

Governing Artificial Intelligence and Protecting Human Rights: A Comparative Legal Perspective from the European Union and China

Maida BEĆIROVIĆ-ALIĆ¹, Jelica GORDANIĆ²

Abstract: The global expansion of automated systems, which often operate in a non-transparent manner and beyond effective human control, opens up complex issues related to the protection of basic human rights, which represents a broader context of this research. Particular challenges arise in the field of international law, since there is no single, binding transnational normative framework that would ensure consistent protection of human rights in the digital environment. By applying the comparative law method and analyzing the content of relevant documents, this paper examines the regulatory models of the European Union and the People's Republic of China, as the two most globally influential models of artificial intelligence management, in which liberal and illiberal concepts of human rights protection are clearly reflected, pointing to their legal-political framework, regulatory scopes and limitations. Special emphasis is placed on identifying normative gaps and risks to human rights, including privacy, non-discrimination, freedom of expression and fair treatment. The starting hypothesis of the paper is that the current international legal framework is unable to respond to the systemic challenges generated by the application of artificial intelligence, and that the divergent models of the EU and China deepen the global normative gap and prevent the establishment of universal standards for the protection of human rights. Through a critical comparison of the differences between Western liberal model and more authoritarian approach to technology management, the paper points to the urgent need to establish a globally

¹ Full Professor, University of Novi Pazar – Law Department, Novi Pazar, Serbia. E-mail: maida.becirovic@uninp.edu.rs, ORCID: <https://orcid.org/0000-0002-9738-6581>.

² Senior Research Fellow, Institute of International Politics and Economics, Belgrade, Serbia. E-mail: jelica@diplomacy.bg.ac.rs, ORCID: <https://orcid.org/0000-0001-8364-7405>.

harmonized and politically sustainable framework of international law for the regulation of AI systems.

Keywords: artificial intelligence, human rights, regulatory models, international law, EU, China.

Introduction

In the context of the global competition for technological and normative dominance, different political systems are developing distinct approaches to the regulation of artificial intelligence (AI), with far-reaching consequences for the realization and protection of human rights. The European Union has positioned itself as a global regulatory leader by constructing the first comprehensive legal framework for artificial intelligence, establishing an ambitious and value-based governance model. The regulation translates the constitutional obligation of the Union that technological progress serves people and society, treating AI not only as an economic instrument, but as an area of constitutional relevance that touches on democracy, the rule of law and human dignity. This approach relies on the EU's distinctive method, which combines market integration, rights protection and procedural governance — thus confirming European regulatory leadership in the digital age (Hulok 2025). In contrast, China is developing a sectorally fragmented but politically centralized model of AI regulation, which is deeply connected to strategies of social surveillance and technological governance. Floridi (2018, 2) states that the speed of technological innovation makes it difficult to establish an adequate normative framework, which is why existing legal and ethical mechanisms often fail to respond in a timely manner to the challenges of digital development. Similarly, Calo (2017, 408) emphasizes that the existing framework for governing AI is still underdeveloped, particularly due to its reliance on ethical standards rather than binding legal rules. These approaches suggest that the development of AI is taking place in a regulatory and normative vacuum, as the international system still lacks uniform norms for managing the new risks posed by AI technologies. The paper proceeds from the hypothesis that the current international legal framework is unable to respond adequately to the systemic challenges generated by the application of artificial intelligence, and that the divergent regulatory models of the European Union and China further deepen the global normative gap and hinder the establishment of universal standards for the protection of human rights.

Through a comparative analysis of these two divergent approaches, the paper identifies key normative gaps in international law, points to specific risks to human rights such as the right to privacy, non-discrimination and freedom of expression, and reveals that the absence of a single transnational framework enables the emergence of a global “normative vacuum”. In the concluding part, the need for a politically sustainable and internationally coherent model of global management of artificial intelligence that would ensure predictable and value-based protection of the individual is discussed.

Methodologically, the paper combines comparative legal, descriptive and critical-analysis methods, using a qualitative approach in examining relevant legal acts, strategic documents, and secondary literature. Comparative legal analysis enables the identification of differences between the European liberal model of human rights protection and the Chinese centralized model of technological management, while the critical analysis reveals the consequences of such differences for the international legal order and global standards.

Although comparative literature often highlights the sharp contrast between European and Chinese models of AI regulation, it is important to avoid an overly rigid dichotomy. Despite the fact that the normative framework of the European Union is based on the protection of the dignity, privacy and autonomy of the individual —values widely recognized as constitutive principles of the EU legal order (Huster and Walden 2019; Lynskey 2015) —the European model is not without limitations. Certain forms of biometric surveillance remain permissible, member states continue to apply broad exceptions in the areas of police powers and security, and the implementation of the General Data Protection Regulation remains uneven across the Union. In addition, forms of technological paternalism and market monopolies point to structural weaknesses. Modern digital ecosystems, especially those based on advanced artificial intelligence systems, further intensify these challenges. Thus, the real scope of protection of dignity, autonomy and privacy depends on the ability of the regulatory system to respond to rapidly growing technological risks, despite the EU value-based foundations. On the other hand, although the Chinese model is indeed characterized by political centralization and a strong focus on social stability, it is not monolithic. China’s academic community exhibits a certain degree of pluralism—Zhu (2022), for example, documents concerns expressed by many domestic experts regarding the overuse of surveillance technologies—while in practice initiatives promoting a framework of “responsible artificial intelligence” such as the Beijing AI Principles (2019) are simultaneously being developed. Also,

certain segments of Chinese legislation seek to limit corporate abuses and improve data security, which is confirmed by the new data management regime established by the Personal Information Protection Law and the Data Security Law, which introduce “strict obligations for private actors with the aim of suppressing excessive data collection, ensuring legal processing and strengthening overall data security mechanisms” (Creemers 2022).

The authors therefore believe that the polarization between the EU and China certainly exists, but that it is not absolute: both models contain internal tensions, contradictions and evolutionary processes that make it difficult to present them as completely opposed concepts of human rights in the digital era. Accordingly, the analysis focuses on a comparative examination of the EU and Chinese regulatory models in order to illustrate how their structural differences contribute to existing normative gaps in international law and to assess the implications of such divergence for the future development of a transnational framework for AI governance.

The European Union as a regulatory leader in the field of Artificial Intelligence

The legal framework of the European Union for the regulation of artificial intelligence represents the most developed horizontal model in contemporary international law. It is based on a combination of the General Data Protection Regulation (GDPR 2016), the EU Charter of Fundamental Rights (EU Charter 2000), the digital regulation of platforms Digital Services Act (DSA 2022), Digital Markets Act (DMA 2022) and, finally, the first global comprehensive law on artificial intelligence — the EU Artificial Intelligence Act (AI Act 2024). The European Commission established the European Artificial Intelligence Office by Decision of 24 January 2024. The binding legal force of the Artificial Intelligence Act, however, derives from its legal nature as a regulation of the European Union that is directly applicable and binding on all persons and entities that develop, use, import or place on the market artificial intelligence systems within the Union. In this framework, the European Artificial Intelligence Office was established as a specialized body with the task of monitoring, coordinating and supporting the implementation of the AI Act, in particular with regard to general-purpose artificial intelligence systems. Therefore, it is expected to play a

significant technical and institutional role in the emergence of a European artificial intelligence governance system (Dabić 2025, 235).

This normative architecture reflects the Union's ambition to position itself as a European and global protector of fundamental rights in the digital era, which Palmiotto describes as "the constitutional orientation of the entire policy of artificial intelligence" (Palmiotto 2025, 770). The EU therefore does not treat artificial intelligence primarily as a technological challenge, but as a legal problem that requires the application of standards of dignity, autonomy and protection of the individual.

The AI Act introduces a risk-based regulatory paradigm, under which obligations and restrictions are not applied equally to all AI systems, but instead vary according to the level of risk posed to security and fundamental rights. Thus, the legal permissibility of technology development is conditioned by measurable and preventive guarantees, which Pirozzoli identifies as the most significant institutional step of the Union in the management of technological risks (Pirozzoli 2024, 106). This approach builds on existing regulations in which the GDPR remains the foundation of privacy protection and data control, whereby, according to Ufert (2020, 1089–1091), it is capable of functioning as an "adaptive and disruptive instrument" in relation to risky forms of automated decision-making, although it does not provide sufficient mechanisms of explainability and transparency for complex AI models. The EU Charter of Fundamental Rights provides constitutional limits to the use of AI, particularly in the areas of privacy, non-discrimination, fair treatment and procedural guarantees, thereby protecting fundamental freedoms regardless of technical developments (Cosentini et al. 2024, 2–3).

Within the AI Act, the EU introduces a multi-layered categorization that includes prohibited practices, high-risk systems with strict obligations, limited-risk systems and minimal-risk systems (AI Act 2024). Prohibited practices include social ranking, manipulation of human behavior and certain forms of biometric surveillance, as they are considered incompatible with the dignity and autonomy of the individual (Pirozzoli 2024, 106). The most important and complex segment is the regulation of high-risk systems, for which strict requirements are introduced regarding data management, technical robustness, human supervision, transparency, security and documentation of the decision-making process (AI Act 2024). In this segment, the EU is introducing the biggest novelty: Fundamental Rights Impact Assessment for High-Risk AI Systems — FRIA, as a preventive instrument that allows assessing the risks to individual rights even before the application of technology. Cosentini et al. (2024, 2) refer to the FRIA

as “the first legally binding mechanism for the anticipatory protection of fundamental rights in the field of artificial intelligence”. Bertaina et al. (2025, 2–3) complement this framework by developing a quantitative risk assessment matrix, which enables the identification of the intensity and probability of violation of each individual right from the EU Charter, including privacy, equality, child protection, freedom of expression and the right to a legal remedy. Thus, for the first time, the EU operationalizes the protection of fundamental rights in a way that enables measurability, verifiability and regulatory oversight.

The impact of the European model on the protection of human rights is most evident in the area of privacy and data protection. The GDPR, as a central pillar, provides a wide range of data subject rights, but the EU through the AI Act further strengthened the mechanisms of transparency, explainability and data management for high-risk AI systems (GDPR 2016; AI Act 2024). This is particularly important in AI applications involving behavioral monitoring, biometric identification or profiling, where the risk of privacy violations is greatest. Ufert (2020, 1090–1091), however, warns that the GDPR does not solve the problem of algorithmic non-transparency, which is why it is necessary to supplement it with sectoral regulatory obligations. It is precisely this normative gap that the Artificial Intelligence Act (AI Act) seeks to fill, introducing mandatory requirements regarding transparency, technical documentation, human oversight, and explainability for high-risk AI systems. In doing so, the Act operationalize core data protection principles within the specific context of automated decision-making process (AI Act 2024).

Another key area of protection relates to non-discrimination and fairness. AI systems can reproduce or intensify existing social inequalities, so the EU introduces mandatory bias mitigation measures, data quality standards, the obligation to test and monitor the effects of models on vulnerable groups, and requirement of human oversight over critical decisions. Pirozzoli (2024, 107–108) emphasizes that the European human-centric approach is fundamentally based on preventing structural discrimination through ex ante normative interventions, rather than relying solely on ex post corrective mechanisms.

The third area — the protection of procedural rights and fair treatment — is most clearly seen in the regulation of algorithmic decision-making in public services, police systems, employment, social benefits and the judiciary. Palmiotto shows that during the legislative process the AI Act underwent a “roller coaster” of changes, whereby the protection of fundamental rights, despite political compromises, remained one of its supporting pillars (Palmiotto 2025, 772–773).

The mandatory FRIA procedure requires AI system owners to assess in advance the possibility of unequal outcomes, inexplicable decisions or violation of the right to a legal remedy, thereby embedding procedural fairness into the very structure of the development of AI technologies.

Taken as a whole, the European model of artificial intelligence regulation represents a unique attempt in contemporary law to harmonize technological progress with firmly institutionalized protection of human rights. By introducing preventive mechanisms such as FRIA, standards of explainability, strict data control and prohibited practices, the EU has created a normative framework that sets clear boundaries for technological innovation to protect dignity, privacy, non-discrimination and fair treatment. This model remains imperfect and dynamic, but it undoubtedly represents the most advanced attempt to translate the principles of human rights embedded in the European legal heritage into an artificial intelligence management system.

