

Konstantinos Kouroupis, PhD, Associate Professor
Depart. of Law, Frederick University, Cyprus
email: law.ck@frederick.ac.cy
0035799114989

George Christofides, PhD candidate
Depart. of Law, Special Teaching Staff, Frederick University, Cyprus
email: gio@legalpartners.cy.net
0035799689355

AI FOR JUSTICE OR JUSTICE FOR AI? A CRITICAL ANALYSIS OF ALGORITHMIC DECISION-MAKING PROCESS IN JUDICIAL SYSTEMS

Abstract:

This paper explores the evolving interplay between Artificial Intelligence (AI) and justice, critically examining whether AI serves justice or if justice is being reshaped to accommodate AI. Central to this analysis is the European Union's AI Act, which establishes a comprehensive regulatory framework aimed at governing the use of AI technologies, particularly within high-stakes domains like the judicial system. The paper investigates the transformative impact of AI on the deliverance of justice, assessing both the efficiency gains and the significant risks associated with its deployment in legal decision-making processes. Particular attention is given to ethical concerns arising from algorithmic bias, opacity in decision-making, and the potential erosion of fundamental legal principles such as fairness, accountability, and human oversight. A comparative methodology is employed to analyze the approaches of the European Union and the United States, highlighting differences in regulatory philosophy, institutional safeguards, and public discourse around AI in justice. Special emphasis is given on the study of two significant cases in which justice is delivered through the deployment of AI tools: the first one refers to Rowicz case (EU law) and the second one to Loomis case (international law). While the EU adopts a precautionary and rights-based approach, the U.S. leans toward innovation-driven, market-led regulation. This transatlantic comparison underscores the urgent need for harmonized legal standards that uphold justice in an era increasingly mediated by algorithms, questioning whether AI can truly serve justice—or if legal systems must adapt to ensure justice remains human-centered.

Key Words : Justice, Artificial Intelligence, transparency, accountability, Loomis case, Rowicz.

1. INTRODUCTION

The rapid integration of Artificial Intelligence into judicial systems worldwide presents a profound question that cuts to the heart of legal philosophy: Are we deploying AI to serve justice, or are we gradually reshaping justice to accommodate AI? This fundamental tension between technological capability and judicial integrity forms the central thesis of our contemporary legal landscape's most pressing challenge.

Artificial Intelligence has undoubtedly emerged as a transformative technology with applications reaching into virtually every aspect of human society. From healthcare diagnostics to educational personalization, AI promises efficiency, consistency, and data-driven decision-making¹. Yet when this technology enters the sacred halls of justice, we must ask whether these promises align with the fundamental principles upon which our legal systems are built².

The justice sector's adoption of AI reflects a broader societal belief in technology's ability to solve complex problems, but unlike other fields focused on efficiency, justice systems bear the unique responsibility of upholding rights, due process, and public trust. This analysis explores whether AI in judicial contexts truly advances "AI for Justice"—technology serving fair and effective legal administration—or instead signals "Justice for AI," where legal norms are reshaped to fit technological limits and commercial pressures. By examining key cases like *State v. Loomis* (2016) in the U.S. and the EU's pending *Case C-159/2025 (Rowicz)*, we assess whether courts are adapting AI to meet justice standards or compromising justice to accommodate AI.

SECTION I: THE LOOMIS CASE.

THE FACTS AND CONSTITUTIONAL CHALLENGE.

The case of *State v. Loomis*³ presents a stark example of how the question "AI for Justice or Justice for AI?" plays out in practice. Brandon Loomis, convicted of a felony in Wisconsin in 2014, became the unwitting subject of a legal precedent that would fundamentally alter the relationship between algorithmic systems and judicial decision-making in American criminal justice⁴.

The facts reveal how AI systems can fundamentally alter the trajectory of individual lives through opaque decision-making processes. Loomis, despite denying involvement

1 AI technology may offer multiple services. See Negnevitsky, M. (2020). *Artificial intelligence: A guide to intelligent systems* (3rd ed.). Addison Wesley.

2 Regarding the relationship between AI and law see Byuers, J. (2018). *Artificial intelligence – The practical legal issues*. Legal Brief Publishing. Ashley, K. (2017). *Artificial intelligence and legal analytics: New tools for practice in the digital age*. Cambridge University Press.

