components of AI systems, often integrated as a constitutive part in these systems, without being AI systems in themselves. In fact, to become AI systems, they require the addition of other components, such as a user interface. If the AI model for general purposes is integrated into an AI system, the latter is qualified as a ‘general-purpose AI system’, meaning “an AI system which is based on a general-purpose AI model and which has the capability to serve a variety of purposes, both for direct use as well as for integration in other AI systems” (Art. 3, para 1, no. 66).

This link and the resulting qualification are of considerable importance because this integration, and the very fact that an AI system has the capacity to pursue various objectives, ultimately allows its qualification as a general-purpose AI system, with the application of the corresponding regulation.

Within the category of AI models for general purposes, there are ‘AI models for general purposes with systemic risk’. The concept of ‘systemic risk’ is related to the particularly high computational capacity able to have a significant impact on “any actual or reasonably foreseeable negative effects in relation to major accidents, disruptions of critical sectors and serious consequences to public health and safety” or on democratic processes and public and economic security, the spread of illegal, false or discriminatory content, which endows the model with a highly evolving content (*Recital* 110)¹³.

The regulation sets an ‘initial threshold’ of “high impact capabilities evaluated on the basis of appropriate technical tools and methodologies, including indicators and benchmarks”. This threshold is reached when “the cumulative amount of computation used for its training measured in floating point operations is greater 10^{25} ” (Art. 51, para. 1, lett. a, and para. 2). This ‘initial threshold’ may be adjusted

¹³ Indeed, “systemic risks should be understood to increase with model capabilities and model reach, can arise along the entire lifecycle of the model”. These risks are influenced, among other factors, by “conditions of misuse, model reliability, model fairness and model security, the level of autonomy of the model, its access to tools”, as well as “potential intentional misuse or unintended issues of control relating to alignment with human intent” (*Recital* 110).

over time and it is up to the Commission to assess, *ex officio* or upon a qualified report from the expert group, whether the AI system has high-impact capacity equivalent to the ‘initial threshold’, considering the criteria indicated in Annex XIII (including, the amount and size of the data set and the input/output modes of the model) (Art. 51, para. 2, lett. b).

These two different types of AI models correspond to distinct regulations.

First of all, when an AI model for general purposes meets the condition specified in Art. 51, para. 1, lett. a, the provider must inform the Commission without delay, providing the necessary information to demonstrate compliance with the requirements. Similarly, if the Commission learns of a general-purpose AI model with systemic risks of which it was not informed, it may decide to designate it as a model with systemic risk. In both cases, a procedural process begins, during which the provider may demonstrate “that, exceptionally, although it meets that requirement, the general-purpose AI model does not present, due to its specific characteristics, systemic risks and therefore should not be classified as a general-purpose AI model with systemic risk”. The Commission may reject the observations made if they are not sufficiently substantiated and qualify the general-purpose AI model as one with systemic risk (Art. 52, para. 1-3). Upon a justified request by the provider, the Commission may always reassess the existence of ‘systemic risk’ (Art. 52, para. 5).

The regulation provides specific obligations for providers of ‘AI models for general purposes’, primarily related to their specific peculiarity: these models, especially generative AI models for generating text, images and other content, need to be developed and trained using access to large amounts of data. The extraction of these data raises two issues: transparency for deployers and protection of the copyright of the data owners. To address this, the regulation requires that providers of AI models for general purposes “draw up and keep up-to-date the technical documentation of the model, including its training and testing process and the results of its evaluation”, develop, update and make available “information and documentation to providers of AI systems who intend to integrate the general-purpose AI model into their AI systems” (so-called ‘downstream providers’), draft and make available to the public “a sufficiently detailed summary about the content used for training” and implement “a policy to comply with Union law on copyright and related rights” (Art. 53, para. 1, lett. a-d).

For providers of AI models for general purposes with systemic risk, additional obligations are foreseen in relation to the possible negative consequences, including conducting evaluations of the models in accordance with standardized protocols and tools to identify and mitigate systemic risks (Art. 55, para. 1, lett. a-d). Until the adoption of the related European regulation, compliance with the above obligations by providers of AI models for general purposes with ‘systemic risk’ “may rely on codes of good practice” at “the Union level”, the preparation

of which is encouraged and facilitated by the competent authorities (Arts. 55, para. 2 and 56, para. 1)¹⁴.