However, the structural weaknesses that characterize this system limit its effectiveness. Palmiotto (2025, 772–773) indicates that the AI Act is the product of political compromises that during the legislative process led to the dilution of some obligations and the retention of areas of regulatory flexibility that may weaken individual protection in practice. The legislative process in the EU is notably slow, which means that the law often lags behind technological changes, leaving long-lasting regulatory gaps where AI systems are implemented without adequate oversight. An additional problem is the uneven implementation of regulatory obligations across Member States. While some countries develop and implement technologically advanced surveillance systems, including biometric identification and mass data processing in the context of policing and border control, others lack comparable technical and institutional capacity, leading to significant differences in the actual level of human rights protection across the Unio. Ufert (2020, 1090–1091) points out that, despite strong privacy guarantees, the GDPR still does not solve the problem of non-transparency of complex models of high autonomy, while Pirozzoli (2024, 107–108) notes that the human-centric approach often remains declarative when faced with the economic interests of the market and the political interests of member states. The lack of technical capacity of supervisory authorities and limited resources also weaken the EU's ability to effectively oversee large technology companies, while the fragmentation of sector-specific regulation (DSA, DMA, Data Act and AI Act) creates normative complexity that complicates enforcement in practice. These shortcomings may be partially mitigated by the establishment of the European Artificial Intelligence Office, which

is intended to concentrate technical expertise, coordinate national supervisory bodies and support the uniform application of the AI Act, particularly with regard to general-purpose AI systems. However, its effectiveness will ultimately depend on the scope of its formal powers, the adequacy of its resources and the quality of institutional cooperation with national authorities.

These weaknesses have direct consequences on the enjoyment of human rights of EU citizens. Regulatory delays means that citizens are exposed to risks to privacy, non-discrimination and autonomy before systems are adequately evaluated and limited, while uneven implementation threatens equality of protection — a person in one state may have significantly weaker procedural guarantees than in another. Ambiguities surrounding biometric surveillance open the way for excessive collection of sensitive data, which threatens the right to privacy and may lead to the normalization of mass surveillance. Poorly developed mechanisms of explainability continue to enable non-transparent algorithmic decisions affecting employment, credit, social rights or law enforcement procedures, thereby challenging the right to a reasoned decision, equality and an effective legal remedy. Finally, the insufficient capacity of supervisory bodies allows large platforms to circumvent regulations, which threatens consumer protection, pluralism of information and freedom of expression. Taken together, these shortcomings do not undermine the normative value of the European system, itself. However, they significantly complicate the practical realization of the level of human rights protection that the EU formally and declaratively seeks to guarantee.

The authors consider that an additional challenge is the uneven implementation of regulatory obligations in the member states, as well as the high administrative burden for small and medium-sized enterprises. Despite this, the strength of the EU model is reflected in its fact that it is the only one at the global level to systematically integrate the protection of human rights into the core of technological regulation through Fundamental Rights Impact Assessment for High-Risk AI Systems (FRIA), prohibition of the most dangerous practices and clearly defined guarantees for high-risk systems (AI Act 2024). Further development of this model would require strengthening mandatory explainability requirements for algorithmic decision-making in order to reduce existing transparency gaps. It would also require enhancing the institutional role and operational capacities of the European AI Office as a central supervisory authority capable of ensuring a more consistent interpretation and implementation of AI regulation across Member States. In addition, the FRIA

mechanism should be expanded to encompass complex generative AI models, accompanied by clearer standards for assessing risks related to freedom of expression, privacy, and non-discrimination.

Europe is showing initiative and leadership when it comes to the regulation of artificial intelligence. In addition to the EU, particular attention should be given to the Council of Europe and its Framework Convention on Artificial Intelligence and human rights, democracy and the rule of law. This Convention represents first legally binding international treaty in the field of artificial intelligence. The provisions of the Convention aim to ensure that activities within the lifecycle of artificial intelligence systems are fully consistent with human rights, democracy and the rule of law (Council of Europe 2024). The Convention has been opened for signature since September 2024, and more than 50 states, including the EU members, have signed it thus far.

The Chinese model of regulating Artificial Intelligence

The Chinese model of regulating artificial intelligence rests on a strategic combination of normative fragmentation, political centralization and technological instrumentalization, which creates a highly adaptive, but at the same time deeply authoritarian system of managing the digital space. Unlike the European Union, which introduces a horizontal and comprehensive approach to the protection of fundamental rights in the field of artificial intelligence, China is developing a vertical model composed of specific laws and by-laws, the most important of which are the Cybersecurity Law (CSL CN 2017), Data Security Law (DSL CN 2021) and Personal Information Protection Law (PIPL CN 2021). These laws form the basic legal framework for data governance in China and, indirectly, for the operation of AI systems, in a manner comparable to the role played by the GDPR in the European Union, which likewise was not adopted as an AI-specific instrument but functions as a foundational layer for regulating data-intensive technologies, including artificial intelligence. This regulatory specificity enables more targeted governance of risks associated with particular sectors. For example, the Deep Synthesis Regulation (DSR CN 2022) explicitly introduces the obligation to label manipulated visual content, which is an important tool in the fight against misinformation, protection of public interest and reduction of citizens' reputational damage. In this sense, the Chinese model provides transparency to users (eg labeling "synthetic content"), which is in line with the protection of the right to

information and the reduction of manipulative content. However, in the normative sense, exceptions in favor of the state are broadly defined, which in practice enables unlimited access to data for the purposes of “national security” and “social stability”. As a result, this legislative configuration makes the protection of human rights dependent on the political assessment of state authorities, opening up space for systemic and permanent digital surveillance.

In this context, Roberts et al. (2021, 4–7) show that China’s AI policy, although outwardly appearing coherent, actually represents a broad but strategically inconsistent combination of ethical principles, bureaucratic guidelines and political directives, with the formal call for accountability and transparency mainly serves to legitimize an already established system of control. This normative dualism is particularly evident when looking at the relationship between declarative protection against algorithmic bias and factual permissiveness regarding state access to data. The Cybersecurity Law and the Data Security Law can reshape the concept of privacy, turning it from an individual right into a surveillance function, while the Personal Information Protection Law, despite establishing formal guarantees such as the right to an explanation of automated decisions, leaves state authorities wide discretion in applying exceptions related to public security.

China’s normative acts in this area are characterized by strong ethical frameworks — such as the Beijing AI Principles, Joint Pledge on Self-Discipline in the AI Industry (Joint Pledge CN 2019) and (RAIP CN 2019) — which, although they arise in the state environment, emphasize the principles of transparency, inclusiveness, security and responsibility. Studies of Chinese academic discourse (Zhu 2022) show that Chinese scholars strongly advocate binding rather than soft regulations, precisely because of the need to limit uncontrolled technological consequences that can affect privacy, freedom of expression and equal treatment of citizens. It is precisely because of these attitudes that the ability to quickly move from ethical to legally binding frameworks has developed. While in the West numerous ethical documents remain non-binding, China has rapidly translated several ethical frameworks into binding technical and administrative regulations, most notably through the *Provisions on the Administration of Algorithmic Recommendation Services* (ARR CN 2021), the *Provisions on the Administration of Deep Synthesis Internet Information Services* (DSR CN 2022), and the *Interim Measures for the Administration of Generative Artificial Intelligence Services* (GAIR CN 2023), which operationalize principles such as transparency, accountability, content responsibility and security.

Although this system is politically problematic, at the same time it provides a certain level of protection of citizens from technological abuses — especially in terms of deepfake manipulations, economic frauds and destabilizing disinformation, through clearly prescribed prohibitions and rules.

Contrary to this opinion, Wu, Huang and Gong (2020) believe that China's ambition to develop a formally "responsible" model of artificial intelligence management in practice shows that ethical principles are often instrumentalized to enable the wide application of surveillance technologies. The development of specific by-laws further demonstrates how China's regulatory model affects human rights. The Regulation on Algorithmic Recommendations from 2021 introduces the obligation to register algorithms in the state database, which enables a detailed insight into the logic of content ranking and assessment of the social and political acceptability of digital material. From critical perspective, this instrument not only allows the state to influence the flow of information, but also to directly direct the behavior patterns of users, shaping „permissible“ patterns of expression and limiting pluralism of opinion. The 2022 Deep Synthesis Regulation is formally presented as an anti-disinformation mechanism, but combined with existing biometric surveillance systems, it strengthens the state's capacity to monitor the identity, movement and digital presence of individuals. We are of the opinion that the obligation to label synthetic content, although technically justified, actually enables improved surveillance of politically sensitive discourses and anonymous actors.

The latest Regulation on Generative Artificial Intelligence of 2023 further narrows the scope of digital freedoms, introducing the obligation that the generated content be "objective", "truthfully presented" and "consistent with basic socialist values". This normative requirement represents an institutionalized form of censorship, as political loyalty is built into an algorithmic architecture, thus limiting freedom of expression and artistic freedom, while pluralism of opinion is suppressed under the pretext of protection from harmful content. This is confirmed by Shen and Liu (2021), who indicate that China is building the so-called „normative ecosystem of responsible AI“, in which ethical standards, technical protocols and regulatory instruments function as a unified support system for the centralized management of digital technologies.

The political logic of the Chinese regulatory model is perhaps most clearly illustrated by Sheehan (2023), who shows that the key by-laws – the Algorithmic Recommendation Regulation, the Deep Synthesis Regulation and the Generative AI Regulation – are not just technical regulation, but part of a politically structured

process in which technologies are aligned with the ideological priorities of the state. According to his interpretation, the Chinese regulation of AI arises within the framework of the so-called policy funnel model, in which regulatory decisions are formed after ideological filtration of technological and social challenges, which ensures that regulation remains in the function of preserving order, and not protecting the individual. From critical approach, this instrument, especially when viewed in connection with the central role of the Cyberspace Administration of China (CAC n.d.), forms the core of the authoritarian surveillance apparatus, as it turns algorithmic transparency into a privilege of the state, while leaving citizens without real mechanisms of legal protection in terms of privacy, freedom of expression and non-discriminatory treatment.

The institutional configuration of the Chinese model may suggest that the regulation of artificial intelligence is not designed as a mechanism for the protection of human rights, but rather as a means of political control and management of society. However, this aspect can also be seen as an advantage of this system, the fact that there is a high degree of technical expertise and integration between the regulator and the industry where the Cyberspace Administration of China (CAC n.d.) and the ministries responsible for technology work closely with technology companies, can be interpreted to ensure a deep understanding of the technical aspects of the system, including algorithm architecture, data management and security protocols (Murnisari 2024, 169). This model allows the state to understand how algorithms work “from the inside” and to establish detailed technical standards. From a human rights perspective, technical sophistication can contribute to better data quality control, greater system security, and prevention of technical abuses, including compromising individual data or leaking sensitive information.

The Chinese model of regulating artificial intelligence shows certain advantages, among which is certainly the rapid normative response to technological innovations, because China adopts regulations on AI much faster than the EU and most Western countries. Such regulatory agility allows the legal system to keep up with innovation in near-real time, preventing long-lasting “regulatory gaps” where technology advances faster than the law. In terms of the protection of human rights, this speed has a certain value: it enables a quick response to new technological risks, such as deepfake manipulation techniques, the misuse of biometrics or the destabilizing effects of viral content. In addition, this system is characterized by developed technical sophistication and the ability to address specific risks sectorally and operationally, its key shortcomings remain

of a structural nature. The biggest shortcomings stem from the lack of independent oversight, broadly defined state exceptions to data access, limited pluralism in shaping regulatory policies, and the systemic dependence of legal protection on political will. We are of the opinion that in this configuration basic human rights – especially the right to privacy, freedom of expression, access to information and protection from discriminatory automated decisions - are in an inferior position. Nevertheless, it is undeniable that the Chinese model has the capacity for rapid adaptation and the technical expertise that can be the basis for improvement.

The Cyberspace Administration is mostly characterized as “an intrusive regime” (Lee 2022, 29). It is a part of the Communist Party of China and, as such, it does not possess institutional independence and functions as a centralized techno-political body. From a human rights perspective, more effective protection would require the introduction of several key mechanisms. Establishment of institutionally independent supervision that would limit the discretion of state authorities in data processing and use of AI systems. Normative strengthening of procedural guarantees, including the right to explain algorithmic decisions, the right to object and the prohibition of algorithmic discrimination, as well as increasing transparency in the work of regulatory bodies, as well as the inclusion of independent academic and civil actors in the processes of drafting regulations. A combination of technical progress, institutional responsibility and normative protection of the individual can provide a model of regulation that is both effective and harmonized with the fundamental values of human rights.