3 *State v. Loomis*, 881 N.W.2d 749 (Wis. 2016). <https://www.wicourts.gov/sc/opinion/DisplayDocument.pdf?content=pdf&seqNo=171690>

4 See a detailed presentation of Loomis Case in : Harvard Law Review. (2017, March 21). Wisconsin Supreme Court requires warning before use of algorithmic risk assessments in sentencing. *Harvard Law Review*, 130. <https://harvardlawreview.org/print/vol-130/state-v-loomis/>, Smith, M. (2016, June 22). In Wisconsin, a backlash against using data to foretell defendants' futures. *The New York Times*. <https://www.nytimes.com/2016/06/23/us/backlash-in-wisconsin-against-using-data-to-foretell-defendants-futures.html>

in the crime, accepted a plea deal that left his sentence to judicial discretion. The judge, seeking to make an informed sentencing decision, ordered a Pre-Sentence Investigation that included the COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) risk assessment. This algorithmic system generated scores indicating that Loomis presented high pre-trial risk, high risk of recidivism, and high risk of violent recidivism.

The impact of these algorithmic predictions was dramatic. Instead of the one-year county jail sentence with probation that both prosecution and defense had negotiated, the court imposed seven years with four years of confinement. This substantial departure from the agreed-upon sentence demonstrates how AI systems can effectively override human judgment and negotiated agreements, raising fundamental questions about who or what ultimately controls judicial outcomes.

Loomis's legal challenge to the use of COMPAS in his sentencing presents a clear test of whether the legal system would prioritize traditional constitutional protections or accommodate AI system limitations. His defense argued that relying on the COMPAS algorithm violated his Sixth Amendment rights, particularly his right to confront witnesses and his right to due process, because the algorithm's proprietary nature made it impossible to scrutinize or challenge its assessments.

This challenge represented a crucial moment where the legal system had to choose between two paths: either requiring AI systems to meet established constitutional standards or modifying constitutional interpretations to accommodate AI system limitations. The defense's argument was fundamentally about maintaining traditional legal principles that ensure defendants can understand and challenge evidence used against them.

The Wisconsin Supreme Court's decision reveals which path the legal system chose. The court ruled that the use of COMPAS did not violate constitutional rights, but the reasoning behind this decision is telling. Rather than requiring the AI system to meet transparency standards that would enable meaningful challenge, the court instead modified the constitutional analysis to accommodate the AI system's opacity.

The court's reasoning that COMPAS was "*just one factor among many*" and that "*the Court needed all the help it could get*" suggests a prioritization of AI assistance over constitutional rigor. The Court additionally gave a general overview of the acceptable COMPAS assessments uses (diverting low-risk prison-bound offenders to a non-prison alternative, assessing whether an offender can be supervised safely and effectively in the community and imposing terms and conditions of probation, supervision, and responses to violations)⁵ and highlighted the caution a judge should possess when treating risk assessments, mandating that pre-sentencing investigation (PSI) reports containing COMPAS risk assessments must make certain disclosures to sentencing courts. This "written advisement of its limitations" should explain that:

- i) COMPAS is a proprietary tool, which has prevented the disclosure of specific information about the weights of the factors or how risk scores are calculated;
- ii) COMPAS scores are based on group data, and therefore identify groups with characteristics that make them high-risk offenders, not particular high-risk individuals;

Several studies have suggested that the COMPAS algorithm may be biased in how it classifies minority offenders;

5 881 N.W.2d 768.

- iii) COMPAS compares defendants to a national sample, but has not completed a cross-validation study for a Wisconsin population, and tools like this must be constantly monitored and updated for accuracy as populations change; and
- iv) COMPAS was not originally developed for use at sentencing. On the contrary, it was destined to assist post-sentencing.

This decision effectively creates a legal framework where AI systems can influence judicial decisions without meeting the same transparency and accountability standards required of human testimony or evidence.

The Wisconsin Supreme Court's decision, allowed to stand by the U.S. Supreme Court's denial of certiorari, established a precedent that appears to prioritize "Justice for AI" over "AI for Justice", which is the core question in our case.

