

Autonomous Decision-Making and Agentic AI: Challenges and Prospects for Cyber Law

Sayid Muhammad Rifki Noval ¹ ✉ , Irma Rachmawati Maruf ² ,
Ahmad Jamaludin ³ , Deden Sumantry ⁴ ,
Mohd Zakhiri Md Nor ⁵ 

¹ Graduate Program, Doctoral Program in Law, Universitas Pasundan,
Bandung, Indonesia

² Graduate Program, Master's Program in Notarial Studies, Universitas
Pasundan, Bandung, Indonesia

³ Department of Law, Faculty of Law, Universitas Islam Nusantara, Bandung,
Indonesia

⁴ Graduate Program, Master's in Notarial Studies, Universitas Pasundan,
Bandung, Indonesia

⁵ UUM College of Law, Government and International Studies, Universiti
Utara Malaysia, Kedah, Malaysia

✉ Corresponding email: sayidrifqi@unpas.ac.id

Abstract

Artificial intelligence's (AI) explosive growth, particularly in the form of Automated Decision Making (ADM) and Agentic AI, has brought significant changes across various sectors of life, while simultaneously posing complex legal

and ethical challenges. This paper evaluates AI regulation and governance from the viewpoint on protection of consumers and cyber legislation, with an emphasis on Indonesia, which is striving to integrate this technology into the legal system and public policy. The findings reveal that Indonesia's current AI-related regulations remain fragmented and insufficient to handle the regulatory dangers that progressively autonomous AI systems bring. This study examines several cases that illustrate the negative impacts of AI, such as algorithmic errors in the credit system in Germany that resulted in injustices for many individuals, as well as the social assistance distribution scandal in the Netherlands that had serious social and political implications. Concurrently, AI Act was adopted by the European Union as a significant regulatory advance that introduced a risk-based framework to increase accountability, transparency and human oversight in AI governance. Additionally, the controversy surrounding moral problems use AI in the US judiciary related to mass surveillance were also critically analyzed. Analysis of the existing regulations, including the Data Protection Law, the Digital Information and Transactions Law, and the Indonesian government's ethical policies, identifies deficiencies that need to be addressed through adaptive and holistic regulations. The concept of computational accountability and the adoption of international regulations as the Product Liability Directive (PLD) and the AI Liability Directive (AILD) are proposed through a normative juridical analysis, as mechanisms to strengthen accountability and legal protection. This paper also highlights the significance of clarification, openness, and awareness of new rights in the setting of increasingly autonomous and adaptive AI. By integrating technical, legal, social, and ethical aspects, this research contributes to the creation of an equitable and effective regulatory system for AI, particularly for developing countries like Indonesia. These findings are expected to act as a guide for those in power on tackling the legal issues of AI in the digital age.

KEYWORDS *Agentic Artificial Intelligence; Algorithmic Bias; Autonomous Decision-Making; Cyber Law.*

Introduction

In recent years, artificial intelligence (AI) has grown exponentially, creating economic and social value while also posing significant risks that require serious attention. Modern AI systems are characterized by high complexity, operational opacity, and an increasing level of autonomy in decision-making (Li et al., 2022), thereby posing complex regulatory and ethical challenges. As recognized in international regulatory and policy frameworks such as the EU AI Act and the Organization for Economic Co-operation and Development AI Principles, which highlight the challenges of transparency, explainability, and accountability in advanced AI systems.¹ Quoted from Luye Bao et al., a global survey in 2022 involving 2,700 AI experts indicated that the majority of them estimated a 5% chance that super intelligent machines could threaten the existence of humanity.² These concerns reflect the potential of AI to deepen socio-economic inequalities, as automated systems trained on historically biased data may reproduce structural discrimination, while unequal access to digital infrastructure and data resources can marginalize disadvantaged communities. At the same time, the concentration of AI development within dominant technology firms risks amplifying power asymmetries, thereby diminishing human values, stifling creativity, discriminating against vulnerable groups, and violating individual privacy. Therefore, there is a strong argument that the development and application of AI cannot be fully entrusted to the technology industry without effective government intervention and regulation.³ The launch of the AI Incidents Monitor by the Organization for Economic Co-operation and Development (OECD) in November 2023 further substantiates this, having documented over 11,000 incidents pertaining to the hazards of AI utilization across diverse sectors.⁴

¹ Alessio Tartaro, "Towards European Standards Supporting the AI Act: Alignment Challenges on the Path to Trustworthy AI," in *Proceedings of the AISB Convention*, 2023, 98–106.

² Luye Bao et al., "Whose AI? How Different Publics Think about AI and Its Social Impacts," *Computers in Human Behavior* 130 (2022): 107182, <https://doi.org/10.1016/j.chb.2022.107182>.

³ Niloufer Selvadurai, "Seeing the Full Picture: An Integrated Regulatory Model for AI," *International Journal of Law and Information Technology* 33 (January 2025): eaaf005, <https://doi.org/10.1093/ijlit/eaaf005>.

⁴ Guido Noto La Diega and Leonardo C T Bezerra, "Can There Be Responsible AI without AI Liability? Incentivizing Generative AI Safety through Ex-Post Tort Liability under the EU AI Liability Directive," *International Journal of Law and Information Technology* 32 (June 2024): eaae021, <https://doi.org/10.1093/ijlit/eaaf005>.

In awareness of potential risks linked to the implementation of AI, various governments and regulatory bodies at both national and regional levels have taken significant legal and regulatory measures. In the United States, several lawsuits have been filed against major technology companies such as Google and OpenAI regarding alleged privacy violations committed by their generative AI models. In Europe, a unique task team has been set up by European Data Protection Board (EDPB) and the European Union to monitor and regulate the utilization of technologies such as ChatGPT. In 2023, the Guarante, the Italian Data Protection Authority filed a lawsuit to OpenAI, leading to an interim nationwide ban on the ChatGPT service. Additionally, similar investigations are underway in Poland, while authorities in Germany and France are actively requesting information and developing action plans to monitor the impact of AI on personal data protection.⁵ These measures reflect a global awareness of the need for strict AI governance to safeguard the confidentiality of human data and defend individual freedoms in the digital age.

The regulatory and supervisory measures taken by various countries on AI reflect a global awareness of the importance of governing this technology. However, public responses and perceptions of AI vary significantly across regions, which also influences the dynamics of regulatory implementation. In Asia, for example, there is a greater tendency towards optimism about the benefits of AI, which could make the public less vigilant about potential risks. Nevertheless, such comparisons should be contextualized in light of methodological limitations, including sample representativeness, cultural response bias, and margins of error that may influence the interpretation of regional differences. The 2023 survey shows that more than 67% of Singaporeans have used Generative AI (Gen AI) technology, with 69% reporting positive experiences. However, the survey's sampling methodology, sample size, and representativeness were not fully disclosed, which may limit the reliability of cross-regional comparisons, such as those between Asia and Europe or the United States, and is very important. The Kameke study in Southeast Asia also

⁵ Hannah Ruschemeier, "Generative AI and Data Protection," *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e6, <https://doi.org/10.1017/cfl.2024.2>.

confirms that the majority of respondents expect greater benefits than losses from AI. In contrast, in Europe and the United States, the prevailing narrative is more pessimistic, driven by decades-long negative social and economic experiences with automation and globalization. The 2023 survey in the United States found that about 75% of citizens are worried about job losses associated with artificial intelligence. This difference in perception underscores the need for regulatory and educational approaches tailored to the social, economic, and cultural contexts of each region to ensure effective and inclusive AI management.⁶

Although several regulatory frameworks have been developed in various developed countries to govern AI technology, there remain significant gaps in the governance of Agentic AI, particularly regarding legal liability and consumer protection in developing countries like Indonesia. Existing regulations tend to focus on traditional AI models and have not fully addressed the complexities and risks posed by increasingly advanced autonomous AI systems. This research aims to fill that gap with a comprehensive normative legal approach that integrates analysis of international regulations and local legal contexts, while highlighting the unanswered challenges in AI regulation within the realms of cyber law and consumer protection.

In fact, a wide range of technological devices fall under the umbrella of artificial intelligence. AI is commonly considered to be systems that exhibit intelligent behavior. The issue that often emerges is that while these advancements exhibit behaviors that seem intelligent through learning and the application of knowledge, these behaviors are ultimately the result of cyber-physical systems' creation, programming, and administration. The issue arises when liberty and flexibility clash with more dynamic and unpredictable environments; the range of possible behaviors generated by AI systems becomes so vast that they can no longer be fully anticipated or controlled through conventional systems engineering practices. This technical unpredictability has significant legal implications, as it complicates liability attribution, weakens

⁶ Jason Grant Allen, Jane Loo, and Jose Luis Luna Campoverde, "Governing Intelligence: Singapore's Evolving AI Governance Framework," *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e12, <https://doi.org/10.1017/cfl.2024.12>.

foreseeability standards in negligence-based regimes, and challenges traditional doctrines that presume human intention, control, and fault as the basis for legal responsibility. In other words, a technology's reactions become unpredictable to human operators. One of the foundational technologies enabling autonomy and self-adaptation is machine learning, a subfield of AI that focuses on developing algorithms that automatically enhance the capabilities of autonomous systems.⁷

The differences in society's perception and response to AI further clarify the complexity of this technology, which encompasses a wide spectrum of systems that behave as if they are intelligent. A framework that can simulate human cognitive behavior via knowledge analysis and implementation is often referred to as artificial intelligence. However, the creation, programming, and execution of intricate digital and physical systems are responsible for such behavior. The main problem is when the degree of flexibility and adaptability of the AI system is combined with a dynamic and unpredictable environment, resulting in a vast number of possible behaviors that are difficult to predict or control by human operators. This phenomenon poses significant operational and legal risks. One of the core technologies supporting this autonomy and adaptation is machine learning, a branch of AI that develops algorithms to enhance the capabilities of autonomous systems automatically through experience and data.⁸ A deep understanding of these characteristics is crucial for formulating a legal structure that is both swift and attuned to the challenges posed by AI.

The development of autonomy in AI systems has driven a significant transition from human-made choices to machine-generated ones. This shift has elicited a variety of responses, particularly concerning the risks to the respect and protection of human rights. ADM raises serious concerns when decisions affecting individuals are made without direct human involvement. In this context, there is a risk that humans grant excessive authority to machines, which could potentially shift personal responsibility and diminish the human factor in

⁷ Jacques Hartmann et al., "Artificial Intelligence, Autonomous Drones and Legal Uncertainties," *European Journal of Risk Regulation* 14, no. 1 (March 2023): 31–48, <https://doi.org/10.1017/err.2022.15>.

⁸ Hartmann et al.

the technique of making decisions. Research shows that humans, both individually and collectively, often struggle to be objective and rational in decision-making; they are susceptible to emotional influences, subconscious biases, and complex social factors, ranging from physiological conditions like hunger or fatigue to social pressure and the environmental context at the time of decision-making.⁹ Therefore, although ADM offers the potential for increased efficiency and objectivity, the ethical and legal challenges related to responsibility and fairness in computerized decision-making have to be seriously addressed.

Concerns regarding risk Making Decisions Automatically (ADM) are also strongly felt by the European Union, particularly within banking, a sector which greatly depends on big data processing to optimize risk management, expedite routine tasks, and personalize services and goods for clients. The potential use of ADM in decision-making processes has expanded owing to the fast development of AI technology, which may both supplement and replace traditional human judgment. In response to the growing amount of standard information and algorithmic methods, the European Union is proposing the Financial Data Access Regulation (FIDA), which intends to govern data access and use within the framework of Open Money in this industry. This regulation is designed to balance consumer rights protection with business interests while ensuring transparency and accountability in the use of ADM in the financial sector.¹⁰ This initiative reflects proactive regulatory efforts to address the legal complications brought on by AI-based judgment-making automation.

In addition to the general risks of ADM, specific concerns arise when AI is used to exercise administrative discretion, which has traditionally been the exclusive domain of humans. AI advancements enable the automation of certain bureaucratic tasks, allowing machines to perform functions that were previously only carried out by public officials. In some cases, properly programmed algorithms can produce decisions comparable to human decision outcomes. In

⁹ Elena Abrusci and Richard Mackenzie-Gray Scott, "The Questionable Necessity of a New Human Right against Being Subject to Automated Decision-Making," *International Journal of Law and Information Technology* 31, no. 2 (August 2023): 114–43, <https://doi.org/10.1093/ijlit/eaad013>.

¹⁰ Felix Pflücke, "Data Access and Automated Decision-Making in European Financial Law," *European Journal of Risk Regulation* 16, no. 1 (March 2025): 11–23, <https://doi.org/10.1017/err.2024.60>.

response to this phenomenon, several European countries have taken regulatory steps; for example, Spain, through Law No. 40 of 2015, based on the Digital Freedom Charter, the governing sector's legal framework, and Estonia, which suggested changes to the Administrative Procedures Act to allow for the use of AI in government. Nevertheless, the Standard Act does not clearly control or prohibit automated systems from making arbitrary choices at the European Union level. Instead, this regulation establishes important protections such as human agency involvement, transparency, and impact assessments on human rights, especially for high-risk AI systems listed in Annexe III.¹¹ This approach reflects a harmony between the use of technology and the defense of fundamental human rights in general administration governance.

