
From Repression to Prevention: Reforming Indonesian AI Governance in the Shadow of the EU AI Act

Nur Hasanah¹, Muhammad Fauzan Ziefran²

¹Politeknik STIA LAN Jakarta

²Sekretariat Jenderal Kementerian Keuangan

e-mail: ¹nur.2514010008@stialan.ac.id, ¹nurhasanah28696@gmail.com, ²fauzanziefran@gmail.com

*Corresponding author: nur.2514010008@stialan.ac.id

ABSTRAK

Informasi Artikel:

Terima: 01-07-2026

Revisi: 25-07-2026

Disetujui: 02-08-2026

Publish online: 11-08-2026

Perkembangan pesat deepfake yang dihasilkan oleh kecerdasan buatan (AI) menghadirkan tantangan kritis terhadap integritas informasi dan stabilitas publik. Sementara Uni Eropa telah mengadopsi paradigma tata kelola yang proaktif dan berbasis risiko melalui EU AI Act—yang mewajibkan transparansi, watermarking digital, dan perlindungan hak-hak fundamental—lanskap regulasi Indonesia masih sangat kurang berkembang dan bersifat reaktif. Penelitian ini menggunakan metodologi deskriptif-komparatif, dengan EU AI Act sebagai tolok ukur, untuk mengkaji kebutuhan mendesak akan kerangka hukum khusus untuk media sintetis di Indonesia. Analisis menunjukkan bahwa ketergantungan Indonesia pada UU ITE dan UU PDP tidak memadai karena instrumen-instrumen tersebut gagal mengatasi tantangan ontologis unik dari konten yang dihasilkan oleh AI. Penelitian ini mengidentifikasi dua temuan kebijakan yang penting: perlunya mengadopsi “model regulasi hibrida” yang menyintesis standar teknis internasional dengan nilai-nilai konstitusional domestik dan penerapan “regulatory sandbox” untuk memfasilitasi pengujian kebijakan berbasis bukti. Oleh karena itu, penelitian ini mengusulkan transisi dari tindakan penghukuman *ex post* menuju arsitektur tata kelola *ex ante*. Temuan ini memberikan peta jalan strategis bagi pembuat kebijakan Indonesia untuk membangun kerangka tata kelola yang akuntabel, berpusat pada manusia, dan kuat secara teknis yang dapat memitigasi risiko eksistensial yang ditimbulkan oleh manipulasi algoritmik.

Kata Kunci: Kecerdasan Buatan, Deepfake, EU AI Act, Regulatory Sandbox, Reformasi Hukum

ABSTRACT

The rapid proliferation of artificial intelligence (AI)-generated *deepfakes* presents a critical challenge to information integrity and public stability. While the European Union has adopted a proactive, risk-based governance paradigm through the *EU AI Act*-mandating transparency, digital watermarking, and fundamental rights protections-Indonesia's regulatory landscape remains significantly underdeveloped and reactive. This study employs a descriptive-comparative methodology, benchmarked against the *EU AI Act*, to examine the urgent need for a dedicated legal framework for synthetic media in Indonesia. The

analysis reveals that Indonesia's reliance on the ITE and PDP laws is insufficient, as these instruments fail to address the unique ontological challenges of AI-generated content. The study identifies two critical policy findings: the necessity of adopting a "hybrid regulatory model" that synthesizes international technical standards with domestic constitutional values and the implementation of a "regulatory sandbox" to facilitate evidence-based policy testing. Consequently, this research argues for a transition from *ex post* punitive measures to an *ex ante* governance architecture. These findings provide a strategic roadmap for Indonesian policymakers to construct an accountable, human-centric, and technically robust governance framework that mitigates the existential risks posed by algorithmic manipulation.

Keywords: Artificial Intelligence, Deepfakes, EU AI Act, Regulatory Sandbox, Legal Reform

INTRODUCTION

The development of artificial intelligence (AI) technology has become one of the milestones of the 21st-century digital revolution. One manifestation of this technology is deepfakes, a highly realistic image, audio, and video manipulation technique using deep learning algorithms (Kristiyenda et al., 2025). The threat of large-scale manipulation through deepfakes is the most impactful. Within minutes, a manipulated video can go viral, creating a severe distortion of reality (Mellaz, 2026). Deepfakes are one of the negative impacts of generative intelligence that is still very difficult to control, because legal instruments are still not available. Jurisprudential efforts or judicial interpretation remains the only way to overcome this kind of legal deadlock (Afif, 2025).

The rapid proliferation of artificial intelligence (AI)-generated *deepfakes* presents a critical challenge to the integrity of the public sphere, creating an environment where the boundary between authentic discourse and synthetic manipulation has become increasingly porous. This phenomenon mirrors Plato's *Allegory of the Cave*, where citizens—like prisoners chained in the dark—are presented with shadows of reality, unable to discern the original object from the reflection. In the contemporary digital age, this gap between reality and perception has evolved into what Jean Baudrillard termed *simulacra*: a state where the representation of reality has replaced reality itself. The *deepfake* incident involving Finance Minister Sri Mulyani Indrawati in August 2025 serves as a poignant empirical manifestation of this digital *simulacrum*. The manipulation of her speech did not merely deceive the public; it eroded the fundamental communicative trust essential for a democratic society. As established in communication ethics, such acts violate the core tenets of honesty and integrity, undermining the social justice required for legitimate governance (Pearson, 2003).

However, the inadequacy of the current Indonesian legal response—which relies on reactive, fragmented applications of the ITE and PDP Laws—highlights an ontological limitation

in our regulatory framework. Viewed through the lens of John Austin's legal positivism, while Indonesia's digital criminal law satisfies the technical requirements of command and sanction, it fails to capture the unique essence of synthetic media, thereby leaving a void in definitions that only exacerbate the public's inability to distinguish truth from fabrication (Al Farizi, 2025). To navigate this challenge, this study adopts Lawrence Lessig's Modalities of Regulation as the primary analytical framework. Lessig posits that human behavior is regulated not only by law but also by architecture (technology), social norms, and the market. By applying this framework, this paper argues that effective AI governance in Indonesia cannot rely solely on the *ex post* enforcement of law but must integrate *ex ante* technical architectures – such as mandatory digital watermarking – and adaptive policy mechanisms like *regulatory sandboxes*. This research provides a strategic roadmap for shifting Indonesia's AI governance toward a preventive, human-centric model that harmonizes international technical standards with domestic constitutional values.

The spread of deepfakes is a real threat. Indonesia actually has Law Number 11 of 2008 concerning Electronic Information and Transactions (ITE), which has been amended to Law Number 19 of 2016. This law contains an article prohibiting the spread of false news that can cause harm to others. However, in practice, proving and taking action against perpetrators of deepfake distribution is often not easy. This is because deepfake technology has a very high level of similarity to the original video, so technical expertise is required to distinguish them (Kurniawan, 2025). The existence of legal loopholes and special technical capabilities of AI deepfakes makes law enforcement in Indonesia against perpetrators of deepfakes or fake digital content face challenges and obstacles.

