

Artificial intelligence algorithms in legal practice: Legal boundaries, socio-legal implications and liability

Yurii Lukashchuk*

PhD in Computer Science, Assistant
Lviv Polytechnic National University
79000, 12 Stepan Bandera Str., Lviv, Ukraine
<https://orcid.org/0000-0002-8933-8635>

Uliana Tsmots

Senior Researcher
Lviv State University of Internal Affairs
79007, 26 Horodotska Str., Lviv, Ukraine
<https://orcid.org/0000-0002-8423-8272>

Abstract. The relevance of this study stems from the rapid digitalisation of legal practice and the growing role of artificial intelligence algorithms in legal decision-making processes, which gives rise to new legal and ethical challenges. The aim of the study was to examine the intersection of legal, ethical and institutional aspects of the application of artificial intelligence in the legal sphere, as well as to define the limits of permissible use of algorithmic systems and approaches to the allocation of liability. The study employed systematic, comparative legal and formal legal methods, as well as an analysis of judicial practice. The risks of automated decision-making were examined: algorithmic bias, the “black box” problem, the reproduction of discrimination and threats to judicial independence. These risks are confirmed by empirical research in the field of algorithmic justice and an analysis of the practice of using risk prediction systems in criminal proceedings (e.g., COMPAS in the US). The regulation of AI is analysed using the example of the EU Artificial Intelligence Act and the GDPR, as well as the ECHR’s approaches to ensuring transparency and procedural safeguards. Approaches to the allocation of responsibility between developers, suppliers, users and the state are summarised, and a model of algorithmic transparency with human oversight and the parties’ right to be informed is proposed. International experience from the US, France, Singapore and China is examined, as well as the prospects for implementing AI in Ukraine within the framework of the EU-ITC and case law analysis systems. Particular attention is paid to the concept of distributed algorithmic liability and the application of the analogy of liability for a source of increased danger to high-risk AI systems. A model of algorithmic transparency is proposed, incorporating the parties’ right to be notified of AI use, algorithmic expertise, and the court’s obligation to justify its consideration of AI recommendations. The principle of “human error” is outlined as a guarantee of judicial discretion and the independence of the judge. The results of the study can be used to improve national legislation in the field of digital justice and to implement standards of algorithmic transparency

Keywords: algorithmic accountability; ethics; discrimination; digital justice; EU Artificial Intelligence Act; General Data Protection Regulation

Introduction

The rapid development of digital technologies and artificial intelligence is transforming the legal profession by automating processes and optimising decision-making. Algorithmic systems are influencing the analysis of case law, risk assessment, document preparation and the prediction of dispute outcomes, thereby altering the role of lawyers, standards of professional responsibility and institutional dynamics. This poses challenges regarding the balance between efficiency and the principles of the rule of law, equality of the parties and judicial independence. Particular attention is

being paid to regulating the use of algorithms, preventing bias, ensuring transparency of decisions and harmonisation with EU standards, particularly regarding the protection of human rights.

Current research at the intersection of law and artificial intelligence is giving rise to several interrelated theoretical strands, centred on algorithmic governance, the explainability of automated decisions, the transformation of justice, and the evolution of models of legal liability. Within the framework of algorithmic governance theory, K. Yeung &

Suggested Citation — **Article’s History:** Received: 24.12.2025 Revised: 20.03.2026 Accepted: 27.05.2026 Published: 01.06.2026

Lukashchuk, M., & Tsmots, U. (2026). Artificial intelligence algorithms in legal practice: Legal boundaries, socio-legal implications and liability. *Social & Legal Studios*, 9(2), 67-78. doi: 10.32518/sals2.2026.67.

*Corresponding author (ndrv@i.ua)



L. Ulbricht (2022) view algorithmic regulation as a form of public administration that relies on data and automated systems to perform regulatory functions. This approach emphasises improved governance efficiency, yet simultaneously raises issues of democratic accountability, oversight, and the legitimacy of delegating governmental functions to algorithms. From a philosophical-legal perspective, M. Hildebrandt (2020) conceptualises law as an interpretative system based on textual openness, proceduralism and the possibility of appeal, which enters into structural tension with automated forms of decision-making, as the latter reduce the scope for legal interpretation and individualised discretion. The issue of transparency and explainability of algorithmic decisions in European Union law is analysed in detail in the works of S. Wachter *et al.* (2017a), who question the normative certainty of the “right to an explanation” within the framework of General Data Protection Regulation (2016). The authors demonstrate that its current provisions do not establish a clear obligation to disclose the internal logic of algorithmic systems, leading to doctrinal uncertainty regarding the scope of data subjects’ information rights.

Further research clarifies that there is a fundamental tension between the legal requirements for the justification of decisions and the technical complexity of modern machine learning models, which makes it impossible to equate explainability strictly with the disclosure of algorithmic logic. In this context, an approach to functional explainability is emerging, within which priority is given to procedural verifiability, the comprehensibility of results and the possibility of judicial review, allowing the requirements of legal certainty to be adapted to the technical limitations of algorithmic systems. In the field of the transformation of justice, R. Susskind (2019) argues that digital technologies do not replace the legal profession, but transform its structure, reducing the proportion of routine tasks and enhancing the importance of analytical and strategic functions. This transformation is taking place in parallel with the development of regulatory frameworks for artificial intelligence in the European Union, which are based on a risk-based approach conceptualised in the research of L. Edwards (2022) and institutionally enshrined in the Artificial Intelligence Act (2024). According to this model, artificial intelligence systems are classified according to their level of risk to fundamental rights and security, which allows for the combination of innovative development with preventive legal control.

Ukrainian legal doctrine, as presented in the works of O.A. Baranov (2023), highlights the absence of a unified legislative definition of artificial intelligence, which complicates the formation of a coherent system of legal regulation and harmonisation with European standards of digital governance. This indicates a transitional state of the national doctrine, in which conceptual understanding precedes normative institutionalisation. The issue of legal liability for harm caused by autonomous systems is traditionally analysed through the prism of classical tort law constructs, particularly in the works of U. Pagallo (2013), where the application of traditional fault-based models to robotic systems is considered to be limited. Further generalisations, presented in *The Cambridge Handbook of the Law of Algorithms* (Barfield, 2020), emphasise the need to adapt legal constructs to the conditions of autonomous decision-making. Contemporary doctrine is seeing a shift towards models of distributed liability, which take into account the

involvement of all parties in the life cycle of an algorithmic system – developers, data providers, operators and users – thereby altering the traditional understanding of causation and individual fault.

A review of the current academic literature allows to identify three interrelated areas of research: algorithmic governance as a form of transformation in public administration; the explainability and procedural accountability of automated decisions; and the evolution of models of legal liability for the actions of autonomous systems. At the same time, conceptual fragmentation persists between these areas, as regulatory, technical-legal and tort approaches develop in parallel but are insufficiently integrated, creating a key research gap in contemporary legal doctrine on artificial intelligence.