A summary consideration: the regulation of ‘AI models for general purposes’ is not only limited to general prescriptions but is primarily concerned with not interfering with the research, both fundamental and applied, that characterizes this field and is capable of generating significant economic value and has an extraordinary geopolitical importance. Moreover, it addresses the field of generative AI, machine learning, deep learning, chatbots, all applications using algorithms that automatically learn from past interactions to refine and respond better to current queries, thus without any human mediation. Not only that: chatbots based on generative AI can, among other things, cause risks to security, issues related to privacy and confidentiality and the constant risk of ‘hallucinations’ due to incorrect input data leading to inaccurate responses. In view of all this, the scope and content of the regulation appear ‘weak’, reduced to encouraging the adoption of good practices, with transparency duties in place, whose effectiveness is, moreover, dependent on a pre-existing trust relationship with the party upstream in the value chain.

9. ‘ALGORITHMIC LEGALITY’ BETWEEN THE ‘MYTH’ OF HUMAN OVERSIGHT AND THE ‘WEAK’ INDIVIDUAL EMANCIPATION

Finally, the AI Act does not forget to (re)formulate and (re)assert some of the principles already present in individual national legal systems through jurisprudence (with particular reference to first-generation algorithmic administrative activity), which are destined to form the core of what has been defined as ‘algorithmic legality’, that is, the set of principles aimed at ensuring individual guarantees whenever AI models or systems are involved in a particular activity or decision.

Individual guarantees take shape, first of all, in the ‘right to know’ whether a certain activity or decision has involved (or been influenced by) any form of artificial intelligence, in the ‘right to understand’ both the decision and the reasons underlying it, as well as in the ‘right to preserve the human element’, meaning “not to be subject to a decision based solely on automated processing, including profiling, which produces legal effects concerning him or her or similarly significantly affects him or her” (Art. 22, para. 1, Regulation (EU) 2016/679 of 27 April

¹⁴ To this end, the codes of conduct, in addition to addressing the obligations mentioned, must also provide that “the means to ensure that the information referred to in Article 53(1), points (a) e (b), is kept up to date”, a “adequate level of detail for the summary about the content used for training”, the “identification of the type and nature of systemic risks at Union level” and the measures, procedures and methods “for the assessment and management of systemic risks at Union level”, which shall be “proportionate to the risks, take into consideration their severity and probability and take into account the specific challenges of tackling those risks in light of the possible ways in which such risks may emerge and materialise along the AI value chain” (Art. 56, para. 2, lett. a-d).

2016). In this regard, with specific reference to the regulation of ‘high-risk’ AI systems, but with a scope that goes beyond this specific area, the AI Act reiterates the principles of ‘transparency’ (Arts. 13 and 50), ‘human oversight’ (Art. 14) and the ‘right to explanation of individual decision-making processes’ (Art. 86). Additionally, there is the duty of ‘accuracy, robustness and cybersecurity’, under which AI systems shall be “designed and developed in such a way that they achieve an appropriate level of accuracy, robustness, and cybersecurity, and that they perform consistently in those respects throughout their lifecycle” (Art. 15, para. 1).

This complex structure highlights one of the fundamental issues posed by AI: what is the most appropriate way to engage with new digital technologies and their applications, in order to make use of the extraordinary opportunities they offer to humankind, without humans becoming mere appendages to the technology? This seems to be the most profound reason behind the formulation of ‘algorithmic legality’: ensuring the centrality of the human subject so that the digital revolution does not become an autonomous process without a subject. This ‘new’ algorithmic legality has its roots in tradition: the idea underlying it is that the algorithm is a verifiable product and “traceable to the predetermined criteria that guide its operation” so that it becomes possible to imagine the possibility “for an expert user to understand its logic of operation” access the algorithm to verify its operational methods “and thus identify potential errors, try to identify its biases, and apply human oversight”¹⁵ to verify its functionality.

The model in the application so far adopted by administrative jurisprudence attempts to replicate in the algorithmic domain that set of tools and means typical of procedural protection: transparency, non-discrimination in selection, human supervision. Essentially, the attempt so far has been to replicate in the computational logic domain, the tools typical of classical logic: this operation becomes more difficult as the ‘algorithmic logic’ becomes increasingly autonomous from human predeterminations.

10. TRANSPARENCY AS THE EFFECTIVE KNOWLEDGE OF ‘SIGNIFICANT INFORMATION ON THE LOGIC USED’

The use of algorithmic technologies and the involvement of AI systems and models has always raised concerns related to their ‘opacity’ and ‘complexity’: the effectiveness of the ‘right to know’ becomes a priority for all parties involved in the AI value chain, and its first interpretation directly relates to the principle of transparency.