A comparative analysis of the EU and Chinese models of AI regulation: key differences, normative gaps and international legal implications

A comparative analysis of the European and Chinese models of artificial intelligence regulation shows profoundly different starting assumptions, normative goals and consequences for the international legal order. The European Union has built a horizontal, systemic model based on the AI Act, which builds on the GDPR, the EU Charter of Fundamental Rights and the Digital Regulation Platform. This model starts from the assumption that the legitimacy of the use of artificial intelligence is conditioned by the protection of the dignity,

autonomy and fundamental rights of the individual, whereby the risk-based approach and mechanisms like FRIA are used as instruments of preventive control over technological development (Ufert 2020; Pirozzoli 2024). In contrast, the Chinese model is based on a network of laws and regulations—such as the Cybersecurity Law, the Data Security Law, the Personal Information Protection Law, and a series of ordinances on algorithms, deep synthesis, and generative AI—all of which are aligned with the broader political goals of preserving social stability, technological sovereignty, and strengthening the state's capacity for digital surveillance (Roberts et al. 2021, 4–7; Wang et al. 2025). We are of the opinion that the EU starts from the paradigm of “human rights as a limitation of technique”, while China starts from the paradigm of „technique as a means of political management”, which already produces two irreducible normative regimes in terms of goals.

In terms of regulatory structure, the differences between two models are equally pronounced. The European AI Act seeks to encompass all AI systems in a single, horizontal framework, in which prohibited practices are clearly limited, high-risk systems are subject to strict ex ante and ex post control, and rights protection is institutionalized through FRIA and other procedural guarantees. The Chinese model develops vertically: each new technology - algorithmic recommendations, deepfake, generative AI - receives a special regulation, often with the obligation of registration, ex ante censorship or compliance with „basic socialist values”, where the question arises whether the target criterion is political security or the protection of individual rights. Certainly, the expediency of the improvement of legal regulations as well as the harmonization of normative acts with real life circumstances is a great advantage that characterizes the Chinese legal system, however, the question of the effectiveness of those legal norms arises. We can certainly say that in practice the system must create a structure in which algorithmic transparency exists towards the state but also towards citizens, because it must serve as a protective mechanism for every individual.

When examining the impact of these models on human rights, not only the instruments differ, but also the understanding of the very concept of rights. In the EU, the EU Charter of Fundamental Rights functions as a constitutional foundation, the legitimacy of AI systems is measured by the extent to which they respect the rights to privacy, equal treatment, fair procedure and effective legal remedy (EU Charter 2000). Privacy is protected through a combination of the GDPR and the AI Act, whereby data processing in AI systems is subject to the

principles of purposeful limitation, minimization and transparency. Non-discrimination is protected by requirements for data quality and representativeness, model testing and bias mitigation, while the FRIA should identify risks to vulnerable groups and process safeguards before the system is put into operation (Pirozzoli 2024; Cosentini et al. 2024). In the Chinese model, rights such as privacy, freedom of expression and non-discrimination formally exist in certain laws, however their effectiveness and efficiency is reduced by the existence of exceptions, such as provisions where it is possible to deviate from the application of legal norms in the case of issues related to “national security” and “social stability”, which creates the impression that the protection of human rights in China is a variable that can be limited when it comes into conflict with state priorities.

When such divergent models are projected into the international legal sphere, a number of normative gaps arise. On the one hand, the EU is trying to project its human-centric approach globally through the so-called “Brussels effect”, expecting third countries, business entities and international organizations to accept AI Act standards and related rights protections, in order to maintain access to the European market (Palmiotto 2025; Pirozzoli 2024). On the other hand, China promotes its own model through the “export” of technological infrastructure and normative solutions in the countries of the Global South, in which AI systems are used for surveillance, population management and information control, which creates a kind of “Beijing effect” of the normative diffusion of authoritarian digital standards. As a consequence, this situation leads to the segmentation of the global space into spheres of different normative influences, whereby international public law remains without a single, binding regime for the regulation of artificial intelligence. This is the cause of the lack of international public law regulation of artificial intelligence.

Normative gaps are particularly evident in the areas of cross-border data flows, transnational use of surveillance technologies, and liability standards for human rights violations caused by AI systems. At the global level, there are primarily fragmentary initiatives – such as OECD (2019) guidelines or UNESCO (2021) recommendations – which do not have a binding character and do not resolve the conflict between the European model of rights protection and the Chinese model of digital sovereignty (Ufert 2020, 1089–1091). We are of the opinion that this creates a “normative vacuum” in which individuals whose rights are violated by the use of AI systems developed in one regime and implemented in another regime often have no clear legal basis for protection at the

international level, while states use regulatory differences as an instrument of geopolitical competition. This has direct consequences for the universality of human rights, because basic guarantees are increasingly viewed through the prism of regional and state models, instead of as a common, universal standard.

A comparative analysis of the EU and China demonstrates that the international legal response cannot be a simple adoption of one of the existing models, but that it is necessary to look for a minimum of common denominators: binding prohibitions of the most dangerous AI practices, minimum standards of transparency and responsibility, as well as procedural guarantees that would apply regardless of the political order. While the European model offers a value-based and legally elaborated framework that, in addition to all its advantages, also has numerous disadvantages, from the uneven application of the GDPR and different regulatory practices of the member states, through the lack of technical capacity to the slow and complex procedure of adopting regulations, which is why the AI Act, which has been developed for years, in some segments already seems normatively “outdated” in relation to the speed of technological innovations, the Chinese model, despite the formality of ethical principles, the speed of regulatory reaction, but also effective regulatory bodies, is characterized by non-transparency, broadly defined exceptions in favor of the state, which turn the law into an instrument of surveillance rather than individual protection. In this split, international law is currently in a late reaction phase, and the normative vacuum is being filled by bilateral agreements, soft-law instruments and unilateral projections of standards. We believe that this gap between the speed of technological changes and the slowness of the international legal norm-making function is the central challenge of the future development of the regulation of artificial intelligence in the context of the protection of human rights.

Perspectives and the need for a transnational legal framework

A comparative analysis of the European and Chinese models clearly shows that the global management of artificial intelligence is developing in the direction of normative fragmentation: on the one hand, there is a value-based, but regionally limited and institutionally burdened European model, and on the other, a politically centralized and authoritarian Chinese approach, with a strong capacity for technological and regulatory export. The European Union is

struggling to assume the role of a strong global actor. On the one hand, the global role of China is in constant expansion. Numerous countries are depending on Chinese investments and technology. Those circumstances raise concerns that Chinese technological infrastructure could enable data collection and government-led surveillance. Issues of digital sovereignty are also emerging, as some countries fear that dependence on Chinese technology will threaten their ability to control their own digital domains and manage national data. This opens the issue of technological dependency. Countries participating in the development of Chinese technologies could become overly dependent on them, which may limit their future freedom of choice and autonomy, as well as their capacity to consider other options or develop their own technologies, leaving them exposed to the influence and control of China (Stekić 2025, 79).

It is precisely this asymmetry between the “European standard” and the “Chinese standard” produces global normative gaps, because there is no transnational legal framework that would define a minimum set of obligations of states and private actors regarding the protection of human rights in the context of AI.

Previous research confirms this conclusion: Jobin, Ienca and Vayena (2019) show that at the global level there is a convergence around several ethical principles (transparency, fairness, harmlessness, responsibility, privacy), but that there is a deep divergence regarding their concrete implementation and the responsibility of the actors, with most documents remaining in the domain of soft law. Hagerty and Rubinov (2019) additionally point out that the global ethical architecture of AI predominantly reflects the perspectives of the global North and that the social consequences for the marginalized and states of the global South are systematically underestimated. On the international level, although there are significant steps, such as the Framework Convention of the Council of Europe on artificial intelligence, human rights, democracy and the rule of law, that instrument remains regionally limited and cannot by itself fill the global normative vacuum. So far, about 50 countries have signed the Framework Convention, but none have ratified it.

Artificial intelligence has become an important area of interest for the United Nations, particularly regarding its potential to accelerate progress towards the Sustainable Development Goals. It is not possible to mention all the UN’s initiatives for digital technologies and artificial intelligence, so we will mention the most significant ones.

Aware of the impact that new technologies have on all areas of life and functioning, the Secretary-General decided to establish the High-Level Panel on Digital Cooperation (in 2018). Its purpose was to enable stakeholders to agree and cooperate on harnessing the potential of digital technologies, including artificial intelligence, while at the same time addressing the unintended consequences of their application. The next significant step occurred in 2020 when the Secretary-General presented the report: "Road map for digital cooperation: implementation of the recommendations of the High-level Panel on Digital Cooperation." The Secretary-General expressed the intention to establish "a multi-stakeholder advisory body on global artificial intelligence cooperation to provide guidance to the international community on artificial intelligence that is trustworthy, human-rights based, safe and sustainable and promotes peace". Such a body should serve as a diverse forum with the aim of sharing and promote best practices, as well as exchange views on artificial intelligence standardization and compliance efforts, while taking into account existing mandates and institutions (UNGA 2020, 18). The report emphasized a noticeable lack of representation and inclusiveness in global AI discussions. Developing countries are largely absent from or not well-represented in most prominent forums on AI, despite having a significant opportunity to benefit from it for their economic and social development (UNGA 2020, 13).

The United Nations continued their work on AI and digital technology in the newly adopted Pact for the Future, a strategy for the reform and future functioning of the organization. Its integral part is the Global Digital Compact which has the aim to advance equitable and inclusive approaches to harnessing artificial intelligence benefits and mitigating risks in full respect of international law, including international human rights law, and taking into account other relevant frameworks. According to the UN, international governance of AI requires an agile, multidisciplinary and adaptable multi-stakeholder approach. The Global Digital Compact has an opportunity to advance international governance of artificial intelligence in ways that complement international, regional, national, and multi-stakeholder efforts (GDC 2025, 12–14).

As a part of activities envisaged under the Global Digital Compact, in 2025 the UN General Assembly established two new mechanisms to promote international cooperation on the governance of AI- the UN Independent International Scientific Panel on AI and the Global Dialogue on AI Governance. The Global Dialogue on AI governance should serve as a platform to discuss international cooperation, share best practices and lessons learned, and facilitate

open, transparent and inclusive discussions on artificial intelligence governance with a view to enabling AI to contribute to the implementation of the Sustainable Development Goals and to close digital divides between and within countries (UNGA 2025). On the other hand, the Independent International Scientific Panel on Artificial Intelligence aims to provide independent scientific assessments that will help the international community to anticipate emerging challenges and make informed decisions about how to govern this challenging technology.

In light of these initiatives and activities, it is clear that the Organization of the United Nations remains “irreplaceable in coordinating truly global efforts to build a sustainable regulatory framework for artificial intelligence and promote its positive impact in global society” (Dabić 2024, 178). At the same time, however, it is necessary to mention that all of the UN initiatives, including the Global Digital Compact, have not yet resulted in a binding regime. These initiatives remain at the declarative level and do not introduce clear, enforceable obligations of states and corporations. All of the above supports our initial thesis that existing international mechanisms are not sufficient to ensure systemic and universal protection of fundamental rights in the context of rapidly accelerating development of AI technologies.

Based on this, we believe that the future transnational legal regime for artificial intelligence should include at least three normative approaches. The first approach would be a set of universal bans on certain AI practices that are incompatible with international human rights standards, such as mass biometric surveillance without judicial oversight, social lending and systematic political profiling of citizens. The second approach would be procedural standards, inspired by the EU experience, which would oblige states and private actors to ex ante assess the impact of AI systems on human rights, the transparency of key algorithmic decisions and the existence of effective legal protection mechanisms, including the right to explanation, objection and access to court. The third approach would be the establishment of an international monitoring mechanism, which could combine elements of existing contractual control in the field of human rights with specific technical expertise, with a global mandate. In an ideal case, it would be an institution functioning within the broader UN framework. Having in mind current global geopolitical situation, establishment of such an institution could happen in next decades. Such an institution could face many functional challenges.

Existing research on the global governance of AI indicates that a human rights approach, which starts from the Universal Declaration of Human Rights and the

Covenant on Civil and Political Rights, is a necessary condition for the legitimacy of the future regime, but that it is not sufficient in itself if it is not connected to concrete regulatory instruments and accountability mechanisms (e.g. Jobin, Ienca and Vayena 2019; Hagerty and Rubinov 2019). In this sense, our recommendations go in the direction of global harmonization of human rights protection standards through a combination of: a minimum binding international treaty on AI based on human rights, mutual harmonization of regional regimes (such as the AI Act and future Asian or American standards) and the institutionalized participation of various actors – states, international organizations, the private sector, the academic community and civil society – in shaping and monitoring the application of those standards. We believe that such a transnational framework represents the only realistic perspective to limit the negative effects of normative competition between the EU and China and to build a minimum “*ius commune*” of human rights protection in the era of artificial intelligence.