The aim of this article was to conduct a comprehensive analysis of the legal boundaries of the application of artificial intelligence algorithms in legal practice, taking into account their socio-legal consequences, identifying key risks, and formulating conceptual approaches to the allocation of liability for the consequences of algorithmic decisions. To achieve this aim, the article examines the regulatory approaches of the European Union, the case law of the European Court of Human Rights, as well as contemporary socio-legal and interdisciplinary approaches to the analysis of algorithmic systems. The research hypothesis was that the effective integration of artificial intelligence algorithms into legal practice is possible only if a distributed algorithmic liability and a procedural model of algorithmic transparency are introduced, thereby ensuring the preservation of the principle of human judicial discretion.

Materials and methods

Conceptually, the study is based on a combination of the doctrine of the rule of law, the theory of digital constitutionalism, and a risk-based approach to technology regulation. Within the scope of the study, algorithmic systems are considered as an element of the transformation of public law governance, combining regulatory, procedural and tort aspects of legal regulation. The methodological basis of the article was a set of general scientific and specialised legal methods of inquiry, applied within the framework of contemporary concepts of the rule of law, algorithmic accountability, and risk-based regulation of artificial intelligence. The theoretical basis of the study was provided by the doctrines of digital constitutionalism, the concept of due process, and approaches to the regulation of high-risk technologies in European Union law. A range of scientific research methods were employed in the study. In particular, the dogmatic method was used to analyse the normative content of the provisions of the Artificial Intelligence Act (2024) and the General Data Protection Regulation (2016), as well as to interpret the standards of a fair trial established in the case law of the European Court of Human Rights. The application of this method made it possible to identify the normative requirements for transparency, accountability and the lawfulness of the use of algorithmic systems in the field of justice. This enabled the content of legal obligations regarding the transparency and accountability of algorithmic systems to be established and the limits of judicial review of algorithmic decisions to be determined.

The comparative legal method was used to compare the European Union’s approaches to the regulation of high-risk

artificial intelligence systems, as provided for in the Artificial Intelligence Act (2024), with the doctrinal models of algorithmic liability established in contemporary European legal doctrine and regulatory practice concerning artificial intelligence systems. The systemic method allowed for the consideration of algorithmic decision-making as a cross-sectoral legal phenomenon combining substantive, procedural and public law aspects, which ensured the formation of a holistic view of algorithmic liability as a component of the modern legal system and the establishment of interrelationships between its substantive, procedural and institutional mechanisms. Functional analysis was applied to assess the impact of algorithmic systems on the implementation of the principles of the rule of law, adversarial proceedings and legal certainty, taking into account the guarantees enshrined, in particular, in Article 6 of the European Convention on Human Rights (1950). The concept of counterfactual explainability by S. Wachter *et al.* (2017b) has also been adapted to analyse the requirements for algorithmic transparency in judicial proceedings.

Results and Discussion

Theoretical and legal foundations and risks of algorithmic decision-making in the legal sphere

In contemporary legal doctrine, algorithmic decision-making is primarily viewed through the prism of ethical principles and technological risks; however, its procedural and legal dimension remains significantly less developed. First and foremost, this concerns determining the status of algorithmic conclusions in court proceedings, the standards of their probative value, and the limits of their permissible use. Algorithmic decision-making in the legal sphere is the process of applying software systems which, based on mathematical models, analyse data and generate recommendations or decisions on legal matters. A distinctive feature of such systems is the use of machine learning; autonomy in generating results; and limited explainability. As F. Pasquale (2015) emphasises, the opacity of algorithms creates a “black box” risk, where a person cannot understand the logic behind a decision that directly affects their rights. At the same time, the “black box” problem should be viewed not merely as a technical limitation, but as an independent procedural risk that affects the admissibility of algorithmic results as a source of evidence, requiring new approaches to the explainability of algorithmic conclusions (Wachter *et al.*, 2017).

Furthermore, it should be emphasised that the “black box” problem takes on not only a technical but also a procedural dimension, as in a judicial context it directly affects the standards of proof. The impossibility of fully reconstructing the logic of an algorithmic conclusion creates a situation of information asymmetry between the parties to the proceedings, which potentially violates the principle of equality of arms. In such circumstances, traditional standards of admissibility and evaluation of evidence require adaptation through the introduction of a presumption of verifiability of algorithmic results, which entails an obligation on the party using the algorithm to prove its reliability, the representativeness of the data, and the absence of discriminatory effects. The procedural implementation of the presumption of verifiability must be based on requirements for transparency and explainability, which, within the framework of the Artificial Intelligence Act (2024) are transformed from purely

ethical categories into strict legal obligations for developers of high-risk systems and implemented through the introduction of a mandatory “digital algorithm passport” in the case file. Such a document must contain not only technical specifications but also the results of stress testing the model for discriminatory deviations. In the event that a party to the proceedings reasonably questions the objectivity of an algorithmic conclusion, the burden of proving the system’s reliability must fall on the entity that implemented it. This will enable the transformation of the abstract principle of transparency into a concrete instrument of judicial review, where the “right to an explanation” becomes an integral part of the right to a fair trial.

Among the fundamental principles governing the use of artificial intelligence in legal practice, the principles of legality, transparency, non-discrimination, accountability and the assurance of human oversight over the functioning of AI systems (human in the loop) are highlighted. This approach has been widely discussed in the academic literature on the ethical aspects of AI use (Davenport & Ronanki, 2018). At the same time, these principles in their current form are largely declaratory in nature and do not specify the mechanisms for their procedural implementation. In this regard, the study proposes their concretisation through the categories of evidence law and procedural safeguards.

These principles are reflected in the approaches of the European Commission’s High-Level Expert Group on AI (2019) to establishing ethical frameworks for artificial intelligence, as well as in the case law of the European Court of Human Rights regarding guarantees of a fair trial in the context of digitalisation. One of the most dangerous risks is the reproduction of systemic bias through training data. Research by J. Angwin *et al.* (2016) on the COMPAS system demonstrated the potential for racial bias in predicting recidivism. C. O’Neil (2016) emphasises that algorithms can exacerbate social inequality. Algorithms can also reproduce structural discrimination (Barocas & Selbst, 2016). This is particularly critical in creditworthiness assessments, the determination of preventive measures, recidivism prediction, and automated recruitment. Recent research confirms that the problem of algorithmic bias in judicial predictive tools remains a systemic challenge, requiring constant auditing and verification of training data to prevent discriminatory effects. This concludes that algorithmic conclusions cannot be presented as neutral; consequently, their use in law enforcement requires mandatory verification of data representativeness and the absence of discriminatory effects, which should be considered as an element of the procedural standard of proof.

The use of AI in judicial proceedings may affect the principle of judicial independence. Excessive reliance on algorithmic recommendations creates the risk of “de facto automated justice”, where statistical probability supplants legal reasoning. In this context, it is important to consider the ethical limits of predictive justice, as a judicial decision requires not only the accuracy of the prediction but also value-based reasoning.