In 2024, the European Union officially passed the EU AI Act, the world's first regulation specifically governing artificial intelligence (AI) with a risk-based approach. This regulation aims to ensure that AI is developed and used safely and ethically and respects human rights and democratic values.

As a country developing its digital ecosystem, including the use of AI, Indonesia needs to pay attention to the legal implications of this regulation (Fibrianti, 2025). The 2025 *deepfake* incident involving Finance Minister Sri Mulyani Indrawati—which falsely depicted her criticizing national teacher policies—serves as an empirical wake-up call, highlighting the systemic vulnerability of Indonesia's institutional credibility to synthetic media manipulation.

Nowadays, Indonesia has adopted a reactive rather than a proactive approach. If Indonesia simply waits for new problems to arise before enacting regulations, the consequences of this deepfake case could impact other public officials, not just the Ministry of Finance, but all government officials and all levels of society. Without a clear legal framework, Indonesian society will remain vulnerable to the dangers of AI misuse, particularly in digital public spaces like social media. The government, as the regulator, needs to act swiftly to establish regulations with strong legal force. Technology-related regulations should continuously adapt and keep pace with

technological developments to maintain social stability in the digital realm. Regulatory development can also encourage innovations aligned with Pancasila values. This commitment can lead Indonesia to become a pioneer in the development of sustainable, ethical, and human rights-based future technology (Badan Pembinaan Hukum Nasional., 2026).

Lessig's Modalities of Regulation (also known as the Pathetic Dot Theory) explains that human behavior, especially in cyberspace, is governed by four main interacting forces. Lessig warns that important decisions are being made all around us on crucial issues, based on society's shift to digital technology. Through this theory, we can identify the values affected and determine how best to address them. One of Lessig's key insights is that "we can break these considerations down into four broad modalities of regulation or constraints". Lessig's modalities of regulation is displayed in Table 1.

Table 1. Lessig's Modalities of Regulation

Modality	Definition	Application to AI Governance
Architecture	The technical constraints or "code" that govern digital behavior (Lessig, 1999, p. 24).	AI design, watermarking, and algorithmic filters.
Law	Formal state-imposed sanctions and legal mandates (Lessig, 1999, p. 30).	Legislative acts (e.g., ITE Law, EU AI Act).
Market	Economic forces that constrain or enable behavior (Lessig, 1999, p. 27).	Compliance costs and market incentives for ethical AI.
Social Norms	Community-enforced expectations and stigmatization (Lessig, 1999, p. 28).	Ethical standards for developers and public trust.

Source: Elaborated by the author (2026); definition framework adapted from Lessig (1999, pp. 24–30)

The theoretical framework applied here follows (Lessig, 1999) thesis that the regulation of behavior in cyberspace—and by extension, the governance of artificial intelligence (AI)—is a function of four integrated constraints. (Lessig, 1999) famously posits that "code is law," emphasizing that the technical architecture of digital platforms is not merely a tool but a fundamental regulator that can render prohibited actions technically impossible. This concept of architecture as a regulatory modality is particularly critical in the context of *deepfakes*, where technical interventions like digital watermarking act as *ex ante* constraints that preempt the need for *ex post* legal intervention.

Furthermore, (Lessig, 1999) warns of "latent ambiguities" in governance, noting that regulators often fail to recognize how these four modalities interact, sometimes allowing architecture to "override" law or social norms in ways that are opaque to the public. For instance, while Indonesian law (the ITE Law) attempts to regulate digital conduct, the underlying architecture of global social media platforms—often governed by private corporate code—frequently dictates the visibility and virality of manipulated content regardless of local legal sanctions (Lessig, 1999).

Finally, Lessig (Lessig, 1999) underscores the necessity of a "constitutional" approach to AI governance, where the state ensures that the modalities of market and architecture are not left entirely to the private sector. By applying these modalities, this study evaluates whether Indonesia's current dependence on *law* (ITE and PDP laws) is sufficient or if it must shift toward a model that integrates *architecture* (technical standards) and *social norms* (ethical communication) to successfully address the challenges posed by generative AI.

Recognizing the transitions taking place should hopefully provide us with some guidance on "how we can reclaim the values that matter in this [cyber]space and how we can insist on bringing into it those that are now absent." Lessig points out that these modalities are malleable. Contrary to much contemporary reporting on the "information revolution," very little of what we think of as today's digital environment is actually inherently or naturally determined.

Technology is a human creation and can be reconstructed to take on almost any attribute we value enough to build into it. Existing laws are also open to interpretation, producing "latent ambiguities" that emerge when trying to apply them to new situations. Given a given set of values, one can use Lessig's four modalities as a means to define our landscape and strategies for action. When we are able to recognize radical changes in one or more modalities, it is time for us to decide how to resist or adapt to these changes in order to continue to effectively promote our values. While "latent ambiguities" in existing laws, policies, and principles are often challenging, they can also provide us with open doors to new opportunities. When ambiguity arises in the traditional version of a modality, it gives us an opportunity to attempt to resolve that ambiguity in a way that promotes our values.

Lessig shows us that these modalities interact with each other in dramatic ways, often winning out relatively to one modality, which can help promote our values relative to another. A key point in Lessig's theory is that "code is law," meaning that digital architecture (technological design) is often more effective at regulating behavior than formal law itself. These four factors work together to constrain our actions. AI systems deemed high-risk must undergo a rigorous series of assessments, including impact evaluations and risk mitigation measures. Transparency is a priority—including the obligation to inform users when interacting with AI—to prevent manipulation.

National authorities are also empowered to oversee implementation and impose sanctions on violators. To encourage innovation, the EU also provides a "regulatory sandbox," a controlled environment for businesses to test AI technology without the risk of full sanctions. This sandbox concept helps innovators mitigate legal risks and gives regulators the opportunity to understand the technology during the trial process, enabling regulations to support innovation without compromising security (Perdana, 2024).

AI systems are considered high-risk under EU AI law if they are products or security components of a product regulated under specific EU law referred to by that law that lists specific uses generally considered high-risk, including the AI systems used. High-risk AI systems must comply with specific requirements (Kosinski, 2024). The rapid proliferation of artificial intelligence (AI) has introduced *deepfake* technology as a significant global challenge, one that fundamentally blurs the boundaries of objective reality and poses substantial threats to information integrity and human rights. Amidst these emerging risks, the urgency for a dedicated regulatory framework in Indonesia has become increasingly critical, transcending the current reliance on reactive interpretations of the Electronic Information and Transactions (ITE) Law.