Concluding considerations

Starting from the hypothesis that the simultaneous existence of a value-based but regionally limited European model of artificial intelligence regulation and a politically centralized, authoritarian Chinese model, with the absence of a coherent international legal regime, produces a global “normative vacuum” in the protection of human rights, the authors believe that the conducted analysis confirms this assumption. We have shown that the EU, through the AI Act, the GDPR and the EU Charter of Fundamental Rights, is building a developed horizontal framework for the protection of the individual, in which privacy, non-discrimination, dignity and process guarantees are integrated into the very structure of the regulation of the AI system, it has many shortcomings that reduce the effectiveness of the norm as well as the mechanisms for the protection of basic human rights. At the same time, the Chinese model, based on a combination of systemic laws on cyber security, data and personal information, and specific regulations on algorithmic recommendations, deep synthesis and generative AI, is developing a highly sophisticated but authoritarian surveillance architecture, in which human rights are subordinated to the priorities of political stability, social governance and technological sovereignty. The analysis showed that both models produce extraterritorial effects – European through the

“Brussels effect” and normative export of standards, Chinese through technological expansion and regulatory influence – but that none of them, by itself, succeeds in providing a universal and transnationally accepted minimum protection of human rights in the sphere of AI.

The key findings of the paper can be summarized in several points. The existing international regulation is fragmented, predominantly “soft” in nature and does not contain binding standards that would address the specific risks of artificial intelligence to human rights. The European model represents the most advanced attempt to integrate human rights into technical regulation, but it is burdened by political compromises, uneven implementation and limited geographical scope. The Chinese model demonstrates exceptional normative and technical adaptability, but is essentially focused on the consolidation of state power, which generates serious risks for the rights to privacy, freedom of expression, non-discrimination and fair treatment. Existing empirical and theoretical research on the global ethics of AI confirms our assessment that a transnational legal framework is needed that would translate the normative principles of human rights into binding standards and operational protection mechanisms.

On this basis, the directions of future research and international initiatives should develop in three main directions. First, it is necessary to follow in detail the implementation of the European and Chinese models in practice, including their extraterritorial effects, in order to identify the real consequences for human rights in different social contexts. Second, it is necessary to further develop theoretical and normative models of transnational artificial intelligence management, with special reference to the role of international organizations, regional human rights courts and new contractual regimes. Third, future research must include the perspectives of countries in the Global South, marginalized groups, and different cultural contexts, in order to avoid reproducing existing inequalities in the global order. Ultimately, only an integrated approach, which connects comparative law, international human rights law and science and technology studies, can offer an adequate answer to the challenges that artificial intelligence poses to the modern international legal order.

It will be particularly important to observe further ideas and initiatives of the United Nations regarding artificial intelligence. It is necessary to point out that, so far, the United Nations is focused on the relation between artificial intelligence and the fulfillment of the Sustainable Development Goals. The issue of a transnational legal framework regarding AI is not the main focus of the United Nations. At the same time, in the era of extremely challenging geopolitical

tensions, major forces are not encouraging collective efforts to establish a clear legal framework for governing artificial intelligence. Quite the opposite, they are more suspicious of the intentions of other states and global actors, considering them as competitors in a geopolitical game. In the future, the perception of significant geopolitical benefits from the rapid development of artificial intelligence can be considered as a major obstacle in reaching an agreement on a transnational legal framework on artificial intelligence.

Управљање вештачком интелигенцијом и заштита људских права: упоредноправна перспектива Европске уније и Кине

Маида БЕЋИРОВИЋ-АЛИЋ¹, Јелица ГОРДАНИЋ²

Апстракт: Глобално ширење аутоматизованих система, који често функционишу нетранспарентно и изван ефективне људске контроле, отвара сложена питања у вези са заштитом основних људских права, што представља шири контекст овог истраживања. Посебни изазови настају у домену међународног права, с обзиром на то да не постоји јединствен, обавезујући транснационални нормативни оквир који би обезбедио доследну заштиту људских права у дигиталном окружењу. Овај рад, применом упоредноправне методе и анализом садржаја релевантних докумената, испитује регулаторне modele Европске уније и Народне Републике Кине као два глобално најутицајнија модела управљања вештачком интелигенцијом, у којима се јасно огледају либерални и нелиберални концепти заштите људских права, указујући на њихов правно-политички оквир, регулаторне домете и ограничења. Посебан акценат стављен је на идентификацију нормативних празнина и ризика по људска права, укључујући приватност, недискриминацију, слободу изражавања и правично поступање. Полазна хипотеза рада јесте да постојећи међународноправни оквир није у стању да одговори на системске изазове које генерише примена вештачке интелигенције, те да дивергентни модели ЕУ и Кине продубљују глобални нормативни јаз и онемогућавају успостављање универзалних стандарда заштите људских права. Кроз критички приказ разлика између западног либералног модела и ауторитарнијег приступа управљању технологијом, рад указује на ургентну потребу за успостављањем глобално усклађеног и политички одрживог оквира међународног права за регулисање система вештачке интелигенције.

Кључне речи: вештачка интелигенција, људска права, регулаторни модели, међународно право, ЕУ, Кина.

¹ Редовни професор, Универзитет у Новом Пазару – Правни факултет, Нови Пазар, Србија. Е-пошта: maida.becirovic@uninp.edu.rs, ORCID: <https://orcid.org/0000-0002-9738-6581>.

² Виши научни сарадник, Институт за међународну политику и привреду, Београд, Србија. Е-пошта: jelica@diplomacy.bg.ac.rs, ORCID: <https://orcid.org/0000-0001-8364-7405>.

Увод

У контексту глобалног такмичења за технолошку и нормативну доминацију, различити политички системи обликују различите приступе регулисању вештачке интелигенције. Ово има далекосежне последице по остваривање и заштиту људских права. Европска унија се позиционирала као глобални регулаторни лидер стварајући први правни оквир за вештачку интелигенцију. Тај оквир је свеобухватан, амбициозан и вредносно утемељен. ЕУ полази од премисе да технолошки напредак служи људима и друштву, па вештачку интелигенцију не третира само као економски инструмент, него и као питање од уставног значаја чије постојање има утицај на области демократије, владавине права и људског достојанства. Овај приступ се ослања на препознатљив метод ЕУ који комбинује интеграцију тржишта, заштиту права и процедурално управљање — чиме се потврђује европско регулаторно лидерство у дигиталном добу (Hulok 2025). Насупрот томе, Кина развија секторски фрагментисан, али политички централизован модел регулисања вештачке интелигенције који је дубоко повезан са стратегијама друштвеног надзора и технолошког управљања. Флориди (Floridi 2018, 2) наводи да брзина технолошких иновација отежава успостављање адекватног нормативног оквира, због чега постојећи правни и етички механизми често не успевају благовремено да одговоре на изазове дигиталног развоја. Слично томе, Кало (Calo 2017, 408) истиче да је постојећи оквир за управљање вештачком интелигенцијом још увек неразвијен, посебно због ослањања на етичке стандарде, а не на обавезујућа правна правила. Ови приступи сугеришу да се развој вештачке интелигенције одвија у нормативном вакууму, јер систему међународног права још увек недостају јединствене норме за управљање ризицима које са собом носи вештачка интелигенција. Рад полази од хипотезе да тренутни међународноправни оквир није у стању да адекватно одговори на системске изазове који настају применом вештачке интелигенције и да различити регулаторни модели Европске уније и Кине продубљују глобални нормативни јаз и ометају успостављање универзалних стандарда заштите људских права.

Кроз упоредноправну анализу ова два модела, рад има за циљ да идентификује кључне нормативне празнине у међународном праву, укаже на ризике по људска права (по право на приватност, недискриминацију и слободу изражавања), и покаже да одсуство јединственог транснационалног

оквира омогућава настанак глобалног „нормативног вакуума“. У закључном делу, разматра се потреба за политички одрживим и универзално кохерентним моделом глобалног управљања вештачком интелигенцијом који би обезбедио јасну и вредносно утемељену заштиту појединца.

Методолошки, рад се заснива на комбинацији упоредноправних, дескриптивних и критичко-аналитичких метода. Ауторке користе квалитативни приступ у испитивању релевантних правних аката, стратешких докумената и секундарне литературе. Упоредноправна анализа омогућава јасно уочавање разлика између европског либералног модела фокусираног на заштиту људских права и кинеског централизованог модела технолошког управљања. Критичка анализа открива последице таквих разлика по међународноправни поредак и глобалне стандарде.

Упоредноправна литература често истиче уочљив контраст између европских и кинеских модела регулисања вештачке интелигенције. Међутим, важно је избећи дихотомију по овом питању. Упркос чињеници да је нормативни оквир Европске уније заснован на заштити достојанства, приватности и аутономије појединца - што јасно потврђују анализе европског права о заштити података које препознају управо те вредности као фундаменталне конститутивне принципе правног поретка ЕУ (Huster and Walden 2019; Lynskey 2015) – овај модел није без ограничења. Тако нпр. одређене праксе биометријског надзора нису у потпуности забрањене, државе чланице користе широке изузетке у областима полицијских овлашћења и безбедности, спровођење Опште уредбе о заштити података остаје неуједначено, док појаве технолошког патернализма и тржишних монопола указују на структурне слабости. Савремени дигитални екосистеми, посебно они засновани на напредним системима вештачке интелигенције, додатно истичу ове изазове јер стварни обим заштите достојанства, аутономије и приватности зависи од способности регулаторног система да одговори на брзо растуће технолошке ризике. У случају ЕУ, ово је проблематично упркос вредносном приступу на којем је изграђен европски модел. С друге стране, иако кинески модел карактерише политичка централизација и снажан фокус на друштвену стабилност, не може се окарактерисати као стриктно хомоген. Кинеска академска заједница показује одређени степен плурализма. Тако нпр. Џу (Zhu 2022) документује забринутост многих кинеских стручњака у вези прекомерне употребе технологија надзора – док се у пракси истовремено развијају иницијативе које промовишу оквир „одговорне вештачке интелигенције“

попут Пекиншких принципа вештачке интелигенције (Beijing AI Principles 2019). Поред тога, одређени сегменти кинеског законодавства настоје да ограниче корпоративне злоупотребе и побољшају безбедност података, што потврђује нови режим управљања подацима успостављен Законом о заштити личних података и Законом о безбедности података. Ови акти уводе „строге обавезе за приватне актере са циљем сузбијања прекомерног прикупљања података, обезбеђивања адекватне обраде и јачања укупних механизма безбедности података“ (Creemers 2022).

Ауторке сматрају да поларизација између ЕУ и Кине свакако постоји, али да није апсолутна. Оба модела садрже унутрашње тензије, контрадикције и еволутивне процесе који отежавају њихово представљање као супротстављених концепата људских права у дигиталној ери. Сходно томе, фокус је на упоредном испитивању модела ЕУ и Кине како би се указало у којој мери њихове структурне разлике доприносе постојећим нормативним празнинама у међународном праву и како би се процениле импликације таквог одступања за будући развој транснационалног оквира за управљање вештачком интелигенцијом.

Европска унија као регулаторни лидер у области вештачке интелигенције

Правни оквир Европске уније за регулисање вештачке интелигенције најразвијенији је хоризонтални модел у савременом међународном праву. Заснован је на комбинацији Опште уредбе о заштити података (GDPR 2016), Повеље ЕУ о основним правима (EU Charter 2000), Закона о дигиталној регулацији платформи (DSA 2022), Закона о дигиталним тржиштима (DMA 2022) и, коначно, првог глобалног свеобухватног закона о вештачкој интелигенцији – Закона о вештачкој интелигенцији ЕУ (AI Act 2024). Поред тога, Европска комисија је основала Европску канцеларију за вештачку интелигенцију 24. јануара 2024. године. Обавезујућа правна снага Закона о вештачкој интелигенцији ЕУ произилази из његове правне природе као прописа ЕУ који је директно применљив и обавезујући за сва лица и субјекте који развијају, користе, увозе или стављају на тржиште системе вештачке интелигенције унутар Уније. У овом оквиру, Европска канцеларија за вештачку интелигенцију основана је као специјализовано тело са задатком праћења, координације и подршке спровођењу Закона о

вештачкој интелигенцији ЕУ, посебно у погледу система вештачке интелигенције опште намене. Очекује се да ће играти значајну техничку и институционалну улогу у настанку европског система управљања вештачком интелигенцијом (Dabić 2025, 235).

