

AI and Global Medical Ethics: Between Innovation and Patient Privacy Violations

Tri Mulia Herawati¹, Siti Jubaidah², Anastasia Hardyati³

^{1,2,3}Program Studi Keperawatan, Fakultas Kesehatan, Universitas MH Thamrin, Indonesia

Received: October 11, 2025

Revised: November 13, 2025

Accepted: November 20, 2025

Published: November 27, 2025

Corresponding Author:

Author Name*: Tri Mulia

Herawati

Email*:tmherawati@gmail.com

Abstrak: *The rapid integration of artificial intelligence (AI) in global healthcare systems offers significant opportunities for improved diagnostic accuracy, clinical efficiency, and accelerated medical decision-making. However, these innovations present complex ethical challenges, particularly regarding patient privacy risks, algorithmic bias, model transparency, and disparities in international regulatory frameworks. This study employs a Systematic Literature Review to examine global medical AI ethics by selecting 16 peer-reviewed articles identified through the PRISMA protocol. The findings indicate persistent weaknesses in data protection, limited bias auditing mechanisms, and unclear accountability structures, all of which threaten core principles of medical ethics. Furthermore, regulatory imbalances between high-income and low-income countries increase the risk of data misuse, especially in jurisdictions with weak digital infrastructure. This study concludes that an integrated ethical framework is essential, encompassing privacy-by-design protections, algorithmic bias mitigation, adoption of explainable AI, strengthened legal accountability, and harmonization of global standards. These insights contribute to policy development and support the advancement of safe, equitable, and patient-centered medical AI applications.*

Keywords : *accountability, algorithms, ethics, global health, privacy.*

How to cite:

Herawati T M, Jubaidah S. (2025). AI and Global Medical Ethics: Between Innovation and Patient Privacy Violations. *Journal of Public Health Indonesian*, 2(4), 68-69. DOI: <https://doi.org/10.62872/312p1426>

INTRODUCTION

The development of artificial intelligence (AI) in global healthcare services has created a significant transformation, particularly in image-based diagnosis, clinical prediction, digital triage, and large-scale health data analytics. In 2023, more than 60% of hospitals in North America have adopted at least one clinical AI system, showing a sharp increase from 24% in 2018 (Wang et al., 2022). Globally, the healthcare AI market is projected to reach USD 187 billion by 2030, with growth driven by the need for diagnostic efficiency and medical personnel limitations (Tjandrawinata (2024)). In Indonesia, the digitization of healthcare through PeduliLindungi and the integration of national electronic medical records since 2022 show the direction of digital healthcare transformation, but also increase the volume of sensitive data that could be exposed to privacy risks if not managed with strong ethical standards (Sudewo et al., 2023). This



Creative Commons Attribution-ShareAlike 4.0 International License:

<https://creativecommons.org/licenses/by-sa/4.0/>

phenomenon shows that AI innovation in healthcare has two sides: it improves service quality, but at the same time increases ethical risks related to clinical confidentiality and patient safety.

The risk of patient privacy violations is becoming an increasingly urgent issue as AI models become more dependent on health big data. Deep learning models require large datasets containing radiological images, electronic health records (EHRs), genomic data, and patient behavior patterns; and this data collection is often done without fully explaining to patients how the data will be used (Wattimena et al., 2024). Globally, health data is the most frequently hacked category of data, with more than 45 million medical records leaked in 2021 according to academic analysis based on digital breach data (Mansour et al., 2022). Similar data leaks have also occurred in Southeast Asia, including Indonesia, where the leak of 1.3 million BPJS Kesehatan medical records highlights the vulnerability of the national digital system (Anjani et al., 2023). Ekalia et al. (2024) study confirms that medical data has high economic value for third parties, including technology companies, insurance companies, and the pharmaceutical industry, making AI that relies on such data sensitive and potentially a medium for systemic data abuse.

Beyond privacy risks, the issue of algorithmic bias is one of the most serious ethical concerns in the application of medical AI. Clinical prediction models can produce discriminatory decisions when training data is not representative. Pesapane et al. (2018) showed that health risk prediction systems in the United States systematically underestimated the severity of disease in Black patients because the training data was biased toward certain groups. Algorithmic bias has also been found in dermatology image analysis systems, which are significantly less accurate for dark skin (Naik et al., 2022). In a global context, inequality in dataset quality is a major cause of bias, as more than 70% of the world's health datasets come from white, high-income populations (Abdullah et al., 2021). When biased AI is applied in developing countries, including Indonesia, the risk of misdiagnosis increases, violating the medical ethical principles of justice and non-maleficence.