In addition to administrative legal challenges, the utilization of methods in artificial intelligence (AI) creates significant social, economic, and ethical issues. Algorithms trained on historical data can perpetuate and amplify existing societal biases. When this biased dataset is used to train AI models, it can perpetuate structural discrimination and deepen social inequality. The issue is exacerbated by a shortage of adequate due diligence mechanisms to ensure transparency and accountability in AI outputs. As a result, discrimination can occur in various sectors, including labor recruitment, credit allocation, and law enforcement.¹² Another management issue with standard simulated data is equally major: it is usually not categorized as personal data because it is used for training and has been anonymized. In this situation, when the data subjects cannot be identified, legal protections such as those regulated by the GDPR (General Data Protection Regulation) of the European Union are no longer effective. Often, data can be stored, copied, transferred, or even traded and reused in different contexts without adequate legal protection, especially when

¹¹ Juan Carlos Covilla, "Artificial Intelligence and Administrative Discretion: Exploring Adaptations and Boundaries," *European Journal of Risk Regulation* 16, no. 1 (March 2025): 36–50, <https://doi.org/10.1017/err.2024.76>.

¹² Guillem Izquierdo Grau, "The Development Risks Defence in the Digital Age," *European Journal of Risk Regulation* 16, no. 1 (March 2025): 197–216, <https://doi.org/10.1017/err.2024.43>.

the purpose of its use is not explicitly defined from the outset. It is hidden behind reasons of anonymity and general consent clauses.¹³

A concrete introduction of an AI is just one instance of this risk model initially developed to diagnose mental illnesses based on behavioral data in the context of medical research and healthcare services, but later used secondarily in job applicant screening systems, which can have detrimental effects. Additionally, the problem of fresh data collected through an evaluation operation —information that infers characteristics or profiles of individuals based on the available data. Although personal data is generally understood as information already known by the individual about themselves, this new data is insights generated by the data controller and can have practical or legal impacts on the data subjects.¹⁴ This issue becomes increasingly complex when data subjects lack access or the ability to oversee and seek redress for decisions made based on data they are unaware of, thereby increasing the potential for bias and injustice. This condition creates an urgent need for full transparency, responsibilities, and rigorous supervision procedures that are essential for machine learning-driven automated decision-making.

A real illustration of the risks of data processing that harm individuals can be found in Victoria Hendrickx's study on the deployment of AI in the legal system, especially by judges. AI has been used in the ordinary law judicial system to automate routine administrative tasks, including scheduling, case assignment, digital filing, precedent and proof analysis, case management, and acceptability evaluation. But developments in AI and machine learning have enabled more sophisticated applications capable of learning from big data and historical experience. Current AI systems not only support administrative functions but also provide substantive support through average sentence calculations, recidivism risk assessments, and case outcome predictions. Generative AI technology can even produce high-quality texts that assist judges in drafting

¹³ Rainer Mühlhoff and Hannah Ruschemeier, "Updating Purpose Limitation for AI: A Normative Approach from Law and Philosophy," *International Journal of Law and Information Technology* 33 (January 2025): eaaf003, <https://doi.org/10.1093/ijlit/eaaf003>.

¹⁴ Bart Custers and Helena Vrabec, "Tell Me Something New: Data Subject Rights Applied to Inferred Data and Profiles," *Computer Law & Security Review* 52 (April 2024): 105956, <https://doi.org/10.1016/j.clsr.2024.105956>.

decisions, summarising documents, and providing legal advice.¹⁵ This development tended to substantially improve judicial practices. In 2024, the Court of Justice of the European Union (CJEU) affirmed that, in the future, all court staff will be able to use AI-powered virtual assistants, including AI avatars for internal training and the creation of didactic materials, as well as case management systems equipped with case classification and correlation analysis capabilities.¹⁶ This situation poses new challenges related to honesty, accountability, and the defense of personal liberties in implementing AI in the criminal justice sector, which must be addressed through appropriate regulations.

This writing rejects the notion that technology developers can achieve perfect objectivity—or aperspectival objectivity—based on the outcomes that are produced via AI systems, which are completely neutral. In fact, every individual, whether consciously or unconsciously, is influenced by the views, beliefs, experiences, and social context surrounding them.¹⁷ Therefore, expecting Automated Decision Making (ADM) to be free from bias and fully respect individuals' fundamental rights is highly questionable. The presence of algorithmic bias inherently threatens the realization of justice, as such bias often leads to unfair decisions. The existence of various interpretations of justice further exacerbates the complexity of justice. A computer scientist might assess justice in terms of equal data processing by algorithms, while a sociologist highlights the impact of systems that perpetuate social inequality. On the other hand, legal experts assess fairness based on the system's compliance with individual rights and applicable legal norms.¹⁸ In fact, there is a view that modern AI is not a normative entity, and therefore cannot be intrinsically

¹⁵ Victoria Hendrickx, “Rethinking the Judicial Duty to State Reasons in the Age of Automation?,” *Cambridge Forum on AI: Law and Governance* 1 (June 2025): e26, <https://doi.org/10.1017/cfl.2025.11>.

¹⁶ Sophie Weerts, “Generative AI in Public Administration in Light of the Regulatory Awakening in the US and EU,” *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e3, <https://doi.org/10.1017/cfl.2024.10>.

¹⁷ Sam Wrigley, “Bring in the Sceptics: Using Science and Technology Studies in Law,” *European Journal of Risk Regulation* 16, no. 1 (March 2025): 51–62, <https://doi.org/10.1017/err.2024.103>.

¹⁸ Eduard Fosch-Villaronga et al., “Exploring Fairness in Service Robotics,” *Cambridge Forum on AI: Law and Governance* 1 (June 2025): e28, <https://doi.org/10.1017/cfl.2025.10013>.

considered fair or unfair.¹⁹ This multidisciplinary understanding is important for formulating a realistic legal and ethical framework in controlling the use of AI.

Furthermore, beyond the previously emphasized issues of bias and fairness, it's critical to recognize that AI technology has also evolved significantly, adding complexity to its regulation. Initially, AI was known as predictive artificial intelligence (PAI), which focused on analyzing historical data to predict future outcomes through statistical models, enabling more proactive decision-making across various domains.²⁰ Then came Generative AI (GAI), which is capable of creating new and realistic material, including text, photos, audio, video, and code, by using patterns revealed via extensive datasets. Generative image systems like DALL-E and linguistic frameworks like ChatGPT are notable instances of GAI.²¹ However, the latest development currently in the spotlight is Agentic AI, sometimes referred to as Reasoning AI, which is the third phase of AI. For this research, the term "predictive AI" refers to programs that use historical data to predict events to come; "generative AI" refers to systems that can produce new content, like text, images, or audio, utilizing patterns discovered in large datasets; and "agentic AI" refers to autonomous agents with the ability to make decisions and act independently in complex environments.

According to Jensen Huang, CEO of Nvidia, AI at this stage is not only capable of perception and content generation but also understanding, problem-solving, and recognizing previously unseen situations. Reasoning AI enables the creation of digital robots and autonomous agents that can act independently and adaptively in complex environments. One of the important subcategories in this evolution is the Large Language Model (LLM), which uses

¹⁹ Marcin Korecki et al., "The Man Behind the Curtain: Appropriating Fairness in AI," *Minds and Machines* 34, no. 1 (April 2024): 7, <https://doi.org/10.1007/s11023-024-09669-x>.

²⁰ Nir Kshetri, "From Predictive and Generative to Agentic AI: Shaping the Future of Marketing Operations and Strategies," *Computer* 58, no. 4 (April 2025): 121–29, <https://doi.org/10.1109/MC.2025.3530304>; Siddhartha Roy and Runa Ganguli, "Comparative Analysis of Machine Learning Classification Algorithms for Predictive Models Using WEKA," in *Emerging Technologies in Data Mining and Information Security*, In P. Dutt (Springer Nature Singapore, 2023), 173–83, https://doi.org/10.1007/978-981-19-4052-1_19.

²¹ Elana Zeide, "Expanding the Paradigm: Generative Artificial Intelligence and U.S. Privacy Norms," *Cambridge Forum on AI: Law and Governance* 1 (March 2025): e16, <https://doi.org/10.1017/cfl.2024.15>.

billions of parameters to generate outputs based on probabilistic inference, rather than causal understanding. LLMs have revolutionized the fields of AI and natural language processing (NLP) by enabling unprecedented understanding and generation of human language.²²

Understanding the fundamental differences between humans and Large Language Models (LLMs) from an epistemological perspective is crucial for analyzing the effects of using AI for decision-making. Semantic talents, or the capacity to give things meaning, enable humans to function as "semantic motors" that flourish in comprehending and expressing reality around them. In contrast, LLMs function as "syntactic engines" that produce probabilistic predictions of the more likely word order to occur later in a document, using intricate neural network architectures to handle vast quantities of data and standardized parameters, while maintaining knowledge or awareness of the content being processed. Unlike the scientific method, which relies on verifiable predictions and explanatory models, the outputs of AI, such as LLM, are not based on such theoretical foundations.²³ The uncertainty and unpredictability of LLMs are the main drivers behind the development of intelligent agent AI, a novel kind of artificial intelligence that refers to autonomous agents capable of making decisions and acting independently.

Agentic AI is a qualitative leap in the development of artificial intelligence, defined by its ability to set complex goals in changing and uncontrollable situations, and to achieve them through autonomous resource management.²⁴ Historically, the term agent has been used in AI studies to refer to systems capable of making autonomous decisions.²⁵ In the broader realm of Agent AI, there is a distinction between entirely independent and self-sufficient systems.

²² Ugo Pagallo, "LLMs Meet the AI Act: Who's the Sorcerer's Apprentice?," *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e1, <https://doi.org/10.1017/cfl.2024.6>.

²³ Ludovica Paseri and Massimo Durante, "Examining Epistemological Challenges of Large Language Models in Law," *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e7, <https://doi.org/10.1017/cfl.2024.7>.

²⁴ Deepak Bhaskar Acharya, Karthigeyan Kuppan, and B. Divya, "Agentic AI: Autonomous Intelligence for Complex Goals—A Comprehensive Survey," *IEEE Access* 13 (2025): 18912–36, <https://doi.org/10.1109/ACCESS.2025.3532853>.

²⁵ Noto La Diega and Bezerra, "Can There Be Responsible AI without AI Liability? Incentivizing Generative AI Safety through Ex-Post Tort Liability under the EU AI Liability Directive."

Non-autonomous agents are typically employed in contexts that are more straightforward or regulated, where activities are predefined and require less adaptive decision-making. This system is not built to change on its own over time; instead, it depends on human oversight or fixed regulations. Conversely, autonomous Agent AI systems operate independently, responding to complex, dynamic environments with minimal human assistance. Based on a basic adaptive framework that carries activities, sensory inputs, and data acquired in a fluid and diverse environment, these autonomous agents are driven by objectives. They adapt their tactics in response to shifting conditions without constant human control. The fundamental principle of autonomy enables Agent AI to make decisions autonomously, adapt through various learning paradigms, and manage real-time tasks effectively, thereby reducing the need for continuous supervision in complex environments.²⁶ In a Capgemini poll of 1,100 executives conducted in the second half of 2024, 10% of major companies have implemented AI agents. More than 50 percent plan to adopt them in the next year, and 82 percent anticipate their integration within three years.²⁷ This trend shows the accelerated adoption of Agentic AI technology across various sectors, which demands serious devotion to elements of legal, ethical, and regulatory concerns.

Given the complexity and dynamics of AI development, particularly Agentic AI, which introduces new challenges in automated decision-making, in-depth and comprehensive legal studies have become crucial. The shift from predictive and generative systems to autonomous agents demands a broader understanding not only of the technology but also of law, ethics, and consumer protection. Therefore, this research will examine the legal aspects relevant to the development of AI, with the aim of providing regulatory recommendations to address issues with existing and future artificial intelligence regulations in the realm of cyber law.

²⁶ Francesco Piccialli et al., “AgentAI: A Comprehensive Survey on Autonomous Agents in Distributed AI for Industry 4.0,” *Expert Systems with Applications* 291 (October 2025): 128404, <https://doi.org/10.1016/j.eswa.2025.128404>.

²⁷ Kshetri, “From Predictive and Generative to Agentic AI: Shaping the Future of Marketing Operations and Strategies.”

Several previous studies have examined the regulatory and ethical implications of Artificial Intelligence from different perspectives. McGraw proposed foundational AI ethical principles emphasizing beneficence, non-maleficence, autonomy, justice, and explicability, which later influenced global AI governance discourse.²⁸ Gacutan and Selvadurai analyzed the “right to explanation” under the GDPR and highlighted the limitations of existing data protection frameworks in addressing automated decision-making.²⁹ Yazdanpanah et al. explored the concept of robotics law and argued that emerging autonomous systems challenge traditional doctrines of liability and foreseeability.³⁰ Tsoukalas critically rejected the idea of granting legal personality to AI, asserting that responsibility must remain attributable to human actors and institutions.³¹ More recently, studies on algorithmic governance and accountability have emphasized the need for risk-based regulatory models and stronger oversight mechanisms for high-risk AI applications.³²

Despite these significant contributions, most existing studies focus on predictive or generative AI in developed jurisdictions. There remains limited doctrinal analysis of Agentic AI systems with greater autonomy, especially in developing countries where regulatory infrastructure and consumer protection regimes are still evolving. Furthermore, the rapid deployment of autonomous AI agents in commercial and administrative sectors has outpaced the development of clear liability standards, leaving unresolved questions regarding fault attribution, due diligence obligations, and consumer remedies.

Therefore, this research is urgent and significant for several reasons. First, it addresses the regulatory vacuum concerning Agentic AI in relation to consumer protection and Indonesian cyberlaw. Second, by fusing local legal concepts with international regulatory changes, it adds to the norms of legal discourse. Third, it provides conceptual clarity on legal accountability in highly autonomous systems, which is crucial for sustaining public trust in AI-driven technology,

²⁸ David K McGraw, “Ethical Responsibility in the Design of Artificial Intelligence (AI) Systems,” *International Journal on Responsibility* 7, no. 1 (2024): 4, <https://doi.org/10.62365/2576-0955.1114>.