In accordance with Article 22 of the General Data Protection Regulation (2016), a person has the right not to be subject to a decision based solely on automated processing which produces legal effects. This provision forms a fundamental safeguard against the complete algorithmisation of justice. At the same time, it is important to bear in mind that

algorithmic systems in the legal sphere are not neutral tools: they reflect the quality of the data on which they are trained and the assumptions embedded in by developers. This means that even technically sound models can produce legally and ethically problematic results. In legal practice, this creates a risk of “legitimising error through technology”, where a decision is perceived as objective solely because of its algorithmic origin, even though it may in fact reproduce hidden biases or methodological limitations of the model.

Furthermore, it should be noted that algorithmic decision-making transforms the very nature of evidence and legal argumentation. Whilst traditional justice is based on the open assessment of evidence and the reasoning behind decisions, algorithmic models operate on statistical correlations that do not always lend themselves to legal interpretation. This gives rise to the problem of the “explanation gap”, where a conceptual mismatch arises between the technical logic of the model and the legal logic of the judicial decision. In this context, the Artificial Intelligence Act (2024) significantly impacts traditional methods of constructing legal reasoning, requiring the adaptation of classical models of legal argumentation to the specificities of algorithmic systems. Consequently, the importance of procedural safeguards for verifying algorithmic conclusions is growing, as is the need for their critical assessment by the judge as an independent law enforcement authority.

An important element of the theoretical and legal analysis of algorithmic decision-making is the classification of risks, which allows for the differentiation of the level of legal regulation depending on the system’s potential impact on human rights. In this context, it is advisable to adopt the risk-based approach set out in the European Union’s Acts on artificial intelligence, which classifies systems into categories of unacceptable, high, limited and minimal risk. The European Parliament’s official position confirms that this approach aims to create a safe environment for innovation whilst protecting citizens’ fundamental rights. Algorithms used in the justice system are, by their very nature, high-risk systems, as they directly impact access to justice and the equality of the parties, which requires specific navigation within the framework of EU regulatory requirements. Such a classification has not only theoretical but also practical significance, as it determines the scope of mandatory requirements for transparency, auditing, risk management and human oversight that must be integrated into legal practice when using AI.

ECHR case law on standards of technological interference with human rights and their relevance to AI

The case law of the European Court of Human Rights includes judgments that establish legal standards regarding the interference of technical systems with fundamental rights, which can be extrapolated to the context of the use of AI algorithms. In particular, in the case of *Zakharov v. Russia* (2015), the Court found that the system of covert interception of mobile communications in Russia did not provide adequate and effective safeguards against arbitrary interference with the right to respect for private life, and therefore violated Article 8 of the European Convention on Human Rights (the right to private and family life). This decision highlights the need for clear procedural restrictions and judicial oversight of technological systems, which directly resonates with the requirements for transparency and accountability in the application of high-risk algorithmic tools in legal practice.

In the case of *Big Brother Watch and Others v. the United Kingdom* (2021), the Grand Chamber of the European Court of Human Rights examined comprehensive electronic surveillance regimes, in particular mass data interception, and concluded that “end-to-end” safeguards were insufficient to protect the right to respect for private life (Article 8 of the Convention). Although this judgment concerns surveillance regimes, it sets out standards of legality, foreseeability and independent oversight of technological tools that must be established to ensure effective legal protection for individuals – and these standards can be adapted to assess the lawfulness of the use of AI algorithms in the legal sphere. Summarising the ECtHR’s approach, the study argues that the standards of “quality of law”, predictability and effective control can be transformed into criteria for the admissibility of algorithmic intervention in legal activities, even in the absence of direct case law from the Court regarding artificial intelligence.

At the same time, although the ECtHR has not explicitly considered algorithmic artificial intelligence systems, its approach to technological interference with fundamental rights allows for the formulation of criteria for assessing algorithmic law enforcement. In particular, what is key is not only the existence of a legal basis for the interference, but also the quality of the procedural safeguards accompanying such interference. In this context, the requirements of “quality of the law” and “effective control” may be interpreted as also encompassing the duty to ensure the comprehensibility of algorithmic logic to the extent necessary to protect the procedural rights of the parties. Thus, the standards of the Convention for the Protection of Human Rights effectively form the basis for a future doctrine of algorithmic rule of law.

Thus, the aforementioned rulings do not directly concern algorithmic decision-making, but they establish standards of procedural predictability, accountability and effective control of technological systems, which may be applied *mutatis mutandis* to artificial intelligence systems in legal practice. In this context, an approach is proposed to understanding algorithmic conclusions as a specific form of evidential information, combining the characteristics of an expert opinion and digital evidence, yet requiring a special procedural regime for verification.

Algorithmic evidence in legal practice requires a rethinking of traditional procedural standards of relevance, admissibility and reliability of evidence. Unlike classical evidence, algorithmic inferences are based on statistical models, which complicates their verification within the judicial process. In this regard, there is a need to introduce special procedural mechanisms for verifying algorithmic results, including algorithmic expert assessment, training data audits and verification of the reproducibility of results. This ensures that algorithmic systems comply with the principles of adversarial proceedings and equality of the parties in the process.

Practical application, legal regulation and liability in the field of AI use

Contemporary examples of the use of artificial intelligence in legal practice demonstrate that algorithms can significantly improve the efficiency and accuracy of legal work (McGinnis & Pearce, 2014; Remus & Levy, 2017; Susskind, 2019). For instance, in the US, the ROSS Intelligence platform, powered by IBM Watson, helps lawyers quickly identify relevant

legislative and judicial provisions, reducing legal research time by tens of hours. Similarly, Ravel Law analyses judges' decisions and identifies trends in rulings, enabling law firms to prepare more well-founded arguments and reduce the influence of human bias on case assessments. In judicial practice, the use of systems for predicting case outcomes and automating document management has proven successful. In France, tools for analysing judicial practice and predicting court decisions using artificial intelligence are employed, based on statistical analysis of court data and machine learning (Sourdin, 2018). In Singapore and China, the eLitigation and Smart Court systems automate the filing of claims, classify cases and operate with mandatory human involvement in the decision-making process, which helps to speed up court proceedings without infringing on human rights.

Law firms are also successfully using AI to analyse contracts and assess risks. For instance, the Kira Systems and Luminance platforms automatically review large volumes of corporate documents, automate the analysis of contracts and corporate documents, and identify key provisions and potential risks, thereby reducing audit time and improving the accuracy of assessments. All these cases demonstrate that, provided human oversight and procedural safeguards are properly integrated, artificial intelligence algorithms can become a powerful tool for effective and responsible legal practice.

The digitalisation of Ukraine's judicial system is a key element in the modernisation of the justice system and paves the way for the gradual integration of algorithmic technologies into legal practice. One of the key tools of this transformation has been the Unified Judicial Information and Telecommunications System, which facilitates electronic document management in courts, remote access to case files, electronic communication between the court and parties to proceedings, and the automated allocation of cases among judges. The introduction of this system contributes to increasing the transparency of the judicial process, optimising administrative procedures, and reducing the risks of human interference in the allocation of court cases.