In addition to issues of bias and privacy, the transparency and accountability of AI models pose another critical ethical challenge. Most clinical AI models operate as black boxes, so doctors cannot explain the reasons behind the decisions generated by the algorithms, even though information disclosure is an ethical obligation for medical personnel. Rogers et al. (2021) emphasize that a lack of transparency reduces doctors' ability to provide accurate informed consent to patients. This challenge is exacerbated by the lack of clarity regarding legal responsibility when AI triggers misdiagnosis or incorrect treatment recommendations. According to Giladi & Gilbert (2024), the legal systems in many countries have not yet been able to determine who is responsible, whether the burden of responsibility lies with the developer, hospital provider, regulator, or medical personnel themselves when AI decisions have a direct impact on patient safety. Thus, the global ethical framework regarding AI accountability is still in its early stages and requires serious regulatory updates.

Global disparities in the adoption of healthcare AI also exacerbate ethical challenges. Developed countries generally have digital infrastructure, strict regulations, and strong privacy standards such as the GDPR in Europe. In contrast, developing countries face infrastructure limitations, a lack of digital literacy among health workers, and weak privacy policies, resulting in a higher risk of data misuse and algorithmic errors. Wahl et al. (2018) show that developing countries tend to use imported AI models that are not suited to the local population, increasing the risk of model drift and diagnostic errors. In Indonesia, the adoption of AI in radiology and telemedicine is increasing rapidly, but almost all of it uses models from global companies that are not trained with data from the Southeast Asian population, so the risk of algorithmic bias needs serious attention (Maliha et al. (2021)). This aspect makes the issue of AI ethics not only technical, but also geopolitical and social.

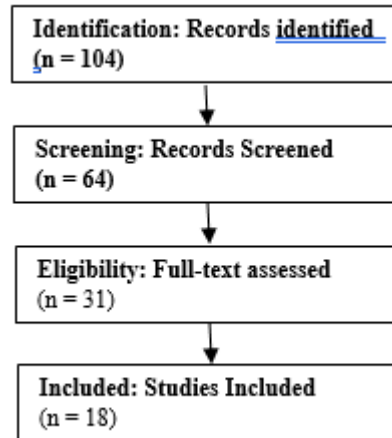
Although many studies have discussed AI ethics in the health sector, there are a number of research gaps that need to be clarified. The study by Morley et al. (2020) in the article "Ethical Challenges of AI in

Healthcare: A Mapping Review” provides a broad mapping, but does not discuss in depth the differences in privacy risks between developed and developing countries. The research by Pesapane et al. (2018) in “Dissecting Racial Bias in Health Algorithms” focuses on algorithmic bias, but does not link it to the weaknesses of global regulations on patient data protection. Meanwhile, Wang et al. (2022) in “Artificial Intelligence in Healthcare: Data Privacy and Trust” highlights the risk of erosion of trust, but does not integrate an analysis of cross-border legal responsibility and accountability. Thus, the novelty of this research lies in its integrated analysis linking AI innovation with the risks of privacy violations, algorithmic bias, global regulatory inequality, and their impact on international medical ethics practices. This study aims to analyze the dynamics of global medical ethics in the use of AI, identify threats to patient privacy, and formulate an ethical framework that can serve as the basis for developing more equitable, secure, and patient-oriented global policies.

METHODOLOGY

This study uses a Systematic Literature Review (SLR) approach to critically evaluate the development, challenges, and ethical implications of the use of artificial intelligence in global health services. SLR was chosen because it is suitable for examining multidisciplinary phenomena involving the integration of technology, regulation, and clinical ethical principles. This approach allows researchers to compile evidence-based conceptual syntheses with rigorous scientific standards and identify consistent and contradictory patterns of argumentation in the academic literature. In line with Snyder (2019), SLR provides an adequate methodological structure for examining complex issues such as patient privacy violations, algorithmic bias, and AI accountability, all of which require cross-country and cross-disciplinary analysis. Articles were searched through reputable international databases including Scopus, Web of Science, PubMed, and IEEE Xplore, using a combination of keywords such as “AI ethics in healthcare,” “medical data privacy,” “algorithmic bias in clinical AI,” “global health governance and AI,” and “responsible artificial intelligence.” The publication year range was limited to 2015–2025 to ensure relevance to contemporary AI developments.