¹⁵ Enrico Carloni, “Dalla legalit  algoritmica alla legalit  (dell’amministrazione) artificiale. Premesse ad uno studio”, *Rivista italiana di informatica e diritto* 2/2024, 456.

In this regard, it is expressly stated that the design and development of AI systems, particularly those ‘high-risk’ systems, must occur “in such a way as to ensure that their operation is sufficiently transparent to enable deployers to interpret a system’s output and use it appropriately”. On one hand, the type and level of transparency must be adequate to meet the obligations imposed by the regulation on the supplier and the deployer. On the other hand, AI systems must be accompanied by user instructions, “that include concise, complete, correct and clear information that is relevant, accessible and comprehensible to deployers” (Art. 13, para. 1 and 2).

Essentially, the principles of reasonableness and proportionality define the scope of the informational and transparency obligations, whose intensity changes according to the risks associated with the AI systems and models. Furthermore, the interpretation of the right to know in terms of transparency seems to advocate for a sort of ‘digital good faith collaboration’ principle, recommending ‘loyalty and fairness’ from all parties involved in the technological supply chain. All of this will only acquire full effectiveness when it involves the recipient of activities related to the use of AI systems and models: only knowledge of the algorithmic origin of the decision enables the subject to consciously make their subsequent determinations, as it is also evident that only by knowing that a decision has been made with the help of or exclusively by an algorithm can the legal principles and rules contemplated in such cases by the legal system be applied.

The problem of transparency involves its content, that is, the scope of the ‘full knowledge of the algorithm’, because, evidently, the ‘algorithmic chain’ is the source of the decision (Cons. Stato, sec. VI, 13 December 2019, no. 8472). In general, it has been stated that “an algorithm is transparent for a group of users if they can understand it, memorize it, teach it, and use it”¹⁶, which precisely specifies the content of an ‘effective knowledge’ of the ‘significant information on the logic used’. But conversely, this brings attention to the opacity of the algorithm itself, to its almost natural inscrutability. This opacity is only partly ‘intentional’, related to the creative process, but is primarily ‘technical’ and ‘ontological’.

It is certainly ‘intentional’ because it is linked to intellectual property rights and thus due to the need to protect one’s intellectual property both for potential economic benefits that can result from its economic exploitation and to safeguard one’s privacy.

It is also a ‘technical opacity’ because the algorithm is written in machine language, the only language compatible with current computers, and uses a syntax that is completely separate from that of human language. Therefore, even when it is made available, consultable, or visible, it remains unreadable and unknowable in its contents, unless the subject possesses advanced mathematical-informatic knowledge that goes far beyond initiatives “aimed at promoting the digital literacy

¹⁶ Gerd Gigerenzer, *Perché l’intelligenza umana batte ancora gli algoritmi*, Raffaello Cortina, Milano 2024, 173.

of citizens” (Art. 8, Legislative Decree 7 March 2005, No. 82) and training programs for public employees.

Finally, it is an ‘ontological opacity’ because generative AI is subject to continuous and uninterrupted evolution, which ultimately prevents any form of control over the decision-making processes carried out by the algorithm in action, even by the programmer. Since AI feeds on data and learns from it, in this generative sense, it challenges the knowability and traceability that is, the ability of the system to explain the decision made and to retrace the steps taken.

So, what transparency? It has been correctly observed that “the goal of transparency can be satisfied by ensuring the actual intelligibility of the algorithmic process in all its phases”¹⁷, where by ‘intelligibility of the algorithmic process’, it should be understood as the ability to ensure the passage from formal knowledge of the data to the understanding of the phenomenon, with the reconstruction of the artificial cognitive processes that led to the final choice.

11. FROM ‘TRANSPARENCY’ TO THE ‘RIGHT TO AN EXPLANATION OF INDIVIDUAL DECISION-MAKING PROCESSES’: AN EFFECTIVE ‘DIGITAL GOOD FAITH COLLABORATION’

It is clear that the ‘intelligibility of the algorithmic process’ calls for, in fact, effective knowledge, demonstrated and falsifiable, and constitutes a critically important step addressed by the Regulation itself, where, with reference to ‘high-risk’ AI systems, it expressly recognizes the ‘right to an explanation of individual decision-making processes’. Indeed, it is provided that “any affected person subject to a decision which is taken by the deployer on the basis of the output from a high-risk AI system” and who “produces legal effects or similarly significantly affects that person in a way that they consider to have an adverse impact on their health, safety or fundamental rights shall have the right to obtain from the deployer clear and meaningful explanations of the role of the AI system in the decision-making procedure and the main elements of the decision taken” so that the concerned individual can exercise their rights (Art. 86).

