

Ethical Disputes and Responsible use of Artificial Intelligence in Healthcare Systems

Anubhav Agrawal¹, Akshat Gupta², Akash Yadav³, Aman Jain⁴, Alok Maurya⁵ and Kajal Dubey⁶

¹⁻⁶Department of Information Technology JSS Academy of Technical Education Noida, India

Email: agrawalanubhav283@gmail.com, akshatgupta110203@gmail.com, ay293044@gmail.com, amanstpaul16@gmail.com, alokmaurya20052003@gmail.com, kajal.dubey@jssaten.ac.in

Abstract— Artificial Intelligence (AI) has become an important technology in advanced healthcare, assisting in clinical decision-making, patient monitoring and constantly improving operational efficiency. However, the rapid integration of AI systems also raises significant concerns related data protection, data integrity, and ethical responsibility. This study presents a comparative analytical framework for evaluating international ethical guidelines. A qualitative research approach is adopted, incorporating policy analysis, comparison of regulatory framework, and assessment of key ethical parameters influencing responsible AI adoption. The analysis notices a gap between widely accepted ethical principles and their effective implementation in real healthcare systems. The study also points on the fact that how a lack of human-based contextual data can distort the ability for decision making. So by combining ethical, social, and technical perspective, this research highlights the need for transparent governance structure in AI-driven healthcare applications.

Index Terms— Artificial Intelligence Ethics, Clinical Decision Support Systems, Healthcare Data Privacy, Responsible AI.

I. INTRODUCTION

For a long period, we have tried to understand the corresponding correlation between the Computer science and Psychological prospects through Artificial Intelligence (AI). Through our society, which is greatly bound to socio-political and legal systems as well as moral lenses, it is actually important to study the impacts and effects it can have on AI. Because in the end, AI is just a computational model/system that is trying to simulate human cognitive function [2].

What we need to understand is that AI is neither something to be scared of nor something that is magical that came out of this world. There are tons and tons of data and complex maths and algorithms which bounds the reach of A.I. In the end, Humans have to take sole responsibility for AI and its after-effects. Machine learning, which is a subset of AI, is intentionally designed based on the data selection, model choice, repetitive training, and deployment, which are all ethical decisions, and at each and every stage of AI/ML, we have to understand that there must be some sort of accountability that not only ensures the positives of AI but also transcends the progress in much higher level. Even though the AI produced in today's world is advanced enough to solve complex problems and reason, it is still not conscious and morally correct. We cannot enforce moral responsibility and decency on AI and turn ourselves blind. Since we cannot provide the noiseless or the cleanest of the data to train the models, we can ask people to audit and challenge the automated decisions [11].

That is why a need for trustworthy AI is needed to ensure no harm to our society, and to make sure of it, we need to focus on three components: being lawful, ethical, and robust. But there is still a limit to it, one more thing that it should be is being spiritual. Because laws and morals are still society and region-driven, it doesn't guarantee a complete ethical ground for AI [11].

II. BACKGROUND OF ARTIFICIAL INTELLIGENCE

Murphy et al. argued that biased clinical data produces biased clinical outcomes, as ML in healthcare is only inclined to the data it learns from. The foundation for any ethical deployment of AI/ML in the medical domain should be privacy and data protection [15].

Bostrom and Yudkowsky addressed that the creation, de-ployment and termination of AI as a form of moral weight which is equivalent to animal welfare and human rights [9]. Artificial Intelligence is made to replicate human beings in all forms, whether it's intelligence, reasoning, maths, or data accumulation and analysis. What makes it easier is that it can store and process much more amount of data better than any average human being. Since it was made by humans through software and processors, it is called Artificial Intelligence. Similarly, Machine learning, which is a subgroup of AI depend on complex data for training, also known as Big Data [4].

One of the major developments which happen is the in-troduction of AI-based models in healthcare, this allowed experts to not only explore the several possibilities of AI helping humankind but also serving as a great mind for future perspective. This also proved the urge to shift from instruction-based algorithmic systems to learn and adapting systems [5].