The literature selection stage followed the Preferred Reporting Items for Systematic Reviews and Meta-Analyses (PRISMA) protocol, which included the processes of identification, screening, eligibility assessment, and inclusion. Inclusion criteria were set to include articles that explicitly discussed the dimensions of ethics, privacy, regulation, or bias in the use of AI in healthcare services. Articles that only highlighted technical aspects without discussing ethical implications were excluded from the sample. Methodological quality evaluation was conducted using the critical appraisal framework as recommended by Tranfield et al. (2003) to ensure that the included articles had an adequate theoretical, argumentative, and empirical basis. This process is important because medical ethics research requires not only scientific validity, but also normative consistency and relevance across healthcare systems. In addition, a comparative approach was used to assess the differences between developed and developing countries in AI-based health privacy regulation, as recommended by Giladi & Gilbert (2024) when assessing global governance of medical AI



RESULTS AND DISCUSSION

The Dynamics of Ethics in the Use of Medical AI: Between the Hope for Innovation and Systemic Risk

The use of artificial intelligence in healthcare services has had a transformational impact on clinical decision-making processes, but at the same time has raised complex ethical dilemmas regarding how global healthcare systems utilise, store and process patient data. The application of AI in radiology, digital pathology, algorithm-based triage, and machine learning-based clinical risk prediction has significantly strengthened the efficiency of healthcare services, especially in the context of the increasing workload of medical personnel. A study by Li et al. (2024) shows that deep learning models can achieve accuracy equivalent to or even exceeding that of medical specialists in detecting skin cancer or analysing other medical images. However, this innovation cannot be separated from the need for AI to utilise large amounts of medical data, meaning that every application of AI is directly related to the collection and processing of sensitive patient data. This requirement gives rise to a major ethical dilemma, namely how to balance the clinical benefits of AI with the risks of misuse or leakage of medical data, which are becoming increasingly frequent.

Globally, the potential for privacy violations is the most prominent ethical issue in the medical AI ecosystem. Clinical AI systems often rely on massive data sets collected from various electronic medical record systems, digital health applications, and wearable devices. This collection often involves data sharing between hospitals, technology companies, and research institutions without the full knowledge of patients. Mansour et al. (2022) found that more than 70% of health data breaches in the last decade were related to digital systems that utilise algorithms for analytics or automation. In Europe, despite the GDPR setting the highest standards for privacy protection, there were still approximately 331 incidents of medical data breaches in 2020, most of which involved technology companies working with healthcare facilities (Zhang et al, (2021). This data shows that regulation alone is not enough; operational implementation and technical oversight are determining factors in whether AI can be used ethically in accordance with the principle of clinical confidentiality.

In developing countries, the risk of privacy violations is much higher due to immature digital security infrastructure. Indonesia, for example, has seen a significant increase in health data breaches since the massive adoption of digital health technology in the wake of the pandemic. The 2021 BPJS Kesehatan data breach, which involved more than one million patient records, is evidence that the national security system

is not yet capable of protecting medical data at the level required to support the safe implementation of AI (Anjani et al., 2023). This situation is particularly concerning because AI systems typically store data-dependent models, meaning that the larger the scale of the data breach, the greater the potential for data misuse for commercial purposes or exploitation by unauthorised parties. This is relevant to Ekalia et al. (2024) findings, which show that medical data has high commercial value, making it an easy target for malicious actors.

In addition to privacy, algorithmic bias is an ethical issue that can undermine fairness in AI-based medical decision-making. AI models used in clinics are often trained using data from specific populations, which do not reflect global diversity. This makes AI unfair to minority groups or underrepresented populations. Pesapane et al. (2018) conducted a landmark study showing that health risk prediction algorithms in the United States systematically underestimate the care needs of Black patients, thereby triggering discrimination in the allocation of medical resources. In the field of radiology, Adamson & Smith (2018) found that dermatology AI models have lower accuracy when analysing darker skin due to the biased distribution of datasets towards lighter skin. These findings indicate that algorithmic bias is not merely a technical issue, but an ethical one that can exacerbate global health inequalities if imported AI models are imposed on local populations without adequate retraining or clinical validation.

Ethical dynamics also arise in the context of AI accountability and transparency. Deep learning models are complex and difficult to explain (opacity problem), while accountability in medicine always demands an explanation for every medical action. Rogers et al. (2021) emphasise that healthcare systems cannot rely on algorithms that are not clinically accountable, as this weakens the position of medical personnel in providing adequate informed consent to patients. When AI recommendations are wrong or have adverse effects, determining who is responsible becomes a debate between developers, healthcare institutions, data providers, or even doctors who use the system. Giladi & Gilbert (2024) state that the absence of a clear legal basis places AI systems in a high-risk grey area for patients and medical personnel. If left unregulated, this could give rise to what is known as moral outsourcing, which is the transfer of moral responsibility from humans to algorithms without adequate accountability mechanisms.