Another key element of the digital justice infrastructure is the Unified State Register of Court Decisions, which provides open access to a vast body of case law. The existence of a structured database of court decisions creates opportunities for the application of big data analysis and machine learning technologies to identify trends in judicial practice, predict likely outcomes of disputes, and improve the quality of legal analysis. In this context, the use of algorithmic tools to analyse judicial practice can significantly improve the efficiency of drafting procedural documents, contribute to the formation of more well-founded legal positions, and improve access to justice.

At the same time, the use of analytical algorithms in the field of judicial practice requires compliance with clear legal and ethical restrictions. Algorithmic systems may be used for the statistical analysis of court decisions, the identification of patterns in the application of the law, and the auxiliary assessment of litigation risks; however, they must not replace the judge's discretion or create the conditions for automated justice. Therefore, the integration of algorithmic analytics into the Ukrainian judicial system must be carried out in a manner that ensures the principles of human oversight, algorithmic transparency, and the possibility of procedural challenges to the algorithmic influence on the legal assess-

ment of a case. The inability to fully explain an algorithmic conclusion contradicts the requirements for the reasoned nature of judicial decisions. S. Wachter *et al.* (2017) emphasise that the right to an explanation is complex and does not always follow directly from the General Data Protection Regulation (2016), yet it is logically linked to the principles of transparency and accountability.

One of the key issues is determining the party legally liable for harm caused by an algorithmic decision, in particular the software developer, the system provider, the court, the lawyer or the state (Calo, 2015; Veale & Borgesius, 2021). In this regard, the study justifies the applicability of a model of functional (distributed) liability, which is based not on formal fault but on the degree of control over algorithmic risk. M. Veale & F.Z. Borgesius (2021) emphasise the need for a clear demarcation of liability between the entities in the chain of algorithm creation and application. The model of distributed liability requires regulatory enshrinement at the level of legislation and professional standards.

The issue of liability for harm caused by algorithmic systems in legal practice is multi-layered and cannot be resolved within the framework of the classical tort model. In this context, the concept of "functional liability" requires particular attention, as it provides for liability to be allocated not according to the formal criterion of fault, but according to the degree of control over algorithmic risk. This approach allows for the reflection of the actual architecture of modern digital systems, in which no single entity has complete control over the final outcome, yet each influences a specific stage of the algorithm's lifecycle. Accordingly, legal liability must shift from an individualised model to a model of distributed risk, where the key criterion is the degree of technical and organisational capacity to prevent harm.

Contemporary legal doctrine identifies several possible models. In particular, classical tort liability, which is based on proving unlawful conduct, the existence of harm, a causal link between the action and the consequences, as well as fault. However, in the case of algorithmic systems, difficulties arise in establishing a defect in the model, proving a causal link between the algorithmic conclusion and the judge's decision, and determining the party at fault. As noted by M. Veale & F.Z. Borgesius (2021), the regulatory framework of the Artificial Intelligence Act (2024) does not create a separate regime of tort liability, but lays the groundwork for its specialisation.

Given the autonomy and complexity of high-risk systems, it is possible to draw an analogy with liability for sources of increased danger. Such an approach would minimise the burden of proof on the victim, establish a presumption of liability on the part of the system operator, and ensure more effective protection of rights, as well as strike a balance between technological innovation and the effective protection of individual rights, particularly in cases where an algorithmic decision directly affects the realisation of the right to a fair trial.

This is particularly relevant in the judicial sphere, where the state may act as the operator of a high-risk system. The chain of creation and use of the algorithm involves the developer, supplier, integrator, user (court, lawyer) and the state as the administrator of justice. The model of distributed liability allows for the specific nature of the algorithmic environment to be taken into account and is consistent with the approach being developed

within European digital policy (Floridi *et al.*, 2018). If an algorithm is used in judicial proceedings, the question arises of the state's liability for violations of the right to a fair trial, as guaranteed by the practice of the European Court of Human Rights. In the event of algorithmic opacity, the impossibility of challenging algorithmic influence, and excessive reliance by the court on automated recommendations, there may be a breach of the standards of Article 6 of the European Convention on Human Rights (1950). Thus, an algorithmic error in the field of justice potentially transforms into a constitutional and legal problem. This section employs the following approaches to legal analysis: identification of regulatory gaps, analysis of legal certainty, inconsistencies between legislation and international standards, and the specifics of law enforcement.

The use of artificial intelligence in the legal sphere requires a clear definition of the limits of admissibility, taking into account the principles of the rule of law, fair trial and professional responsibility. Such limits are formed through a combination of a functional approach (the auxiliary nature of AI), risk-based regulation and procedural guarantees for the protection of human rights. In this context, a regulatory gap is evident, as the legislation does not define the procedural status of algorithmic conclusions or the limits of their use. At the same time, there is a problem of legal certainty, as the absence of clear criteria for the application of AI complicates the predictability of legal consequences. The use of algorithmic systems in judicial activities is permissible within the scope of auxiliary functions. In particular, artificial intelligence technologies may be used for the analysis of judicial practice, the automation of document management, the search for relevant court decisions, and the statistical analysis of the workload on the courts. Such tools contribute to improving the efficiency of the judicial system, reducing the time taken to hear cases, and optimising administrative processes. At the same time, the final judicial decision must remain the exclusive prerogative of a human judge. This approach is consistent with the concept of human oversight enshrined in the Artificial Intelligence Act, which provides for mandatory human control over high-risk systems. Delegating the function of justice to an algorithm would contradict the principles of judicial independence, individualisation of liability and guarantees of a fair trial (Artificial Intelligence Act, 2024).

In the legal profession, the use of AI is permissible as a tool for professional support. This includes the automated analysis of contracts, the prediction of legal risks, the conduct of legal research, and the verification of a client's activities against established regulatory requirements. Such technologies are capable of significantly accelerating the processing of large volumes of information and improving the quality of preparatory work. However, the complete delegation of professional legal assessment to an algorithmic system without verification by a lawyer is unacceptable. This contravenes standards of due professional care, the principle of a lawyer's personal responsibility, and the requirements of legal ethics. AI can serve only as an auxiliary tool, but not as a substitute for professional judgement. This also highlights a regulatory gap regarding the definition of the limits of professional liability when using AI.

The Artificial Intelligence Act introduces a risk-based model for regulating artificial intelligence systems. It provides for the classification of systems into four categories:

unacceptable risk, high risk, limited risk and minimal risk. This approach allows for the differentiation of regulatory requirements depending on the potential impact of the technology on human rights and freedoms. AI systems used in the justice sector fall into the high-risk category. This means that a compliance assessment must be carried out, risk management procedures must be implemented, technical and organisational transparency must be ensured, and audits and continuous human oversight must be in place. At the same time, this highlights the inconsistency of national legislation with international standards, in particular with the approaches enshrined in European Union law (Artificial Intelligence Act, 2024).

At the same time, the regulatory framework of the Artificial Intelligence Act (2024), despite its progressive nature, does not resolve a number of conceptual and practical uncertainties. In particular, the question of the limits of the actual autonomy of high-risk systems in the context of judicial activities remains open, as the regulation primarily sets out requirements for design and risk management in advance, but does not establish a sufficiently clear mechanism for retroactive liability for algorithmic errors. Furthermore, the risk-based approach enshrined in the Regulation does not always take into account the cumulative effect of the interaction of several algorithmic systems within a single process, which may lead to the fragmentation of liability and complicate its legal attribution.