As a comparative benchmark, the EU AI *Act* has pioneered a risk-based governance paradigm, establishing rigorous standards for AI systems classified as high-risk. Within this framework, AI systems—whether integrated as safety components of regulated products or identified for specific high-risk applications—are mandated to comply with comprehensive obligations designed to mitigate foreseeable risks throughout their operational lifecycle. Adherence to the stringent standards of the EU AI *Act* necessitates the implementation of continuous risk management systems, requiring providers to proactively monitor and minimize the adverse impacts of their technologies. Beyond risk mitigation, this framework emphasizes robust data governance, ensuring that training, validation, and testing datasets meet high-quality benchmarks, including structured oversight to prevent and reduce algorithmic bias. Furthermore, operational transparency is bolstered through the mandatory maintenance of comprehensive technical documentation, which must detail system design specifications, functional capabilities, inherent limitations, and evidence of regulatory compliance. The integration of such transparency obligations, including the application of machine-readable labeling for generative content, serves as a vital preventive measure to ensure technical accountability within an ecosystem increasingly vulnerable to synthetic manipulation.

In the domestic context, the limitations of current legal instruments to effectively address incidents such as the manipulated media targeting public officials highlight a persistent *legal deadlock* that necessitates the development of novel regulatory constructs. Drawing upon Lawrence Lessig's modalities of regulation, the architecture of AI governance in Indonesia must integrate not only formal legal sanctions but also technical architectures—or "code"—that inherently prioritize transparency and accountability. Consequently, this research aims to analyze the necessity of transitioning Indonesia's regulatory model toward a hybrid approach, one that synthesizes international technical standards with domestic constitutional values. Through mechanisms such as *regulatory sandboxes*, it is posited that Indonesia can foster an innovative digital ecosystem that simultaneously safeguards national digital sovereignty and protects human rights in an era of algorithmic totalization.

Beyond general compliance obligations, the emerging regulatory framework mandates specific transparency requirements for certain artificial intelligence (AI) systems. First, AI systems designed for direct interaction with human subjects must adhere to a principle of proactive transparency. This necessitates that the system explicitly notify the user regarding the AI-driven nature of the interaction, unless such identity is self-evident from the immediate context. By way of illustration, interface implementations such as chatbots are required to provide an explicit disclosure identifying the non-human entity to the user.

Second, AI systems with generative capabilities—specifically those producing text, imagery, or other multimedia content—are obligated to employ machine-readable formatting to identify the output as AI-generated or synthetically manipulated. This labeling mechanism serves to ensure transparent identification of generative or technically altered content. This mandate encompasses the detection and labeling of *deepfake* media, defined as content manipulated through AI technologies to portray a subject performing actions or uttering statements that did not factually occur.

The Risk-Based Paradigm in Artificial Intelligence Governance with digital transformation propelled by artificial intelligence (AI) has instigated a fundamental shift in public policy, where technological innovation frequently outpaces existing legal frameworks. Amidst the urgent necessity to manage potential systemic risks, the EU *AI Act* has introduced a pioneering risk-based governance paradigm, establishing a global benchmark for AI oversight. Rather than adopting a uniform regulatory approach, this framework segments AI systems based on the severity of risks they pose to fundamental rights and public safety. AI systems classified as high-risk—including security components of products regulated under European Union law—are subject to rigorous compliance obligations. These mandates encompass the maintenance of comprehensive technical documentation, stringent data governance to mitigate algorithmic bias, and the implementation of continuous risk management systems throughout the technology's operational lifecycle (The Conversation, 2024).

Integrating this risk-based approach requires technology providers to transcend mere administrative compliance by embedding the principles of transparency and accountability directly into the system's architecture. Requirements for AI systems that interact with human subjects to explicitly disclose their non-human identity, alongside the use of machine-readable formatting to label generative or manipulative content such as *deepfakes*, are critical measures for preserving the integrity of public information. By adopting such a framework, the state functions not merely as a repressive regulator but as a facilitator of responsible innovation through preventive legal architecture. Consequently, for nations like Indonesia, the adoption of technical standards aligned with the EU *AI Act* is not merely a policy preference but an urgent imperative to bridge the regulatory gap, thereby safeguarding digital sovereignty and human dignity amidst the escalating threat of algorithmic manipulation.

The rapid evolution of artificial intelligence (AI) has consistently outpaced traditional legislative cycles, creating a "pacing problem" that necessitates innovative governance mechanisms. To address this, the EU *AI Act* has institutionalized the "regulatory sandbox"—a controlled environment that enables businesses to develop and test high-risk AI technologies under regulatory supervision without the immediate threat of full sanctions. This mechanism serves a dual purpose: it allows innovators to mitigate legal uncertainties while providing regulators with critical, real-time insights into emerging technologies. By fostering a collaborative space for experimentation, the sandbox ensures that regulatory frameworks remain adaptive and evidence-based, effectively balancing the drive for technological advancement with the imperative to protect fundamental rights.

Crucially, this approach does not imply unrestricted freedom; it operates alongside strict prohibitions on harmful practices, such as social scoring and exploitative behavioral manipulation, thereby ensuring that innovation remains anchored in public trust and safety (Perdana, 2024).

For Indonesia, the implementation of a regulatory sandbox represents a vital strategic imperative to bridge the widening gap between rapid digital innovation and the often-protracted legislative process (Elviandri., 2026). Given the existing legal vacuum surrounding generative AI technologies—exemplified by the recent rise of sophisticated *deepfakes* that threaten information integrity and public stability—a sandbox mechanism offers a pragmatic pathway for policy testing. By facilitating small-scale, controlled trials before the formal enactment of comprehensive AI legislation, Indonesia can synthesize global technical standards with its unique constitutional values. This experimental approach to policymaking is essential for navigating the complex digital landscape, as it allows for the development of a "hybrid" regulatory model that safeguards digital sovereignty and social harmony while nurturing an ethical, robust, and competitive AI ecosystem.

The implementation of a regulatory sandbox represents a vital strategic imperative for Indonesia in addressing the "pacing problem"—the widening gap between the rapid acceleration of artificial intelligence (AI) innovation and the sluggish adaptation of national legal frameworks. This phenomenon frequently precipitates the "Collingridge Dilemma," wherein the potential impacts of a new technology remain elusive during the initial stages of development yet become nearly impossible to mitigate once the technology is deeply entrenched within the social and economic fabric (Elviandri., 2026). Through a sandbox mechanism, the government can facilitate early risk detection and apply regulatory flexibility, thereby balancing the exigencies of innovation without triggering a "chilling effect" or, conversely, fostering a normative void that jeopardizes public interest. Consequently, the sandbox functions as an empirical testing ground that enables evidence-based policy formulation, replacing approaches that rely solely on political intuition.