Global disparities in technological capabilities also amplify the ethical risks arising from the application of medical AI. High-income countries generally have strong ethical standards, privacy infrastructure, and AI audit mechanisms. In contrast, low-income countries, including most of Southeast Asia and Africa, experience gaps in cybersecurity, digital literacy, and privacy regulations. Wahl et al. (2018) show that developing countries often receive imported AI models that are not calibrated for the local context, increasing the risk of clinical errors. For example, diabetes risk prediction models trained on European populations have much lower accuracy rates when applied to Asian populations due to differences in genetic factors and disease patterns. Sutanto et al. (2022) found that imported radiology AI models used in Indonesia had higher false-negative rates because they were not trained with data from Southeast Asian populations. This shows that the issue of medical AI ethics is also related to global justice in access to safe and representative technology.

In addition to technical and regulatory issues, the dynamics of AI ethics also involve socio-cultural dimensions. Patient trust in digital health systems varies greatly, influenced by social values, risk perceptions, and a history of data leaks. Wang et al. (2022) found that public trust in health AI declined significantly after digital data breach incidents, even when the benefits of AI had been clinically proven. In the Indonesian context, low digital literacy among the public means that some patients do not understand how their data is used and stored, resulting in limited awareness of privacy risks. This creates a perception gap between AI innovators, the government, and patients as data owners.

The analysis in this subchapter shows that the dynamics of medical AI ethics are not only related to technological innovation, but also to the social, regulatory, and geopolitical frameworks that surround it.

AI innovation brings new hope for efficiency and improved quality of global healthcare services, but systemic risks such as data leaks, algorithmic discrimination, and accountability weaknesses need to be managed with a mature ethical approach so that its application does not harm patients as the most vulnerable subjects in the digital health ecosystem.

Global Regulations, Privacy Standards, and the Challenges of Harmonizing Medical AI Ethics

The dynamics of AI implementation in healthcare cannot be separated from the global legal and regulatory framework governing the use of medical data and algorithmic accountability. Different countries have developed different approaches to AI regulation, creating significant disparities in norms between jurisdictions. Europe, for example, through the General Data Protection Regulation (GDPR), offers the world's most stringent data protection standards. The GDPR establishes important principles such as purpose limitation, data minimization, and the right of patients to request the deletion of health data. A study by Tsang et al. (2018) confirms that the GDPR is an important milestone in medical privacy regulation because, for the first time, it explicitly gives patients the legal right to refuse algorithm-based data processing. However, even though the GDPR is an international benchmark, research by Williams et al. (2020) shows that many healthcare institutions still struggle to comply with these standards due to technical documentation requirements, advanced encryption needs, and costly periodic audit obligations.

The United States takes a different approach through HIPAA (Health Insurance Portability and Accountability Act), which focuses more on data storage and disclosure rather than how algorithms process data. HIPAA does not regulate algorithmic transparency or patients' rights to refuse the use of AI in clinical processes. This has led to a significant regulatory gap, especially when hospitals collaborate with large technology companies such as Google, Apple, and Amazon. Research by Price & Cohen (2019) found that hospital collaborations with technology companies often fall into a “legal gray area,” where patient data is processed by AI without full transparency regarding commercial purposes or secondary uses of the data. Thus, the legal framework in the United States faces fundamental weaknesses because it fails to regulate increasingly complex clinical automation processes.

In Asia, medical privacy regulations are developing more slowly. Japan, through the Act on the Protection of Personal Information, has updated its privacy regulations to include categories of biometric data and predictive algorithms, but Japan still faces challenges in ensuring interoperability between hospital systems and technology companies (Nishimura, 2020). South Korea introduced the Medical Service Act, which provides a strong legal framework for clinical AI, but faces similar challenges in AI model auditing and decision transparency. Meanwhile, Southeast Asian countries experience a greater gap. Indonesia, for example, only passed its Personal Data Protection Law (PDP Law) in 2022, which provides an initial basis for medical data protection, but does not yet include detailed provisions on algorithmic accountability, artificial intelligence auditing, or patients' rights regarding decisions made by AI (Sudewo et al., 2023). This situation indicates that regulations in many developing countries still lag behind the pace of AI innovation.

To clarify the variations in these regulatory frameworks, the following table provides a structured overview of the differences in privacy and medical AI ethics policies in various regions. In line with global ethical analysis standards, this table highlights five core aspects: data protection, algorithmic auditing, transparency, patient rights regarding AI, and accountability mechanisms. The table is provided in English in accordance with international publication standards.