The integration of algorithms into legal practice requires not only substantive but also procedural regulation. This involves determining the procedural status of algorithmic conclusions, guarantees of the parties' access to information on the use of AI, and mechanisms for challenging the algorithmic influence on the outcome of a case. The use of algorithmic systems in court proceedings raises a number of fundamental questions: does a party have access to information about the algorithm, can they request an expert examination of the algorithmic model, and does the algorithmic conclusion form part of the evidence in the case? This highlights problems with the application and implementation of legislation, as the relevant procedural mechanisms remain insufficiently regulated.

Under Article 22 of the General Data Protection Regulation (2016), a person has the right not to be subject to a decision based solely on automated processing which produces legal effects concerning them. However, in the context of judicial proceedings, this provision is insufficient. It is necessary to establish specific procedural safeguards: the right to be informed of the use of AI in the case, the right to challenge algorithmic influence, and the right to receive an explanation of the system's logic within limits compatible with the protection of trade secrets.

The academic literature also emphasises that the "right to an explanation" is not absolute, but it follows from the general principles of transparency and accountability of algorithmic systems (Wachter *et al.*, 2017). The integration of artificial intelligence algorithms into legal practice entails not only a transformation of regulatory frameworks but also significant changes in the behaviour of law enforcement actors, the structure of legal institutions, and the nature of interaction between parties to legal relationships. In this context, artificial intelligence should be viewed as a socio-technical phenomenon, the impact of which extends beyond the purely technological or legal sphere and encompasses

behavioural, institutional and cultural aspects of the functioning of the legal system.

Further empirical evidence of the impact of algorithmic systems on the justice sector is provided by the use of the COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) system in the US criminal justice system, which is used to assess the risk of reoffending and to assist judges in making decisions regarding preventive measures and sentencing. Research has shown that algorithmic risk assessments can demonstrate both potential effectiveness in standardising decisions and problems related to the opacity of the model and possible biases in predicting recidivism, which calls into question their neutrality and objectivity (Angwin *et al.*, 2016; Dressel & Farid, 2018). At the same time, empirical analysis indicates that judges do not always critically re-evaluate such recommendations, confirming the existence of an algorithmic trust effect and the potential influence of algorithms on judicial discretion. Thus, experience with COMPAS demonstrates that the integration of algorithms into criminal justice can simultaneously increase the predictability of decisions and give rise to new risks to the principles of transparency, equality and fair trial. Empirical studies also confirm that the use of algorithmic tools in criminal justice can influence judges' perception of risk and the severity of punishment, altering the structure of decision-making (Dressel & Farid, 2018).

One of the key socio-legal effects is the transformation of judges' behaviour in the decision-making process (Sourdin, 2018; Dressel & Farid, 2018). The use of algorithmic systems that generate recommendations or predict case outcomes can lead to automation bias (Parasuraman & Manzey, 2010), i.e. a tendency to favour machine-generated conclusions even when there are grounds for critical evaluation (Parasuraman & Riley, 1997; Goddard *et al.*, 2012). This tendency is explained by the perception of algorithms as objective and neutral tools, which, however, does not always correspond to reality given that results depend on the quality of training data and the underlying models. Consequently, there is a risk of a gradual shift in a judge's internal conviction from autonomous analysis to the indirect adoption of algorithmic recommendations, which may affect the implementation of the principle of judicial independence. Separately, the phenomenon of algorithmic authority (Logg *et al.*, 2019) should be highlighted, which consists of the tendency of law enforcement actors to place heightened trust in results generated by algorithmic systems, even in the absence of a full understanding of their methodology (Citron, 2008).

Furthermore, one must take into account the emergence of a new phenomenon, "institutional algorithmicisation of trust", whereby the legitimacy of law enforcement decisions begins to shift in part from judicial institutions to the technological systems involved in preparing such decisions. This may lead to a gradual shift in the source of authority within the legal system: from personified judicial discretion to the technically mediated expertise of algorithms. Such a transformation is not neutral, as it influences perceptions of justice and may alter the behavioural patterns of those involved in the process, in particular by reducing the level of critical scrutiny of automated recommendations.

No less significant is the impact of artificial intelligence on the professional activities of solicitors and other legal practitioners. Algorithmic tools that provide automated

analysis of contracts, case law or the assessment of legal risks are changing the nature of legal work, reducing the proportion of routine tasks and increasing the role of strategic thinking. At the same time, there is a risk of excessive reliance on technological solutions, which may lead to a decline in the level of critical analysis and professional prudence. In such a case, the content of the standard of due legal care also changes, as the lawyer must not only apply legal norms but also assess the reliability and limitations of the algorithmic tools used (Remus & Levy, 2017).

The socio-legal dimension is also evident in the impact of artificial intelligence on access to justice. On the one hand, the automation of legal processes and the use of algorithmic systems help to reduce the cost of legal services, simplify procedures and speed up case processing, thereby expanding access to legal aid for broad sections of the population. On the other hand, there is a risk of new forms of inequality emerging, linked to the digital divide, unequal access to technology and differences in digital literacy levels. Thus, the introduction of AI can simultaneously act as both a tool for inclusion and a factor of social differentiation in the justice sector (Susskind, 2019; OECD, 2020).

The issue of transforming the institutional structure of legal practice warrants particular attention. The introduction of algorithmic systems facilitates the emergence of new participants in legal relations, including software developers, legal technology providers, algorithmic audit specialists and regulatory bodies responsible for overseeing the use of AI. This leads to an expansion of the traditional circle of participants in legal practice and the formation of a multi-level system of liability, in which legal, technical and organisational aspects are closely intertwined. At the same time, such a transformation complicates the definition of the boundaries of liability and requires new mechanisms for coordination between different institutional actors.

Furthermore, the use of artificial intelligence affects the level of trust in the judicial system and legal institutions in general. The use of algorithmic tools may increase trust through greater consistency and predictability of decisions; however, in the absence of transparency or clarity regarding the algorithms, the opposite effect is possible – a reduction in the legitimacy of justice. In this context, ensuring the explainability of algorithmic decisions, access to information about their use, and the possibility of effectively challenging their impact on the outcome of a case takes on particular significance.

Thus, the socio-legal consequences of applying artificial intelligence algorithms in legal practice manifest themselves in changes to the behavioural patterns of legal actors, the transformation of professional standards, a rethinking of access to justice, and the increasing complexity of the institutional structure of the legal system. This points to the need to integrate a sociological approach into the analysis of algorithmic systems, which allows for consideration not only of their formal legal characteristics but also of their actual impact on the functioning of law as a social institution.