Оваква нормативна архитектура одражава амбицију ЕУ да се позиционира као европски и глобални заштитник основних права у дигиталној ери, што Палмиото описује као „уставну оријентацију целокупне политике вештачке интелигенције“ (Palmioto 2025, 770). ЕУ не третира вештачку интелигенцију само као технолошки изазов, него као правни проблем који захтева примену стандарда достојанства, аутономије и заштите појединца.

Закон о вештачкој интелигенцији ЕУ уводи регулаторну парадигму засновану на процени ризика, према којој се обавезе и ограничења не примењују подједнако на све системе вештачке интелигенције. Постоје разлике према степену претње по безбедност и основна права. Дакле, законска дозвољеност развоја технологије условљена је мерљивим и превентивним гаранцијама, које Пироцоли одређује као најзначајнији институционални корак ЕУ у управљању технолошким ризицима (Pirozzoli 2024, 106). Овај приступ се надовезује на постојеће прописе у којима GDPR остаје темељ заштите приватности и контроле података, при чему је, према Уферту (Ufert 2020, 1089–1091), способан да функционише као „прилагодљиви и ометајући инструмент“ у односу на ризичне облике аутоматизованог доношења одлука, иако не пружа довољне механизме разумљивости и транспарентности за сложене modele вештачке интелигенције. Повеља ЕУ о основним правима пружа уставна ограничења за употребу вештачке интелигенције, посебно у питањима приватности, недискриминације, равноправног третмана и процедуралних гаранција. На овај начин штити основне слободе без обзира на технички развој (Cosentini et al. 2024, 2–3).

У оквиру Закона о вештачкој интелигенцији ЕУ уводи се вишеслојна категоризација која укључује забрањене праксе, системе високог ризика са строгим обавезама, системе ограниченог ризика и системе минималног ризика (AI Act 2024). Забрањене праксе обухватају друштвено рангирање, манипулацију људским понашањем и одређене облике биометријског надзора. Сматрају се неспојивим са достојанством и аутономијом појединца (Pirozzoli 2024, 106). Најважнији и најсложенији сегмент је регулисање система високог ризика, за које се уводе строги захтеви у

погледу управљања подацима, оперативне поузданости, надзора, транспарентности, безбедности и документовања процеса доношења одлука (AI Act 2024). У овом сегменту, ЕУ уводи највећу новину- процену утицаја на основна права (FRIA), као превентивни инструмент који омогућава процену ризика по права појединца чак и пре примене технологије. Косентини и сарадници (Cosentini et al. 2024, 2) називају FRIA „првим правно обавезујућим механизмом за антиципаторну заштиту основних права у области вештачке интелигенције“. Бертаина и сарадници (Bertaina et al. 2025, 2–3) допуњују овај оквир развојем квантитативне процене ризика, која омогућава утврђивање озбиљности и вероватноће кршења сваког појединачног права из Повеље ЕУ – укључујући приватност, једнакост, заштиту деце, слободу изражавања и право на правни лек. Тиме, по први пут, ЕУ ствара услове за заштиту основних права на начин који омогућава мерљивост, проверљивост и контролу.

Утицај европског модела на заштиту људских права највише се види у области приватности и заштите података. GDPR, као основни стуб, пружа широк спектар права за субјекте података. ЕУ је, кроз Закон о вештачкој интелигенцији, додатно ојачала механизме транспарентности, разумљивости и управљања подацима за системе вештачке интелигенције високог ризика (GDPR 2016; AI Act 2024). Ово је посебно важно када је реч о апликацијама вештачке интелигенције које укључују праћење понашања, биометријску идентификацију или профилисање, области у којима је ризик од кршења приватности највећи. С друге стране, Уферт (Ufert 2020, 1090–1091) упозорава да GDPR не решава проблем алгоритарске нетранспарентности, због чега је потребно допунити га секторским регулаторним обавезама. Баш ту нормативну празнину Закон о вештачкој интелигенцији ЕУ (AI Act) настоји да попуни, уводећи обавезне захтеве у погледу транспарентности, техничке документације, надзора и разумљивости за системе вештачке интелигенције високог ризика. На овај начин реализује принципе заштите података у специфичном контексту аутоматизованог доношења одлука (AI Act 2024).

Још једно кључно подручје заштите односи се на недискриминацију и правичност. Системи вештачке интелигенције могу појачати постојеће друштвене неједнакости, па ЕУ уводи обавезне мере за ублажавање пристрасности, стандарде квалитета података, обавезу тестирања и праћења ефеката модела на рањиве групе, као и услов људског надзора над осетљивим одлукама. Пироцоли (Pirozzoli 2024, 107–108) истиче да се

европски приступ усмерен на човека заснива на спречавању структурне дискриминације кроз *ex ante* нормативне интервенције, а не само кроз накнадне исправке.

Трећа област – заштита процедуралних права и правичан третман – најјасније се види у регулисању алгоритамског одлучивања у јавним службама, полицијским системима, запошљавању, социјалним давањима и правосуђу. Палмиото (Palmiotto 2025, 772–773) указује да је током усвајања Закон о вештачкој интелигенцији ЕУ прошао кроз „ролеркостер“ промена, при чему је заштита основних права, упркос политичким компромисима, остала један од његових носећих стубова. Обавезна FRIA процедура захтева од власника вештачких система да унапред процене могућност неједнаких исхода, необјашњивих одлука или кршења права на правни лек, чиме се процедурална правичност уграђује у саму структуру развоја вештачких технологија.

Коначно, европски модел регулисања вештачке интелигенције представља јединствен покушај у савременом праву да се технолошки напредак усклади са чврсто утемељеном заштитом људских права. Увођењем превентивних механизма попут FRIA, стандарда објашњивости, строге контроле података и забрањених пракси, ЕУ је створила нормативни оквир који поставља јасне границе за технолошке иновације у циљу заштите достојанства, приватности, недискриминације и правичног третмана. Овај модел можда није у потпуности савршен, али несумњиво представља најнапреднији покушај да се принципи људских права уграђени у европско правно наслеђе преточе у систем управљања вештачком интелигенцијом.

Треба напоменути да структурне слабости које карактеришу овај систем ограничавају његову ефикасност. Палмиото (Palmiotto 2025, 772–773) указује да је Закон о вештачкој интелигенцији ЕУ производ политичких компромиса који су током процеса усвајања довели до разводњавања неких обавеза и задржавања регулаторне флексибилности. Ово може ослабити заштиту појединца у пракси. Процес доношења прописа у ЕУ је спор, што значи да правна регулација често заостаје за технолошким променама, остављајући дуготрајне правне празнине тамо где се системи вештачке интелигенције примењују без адекватног надзора. Додатни проблем представља и неуједначена примена регулаторних обавеза у државама чланицама. Неке државе увелико развијају и примењују технолошки напредне системе надзора, укључујући биометријску идентификацију и масовну обраду података у контексту полицијског рада

и граничне контроле. С друге стране, неким чланицама недостају одговарајући технички и институционални капацитети, што доводи до значајних разлика у стварном нивоу заштите људских права у ЕУ. Уферт (Ufert 2020, 1090–1091) истиче да, упркос снажним гаранцијама приватности, GDPR и даље не решава проблем нетранспарентности сложених модела високе аутономије, док Пироцоли (Pirozzoli 2024, 107–108) упозорава да хомоцентрични приступ често остаје декларативан када се суочи са економским интересима тржишта и политичким интересима држава чланица. Недостатак техничких капацитета надзорних органа и ограничени ресурси слабе способност ЕУ да ефикасно надгледа велике технолошке компаније. Фрагментација секторских прописа (DSA, DMA, Data Act, AI Act) ствара нормативни неред која отежава практичну примену. Ови недостаци могу делимично бити ублажити деловањем Европске канцеларије за вештачку интелигенцију, која треба да развија техничке капацитете, врши координацију националних надзорних тела и ствара услове за примену Закона о вештачкој интелигенцији ЕУ, посебно у погледу система вештачке интелигенције опште намене. Коначна ефикасност Европске канцеларије за вештачку интелигенцију ће највише зависити од њених формалних овлашћења, ресурса и квалитета институционалне сарадње са националним властима.

Све наведене слабости имају директне последице по остваривање људских права у ЕУ. Регулаторна спорост значи изложеност грађана ризицима у погледу приватности, недискриминације и аутономије. Неуједначена примена угрожава принцип једнакости. Тако нпр. особа у једној држави може имати знатно слабије процедуралне гаранције него у другој. Нејасноће око биометријског надзора отварају пут прекомерном прикупљању осетљивих података, што угрожава право на приватност и може довести до нормализације масовног надзора. Слабо развијени механизми објашњивости и даље омогућавају нетранспарентне алгоритамске одлуке које имају утицај на запошљавање, кредите, социјална права или полицијске поступке. Овим се доводи у питање право на образложену одлуку, једнакост и ефикасан правни лек. На крају, недовољан капацитет надзорних тела омогућава великим платформама да заобилазе прописе, што угрожава заштиту потрошача, процес информисања и слободу изражавања. Слабости европског система ни у ком случају не поткопавају његову нормативну вредност, али у значајној мери отежавају практичну заштиту људских права.

Ауторке сматрају да је додатни изазов неуједначена примена регулаторних обавеза у државама чланицама, као и висока административна оптерећења за мала и средња предузећа. Упркос томе, снага модела ЕУ огледа се у чињеници да је једини на глобалном нивоу који повезује заштиту људских права и технолошку регулацију кроз процену утицаја на основна права за системе вештачке интелигенције високог ризика (FRIA), забрану најопаснијих пракси и јасно дефинисане гаранције за системе високог ризика (AI Act 2024). Унапређење овог модела захтевало би јачање обавезних захтева објашњивости у погледу алгоритаМСког доношења одлука како би се елиминисали постојећи недостаци транспарентности и обезбедило јачање институционалне улоге и оперативних капацитета Европске канцеларије за вештачку интелигенцију. На овај начин би се осигурало доследно тумачење и примена регулативе вештачке интелигенције у свим државама чланицама и проширење механизма FRIA ка сложеним генеративним моделима вештачке интелигенције, праћеним јасним стандардима за процену ризика по слободу изражавања, приватност и недискриминацију.

Европа показује иницијативу и лидерство када је у питању регулисање вештачке интелигенције. Поред ЕУ, неопходно поменути и Савет Европе и његову Оквирну конвенцију о вештачкој интелигенцији и људским правима, демократији и владавини права. То је први међународноправно обавезујући уговор у области вештачке интелигенције. Одредбе поменуте Конвенције имају за циљ да активности система вештачке интелигенције буду у потпуности у складу са људским правима, демократијом и владавином права (Council of Europe 2024). Конвенција је отворена за потписивање од септембра 2024. године. До сада ју је потписало више од 50 држава, укључујући и чланице ЕУ.

Кинески модел регулисања вештачке интелигенције

Кинески модел регулисања вештачке интелигенције почива на стратешкој комбинацији нормативне фрагментације, политичке централизације и технолошке инструментализације. Ово ствара веома прилагодљив, али истовремено дубоко ауторитаран систем управљања дигиталним простором. За разлику од Европске уније која уводи хоризонтални приступ заштити основних права у области вештачке

интелигенције, Кина развија вертикални модел састављен од специфичних закона и подзаконских аката, од којих су најважнији Закон о сајбер безбедности (CSL CN 2017), Закон о безбедности података (DSL CN 2021) и Закон о заштити личних података (PIPL CN 2021). Ови закони чине основни правни оквир за управљање подацима у Кини и, индиректно, за рад система вештачке интелигенције, на начин упоредив са GDPR у Европској унији. GDPR није усвојен као инструмент искључиво за вештачку интелигенцију, већ као систем за регулисање технологија које интензивно користе податке, укључујући вештачку интелигенцију. Ова прецизност омогућава боље управљање ризицима својственим одређеним секторима. Тако нпр. Уредба о дубокој синтези (DSR CN 2022) експлицитно уводи обавезу означавања манипулисаног визуелног садржаја, што је важан алат у борби против дезинформација, заштити јавног интереса и смањењу штете по грађане. У том смислу, кинески модел пружа транспарентност корисницима (нпр. означавање „синтетичког садржаја“), што је у складу са заштитом права на информације и смањењем манипулативног садржаја. Међутим, у нормативном смислу, изузеци у корист државе су широко дефинисани, што у пракси омогућава неограничен приступ подацима у сврху „националне безбедности“ и „друштвене стабилности“. Овакво законодавно решење чини заштиту људских права зависном од политичке процене државних власти, отварајући простор за различите врсте дигиталног надзора.