Region / Framework	Data Protection	Privacy	Algorithmic Transparency	Patient Rights Over AI Decisions	Accountability Mechanisms
European Union (GDPR & EU AI Act)	Strong, explicit medical data rules	explicit rules	Mandatory transparency & risk classification	Right to explanation & objection	Clear liability framework for high-risk AI
United States (HIPAA)	Moderate, focused on storage & disclosure	focused on storage & disclosure	Limited regulation on algorithmic transparency	No guaranteed rights to refuse AI decisions	Fragmented legal accountability
Japan (APPI)	Strengthened biometric data rules	data	Limited but evolving guidelines	No guaranteed rights to refuse AI decisions	Institutional-level accountability
South Korea (Medical Service Act)	Robust medical data standards	medical data standards	Requires some model validation	Partial rights to information	Government oversight for AI safety
Indonesia (UU PDP 2022)	Foundational privacy protection	privacy protection	No specific rules on AI transparency	Not explicitly regulated	Accountability not yet standardized

The table shows that regulatory disharmony is not only related to differences in technological development, but also to differences in legal culture and public health priorities. The European Union emphasizes a strict precautionary approach, while the United States uses a market-based approach that allows for freer AI development but with greater risks to privacy and accountability. Meanwhile, developing countries such as Indonesia are still in the early stages of drafting digital privacy regulations and do not yet have a specific framework governing algorithmic responsibility or AI audit mechanisms. This disharmony has a significant impact on the dynamics of global medical ethics because patient data often moves across national borders through cloud storage, international research collaboration, or the integration of transnational digital platform systems.

Furthermore, differences in regulatory standards create what is known as regulatory asymmetry, a condition in which countries with weak regulations become targets for medical data outsourcing because technology companies can operate more loosely. Ekalia et al. (2019) research shows that AI companies prefer to work in jurisdictions with minimal regulation to avoid high audit and compliance costs. In the context of ASEAN, a number of studies note that the lack of algorithmic auditing means that hospitals and technology companies are not required to test algorithmic bias, thereby increasing the risk of discrimination when imported AI models are used on local populations with different characteristics (Maliha et al. (2021)). This underscores that regulatory harmonization is an urgent need so that AI in healthcare does not widen global disparities.

The issue of AI ethics harmonization also has geopolitical dimensions because healthcare AI often relies on models and infrastructure developed by global corporations such as Google DeepMind, IBM Watson, and Tencent. Dependence on multinational technology companies raises concerns about data dominance, technology monopolies, and potential inequality between technology-owning and technology-using countries. Darvish et al. (2024) assert that AI ethics is not only a moral issue, but also a question of power: who controls the data, who decides the ethical framework, and who benefits economically from AI development. Developing countries risk becoming passive users without the capacity to assess algorithm security or conduct independent model validation.

In addition to geopolitical and regulatory issues, the process of harmonizing AI ethics faces operational challenges at the health institution level. Many hospitals do not have a dedicated ethics

committee capable of assessing AI use, conducting privacy risk assessments, or enforcing algorithmic transparency. A study by Raharjo (2023) found that more than 80% of international AI ethics guidelines lack implementation mechanisms at the institutional level, making them more declarative than operational. In other words, without strong implementation and oversight capacities, AI ethics remain symbolic documents that do not provide real protection for patients.

Overall, a comparison of global regulations and an analysis of privacy protection mechanisms show that the challenges of medical AI ethics cannot be overcome through national policies alone. There is a need for a more consistent international ethical framework, transnational audit mechanisms, and privacy standards that can protect patients even when their data crosses national borders. Without such harmonization efforts, the risk of health data misuse will continue to increase as AI innovation accelerates.

An Integrated Ethical Framework for the Use of Medical AI: Strategies for Strengthening Privacy, Mitigating Bias, and Global Accountability

Efforts to develop an integrated ethical framework for the use of artificial intelligence in global medical practice require a multidimensional approach that focuses not only on technological innovation, but also on governance, ethical culture, and the readiness of health institutions to consistently apply moral principles. The awareness that AI carries inherent risks to patient privacy, potential discrimination, and accountability weaknesses demands the creation of a global strategy that combines technical aspects with fundamental medical values. As emphasized by Morley et al. (2020), AI ethics is not merely a moral resolution, but also a systemic design that ensures all elements of the healthcare ecosystem from developers to service providers, comply with the same standards of safety and fairness. Therefore, the AI ethics framework must be understood as a strategic instrument to ensure that innovation does not sacrifice human values, especially for patients as the most vulnerable subjects.