The transition of AI governance from non-binding ethical guidelines—as currently reflected in the Ministry of Communication and Information Technology Circular Letter No. 9 of 2023—toward legally binding regulation requires a calibrated approach (Djafar, 2024). In this context, the sandbox serves as a "policy bridge," allowing for the testing of values such as transparency and accountability in real-world practice before they are formalized into presidential regulations or statutory laws. This approach aligns with the concept of "experimental legislation," which integrates a "sunset clause" into the Indonesian legal system (Elviandri., 2026). Under this mechanism, regulations are granted a limited validity period—for instance, two years—with their permanence contingent upon proven effectiveness during the trial phase; should a regulation prove ineffective, it expires automatically, circumventing the complex and protracted revocation processes typical of the House of Representatives (DPR). Furthermore, innovators and startups can be granted temporary regulatory relief under stringent supervision, facilitating legal experimentation without compromising public safety. To ensure efficacy, the institutional structure of an AI regulatory sandbox in Indonesia should be designed through systematic stages: an exploration phase to identify critical legal issues, a pilot testing phase within a limited environment or specific sector, and an evaluation phase to determine the feasibility of national-scale implementation (Baskoro, 2024). Indonesia's empirical experience with regulatory sandboxes in the financial services sector, overseen by the Financial Services Authority (OJK) and Bank Indonesia, as well as in disruptive health technology, provides a valuable precedent for AI governance. Consistent with global standards—such as the EU *AI Act*, which mandates that EU member states establish AI sandboxes by 2026 to support SMEs and startups in achieving compliance with high-risk standards—Indonesia has a strategic opportunity to adopt these international best practices (Ramli, 2025). The integration of this model is expected to mitigate legal risks while simultaneously fostering a resilient, ethical, and competitive innovation ecosystem that remains firmly aligned with the nation's constitutional values.

RESEARCH METHODE

A qualitative comparative legal study between the EU AI Act and Indonesia's regulatory framework regarding AI deepfakes reveals significant differences in regulatory maturity, transparency obligations, and enforcement mechanisms. This gap highlights the urgency for Indonesia to develop ideal AI regulations (*ius constituendum*) to provide legal certainty in addressing the rapidly evolving threat of deepfakes (Utami, 2025). Table 2 shows the comparison between the EU AI Act and Indonesia's regulatory framework.

Table 2. the Comparative Legal Study Between the EU AI Act and Indonesia's Regulatory Framework

Comparative Aspect	EU AI Act	Indonesia
Philosophical Foundations and Regulatory Approach	Using a risk-based approach that explicitly regulates AI systems based on their potential harm to fundamental rights (Barrister Dr. Anwar Baig, 2025), deepfakes are placed in a category requiring special transparency obligations due to the risk of fraud and identity manipulation (Nicolaj Feltes, 2025).	Currently, there is a legal vacuum specific to AI (Ni Putu Gita Sri Utami et al., 2025). Existing regulations are "soft" regulations, such as Circular Letter (SE) of the Minister of Communication and Information Technology Number 9 of 2023 concerning the Ethics of Artificial Intelligence, which lacks imperative force or binding sanctions (Wahyu Sudoyo, 2024).
Definition and Identification of Deepfakes	Provides a clear legal definition in Article 3(60) as audio, image, or video content generated or manipulated by AI that resembles a real person, object, or event so that it appears authentic to others (Nicolaj Feltes, 2025).	There is no specific legal definition for deepfakes in the law. Deepfakes are generally viewed as a form of technological misuse that can cause material and immaterial harm, but their legal limitations still refer to the general definition of an electronic system in the ITE Law (Ni Putu Gita Sri Utami et al., 2025).
Synergy and Regulation of Digital Platforms	There is an interaction between the AI Act and the Digital Services Act (DSA). While labeling violations do not automatically render content "illegal," large platforms (VLOPs) are required to mitigate systemic risks, including prominent deepfake labeling (Nicolaj Feltes, 2025).	The 2020-2045 National Strategy (Stranas KA) and the minister of communication and information's circular letter emphasize quadruple-helix collaboration (government, industry, academia, and community) but have not yet integrated strict AI content moderation obligations for platform providers (BPPT, 2020).
Obligations of Transparency and Disclosure	EU AI Act (Article 50): a. Disclosure Obligation: Deployers (users) of AI systems that generate deepfakes are required to disclose that the content is artificially manipulated (Nicolaj Feltes, 2025) b. Format: Disclosure must be made in a "clear and distinguishable" manner (Nicolaj Feltes, 2025). c. Exception: There is leeway for content that is artistic, creative, satirical, or fictional, but there must still be minimal disclosure so as not to interfere with the enjoyment of the work (Nicolaj Feltes, 2025).	a. The responsibility for disclosure is delegated to Electronic System Providers (PSE) based on ethics, but there are no legally binding technical standards regarding the labeling of AI content (Ni Putu Gita Sri Utami et al., 2025). b. The ITE Law and the PDP Law are used indirectly to process the impacts of deepfakes (such as data theft or defamation) but do not regulate the obligation of labeling in advance (ex ante) like the EU AI Act (Ni Putu Gita Sri Utami et al., 2025).

Enforcement Mechanisms and Sanctions	Establishing severe administrative sanctions. Violations of transparency obligations (including deepfakes) could result in fines of up to €15 million or 3% of a company's total annual global turnover (Nicolaj Feltes, 2025).	There are no specific sanctions for AI transparency violations. Law enforcement relies on articles in the ITE Law or the Criminal Code if the deepfake has resulted in a crime but fails to provide legal certainty as a preventive measure (Ni Putu Gita Sri Utami et al., 2025).
---	---	--

Source: Author's own elaboration (2026), adapted from Baig (2025),

Feltes (2025), Sudoyo (2024), and Utami et al. (2025)

From a *de jure* perspective, Indonesia currently confronts a significant "pacing problem," characterized by the disparity between the exponential trajectory of artificial intelligence (AI) innovation and the inherently linear pace of the national legislative process. Legally, the Indonesian framework lacks a specialized regulation that defines *deepfakes* as a distinct *sui generis* criminal offense, forcing law enforcement to rely on extensive, and often contentious, interpretations of existing legal instruments (Al Farizi, 2025).

The primary reliance is placed upon the Law on Information and Electronic Transactions (ITE Law), which serves as the bedrock for prosecuting cases involving the manipulation of electronic data – pursuant to Article 35 – and the dissemination of fabricated information deemed capable of triggering public unrest under Article 28, paragraph 3 (Nurwahid, 2026; Khaeron, 2025). Complementing this, the Personal Data Protection (PDP) Law provides a layer of protection by categorizing facial and vocal features as sensitive biometric data under Article 4, paragraph 2, whereby unauthorized exploitation may entail severe criminal sanctions (Khaeron, 2025).

Despite the application of these statutes, the reliance on existing legal instruments reveals a critical ontological limitation when viewed through the lens of legal theory. Referencing John Austin's positivist framework, while the validity of Indonesia's digital criminal law satisfies the essential elements of command, duty, and sanction, it struggles to bridge the gap between abstract legal norms and the concrete reality of AI-generated synthetic imagery (Al Farizi, 2025). The current legislative reliance on pre-existing instruments fails to account for the unique characteristics of synthetic media, which fundamentally challenges the traditional definition of "bodies" or "identity" in criminal law. Consequently, this reliance on interpretative flexibility highlights the urgency for a specialized regulatory construct that acknowledges the distinct nature of generative AI, ensuring that law enforcement remains grounded in clear, normative definitions rather than relying on the potential overreach of broader, general provisions.