This precedent raises profound questions about the direction of legal evolution. On the one hand, it has been proved that there is a 70% chance that any randomly selected higher-risk individual is classified as higher risk than a randomly selected low-risk individual⁶. On the other hand, it is also generally confirmed that risk assessment instruments can predict who is at risk to recidivate with at least some degree of accuracy within societies that are marked by a high degree of diversity – like that of the United States⁷.

The Loomis precedent suggests that when AI capabilities conflict with traditional legal protections, courts may be more willing to adapt legal standards than to require AI systems to meet established constitutional requirements. This adaptation represents a fundamental shift in the balance of power between technology and legal principles, with potentially far-reaching implications for the future of constitutional protections in an AI-driven legal system.

The Loomis decision essentially resolved the transparency paradox by accepting AI decisions without requiring full understanding. This approach prioritizes AI capability over human understanding, suggesting that the benefits of AI decision-making can justify the acceptance of opaque processes. However, this acceptance comes at the cost of traditional accountability mechanisms and may create a two-tiered justice system where some decisions are subject to scrutiny and challenge while others are protected by claims of algorithmic complexity.

SECTION II: THE LOOMIS CASE UNDER THE LIGHT OF THE EU AI ACT

REGULATORY PHILOSOPHY: JUSTICE AS THE STANDARD

The European Union's approach to AI regulation, as embodied in the EU AI Act, represents a markedly different response to the fundamental question of AI's role in justice systems. Rather than adapting legal principles to accommodate AI limitations, the EU AI Act attempts to require AI systems to conform to established legal and ethical principles. This approach suggests a commitment to "AI for Justice" rather than "Justice for AI."

⁶ Lightbourne, J. (2017). Damned lies and criminal sentencing using evidence-based tools. *Duke Law & Technology Review*, 15(1), 1–25.

⁷ James, N. (2015). *Risk and needs assessment in the criminal justice system* (CRS Report No. R44087). Congressional Research Service.

The EU AI Act's approach to high-risk AI systems, particularly those used in law enforcement and judicial decision-making⁸, establishes justice and fundamental rights as the benchmark against which AI systems must be measured. Rather than asking how legal systems can accommodate AI limitations, the Act asks how AI systems can be designed and deployed to serve justice effectively while respecting fundamental rights.

The EU AI Act's transparency requirements directly address the core problem identified in the *Loomis* case: the use of opaque AI systems in contexts where transparency is essential for justice. Article 13's requirement that high-risk AI systems be "sufficiently transparent" to enable users to "interpret the system's output and use it appropriately" represents a clear statement that AI systems must serve human understanding rather than requiring humans to accept AI outputs without comprehension.

The EU AI Act prioritizes the needs of justice over the commercial interests of AI developers by requiring transparency and human oversight in high-risk contexts, including judicial settings. Rather than modifying justice systems to fit opaque AI tools, the Act insists that only transparent and accountable systems be used, addressing concerns raised in cases like *Loomis*, where judges relied on opaque algorithms. Article 14 mandates effective human supervision to preserve human agency in decisions affecting fundamental rights, ensuring AI serves rather than replaces human judgment. Unlike the *Loomis* court, which adapted legal standards to fit AI limitations, the EU AI Act upholds legal and ethical norms, requiring contestability, bias mitigation, and transparency—signaling that AI must adapt to justice, not the other way around.

According to the European leaders, the EU AI Act aims to establish harmonised rules on artificial intelligence. It is part of a wider package of policy measures to support the development of trustworthy AI, which also includes the AI Innovation Package, the launch of AI Factories and the Coordinated Plan on AI⁹. In order to build a powerful and trustful AI it is an imperative need to respect certain legal and ethical guidelines. This approach treats justice and fundamental rights as non-negotiable standards that AI systems must meet.

SECTION III: THE EU AI ACT UNDER EXAMINATION IN THE EU LEGAL SYSTEM: THE CASE C-159/2025 (ROWICZ).