This, however, does not resolve the question underlying many reflections on the relationship between AI, algorithms, and comprehensibility: is this demand for comprehensibility achievable or an utopia, if not even a misrepresentation? Is it imaginable, here and now, that every aspect of the algorithm and the ‘logical’ development of the decision be made comprehensible? Can the procedural sequence used to elaborate it, the decision-making mechanism, and the priorities

¹⁷ Antonella Mascolo, “*Gli algoritmi amministrativi: la sfida della comprensibilità*”, *Giornale di diritto amministrativo* 3/2020, 372.

assigned during the procedure, up to the data used in the formulation of the decision, be effectively reconstructed and made understandable?

The likely outcome given the current state of technology seems to approach a kind of simulacrum of knowability, an explicability reduced to the technical specifications provided by the provider in relation to the logical process that characterizes a particular AI system. This tension appears inevitable, and there is currently no technical solution to resolve it. Without any hypocrisy, the legal system only has the option to reduce the risk, a path the legal system has started down with the provision of appropriate ‘human oversight’ measures.

12. ALGORITHMIC EXPLICABILITY IN THE PRISM OF ‘HUMAN OVERSIGHT’: BRINGING ‘ARTIFICIAL THINKING’ BACK TO CLASSICAL LOGIC BY ATTRIBUTING THE CHOICE TO THE HOLDER OF POWER

In light of the uncertain practical implications of AI research, ‘human oversight’ currently seems to be the safest approach concerning the actual reality of algorithmic development: ‘high-risk’ AI systems “shall be designed and developed in such a way, including with appropriate human-machine interface tools, that they can be effectively overseen by natural persons during the period in which they are in use” (Art. 14, para. 1). Oversight measures must be in place from the very beginning: “appropriate human oversight measures should be identified by the provider of the system before its placing on the market or putting into service” (*Recital 73*). These measures must continue throughout the entire life cycle of the system and along the entire value chain: “human oversight shall aim to prevent or minimise the risks to health, safety or fundamental rights” that may arise when an AI system “is used in accordance with its intended purpose or under conditions of reasonably foreseeable misuse, in particular where such risks persist” despite the application of other guarantee and safety requirements set forth by the regulation (Art. 14, para. 2).

‘Human oversight’ attributed to “at least two natural persons with the necessary competence, training and authority” to perform this role (Art. 14, para. 5), subjects the AI system to operational constraints that the system itself cannot nullify and forces it to answer to the human operator. Oversight measures are not uniform and standardized, but are differentiated and commensurate with the risks, the level of autonomy, and the context of use of the AI systems (Art. 14, para. 3).

The common denominator of these measures is to ensure that the AI system (specifically ‘high-risk’) does not undermine human autonomy or produce other negative effects and allows, in real-time and retrospectively, verification of the logicity and correctness of the decision-making process through the provision of mechanisms that enable human intervention (human-in-the-loop), human super-

vision (human-on-the-loop), or human control (human-in-command). As already stated by administrative jurisprudence, “there must nonetheless be a human contribution in the decision-making process, capable of controlling, validating, or disproving” and we might add, of not relying on or stopping the automatic decision-making procedure. All of this “guarantees the attribution of the choice to the holder of the authoritative power, identified based on the principle of legality, as well as verification of the identification of the responsible subject, both in the interest of the public administration and of the individuals involved and affected by the administrative action entrusted to the algorithm” (Cons. Stato, sec. VI, 4 February 2020, No. 881). The provision of a particular outcome by the algorithm should always and under all circumstances lead to a renewed examination of the matter by a human. Clearly, this eventuality in itself negates any benefits, in terms of costs, speed, and efficiency, derived from the inclusion of the algorithm in the decision-making process, ultimately ending up duplicating, if not the time, at least the activity. However, even if one were to limit human intervention to the assessment of the non-manifest illegality of the outcome produced by the ‘*sapiens machine*’, the human attitude of submission and indifference towards algorithmic outcomes is undoubtedly the weak point of ‘human oversight’. The constant interaction between human and system thus risks disguising an approach where tasks are entirely delegated to the machine, and decisions are made autonomously by the system once it has reached the required level of accuracy in testing, moving from human-in-the-loop to human-out-of-the-loop.

