



Regulatory Approaches to AI-Generated Video Content

Yatama Zahra ¹, Amirah ²

¹ Faculty of Law, Sriwijaya University, Palembang, Indonesia,

² Department of Publisher, Lentera Ilmu Publisher

ARTICLE INFO

Article history:

Received 25 June 2026

Revised 11 August 2026

Accepted 11 August 2026

Keywords:

Artificial Intelligence

Video Content

Regulatory Approaches

ABSTRACT

The rapid advancement of artificial intelligence (AI) has transformed the creation and distribution of video content, enabling the production of highly realistic synthetic media at unprecedented speed and scale. While AI-generated videos offer significant benefits in education, entertainment, marketing, and communication, they also create complex legal and regulatory challenges involving misinformation, privacy violations, intellectual property infringement, impersonation, and public trust. This study examines existing regulatory approaches to AI-generated video content and assesses their effectiveness in addressing emerging legal and governance challenges. Using a qualitative literature review and comparative policy analysis, the research examines regulatory frameworks across jurisdictions, focusing on transparency requirements, content labeling, accountability mechanisms, content moderation, privacy protection, and liability provisions. The findings indicate that existing regulations often struggle to keep pace with technological developments, creating gaps in enforcement, transparency, and legal certainty. The study proposes a balanced regulatory framework that promotes innovation while strengthening accountability, transparency, and protection of individual rights.

This is an open access article under the [CC BY-SA](#) license.



1. Introduction

Artificial intelligence (AI) has rapidly become an integral part of Indonesia's digital transformation journey, powering innovations in sectors such as finance, education, health, and public services [1,2]. At the heart of many AI systems lies the collection and processing of personal data, which, if not properly managed, can pose serious risks to privacy and individual rights [3,4]. The benefits of AI—such as improved efficiency, personalized services, and data-driven decision-making—must be weighed

against the potential for misuse, especially in a landscape where data governance is still maturing [5].

In response to these growing concerns, Indonesia enacted the Personal Data Protection (PDP) Law in 2022, marking a significant milestone in the country's effort to strengthen digital rights and data privacy [6]. The PDP Law aims to provide legal clarity around the collection, processing, and protection of personal data, including data used by AI systems [6,7]. However, despite this progress, challenges remain in ensuring that AI technologies comply with the principles outlined in the law, such as transparency, purpose limitation, and accountability [8,9]. Many AI systems operate as black boxes,

* Corresponding author: Amirah
email: amirah@lenterailmu.com

making it difficult for individuals to know how their data is used or to exercise their rights under the law [4,10].

As AI applications become more embedded in daily life, the tension between technological innovation and data protection becomes more pronounced. For example, AI-based facial recognition in public spaces or algorithmic profiling in e-commerce raises important questions about informed consent, data minimization, and fairness [3,9,11]. Enforcement mechanisms under the PDP Law are still developing, and there is a need for stronger collaboration between government agencies, technology companies, and civil society to ensure AI systems are both effective and legally compliant [7,12,13].

2. Literature Review

The rapid development of generative artificial intelligence (AI) has transformed the production and distribution of video content, enabled the creation of highly realistic synthetic media while simultaneously generated significant legal, ethical, and social concerns [10,14]. Previous studies have identified deepfakes, misinformation, impersonation, privacy violations, lack of consent, copyright infringement, and limited content authenticity as major challenges associated with AI-generated video [10,14,15]. In response, regulatory approaches have increasingly

emphasized transparency and disclosure, including AI-content labeling, digital watermarking, metadata, and content provenance mechanisms to help users distinguish synthetic content from authentic content [14,15]. Other regulatory measures focus on privacy and consent, particularly when AI is used to reproduce an individual's face, voice, or likeness without authorization, while copyright frameworks seek to address the use of protected creative works in AI development and content generation. Platform responsibility has also become an important regulatory dimension, requiring digital platforms to establish mechanisms for content detection, reporting, moderation, and removal of harmful synthetic media. However, existing approaches face limitations because AI technologies evolve rapidly, cross-border distribution complicates enforcement, and technical detection mechanisms may become less effective as generation techniques improve [10,15]. Therefore, the literature increasingly supports a risk-based and multi-stakeholder regulatory approach that combines preventive safeguards, transparency mechanisms, accountability measures, legal remedies, and continuous regulatory adaptation. Such an integrated approach is considered necessary to balance technological innovation and freedom of expression with the protection of privacy, intellectual property, individual rights, and public interests.