The effective and safe implementation of artificial intelligence algorithms in legal practice is impossible without a comprehensive combination of key principles, institutional and regulatory mechanisms, procedural safeguards and shared responsibility. Table 1 systematises a model for the responsible use of AI, demonstrating how technological tools integrate with legal standards and professional ethics. Each level of the model reinforces the previous one: from ensuring

human oversight and transparency to implementing regulatory and institutional procedures, establishing procedural safeguards, and delineating responsibility between the developer, supplier, integrator, user and the state. This structure allows for the minimisation of risks of algorithmic errors, ensures compliance with the General Data Protection Regulation (2016) and the Artificial Intelligence Act, and creates

the conditions for the responsible, ethical and lawful application of AI in the justice sector. At the same time, current regulation remains fragmented and does not ensure systematic coherence between the substantive, procedural and technological aspects of AI use. A synthesis of these approaches allows for the formulation of a model for the responsible use of AI in legal practice (see Table 1).

Table 1. Model for the responsible use of AI in legal practice

Level	Actors / Components	Main role / Responsibility
Key principles	Human oversight, transparency, accountability, algorithmic audit, ethical standards	Guarantee that decisions remain under human control, transparent algorithm, responsible parties are known, ethics are upheld
Institutional and regulatory mechanisms	EU regulation, national mechanisms, system register, certification, professional standards (codes, training, insurance)	Ensuring high-risk systems comply with legal standards, standardisation of processes, professional training
Procedural safeguards	Parties to the proceedings (court, lawyer, participants), appeals, court decisions	Right to information, access to expert evidence, possibility of challenging algorithmic influence, obligation to give reasons for a court decision
Shared responsibility	Developer → supplier → integrator → user (court/lawyer) → state	Responsibility for algorithm quality, compliance with the General Data Protection Regulation (2016), integration, ethical use, the administration of justice and oversight

Source: developed by the authors based on the General Data Protection Regulation (2016), the Artificial Intelligence Act (2024), the case law of the European Court of Human Rights (Judgment of the European Court of Human Rights in Case No. 47143/06, 2015; Judgment of the European Court of Human Rights in Cases Nos. 58170/13, 62322/14, 24960/15, 2021) and academic approaches to algorithmic governance and responsible AI (Wachter *et al.*, 2017; Floridi *et al.*, 2018; Veale & Borgesius, 2021)

The proposed model is original in nature and synthesises a combination of normative, procedural and ethical approaches to the use of artificial intelligence in legal practice. The socio-legal implications and institutional transformations of AI use. The issue of proving algorithmic error requires clear regulatory framework. First and foremost, it is necessary to determine the procedural status of an algorithmic conclusion – whether it will be considered as independent evidence, supporting material, or merely an analytical tool. At the same time, criteria for its admissibility and relevance must be established, as well as the procedure for appointing and conducting technical expert assessments of algorithmic systems. In the law of evidence, an algorithmic result may take three procedural forms: supplementary analytical material, an expert opinion subject to technical verification, or digital evidence provided its status is specifically enshrined in law. In this context, the burden of proof in disputes concerning algorithmic decisions should shift to the party implementing or operating the system, as it is this party that possesses the technical information regarding its functioning. Procedurally, this means that an algorithmic opinion can be equated with electronic evidence only provided it is reproducible, verifiable and capable of independent verification by a third party.

It would be advisable to introduce the institution of algorithmic forensic examination, aimed at a comprehensive review of the system's functioning. In practical terms, such an examination involves an algorithmic audit, which includes checking training data, identifying statistical biases, assessing the model's stability, and testing for discriminatory effects. As an alternative to forensic examination, an independent algorithmic audit could be carried out by certified private or regulatory bodies during the preliminary system review stage.

A court decision must clearly state the fact that an algorithmic tool was used and describe its role in forming the judge's internal conviction. Furthermore, it is necessary to justify why the algorithmic recommendation was accepted or rejected, in order to ensure the clarity and accountability of the judicial process. This approach is consistent with the requirements for the reasoning of court decisions in the practice of the ECtHR and ensures the traceability of the influence of algorithmic systems on judicial discretion. In the practice of the ECtHR, the requirement for a reasoned judicial decision includes the duty to ensure the clarity of the sources of evidential information; however, there is currently no direct case law regarding algorithmic systems, which creates a gap in the standards of a fair trial.

Regulating the use of artificial intelligence in the legal sphere requires a combination of European standards, national regulations and professional initiatives. It encompasses both requirements for the technical and procedural security of algorithmic systems and the safeguarding of human rights, ethical standards and the accountability of professional communities. At EU level, such requirements include fundamental rights impact assessments, risk management and technical documentation for high-risk systems. The application of a risk-based approach entails not only the classification of systems but also the differentiation of users' procedural obligations, in particular the obligation to document and log algorithmic decisions.

A comprehensive regulatory framework is taking shape within the European Union, which includes requirements regarding data quality and representativeness, mandatory technical documentation, fundamental rights impact assessments, and mechanisms for public oversight (Cath *et al.*, 2018). The Council of Europe's recommendations emphasise the priority of human dignity and non-discrimination

in the application of algorithmic systems, highlighting the need to combine technological efficiency with legal safeguards. However, these standards are not directly binding rules and require transformation into procedural safeguards, in particular in the form of the parties' right of access to algorithmic information. At the same time, these recommendations are of a programmatic nature and require implementation through national procedural codes and subordinate legislation defining algorithmic admissibility.

National legal systems must adapt procedural codes, rules on electronic evidence, rules of professional conduct for lawyers, and mechanisms of civil liability for harm caused by algorithmic decisions. Such adaptation entails the procedural enshrinement of algorithmic evidence, the parties' right to access information about the model, and the possibility of initiating its expert examination. At the same time, the effectiveness of procedural safeguards depends on a clear definition of the standard of proof for algorithmic error, which is currently absent in most legal systems. The effectiveness of such changes depends on the establishment of a procedural mechanism for the allocation of the burden of proof in disputes concerning algorithmic decisions.

Professional communities can develop their own codes of conduct for the use of legal technologies, standards for algorithmic transparency, internal audit procedures, and mechanisms to enhance lawyers' digital competence. Such an approach fosters a culture of responsible AI use and provides additional safeguards for clients and judicial bodies. From a comparative law perspective, such ethical standards have already been implemented in the practice of certain European bar associations, which require mandatory disclosure of the use of AI tools in professional practice. From a doctrinal perspective, such codes serve a compensatory function in an area where state regulation is insufficiently detailed.

For the practical implementation of the provisions of the Artificial Intelligence Act (2024), a systematic adaptation of the national legal environment is required, encompassing legislative, institutional and professional mechanisms. It is advisable to provide for amendments to procedural codes regarding the procedure for using AI, the recording of algorithmic intervention and guarantees of the parties' rights, as well as to regulate the status of algorithmic conclusions. Furthermore, it is necessary to introduce a mandatory assessment of algorithmic impact for high-risk systems, which will ensure that technologies comply with safety and human rights standards. In particular, the principle of "risk classification + human oversight" is key, as it provides for the differentiation of responsibilities depending on the degree of the system's potential impact on an individual's rights.