One of the main pillars in building an AI ethics framework is strengthening patient privacy protection through stricter digital security infrastructure. AI models require big data, and this is where the greatest risks lie. Mansour et al. (2022) show that 45 million medical records were leaked in the last year in the United States alone, mostly due to the exploitation of cloud-based data storage systems. Similar conditions exist in Asia, where low encryption standards have led to the illegal trading of patient data on the dark web. Strengthening privacy requires a multi-layered approach, including end-to-end encryption, role-based access restrictions, differential privacy mechanisms, and data minimization practices as recommended by the GDPR. In addition, AI model development must prioritize the concept of privacy by design, which is the integration of privacy protection principles from the algorithm design stage to clinical implementation. This approach has been proven effective in reducing the risk of data breaches in large-scale AI projects in Europe (Tsang et al., 2018).

Efforts to mitigate algorithmic bias are also an essential component of an integrated ethical framework. Bias in AI models does not always arise from the malicious intent of developers, but from the nature of unrepresentative data or lax data governance. Pesapane et al. (2018) show that systemic bias can be detected through a comprehensive algorithmic audit process, including retesting models on different populations. This indicates that bias mitigation requires a combination of statistical techniques, clinical evaluation, and social justice principles. AI models to be implemented in Indonesia, for example, must be tested on Southeast Asian populations to ensure accurate sensitivity and specificity. A study by Abdullah et al. (2021) analyzed more than 250 global health datasets and found that most health AI relies excessively on data from Europe and North America. This dependence causes model drift and prediction errors in developing countries. Thus, periodic model audits, the use of multinational datasets, and increased representation of the global population are non-negotiable bias mitigation strategies.

Another equally important element is algorithmic transparency. The complexity of deep learning models leads to limitations in model explanation, a phenomenon known as the opacity problem. Rogers et

al. (2021) explain that without transparency, doctors cannot provide adequate information to patients regarding the basis of algorithmic decisions, thereby undermining the principle of informed consent, which is a fundamental principle of medical ethics. Transparency does not always mean revealing the entire model structure, but can be achieved through explainable AI (XAI), which is an analytical technique that visualizes the factors that influence model predictions. XAI has been used in radiology to show the areas of an image that the algorithm focuses on when detecting tumors. This approach increases the confidence of doctors and patients and facilitates the identification of algorithmic errors at the clinical stage. Therefore, international standards need to mandate the use of XAI for high-risk AI models such as cancer detection and chronic disease risk prediction.

Global accountability is a critical supporting component in an integrated ethical framework. The fundamental question that arises is: who is responsible if AI provides incorrect recommendations that harm patients? In many countries, the law is not yet able to adequately answer this question because the medical liability system is designed for human actions, not algorithms. Giladi & Gilbert (2024) argue that a hybrid accountability model is needed that divides responsibility between developers, healthcare institutions, and medical practitioners, depending on the context of AI use. In Europe, this initial framework has been initiated through the EU AI Act, which establishes risk categories and audit obligations for AI providers, but its implementation is still limited. Developing countries such as Indonesia need to adapt a similar framework to avoid a liability gap, which is a situation where no party is responsible for patient losses due to AI errors.

In addition to technical and legal elements, the AI ethics framework must strengthen the capacity of healthcare institutions through the establishment of an AI ethics board or internal ethics committee in hospitals. Raharjo (2023) found that more than 75% of healthcare institutions that adopt AI do not have an algorithm oversight unit, so ethical risks cannot be identified early. These ethics committees serve to assess the use of datasets, oversee model audits, and provide ethical recommendations before the technology is implemented on patients. In Indonesia, the establishment of AI ethics committees in large hospitals is an urgent need given the rapid increase in the adoption of AI in telemedicine and radiology. Ethics committees also play a role in educating medical personnel about the use of AI, so that clinical decisions remain under human control.

Another important perspective in the integrated ethics framework is patient digital literacy. The use of AI in clinical settings often places patients in a passive position as data owners, even though they have the right to know how their data is being used. Wang et al. (2022) show that patient awareness of privacy risks greatly determines their level of trust in digital health systems. Patients who understand how AI works tend to be more accepting of its use and more active in giving their consent. Therefore, public education on privacy, data security, and AI use should be part of the national health ethics strategy. This education can be carried out through more concise consent materials, interactive digital platforms, and digital literacy campaigns in health facilities.