3. Methods

3.1 Research Design

This study employs a qualitative comparative research approach to examine regulatory approaches to AI-generated video content (Figure 1). The research focuses on identifying, analyzing, and comparing legal and regulatory frameworks governing the creation, modification, distribution, and use of AI-generated and AI-manipulated audiovisual content. A qualitative approach was selected because the study aims to interpret regulatory principles, institutional responsibilities, legal obligations, and governance mechanisms rather than quantitatively measure the performance of a specific technological system. The research process consists of four major stages: source identification, source screening and selection, thematic analysis, and comparative evaluation. Academic literature and regulatory documents are examined to identify the principal risks associated with AI-generated video and the mechanisms adopted by different regulatory approaches to address those risks. The results are subsequently synthesized to identify similarities, differences, regulatory gaps, and potential directions for a more comprehensive governance framework.

3.2 Data Sources and Literature Selection

The study uses secondary data obtained from peer-reviewed journal articles, conference proceedings, books, legislation, regulatory guidelines, government publications, and institutional reports. The literature search focuses on studies addressing generative AI, AI-generated video, deepfakes, synthetic media, content transparency, digital watermarking, content provenance, privacy, consent, copyright, misinformation,

and platform responsibility. Sources are selected according to their relevance to the research objectives, credibility, regulatory significance, and contribution to the understanding of AI-generated audiovisual content. Greater emphasis is placed on recent publications because generative AI technologies and regulatory responses continue to evolve rapidly. Nevertheless, influential earlier studies are retained when they provide important theoretical or conceptual foundations (Table 1).

Table 1. Literature and Regulatory Source Selection Criteria

Criterion	Description
Research relevance	The source directly discusses AI-generated video, deepfakes, synthetic media, or AI regulation
Source credibility	Peer-reviewed publications or authoritative legal, governmental, and institutional sources
Regulatory relevance	The source provides information about laws, policies, governance mechanisms, or regulatory principles
Publication period	Priority is given to recent literature while retaining foundational studies
Content scope	Sources addressing transparency, privacy, consent, copyright, misinformation, accountability, or platform responsibility
Accessibility	Full-text sources that provide sufficient information for systematic analysis