²⁹ Joshua Gacutan and Niloufer Selvadurai, “A Statutory Right to Explanation for Decisions Generated Using Artificial Intelligence,” *International Journal of Law and Information Technology* 28, no. 3 (2020): 193–216, <https://doi.org/10.1093/ijlit/eaad016>.

³⁰ Vahid Yazdanpanah et al., “Reasoning about Responsibility in Autonomous Systems: Challenges and Opportunities,” *AI & Society* 38, no. 4 (2023): 1453–64.

³¹ Anastasios Tsoukalas, “Artificial Intelligence and Legal Personality: Obligations, Rights, and Liability in the Age of Autonomous Systems,” *Επιθεώρηση Δικαίου Πληροφορικής* 6, no. 1 (2025), <https://doi.org/10.26262/infolawj.v1i1.10681>.

³² Veve Fry, “An Overview of the Risk-Based Model of AI Governance,” *ArXiv Preprint ArXiv:2507.15299*, 2025.

ensuring legal certainty, and safeguarding basic rights. It intends to provide a forward-looking framework that may address the issues through the inclusion of Agentic AI inside an organized legal examination, accelerating the transformation of digital governance.

It uses an established legal framework with a descriptive-analytical specification, which is appropriate for examining legal doctrines, statutory frameworks, and principles such as autonomy, accountability, and liability in AI governance. The analytical approach was applied to examine doctrinal concepts such as autonomy, accountability, and liability in AI governance. The comparative approach was used to contrast Indonesian laws along with worldwide frameworks, particularly developments within the European Union. The case approach analysed selected judicial and regulatory cases involving automated decision-making to identify practical legal gaps, while the statute approach examined relevant legislative instruments governing electronic transactions, data protection, and emerging AI regulation. All legal materials were analyzed qualitatively to assess their suitability for addressing the challenges posed by self-governing AI systems.

Secondary data were used in this study. Research from libraries was used to gather data. There are two types of legal materials utilized: main and secondary. Primary materials include laws and regulations governing electronic transactions, personal data protection, and AI-related policies, as well as relevant court decisions on automated decision-making. Secondary materials include scholarly articles, comparative legal studies, and policy documents issued by authorities. Data analysis is conducted qualitatively juridically, and the research specification uses descriptive-analytical.

Automated Decision Making (ADM) and Legal Risks in Automated Decision Making

Automated Decision Making (ADM) is a software component that classifies or scores individuals or objects based on available data and attributes. Generally, there are two types of ADM systems: first, systems that have two

types of expert systems to consider: rule-based, where decision rules are produced by experts using empirical data and common sense, and data-driven, whose decision rules are learned from prior decision data. In rule-based systems, if the data subject feels aggrieved by the decision made, experts can justify the rules applied in the decision-making process.³³

However, the main risk inherent in ADM arises from the limitations of AI systems in applying common sense and causality, which is the ability to infer cause-and-effect relationships. Artificial intelligence systems are unable to understand concepts deeply or adapt solutions to problems they have never encountered before, unlike human thinking. Therefore, when it is said that AI systems make unfair decisions, it refers to their inability to provide equal and nondiscriminatory treatment to all individuals.³⁴ Case is a particular example of Google mistakenly locking the account of a user named Mark because it was deemed to contain child sexual abuse material. In fact, the uploaded photo was a medical documentation related to the swelling of his child's genitals, which was sent for a doctor's consultation via video call during the pandemic. Google uses the AI Content Safety API, which is trained to detect child abuse material based on unique digital fingerprints (hashes). Although the San Francisco Police Department stated there was no violation, Mark's account remained blocked.³⁵ This case illustrates the limitations of AI reasoning and the potential losses and violations of individual rights due to erroneous automated decisions.

From a regulatory perspective, this case raises important concerns regarding due process, proportionality, and the right to effective remedy in automated decision-making systems. In the European context, for example, Article 22 of the General Data Protection Regulation (GDPR) provides safeguards against decisions based solely on automated processing that

³³ Tobias D. Krafft, Marc P. Hauer, and Katharina Zweig, "Black-Box Testing and Auditing of Bias in ADM Systems," *Minds and Machines* 34, no. 2 (May 2024): 15, <https://doi.org/10.1007/s11023-024-09666-0>.

³⁴ Markus Kattinig et al., "Assessing Trustworthy AI: Technical and Legal Perspectives of Fairness in AI," *Computer Law & Security Review* 55 (November 2024): 106053, <https://doi.org/10.1016/j.clsr.2024.106053>.

³⁵ Fiona Jackson, "Stay-at-Home Dad's 'Life Ruined' by Google after He Was Investigated When Photos He Took of His Sick Toddler Son's Genitals for a Doctor to Review during the Pandemic Were Flagged by AI as Potential Child Sexual Abuse Material," *Daily Mail*, August 2022, <https://www.dailymail.co.uk/profile-223/fiona-jackson.html?page=7>.

significantly affect individuals.³⁶ Similarly, evolving AI for highly dangerous AI systems, governance systems like the EU AI Norm Act prioritize human supervision, responsibilities, and transparency. In the United States, while sectoral regulation applies, discussions around algorithmic accountability have intensified, including proposals for stronger oversight of automated content moderation systems.³⁷

In addition to decision-making errors, AI systems' suggestions can also pose dangers that jeopardize users' physical safety. For example, Amazon's voice assistant, Alexa, once gave dangerous instructions to a 10-year-old child who asked for a challenge. Alexa urged the youngster to squeeze a penny to the visible end of the phone battery after plugging it into the socket. This is an extremely risky move because metal conducts electricity, which can cause electric shocks, fires, and serious damage. After this incident was revealed, Amazon made improvements to the Alexa system to prevent similar occurrences.³⁸ From an algorithmic safety perspective, the case demonstrates the necessity of robust risk-filtering mechanisms, context-sensitive safeguards, and continuous system monitoring to prevent AI systems from generating harmful instructions, especially when interacting with children. From a legal standpoint, it raises the issue of whether AI developers and platform operators bear a duty of care to anticipate foreseeable risks and implement preventive measures. The failure to adequately filter dangerous content may indicate a lack of safety-by-design and insufficient human oversight, both of which are increasingly recognized as core principles in emerging AI governance frameworks.

Cases involving losses due to Automated Decision Making (ADM) have emerged in various jurisdictions, including Germany. In the case *OQ v Land Hessen* (Case C-634/21),³⁹ an individual contested a loan rejection decision based on a negative credit score issued by the *Schutzgemeinschaft für allgemeine*

³⁶ Gianclaudio Malgieri, "Automated Decision-Making and Data Protection in Europe," in *Research Handbook on Privacy and Data Protection Law* (Edward Elgar Publishing, 2022), 433–48.

³⁷ Federico Galli, Andrea Loreggia, and Giovanni Sartor, "The Regulation of Content Moderation," in *International Conference on the Legal Challenges of the Fourth Industrial Revolution* (Springer, 2022), 63–87.

³⁸ BBC, "Alexa Tells 10-Year-Old Girl to Touch Live Plug with Penny. BBC," *BBC*, December 2021, <https://www.bbc.co.uk/news/technology-59810383>.

³⁹ Angela Maria Noguera, "Case C-634/21, *OQ v. Land Hessen* (C.J.E.U.)," *International Legal Materials* 63, no. 5 (October 2024): 946–61, <https://doi.org/10.1017/ilm.2024.23>.

Kreditsicherung (SCHUFA), the German credit rating agency. The plaintiff requested access to the personally identifiable data used to determine the score and asked that any information considered erroneous be removed. SCHUFA provides general information about how the score is calculated but refuses to disclose specific factors and weights used, citing trade secrecy. In addition, SCHUFA states that the bank makes the final decision based on the provided score. This case raises important questions. Considering the freedoms of data subject matter, Article 22 of the GDPR upholds the right to be free from decisions made only by automated procedures, unless certain legal conditions permit it. The Court of Justice of the European Union (CJEU) emphasized that the GDPR's overarching objectives of ensuring a high level of defense for fundamental liberties, including the rights to privacy and data protection, must be taken into consideration when interpreting Article 22. According to the Supreme Court of Justice, a “decision” under Article 22 does not require formal legal effects alone; it is sufficient that the automated processing significantly affects the individual's circumstances, behavior, or opportunities. The Court also stressed that Protections such the ability to voice one's opinions, seek human assistance, and challenge the judgment are not merely procedural formalities, but essential mechanisms to prevent arbitrary or opaque algorithmic determinations. Moreover, the Court underscored the principle of transparency, requiring that data provide participants with helpful details about the logic underlying the automated decision-making process. The logic illustrates the CJEU's concern that automated systems, particularly those relying on complex algorithms or profiling, may create asymmetries of power and information that undermine individual autonomy. Therefore, the Court's interpretation strengthens Article 22 as a substantive protective norm rather than a narrow procedural exception, reinforcing accountability in algorithmic governance.

Articles 13, 14, and 15 of the GDPR also regulate the right to clear information regarding the logic of processing, significance, and consequences of personal data processing. On December 7, 2023, the Court of Justice of the European Union (CJEU) issued a preliminary ruling interpreting Article 22

GDPR in the context of credit scoring, specifically regarding the evaluation and prediction of an individual's future payment capacity.⁴⁰

The case in Germany is not the only example of how ADM can cause controversy and discriminatory impacts. In England, the Home Office faced a trial over the use of streaming visa technology between 2015 and 2020, which classified visa applicants by risk level for visa decision-making. In 2020, the legal firm FoxGlove Corporation and the Joint Council for the Welfare of Immigrants (JCWI) opposed this tool because it was considered discriminatory and lacked transparency. They argued that the algorithm assigned higher risk scores to applicants from suspected countries, thereby disproportionately increasing surveillance. As a result of this pressure, the use of streaming visas was halted in August 2020.⁴¹ A similar case arose in the European Travel Information and Authorization System (ETIAS). This European Union information system provides travel authorization to visa-free third-country nationals using profiling algorithms. At the beginning of 2022, around 23.8 million non-EU citizens lived in EU member states, representing 5.3% of the total population. Every year, tens of thousands of third-country nationals are denied entry, found to be in the country illegally, or deported.⁴² ETIAS is part of the European Union's massive digital surveillance infrastructure. Still, its algorithmic processes have the potential to create differential exclusions that disadvantage certain groups of tourists.⁴³

Beyond the controversies surrounding visa surveillance and travel authorization, similar challenges arise in the criminal justice system, which uses algorithms for decision-making. The United States of America case *State v. Loomis* illustrates how the COMPAS methodology is used to determine a

⁴⁰ Emmanuel Vargas Penagos, "Platforms on the Hook? EU and Human Rights Requirements for Human Involvement in Content Moderation," *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e23, <https://doi.org/10.1017/cfl.2025.3>.

⁴¹ Francesca Palmiotto, "When Is a Decision Automated? A Taxonomy for a Fundamental Rights Analysis," *German Law Journal* 25, no. 2 (March 2024): 210–36, <https://doi.org/10.1017/glj.2023.112>.

⁴² Erzsébet Csatlós, "Prospective Implementation of Ai for Enhancing European (in)Security: Challenges in Reasoning of Automated Travel Authorization Decisions," *Computer Law & Security Review* 54 (September 2024): 105995, <https://doi.org/10.1016/j.clsr.2024.105995>.

⁴³ Charly Derave, Nathan Genicot, and Nina Hetmanska, "The Risks of Trustworthy Artificial Intelligence: The Case of the European Travel Information and Authorisation System," *European Journal of Risk Regulation* 13, no. 3 (September 2022): 389–420, <https://doi.org/10.1017/err.2022.5>.

defendant's probability of recidivism. Although the evaluation technique is not fully disclosed, the Wisconsin Supreme Court ruled that the use of this algorithm does not violate due process rights. An independent study revealed significant racial bias. Black defendants are 45% more likely to receive higher recidivism scores compared to white defendants, raising questions about the validity and fairness of this algorithm. The study concluded that COMPAS is no more accurate or fair than predictions made by individuals with little or no expertise in criminal justice, raising serious doubts about the use of algorithms in the criminal justice system.⁴⁴ This case, along with the issue of discrimination in the visa system, underscores the need for stringent oversight and regulation that are responsive to the risk of algorithmic bias across sectors.

In addition to the prospect of prejudice in automated judgment affecting individual rights, there are serious concerns about the recommendations provided by ADM systems based on user data. An important example is the case of *Dyroff v. Ultimate Software Group, Inc.*, which began with a lawsuit against Ultimate Software because its algorithm generated recommendations that contributed to a child's death due to drug consumption. Ultimate Software used user activity data from the Experience Project platform to mine insights from their posts, then made actionable recommendations, including directing users to heroin-related discussion groups that facilitated the purchase of illegal drugs.⁴⁵ Additionally, a 2016 study by Lum and Isaac examined the use of predictive policing algorithms in the United States that analyze drug crime data by type, location, and time, without using personal information. The research results indicate that the algorithm actually reinforces existing biases in historical data, more frequently directing police to Black neighborhoods compared to White neighborhoods, deepening injustices in law enforcement.⁴⁶

⁴⁴ Niamh Kinchin, "Voiceless": The Procedural Gap in Algorithmic Justice," *International Journal of Law and Information Technology* 32 (June 2024): eaae024, <https://doi.org/10.1093/ijlit/eaae024>.

⁴⁵ Elif Kiesow Cortez and Nestor Maslej, "Adjudication of Artificial Intelligence and Automated Decision-Making Cases in Europe and the USA," *European Journal of Risk Regulation* 14, no. 3 (September 2023): 457–75, <https://doi.org/10.1017/err.2023.61>.