TABLE I: DATA BREACH IMPACT STATISTICS FOR HEALTHCARE (2023)

Metric	Description
Average Total Cost of Breach	USD 10.993 million per incident (highest of all sectors).
Data Records Stolen (Approximate)	12 million patient records (specific major incident example like HCA or Change Healthcare).
Key Security Failure (Root Cause)	Phishing or stolen credentials in 67% of cases.
Key Security Failure (Root Cause)	Phishing or stolen credentials in 67% of cases.
Total Data Breach Incidents	Over 540 reported incidents (based on HIPAA data).
Patient Trust Impact	Over 65% of surveyed patients report reduced trust.

Table I summarizes the primary advantages of integrating blockchain with generative AI systems.

AI has the potential to bring revolution in the healthcare sector by working together with the doctors and the medical experts, by being the backbone of the medical field, can generate outcomes, examine X-rays hold onto the progress and can certainly be a new artificial intelligence-based doctor [5].

This does not mean replacement of doctors and experts, but it is there to act as a second brain with enormous levels of knowledge and data put together to save a human life and evolve itself based on from the data at the same time [5].

III. ETHICAL ISSUES

Cavoukian suggested proactive prevention, privacy as the default setting, privacy-based design, complete functionality, end-to-end security, visibility and transparency, and protection of user privacy [17]. In the healthcare and medical field, we are stuck in a moral dilemma which are, one being too pessimistic and precautionary, while the other is pushing something until it gets faulty. So there is a need for our data to be more complex and not just another performance overview [20].

Being too pessimistic often hinders our ultimate goal of knowing the limits of a model, but it also makes us realise that it should not be triggered in a way that it can cause harm to people or loss of money. On the other hand, people want to exploit the very model so much that it can cause some serious issues in terms of daily lives [20].

That is why we need something that can supervise the model and remind it of its job as a healthcare expert. It can act as a performance enhancer as well as a guardian, which inhibits certain things that have the potential of going out of hand [20].

IV. LITERATURE REVIEW

Hagendroff addressed the absence of trained engineers and developers who are not systematically educated about ethi-cal issues, nor are they empowered enough to raise ethical concerns. He discussed lack of intervention in business and medical science is limited and only operates on economic logic and not on ethical ones [8]. AI is

being made with certain robustness and carefully driven limits where it is required to behave minimally where there is a risk of undesirable harm to others. This also includes AI-driven medical hardware and equipment, which have a tendency to malfunction for whatever reason, either due to poor or a lack of data training or due to a bug. Undermining patient health due to AI negligence can lead to loss of life and people’s trust in modern software [11].

TABLE II: ETHICAL CHALLENGES ASSOCIATED WITH GENERATIVE AI

Challenge	Description
Bias and Fairness	Potential for biased behavior leading to unfair or discriminatory outcomes.
Privacy	Risk of misuse or exposure of sensitive and personal data.
Accountability	Absence of clearly defined responsibility for AI-generated decisions.

Table II highlights the major ethical concerns related to generative AI systems.

There is also a vast majority of mental health problems that have been reported in this century, and there have been efforts to curb this psychological epidemic with the help of AI. Whether this happens due to a poor AI agent or irrational prompting is another case, but what it signifies is that there is a lack of more context of the user, which is required to get the overall mental model and provide better results [11].

Similarly, machine learning models also have a tendency to misdiagnose a patient due to non-optimal reasons, which could be life-threatening. Removal of design flaws at the software and hardware level is a crucial part of medical machines [11]. So, although prior studies provided a wide variety of ways to make AI more optimised and hence more ethical using different data models and algorithms, and most of that work relies on domain-specific knowledge, in contrast, this work focuses on a user-based contextual model, which is subjectively ethical rather than data-driven, trained, and predictive models. This means that a model should be able to recognise and understand the need and background of the user before providing the output. This can allow AI to get deeper into the factors of data and decide the best outcome corresponding to its user [16].