The vulnerability of the public sphere to artificial intelligence (AI) manipulation was empirically underscored in August 2025, when a *deepfake* video surfaced featuring Finance Minister Sri Mulyani Indrawati. By utilizing synthetic media to manipulate a genuine speech delivered at the Bandung Institute of Technology (ITB), perpetrators fabricated a false narrative

regarding teacher policies, thereby violating fundamental principles of ethical communication and infringing upon human dignity (Sulistya, 2026). This case provides critical insights into the limitations of current legal redress; it demonstrates that the immaterial losses – manifesting as profound psychological trauma and reputational damage – often far exceed material damages. Consequently, such incidents necessitate the robust application of rights-redress mechanisms, most notably the "Right to be Forgotten" as stipulated in Article 26 of the ITE Law, though its efficacy remains a subject of ongoing legal debate (Al Farizi, 2025; Kamasi, 2026).

The Indonesian response to such incidents currently leans toward a repressive, *ex post facto* punitive approach, which stands in stark contrast to the preventive, risk-based paradigm championed by the European Union. Under the EU AI Act (2024), AI governance is predicated on the principle that the stringency of regulation must be commensurate with the level of risk posed by the system (Fibrianti, 2025). Within this framework, while *deepfakes* are generally classified as limited-risk, they escalate to high-risk categories when utilized to manipulate political discourse or electoral processes. To mitigate these risks, the EU AI Act mandates that providers ensure AI outputs are human-identifiable, explicitly requiring the implementation of digital watermarking mechanisms to ensure provenance and authenticity (Legal 500., 2024)

For Indonesian policymakers, adopting the European principle of "the higher the risk, the stricter the rules" offers a strategic pathway toward developing more accountable technical standards for synthetic media (Fibrianti, 2025; Tarigan, 2025). By shifting from a purely reactive stance to a proactive, risk-informed regulatory architecture, Indonesia could better safeguard its information ecosystem. Such an integrated approach not only addresses the specific technical threats posed by generative AI but also aligns national digital governance with emerging international standards, ensuring that human dignity and the integrity of public discourse are protected against the sophisticated threats of algorithmic manipulation.

The lack of technical regulations in Indonesia has left the handling of deepfake crimes trapped in a gray area that relies on general articles. To achieve optimal legal certainty, modernization of investigative instruments and the establishment of regulations at the level of government concerning the standardization of AI-generated media are needed to maintain the integrity of public information in the era of artificial intelligence (Al Farizi, 2025). A comparison of policy instruments between Indonesia and the European Union in addressing the AI deepfake phenomenon reflects fundamental differences in legal philosophies and digital risk mitigation strategies. The following is a comparative analysis of AI policy instruments from both jurisdictions is displayed in Table 3.

Table 3. Policy Instruments Between Indonesia and the European Union

Policy Instruments	Indonesian Model (Laissez-faire/Reactive)	European Union Model (Precautionary/Proactive)
Primary Logic	Post-Violation (Ex Post) enforcement focuses on criminal offenses after content is distributed.	System Design-Based Prevention (Ex Ante) focuses on technical compliance (ante) before the system is launched.
Key Actors	Law enforcement agencies , especially the Republic of Indonesia National Police (Polri) and the Ministry of Communication and Digital (Komdigi) as the authorities for taking action and terminating access.	The AI ecosystem consists of AI developers, platforms, and the AI oversight body (AI office) through transparency obligations.
Strategic Objectives	Public Order aims to maintain stability through the ITE Law and prevent disinformation.	Fundamental rights aim to protect human rights, digital security, and democratic values.

Source: Author's own elaboration (2026), adapted from the EU AI Act (2024) and Indonesia's

Law No. 1/2024 on the Second Amendment to the ITE Law

RESULTS AND DISCUSSION

An analysis of deepfake cases involving public figures (state officials) such as Sri Mulyani Indrawati reveals a significant regulatory lag in the Indonesian legal system (Habaib, 2026). The current regulatory failure stems from the absence of jurisdictions that specifically classify deepfakes as a separate crime, creating a legal gap that contributes to legal uncertainty. Current law enforcement remains reactive and relies entirely on the authorities' normative interpretation of existing legal instruments, which are not designed to address the complexities of AI manipulation (Wicaksono, 2026).

Regarding the effectiveness of the ITE Law, this instrument is considered inadequate to comprehensively ensnare AI creators or platform providers. Although Articles 35 and 28 of the ITE Law can be used to ensnare the practice of information manipulation and the spread of fake news, the focus of these regulations is on the downstream (distribution and results of content), not the upstream (the technical process of manipulating AI technology). There is a fundamental problem in criminal liability; Indonesian positive law has not been able to firmly determine who should be held responsible in the AI ecosystem—whether the algorithm creator, the system operator, or the technology provider corporation.

Without regulations that explicitly regulate legal subjects in the AI ecosystem, platform providers and technology creators tend to be free from legal entanglements, while law enforcement can only reach the end-users who distribute the content (Habaib, 2026). Therefore, a new legal construction is needed that specifically regulates the misuse of AI technology to provide optimal legal protection for victims (Arvitto, 2025).

Drawing from the experience of Indonesia, we must approach the governance of artificial intelligence through a multi-dimensional strategy that balances rigorous oversight with the necessity for industrial innovation. Central to this approach is the adoption of a phased implementation framework, which allows industries sufficient time to adapt to new requirements—such as those concerning data rights, privacy, and algorithmic transparency—without precipitating economic shocks.

This must be complemented by a flexible regulatory structure capable of evolving alongside rapid technological advancements, ensuring that oversight mechanisms remain relevant and adaptive without requiring total legislative overhauls. To ensure the integrity of this process, Indonesia should institutionalize cross-sectoral collaboration, actively engaging government agencies, industry stakeholders, academia, and civil society to establish a culture of transparency and accountability, particularly in sectors where the social impact of AI is most profound.

Rather than attempting to regulate the entire AI spectrum simultaneously, Indonesian policymakers would benefit from a focused approach that prioritizes areas most critical to the national context. This strategy should be bolstered by regional cooperation to harmonize standards, thereby enhancing collective competitiveness and minimizing the fragmentation of the digital market, while simultaneously maintaining alignment with international regulatory benchmarks. Furthermore, the regulatory architecture must incorporate incentive-based mechanisms to encourage ethical AI adoption and foster a collaborative business environment. In this regard, it is imperative to implement safeguards against the emergence of "super policy entrepreneurs"—corporations that exert disproportionate influence over the legislative process—to ensure that AI development serves the broader public interest rather than the interests of specific groups or private entities.

Ultimately, the sustainability of AI governance rests upon the continuous development of regulatory capacity and the establishment of an experimental space through "regulatory sandboxes," which allow for the controlled testing of new technologies before national deployment. When designing these regulations, particular attention must be paid to the limitations and potentials of local small and medium enterprises (SMEs) and startups, ensuring that they are supported by technical and financial assistance to strengthen the domestic innovation ecosystem.

By learning from these lessons, Indonesia can create a regulatory environment that encourages AI innovation while protecting the public interest. A balance between regulation and flexibility to encourage innovation is key (Perdana, 2024).