⁴⁶ Siddharth Peter de Souza, "The Spread of Legal Tech Solutionism and the Need for Legal Design," *European Journal of Risk Regulation* 13, no. 3 (September 2022): 373–88, <https://doi.org/10.1017/err.2022.4>.

The phenomenon of algorithmic bias and discriminatory impacts arising from the broader hazards of deploying Automated Decision Making (ADM) across government agencies is evident in law enforcement, as shown by the literature on predictive police algorithms in the US. One concrete example is the education sector during the COVID-19 pandemic, which illustrates how ADM algorithms can significantly affect individual rights. The COVID-19 pandemic forced the UK government to suspend face-to-face teaching, creating significant challenges in academic assessment. As a solution, the Office of Qualifications (OfQual) was charged by then-Minister of Education Gavin Williamson with creating a substitute marking scheme for pupils, based on the requirements they would have met had the examinations not been canceled. OfQual then developed an algorithm known as the Center's Assessed Grade (CAG), which combined student rankings from their respective schools and the historical performance of both students and schools. However, when the algorithm results were announced on August 13, 2020, about 40% of students' results were lower than expected, sparking massive protests in London. This controversy led to the dismissal of Gavin Williamson as Secretary of State for Education on September 15, 2020.⁴⁷ This case illustrates the real risks of using ADM in the public context, which can significantly affect individual rights, while also highlighting the importance of transparency, responsibility, and stringent monitoring of algorithmic systems, especially when used in decision-making with far-reaching implications for society.

In addition to the significant risks that have emerged in the education sector, errors in automated decision-making have also caused serious socio-economic impacts in public administration. One of the most striking examples is the childcare benefits scandal in the Netherlands in 2021, which not only caused financial losses but also a social and political crisis. In this incident, around 26,000 people, mostly from immigrant backgrounds, were wrongly accused of child benefit fraud by the tax authorities. They were required to repay benefits they had received over several years, even though many of the

⁴⁷ Merlin Tieleman, "Fairness in Tension: A Socio-Technical Analysis of an Algorithm Used to Grade Students," *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e19, <https://doi.org/10.1017/cfl.2025.6>.

accusations were based on administrative errors and inaccurate evidence. As a result of the racially biased algorithm assessments, the victims faced severe social and economic impacts, including job loss, bankruptcy, and even divorce. One of the victims had to pay a tax of €48,000, while 11,000 others underwent additional scrutiny simply because they had dual citizenship or names that sounded foreign.⁴⁸ This case underscores that an Automated Decision Making (ADM) system is not constantly controlled, it might perpetuate systematic prejudice and structural inequities. Therefore, to safeguard individual rights against the detrimental effects of automated technology, especially in the context of social assistance and public administration, extensive rules and efficient supervision systems are required.

Robodebt, the Australian government's debt collection scheme, is another unfortunate effect of using Automated Decision Making (ADM). This system was created to automatically compute overpayments and, if required, seek refunds by comparing data from social assistance beneficiaries with tax office income records. However, the algorithm used produced invalid debt estimates, causing thousands of citizens to experience unfair collections and significant suffering, including alleged cases of suicide. After increased criticism and litigation, the court declared the debt-calculation method illegal, ordering the government to halt the program and refund unlawfully collected funds.⁴⁹ The Robodebt case highlights the importance of strict regulations and effective oversight mechanisms in the use of ADM, especially when automated decisions directly affect citizens' welfare and fundamental rights.

The cases discussed highlight two main issues. First, the occurrence of injustice due to unequal treatment of individuals leads to discriminatory actions. Second, there is reliance on automated systems that cannot fully eliminate bias, considering that technology is never truly neutral. In this

⁴⁸ BBC, "Dutch Rutte Government Resigns over Child Welfare Fraud Scandal," BBC, January 2021, <https://www.bbc.com/news/world-europe-55674146>; Jon Henley, "Dutch Government Resigns over Child Benefits Scandal," *The Guardian*, January 2021, <https://www.theguardian.com/world/2021/jan/15/dutch-government-resigns-over-child-benefits-scandal>.

⁴⁹ Tapani Rinta-Kahila et al., "Algorithmic Decision-Making and System Destructiveness: A Case of Automatic Debt Recovery," *European Journal of Information Systems* 31, no. 3 (May 2022): 313–38, <https://doi.org/10.1080/0960085X.2021.1960905>.

context, it is important to distinguish between explicit and implicit bias. When a person or group is treated less favorably because of traits such as gender or color, this is known as direct bias, which is legally prohibited without exception. Conversely, indirect discrimination arises when rules or policies that appear formally neutral actually place certain groups in a less favorable factual position. In the realm of AI, this indirect discrimination is often difficult to detect. It is known as proxy discrimination, the use of seemingly neutral attributes as substitutes for legally prohibited ones, such as using an address to infer race or ethnicity.⁵⁰ A deep understanding of these types of discrimination is crucial in designing effective regulations to prevent hidden injustices in ADM systems.

The entire discussion in this subsection emphasizes that it carries significant legal risks, particularly related to injustice, discrimination, and algorithmic bias that are difficult to eliminate. Various cases from the financial sector, criminal law, education, and public administration show that automated decisions can seriously impact the rights of individuals and vulnerable groups. The limitations of transparency, lack of accountability, and technical complexity of ADM systems exacerbate these risks, creating an urgent need for strict regulations and effective oversight mechanisms. A deep understanding of the types of discrimination and their social impacts is a crucial foundation for designing a legal framework that ensures the defense of justice and rights when using such technology.

Legal Responsibility and AI Governance in Indonesia and Internationally

To manage the associated legal and societal risks, some nations have developed flexible regulatory and governance systems in response to advances in Artificial Intelligence (AI), notably in the form of Automated Decision Making (ADM) and Agentic AI. The AI Liability Directive (AILD) is an international

⁵⁰ Carlotta Rigotti and Eduard Fosch-Villaronga, “Fairness, AI & Recruitment,” *Computer Law & Security Review* 53 (July 2024): 105966, <https://doi.org/10.1016/j.clsr.2024.105966>; Kattinig et al., “Assessing Trustworthy AI: Technical and Legal Perspectives of Fairness in AI.”

proposal from the European Union. To hold AI manufacturers and service suppliers accountable for any harm the technology causes, the EU passed the Product Liability Directive (PLD). In the meantime, Congress is being asked to consider the Algorithmic Accountability Act (AAA) in the US to balance the standard's advantages and disadvantages through more stringent oversight and transparency.⁵¹

However, Indonesia currently lacks a comprehensive and precise legal framework that specifically regulates artificial intelligence (AI). As a result, legal regulation of AI is still scattered over a number of current sectoral legislation, including those pertaining to consumer protection, electronic information and transactions, and personal data protection. This disjointed approach may make it difficult to maintain legal clarity, regulatory consistency, and efficient supervision, especially when handling the unique traits and risks posed by more autonomous AI systems. Law Number 1 of 2024, the 2nd Amendment to Law Number 11 of 2008 on Electronic Information and Transactions (ITE Law), marked the beginning of Indonesia's efforts to govern AI. AI systems are included under the term "electronic agent" in the ITE Law. Senders and receivers may undertake electronic transactions directly, via designated individuals, or through electronic agents, according to Article 21 paragraph (1) of the ITE Law. Furthermore, paragraph (2) stipulates that legal responsibility for the execution of electronic transactions through electronic agents rests with the organizers of the electronic agents, except in cases of force majeure, error, or negligence by users of the electronic system, as stipulated in paragraph (5). Article 22 of the ITE Law also requires electronic agent organizers to provide features that allow users to make changes to information still in the transaction process.

⁵¹ Christiane Wendehorst, "Liability for Artificial Intelligence: The Need to Address Both Safety Risks and Fundamental Rights Risks," in *The Cambridge Handbook of Responsible Artificial Intelligence*, In S. Voen (Cambridge University Press, 2022), 187–209, <https://doi.org/10.1017/9781009207898.016>; Jakob Mökander et al., "The US Algorithmic Accountability Act of 2022 vs. The EU Artificial Intelligence Act: What Can They Learn from Each Other?," *Minds and Machines* 32, no. 4 (December 2022): 751–58, <https://doi.org/10.1007/s11023-022-09612-y>; Izquierdo Grau, "The Development Risks Defence in the Digital Age."

In addition to the Implementation of Electronic Systems and Transactions (ITE) Law, the use of AI in electronic systems is regulated by Government Regulation Number 71 of 2019, specifically in Chapter III, which pertains to electronic agent organizers. Articles 36 through 40 regulate organizers' obligations to ensure system security and integrity. The Minister of Communication and Informatics' KBLI code 62015, Regulation Number 3 of 2021, also establishes AI-based business standards. It requires electronic system organizers and business actors to create internal guidelines for data management and AI ethics. The Ministry of Research and Technology (Kemristek) and the National Research and Innovation Agency (BRIN) unveiled the National Artificial Intelligence Strategy (Stranas KA) 2020–2045 on August 10, 2022. This plan outlines focus areas and priority domains for AI technology for ministries, agencies, regions, and other stakeholders. The five major priority programs are healthcare services, bureaucratic reform, education and research, food security, and mobility and smart cities.

The Personal Data Protection Law (PDP Law), Law Number 27 of 2022, was thereafter approved. According to this law's Article 10, data subjects have the right to object to decisions made solely on the basis of automated processing, including profiling, that significantly affect them or have legal repercussions. Additionally, Article 34 paragraph (1) requires Personal Data Controllers to conduct a data protection impact assessment where processing personal data poses a serious danger to the data subject. The processing of personal data for systematic evaluation, scoring, or monitoring activities related to the data subject, as well as automated decision-making that has legal ramifications or significant effects on the data subject, are among the potential risks listed in Article 34 paragraph (2) letters a and d. The administrative consequences for violating Article 34 are governed by Article 57 and include administrative fines, written warnings, temporary suspension of the processing of personal data, and the deletion or destruction of personal data.

In 2024, the Minister of Communication and Information of the Republic of Indonesia published Circular Letter No. 9 of 2023 on Artificial Intelligence Ethics. It includes specific regulations on broad definitions, basic

ethical and values principles, and oversight of corporate actors' consultation, analysis, and AI-based programming activities. This SE has provided an explanation of the ethics of artificial intelligence, which are founded on the principles of inclusion, transparency, humanism, and security in the administration of accessible data resources. They act as the cornerstone for governing moral standards and values in the use of AI-based programming. It also dictates the need to prevent racism and other harmful behaviors, as well as to ensure that human judgments and policies are not made by artificial intelligence. This Circular presents Indonesia as a country that understands the potential problems ADM may pose. Similar to discriminatory behaviors, these AI decisions might be harmful and perhaps algorithmically biased. However, since it is a circular, it is not legally enforceable and those who break it face no consequences. Law No. 12 of 2011 on the Formation of Laws and Regulations states that (UU PPP), as modified by Law No. 13 of 2022. Under this statutory hierarchy, only instruments expressly recognized, such as statutes, government regulations, and presidential regulations, possess binding normative force, whereas circulars function merely as internal administrative guidance. This SE ultimately serves only as a guideline, opening up opportunities for violations against it, because Circular Letters do not fall under the category of legislation according to the Law on the Formation of Legislation in Indonesia. As this article is being written, the Ministry of Communication and Digital is drafting a roadmap and governance for the utilization of artificial intelligence and assisting in the preparation of a Presidential Regulation on AI.

In facing the emerging complexities and risks, the risk-based regulation paradigm has now become the main approach in the governance of artificial intelligence in many countries, including Indonesia. Unlike reactive (ex-post) regulation, which only responds to problems after they occur, risk regulation allows for the identification and mitigation of potential dangers that can be anticipated beforehand, thus providing more effective protection for individuals

and society. This approach is highly relevant in the context of AI, where many risks can be predicted before harmful impacts arise.⁵²

Indonesia also feels the impact of AI development, including the government, which is striving to benefit from it. One of the agencies that has utilized AI technology is the Ministry of Finance (Kemenkeu). Since 2021, data analytics has become a strategic initiative. Several data analytics projects have been developed, such as the Executive Information System (EIS), an information system for officials within the Ministry of Finance that provides aggregated and structured information/a single source of truth. The PNBPDashBoard, developed by the Directorate General of Budget to monitor and predict Non-Tax State Revenue (PNBP), then the SetPP Decision Prediction, a system elaborated with the e-Tax Court application from the Tax Court Secretariat that can provide case profiling predictions and case decision outcome predictions, and finally the Expert Locator (ExLoc), an application aimed at finding employees with expertise in certain fields from all Ministry of Finance employees. ExLoc uses as many as 6.2 million data indices from various sources. Not only that, but the Ministry of Finance is currently developing Law Analyzer, an application that can perform automatic legal searches and analyses using generative AI to assist rule makers. Then there's Nadine AI, an automation of correspondence using Gen AI with planned features such as summarising important information from incoming letters, semi-automatically replying to letters, creating letters more quickly and accurately, and performing intelligent inbox priority scoring. It is planned that AI will assist the government in monitoring and matching the assets owned by taxpayers. Through social media, the DJP will identify discrepancies if there are assets that have not been reported or differ from the SPT (Tax Return) or the LHKPN (State Officials' Wealth Report).⁵³

⁵² Beatriz Botero Arcila, "AI Liability in Europe: How Does It Complement Risk Regulation and Deal with the Problem of Human Oversight?," *Computer Law & Security Review* 54 (September 2024): 106012, <https://doi.org/10.1016/j.clsr.2024.106012>; Gianclaudio Malgieri and Frank Pasquale, "Licensing High-Risk Artificial Intelligence: Toward Ex Ante Justification for a Disruptive Technology," *Computer Law & Security Review* 52 (April 2024): 105899, <https://doi.org/10.1016/j.clsr.2023.105899>.