У овом контексту, Робертс и сарадници (Roberts et al. 2021, 4–7) показују да кинески систем вештачке интелигенције, иако споља делује кохерентно, суштински представља широку, али стратешки недоследну комбинацију етичких принципа, бирократских смерница и политичких директива. Формални позив на одговорност и транспарентност се углавном користи као легитимизација већ успостављеног система контроле. Овај нормативни дуализам је посебно уочљив када се посматра однос између декларативне заштите од алгоритамске пристрасности и фактичке толерантности у случају приступа државе подацима. Закон о сајбер безбедности и Закон о безбедности података могу преобликовати концепт приватности, претварајући га из индивидуалног права у функцију надзора, док Закон о заштити личних података, упркос успостављању формалних гаранција као што је право на објашњење аутоматизованих одлука, оставља државним органима широку листу изузетака везаних за јавну безбедност.

Кинеске нормативне акте у овој области карактеришу снажни етички оквири — као што су Пекиншки принципи вештачке интелигенције, Заједничка обавеза о самодисциплини у индустрији вештачке интелигенције (Joint Pledge CN 2019) и (RAIP CN 2019) — који, иако настају у државном окружењу, наглашавају принципе транспарентности, инклузивности, безбедности и одговорности. Студије кинеског академског дискурса (Zhu 2022) показују да кинески научници снажно заговарају обавезујуће, а не меке прописе, управо због потребе да се ограниче неконтролисане технолошке последице које могу утицати на приватност, слободу изражавања и једнак третман грађана. Управо због ових ставова развила се способност брзог преласка са етичких на правно обавезујуће оквири. Док на Западу бројни етички документи остају необавезујући, Кина је брзо превела неколико етичких оквира у обавезујуће техничке и административне прописе. У овом контексту требамо поменути Одредбе о администрацији услуга алгоритамских препорука (ARR CN 2021), Одредбе о администрацији интернет информационих услуга дубоке синтезе (DSR CN 2022) и Привремене мере за администрацију генеративних услуга вештачке интелигенције (GAIR CN 2023), које операционализују принципе као што су транспарентност, одговорност, одговорност за садржај и безбедност. Иако је овај систем политички проблематичан, пружа и одређени ниво заштите грађана од технолошких злоупотреба — посебно у смислу манипулација дипфејком, економских превара и дестабилизујућих дезинформација, кроз јасно прописане забране и правила.

Супротно овом мишљењу, Ву, Хуанг и Гонг (Wu, Huang and Gong 2020) сматрају да амбиција Кине да развије формално „одговоран“ модел управљања вештачком интелигенцијом у пракси показује да се етички принципи често инструментализују како би се омогућила широка примена технологија надзора. Развој специфичних подзаконских аката додатно показује како кинески регулаторни модел утиче на људска права. Уредба о алгоритамским препорукама из 2021. године уводи обавезу регистрације алгоритама у државној бази података, што омогућава детаљан увид у процес рангирања садржаја и процену друштвене и политичке прихватљивости дигиталног материјала. Овај инструмент не само да омогућава држави да утиче на ток информација, већ и да директно усмерава обрасце понашања корисника, обликујући „дозвољене“ обрасце изражавања и ограничавајући плурализам мишљења. Уредба о дубокој синтези из 2022. године формално је представљена као механизам против дезинформација, али у комбинацији са постојећим системима

биометријског надзора, јача капацитет државе да прати идентитет, кретање и дигитално присуство појединаца. Ауторке овог рада сматрају да обавеза означавања синтетичког садржаја, иако технички оправдана, омогућава побољшани надзор политички осетљивих дискурса.

Уредба о генеративној вештачкој интелигенцији из 2023. године додатно сужава обим дигиталних слобода, уводећи обавезу да генерисани садржај буде „објективан“, „истинито представљен“ и „у складу са основним социјалистичким вредностима“. Овај нормативни захтев представља институционализовани облик цензуре, јер је политичка лојалност уграђена у алгоритамску архитектуру. На овај начин се ограничава слобода изражавања и уметничка слобода, а плурализам мишљења ограничава под изговором заштите од штетног садржаја. То потврђују Шен и Лиу (Shen and Liu 2021), који указују да Кина гради такозвани „нормативни екосистем одговорне вештачке интелигенције“, у коме етички стандарди, технички протоколи и регулаторни инструменти функционишу као јединствени систем подршке за централизовано управљање дигиталним технологијама.

Најјаснију илустрацију политичке логике кинеског модела пружа Шихан (Sheehan 2023), који показује да кључни подзаконски акти – Уредба о алгоритамским препорукама, Уредба о дубокој синтези и Уредба о генеративној вештачкој интелигенцији – нису само техничка регулација, већ део политичког процеса у коме су технологије усклађене са идеолошким приоритетима државе. Према његовом тумачењу, кинеска регулација вештачке интелигенције настаје у оквиру тзв. модела политичког левка, у коме се регулаторне одлуке формирају након идеолошке филтрације технолошких и друштвених изазова, што осигурава да регулација остаје у функцији очувања поретка, а не заштите појединца. Мишљења смо да овај инструмент, посебно када се посматра у вези са централном улогом Кинеске администрације за сајбер простор (CAC n.d.), чини језгро ауторитарног апарата за надзор. Алгоритамска транспарентност постаје привилегија државе, а грађани остају без стварних механизма заштите у погледу приватности, слободе изражавања и недискриминаторних поступака.

Кинески модел се може тумачити на начин да регулисање вештачке интелигенције није осмишљено као механизам за заштиту људских права, већ као средство политичке контроле и управљања друштвом. С друге стране, овај аспект се може посматрати и као предност овог система. Чињеница је да постоји висок степен техничке стручности и интеграције

између регулатора и индустрије. Кинеска администрација за сајбер простор (CAC n.d.) и министарства надлежна за технологију тесно сарађују са технолошким компанијама, што се може тумачити као постојање доброг разумевања техничких аспеката система, укључујући архитектуру алгоритама, управљање подацима и безбедносне протоколе (Murnisari 2024, 169). Овај модел омогућава држави да разуме како алгоритми функционишу „изнутра“ и да успостави одговарајуће техничке стандарде. Са становишта људских права, техничка софистицираност може допринети бољој контроли квалитета података, већој безбедности система и спречавању техничких злоупотреба, укључујући угрожавање индивидуалних података или цурење осетљивих информација.

Кинески модел регулисања вештачке интелигенције показује одређене предности, међу којима је свакако брз нормативни одговор на технолошке иновације. Кина усваја прописе о вештачкој интелигенцији много брже од ЕУ и већине западних држава. Таква брзина омогућава правном систему да прати иновације у скоро реалном времену, спречавајући дуготрајне „регулаторне празнине“ где технологија напредује брже од закона. У погледу заштите људских права, ова брзина има одређену вредност: омогућава брз одговор на нове технолошке ризике попут манипулација дипфејком, злоупотребе биометрије или дестабилизујућих ефеката виралног садржаја. Поред тога, овај систем карактерише развијена техничка софистицираност и способност да се секторски и оперативно бави специфичним ризицима. Његови кључни недостаци остају структурне природе и произилазе из недостатка независног надзора, широко дефинисаних државних изузетака од приступа подацима, ограниченог плурализма у обликовању регулаторних политика и системске зависности правне заштите од политичке воље. Ауторке сматрају да су у овој конфигурацији основна људска права – посебно право на приватност, слободу изражавања, приступ информацијама и заштиту од дискриминаторних аутоматизованих одлука – у инфериорном положају. Поред свих мањкавости, неспорно је да кинески модел поседује капацитет за брзе промене, као и техничку стручност. Ови фактори могу бити основа за побољшање.

Управа за сајбер простор се углавном карактерише као „наметљив режим“ (Lee 2022, 29). Део је Комунистичке партије Кине и, као таква, нема сопствену независност. У питању је централизовано техно-политичко тело. Верујемо да би ефикасна заштита људских права захтевала увођење

неколико кључних механизма. Као прво, успостављање институционално независног надзора који би ограничио дискреционо право државних органа у обради података и коришћењу система вештачке интелигенције. Затим, нормативно јачање процедуралних гаранција, укључујући право на објашњење алгоритамих одлука, право на приговор и забрану алгоритамих дискриминације. Коначно, неопходно је и повећање транспарентности у раду регулаторних тела, као и укључивање независних академских и недржавних актера у процесе израде прописа. Комбинација техничког напретка, институционалне одговорности и нормативне заштите појединца може пружити модел регулације који је и ефикасан и усклађен са најважнијим тековинама људских права.

Упоредноправна анализа модела регулације вештачке интелигенције у ЕУ и Кини: кључне разлике, нормативне празнине и међународноправне импликације

Упоредноправна анализа европског и кинеског модела регулисања вештачке интелигенције показује различите почетне претпоставке, нормативне циљеве и последице по међународноравни поредак. Европска унија је изградила хоризонтални, системски модел заснован на Закону о вештачкој интелигенцији, који се надовезује на Општу уредбу о заштити података о заштити података (GDPR), Повељу ЕУ о основним правима и Платформу за дигиталну регулацију. Овај модел полази од претпоставке да је легитимитет употребе вештачке интелигенције условљен заштитом достојанства, аутономије и основних права појединца, при чему се приступ заснован на ризику и механизми попут FRIA користе као инструменти превентивне контроле над технолошким развојем (Ufert 2020; Pirozzoli 2024). Насупрот томе, кинески модел се заснива на мрежи закона и прописа – као што су Закон о сајбер безбедности, Закон о безбедности података, Закон о заштити личних података, и низа уредби о алгоритмима, дубокој синтези и генеративној вештачкој интелигенцији. Сви ови акти су усклађени са ширим политичким циљевима очувања друштвене стабилности, технолошког суверенитета и јачања капацитета државе за дигитални надзор (Roberts et al. 2021, 4–7; Wang et al. 2025). Мишљења смо да ЕУ полази од парадигме „људских права као ограничења технике“, док Кина

полази од парадигме „технике као средства политичког управљања“. То су два потпуно другачија нормативна режима у смислу циљева.

Разлике су присутне и на нивоу правне технике. Европски Закон о вештачкој интелигенцији тежи да обухвати све системе вештачке интелигенције у јединственом, хоризонталном оквиру, у којем су забрањене праксе јасно ограничене, а системи високог ризика подлежу строгој *ex ante* и *ex post* контроли, док су права заштићена кроз FRIA и остале процедуралне гаранције. Кинески модел се развија вертикално: свака нова технологија – алгоритамске препоруке, дипфејк, генеративна вештачка интелигенција – добија посебну регулативу, често уз обавезу регистрације, претходне цензуре или поштовања „основних социјалистичких вредности“. Остаје простора за питање да ли је крајњи циљ политичка безбедност или заштита индивидуалних права. Свакако, сврсисходност унапређења законских прописа, као и усклађивање нормативних аката са реалним животним околностима, велика је предност која карактерише кинески правни систем. С друге стране, поставља се питање ефикасности и делотворности тих правних норми. Сваки систем у пракси мора створити структуру у којој постоји алгоритамска транспарентност према држави, али и према грађанима, јер мора служити као механизам заштите сваког појединца.

Када се посматра утицај ових модела на људска права, не разликују се само инструменти, већ и разумевање концепта права. У ЕУ, Повеља ЕУ о основним правима функционише као уставни темељ. Легитимитет система вештачке интелигенције мери се критеријумима који поштују права на приватност, једнак третман, праведан поступак и делотворан правни лек (EU Charter 2000). Приватност је заштићена комбинацијом GDPR-а и Закона о вештачкој интелигенцији ЕУ, при чему је обрада података у системима вештачке интелигенције подложна принципима сврсисходног ограничавања и транспарентности. Недискриминација је заштићена захтевима за квалитет и репрезентативност података, тестирањем модела и ублажавањем пристрасности, док FRIA треба да идентификује ризике за рањиве групе и обрађује заштитне мере пре него што се систем пусти у рад (Pirozzoli 2024; Cosentini et al. 2024). У кинеском моделу, права као што су приватност, слобода изражавања и недискриминација формално постоје у одређеним законима, међутим њихова ефикасност и делотворност је смањена постојањем изузетака везаних за „националну безбедност“ и „друштвену

стабилност“. Ово ствара утисак да је заштита људских права у Кини питање која се може ограничавати када дође у сукоб са приоритетима државе.