The European Union, through the EU AI Act (2024), has established the first legally binding global standard with a risk-based approach (Kurniati, 2025). For limited risk categories, such as generative AI that generates text, images, or videos (deepfakes), this regulation requires explicit transparency (HCLTech, 2026). A key point of the EU standard is the obligation for AI system providers to ensure that synthetically generated content is technically recognizable. This includes the implementation of machine-readable watermarking and hard-to-remove permanent marks to ensure content integrity and source attribution (Nemecek, 2025). Failure to comply with

this transparency obligation can result in severe administrative sanctions, up to €7.5 million or 1.5% of the company's total global annual turnover (HCLTech, 2026). Table 4 shows the regulatory aspects between Indonesia and the European Union.

Table 4. Regulatory Aspects Between Indonesia and the European Union

Regulatory Aspects	Indonesia (Current Status)	European Union (EU AI Act)
Legal Definition	Broadly categorized as an "Electronic System" under the ITE Law.	Technically precise; defined by autonomous capabilities and environmental impact.
Transparency (Labeling)	Voluntary/Ethical; reliant on non-binding circulars (SE Menkominfo 9/2023).	Mandatory; absolute requirement for digital watermarking and synthetic content labeling.
Enforcement Logic	<i>Ex post</i> (Reactive); focuses on punishing illegal content dissemination.	<i>Ex ante</i> (Proactive); focuses on technical compliance and risk management.
Primary Sanction Basis	Illegal content/individual defamation.	Systemic failures, algorithmic bias, and audit non-compliance.
Regulatory Actors	Ministry of Communication and Digital (general oversight).	Specialized AI Office and independent national supervisory bodies.

Source: Author's own elaboration (2026), adapted from the EU AI Act (2024) and SE Menkominfo No. 9/2023

The absence of binding, comprehensive artificial intelligence (AI) legislation in Indonesia can be attributed to a confluence of institutional, political, and economic factors that necessitate a cautious approach to governance. Primarily, the nation confronts significant institutional capacity limitations, particularly regarding the availability of human resources equipped to perform independent algorithmic audits. While the National AI Strategy acknowledges this deficit, there remains a persistent shortage of competent technical experts, researchers, and specialized laboratories required to monitor the highly complex and evolving development of AI systems. This lack of technical infrastructure complicates the state's ability to transition from passive observation to active, evidence-based regulation.

Politically, the legislative agenda of the Indonesian government currently prioritizes the reinforcement of "soft law" frameworks over the immediate codification of a comprehensive AI act. While preliminary efforts toward specific AI governance are underway, the primary focus of the national legislature remains directed toward the implementation of the Personal Data Protection (PDP) Law and the deliberation of the Cybersecurity Bill. This prioritization suggests a strategic focus on foundational digital security before addressing the nuanced governance requirements of AI. Consequently, the development of binding AI norms is currently subordinated to broader cybersecurity and data privacy objectives, reflecting an incremental approach to digital policymaking.

Furthermore, there are profound concerns regarding the potential economic impact of imposing stringent, EU-style regulatory frameworks on the domestic market.

Policymakers are wary that rigid compliance requirements could impose unsustainable costs on local AI startups, thereby creating significant barriers to entry that might stifle the growth of the burgeoning domestic innovation ecosystem. Balancing the need for consumer protection with the necessity of maintaining a competitive environment for local innovators remains a primary challenge. As a result, the current legislative environment reflects a deliberate attempt to avoid stifling domestic technological progress while concurrently working to build the institutional and normative capacity required for more comprehensive, binding future regulation. To close regulatory gaps without stifling innovation, the strategic steps required for Indonesia are outlined in the following Table 5.

Table 5. Tactical Steps Towards Effective AI Governance in Indonesia

Mandatory Labeling through Platform Policies	The government should require global digital platforms (such as Meta, X, and TikTok) to provide free AI content detection features to the public. This would shift the detection burden from users to platforms with high computing power (Komdigi, 2025).
Public Literacy Campaign	Non-regulatory policies should be prioritized through digital education to build public resilience against AI disinformation and deepfakes. This literacy is crucial given the wide disparity in technology access across various regions (Safer Internet Lab, 2025).
Collaborative Governance	The implementation of the Quadruple-Helix model involves collaboration between government, academia, industry, and the community through platforms like KORI-KA. This approach allows for the creation of demand-driven standards that adapt to technological change without having to wait for a lengthy legislative process (BPPT, 2020).

Source: Author's own elaboration (2026), synthesized from BPPT (2020), Komdigi (2025), and Safer Internet Lab (2025)

The analytical framework used to interpret the findings in this study is based on Lawrence Lessig's Modalities of Regulation. This approach allows for a comprehensive evaluation of Indonesia's AI governance by assessing how law, technology (architecture), social norms, and market forces interact to either mitigate or exacerbate the risks posed by deepfake manipulation.

Within this framework, the August 2025 ministerial *deepfake* case functions as the primary unit of analysis. In this instance, perpetrators utilized sophisticated generative AI to manipulate a genuine speech delivered by Finance Minister Sri Mulyani Indrawati at the Bandung Institute of Technology (ITB). By re-synthesizing her voice and likeness, the actors fabricated a narrative – purporting that the minister labeled "teachers as a burden on the state" – which was subsequently disseminated across social media. Beyond the legal dimension of individual defamation, this case

demonstrates how synthetic media acts as a *simulacrum* (Baudrillard, 2025), replacing the minister's authentic policy stance with a synthetic fabrication that directly jeopardizes public trust in fiscal governance.

The incident revealed that existing legal instruments, specifically the ITE and PDP Laws, were insufficient in providing rapid, *ex ante* mitigation, as they are inherently designed for *ex post* reactive litigation.

CONCLUSION

The escalating threat of artificial intelligence, personified through synthetic media manipulation, confirms that without specific legislation, Indonesia's digital sovereignty is at existential risk. This study highlights the urgency of transitioning from partial sectoral regulations to a "hybrid model" that integrates global security standards such as risk classification in the EU AI Act and the data protection principles of the General Data Protection Regulation (GDPR) with the foundation of Indonesian "cultural constitutionalism." This model offers a collective accountability approach rooted in social ethics and customary law, where responsibility is not limited to technical-legal aspects but also to maintaining moral harmony and human intensity (Pancasila by design). The significance of this model is validated by the chronology of the case of state official and minister of finance Sri Mulyani Indrawati in 2025, which became a crucial turning point for the digital landscape of the Global South.