THE CASE AND ITS SIGNIFICANCE

The preliminary ruling request in Case C-159/2025 (Rowicz)¹⁰ presents a crucial test of whether European courts will follow the regulatory philosophy established by the EU AI Act or succumb to the same pressures that led to the *Loomis* decision. The case

8 According to art.6.2 of the EU AI Act and the Annex III, AI systems intended to be used by a judicial authority serving in the administration of justice as well as any AI tools used for purposes of law enforcement are classified as high-risk AI systems. See European Parliament. (n.d.). *Annex III: High-risk AI systems (Artificial Intelligence Act)*. Retrieved June 27, 2025, from <https://artificialintelligenceact.eu/annex/3/>.

9 European Commission. (n.d.). *Regulatory framework for artificial intelligence*. Retrieved June 27, 2025, from <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>.

10 <https://curia.europa.eu/juris/document/document.jsf?jsessionid=AF710376452DA4C6D70CE45E6AF710376452DA4C6D70CE45E6461B292?text=&docid=301557&pageIndex=0&doclang=en&mode=req&dir=&occ=first&part=1&cid=3854517>.

involves Poland's Random Case Allocation System (SLPS), an AI system developed by the Ministry of Justice to assign judges to cases.

The *Rowicz* case highlights a core conflict between the efficiency offered by AI and the constitutional principle of judicial independence. While the SLPS system, developed by Poland's Ministry of Justice, may enhance administrative efficiency and consistency, its executive origin raises serious concerns about the separation of powers. This creates a constitutional dilemma: should the operational benefits of AI justify potential compromises to judicial autonomy?

By allowing the executive branch to control judicial assignments through an AI system, the SLPS blurs institutional boundaries and risks indirect executive influence over the judiciary. Though the algorithm may appear neutral, its development, implementation, and oversight by the executive open the door to impermissible forms of control, undermining traditional safeguards of judicial independence. The *Rowicz* case will test whether European courts will uphold EU AI Act principles, even when doing so could mean restricting widely used AI tools.

More broadly, the case raises concerns about democratic accountability and the concentration of institutional power in the executive. When AI systems impacting justice outcomes are developed outside of judicial oversight, they shift power dynamics and erode traditional checks and balances. The outcome of *Rowicz* will indicate whether the EU's commitment to "AI for Justice" is substantive or merely rhetorical in the face of institutional convenience and executive influence.

SECTION IV: THE LOOMIS CASE VS. ROWICZ CASE.

The comparison between the American approach exemplified by *Loomis* and the European approach embodied in the AI Act and tested in *Rowicz* reveals two fundamentally different philosophies about the relationship between AI and justice.

The *Loomis* decision may be an ideal example of adaptation of legal standards and constitutional interpretations to AI systems, on the basis of the respect of specific requirements. This approach prioritizes the benefits that AI systems can provide to judicial administration, even when those benefits come at the cost of traditional legal protections.

Therefore, AI systems as beneficial tools that should be integrated into legal processes with minimal interference. When conflicts arise between AI capabilities and legal requirements, this approach tends to resolve those conflicts by modifying legal standards rather than requiring changes to AI systems.

This philosophy reflects a faith in technological solutions and a willingness to adapt legal institutions to take advantage of technological capabilities. However, it also raises concerns about the gradual erosion of legal protections and the increasing dependence of legal institutions on commercial AI products that may not be designed with legal principles in mind.

The EU AI Act corresponds to a more concrete, pure legal and human-centric approach in the sense that AI systems shall be adapted to legal and ethical principles rather than expecting legal systems to adapt to AI limitations.

This approach treats justice and fundamental rights as non-negotiable standards that AI systems must meet.¹¹

The aforementioned approach recognizes that AI systems can provide valuable benefits to legal institutions but insists that these benefits cannot come at the expense of fundamental rights or core legal principles. When conflicts arise between AI capabilities and legal requirements, this approach requires changes to AI systems rather than modifications to legal standards.

The contrast between the *State v. Loomis* decision and the European Union's regulatory framework, exemplified in cases like *Rowicz*, reveals fundamentally different philosophies regarding AI deployment in judicial and administrative systems. These differences manifest across four critical dimensions that define the relationship between artificial intelligence and democratic governance.

i) the principle of transparency.