Unless the role of the algorithm is restricted to being merely auxiliary and instrumental, the man-algorithm, official-machine confrontation thus risks leading to uncritical and passive acceptance of the ‘algorithm’s behaviour’, to a superficial and façade-like human oversight, reduced to a ritual. However, if “the use of ‘robotized’ procedures cannot be a reason for evading the principles that shape our legal system and regulate the conduct of administrative activity” (Cons. Stato, sec. VI, 8 April 2019, No. 2270), both organizational solutions that progressively make human oversight effective and the scope of the intervention of the ‘*sapiens machine*’ must be identified, not so much to exclude it but to specify its position of autonomy or auxiliary role.

13. THE LIMITS OF ALGORITHMIC EXPLAINABILITY AND THE ‘RETURN’ TO THE CENTRALITY OF THE SUBJECT: THE NON-EXCLUSIVITY OF THE ALGORITHMIC DECISION

This complex system of principles and rules defined by Regulation (EU) 2024/1689, in precarious balance between humanizing technology and not mechanizing the human being, implicitly recalls the so-called principle of the ‘non-ex-

clusivity of the algorithmic decision' as provided by Art. 22 of Reg. (EU) 2016/679. Indeed, when a decision based solely on automated processing, including profiling, produces legal effects that concern or significantly affect a person, that person has the right to ensure that such a decision is not solely based on that automated process (Art. 22 GDPR). Conversely, the individual becomes the leverage capable of legitimizing the replacement of human intelligence with the predictive algorithm: a prediction that therefore defines the person's role as both strong and weak.

Strong because, at least theoretically, it places the individual's self-determination at the center of the system, granting the individual a specific primacy. In the face of the capabilities of the '*sapiens* machine' and the deterministic drift associated with the use of algorithms in public and private activities, the rule re-centers the subject with an original and specific interpretation of the principle of self-determination in the digital space, understood as the possibility of at least partially avoiding the use of algorithmic forms in decision-making procedures.

The weakness certainly lies in the exceptions that appear potentially very broad: this provision does not apply, in fact, when the decision (a) is necessary for the conclusion or execution of a contract between the data subject and the data controller, (b) is authorized by Union law or the law of the Member State to which the data controller is subject, (c) is based on the explicit consent of the data subject (Art. 22, para. 2). In the cases under a) and c), "the data controller shall implement suitable measures to safeguard the data subject's rights and freedoms and legitimate interests, at least the right to obtain human intervention on the part of the controller, to express his or her point of view and to contest the decision" (Art. 22, para. 3).

But above all, the strength of this provision (self-determination) is, at the same time, its weakness, because it ultimately relies solely on the caution of the subject, i.e., on a configuration that is exclusively defensive.

14. THE DIFFICULT SUBSUMPTION OF THE SUBJECT INTO THE 'CALCULATING MACHINE'

We arrive here at the initial question, which is of a political and constitutional nature: is this instrumental-calculating rationality, built on the subsumption of each individual through digitalization and big data, desirable? Is it consistent with the constitutional centrality of the 'person'?

The algorithm processes input data according to a set of entirely formal rules, which, due to their formalism, do not need to refer to any meaning. The transformation of natural languages into alphanumeric languages allows for the possibility of calculating signs, ultimately mathematical ones.

In this way, information is transformed into strings of signs that are processed according to the primary principle of non-contradiction and the formal rules of mathematical calculation. These rules constitute the syntax of computational processes, which may vary in terms of the type of computational language chosen, but all necessarily have syntax without semantics, as their mode of operation responds only to the basic rule of composing links and operations of calculation, avoiding contradiction.

The ‘calculating machine’, once a tool available to humanity, risks becoming a tool for the governance of numbers, which, instead of opening up or facilitating areas of democracy and dialogical discussion, becomes a means of automatism that, stemming from its mathematical-computational nature, claims to guarantee objectivity in decision-making, as well as speed and decisiveness in behaviour.

If the observation still holds that “the nature of the human being is the sum of their social relations” which means that it depends much more on the context of social relationships than on the processes of algorithmic computation, the subsumption of the subject into the data seems difficult to reconcile with the idea of ‘personhood’. AI is always “a reminiscence of the past; conversely, constitutional law is always a projection toward a desired future, which, moreover, cannot be reduced to and framed within preset schemes and can always be implemented in light of historical contingencies”¹⁸. Therefore, it is essential to always remember that the algorithm is an expression and symbol of a mechanism of undifferentiation that flattens individual identities onto data, no matter how numerous.