The regulatory void in Indonesia was starkly illustrated by the February 4, 2025, arrest of a perpetrator by the National Police's Cyber Crime Directorate. The individual was charged with disseminating *deepfake* content that manipulated the visuals and voices of President Prabowo Subianto and Finance Minister Sri Mulyani Indrawati to propagate the false narrative that 'Teachers Are a Burden to the State.' This incident serves as a critical unit of analysis, demonstrating how synthetic media has evolved into an ideological weapon of anti-elite populism, capable of delegitimizing state authority within hours. The case underscores a systemic failure: the current legal framework lacks mandatory algorithmic audit mechanisms and biometric verification requirements, allowing the content to achieve widespread viral exposure on platforms like Instagram before law enforcement could intervene. Ethically, Indonesia must adopt international human rights standards, particularly the data subject's right to transparency and explanation (right to explanation) as stipulated in the GDPR, to mitigate algorithmic black-box behavior. Integrating Pancasila values ensures that technology remains human-centric, protects privacy as an inviolable human right, and prevents "data colonization" by foreign infrastructure. In conclusion, harmonization between global security standards and local Indonesian legal wisdom is an absolute prerequisite for building a digital ecosystem that is not only innovative and competitive but also just, sovereign, and socially legitimate in the era of algorithmic totalitarianism. As evidenced by the 2025 ministerial *deepfake* case, the current reliance on punitive, reactive legal measures is inadequate to safeguard the integrity of the digital public sphere. This incident underscores that the threat posed by algorithmic manipulation is not merely a matter of individual harm but a systemic challenge that requires shifting from a model

of *ex post* punishment to an *ex ante* architecture of transparency and technological verification, as advocated by Lessig's regulatory modalities.

REFERENCES

- Afif, M. (2025). Mengatasi deadlock hukum atas teknologi kecerdasan buatan generatif melalui interpretasi yudisial. *Jurnal Hukum dan Kebijakan Digital*, 6(1), 45–58.
- Ahmad M. Ramli. (2025). *Adaptasi Standar Regulasi AI Global ke dalam Hukum Nasional*. Bandung: Refika Aditama.
- Aji Baskoro. (2024). *Kerangka Kelembagaan Regulatory Sandbox untuk Inovasi Digital di Indonesia*. Jakarta: References Lembaga Penerbitan Hukum.
- Al Farizi, F. S. (2025). Penegakan hukum terhadap manipulasi media digital deepfake berdasarkan perspektif positivisme hukum John Austin. *Jurnal Hukum Siber dan Hukum Pidana Digital*, 3(1), 78–92.
- AnyFlip. (n.d.). *Panduan Sandbox Regulasi Kecerdasan Artifisial*. <https://anyflip.com>
- Arvitto, R. S. (2025). Implikasi hukum deepfake: telaah terhadap UU ITE dan UU PDP. *Jurnal Ilmiah Hukum dan Hak Asasi Manusia*, 4(2), 73-82. <https://doi.org/10.35912/jihham.v4i2.3937>
- Athaillah, C. D., Tupen, D. L., Sari, A. J., Adzzagra, F. C., & Indrawan, J. (2026). Tantangan Cyber Defense dalam Menghadapi Serangan Siber terhadap Infrastruktur Pemerintah di Indonesia. *Governance: Jurnal Ilmiah Kajian Politik Lokal dan Pembangunan*, 8(1). <https://doi.org/10.33753/governance.v8i1.412>
- Badan Pembinaan Hukum Nasional. (2025). *Kesenjangan Regulasi Artificial Intelligence (AI) di Indonesia: Teknologi Melesat, Hukum Jalan di Tempat*. <https://www.bphn.go.id>
- Badan Pengkajian dan Penerapan Teknologi. (2020). *Strategi Nasional Kecerdasan Artifisial Indonesia 2020–2045*. Jakarta: BPPT.
- Baudrillard, J. (2025). Hyperreality and synthetic media: Analyzing the Sri Mulyani deepfake as a baudrillardian simulacrum in digital governance. *Jurnal Teori Sosial dan Komunikasi Kontemporer*, 9(2), 210–225.
- Center for Governance and Digital Society. (2024). *Digital trust landscape in Indonesia 2023 (DTLI 2023)*. <https://cgd.ibc-institute.id/wp-content/uploads/2024/10/DTLI-2023.pdf>
- DNV. (n.d.). *Introduction to the EU's AI Act: What you should know*. <https://www.dnv.com>
- ELSAM. (n.d.). *National AI Symposium*. <https://elsam.or.id>
- Elviandri, E., Putri, O., Azqmi, U., et al. (2026). Legislasi di era transformasi digital. *RIGGS*, 4(4), 8466-8473. <https://journal.ilmudata.co.id/index.php/RIGGS/article/view/4870>
- European Commission. (2025). *Guidelines and code of practice on transparent AI systems*. <https://ec.europa.eu>

- European Parliamentary Research Service. (2022). *Artificial intelligence act and regulatory sandboxes*. [https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI\(2022\)733544](https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI(2022)733544)
- European Union. (n.d.). Article 50: Transparency obligations for providers and deployers of certain AI systems. In *the Artificial Intelligence Act*. <https://artificialintelligenceact.eu/>
- Farihin Salman Al Farizi. (2025). *Tinjauan Yuridis terhadap Penyalahgunaan Teknologi Deepfake AI (Artificial Intelligence) sebagai Tindak Pidana*. <https://digilib.uin-suka.ac.id/id/eprint/75383/>
- Feltes, N. (2025). Article 50 AI Act: Do the transparency provisions improve upon the Commission's draft? *JIPITEC*, 16(2), 222-237. <https://www.jipitec.eu/jipitec/article/view/438>
- Fibrianti, S. (2025). *Dampak hukum EU AI Act bagi Indonesia*. JDIH Kementerian Komunikasi dan Digital. https://jdih.komdigi.go.id/storage/file-artikel-hukum/1748403902-FIX_Dampak_Hukum_EU_AI_Act_bagi_Indonesia.pdf
- Habaib, N. J. (2026). An analysis of deepfake cases involving public figures: Regulatory lag and legal certainty in the Indonesian legal system. *Jurnal Penegakan Hukum dan Keadilan Digital*, 8(1), 74-89.
- HCLTech. (2026). *EU AI Act guide v2*. <https://www.hcltech.com>
- IBM. (n.d.). *Apa itu EU AI Act?* <https://www.ibm.com>
- InfoPublik. (2024, January 21). *Wamenkominfo: Surat Edaran Tata Kelola AI becomes soft regulation*. <https://infopublik.id>
- Ismail, M. N. (2025). Pengaruh Teknologi AI terhadap Evolusi Modus Kejahatan Siber di Indonesia Tahun 2024-2025 dan Implikasinya terhadap Penegakan Hukum. *Jurnal Intelek dan Cendikiawan Nusantara*. <https://jicnusantara.com/index.php/jicn/article/view/6108>
- Jurnal RechtsVinding. (n.d.). *Transformasi Kualitas Legislasi: Regulatory Sandbox sebagai Sarana Partisipasi Publik*. <https://rechtsvinding.bphn.go.id>
- Kamasi, A. I., Timomor, A., & Lasut, M. M. W. (2026). Perlindungan Korban Rekayasa Digital Berbasis Artificial Intelligence terhadap Data Pribadi dalam Perspektif Hukum Pidana Indonesia. *Jurnal Inovasi*, 5(1). <https://ejournal.unima.ac.id/index.php/governance/article/view/14119>
- Kementerian Komunikasi dan Digital Republik Indonesia. (2025). *Deepfake naik 550%; Kemkomdigi meminta platform global menyediakan fitur cek konten AI*. <https://www.kominfo.go.id>
- Khaeron, R. A. (2025, August 20). Sri Mulyani diterjang deepfake, pembuat dan penyebar bisa dijerat hukum pakai aturan ini. *Metro TV News*. <https://metrotvnews.com>
- Kompas.com. (2025, March 5). *Regulatory sandbox pendukung startup menghasilkan AI tepercaya*. <https://kompas.com>
- Kosinski, M. (2024). *What is the EU AI Act?* IBM.