The Wisconsin Supreme Court in *Loomis* effectively normalized algorithmic opacity in judicial decision-making. The court acknowledged that the COMPAS risk assessment tool's proprietary algorithms remained entirely opaque to defendants, judges, and even court officials, yet deemed this acceptable provided certain procedural safeguards were observed. This approach treats algorithmic inscrutability as an inevitable cost of technological advancement, requiring judicial actors to work around rather than through AI systems.

The *Loomis* framework established a troubling precedent by suggesting that due process requirements could be satisfied even when neither the defendant nor the judge could understand how critical risk assessments were generated. The court's reasoning implicitly accepted that commercial interests in protecting trade secrets could outweigh fundamental transparency requirements in criminal justice proceedings¹².

In stark contrast, the EU's regulatory framework, as demonstrated in *Rowicz* and codified in the AI Act, treats explainability not as a desirable feature but as a fundamental prerequisite for AI deployment in high-risk contexts. This approach recognizes that algorithmic transparency is essential for maintaining the rule of law and ensuring that automated systems remain subject to meaningful legal oversight.

The EU framework requires that individuals subject to AI-driven decisions must be able to understand the logic, significance, and consequences of the automated processing. This goes beyond mere notification of AI use to demand substantive explainability that enables meaningful challenge and review. The *Rowicz* case exemplifies this principle by establishing that automated decision-making systems must provide sufficient transparency to allow for effective judicial review and individual redress.

ii) the element of human agency.

11 See a detailed commentary of the EU AI Act in Ceyhun Necati Pehlivan, Forgó, N., & Valcke, P. (2024). *The EU artificial intelligence (AI) act: A commentary*. Kluwer Law International.

12 See about the impact of the principle of transparency in the deliverance of justice in Ryberg, J. (2022). *Sentencing and algorithmic transparency*. Oxford Academic Books., Zerilli, J. (2022). *Algorithmic sentencing: Drawing lessons from human factors research*. Oxford Academic Books., Begby, E. (2021). *Automated risk assessment in the criminal justice process: A case of 'algorithmic bias'?* Oxford Academic Books., Monahan, J., & Skeem, J. (2016). Risk assessment in criminal sentencing. *Annual Review of Clinical Psychology*, 12, 489–513. <https://doi.org/10.1146/annurev-clinpsy-021815-092945>

The *Loomis* decision effectively legitimized a model where AI systems can substantially influence or even override human judgment in critical decisions. While the court nominally required judges to consider additional factors beyond the COMPAS score, the practical reality created a system where algorithmic recommendations carried disproportionate weight precisely because their basis could not be interrogated or challenged.

This approach fundamentally alters the relationship between human decision-makers and automated systems, positioning judges as consumers rather than controllers of algorithmic outputs. The inability to examine the reasoning behind risk assessments means that human oversight becomes largely ceremonial, reduced to either accepting or rejecting recommendations they cannot fully evaluate ¹³.

The EU Framework: Human Agency Principle.

The EU approach, exemplified in *Rowicz*, insists on maintaining meaningful human agency throughout automated decision-making processes. Thus, human operators retain genuine authority over AI systems and possess the information necessary to exercise that authority effectively.

Under this framework, automated systems must be designed to support rather than supplant human judgment. The *Rowicz* case established that meaningful human oversight requires not only the formal authority to override AI recommendations but also the practical capability to do so based on comprehensible information about how those recommendations were generated. This ensures that human decision-makers remain the ultimate arbiters of consequential choices affecting individual rights.

iii) the constitutional approach.

The Loomis Accommodation Model

The Wisconsin Supreme Court in *Loomis* constituted an extraordinarily important decision, because it was the first ever to address the constitutionality of the use of algorithms in sentencing¹⁴. Rather than requiring AI systems to meet established constitutional benchmarks, the court adapted constitutional requirements to fit around technological constraints.

This approach represents a concerning inversion of the traditional relationship between technology and constitutional rights. Instead of technology serving constitutional values, constitutional interpretation was modified to accommodate technological limitations. The court's reasoning suggested that the benefits of AI-assisted decision-making could justify accepting reduced procedural protections, fundamentally altering the constitutional landscape without explicit democratic authorization ¹⁵.