Када су тако различити модели присутни у међународноправној сфери, јавља се велики број нормативних празнина. С једне стране, ЕУ покушава да глобално усмери свој хуманоцентрични приступ кроз тзв. „Бриселски ефекат“, очекујући да треће државе, компаније и међународне организације прихвате стандарде Закона о вештачкој интелигенцији и сродних права, како би задржале приступ европском тржишту (Palmiotto 2025; Pirozzoli 2024). С друге стране, Кина развија сопствени модел кроз „извоз“ технолошке инфраструктуре и нормативних решења у земље глобалног Југа. Системи вештачке интелигенције се користе за надзор, управљање становништвом и контролу информација, што ствара неку врсту „пекиншког ефекта“ нормативног ширења ауторитарних дигиталних стандарда. Као последица свега овога, настаје подељеност глобалног простора на сфере различитих нормативних утицаја, при чему међународно јавно право остаје без јединственог, обавезујућег режима за регулисање вештачке интелигенције. Ово је узрок недостатка међународноправне регулисаности питања вештачке интелигенције.

Нормативне празнине су посебно уочљиве у областима прекограничног тока података, транснационалне употребе технологија надзора и стандарда одговорности за кршења људских права узрокованих системима вештачке интелигенције. На глобалном нивоу постоје само делимичне иницијативе – попут смерница ОЕЦД-а (OECD 2019) или препорука УНЕСКО-а (UNESCO 2021) – које немају обавезујући карактер и које не решавају сукоб између европског модела заштите права и кинеског модела дигиталног суверенитета (Ufert 2020, 1089–1091). Мишљења смо да се на овај начин ствара „нормативни вакуум“ у којем појединци чија су права повређена употребом система вештачке интелигенције развијених у једном режиму, а примењиваних у другом режиму, често немају јасан правни основ за заштиту на међународном нивоу, док државе користе регулаторне разлике као инструмент геополитичке конкуренције. Ово има директне последице по универзалност људских права, јер се основне гаранције све више посматрају кроз призму регионалних и државних модела, уместо као један утврђени, универзални стандард.

Ауторке су мишљења да упоредноправна анализа модела ЕУ и Кине показује да адекватан међународноправни одговор не може бити усвајање једног од постојећих модела. Неопходно је тражити више заједничких

тачака: обавезујуће забране најопаснијих пракси вештачке интелигенције, минималне стандарде транспарентности и одговорности, као и процедуралне гаранције које би се примењивале без обзира на политички поредак. Европски модел нуди вредносно заснован и правно разрађен оквир који, поред свих предности, поседује и многе недостатке. Ту треба истаћи неуједначену примену GDPR-а и различите регулаторне праксе држава чланица, недостатак техничких капацитета, као и спори, а комплексан поступак усвајања правних аката. Због свега овога, Закон о вештачкој интелигенцији ЕУ, који се развија годинама, у неким сегментима већ делује нормативно „застарело“ у односу на брзину технолошких иновација. С друге стране, кинески модел, упркос формалности етичких принципа, брзини регулаторне реакције, али и ефикасним регулаторним телима, карактеришу нетранспарентност и широко дефинисани изузеци у корист државе. На овај начин се закон претвара у инструмент надзора, а не индивидуалне заштите. У овако хаотичном окружењу, међународно право се налази у фази касне реакције, док се нормативни вакуум попуњава билатералним споразумима, инструментима меког права и једностраним пројекцијама стандарда. Верујемо да је овај раскорак између брзине технолошких промена и неадекватног одговора међународног права централни изазов будућег развоја вештачке интелигенције у контексту заштите људских права.

Перспективе и потреба за транснационалним правним оквиром

Компаративна анализа европског и кинеског модела јасно показује да се глобално управљање вештачком интелигенцијом развија у правцу нормативне фрагментације. С једне стране, постоји вредносно заснован, али регионално ограничен и институционално оптерећен европски модел. С друге стране имамо политички централизован и ауторитаран кинески приступ, са снажним капацитетом за технолошки и регулаторни извоз. Европска унија се бори да преузме улогу снажног глобалног актера. С друге стране, глобална улога Кине је у сталном ширењу. Бројне земље зависе од кинеских инвестиција и технологије. Те околности изазивају забринутост да би кинеска технолошка инфраструктура могла да омогући прикупљање података и надзор који води влада. Појављују се и питања дигиталног

суверенитета, јер се неке државе плаше да ће зависност од кинеских технологија угрозити њихову способност да контролишу сопствене дигиталне домene и управљају националним подацима. Ово отвара питање технолошке зависности. Државе које учествују у развоју кинеских технологија могле би постати превише зависне од њих. То може ограничити њихову будућу слободу избора и аутономије, као и њихов капацитет да размотре друге опције или развију сопствене технологије, остављајући их изложеним утицају и контроли Кине (Stekić 2025, 79).

Ауторке сматрају ова асиметрија између „европског стандарда“ и „кинеског стандарда“ узрокује глобалне нормативне празнине, јер не постоји транснационални правни оквир који би дефинисао минимални скуп обавеза држава и приватних актера у вези са заштитом људских права у контексту вештачке интелигенције.

Постојећа истраживања потврђују овај закључак. Џобин, Иенка и Вајена (Jobin, Ienca and Vayena 2019) показују да на глобалном нивоу постоји сагласност око неколико етичких принципа (транспарентност, правичност, нешкодљивост, одговорност, приватност), али да постоји дубоко размимоилажење у погледу њихове конкретне примене и одговорности актера, при чему већина докумената остаје у домену меког права. Хагерти и Рубинов (Hagerty and Rubinov 2019) додатно истичу да глобална етичка архитектура вештачке интелигенције претежно одражава перспективе глобалног Севера и да се друштвене последице по маргинализоване и државе глобалног Југа систематски потцењују. На међународном нивоу, иако постоје значајни кораци, попут Оквирне конвенције Савета Европе о вештачкој интелигенцији, људским правима, демократији и владавини права, ситуација остаје суштински непромењена. Један споразум не може избрисати глобални нормативни вакуум. До сада је око 50 земаља потписало Оквирну конвенцију, али је ниједна није ратификовала.

Вештачка интелигенција је била тема од интереса за Уједињене нације, посебно када је у питању њен потенцијал у остварењу циљева одрживог развоја. Није могуће поменути све иницијативе УН за дигиталне технологије и вештачку интелигенцију, па ћемо поменути најзначајније.

Свестан утицаја који нове технологије имају на све области живота и функционисања, Генерални секретар је одлучио да оснује Панел високог нивоа за дигиталну сарадњу (2018. године) како би омогућио заинтересованим странама да се договоре и сарађују на искоришћавању потенцијала дигиталних технологија, као што је вештачка интелигенција,

истовремено се бавећи нежељеним последицама њихове примене. Следећи значајан корак догодио се 2020. године када је Генерални секретар представио извештај: „Мапа пута за дигиталну сарадњу: спровођење препорука Панела високог нивоа за дигиталну сарадњу“. Генерални секретар је изразио намеру да оснује „саветодавно тело са више заинтересованих страна о глобалној сарадњи у области вештачке интелигенције како би пружило смернице међународној заједници о вештачкој интелигенцији која је поуздана, заснована на људским правима, безбедна и одржива и промовише мир“. Такво тело би требало да служи као разноврстан форум са циљем размене и промоције најбољих пракси, као и размене мишљења о стандардизацији и напорима за усклађеност вештачке интелигенције, узимајући у обзир постојеће мандате и институције (UNGA 2020, 18). Извештај је истакао приметан недостатак инклузивности у глобалним дискусијама о вештачкој интелигенцији. Земље у развоју су углавном одсутне или нису добро заступљене у већини истакнутих форума о вештачкој интелигенцији, упркос томе што имају значајну прилику да од ње имају користи за свој економски и друштвени развој (UNGA 2020, 13).

Уједињене нације су наставиле свој рад на вештачкој интелигенцији и дигиталној технологији у новоусвојеном „Пакту за будућност“, стратегији за реформу и будуће функционисање организације. Његов саставни део је Глобални дигитални договор који има за циљ унапређење праведних и инклузивних приступа искоришћавању предности вештачке интелигенције и ублажавању ризика уз пуно поштовање међународног права, укључујући међународно право о људским правима, и узимајући у обзир друге релевантне оквире. УН сматрају да међународноправно управљање вештачком интелигенцијом захтева ефикасан, мултидисциплинаран и прилагодљив приступ са више заинтересованих страна. Глобални дигитални договор има прилику да унапреди међународно управљање вештачком интелигенцијом на начине који допуњују међународне, регионалне, националне и напоре са више заинтересованих страна (GDC 2025, 12–14).

Као део активности у оквиру Глобалног дигиталног договора, Генерална скупштина УН је 2025. године успоставила два нова механизма за промоцију међународне сарадње у управљању вештачком интелигенцијом – Независни међународни научни панел УН о вештачкој интелигенцији и Глобални дијалог о управљању вештачком интелигенцијом. Глобални

дијалог о управљању вештачком интелигенцијом требало би да послужи као платформа за дискусију о међународној сарадњи, размену најбољих пракси и научених лекција, као и за олакшавање отворених, транспарентних и инклузивних дискусија о управљању вештачком интелигенцијом са циљем да вештачка интелигенција допринесе спровођењу циљева одрживог развоја и смањењу дигиталних јаза између и унутар држава (UNGA 2025). С друге стране, Независни међународни научни панел о вештачкој интелигенцији има циљ да пружи независне научне процене које ће помоћи међународној заједници да предвиди нове изазове и донесе одлуке о управљању овом изазовном технологијом.

Имајући у виду све ове иницијативе и активности, јасно је да организација Уједињених нација остаје „незаменљива у координирању истински глобалних напора за изградњу одрживог регулаторног оквира за вештачку интелигенцију и промоцију њеног позитивног утицаја у глобалном друштву“ (Dabić 2024, 178). С друге стране, неопходно је напоменути да све иницијативе УН, укључујући Глобални дигитални договор, још увек нису резултирале обавезујућим режимом. Ове иницијативе остају на декларативном нивоу и не уводе јасне обавезе ни државама ни компанијама. Све наведено подржава нашу почетну тезу да постојећи међународни механизми нису довољни да обезбеде системску и универзалну заштиту основних права у условима убрзаног развоја технологија вештачке интелигенције.

На основу овога, сматрамо да би будући транснационални правни режим за вештачку интелигенцију требало да буде комбинација најмање три приступа. Први приступ би био скуп универзалних забрана одређених пракси вештачке интелигенције које су неспојиве са међународним стандардима људских права. Овде спадају масовни биометријски надзор без судског надзора, социјално кредитирање и систематско политичко профилисање грађана. Други приступ били би процедурални стандарди, инспирисани искуством ЕУ, који би обавезали државе и приватне актере да *ex ante* процене утицај система вештачке интелигенције на људска права, транспарентност кључних алгоритамских одлука и постојање ефикасних механизма правне заштите, укључујући право на објашњење, приговор и приступ суду. Трећи приступ би био успостављање међународног механизма за праћење, са глобалним мандатом, који би могао да комбинује елементе постојеће уговорне контроле у области људских права са специфичном техничком експертизом. У идеалном случају, то би била

институција која функционише у оквиру ширег оквира УН. Имајући у виду тренутну светску геополитичку ситуацију, успостављање такве институције могло би се догодити тек у наредним деценијама. Таква институција би се могла суочити са многим функционалним изазовима.

Постојећа истраживања о глобалном управљању вештачком интелигенцијом указују да је приступ људским правима, заснован на Универзалној декларацији о људским правима и Пакту о грађанским и политичким правима, неопходан услов за легитимитет будућег режима, али да сам по себи није довољан ако није повезан са конкретним регулаторним инструментима и механизмима одговорности (Jobin, Ienca и Vayena 2019; Hagerty и Rubinov 2019). У том смислу, наше препоруке иду у правцу глобалне хармонизације стандарда заштите људских права кроз комбинацију минималног обавезујућег међународног уговора о вештачкој интелигенцији заснованог на људским правима. Неопходна је и међусобна хармонизација регионалних режима (као што је Закон о вештачкој интелигенцији ЕУ и будући азијски или амерички стандарди) и институционализовано учешће различитих актера. Држава, међународне организације, приватни сектор, академска заједница и цивилно друштво требају имати неку врсту укључености у обликовање и праћење примене тих стандарда. Верујемо да такав транснационални оквир представља једини реалан начин за ограничавање негативних ефеката нормативне конкуренције између ЕУ и Кине и за изградњу минималног „*ius commune*“ заштите људских права у ери вештачке интелигенције.