- Kristiyenda, Y. S., Faradila, J., & Basanova, C. (2025). Pencegahan kejahatan deepfake: Studi kasus terhadap modus penipuan deepfake Prabowo Subianto dalam tawaran bantuan uang. *ALADALAH: Jurnal Politik, Sosial, Hukum dan Humaniora*, 3(2), 149-164. <https://doi.org/10.59246/aladalah.v3i2.1263>
- Kurniati, A., dkk. (2025). The EU AI Act (2024) as a global standard: Analyzing the risk-based approach in artificial intelligence governance. *Jurnal Hukum Internasional dan Teknologi*, 5(1), 58-73.
- Kurniawan, R. (2025). Tantangan penegakan hukum terhadap kejahatan digital berbasis deepfake di Indonesia. *Jurnal Hukum Siber Nusantara*, 10(2), 99-114.
- Legal500. (2024). *Italy's intellectual property is a hot topic*. <https://www.legal500.com>
- Lessig, L. (1999). *Code and other laws of cyberspace*. Basic Books.
- Ma'unah, I., Musyarofah, S., Zahra, A. F. Z. A., & Agrariyanti, Y. (2025). Pancasila and Artificial Intelligence: Ethical Analysis of AI Regulation in Indonesia. *Jurnal Ilmiah Multidisiplin*. <https://litera-academica.com/ojs/litera/article/view/217>
- Mellaz, A. (2025). *Antisipasi perkembangan artificial intelligence (AI) dan model pengawasan digital*. Komisi Pemilihan Umum & Media Kepemiluan. <https://www.lspr.ac.id/navigating-synthetic-democracy-lspr-ajak-mahasiswa-memahami-masa-depan-demokrasi-digital/>
- Mediodecchi, L. (2022). *Kepastian Hukum Pelindungan Data Pribadi Pasca Pengesahan UU Nomor 27 Tahun 2022*. Jakarta: Rajawali Pers.
- Nemecek, A., Jiang, Y., & Ayday, E. (2025). Watermarking without standards is not AI governance. *arXiv*. <https://arxiv.org/abs/2502.13328>
- Nurwahid, T., Muin, S. A., & Ramadhani, R. (2026). Deepfake sebagai Kejahatan: Manipulasi AI Audio, Video, dan Gambar dalam Hukum Pidana Indonesia. *Jurnal FH UMI*. <https://pasca-umi.ac.id/index.php/ahkam/article/view/1917>
- Pearson, J. C., & Nelson, P. E. (2003). *An introduction to human communication: Understanding and sharing* (8th ed.). McGraw-Hill.
- Perdana, A. (2024). Navigating AI regulation through regulatory sandboxes: Balancing innovation and security in developing digital ecosystems. *Jurnal Tata Kelola Digital dan Teknologi Informasi*, 5(2), 112-125.
- Prasetyo, A. B., & Rahmawati, D. (2025). Tantangan regulasi artificial intelligence dalam perlindungan data pribadi di Indonesia. *Jurnal Ilmu Sosial*. <https://jurnal.unigal.ac.id/index.php/jish/article/view/11831>
- Proceedings of the International Conference on Social, Economic, New Education, and Digital Transformation (ICOSEND 2025). (2025). <https://www.atlantispress.com/proceedings/icosend-25>
- Rahmawati, N., & Setiawan, D. (2025). Etika komunikasi digital dalam menghadapi perkembangan artificial intelligence di Indonesia. *ArtComm: Jurnal Komunikasi dan Desain*. <https://journal.unas.ac.id/artcomm/article/view/3154>

- Risnandar, A., & Zaenudin, K. M. I. S. (2025). Development and Evolution of Indonesian Law from the Perspective of Development Law Theory. *Rechtsvinding*. <https%3A%2F%2Frechtsvinding.bphn.go.id%2Findex.php%2Fjrv%2Farticle%2Fview%2F1294>
- Safer Internet Lab. (2025). *Panel 2 – Alia*. <https://saferinternetlab.org>
- Sagena, U., Tukina, & Sari, T. R. (2025). Perlindungan Data Pribadi sebagai Pilar Keberlanjutan dan Pertumbuhan Ekonomi Digital di Indonesia. *Scripta Economica: Journal of Economics, Management, and Accounting*. <https://jurnal.untidar.ac.id/index.php/scriptaeconomica/article/view/10008>
- Sibernusantara.co.id. (2026). *AI bisa mengaburkan kebenaran: pakar hukum IT soroti ancaman deepfake*. <https://sibernusantara.co.id>
- Sulistya, E., Rudolf, J., Yasin, R. L., & Haryanto, M. F. T. S. (2026). Disinformasi dan krisis etika komunikasi publik: analisis kasus deepfake Menteri Sri Mulyani di media sosial. *Orasi: Jurnal Ilmu Politik dan Sosial*, 2(2), 157-168. <https://doi.org/10.33753/governance.v8i1.412>
- Sutrisno, A., Hidayat, R., & Pramudya, D. (2025). Tantangan regulasi kecerdasan artifisial dalam perlindungan data pribadi di Indonesia. *Journal of Contemporary Administration and Social Science*, 3(2), 115-128.
- Tarigan, E. P., & Rumiarta, I. N. P. B. (2025). Kekosongan Huku AI di Indonesia: Kasus Deepfake terhadap Sri Mulyani dan Perbandingan EU AI ACT. *Jurnal Media Akademik*. <https://jurnal.mediaakademik.com/index.php/jma/article/view/3265>
- The Conversation. (2024). *Dua sisi EU AI Act: 10 hal yang bisa dipelajari Indonesia*. <https://theconversation.com>
- Utami, N. P. G. S., Sutrisni, N. K., Suharyanti, N. P. N., et al. (2025). Pengaturan Hukum Tanggung Jawab Pemanfaatan Artificial Intelligence di Indonesia. *Jurnal Hukum Mahasiswa*, 5(2). <https://e-journal.unmas.ac.id/index.php/jhm/article/view/10041>
- Wahyudi Djafar. (2024). *Laporan Evaluasi Kebijakan AI di Indonesia*. Jakarta: ELSAM.
- Wicaksono, D. T., Prasetya, F., & Lestari, N. P. (2026). Penegakan Hukum terhadap Kejahatan Deepfake dalam Perspektif Hukum Positif Indonesia. *Jurnal Fakta Hukum*. <https://ojsstihpertiba.ac.id/index.php/jfh/article/view/233>
- YouTube. (n.d.). *Video YouTube soal "Guru Beban Negara": Kemenkeu says it's a hoax and the result of a deepfake*. <https://youtube.com>