At the level of state institutions, it is proposed to establish a national supervisory body for high-risk AI systems, maintain a register of algorithmic systems used in the justice sector, and introduce a certification mechanism for legal technology providers. The relevant conformity assessment requirements are set out in the Artificial Intelligence Act. Such institutional models are consistent with the concept of regulatory control, which envisages a combination of state supervision and technical certification as a single control mechanism. Such bodies perform a regulatory control function, but do not replace judicial review, thereby creating a dual model of supervision.

Professional bodies should update their codes of professional conduct, introduce compulsory training in algorithmic

literacy, and provide for liability insurance when using AI. The concept of "trustworthy and reliable artificial intelligence", developed by the European Commission, can serve as a regulatory framework for establishing relevant standards and ensure a balance between technological innovation and compliance with professional and ethical standards. To ensure the effective integration of innovative technologies with the protection of human rights, it is advisable to implement a set of measures, including mandatory algorithmic impact assessments, the principle of human oversight, ensuring algorithmic traceability, enhancing the digital competence of lawyers, professional liability insurance for the use of AI, and independent external audits of high-risk systems. The concept of responsible AI should be based on the integration of legal, technical and ethical standards, which allows for the minimisation of risks and an increase in trust in technologies within the justice sector (Floridi *et al.*, 2018). The regulation of AI in the legal sphere is taking on a multi-layered structure that combines procedural, institutional and ethical components, yet remains fragmented due to the absence of a unified codified approach.

An in-depth analysis demonstrates that the effective use of AI algorithms in legal practice requires the introduction of a specific model of distributed civil liability, the establishment of procedural safeguards for algorithmic transparency, the institutionalisation of algorithmic auditing (Raji *et al.*, 2020), the harmonisation of national legislation with the provisions of the Artificial Intelligence Act (2024) and the General Data Protection Regulation (2016), as well as the constitutionalisation of the principle of human oversight in the justice system. In summary, it can be stated that the current regulation of algorithmic systems lies at the intersection of three models: normative, procedural (judicial proceedings) and ethical (professional codes), yet their integration remains incomplete. Thus, the integration of algorithmic systems into legal practice requires not only adherence to general ethical principles, but also the establishment of a specific procedural framework for their use, which includes the verifiability of algorithmic conclusions, the possibility of challenging them, and a clear division of responsibility among the parties involved.

Conclusions

The subject of the study was the legal, ethical and institutional aspects of the application of artificial intelligence algorithms in legal practice, particularly in the fields of justice, legal practice and legal analysis. The aim of the article was to define the limits of permissible AI use in law and to analyse the key risks of its implementation, including issues of algorithmic liability and the protection of human rights. This aim was achieved through a comprehensive study of the European Union's regulatory framework, the practice of the European Court of Human Rights, and contemporary academic approaches to algorithmic regulation. The research findings are organised into three main areas: a theoretical and legal analysis of algorithmic decision-making, an assessment of the practice of applying AI in legal activities, and the identification of models of legal liability for the consequences of using algorithmic systems.

The study identifies a scientific novelty, which lies in the conceptualisation of algorithmic liability as a multi-party model encompassing developers, users of algorithmic systems,

and the state as the institutional guarantor of human rights in the field of justice. It is argued that such liability goes beyond the scope of classical tort law constructs and requires the application of a comprehensive regulatory approach based on the principles of risk-based regulation, algorithmic accountability and human oversight. It is demonstrated that the European model of artificial intelligence regulation forms a new paradigm of legal regulation of technologies, in which priority is given to preventive risk management rather than post-hoc liability.

It has been established that the application of artificial intelligence algorithms in the legal sphere is accompanied by a number of systemic risks, including algorithmic bias, lack of transparency in decision-making (the “black box effect”), the reproduction of social discrimination, and potential restrictions on the principle of judicial independence. The findings are consistent with current research in the field of algorithmic justice, which confirms the existence of the automation bias effect and growing trust in algorithmic recommendations in the legal decision-making process. An analysis of EU regulatory frameworks, in particular the Artificial Intelligence Act and the General Data Protection Regulation, has revealed the emergence of a risk-based regulatory model and the enshrinement of the principle of human oversight over high-risk systems. The research findings indicate that the practice of the European Court of Human Rights establishes standards of transparency, predictability and procedural safeguards that are relevant to the assessment of the use of algorithmic systems in the administration of justice. These findings point to the need to limit the autonomy of algorithms in judicial proceedings and to preserve the judge’s key role in reaching the final decision. It was also established that the issue of liability for the consequences of algorithmic decisions cannot be fully resolved within the framework of traditional tort

law. The findings suggest that algorithmic liability is of an inter-subject nature and requires consideration of the roles of all participants in the technological chain – from the developer to the state as the organiser of justice. It is also noted that the digitalisation of justice creates conditions for expanding the analytical capabilities of legal practice, yet does not alter the fundamental principle of human decision-making in the courts.

Summarising the research findings, it has been established that the integration of artificial intelligence into legal practice shapes a new legal reality in which classical models of law enforcement are transformed under the influence of digital technologies. Such a transformation does not negate the fundamental principles of the rule of law, but requires their adaptation to the conditions of the algorithmic environment through the introduction of clear legal boundaries for the use of AI, the strengthening of procedural safeguards, and the development of institutions for algorithmic control. A promising approach is the development of a hybrid model of justice, in which algorithmic systems perform a supporting function, whilst the final decision remains with humans as the bearers of legal responsibility and guarantors of justice.

Further research should focus on empirical studies of the impact of algorithmic systems on the behaviour of judges, lawyers and other participants in the legal process, particularly in the context of algorithmic trust and algorithmic bias.

Acknowledgements

None.

Funding

The study was not funded.

Conflict of interest

None.

References

- [1] Angwin, J., Larson, J., Mattu, S., & Kirchner, L. (2016). Machine bias. *ProPublica*. Retrieved from <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>.
- [2] Baranov, O.A. (2023). Definition of the term “artificial intelligence”. *Information and Law*, 1(44), 32-49. doi: 10.37750/2616-6798.2023.1(44).287537.
- [3] Barfield, W. (Ed.). (2020). *The Cambridge handbook of the law of algorithms*. Cambridge: Cambridge University Press. doi: 10.1017/9781108680844.
- [4] Barocas, S., & Selbst, A.D. (2016). *Big data’s disparate impact*. *California Law Review*, 104(3), 671-732.
- [5] Calo, R. (2015). *Robotics and the lessons of cyberlaw*. *California Law Review*, 103(3), 513-563.
- [6] Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). Artificial intelligence and the “good society”. *Science and Engineering Ethics*, 24(2), 505-528. doi: 10.1007/s11948-017-9901-7.
- [7] Citron, D.K. (2008). *Technological due process*. *Washington University Law Review*, 85(6), 1249-1313.
- [8] Davenport, T.H., & Ronanki, R. (2018). *Artificial intelligence for the real world*. *Harvard Business Review*, 96(1), 108-116.
- [9] Dressel, J., & Farid, H. (2018). The accuracy, fairness, and limits of predicting recidivism. *Science Advances*, 4(1), article number eaao5580. doi: 10.1126/sciadv.aao5580.
- [10] Edwards, L. (2022). *Regulating AI in Europe: Four problems and four solutions*. Retrieved from <https://www.adalovelaceinstitute.org/report/regulating-ai-in-europe/>.
- [11] European Commission High-Level Expert Group on AI. (2019). *Ethics guidelines for trustworthy AI*. Retrieved from <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- [12] European Convention on Human Rights. (1950, November). Retrieved from <https://www.echr.coe.int/en/web/echr/european-convention-on-human-rights>.
- [13] Floridi, L., et al. (2018). AI4People – an ethical framework. *Minds and Machines*, 28(4), 689-707. doi: 10.1007/s11023-018-9482-5.
- [14] Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias in clinical decision support systems. *Artificial Intelligence in Medicine*, 57(3), 193-200. doi: 10.1016/j.artmed.2012.01.001.
- [15] Hildebrandt, M. (2020). *Law for computer scientists and other folk*. Oxford: Oxford University Press. doi: 10.1093/oso/9780198860877.001.0001.