⁵³ Nopita Dewi, "DJP Kemenkeu Bakal Gunakan AI Untuk Mengawasi Wajib Pajak Di Medsos," *Metro TV News*, July 2025,

Another institution in Indonesia that has utilised AI is the Riau Islands Provincial Government, which, through the Communication and Information Office, has used AI in the form of creating leadership videos using deepfake and text-to-speech technology. The use of AI within the Riau Islands Provincial Government has been regulated through the Riau Islands Governor's Circular Letter number B/120/672.2/DKI-SET/2023, dated August 8, 2023.⁵⁴ Meanwhile, the East Java Education Office (Dindik) launched an AI-based application to assist the public in accessing information about the New Student Enrollment Selection (SPMB) for the year 2025. The application is named Senopati and is the result of a collaboration between the East Java Education Office (Dindik) and the Sepuluh Nopember Institute of Technology (ITS). The system allows parents or prospective students to inquire about various matters, ranging from the criteria for achievement pathways, the requirements for entering Senior High Schools/Vocational High Schools, to the ranking mechanism. Additionally, Senopati has access to the data on school sites, zoning quota status, and standardized registration processes that correspond with the context of potential new pupils. However, this application has received criticism from the East Java Information Commission, particularly regarding the transparency of information in the selection process. The Information Commission highlighted the issue of transparency in ranking data, especially regarding domicile and achievement pathways, as well as the use of AI technology in the selection process, due to public reports about ranking data not being fully disclosed. Furthermore, the use of AI has not been explained openly to the public. Parameters, assessment weights, and result audit mechanisms need to be communicated clearly.⁵⁵ So that citizen control cannot be fulfilled, to ensure that an individual's rights are fully protected. The East Java representative of the Indonesian Ombudsman also received 15 reports of

<https://www.metrotvnews.com/play/k8oCVgDe-djp-kemenkeu-bakal-gunakan-ai-untuk-mengawasi-wajib-pajak-di-medsos>.

⁵⁴ Pemprov Kepri, "Pemprov Kepri Manfaatkan Kecerdasan Buatan (AI) Untuk Meningkatkan Strategi Informasi Komunikatif," *Pemerintah Provinsi Kepulauan Riau*, August 2025, <https://www.kepriprov.go.id/berita/perangkat-daerah/pemprov-kepri-manfaatkan-kecerdasan-buatan-ai-untuk-meningkatkan-strategi-informasi-komunikatif>.

⁵⁵ Ima Karimah, "Minimnya Transparansi AI Dan Data Pemeringkatan Di SPMB 2025," *Lintasan*, July 2025.

alleged maladministration in the 2025 SPMB process, one of which is the ranking of the sports achievement pathway, which is considered non-transparent.⁵⁶ The Provincial Government has acknowledged several findings regarding the obstacles that occurred, such as discrepancies in report card scores and the system, improvement efforts and the establishment of a call center and temporary complaint posts have been offered as solutions.⁵⁷

There is a concern regarding the implementation and planned application of AI in Indonesian government policies, particularly in the distribution of social assistance funds (Bansos). The first concern is related to the limited information regarding the social assistance (Bansos) currently distributed through Pos Indonesia. After the previous distribution of social assistance in 2022, which involved a rather massive repetition of data recording, it turned out that the data recording, especially photos, had issues with photo quality, photo accuracy, and inaccurate geo-tagging. Finally, in 2023, Pos Indonesia utilized AI for the tapes. If the face is not straight in the photo, it will be immediately rejected. House photos will be immediately rejected if it turns out that what was photographed is a tree or another object that is not a house. Then, the geo-tagging will be compared between before and after. Previously, the recording of the recipient found that the person and the address were the same, but the geo tagging differed by 500 meters to 1 kilometre. However, there is no information on whether, in the end, technological constraints or data bias occurred. Thus, the implementation of AI is not supported by mitigation information or post-decision handling.⁵⁸

⁵⁶ Amaluddin, "Ombudsman Jatim Terima 15 Laporan SPMB 2025, Puluhan Difabel Gagal Lolos Jalur Afirmasi," *Metro TV News*, July 2025, <https://www.metrotvnews.com/read/koGCdq39-ombudsman-jatim-terima-15-laporan-spm-2025-puluhan-difabel-gagal-lolos-jalur-afirmasi>.

⁵⁷ Willi Irawan, "Gubernur Jatim Sebut AI Senopati Permudah Layanan SPMB 2025," *Antara News*, May 2025, <https://jatim.antaranews.com/berita/926381/gubernur-jatim-sebut-ai-senopati-permudah-layanan-spm-2025>.

⁵⁸ Deri Dahuri, "Didukung AI, Pos Indonesia Pastikan 3,2 Juta KPM Terima Bansos BPNT Dan PKH Sebelum Lebaran," *Media Indonesia*, April 2023, <https://mediaindonesia.com/ekonomi/575758/pos-indonesia-bantuan-uang-sem-bako-pensiunan-pelindo-cair-sebelum-idul-fitri>.

Meanwhile, another concern is the plan to implement a new policy conveyed by the Chairman of the National Economic Council (DEN), Luhut Binsar Pandjaitan, that in August 2025, a new system will be implemented, namely the distribution of social assistance determined through face recognition with data available in the government. The data of social assistance recipients, which has so far only consisted of administrative data, will be supplemented with face recognition. Luhut claims that the distribution of social assistance by implementing face recognition can save up to Rp100 trillion. Continuously updated data can determine the number of social assistance recipients. When a number of social assistance recipients get jobs in the following month, the data in the system will be updated, and automatically, those individuals will no longer be entitled to receive social assistance. This system will be able to distinguish between individuals who are entitled to receive social assistance today, next month, or next week, because it can see whether the recipient has found a job or not. This system will build an integrated digital ecosystem, making the distribution of social assistance more targeted and efficient in the use of the national budget.⁵⁹ The plan for this social assistance distribution program prioritizes efficiency issues, with minimal attention to data accuracy. If incorrect, this policy could violate Article 34 paragraphs (1) and (2) of the UUPDP, which controls the processing of personal data for assessment, scoring, or systematic monitoring of personal data subjects, as well as automated decision-making that has legal ramifications or substantial effects on personal data subjects. Because an automated system that still leaves room for mistakes determines a person's life, in this instance, social assistance. If this occurs, the legal system does not yet have precise rules governing who is in charge of the judgments made by AI; instead, it merely offers administrative consequences.

The group of experts within the European Commission's High Level Expert Group (HLEG) has emphasized that the issue of AI must be addressed

⁵⁹ CNN, "Luhut Dorong Distribusi Bansos Dibantu Pakai Teknologi Ini," *CNN*, June 2025, <https://www.cnnindonesia.com/teknologi/20250603053939-185-1235730/luhut-dorong-distribusi-bansos-dibantu-pakai-teknologi-ini>; Elsa Catriana, "Luhut Sebut Penyaluran Bansos Pakai AI Bisa Hemat Anggaran Rp 100 Triliun," *Kompas*, June 2025, <https://money.kompas.com/read/2025/06/02/161607726/luhut-sebut-penyalaran-bansos-pakai-ai-bisa-hemat-anggaran-rp-100-triliun>.

by ensuring that the entity responsible for the decision can be identified and the decision-making process can be explained.⁶⁰ It should be noted that most criminal law studies to date will not criminalize AI-related dangers, as Indonesia still imposes administrative sanctions for violations of AI-related regulations. This seems to be predicated on the idea that AI lacks moral agency and that humans must always bear accountability for machine outputs. Therefore, ensuring that AI does not impede the attribution of responsibility to humans or result in people being unjustly held liable for circumstances beyond their control is the primary concern in the ethical use of AI. This circumstance is called the “responsibility gap”.⁶¹

Although there have been efforts to incorporate criminal sanctions in the field of AI, currently, this view is based on the premise that the losses or impacts associated with AI are grounded in an interaction, such as the interaction between designers, users, and the environment in which AI operates. When the losses are the result of an interaction, it is necessary to delve into the qualifications of that interaction. This interaction can be malicious if intentionally designed or implemented to cause deliberate harm. Conversely, the interaction can be dangerous if designed or implemented legally, but still capable of causing unintentional damage.⁶² Therefore, for these two interactions, both criminal law enforcement and criminal law prevention are needed, considering both ex-ante and ex-post approaches simultaneously. In the case of unintentional interactions, operators cannot be held accountable unless the legal system establishes an obligation to act that requires them to be responsible for preventing harm. However, if the impact is deemed very dangerous, one might consider corporate criminal liability compared to human agent criminal liability in the context of AI.⁶³ As a follow-up to this challenge, it is essential to explore how international regulations can offer concrete solutions that Indonesia can adopt. This system of checks and balances enables

⁶⁰ Ljupcho Grozdanovski and Jerome De Cooman, “Individual Safeguards in the Era of AI: Fair Algorithms, Fair Regulation, Fair Procedures,” *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e18, <https://doi.org/10.1017/cfl.2025.10>.

⁶¹ Robert Sparrow and Gene Flenady, “The Testimony Gap: Machines and Reasons,” *Minds and Machines* 35, no. 1 (February 2025): 12, <https://doi.org/10.1007/s11023-025-09712-5>.

⁶² Beatrice Panattoni, “Generative AI and Criminal Law,” *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e9, <https://doi.org/10.1017/cfl.2024.9>.

⁶³ Panattoni.

government agencies to enforce laws while legislative bodies ensure compliance with fundamental principles.⁶⁴

As a concrete step in strengthening the legal framework, the adoption of regulations that have proven effective at the international level can serve as an important reference. For example, the AI Liability Directive (AILD) and the Product Liability Directive (PLD) that have been implemented in the European Union can provide a strong legal framework for Indonesia in regulating the liability of AI producers and service providers. This product-based and fault-based approach allows for clearer and more effective law enforcement against losses caused by AI technology. Thus, these two approaches have become important elements in building regulatory governance that ensures the responsible and fair use of AI in Indonesia.⁶⁵ However, effective regulation not only requires the adoption of international law but also a deep understanding of the technological and ethical complexities underlying AI itself.

Considering the ethical complexity and the multidisciplinary collaboration required, the aspect of legal regulation becomes very important to ensure that AI technology is not only developed responsibly but also regulated with an adequate legal framework. In Indonesia, although there have been initial guidelines, comprehensive and binding regulations remain urgently needed to fill existing legal gaps. In the context of regulation and legal accountability for AI-generated automatic decisions, a clear and comprehensive legal framework is required. In Indonesia, the current regulations remain guideline-based and do not comprehensively address all aspects of AI. Therefore, the presence of a specific law that integrates all those associated with AI is highly anticipated. A proper understanding of machine autonomy is also important; this autonomy refers to the system's ability to operate without human intervention for certain tasks, while humans still play a crucial role in maintaining the system's overall function. Even the most advanced autonomous systems, supported by large and complex machine learning models, usually still require significant human

⁶⁴ Anis Widiyawati et al., "Strengthening the Correctional System through Electronic Supervision of Prisoners: A Comparative Legal Study for Reforming Indonesia's Penitentiary Law," *Jurnal Hukum Novelty* 16, no. 2 (2025): 264.

⁶⁵ Shannon Vallor and Tillmann Vierkant, "Find the Gap: AI, Responsible Agency and Vulnerability," *Minds and Machines* 34, no. 3 (June 2024): 20, <https://doi.org/10.1007/s11023-024-09674-0>.

support to maintain their performance.⁶⁶ To address the need for comprehensive regulation, several innovative concepts and approaches have been proposed to strengthen accountability mechanisms.

In line with the need for clear regulations, several concepts and approaches have been proposed to address accountability challenges in AI use. One of them is the concept of computational accountability, which offers a collaborative mechanism to ensure that AI systems operate responsibly and transparently. Some ideas that can be applied in Indonesia for AI regulation include computational accountability, assessment of human-machine interactions, and adoption of policies such as the European Union's Product Liability Directive (PLD) and AI Liability Directive (AILD). Three strategies are provided by Joris Hulstijn's concept of computational accountability to guarantee responsible automated decision-making systems: a top-down strategy through ethical and legal guidelines; a bottom-up strategy by developers and activists concentrating on the creation and testing of high-quality AI; and an external strategy by independent auditors who confirm the statements made by developers and system controllers. The collaboration of these three approaches can be facilitated through a dialogue platform hosted by a trusted third party, enabling transparency and effective oversight.⁶⁷ In addition to these accountability mechanisms, the interaction aspects within AI systems must also be specifically considered when determining legal liability.

In addition to those accountability mechanisms, the interaction aspects of AI systems also require special attention when determining legal liability. The losses arising from the use of AI can stem from both intentional and unintentional interactions, so the legal approach must accommodate this complexity. Losses resulting from intentional interactions, designed to cause harm, must receive strict legal treatment. Conversely, losses arising from unintentional interactions, even within a legitimate framework, also pose significant challenges in law enforcement. In such situations, criminal law

⁶⁶ Vallor and Vierkant.

⁶⁷ Joris Hulstijn, "Computational Accountability," in *Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law* (New York, NY, USA: ACM, 2023), 121–30, <https://doi.org/10.1145/3594536.3595122>.

enforcement and prevention must consider both ex-ante and ex-post approaches simultaneously. However, tracing responsibility becomes difficult when many parties are involved, and there is no clear hierarchical structure. Therefore, regulations are needed to govern the obligations of information providers and transparency in the development and application of AI so that accountability can be effectively enforced.⁶⁸ Given this complexity, adopting well-tested international regulations is a strategic choice to strengthen legal governance in Indonesia.