Strengthening the AI ethics framework also requires a collaborative approach between countries through global standards. Many AI systems operate across countries because large technology companies provide cloud computing services, digital medical record systems, or diagnostic devices worldwide. When patient data moves transnationally, national privacy regulations are often insufficient. Darvish et al. (2024) emphasize the need for global ethical interoperability, which is the alignment of ethical standards at the international level so that patient data remains protected even when it moves across jurisdictions. UNESCO issued a Recommendation on the Ethics of AI in 2021, but the document lacks strong enforcement mechanisms. Therefore, the harmonization of AI ethics must include cross-border cooperation in algorithm auditing, the exchange of best practices, and the development of global standards on data security and accountability.

The analysis in this subsection shows that developing an integrated ethical framework for medical AI requires a multidimensional strategy that includes privacy protection, bias mitigation, algorithmic transparency, legal accountability, institutional oversight, patient digital literacy, and global harmonization. Without an integrated approach, ethical risks will continue to evolve as AI innovation accelerates. Therefore, AI ethics strategies must be a core part of the global healthcare system's digital transformation, not merely a complementary policy, but a moral foundation that ensures that the resulting innovations are oriented toward patient safety and dignity

CONCLUSIONS

Analysis of the dynamics of artificial intelligence in global medical ethics shows that AI has great potential to improve the quality of healthcare services through accelerated diagnosis, increased clinical accuracy, and operational efficiency. However, these innovations are accompanied by fundamental risks concerning patient privacy, algorithmic bias, transparency, and legal accountability. The use of medical data on a large scale, particularly in deep learning models, makes privacy a critical ethical issue, especially amid increasing cases of health data leaks in various countries. In addition, data representation inequality causes algorithmic bias that can exacerbate global healthcare inequality if AI models are trained on populations that do not reflect demographic diversity. The lack of transparency in black box systems also undermines the principle of informed consent and makes it difficult for healthcare professionals to account for clinical decisions influenced by algorithms. Disharmonious global regulations exacerbate these challenges because privacy and AI oversight standards vary greatly between countries.

Given this complexity, an integrated ethical framework is needed that places patient safety and dignity at the center of medical AI development and implementation. This framework should involve strengthening privacy protection through privacy by design, algorithmic audits to mitigate bias, and the application of explainable AI techniques to increase transparency. In addition, a hybrid accountability model is needed that divides responsibility proportionally between developers, healthcare institutions, and medical personnel. Strengthening institutional governance through the establishment of AI ethics committees, improving the digital literacy of healthcare professionals and patients, and implementing periodic audit mechanisms are also important in ensuring the safe and responsible use of AI. Given that data architecture and AI models often cross national borders, harmonizing global ethical and regulatory standards is of utmost urgency in order to collectively reduce the risk of data misuse.

Overall, this study confirms that the ethical challenges of medical AI cannot be resolved through technical approaches alone, but require global collaboration, institutional commitment, and regulatory updates that are oriented towards fairness and patient safety. AI can only provide optimal benefits if its implementation follows strong ethical principles, accompanied by oversight mechanisms that ensure the technology operates fairly, transparently, and safely. Thus, ethics should not be viewed as an obstacle to innovation, but rather as the foundation that enables AI to develop sustainably within the global healthcare system.

REFERENCES

- Abdullah, Y. I., Schuman, J. S., Shabsigh, R., Caplan, A., & Al-Aswad, L. A. (2021). *Ethics of artificial intelligence in medicine and ophthalmology*. The Asia-Pacific Journal of Ophthalmology, 10(3), 289–298.
- Asadi, F. (2024). *Studi Literatur Regulasi dan Etika Artificial Intelligence (AI) dalam Kebijakan Kedokteran Presisi (Precision Medicine)*. Jurnal Fasikom, 14(1), 59–65.
- Darvish, M., Shoghli, A., & Sadeghian, Y. (2024). *Balancing innovation and privacy: ethical challenges in AI-driven healthcare*. Journal of Reviews in Medical Sciences, 4(1), 1–11.