- [16] Judgment of the European Court of Human Rights in the Case No. 47143/06 “Zakharov v. Russia”. (2015, December). Retrieved from <https://hudoc.echr.coe.int/eng?i=001-155159>
- [17] Judgment of the European Court of Human Rights in the Case No. 58170/13, 62322/14, 24960/15 “Big Brother Watch and Others v. the United Kingdom”. (2021, May). Retrieved from <https://hudoc.echr.coe.int/eng?i=001-210077>.
- [18] Logg, J.M., Minson, J.A., & Moore, D.A. (2019). Algorithm appreciation. *Organizational Behavior and Human Decision Processes*, 151, 90-103. doi: 10.1016/j.obhdp.2018.12.005.
- [19] McGinnis, J.O., & Pearce, R.G. (2014). *The great disruption*. *Fordham Law Review*, 82(6), 3041-3066.
- [20] O’Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. New York: Crown Publishers. doi: 10.5860/crl.78.3.403.
- [21] OECD. (2020). *OECD principles on AI*. Retrieved from <https://www.oecd.org/en/topics/ai-principles.html>.
- [22] Pagallo, U. (2013). *The laws of robots*. New York: Springer. doi: 10.1007/978-94-007-6564-1
- [23] Parasuraman, R., & Manzey, D.H. (2010). Complacency and bias in automation. *Human Factors*, 52(3), 381-410. doi: 10.1177/0018720810376055.
- [24] Parasuraman, R., & Riley, V. (1997). Humans and automation. *Human Factors*, 39(2), 230-253. doi: 10.1518/00187209778543886.
- [25] Pasquale, F. (2015). *The black box society*. Cambridge: Harvard University Press.
- [26] Raji, I.D., et al. (2020). Closing the AI accountability gap. *arXiv*. doi: 10.48550/arXiv.2001.00973.
- [27] Regulation of European Union No. 2016/679 “General Data Protection”. (2016, April). Retrieved from <https://eur-lex.europa.eu/eli/reg/2016/679/oj>.
- [28] Regulation of European Union No. 2024/1689 “Artificial Intelligence Act”. (2024, June). Retrieved from <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>.
- [29] Remus, D., & Levy, F. (2017). *Can robots be lawyers?* *Georgetown Journal of Legal Ethics*, 30, 501-558.
- [30] Sourdin, T. (2018). *Judge v robot?* *University of New South Wales Law Journal*, 41(4), 1114-1133.
- [31] Susskind, R. (2019). *Online courts and the future of justice*. Oxford: Oxford University Press. doi: 10.1093/oso/9780198838364.001.0001.
- [32] Veale, M., & Borgesius, F.Z. (2021). *Demystifying the AI Act*. *Computer Law Review International*, 22(4), 97-112.
- [33] Wachter, S., Mittelstadt, B., & Floridi, L. (2017a). Transparent AI. *Science Robotics*. doi: 10.1126/scirobotics.aan6080.
- [34] Wachter, S., Mittelstadt, B., & Floridi, L. (2017b). Why a right to explanation does not exist in GDPR. *International Data Privacy Law*, 7(2), 76-99. doi: 10.1093/idpl/ix005.
- [35] Yeung, K., & Ulbricht, L. (2022). Understanding algorithmic regulation. *Regulation & Governance*, 16(1), 3-22. doi: 10.1111/rego.12437.

Алгоритми штучного інтелекту у правничій діяльності: правові межі, соціально-правові наслідки та відповідальність

Юрій Лукашук

Доктор філософії в галузі комп'ютерних наук, асистент
Національний університет «Львівська політехніка»
79000, вул. Степана Бандери, 12, м. Львів, Україна
<https://orcid.org/0000-0002-8933-8635>

Уляна Цмоць

Старший науковий співробітник
Львівський державний університет внутрішніх справ
79007, вул. Городоцька, 26, м. Львів, Україна
<https://orcid.org/0000-0002-8423-8272>

Анотація. Актуальність дослідження зумовлена стрімкою цифровізацією правничої діяльності та зростанням ролі алгоритмів штучного інтелекту у процесах ухвалення юридичних рішень, що породжує нові правові та етичні виклики. Метою роботи було вивчення перетину правових, етичних та інституційних аспектів застосування штучного інтелекту в правничій сфері, а також визначення меж допустимого використання алгоритмічних систем і підходів до розподілу відповідальності. У дослідженні використано системний, порівняльно-правовий, формально-юридичний методи, а також аналіз судової практики. Досліджено ризики автоматизованого ухвалення рішень: алгоритмічну упередженість, проблему «чорної скриньки», відтворення дискримінації та загрози незалежності судді. Ці ризики підтверджуються емпіричними дослідженнями у сфері алгоритмічного правосуддя та аналізом практики використання систем прогнозування ризиків у кримінальному судочинстві (наприклад, COMPAS у США). Проаналізовано регулювання ШІ на прикладі Закону ЄС про штучний інтелект і GDPR, а також підходи ЄСПЛ до забезпечення прозорості та процесуальних гарантій. Узагальнено підходи до розподілу відповідальності між розробниками, постачальниками, користувачами й державою та запропоновано модель алгоритмічної прозорості з людським контролем і правом сторін на поінформованість. Розглянуто міжнародний досвід США, Франції, Сінгапуру та Китаю, а також перспективи впровадження ШІ в Україні в межах ЄСІТС і систем аналізу судової практики. Особливу увагу приділено концепції розподіленої алгоритмічної відповідальності та застосуванню аналогії з відповідальністю за джерело підвищеної небезпеки до високоризикових систем ШІ. Запропоновано модель алгоритмічної прозорості з правом сторін на повідомлення про використання ШІ, алгоритмічну експертизу та обов'язком суду мотивувати врахування рекомендацій ШІ. Окреслено принцип «людської помилки» як гарантію збереження судового розсуду й незалежності судді. Результати дослідження можуть бути використані для вдосконалення національного законодавства у сфері цифрового правосуддя та впровадження стандартів алгоритмічної прозорості

Ключові слова: алгоритмічна відповідальність; етика, дискримінація; цифрове правосуддя; Закон ЄС про штучний інтелект; Загальний регламент про захист даних