The paper proposes the need for having more contextual understanding of AI at an individual level, since the mental model and biological blueprint of every human being can be different. So, to make sure that the AI model delivers the best result, it must carry the idea and background of the user it is dealing with. Background can be composed of factors like mental health, trauma, emotional capacity, and physical behaviour. These factors ensure that AI can grasp overall familiarity with the user and hence try to treat the user in the best way possible. Ethical values also require deeper thoughts of humans, which cannot be accessed until information is collected at a large level [16].

V. METHODOLOGY

Veale and Zuiderveen Borgesius praised the risk-based approach of AI. Instead of regulating all types of AI, it assigns stricter obligations to higher-risk systems- like employment, education, law enforcement and digital infrastructure. They also criticized the inconsistencies in the extent of harm-risks. It is also seen the maximum harmonization of AI across EU could prevent other countries from adopting stronger, more protective national AI policies [19].

TABLE III: COMPARATIVE ANALYSIS OF ETHICAL AI FRAMEWORKS

Ethical Principle	WHO	EU (AI Act)	This Work
Privacy Protection	Focus on data confidentiality and informed patient consent	Alignment with GDPR and data protection regulations	Privacy-aware generative AI with data minimization practices
Accountability	Human oversight with clearly defined responsibility	Legal accountability assigned to AI providers and deployers	Stakeholder-oriented accountability mechanisms
Transparency	Emphasis on explainability and operational safety	Risk-based transparency and system traceability	Intent-level transparency and auditability of AI outputs
Equity and Inclusion	Promotion of global health equity	Non-discriminatory access and inclusive system design	Commitment to fair, inclusive, and unbiased outcomes

Mulligan et al. addressed use of the word “fairness” giving different contextual meaning , and this conflation of ideas cre-ating confusion. They proposed that the AI systems cannot be declared fair without first specifying—the contextual grounds and data and toward whom and what goal [18].

This study adopts a qualitative approach based on re-search methodology which is divided into four stages: data collection, framework analysis, comparative evaluation, and modal formulation. First, data sources including

international policy frameworks, ethical guidelines, academic literature, and regulatory documents were collected from established from organizations such as WHO, EU and peer-reviewed databases. Secondly, an approach should be applied to examine recurring ethical principles including transparency, fairness and accountability. Each framework is again reevaluated against predefined criteria such model contexts, governance support and patient-centric adaptability. Third, a gap analysis was conducted to identify limitations and drawbacks in the existing frameworks, particularly regarding contextual awareness, user-specific adaptability and responsibility. Finally, an ethical model has been proposed that combine stakeholder-centric accountability, contextual data interpretation. This model aims to bridge the gap between theoretical ethical guidelines and real-world clinical deployment

Table III presents a comparative view of major ethical AI frameworks and the proposed approach in this work.

VI. DISCUSSION

The analysis demonstrates that although AI has significantly enhanced healthcare decision-making and automation, major ethical limitations remain unaddressed in current scenarios. A large section of recommendations provided by international frameworks such as WHO, EU guidelines, and UNESCO has remained largely policy-oriented rather than implementation-driven. The current testing and evaluation indicate that most current systems emphasize performance optimization while neglecting different human factors like background checks of patients, psychological state, and other ethical variables. This creates a clear gap between theoretical compliance and real-world clinical applicability. Therefore, the proposed framework and notion contribute by integrating accountability, context-based decision-making model, and privacy protecting mechanisation. This method aligns ethical principles with functional design rather than testing them just to comply by the requirements.

VII. RESULTS

The framework evaluation revealed three major findings: 1. Compliance with the global standards—Security, accountability, and fairness were identified as universal principles across the WHO, EU, and UNESCO frameworks. 2. Implementation gap—a practical mechanism for ethical values and monitoring its progress remains limited. 3. Contextual deficiency—Lack of mechanisms and efforts to embed human-driven contextual variables based on different backgrounds, which are essential for accurate and ethical decision-making.