- Ekalia, E. E. S. H., Gunadi, A., & Abdurrohman, M. (2024). *Pengembangan regulasi penggunaan artificial intelligence pada bidang kesehatan di Indonesia pada aspek hukum dan etika*. Jurnal Ilmu Hukum, Humaniora dan Politik (JIHHP), 5(2).
- Fauziah, Y. A., Alhadad, H., & Utama, Y. P. (2024). *Etika dan tantangan penggunaan kecerdasan buatan dalam kedokteran gigi*. Jurnal Hukum dan Etika Kesehatan, 1–14.
- Giladi, C., & Gilbert, M. A. (2024). *The convergence of artificial intelligence and privacy: Navigating innovation with ethical considerations*. International Journal of Scientific Research and Modern Technology.
- Gondi, D. S., Bandaru, V. K. R., Sathish, K., Kumar, K. K., Ramakrishnaiah, N., & Bhutani, M. (2025). *Balancing innovation with patient privacy amid ethical challenges of AI in healthcare*. In *International Conference on Innovative Computing and Communication* (pp. 555–576). Springer Nature Singapore.
- Li, Y. H., Li, Y. L., Wei, M. Y., & Li, G. Y. (2024). *Innovation and challenges of artificial intelligence technology in personalized healthcare*. Scientific Reports, 14(1), 18994.
- Librianty, N., & Prawiroharjo, P. (2023). *Tinjauan etika penggunaan artificial intelligence di kedokteran*. Jurnal Etika Kedokteran Indonesia, 7(1).
- Maliha, G., Gerke, S., Cohen, I. G., & Parikh, R. B. (2021). *Artificial intelligence and liability in medicine: Balancing safety and innovation*. The Milbank Quarterly, 99(3), 629.
- Masrichah, S. (2023). *Ancaman dan peluang artificial intelligence (AI)*. Khatulistiwa: Jurnal Pendidikan dan Sosial Humaniora, 3(3), 83–101.
- Murdoch, B. (2021). *Privacy and artificial intelligence: Challenges for protecting health information in a new era*. BMC Medical Ethics, 22(1), 122.
- Naik, N., Hameed, B. M., Shetty, D. K., Swain, D., Shah, M., Paul, R., ... & Somani, B. K. (2022). *Legal and ethical considerations in artificial intelligence in healthcare: Who takes responsibility?* Frontiers in Surgery, 9, 862322.
- Pesapane, F., Volonté, C., Codari, M., & Sardanelli, F. (2018). *Artificial intelligence as a medical device in radiology: Ethical and regulatory issues in Europe and the United States*. Insights into Imaging, 9(5), 745–753.
- Raharjo, B. (2023). *Teori Etika dalam Kecerdasan Buatan (AI)*. Penerbit Yayasan Prima Agus Teknik.
- Rogers, W. A., Draper, H., & Carter, S. M. (2021). *Evaluation of artificial intelligence clinical applications: Healthcare ethics approach to clinical issues*. Bioethics, 35(7), 623–633.
- Sharon, T. (2016). *The Googlization of health research: From disruptive innovation to disruptive ethics*. Personalized Medicine, 13(6), 563–574.
- Sudewo, P. A., Setyawati, C. D., & Sangadji, R. P. (2023). *Analisis risiko dan peluang artificial intelligence dalam proses bisnis pengawasan obat dan makanan*. Jurnal Widyaiswara Indonesia, 4(3), 99–114.
- Sylvia Anjani, S. K. M., Abiyasa, M. T., & PK, S. (2023). *Disrupsi digital dan masa depan rekam medis*. Selat Media.
- Tjandrawinata, R. R. (2024). *Pharma 4.0 masa depan manufaktur farmasi*. Jurnal Farmasi Indonesia, 12(2), 123–134.
- Tsang, L., Kracov, D. A., Mulryne, J., Strom, L., Perkins, N., Dickinson, R., ... & Jones, B. (2017). *The impact of artificial intelligence on medical innovation in the EU and US*. Intellect Prop Technol Law J, 29(8), 3–12.
- Wang, C., Zhang, J., Lassi, N., & Zhang, X. (2022). *Privacy protection in using AI for healthcare: Chinese regulation in comparative perspective*. Healthcare, 10(10), 1878.
- Wattimena, F. Y., Renyaan, A. S., Koibur, R., Manurung, H. E., & Koibur, M. E. (2024). *Inovasi digital dalam pemerintahan: AI, IoT, dan Blockchain*. Kaizen Media Publishing.



Journal of Public Health Indonesian

Volume.2 Issue.4, (November, 2025) Pages 68-69

E-ISSN: 3048-1139

DOI : <https://doi.org/10.62872/312p1426>

<https://nawalaeducation.com/index.php/JHH>

- Widyasari, E., Murdiyasa, B., & Supriyanto, E. (2024). *Revolusi pendidikan dengan artificial intelligence: Peluang dan tantangan*. Jurnal Ilmiah Edukatif, 10(2), 302–311.
- Zhang, Y., Wu, M., Tian, G. Y., Zhang, G., & Lu, J. (2021). *Ethics and privacy of artificial intelligence: Understandings from bibliometrics*. Knowledge-Based Systems, 222, 106994.