Considering the complexity and challenges of establishing legal liability for AI, adopting international regulations such as the AI Liability Directive (AILD) and the Product Liability Directive (PLD) can be an effective way to strengthen governance and legal protection. The adoption of the AI Liability Directive (AILD) and the Product Liability Directive (PLD), which have been implemented in the European Union, can provide a strong legal framework for Indonesia in regulating the liability of AI producers and service providers. This product and fault-based approach allows for clearer and more effective law enforcement against losses caused by AI technology. Thus, these two approaches have become important elements in building regulatory governance that ensures the responsible and fair use of AI in Indonesia.

Although regulations like this provide a solid legal framework, real challenges arise when they must be adapted and implemented in developing countries like Indonesia, which have social, economic, and technological contexts that differ from those of developed countries. Developing countries face unique challenges in regulating AI technology, including limitations in technological infrastructure, human resources, and still-developing regulatory capacity. Moreover, existing social and economic gaps can exacerbate the risk of injustice if regulations are not designed in a context-specific and inclusive manner. Therefore, AI regulation in developing countries must consider local conditions and the specific needs of the community, without neglecting

⁶⁸ Roos de Jong, "The Retribution-Gap and Responsibility-Loci Related to Robots and Automated Technologies: A Reply to Nyholm," *Science and Engineering Ethics* 26, no. 2 (April 2020): 727–35, <https://doi.org/10.1007/s11948-019-00120-4>.

international standards and principles. This responsive and adaptive approach is important to ensure that AI technology can provide maximum benefits while minimizing potential social and legal risks.

While the adoption of international instruments such as the AI Liability Directive (AILD) and Product Liability Directive (PLD) is recommended, their feasibility in Indonesia requires careful consideration. Indonesia's civil law system, limited judicial resources, and corporate governance structures necessitate adaptation, including clear statutory authority, defined procedural mechanisms, and practical enforcement strategies to ensure these liability frameworks function effectively.

Overall, the development of AI technology demands an adaptive and comprehensive regulatory framework to ensure effective legal accountability. An approach that combines mechanisms of computational accountability, product and fault-based regulation, and the adoption of international standards such as the AI Liability Directive and the Product Liability Directive is an important foundation in building responsible AI governance. In Indonesia, the need for binding, integrated regulations is becoming increasingly urgent to address existing legal gaps.

Thus, effective AI governance requires structured collaboration between policymakers, technology developers, and society. Such collaboration should translate into the development of a clear and comprehensive AI regulatory framework, the adoption of a risk-based approach that emphasizes transparency, human oversight, and safety-by-design principles, and the clarification of liability standards for autonomous systems. At the same time, consumer protection mechanisms and accessible avenues for legal redress must be strengthened to ensure that technological innovation remains aligned with legal certainty, fairness, and public safety.

The inclusion of explicit statutory mandates, such as a future Presidential Regulation or amendments to the ITE and PDP Laws, is essential to move beyond guideline-based instruments like ministerial circulars and establish enforceable obligations for AI developers, users, and public institutions. Such instruments would improve accountability, reduce the "responsibility gap," and provide legal clarity for autonomous AI operations in Indonesia.

Ethical, Social, and Regulatory Implications of Future Agentic AI

The non-neutrality of technology is a fundamental issue in the development and application of Artificial Intelligence (AI). Lawrence Lessig's views on the four modalities of regulation—law, social norms, market, and architecture—provide a comprehensive framework for understanding. Lessig emphasizes that in the current technological era, positive law alone is not sufficient to regulate societal behavior. Rules in the digital world are not only shaped by formal government laws but also by code in the form of software, protocols, and algorithms that control how technology operates and influences user behavior. These four modalities interact with each other and can strengthen or weaken one another. A concrete example is the modality of architecture or code in the form of algorithms, which directly regulates the behavior of the system and its users.⁶⁹ Understanding the concept of regulatory modality is important because algorithms not only execute commands but also absorb and reflect the social and technical conditions in which they are designed and implemented. This opens up the possibility of algorithmic bias emerging, which can reinforce social injustices.

It is important to clarify that Agentic AI represents an emerging “third wave” of AI technology, characterized by systems capable of autonomous decision-making in complex environments. Currently, such systems are not widely deployed in Indonesia, and most applications are limited to semi-autonomous or rule-based automation. Recognizing this distinction helps separate current regulatory needs from hypothetical future governance challenges.

AI systems' algorithms take in and mirror the circumstances around their selection, development, and use. As a result, algorithmic bias may appear at different phases in the lifespan of an AI system. This bias may come from

⁶⁹ Ahmed Ragib Chowdhury, “Techno-Authoritarianism & Copyright Issues of User-Generated Content on Social-Media,” *Computer Law & Security Review* 55 (November 2024): 106068, <https://doi.org/10.1016/j.clsr.2024.106068>.

societal inequality already present in the data, the computational model itself, or the data used to train the algorithm. Because they identify patterns and connections in historical data, algorithms trained on it can perpetuate preconceived notions. Measurement bias, representation bias, aggregate bias, omitted variable bias, and connection bias are common types of prejudice.(2024, Kinchin) Regulations like the EU AI Act, which will take effect in 2024, require system providers to assess and reduce bias in their training and testing data, given the awareness of this potential bias. However, since sensitive data is collected for de-biasing purposes, these restrictions are also criticized for potential privacy violations.

The EU AI Act's Article 10 requires system suppliers to assess training, validation, and testing datasets to ensure quality and reduce bias. However, the GDPR's ban on collecting sensitive data is contradicted by Article 10, paragraph 5, which permits the gathering of sensitive data for this reason. This has drawn harsh criticism, as the approach may compromise personal information and privacy, posing a conflict between the objective of de-biasing and data protection. In the end, this rule forces people to choose between preserving their privacy and risking prejudice.⁷⁰ In facing this dilemma, transparency is crucial to mitigating the negative impact of algorithmic bias by providing data subjects with a better understanding of the automated decision-making process.

Applying these ethical principles to Indonesia reveals real-world dilemmas, even with semi-autonomous AI systems. In healthcare, some Indonesian hospitals have piloted AI-based diagnostic tools for radiology. Limited transparency in decision-making processes has raised concerns among clinicians about accountability when AI suggestions conflict with medical judgment. In mass surveillance, the use of AI by local governments for traffic monitoring or public security in Jakarta has led to debates on privacy and citizen consent. Similarly, in judicial pilot programs, AI-assisted case management tools in regional courts have raised procedural concerns regarding the explanation of

⁷⁰ Marvin van Bekkum, "Using Sensitive Data to De-Bias AI Systems: Article 10(5) of the EU AI Act," *Computer Law & Security Review* 56 (April 2025): 106115, <https://doi.org/10.1016/j.clsr.2025.106115>.

automated recommendations, though full-scale Agentic AI deployment remains hypothetical.

The transparency principle, sometimes referred to as "citizen control," enables data subjects to comprehend how AI systems make choices that impact them. A number of technical and regulatory mechanisms, such as required algorithmic impact assessments, model documentation (e.g., model cards and data sheets for datasets), and user-facing explanations of the rationale, goals, and implications of automated decision-making, can be used to operationalize this principle in practice. While the Artificial Intelligence Act establishes transparency requirements for some AI systems, including disclosure requirements when users interact with AI-generated content, regulatory frameworks such as the General Data Protection Regulation (GDPR) mandate the provision of meaningful information about the logic underlying automated processing. This idea distinguishes between qualified transparency, which offers a practical and useful degree of control, and absolute transparency, which is challenging to attain because of the intrinsic complexity of AI systems, including millions of parameters and learning processes that are incomprehensible to humans. Although absolute transparency is hard to attain, transparency remains an important prerequisite for individuals to have the information needed to raise objections or challenge decisions they consider detrimental. These efforts encourage the development of Explainable Artificial Intelligence (XAI), which aims to make AI easier to understand and explain, and reinforce the concept of meaningful human control, which places humans as the primary decision-makers in automated decision-making.⁷¹

In line with the importance of transparency and human control, the analysis of contemporary case studies depicting ethical and social dilemmas in the application of autonomous AI is highly relevant for deepening understanding of the real challenges faced. As a concrete illustration, the

⁷¹ Florian J. Boge and Axel Mosig, "Put It to the Test: Getting Serious About Explanation in Explainable Artificial Intelligence," *Minds and Machines* 35, no. 2 (June 2025): 26, <https://doi.org/10.1007/s11023-025-09724-1>; Filippo Santoni de Sio et al., "Realising Meaningful Human Control Over Automated Driving Systems: A Multidisciplinary Approach," *Minds and Machines* 33, no. 4 (July 2022): 587–611, <https://doi.org/10.1007/s11023-022-09608-8>; Palmiotto, "When Is a Decision Automated? A Taxonomy for a Fundamental Rights Analysis."

application of Artificial Intelligence (AI) in the justice system to predict recidivism risk has sparked significant controversy, particularly regarding racial bias and procedural justice.⁷² This case emphasizes that AI regulation should not only focus on technical aspects but also prioritize the protection of human rights and fundamental principles of social justice. On the other hand, the use of AI in mass surveillance raises serious concerns regarding privacy violations and civil liberties, necessitating a strict, transparent, and accountable regulatory framework to prevent potential misuse of the technology.⁷³ Furthermore, the healthcare sector, which is beginning to adopt AI for diagnosis and treatment, presents complex ethical dilemmas, particularly concerning the accuracy of diagnostic results, the protection of patient data privacy, and the legal and social consequences of automated decisions that can directly impact human lives.⁷⁴ These case studies clearly emphasize that effective AI governance must holistically integrate technical, legal, and social aspects to ensure the responsible and equitable use of technology.

The necessity for human control and transparency in automated decision-making is growing in tandem with the complexity of these issues. As a result, there is an urgent need to acknowledge and regulate newly created rights, especially to address the unique traits and risks associated with contemporary AI technology. The advancement of AI technology requires identifying and protecting new rights specific to automated decision-making, as well as addressing issues of transparency and human control. These rights are crucial to

⁷² Michael Mayowa Farayola et al., “Enhancing Algorithmic Fairness: Integrative Approaches and Multi-Objective Optimization Application in Recidivism Models,” in *Proceedings of the 19th International Conference on Availability, Reliability and Security* (New York, NY, USA: ACM, 2024), 1–10, <https://doi.org/10.1145/3664476.3669978>.

⁷³ Giulia Gabrielli, “The Use of Facial Recognition Technologies in the Context of Peaceful Protest: The Risk of Mass Surveillance Practices and the Implications for the Protection of Human Rights,” *European Journal of Risk Regulation* 16, no. 2 (June 2025): 514–41, <https://doi.org/10.1017/err.2025.26>; P. Kitos, “The Limits of Government Surveillance: Law Enforcement in the Age of Artificial Intelligence,” *CEUR Workshop Proceedings* 2844 (2020): 164–68.

⁷⁴ Sourav Madhur Dey and Pushan Kumar Dutta, “From Code to Care and Navigating Ethical Challenges in AI Healthcare,” in *In Human-Centered Approaches in Industry 5.0: Human-Machine Interaction, Virtual Reality Training, and Customer Sentiment Analysis*, 2024, 210–25, <https://doi.org/10.4018/979-8-3693-2647-3.ch009>; Adan Khan, Venus Oseyi Omokhoa, and Safa Inaya Afzal, “Beyond the Hype: The Realities of AI in Clinical Practice,” in *In Clinical Practice and Unmet Challenges in AI-Enhanced Healthcare Systems*, 2024, 1–22, <https://doi.org/10.4018/979-8-3693-2703-6.ch001>.

addressing the issues arising from the transition from human decision-makers with moral agency and intuition to automated systems lacking these qualities. Therefore, the recognition of these new rights can serve as a strong legal foundation to prevent individuals from being subjected to Automated Decision Making (ADM) decisions that are harmful and unfair.⁷⁵

Understanding these new rights needs to be enriched with an in-depth study of computer ethics, particularly regarding the limitations and responsibilities of machine decision-making, which is the main focus of current technology philosophy discourse. In the realm of computer ethics, the work of James H. Moor since 1979 has become a highly relevant fundamental reference. Moor posed a critical question: “Are There Decisions That Should Not Be Made by Computers?” This question is further elaborated into three main issues, namely: Do computers really make decisions? To what extent can the decision-making competence of computers be measured? And how can the validity of that competence be assessed?⁷⁶ Moor concludes that computers should not take over decisions involving fundamental human goals and values. The decision to delegate decision-making authority to computers should be based on considerations of whether the use of computers can effectively promote human values and achieve human goals. Moor formulated three main ethical maxims: computers should not make decisions that humans want to make; computers should not make decisions that humans can make more competently; and computers should not make decisions that cannot be reversed or corrected by humans.⁷⁷ Furthermore, Moor offers three strategic recommendations for the more mature development of ethics in the ever-evolving technology. First, ethics must be dynamic and become a continuous process throughout the technology development cycle, from the conceptualisation stage to implementation and evaluation. Second, close collaboration between researchers, technology developers, ethicists, and social

⁷⁵ Abrusci and Mackenzie-Gray Scott, “The Questionable Necessity of a New Human Right against Being Subject to Automated Decision-Making.”

⁷⁶ Philip A. E. Brey, “The Historical Development of Ethics of Emerging Technologies,” *Minds and Machines* 35, no. 2 (March 2025): 21, <https://doi.org/10.1007/s11023-025-09720-5>.