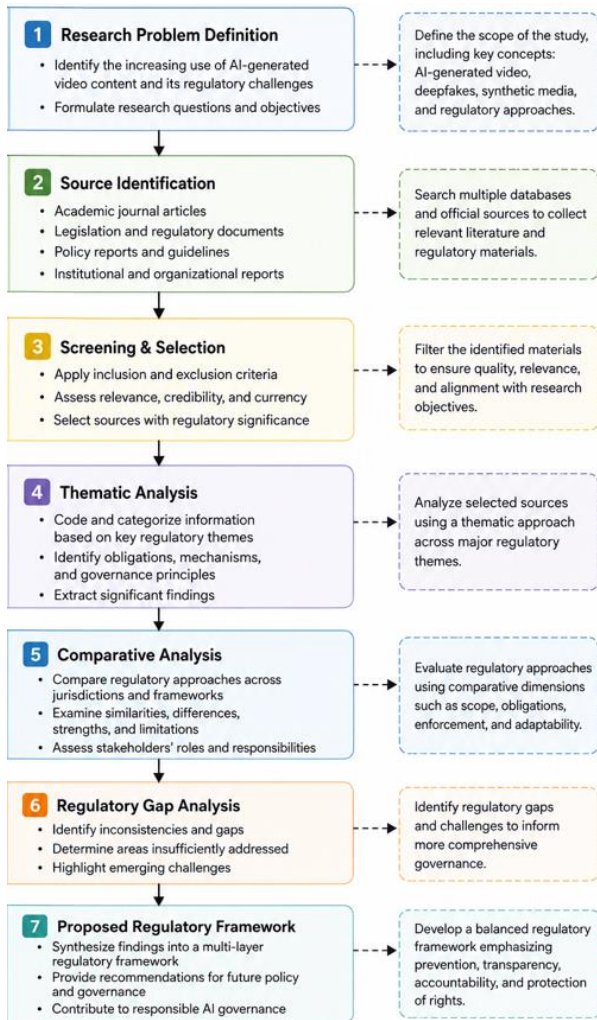


Figure 1 - Research Methodology Framework

3.3 Regulatory Framework Analysis

The selected sources are analyzed using a thematic analysis approach. Each source is examined to identify regulatory principles, stakeholder obligations, governance mechanisms, and enforcement strategies relevant to AI-generated video content. The analysis is organized into eight principal themes: (1) transparency and disclosure, (2) content labeling and watermarking, (3) content provenance and authenticity verification, (4) privacy and consent, (5) copyright and intellectual property, (6) platform responsibility, (7) misinformation and harmful content, and (8) enforcement and accountability (Table 2).

Table 2 - Thematic Analytical Framework

Theme	Key Issues Examined	Regulatory Objective
Transparency & disclosure	AI-generated content identification	Inform users about synthetic content
Labeling & watermarking	Visible labels, invisible watermarks	Increase content traceability
Content provenance	Metadata and authenticity	Establish content origin

Privacy & consent	verification	Face, voice, likeness, personal data	Protect individual rights
Copyright & IP	Training data and generated content	Deepfakes and deceptive content	Protect creators and rights holders
Misinformation	Detection, moderation, reporting	Penalties, removal, legal remedies	Reduce societal and informational harm
Platform responsibility			Increase intermediary accountability
Enforcement			Ensure regulatory compliance

The thematic framework enables the study to systematically compare how different regulatory approaches address similar risks. It also allows regulatory mechanisms to be categorized according to whether they primarily target prevention, transparency, mitigation, or enforcement.

3.4 Comparative Regulatory Analysis

A comparative analysis is subsequently conducted to examine how different jurisdictions and regulatory frameworks address AI-generated video content. The comparison focuses on regulatory scope, affected stakeholders, disclosure requirements, technical safeguards, individual rights, platform obligations, enforcement mechanisms, and adaptability to technological. The comparative analysis is used to identify both common regulatory trends and differences among approaches (Figure 2).

3.5 Analytical Framework

Based on the thematic and comparative analysis, this study categorizes regulatory mechanisms into three complementary layers: preventive, transparency-based, and corrective regulation. The preventive layer focuses on reducing risks before harmful content is created or distributed. It includes consent requirements, risk assessments, safety-by-design principles, and restrictions on high-risk applications. The transparency layer aims to help users distinguish AI-generated or manipulated content from authentic material through labeling, watermarking, metadata, and provenance mechanisms.

4. Results

The analysis indicates that regulatory approaches to AI-generated video content can be broadly categorized into three complementary mechanisms: preventive, transparency-based, and corrective regulation. Preventive mechanisms focus on reducing potential harm before AI-generated content is produced or distributed, particularly through consent requirements, risk assessments, safety-by-design principles, and restrictions on high-risk applications. Transparency-based mechanisms emphasize enabling users to recognize synthetic or manipulated content through AI-content labels, digital watermarking, metadata, and content provenance technologies.