And it is then clear how the question posed in the title lies along that ridge where ‘apocalyptic’ and ‘integrated’ views always clash. But our perspective is different: no triumphalism from those who celebrate a new era, nor the technological luddism of those who fear every form of innovation.

Only prompt attention to transformation, to nourish a critical evaluation of the phenomenon and its effects capable of maintaining the centrality of the person, the subject holding fundamental rights and freedoms, with the awareness that subsumption is mere presumption.

REFERENCES

- Andronico Aberto, Casadei Thomas, “Introduzione”, *Ars interpretandi* 1/2021, 7-11;
Carloni Enrico, “Dalla legalità algoritmica alla legalità (dell’amministrazione) artificiale. Premesse ad uno studio”, *Rivista italiana di informatica e diritto* 2/2024, 452-465;
Cicu Antonio, “Gli automi nel diritto privato”, *Il Filangieri. Rivista giuridica, dottrinale e pratica* 1901, 561-580;

¹⁸ Nicolò Lipari, “Diritto, algoritmo, predittività”, *Rivista trimestrale di diritto e procedura civile* 3/2023, 723.

- Finocchiaro Giusella, “*La regolazione dell’intelligenza artificiale*”, *Rivista trimestrale di diritto pubblico* 4/2022, 1085-1099;
- Frosini Vittorio, *Cibernetica, diritto e società*, Comunità, Roma 1968;
- Garapon Antoine, Lassègue Jean, *La giustizia digitale. Determinismo tecnologico e libertà*, Il Mulino, Bologna 2021;
- Gigerenzer Gerd, *Perché l’intelligenza umana batte ancora gli algoritmi*, Raffaello Cortina, Milano 2024, 173;
- Lipari Nicolò, “*Diritto, algoritmo, predittività*”, *Rivista trimestrale di diritto e procedura civile* 3/2023, 721-739;
- Lo Sapia Germana, “*L’Artificial Intelligence Act e la prova di resistenza per la legalità algoritmica*”, *Federalismi.it* 16/2024, 265-290;
- Losano Mario G., *Giuscibernetica. Macchine e modelli cibernetici nel diritto*, Einaudi, Torino 1969;
- Mascolo Antonella, “*Gli algoritmi amministrativi: la sfida della comprensibilità*”, *Giornale di diritto amministrativo* 3/2020, 366-375;
- Messina Daniela “*La proposta di regolamento europeo in materia di Intelligenza Artificiale: verso una ‘discutibile’ tutela individuale di tipo consumer-centric nella società dominata dal ‘pensiero artificiale’*”, *MediaLaw* 2/2022, 196-231;
- Simoncini Andrea, Suweis Samir, “*Il cambio di paradigma nell’intelligenza artificiale e il suo impatto sul diritto costituzionale*”, *Rivista di filosofia del diritto* 8/2019, 87-106.

Владимиро Гаспарри
Универзитет у Фиренци
Департаман за правне науке – DSG
wladimiro.gasparri@unifi.it
ORCID ID: 0000-0002-9768-5725

Франческа Теси
Универзитет у Фиренци
Департаман за правне науке – DSG
francesca.tesi2@unifi.it
ORCID ID: 0000-0001-9279-9646

Вештачка интелигенција и Акт о вештачкој интелигенцији (AI Act): Од појединца до алгоритма?

Сажетак: Овај рад посвећен је анализи Акта о вештачкој интелигенцији, као модела регулисања вештачке интелигенције који се заснива на ризику, у циљу успостављања равнотеже између употребе за технолошким развојем, који може да обезбеди већу ефикасност и благодинање, и заштити основних права, у односу на која постоје ризици и морају да се утврде, уз одговарајуће мере за њихово ублажавање. Циљ је да се спречи да се доношење јавних одлука преусти искључиво „машини која размишља“, као изразу механизма који, без разликовања, користићи машиничко-рачунарску логику, идентификује појединца (сложености друштва) поравнава и своди на податке, без обзира на то колико су они многобројни.

Кључне речи: вештачка интелигенција, Акт о вештачкој интелигенцији, јавна управа, владавина права.

Датум пријема рада: 17. 03. 2025.

Датум достављања коначне верзије рада: 15. 04. 2025.

Датум прихватања рада: 22. 04. 2025.