13 See De Miguel Beriain, I. (2018). Does the use of risk assessments in sentences respect the right to due process? A critical analysis of the Wisconsin v. Loomis ruling. *Law, Probability and Risk*, 17(1), 45–53. <https://doi.org/10.1093/lpr/mgy001>, as well as Freeman, K. (2016). Algorithmic injustice: How the Wisconsin Supreme Court failed to protect due process rights in State v. Loomis. *North Carolina Journal of Law & Technology Online*, 18, 75–97. http://ncjolt.org/wp-content/uploads/2016/12/Freeman_Final.pdf

14 De Miguel Beriain, I. (2018). Does the use of risk assessments in sentences respect the right to due process? A critical analysis of the Wisconsin v. Loomis ruling. *Law, Probability and Risk*, 17(1), 45–53. <https://doi.org/10.1093/lpr/mgy001>.

15 See a critical approach of the use of AI tools in justice from a constitutional point of view in O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.,

The EU Approach: Technology Must Meet Constitutional Standards

Contrary to *Loomis* case, the European framework asserts that AI systems must comply with existing human rights and constitutional standards, not the other way around. The *Rowicz* case reflects this by affirming that automated systems cannot be used unless they meet the same procedural and substantive safeguards required of human decision-makers. This approach upholds constitutional values over technological convenience, requiring AI to adapt to legal standards rather than compromising protections for the sake of innovation.

iv) issues of democratic accountability.

The *Loomis* decision reflected a troubling deference to commercial AI products, accepting private companies' assertions about algorithmic effectiveness without requiring independent validation or public oversight. This approach essentially outsourced critical aspects of judicial decision-making to private entities while insulating those entities from meaningful public accountability¹⁶.

The court's acceptance of COMPAS without rigorous public evaluation or ongoing oversight created a precedent for incorporating commercial AI products into governmental functions without the transparency and accountability mechanisms typically required for public institutions. This model treats AI systems as consumer products rather than components of democratic governance subject to public scrutiny and control.

The EU Framework: Public Accountability and Democratic Oversight

The European approach, exemplified in *Rowicz* and formalized by the AI Act, requires AI systems used in government to be subject to strong public accountability measures. These include independent audits, transparent documentation, and oversight by democratically accountable institutions. Recognizing that such systems become part of the democratic process, the EU mandates transparency, performance evaluation, and mechanisms for oversight and correction, ensuring AI in public administration adheres to the same standards as other governmental functions.

Implications for Democratic Governance

The contrasting approaches illustrated by the *Loomis* and *Rowicz* cases reflect fundamentally different conceptions of the relationship between artificial intelligence and democratic governance. The *Loomis* model suggests that democratic institutions must adapt to technological constraints, accepting diminished transparency, reduced human agency, and weaker constitutional safeguards as trade-offs for innovation and efficiency. In contrast, the European approach, exemplified by *Rowicz*, maintains that AI must conform to existing democratic and legal standards—ensuring transparency, human oversight, constitutional compliance, and accountability. This divergence underscores that the

Hamilton, M. (2015). Risk-needs assessment: Constitutional and ethical challenges. *American Criminal Law Review*, 52(2), 231–274.

16 Liu, H.-W., Lin, C.-F., & Chen, Y.-J. (2019). Beyond *State v. Loomis*: Artificial intelligence, government algorithmization, and accountability. *International Journal of Law and Information Technology*, 27(2), 122–141. <https://doi.org/10.1093/ijlit/eaz004>. srani, E. (2017, August 31). Algorithmic due process: Mistaken accountability and attribution in *State v. Loomis*. *Harvard Journal of Law & Technology Digest*. <https://jolt.law.harvard.edu/digest/algorithmic-due-process-mistaken-accountability-and-attribution-in-state-v-loomis-1>. Van Meter, M. (2016, February 25). One judge makes the case for judgment. *The Atlantic*. <https://www.theatlantic.com/politics/archive/2016/02/one-judge-makes-the-case-for-judgment/463380/>.

integration of AI into governance is not merely a technical issue but a societal decision about which values to uphold.