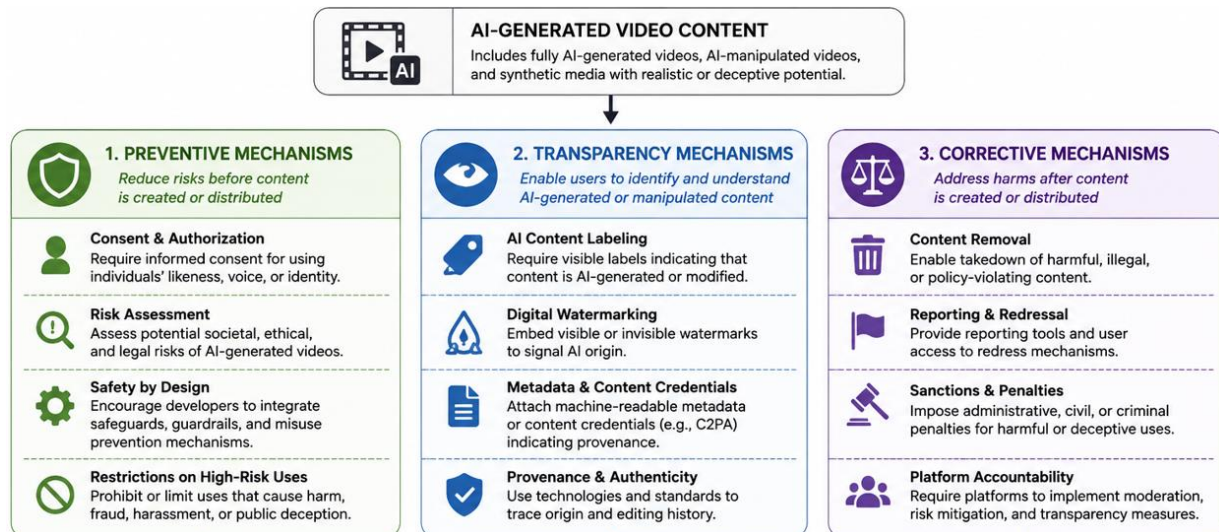


Figure 2. Three-Layer Regulatory Framework

5. Conclusion

This study concludes that the regulation of AI-generated video content requires a comprehensive, risk-based, and multi-stakeholder approach that balances technological innovation with the protection of individual and public interests. The analysis demonstrates that preventive mechanisms, transparency measures, and corrective actions should be implemented as complementary layers to address risks related to deepfakes, misinformation, privacy, consent, copyright, and platform accountability. While technologies such as AI-content labeling, watermarking, and content provenance can improve transparency and authenticity verification, they should be supported by appropriate legal requirements, enforcement mechanisms, and clear responsibilities for AI developers, content creators, digital platforms, and users.

REFERENCES

- [1] OECD. (2019). *Recommendation of the Council on Artificial Intelligence*. OECD.
- [2] UNESCO. (2021). *Recommendation on the Ethics of Artificial Intelligence*. UNESCO.
- [3] European Union. (2016). *Regulation (EU) 2016/679 (General Data Protection Regulation)*.
- [4] Burrell, J. (2016). How the machine “thinks”: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12. <https://doi.org/10.1177/2053951715622512>
- [5] NIST. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1)*. National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.600-1>
- [6] Republic of Indonesia. (2022). *Law No. 27 of 2022 concerning Personal Data Protection*.
- [7] Badan Pemeriksa Keuangan Republik Indonesia. (2022). *Undang-Undang Nomor 27 Tahun 2022 tentang Pelindungan Data Pribadi*.
- [8] Kroll, J. A., Huey, J., Barocas, S., Felten, E. W., Reidenberg, J. R., Robinson, D. G., & Yu, H. (2017). Accountable algorithms. *University of Pennsylvania Law Review*, 165(3), 633–705. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2765268
- [9] Selbst, A. D., Boyd, D., Friedler, S. A., Venkatasubramanian, S., & Vertesi, J. (2019). Fairness and abstraction in sociotechnical systems. *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 59–68. <https://doi.org/10.1145/3287560.3287598>
- [10] Mirsky, Y., & Lee, W. (2021). The creation and detection of deepfakes: A survey. *ACM Computing Surveys*, 54(1), Article 7. <https://doi.org/10.1145/3425780>
- [11] U.S. Copyright Office. (2024). *Copyright and Artificial Intelligence, Part 1: Digital Replicas*. U.S. Copyright Office.
- [12] Citron, D. K., & Chesney, R. (2019). Deepfakes and the new disinformation war. *Foreign Affairs*, 98(1), 147–155.
- [13] Brundage, M., Avin, S., Clark, J., et al. (2018). *The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation*. Future of Humanity Institute, University of Oxford. <https://doi.org/10.17863/CAM.22520>
- [14] European Union. (2024). *Regulation (EU) 2024/1689 (Artificial Intelligence Act)*.
- [15] Kietzmann, J., Lee, L. W., McCarthy, I. P., & Kietzmann, T. C. (2020). Deepfakes: Trick or treat? *Business Horizons*, 63(2), 135–146. <https://doi.org/10.1016/j.bushor.2019.11.006>