⁷⁷ Deborah G. Johnson, “Moor’s ‘Are There Decisions Computers Should Never Make?,’” *Minds and Machines* 35, no. 2 (March 2025): 19, <https://doi.org/10.1007/s11023-025-09719-y>.

scientists is essential to ensure comprehensive and contextual ethical analysis. Third, ethical analysis must become increasingly sophisticated and adaptive, capable of anticipating the complexities and implications of ever-evolving modern technology.⁷⁸

In the context of such multidisciplinary collaboration, the active involvement of civil society and users becomes an important element that cannot be overlooked. The persistence of these issues gravely violates human rights and warrants further consideration of alternative penitentiary laws.⁷⁹ The repeated problems with prisons have been called long-term attacks on human rights. Their role in overseeing and controlling the use of AI technology is crucial to ensure ethical and responsible governance. The involvement of civil society, users, and oversight organizations is a key component in sustainable and democratic AI governance. Through participatory mechanisms such as public forums, AI incident reporting, and independent audits, the community can provide effective oversight of AI development and implementation. This active participation not only enhances transparency and accountability but also strengthens public trust in increasingly complex technology. Thus, empowering civil society becomes an important foundation in creating an AI ecosystem that is not only innovative but also socially and legally fair and responsible.

Conclusion

The rapid development of Artificial Intelligence (AI) technology, particularly in the form of Automated Decision Making (ADM) and Agentic AI, has presented complex and multidimensional regulatory and ethical challenges. In facing this development, the existing regulations in Indonesia are still not fully adequate to comprehensively and integratively regulate AI technology. Although there are several regulations and guidelines, such as the

⁷⁸ John Weckert, "Moor on Ethics for Emerging Technologies: Some Environmental Considerations," *Minds and Machines* 35, no. 2 (March 2025): 18, <https://doi.org/10.1007/s11023-025-09721-4>.

⁷⁹ Anis Widyawati et al., "Crafting an Ideal Penitentiary Law: A Global Comparative Framework for Indonesia's Correctional System," *Legality: Jurnal Ilmiah Hukum* 33, no. 2 (2025): 418, <https://doi.org/10.22219/ljih.v33i2.40358>.

ITE Law, the PDP Law, Government Regulations, and the Minister of Communication and Information's Circular, these regulations have not fully addressed the challenges arising from the increasingly autonomous and adaptive characteristics of AI. As a result, this creates legal uncertainty and potential risks that could impact the fundamental rights of society. As explained by Lawrence Lessig, regulation in the digital era does not only rely on formal law but is also influenced by code in the form of software, protocols, and algorithms that govern the behaviour of technology and its users. The four modalities of regulation—law, social norms, market, and architecture—interact and determine how technology operates and impacts society. Therefore, AI regulation must adopt a proactive and comprehensive paradigm, such as risk-based regulation, which involves risk assessment, mitigation, and oversight in a transparent and accountable manner. Priorities for Indonesia might include creating an independent supervisory body to oversee compliance, enforce laws, and guarantee the protection of individual rights, as well as clearly defining legal liability for damages caused by AI systems and determining who is accountable in cases of automated decision-making errors.

James H. Moor's ethical perspective emphasizes that machines shouldn't make choices about core human values and objectives. When deciding whether to give computers decision-making power, it is important to consider how using computers may advance human ideals and achieve human objectives. To address the complexity of contemporary technology, Moor also emphasizes the importance of dynamic ethics, interdisciplinary collaboration, and nuanced ethical analysis. With that basis, the idea of computational accountability becomes a crucial strategy that involves cooperation among independent auditors, technology developers, regulators, and civil society. This approach enables clear legal accountability while accommodating the complexity of human-machine interactions that can cause harm, whether intentional or unintentional. However, Indonesia faces challenges in designing effective AI regulations, given resource and infrastructure limitations and a still-developing legal capacity. Therefore, regulations must be formulated in a context-specific and inclusive manner, taking into account local needs without neglecting

international standards. Thus, effective regulation needs to integrate technical, legal, social, and ethical aspects to protect consumer rights and the broader community while promoting responsible technological innovation.

Furthermore, empowering civil society and users in AI governance is crucial to ensuring transparency, accountability, and public trust. Active community participation in oversight, incident reporting, and AI audits can strengthen the regulatory framework and prevent the misuse of potentially harmful technology. Therefore, cross-sector collaboration among government, industry, academia, and society is key to building a fair, safe, and sustainable AI ecosystem. Considering these factors, Indonesia has a strategic opportunity to develop and implement comprehensive, adaptive, and socially just AI regulations. However, potential obstacles such as limited institutional capacity, budget constraints, and the need for technical expertise may affect the pace and effectiveness of implementation. This regulation must address the challenges of machine autonomy, algorithmic bias, and the complexity of human-machine interactions, and ensure that AI technology can provide maximum benefits for national development and the protection of human rights. Overall, a holistic and collaborative regulatory approach will serve as the foundation for optimal risk management and the potential of AI in the ever-evolving digital era.

References

- Abrusci, Elena, and Richard Mackenzie-Gray Scott. "The Questionable Necessity of a New Human Right against Being Subject to Automated Decision-Making." *International Journal of Law and Information Technology* 31, no. 2 (August 2023): 114–43. <https://doi.org/10.1093/ijlit/eaad013>.
- Acharya, Deepak Bhaskar, Karthigeyan Kuppan, and B. Divya. "Agentic AI: Autonomous Intelligence for Complex Goals—A Comprehensive Survey." *IEEE Access* 13 (2025): 18912–36. <https://doi.org/10.1109/ACCESS.2025.3532853>.
- Allen, Jason Grant, Jane Loo, and Jose Luis Luna Campoverde. "Governing Intelligence: Singapore's Evolving AI Governance Framework." *Cambridge*

- Forum on AI: Law and Governance* 1 (January 2025): e12.
<https://doi.org/10.1017/cfl.2024.12>.
- Amaluddin. "Ombudsman Jatim Terima 15 Laporan SPMB 2025, Puluhan Difabel Gagal Lolos Jalur Afirmasi." *Metro TV News*, July 2025.
<https://www.metrotvnews.com/read/koGCdq39-ombudsman-jatim-terima-15-laporan-spmb-2025-puluhan-difabel-gagal-lolos-jalur-afirmasi>.
- Arcila, Beatriz Botero. "AI Liability in Europe: How Does It Complement Risk Regulation and Deal with the Problem of Human Oversight?" *Computer Law & Security Review* 54 (September 2024): 106012.
<https://doi.org/10.1016/j.clsr.2024.106012>.
- Bao, Luye, Nicole M Krause, Mikhaila N Calice, Dietram A Scheufele, Christopher D Wirz, Dominique Brossard, Todd P Newman, and Michael A Xenos. "Whose AI? How Different Publics Think about AI and Its Social Impacts." *Computers in Human Behavior* 130 (2022): 107182.
<https://doi.org/10.1016/j.chb.2022.107182>.
- BBC. "Alexa Tells 10-Year-Old Girl to Touch Live Plug with Penny. BBC." *BBC*, December 2021. <https://www.bbc.co.uk/news/technology-59810383>.
- . "Dutch Rutte Government Resigns over Child Welfare Fraud Scandal." *BBC*, January 2021.
<https://www.bbc.com/news/world-europe-55674146>.
- Bekkum, Marvin van. "Using Sensitive Data to De-Bias AI Systems: Article 10(5) of the EU AI Act." *Computer Law & Security Review* 56 (April 2025): 106115. <https://doi.org/10.1016/j.clsr.2025.106115>.
- Boge, Florian J., and Axel Mosig. "Put It to the Test: Getting Serious About Explanation in Explainable Artificial Intelligence." *Minds and Machines* 35, no. 2 (June 2025): 26. <https://doi.org/10.1007/s11023-025-09724-1>.
- Brey, Philip A. E. "The Historical Development of Ethics of Emerging Technologies." *Minds and Machines* 35, no. 2 (March 2025): 21.
<https://doi.org/10.1007/s11023-025-09720-5>.
- Catriana, Elsa. "Luhut Sebut Penyaluran Bansos Pakai AI Bisa Hemat Anggaran Rp 100 Triliun." *Kompas*, June 2025.
<https://money.kompas.com/read/2025/06/02/161607726/luhut-sebut-penyaluran-bansos-pakai-ai-bisa-hemat-anggaran-rp-100-triliun>.
- Chowdhury, Ahmed Ragib. "Techno-Authoritarianism & Copyright Issues of User-Generated Content on Social-Media." *Computer Law & Security Review* 55 (November 2024): 106068.
<https://doi.org/10.1016/j.clsr.2024.106068>.

- CNN. "Luhut Dorong Distribusi Bansos Dibantu Pakai Teknologi Ini." *CNN*, June 2025. <https://www.cnnindonesia.com/teknologi/20250603053939-185-1235730/luhut-dorong-distribusi-bansos-dibantu-pakai-teknologi-ini>.
- Covilla, Juan Carlos. "Artificial Intelligence and Administrative Discretion: Exploring Adaptations and Boundaries." *European Journal of Risk Regulation* 16, no. 1 (March 2025): 36–50. <https://doi.org/10.1017/err.2024.76>.
- Csatlós, Erzsébet. "Prospective Implementation of Ai for Enhancing European (in)Security: Challenges in Reasoning of Automated Travel Authorization Decisions." *Computer Law & Security Review* 54 (September 2024): 105995. <https://doi.org/10.1016/j.clsr.2024.105995>.
- Custers, Bart, and Helena Vrabec. "Tell Me Something New: Data Subject Rights Applied to Inferred Data and Profiles." *Computer Law & Security Review* 52 (April 2024): 105956. <https://doi.org/10.1016/j.clsr.2024.105956>.
- Dahuri, Deri. "Didukung AI, Pos Indonesia Pastikan 3,2 Juta KPM Terima Bansos BPNT Dan PKH Sebelum Lebaran." *Media Indonesia*, April 2023. <https://mediaindonesia.com/ekonomi/575758/pos-indonesia-bantuan-uan-g-semako-pensiunan-pelindo-cair-sebelum-idul-fitri>.
- Derave, Charly, Nathan Genicot, and Nina Hetmanska. "The Risks of Trustworthy Artificial Intelligence: The Case of the European Travel Information and Authorisation System." *European Journal of Risk Regulation* 13, no. 3 (September 2022): 389–420. <https://doi.org/10.1017/err.2022.5>.
- Dewi, Nopita. "DJP Kemenkeu Bakal Gunakan AI Untuk Mengawasi Wajib Pajak Di Medsos." *Metro TV News*, July 2025. <https://www.metrotvnews.com/play/k8oCVgDe-djp-kemenkeu-bakal-gunakan-ai-untuk-mengawasi-wajib-pajak-di-medsos>.
- Dey, Sourav Madhur, and Pushan Kumar Dutta. "From Code to Care and Navigating Ethical Challenges in AI Healthcare." In *Human-Centered Approaches in Industry 5.0: Human-Machine Interaction, Virtual Reality Training, and Customer Sentiment Analysis*, 210–25, 2024. <https://doi.org/10.4018/979-8-3693-2647-3.ch009>.
- Farayola, Michael Mayowa, Malika Bendeche, Takfarinas Saber, Regina Connolly, and Irina Tal. "Enhancing Algorithmic Fairness: Integrative Approaches and Multi-Objective Optimization Application in Recidivism

- Models.” In *Proceedings of the 19th International Conference on Availability, Reliability and Security*, 1–10. New York, NY, USA: ACM, 2024. <https://doi.org/10.1145/3664476.3669978>.
- Fosch-Villaronga, Eduard, Antoni Mut-Piña, Mohammed Raiz Shaffique, and Marie Schwed-Shenker. “Exploring Fairness in Service Robotics.” *Cambridge Forum on AI: Law and Governance* 1 (June 2025): e28. <https://doi.org/10.1017/cfl.2025.10013>.
- Fry, Veve. “An Overview of the Risk-Based Model of AI Governance.” *ArXiv Preprint ArXiv:2507.15299*, 2025.
- Gabrielli, Giulia. “The Use of Facial Recognition Technologies in the Context of Peaceful Protest: The Risk of Mass Surveillance Practices and the Implications for the Protection of Human Rights.” *European Journal of Risk Regulation* 16, no. 2 (June 2025): 514–41. <https://doi.org/10.1017/err.2025.26>.
- Gacutan, Joshua, and Niloufer Selvadurai. “A Statutory Right to Explanation for Decisions Generated Using Artificial Intelligence.” *International Journal of Law and Information Technology* 28, no. 3 (2020): 193–216. <https://doi.org/10.1093/ijlit/ehaa016>.
- Galli, Federico, Andrea Loreggia, and Giovanni Sartor. “The Regulation of Content Moderation.” In *International Conference on the Legal Challenges of the Fourth Industrial Revolution*, 63–87. Springer, 2022.
- Grozdanovski, Ljupcho, and Jerome De Cooman. “Individual Safeguards in the Era of AI: Fair Algorithms, Fair Regulation, Fair Procedures.” *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e18. <https://doi.org/10.1017/cfl.2025.10>.
- Hartmann, Jacques, Eva Jueptner, Santiago Matalonga, James Riordan, and Samuel White. “Artificial Intelligence, Autonomous Drones and Legal Uncertainties.” *European Journal of Risk Regulation* 14, no. 1 (March 2023): 31–48. <https://doi.org/10.1017/err.2022.15>.
- Hendrickx, Victoria. “Rethinking the Judicial Duty to State Reasons in the Age of Automation?” *Cambridge Forum on AI: Law and Governance* 1 (June 2025): e26. <https://doi.org/10.1017/cfl.2025.11>.
- Henley, Jon. “Dutch Government Resigns over Child Benefits Scandal.” *The Guardian*, January 2021. <https://www.theguardian.com/world/2021/jan/15/dutch-government-resigns-over-child-benefits-scandal>.
- Hulstijn, Joris. “Computational Accountability.” In *Proceedings of the*