These cases also challenge the notion of algorithmic neutrality, exposing how AI systems inevitably encode the biases and priorities of their creators. Loomis highlights the risks of commercial AI prioritizing efficiency and risk assessment over due process, while Rowicz demonstrates how government-developed AI may advance administrative or political objectives that undermine judicial independence. The central issue, therefore, is whether justice systems will adapt to these embedded biases or demand that AI technologies be designed to uphold legal and ethical standards. The path chosen will have profound implications not only for innovation but for the legitimacy, fairness, and democratic accountability of legal institutions.

6. RECOMMENDATIONS - CONCLUSIONS

The central question—"AI for Justice or Justice for AI?"—is not just theoretical but urgently practical, as current decisions on AI integration will shape legal institutions for generations. To ensure AI serves justice rather than distorting it, several key recommendations emerge: transparency must be a fundamental requirement to maintain accountability and legitimacy; human agency and oversight must be preserved through meaningful control and expertise; bias testing and mitigation must be continuous, not one-off; and democratic participation must be embedded in AI governance to align systems with public values.

This dilemma is fundamentally one of values: whether to design AI systems that uphold the principles of justice—even at the cost of reduced functionality or increased development expenses—or to reshape legal frameworks to accommodate the inherent limitations and commercial interests of AI technologies. Resolving this issue requires inclusive democratic deliberation, extending beyond the perspectives of technologists and policymakers, as the legitimacy and fairness of justice systems rely on public trust and collective societal values. As AI becomes increasingly integrated into legal institutions and commercial incentives intensify, the challenge of preserving justice system integrity against the pressures of technological accommodation will only grow. The choice before us is critical: to develop AI tools that support and remain subordinate to human judgment and democratic oversight, or to risk creating a system where efficiency is prioritized over fairness, transparency, and accountability—ultimately shaping not only the performance of justice but its ethical foundation and societal legitimacy.



Dr Konstantinos Kouroupis, vanredni profesor
Pravni fakultet, Univerzitet Frederik, Kipar

George Christofides, doktorand
Pravni fakultet, Specijalno nastavno osoblje, Univerzitet Frederik, Kipar

VEŠTAČKA INTELIGENCIJA ZA PRAVDU ILI PRAVDA ZA VEŠTAČKU INTELIGENCIJU? KRITIČKA ANALIZA ALGORITAMSKOG PROCESA DONOŠENJA ODLUKA U PRAVOSUDNIM SISTEMIMA

Apstrakt:

Ovaj rad istražuje razvijajući odnos između veštačke inteligencije (AI) i pravde, kritički ispitujući da li AI služi pravdi ili se pravda preoblikuje kako bi se prilagodila AI. U središtu ove analize je Akt o veštačkoj inteligenciji Evropske unije, koji uspostavlja sveobuhvatan regulatorni okvir za upravljanje korišćenjem AI tehnologija, naročito u oblastima visokog rizika poput pravosudnog sistema. Rad istražuje transformativni uticaj AI na pružanje pravde, ocenjujući kako dobitke u efikasnosti tako i značajne rizike povezane sa njenom primenom u procesima donošenja pravnih odluka. Posebna pažnja posvećena je etičkim nedoumicama koje proizilaze iz algoritamske pristrasnosti, netransparentnosti u odlučivanju i mogućem narušavanju osnovnih pravnih načela kao što su pravičnost, odgovornost i ljudski nadzor. Komparativna metodologija koristi se za analizu pristupa Evropske unije i Sjedinjenih Američkih Država, ističući razlike u regulatornoj filozofiji, institucionalnim merama i javnoj debati o veštačkoj inteligenciji u oblasti pravde. Poseban naglasak stavljen je na proučavanje dva značajna slučaja u kojima se pravda ostvaruje primenom alata veštačke inteligencije: prvi se odnosi na slučaj Rowicz (pravo EU), a drugi na slučaj Loomis (međunarodno pravo). Dok EU usvaja preventivni i na pravima zasnovan pristup, SAD-e teže regulatornom pristupu zasnovanom na inovacijama i vođenom tržišnim principima. Ova transatlantska komparacija naglašava hitnu potrebu za usklađenim pravnim standardima koji podržavaju pravdu u eri algoritamskog posredovanja, postavljajući pitanje da li AI zaista može služiti pravdi ili pravni sistem mora da se prilagodi kako bi pravda ostala usredsređena na čoveka.