The suggested model addresses these limitations by introducing a context-based framework that fills the gap between ethics with system level execution.

VIII. CONCLUSION AND FUTURE WORK

This study described the ethical explorations of artificial intelligence in healthcare systems and identified a critical gap between established ethical principles and their practical usage. Although global bodies have issued certain guide-lines, current systems still cannot comply with the context, accountability, security, and decision-making. The proposed framework demonstrates that combining systems like context-based data, integrity, and system design can significantly improve ethical compliance and trust in AI-driven healthcare solutions. Future research should focus on real-world data and its implications. Data verification and integrity will be used to test ethical frameworks. Introduction to superintelligence and advanced LLMs can enhance the collaboration between technologists, healthcare professionals, and policymakers.

REFERENCES

- [1] J. Bryson, *The Oxford Handbook of Ethics of Artificial Intelligence*. Oxford, UK: Oxford University Press, 2020.
- [2] World Health Organization, *Ethics and Governance of Artificial Intelligence for Health*. Geneva, Switzerland: WHO Press, 2021.
- [3] European Commission, *Ethics Guidelines for Trustworthy AI*. Brussels, Belgium: European Union, 2019.
- [4] S. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 4th ed. Upper Saddle River, NJ, USA: Pearson, 2021.
- [5] L. Floridi, *The Ethics of Information*. Oxford, UK: Oxford University Press, 2013.
- [6] L. Floridi et al., "AI4People—An ethical framework for a good AI society," *Minds and Machines*, vol. 28, no. 4, pp. 689–707, 2018.
- [7] B. Mittelstadt et al., "The ethics of algorithms: Mapping the debate," *Big Data & Society*, vol. 3, no. 2, pp. 1–21, 2016.
- [8] T. Hagendorff, "The ethics of AI ethics: An evaluation of guidelines," *Minds and Machines*, vol. 30, pp. 99–120, 2020.

- [9] N. Bostrom and E. Yudkowsky, "The ethics of artificial intelligence," in *The Cambridge Handbook of Artificial Intelligence*. Cambridge, UK: Cambridge University Press, 2014, pp. 316–334.
- [10] V. Dignum, *Responsible Artificial Intelligence*. Cham, Switzerland: Springer, 2019.
- [11] A. Jobin, M. Ienca, and E. Vayena, "The global landscape of AI ethics guidelines," *Nature Machine Intelligence*, vol. 1, pp. 389–399, 2019.
- [12] R. Benjamins et al., "Toward a human-centered AI," *AI Magazine*, vol. 41, no. 2, pp. 62–75, 2020.
- [13] I. G. Cohen et al., *The Future of Health Care Privacy*. Cambridge, MA, USA: MIT Press, 2020.
- [14] U. Scherer, "Regulating artificial intelligence systems," *Harvard Journal of Law & Technology*, vol. 29, no. 2, pp. 353–400, 2016.
- [15] K. Murphy et al., "Machine learning for healthcare," *IEEE Journal of Biomedical and Health Informatics*, vol. 24, no. 1, pp. 1–3, 2020.
- [16] P. Voigt and A. von dem Bussche, *The EU General Data Protection Regulation (GDPR): A Practical Guide*. Cham, Switzerland: Springer, 2017.
- [17] A. Cavoukian, "Privacy by design," *IEEE Technology and Society Magazine*, vol. 31, no. 4, pp. 18–21, 2012.
- [18] J. K. Mulligan et al., "This thing called fairness," in *Proc. ACM Conf. Fairness, Accountability, and Transparency (FAT)*, pp. 77–87, 2019.
- [19] M. Veale and F. Zuiderveen Borgesius, "Demystifying the draft EU AI Act," *Computer Law Review International*, vol. 22, no. 4, pp. 97–112, 2021.
- [20] UNESCO, "Recommendation on the Ethics of Artificial Intelligence," Paris, France: UNESCO, 2022.