- Nineteenth International Conference on Artificial Intelligence and Law*, 121–30. New York, NY, USA: ACM, 2023. <https://doi.org/10.1145/3594536.3595122>.
- Irawan, Willi. “Gubernur Jatim Sebut AI Senopati Permudah Layanan SPMB 2025.” *Antara News*, May 2025. <https://jatim.antaranews.com/berita/926381/gubernur-jatim-sebut-ai-senopati-permudah-layanan-spmb-2025>.
- Izquierdo Grau, Guillem. “The Development Risks Defence in the Digital Age.” *European Journal of Risk Regulation* 16, no. 1 (March 2025): 197–216. <https://doi.org/10.1017/err.2024.43>.
- Jackson, Fiona. “Stay-at-Home Dad’s ‘Life Ruined’ by Google after He Was Investigated When Photos He Took of His Sick Toddler Son’s Genitals for a Doctor to Review during the Pandemic Were Flagged by AI as Potential Child Sexual Abuse Material.” *Daily Mail*, August 2022. <https://www.dailymail.co.uk/profile-223/fiona-jackson.html?page=7>.
- Johnson, Deborah G. “Moor’s ‘Are There Decisions Computers Should Never Make?’” *Minds and Machines* 35, no. 2 (March 2025): 19. <https://doi.org/10.1007/s11023-025-09719-y>.
- Jong, Roos de. “The Retribution-Gap and Responsibility-Loci Related to Robots and Automated Technologies: A Reply to Nyholm.” *Science and Engineering Ethics* 26, no. 2 (April 2020): 727–35. <https://doi.org/10.1007/s11948-019-00120-4>.
- Karimah, Ima. “Minimnya Transparansi AI Dan Data Pemeringkatan Di SPMB 2025.” *Lintasan*, July 2025.
- Kattnig, Markus, Alessa Angerschmid, Thomas Reichel, and Roman Kern. “Assessing Trustworthy AI: Technical and Legal Perspectives of Fairness in AI.” *Computer Law & Security Review* 55 (November 2024): 106053. <https://doi.org/10.1016/j.clsr.2024.106053>.
- Khan, Adan, Venus Oseyi Omokhoa, and Safa Inaya Afzal. “Beyond the Hype: The Realities of AI in Clinical Practice.” In *In Clinical Practice and Unmet Challenges in AI-Enhanced Healthcare Systems*, 1–22, 2024. <https://doi.org/10.4018/979-8-3693-2703-6.ch001>.
- Kiesow Cortez, Elif, and Nestor Maslej. “Adjudication of Artificial Intelligence and Automated Decision-Making Cases in Europe and the USA.” *European Journal of Risk Regulation* 14, no. 3 (September 2023): 457–75. <https://doi.org/10.1017/err.2023.61>.
- Kinchin, Niamh. “‘Voiceless’: The Procedural Gap in Algorithmic Justice.”

- International Journal of Law and Information Technology* 32 (June 2024): eaae024. <https://doi.org/10.1093/ijlit/eaae024>.
- Kitos, P. "The Limits of Government Surveillance: Law Enforcement in the Age of Artificial Intelligence." *CEUR Workshop Proceedings* 2844 (2020): 164–68.
- Korecki, Marcin, Guillaume Köstner, Emanuele Martinelli, and Cesare Carissimo. "The Man Behind the Curtain: Appropriating Fairness in AI." *Minds and Machines* 34, no. 1 (April 2024): 7. <https://doi.org/10.1007/s11023-024-09669-x>.
- Krafft, Tobias D., Marc P. Hauer, and Katharina Zweig. "Black-Box Testing and Auditing of Bias in ADM Systems." *Minds and Machines* 34, no. 2 (May 2024): 15. <https://doi.org/10.1007/s11023-024-09666-0>.
- Kshetri, Nir. "From Predictive and Generative to Agentic AI: Shaping the Future of Marketing Operations and Strategies." *Computer* 58, no. 4 (April 2025): 121–29. <https://doi.org/10.1109/MC.2025.3530304>.
- Malgieri, Gianclaudio. "Automated Decision-Making and Data Protection in Europe." In *Research Handbook on Privacy and Data Protection Law*, 433–48. Edward Elgar Publishing, 2022.
- Malgieri, Gianclaudio, and Frank Pasquale. "Licensing High-Risk Artificial Intelligence: Toward Ex Ante Justification for a Disruptive Technology." *Computer Law & Security Review* 52 (April 2024): 105899. <https://doi.org/10.1016/j.clsr.2023.105899>.
- McGraw, David K. "Ethical Responsibility in the Design of Artificial Intelligence (AI) Systems." *International Journal on Responsibility* 7, no. 1 (2024): 4. <https://doi.org/10.62365/2576-0955.1114>.
- Mökander, Jakob, Prathm Juneja, David S. Watson, and Luciano Floridi. "The US Algorithmic Accountability Act of 2022 vs. The EU Artificial Intelligence Act: What Can They Learn from Each Other?" *Minds and Machines* 32, no. 4 (December 2022): 751–58. <https://doi.org/10.1007/s11023-022-09612-y>.
- Mühlhoff, Rainer, and Hannah Ruschemeier. "Updating Purpose Limitation for AI: A Normative Approach from Law and Philosophy." *International Journal of Law and Information Technology* 33 (January 2025): eaaf003. <https://doi.org/10.1093/ijlit/eaaf003>.
- Noguera, Angela Maria. "Case C-634/21, OQ v. Land Hessen (C.J.E.U.)." *International Legal Materials* 63, no. 5 (October 2024): 946–61. <https://doi.org/10.1017/ilm.2024.23>.

- Noto La Diega, Guido, and Leonardo C T Bezerra. "Can There Be Responsible AI without AI Liability? Incentivizing Generative AI Safety through Ex-Post Tort Liability under the EU AI Liability Directive." *International Journal of Law and Information Technology* 32 (June 2024): eaae021. <https://doi.org/10.1093/ijlit/eaae021>.
- Pagallo, Ugo. "LLMs Meet the AI Act: Who's the Sorcerer's Apprentice?" *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e1. <https://doi.org/10.1017/cfl.2024.6>.
- Palmiotto, Francesca. "When Is a Decision Automated? A Taxonomy for a Fundamental Rights Analysis." *German Law Journal* 25, no. 2 (March 2024): 210–36. <https://doi.org/10.1017/glj.2023.112>.
- Panattoni, Beatrice. "Generative AI and Criminal Law." *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e9. <https://doi.org/10.1017/cfl.2024.9>.
- Paseri, Ludovica, and Massimo Durante. "Examining Epistemological Challenges of Large Language Models in Law." *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e7. <https://doi.org/10.1017/cfl.2024.7>.
- Pemprov Kepri. "Pemprov Kepri Manfaatkan Kecerdasan Buatan (AI) Untuk Meningkatkan Strategi Informasi Komunikatif." *Pemerintah Provinsi Kepulauan Riau*, August 2025. <https://www.kepriprov.go.id/berita/perangkat-daerah/pemprov-kepri-manfaatkan-kecerdasan-buatan-ai-untuk-meningkatkan-strategi-informasi-komunikatif>.
- Pflücke, Felix. "Data Access and Automated Decision-Making in European Financial Law." *European Journal of Risk Regulation* 16, no. 1 (March 2025): 11–23. <https://doi.org/10.1017/err.2024.60>.
- Piccialli, Francesco, Diletta Chiaro, Sundas Sarwar, Donato Cerciello, Pian Qi, and Valeria Mele. "AgentAI: A Comprehensive Survey on Autonomous Agents in Distributed AI for Industry 4.0." *Expert Systems with Applications* 291 (October 2025): 128404. <https://doi.org/10.1016/j.eswa.2025.128404>.
- Rigotti, Carlotta, and Eduard Fosch-Villaronga. "Fairness, AI & Recruitment." *Computer Law & Security Review* 53 (July 2024): 105966. <https://doi.org/10.1016/j.clsr.2024.105966>.
- Rinta-Kahila, Tapani, Ida Someh, Nicole Gillespie, Marta Indulska, and Shirley Gregor. "Algorithmic Decision-Making and System Destructiveness: A

- Case of Automatic Debt Recovery.” *European Journal of Information Systems* 31, no. 3 (May 2022): 313–38. <https://doi.org/10.1080/0960085X.2021.1960905>.
- Roy, Siddhartha, and Runa Ganguli. “Comparative Analysis of Machine Learning Classification Algorithms for Predictive Models Using WEKA.” In *Emerging Technologies in Data Mining and Information Security*, In P. Dutt., 173–83. Springer Nature Singapore, 2023. https://doi.org/10.1007/978-981-19-4052-1_19.
- Ruschmeier, Hannah. “Generative AI and Data Protection.” *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e6. <https://doi.org/10.1017/cfl.2024.2>.
- Selvadurai, Niloufer. “Seeing the Full Picture: An Integrated Regulatory Model for AI.” *International Journal of Law and Information Technology* 33 (January 2025): eaaf005. <https://doi.org/10.1093/ijlit/eaaf005>.
- Sio, Filippo Santoni de, Giulio Mecacci, Simeon Calvert, Daniel Heikoop, Marjan Hagenzieker, and Bart van Arem. “Realising Meaningful Human Control Over Automated Driving Systems: A Multidisciplinary Approach.” *Minds and Machines* 33, no. 4 (July 2022): 587–611. <https://doi.org/10.1007/s11023-022-09608-8>.
- Souza, Siddharth Peter de. “The Spread of Legal Tech Solutionism and the Need for Legal Design.” *European Journal of Risk Regulation* 13, no. 3 (September 2022): 373–88. <https://doi.org/10.1017/err.2022.4>.
- Sparrow, Robert, and Gene Flenady. “The Testimony Gap: Machines and Reasons.” *Minds and Machines* 35, no. 1 (February 2025): 12. <https://doi.org/10.1007/s11023-025-09712-5>.
- Tartaro, Alessio. “Towards European Standards Supporting the AI Act: Alignment Challenges on the Path to Trustworthy AI.” In *Proceedings of the AISB Convention*, 98–106, 2023.
- Tieleman, Merlin. “Fairness in Tension: A Socio-Technical Analysis of an Algorithm Used to Grade Students.” *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e19. <https://doi.org/10.1017/cfl.2025.6>.
- Tsoukalas, Anastasios. “Artificial Intelligence and Legal Personality: Obligations, Rights, and Liability in the Age of Autonomous Systems.” *Επιθεώρηση Δικαίου Πληροφορικής* 6, no. 1 (2025). <https://doi.org/10.26262/infolawj.v1i1.10681>.
- Vallor, Shannon, and Tillmann Vierkant. “Find the Gap: AI, Responsible Agency and Vulnerability.” *Minds and Machines* 34, no. 3 (June 2024): 20.

<https://doi.org/10.1007/s11023-024-09674-0>.

- Vargas Penagos, Emmanuel. "Platforms on the Hook? EU and Human Rights Requirements for Human Involvement in Content Moderation." *Cambridge Forum on AI: Law and Governance* 1 (May 2025): e23. <https://doi.org/10.1017/cfl.2025.3>.
- Weckert, John. "Moor on Ethics for Emerging Technologies: Some Environmental Considerations." *Minds and Machines* 35, no. 2 (March 2025): 18. <https://doi.org/10.1007/s11023-025-09721-4>.
- Weerts, Sophie. "Generative AI in Public Administration in Light of the Regulatory Awakening in the US and EU." *Cambridge Forum on AI: Law and Governance* 1 (January 2025): e3. <https://doi.org/10.1017/cfl.2024.10>.
- Wendehorst, Christiane. "Liability for Artificial Intelligence: The Need to Address Both Safety Risks and Fundamental Rights Risks." In *The Cambridge Handbook of Responsible Artificial Intelligence*, In S. Voenn., 187–209. Cambridge University Press, 2022. <https://doi.org/10.1017/9781009207898.016>.
- Widyawati, Anis, Ade Adhari, Ridwan Arifin, and Helda Rahmasari. "Crafting an Ideal Penitentiary Law: A Global Comparative Framework for Indonesia's Correctional System." *Legality: Jurnal Ilmiah Hukum* 33, no. 2 (2025): 418. <https://doi.org/10.22219/ljih.v33i2.40358>.
- Widyawati, Anis, Didik Purnomo, Heru Setyanto, and Leony Sondang Suryani. "Strengthening the Correctional System through Electronic Supervision of Prisoners: A Comparative Legal Study for Reforming Indonesia's Penitentiary Law." *Jurnal Hukum Novelty* 16, no. 2 (2025): 264.
- Wrigley, Sam. "Bring in the Sceptics: Using Science and Technology Studies in Law." *European Journal of Risk Regulation* 16, no. 1 (March 2025): 51–62. <https://doi.org/10.1017/err.2024.103>.
- Yazdanpanah, Vahid, Enrico H Gerding, Sebastian Stein, Mehdi Dastani, Catholijn M Jonker, Timothy J Norman, and Sarvapali D Ramchurn. "Reasoning about Responsibility in Autonomous Systems: Challenges and Opportunities." *Ai & Society* 38, no. 4 (2023): 1453–64.
- Zeide, Elana. "Expanding the Paradigm: Generative Artificial Intelligence and U.S. Privacy Norms." *Cambridge Forum on AI: Law and Governance* 1 (March 2025): e16. <https://doi.org/10.1017/cfl.2024.15>.

Acknowledgment

The author(s) extend sincere gratitude to the Ministry of Higher Education, Science, and Technology of the Republic of Indonesia for the support and funding of the research and publication of this article.

Funding Information

None.

Conflicting Interest Statement

The authors state that there is no conflict of interest in the publication of this article.

Publishing Ethical and Originality Statement

All authors declared that this work is original and has never been published in any form and in any media, nor is it under consideration for publication in any journal, and all sources cited in this work refer to the basic standards of scientific citation.

Generative AI Statement Statement

N/A