Ključne reči: *pravda, veštačka inteligencija, transparentnost, odgovornost, slučaj Loomis, slučaj Rowicz*

REFERENCES

1. Ashley, K. (2017). *Artificial intelligence and legal analytics: New tools for practice in the digital age*. Cambridge University Press.
2. Begby, E. (2021). *Automated risk assessment in the criminal justice process: A case of 'algorithmic bias'?* Oxford Academic Books.

3. Byuers, J. (2018). *Artificial intelligence – The practical legal issues*. Legal Brief Publishing.
4. Ceyhun Necati Pehlivan, Forgó, N., & Valcke, P. (2024). *The EU artificial intelligence (AI) act: A commentary*. Kluwer Law International.
5. De Miguel Beriain, I. (2018). Does the use of risk assessments in sentences respect the right to due process? A critical analysis of the Wisconsin v. Loomis ruling. *Law, Probability and Risk*, 17(1), 45–53. <https://doi.org/10.1093/lpr/mgy001>
6. European Commission. (n.d.). *Regulatory framework for artificial intelligence*. Retrieved June 27, 2025, from <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
7. European Parliament. (n.d.). *Annex III: High-risk AI systems (Artificial Intelligence Act)*. Retrieved June 27, 2025, from <https://artificialintelligenceact.eu/annex/3/>
8. Freeman, K. (2016). Algorithmic injustice: How the Wisconsin Supreme Court failed to protect due process rights in State v. Loomis. *North Carolina Journal of Law & Technology Online*, 18, 75–97. http://ncjolt.org/wp-content/uploads/2016/12/Freeman_Final.pdf
9. Hamilton, M. (2015). Risk-needs assessment: Constitutional and ethical challenges. *American Criminal Law Review*, 52(2), 231–274.
10. Harvard Law Review. (2017, March 21). Wisconsin Supreme Court requires warning before use of algorithmic risk assessments in sentencing. *Harvard Law Review*, 130. <https://harvardlawreview.org/print/vol-130/state-v-loomis/>
11. Israni, E. (2017, August 31). Algorithmic due process: Mistaken accountability and attribution in State v. Loomis. *Harvard Journal of Law & Technology Digest*. <https://jolt.law.harvard.edu/digest/algorithmic-due-process-mistaken-accountability-and-attribution-in-state-v-loomis-1>
12. James, N. (2015). *Risk and needs assessment in the criminal justice system* (CRS Report No. R44087). Congressional Research Service.
13. Lightbourne, J. (2017). Damned lies and criminal sentencing using evidence-based tools. *Duke Law & Technology Review*, 15(1), 1–25.
14. Liu, H.-W., Lin, C.-F., & Chen, Y.-J. (2019). Beyond State v. Loomis: Artificial intelligence, government algorithmization, and accountability. *International Journal of Law and Information Technology*, 27(2), 122–141. <https://doi.org/10.1093/ijlit/eaz004>
15. Monahan, J., & Skeem, J. (2016). Risk assessment in criminal sentencing. *Annual Review of Clinical Psychology*, 12, 489–513. <https://doi.org/10.1146/annurev-clinpsy-021815-092945>
16. Negnevitsky, M. (2020). *Artificial intelligence: A guide to intelligent systems* (3rd ed.). Addison Wesley.
17. O’Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
18. Ryberg, J. (2022). *Sentencing and algorithmic transparency*. Oxford Academic Books.
19. Smith, M. (2016, June 22). In Wisconsin, a backlash against using data to foretell defendants’ futures. *The New York Times*. <https://www.nytimes.com/2016/06/23/us/backlash-in-wisconsin-against-using-data-to-foretell-defendants-futures.html>
20. *State v. Loomis*, 881 N.W.2d 749 (Wis. 2016). <https://www.wicourts.gov/sc/opinion/DisplayDocument.pdf?content=pdf&seqNo=171690>
21. Van Meter, M. (2016, February 25). One judge makes the case for judgment. *The Atlantic*. <https://www.theatlantic.com/politics/archive/2016/02/one-judge-makes-the-case-for-judgment/463380/>
22. Zerilli, J. (2022). *Algorithmic sentencing: Drawing lessons from human factors research*. Oxford Academic Books.