Закључна разматрања

Полазећи од хипотезе да истовремено постојање вредносно заснованог, али регионално ограниченог европског модела регулисања вештачке интелигенције и политички централизованог, ауторитарног кинеског модела, уз одсуство кохерентног међународног правног режима, производи глобални „нормативни вакуум“ у заштити људских права, ауторке сматрају да спроведена анализа потврђује ову претпоставку. Рад је показао да Европска унија, кроз Закон о вештачкој интелигенцији ЕУ, Општу уредбу о заштити података и Повељу ЕУ о основним правима, гради развијен хоризонтални оквир за заштиту појединца. У овом систему су приватност, недискриминација, достојанство и процесне гаранције

интегрисане у саму структуру регулације система вештачке интелигенције. Овај модел има и много недостатака који смањују ефикасност правних норми, као и механизме за заштиту основних људских права. С друге стране, кинески модел је заснован на комбинацији системских закона о сајбер безбедности, подацима и личним информацијама, специфичним прописима о алгоритамским препорукама, дубокој синтези и генеративној вештачкој интелигенцији. Овај модел развија високо софистицирану, али ауторитарну архитектуру надзора, у којој су људска права подређена приоритетима политичке стабилности, друштвеног управљања и технолошког суверенитета. Овај рад је показао да оба модела производе екстериторијалне ефекте – европски кроз „Бриселски ефекат“ и нормативни извоз стандарда, а кинески кроз технолошку експанзију и регулаторни утицај. Међутим, ниједан од њих не успева да обезбеди универзалну и транснационално прихваћену минималну заштиту људских права у сфери вештачке интелигенције.

Закључци рада могу се истаћи у неколико тачака. Пре свега, постојећа међународноправна регулатива је фрагментисана, претежно „меке“ природе и не садржи обавезујуће стандарде који би се бавили специфичним ризицима вештачке интелигенције по људска права. Европски модел представља најнапреднији покушај интеграције људских права у техничку регулативу, али је оптерећен политичким компромисима, неуједначеном имплементацијом и ограниченим географским дометом. С друге стране, кинески модел показује изузетну нормативну и техничку мобилност, али је у суштини усмерен на консолидацију државне власти, што генерише озбиљне ризике за права на приватност, слободу изражавања, недискриминацију и праведан третман. Постојећа емпиријска и теоријска истраживања о глобалној етици вештачке интелигенције потврђују нашу полазну тачку о неопходности транснационалног правног оквира који би најважније принципе људских права поставио као главни циљ заштите.

На основу тога, правци будућих истраживања и међународних иницијатива требало би да се крећу у три правца. Прво, потребно је детаљно пратити примену европских и кинеских модела у пракси, укључујући њихове екстериторијалне ефекте, како би се утврдиле стварне последице по људска права у различитим друштвеним контекстима. Друго, потребно је даље развијати теоријске и нормативне моделе транснационалног управљања вештачком интелигенцијом, са посебним освртом на улогу међународних

организација, регионалних судова за људска права и нових уговорних режима. Треће, будућа истраживања морају укључити перспективе држава глобалног Југа, маргинализованих група и различитих културних контекста, како би се избегло продубљивање јаза на глобалном нивоу. Ауторке сматрају да само интегрисани приступ, који повезује упоредно право, међународно право људских права и студије науке и технологије, може понудити адекватан одговор на изазове које вештачка интелигенција поставља савременом међународноправном поретку.

Биће интересно пратити даље идеје и иницијативе Уједињених нација у погледу вештачке интелигенције. Неопходно је истаћи да је рад Уједињених нација фокусиран на однос између вештачке интелигенције и испуњење циљева одрживог развоја. Питање транснационалног правног оквира у вези са вештачком интелигенцијом није најважнији фокус светске организације. С друге стране, у ери изразитих геополитичких тензија, велике силе не подстичу колективне напоре за успостављање одговарајућег правног оквира за управљање вештачком интелигенцијом. Напротив, оне су сумњичаве према намерама других држава и недржавних актера, сматрајући их конкурентима у геополитичкој игри. У будућности, превелики фокус држава на остваривање геополитичких циљева се може сматрати главном препреком у постизању договора о транснационалном правном оквиру који ће регулисати питање вештачке интелигенције.

Dr. Jelica Gordanić's part of this paper present findings of a study developed as a part of the research project "Serbia and challenges in international relations in 2026", financed by the Ministry of Science, Techonological Development and Inovation of the Republic of Serbia, and conducted by Institute of International Politics and Economics, Belgrade, during the year 2026.

Bibliography/Библиографија

[AI Act] Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 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 (Artificial

- Intelligence Act). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>.
- [ARR CN] Provisions on the Administration of Algorithmic Recommendation in Internet Information Services. 31 December 2021. <https://www.chinalawtranslate.com/en/algorithms/>.
- [Beijing AI Principles] Beijing AI Principles. 2019. <https://ai-ethics-and-governance.institute/beijing-artificial-intelligence-principles/>.
- Bertaina, Samuele, Ilaria Biganzoli, Rachele Desiante, Dario Fontanella, Nicole Inverardi, Ilaria Giuseppina Penco, and Andrea Claudio Cosentini. 2025. "Fundamental Rights and Artificial Intelligence Impact Assessment: A New Quantitative Methodology in the Upcoming Era of AI Act" *Computer Law & Security Review* 56: 106101.
- [CAC] Mandatory Algorithm Registry of the Cyberspace Administration of China. n.d. <https://www.cac.gov.cn/index.htm>.
- Calo, Ryan. 2017. *Artificial Intelligence Policy: A Primer and Roadmap*. Seattle, WA: University of Washington, School of Law.
- Cosentini, Andrea, Oreste Pollicino, Giovanni De Gregorio, Andrea Ermellino, Dario Fontanella, Nicole Inverardi, Federica Paolucci, Ilaria Giuseppina Penco, Daniele Regoli, and Silvia Tessaro Trapani. 2024. *Assessing the Impact of Artificial Intelligence Systems on Fundamental Rights*. Vienna: European Union Agency for Fundamental Rights.
- [Council of Europe] Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Council of Europe Treaty Series – No. 225, Vilnius, 25 September 2024. https://www.gov.il/BlobFolder/news/ai2024/en/CETS_225_EN.docx_.pdf.
- Creemers, Rogier. 2022. "China's Emerging Data Protection Framework". *Cybersecurity* 8 (1): 105676.
- [CSL CN] Cybersecurity Law of the People's Republic of China. 2017. <https://digichina.stanford.edu/work/translation-cybersecurity-law-of-the-peoples-republic-of-china-effective-june-1-2017/>.
- Dabić, Dragana. 2024. "Predlozi za uspostavljanje regulatornog okvira veštačke inteligencije na međunarodnom planu". U: *Dinamika savremenog pravnog poretka*, uredio Srđan Radulović, 171–188. Kosovska Mitrovica i Beograd: Univerzitet u Prištini – Pravni fakultet, Institut za kriminološka i sociološka istraživanja i Institut za uporedno pravo.

- Dabić, Dragana. 2025. "Evropska kancelarija za veštačku inteligenciju". *Evropsko zakonodavstvo* 89 (2): 221–238.
- [DMA] Regulation (EU) 2022/1925 of the European Parliament and of the Council of 14 September 2022 on contestable and fair markets in the digital sector and amending Directives (EU) 2019/1937 and (EU) 2020/1828 (Digital Markets Act). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A32022R1925>.
- [DSA] Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC (Digital Services Act). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32022R2065&qid=1778602543598>.
- [DSI CN] Provisions on the Administration of Deep Synthesis Internet Information Services. 2022. <https://www.chinalawtranslate.com/en/deep-synthesis/>.
- [DSL CN] Data Security Law of the People's Republic of China. 2021. https://en.spp.gov.cn/2021-06/10/c_948426_2.htm.
- [EU Charter] Charter of Fundamental Rights of the European Union. 2000. https://www.europarl.europa.eu/charter/pdf/text_en.pdf.
- Floridi, Luciano. 2018. "Soft Ethics and the Governance of the Digital". *Philosophy & Technology* 31: 1–8.
- [GAIR CN] Interim Measures for the Administration of Generative Artificial Intelligence Services. 2023. <https://www.chinalawtranslate.com/en/generative-ai-interim/>.
- [GDC] United Nations Global Digital Compact 2025. <https://www.un.org/digital-emerging-technologies/global-digital-compact>.
- [GDPR] Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32016R0679&qid=1778602643426>.
- Hagerty, Alex, and Igor Rubinov. 2019. "Global AI Ethics: A Review of the Social Impacts and Governance of Artificial Intelligence". Preprint. DOI: <https://doi.org/10.48550/arXiv.1907.07892>.
- Hulok, Martina. 2025. "The EU Model of AI Governance: Regulating Artificial Intelligence Through Law and Policy". *ERA Forum* 26: 527–547.

- Huster, Stefan, and Ian Walden. 2019. "EU Data Protection Law and the Protection of Human Dignity, Autonomy and Privacy". *Common Market Law Review* 56 (4): 1005–1034.
- Jobin, Anna, Marcello Lenca, and Effy Vayena. 2019. "The Global Landscape of AI Ethics Guidelines". *Nature Machine Intelligence* 1 (9): 389–399.
- [Joint Pledge CN] Joint Pledge on Self-Discipline in the AI Industry. 2019. <https://digichina.stanford.edu/work/translation-chinese-ai-alliance-drafts-self-discipline-joint-pledge/>.
- Lee, John. 2022. *"Cyberspace Governance in China: Evolution, Features and Future Trends"*. Paris: IFRI.
- Lynskey, Orla. 2015. *The Foundations of EU Data Protection Law*. Oxford: Oxford University Press.
- Murnisari, Mira. 2024. „China's Cyber Security Governance as Part of Foreign and Security Policy". *Security Intelligence Terrorism Journal* 1 (2): 164–177.
- [OECD] OECD Principles on Artificial Intelligence. 2019. <https://www.oecd.org/en/topics/ai-principles.html>.
- Palmiotto, Francesca. 2025. "The AI Act Roller Coaster: The Evolution of Fundamental Rights Protection in the Legislative Process and the Future of the Regulation". *European Journal of Risk Regulation* 16 (2): 770–793.
- [PIPL CN] Personal Information Protection Law of the People's Republic of China. 2021. <https://personalinformationprotectionlaw.com/>.
- Pirozzoli, Anna. 2024. "The Human-Centric Perspective in the Regulation of Artificial Intelligence". *European Papers* 9 (1): 105–116.
- [RAIP CN] Governance Principles for Responsible AI. 2019. https://www.most.gov.cn/kjbgz/201906/t20190617_147107.html.
- Roberts, Huw, Josh Cows, Jessica Morley, Mariarosaria Taddeo, Vincent Wang, and Luciano Floridi. 2021. "The Chinese Approach to Artificial Intelligence: An Analysis of Policy, Ethics and Regulation". *AI & Society* 36: 59–77.
- Sheehan, Matt. 2023. *China's AI Regulations and How They Get Made*. Washington, DC: Carnegie Endowment for International Peace.
- Shen, Wei, and Huan Liu. 2021. "China's Normative Systems for Responsible AI: From Soft Law to Hard Law". In: *The Oxford Handbook of Ethics of AI*, edited by Markus D. Dubber, Frank Pasquale and Sunit Das, 149–170. Oxford: Oxford University Press.

- Stekić, Nenad. 2025. "Globalno upravljanje u doba digitalnih rizika: kineska inicijativa za veštačku inteligenciju". *Sociološki pregled* 59 (1): 76–100.
- Ufert, Fabienne. 2020. "AI Regulation Through the Lens of Fundamental Rights: How Well Does the GDPR Address the Challenges Posed by AI?". *European Papers* 5 (2): 1087–1097.
- [UNESCO] UNESCO Recommendation on the Ethics of Artificial Intelligence. 2021. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>.
- UNGA [UNGA] United Nations General Assembly. Resolution 74/281. 2020. <https://docs.un.org/en/A/RES/74/281>.
- UNGA [UNGA] United Nations General Assembly. Resolution 79/L.118. 2025. <https://docs.un.org/en/A/79/L.118>.
- Wang, Wayne Wei, Lingfeng Zhu, Xiang Wang, Xingsi Di, and Yue Zhu. 2025. "Artificial Intelligence 'Law(s)' in China: Retrospect and Prospect". *Journal of AI Law and Regulation* 2 (1–2): 29–36 & 139–157.
- Wu, Wenjun, Tiejun Huang, and Ke Gong. 2020. "Ethical Principles and Governance Technology Development of AI in China". *Engineering* 6 (3): 302–309.
- Zhu, Junhua. 2022. "AI Ethics with Chinese Characteristics? Concerns and Preferred Solutions in Chinese Academia". *AI & Society* 37 (4): 1487–1505.