

Ethical Misuse of Ai in Decision-Making Processes in Healthcare

Imamalieva Diyora Imamali qizi

Lecturer of the Department of International Private Law Tashkent State University of Law, E-mail: diyoraimamalieva@gmail.com

Annotation: Artificial intelligence (AI) is increasingly deployed in healthcare decision-making – from diagnostics and treatment recommendations to resource allocation and insurance claims – offering significant benefits in efficiency and accuracy. However, the misuse or uncritical reliance on AI systems poses serious legal and ethical risks. This article provides a comprehensive legal-academic analysis of these risks, anchored in international human rights principles, medical ethics, and emerging regulatory frameworks. We employ doctrinal and comparative methods to examine how different jurisdictions (notably the European Union, United States, and Singapore) are responding to challenges such as bias and discrimination in AI algorithms, threats to patient privacy and autonomy, and the problem of assigning responsibility for automated or semi-automated decisions. A focused case study on Uzbekistan’s developing approach to AI in the medical sector is included, highlighting existing norms, initiatives, and regulatory challenges in that jurisdiction. The discussion synthesizes international guidelines (from bodies like UNESCO, WHO, OECD, and the Council of Europe) and national laws, and evaluates procedural safeguards needed to ensure that AI deployment in healthcare remains aligned with ethical duties and legal accountability. We conclude by emphasizing the need for robust, harmonized legal frameworks and oversight mechanisms that protect patients’ rights and safety while fostering innovation in medical AI.

Keywords: Artificial Intelligence (AI), ethical misuse, healthcare decision-making, patient autonomy, algorithmic bias, transparency and explainability, human oversight, regulatory frameworks, liability and accountability, WHO ethics guidelines, UNESCO AI ethics principles, medical malpractice.



This is an open-access article under the [CC-BY 4.0](https://creativecommons.org/licenses/by/4.0/) license

Introduction

Artificial intelligence technologies have rapidly permeated healthcare systems worldwide, promising improvements in patient outcomes, clinical efficiency, and cost savings. From AI-driven diagnostic tools that detect diseases earlier to predictive algorithms that optimize hospital workflows, the potential benefits are profound. Leading countries are already integrating AI into daily medical practice to enhance decision-making [1]. In parallel, however, there is growing

recognition that misuse or unregulated deployment of AI in healthcare can lead to serious ethical and legal challenges. Improper or biased algorithms may produce discriminatory outcomes, undermine patient privacy, or even cause direct harm through flawed treatment recommendations. For instance, an AI system might disproportionately under-treat certain populations due to biased training data, or a predictive model might wrongfully deny insurance coverage to vulnerable patients [OBJ] [OBJ]. These scenarios underscore that without adequate oversight, AI's vaunted accuracy and efficiency can translate into infringements of fundamental rights and erosion of trust in healthcare.

Against this backdrop, policymakers and scholars worldwide are grappling with how to ensure ethical use of AI in healthcare decision-making. At the heart of the issue lie timeless principles of medical ethics – such as beneficence (“do no harm”), patient autonomy, justice, and confidentiality – now tested by algorithms capable of independent or opaque decision-making. International human rights law adds another crucial lens: the right to privacy, the right to equality and non-discrimination, and the right to health are all implicated when healthcare AI is misused. Unchecked AI systems could, for example, violate privacy through unwarranted data processing or profiling, or contravene equality rights if their outputs reflect and reinforce societal biases [1]. The human dignity of patients and the doctor–patient relationship itself may be at stake when decisions are driven by machines without sufficient human oversight.

The legal landscape responding to these concerns is in flux. International organizations and national governments have begun developing frameworks to govern AI in healthcare, though approaches vary. Many early guidelines take the form of soft law and ethical principles, focusing on values like transparency, accountability, and fairness in AI deployment. For example, the World Health Organization (WHO) and UNESCO have issued guidance urging that AI systems be designed and used in ways that respect patient rights, promote equity, and mitigate biases. In parallel, some jurisdictions are adapting or proposing binding laws: the European Union's proposed AI Act will impose strict requirements on high-risk AI (including most medical AI) to protect safety and fundamental rights, while U.S. regulators like the FDA have been overseeing AI-based medical devices under existing health and safety laws. Other countries, such as Singapore, rely on a combination of sectoral regulations and detailed guidance to ensure responsible AI innovation without stifling progress. Uzbekistan, for its part, is only beginning to formulate legal strategies for AI in healthcare, illustrating the difficulties faced by emerging economies in balancing innovation with ethical safeguards [OBJ] [OBJ].

This article aims to provide a comprehensive analysis of the ethical misuse of AI in healthcare decision-making processes and the attendant legal implications. We examine how international frameworks and human rights principles can guide the governance of medical AI, analyze comparative approaches in key jurisdictions (Singapore), and discuss the current state and future directions of Uzbekistan's regulatory response. The analysis is grounded in academic legal methodology, combining doctrinal analysis of laws and ethical codes with comparative evaluation of different legal systems. By surveying treaties, guidelines, case law, and scholarly commentary, we seek to identify common challenges and best practices for ensuring that AI serves as a tool for improving healthcare without compromising legal and ethical standards. Ultimately, the goal is to highlight how law and policy can address the “responsibility gap” – ensuring that when AI systems assist or make medical decisions, there are clear lines of accountability and robust protections for patients' rights. In doing so, we confront questions such as: Who is liable when an AI recommendation leads to a misdiagnosis or injury? How can we preserve informed consent and transparency when decisions emerge from an algorithmic “black box”? And what regulatory and procedural safeguards are necessary for the safe and equitable integration of AI into healthcare? The following sections tackle these questions, beginning with an overview of our methodological approach and then delving into the substantive findings on law, ethics, and policy.

Methods and Materials

This research adopts an academic legal methodology combining doctrinal, comparative, and formal legal analysis. The doctrinal approach involves a close examination of legal sources – including legislation, regulations, international treaties, case law, and authoritative guidelines – that govern or inform the use of AI in healthcare. We analyzed primary legal texts such as data protection laws, medical device regulations, health care statutes, and emerging AI-specific legislation (e.g., the EU Artificial Intelligence Act) to ascertain the current legal requirements and standards applicable to AI-driven decision-making in medicine. In addition, we reviewed medical ethics codes and human rights instruments (for example, the World Medical Association’s ethical guidelines, the UNESCO Recommendation on the Ethics of AI, and the Council of Europe’s Oviedo Convention on human rights in biomedicine) to identify foundational principles relevant to AI misuse in healthcare.

Comparative analysis was undertaken by selecting a sample of jurisdictions that represent a spectrum of regulatory approaches and contexts. This article focuses on the Singapore as illustrative case studies for its innovative use of guidance and sandboxes in a smaller, tech-forward jurisdiction. Each jurisdiction’s laws, policy documents, and relevant scholarly commentary were surveyed to highlight how they address key issues such as algorithmic bias, liability for AI-caused harm, and data governance in health. A special section is devoted to Uzbekistan, reflecting the authors’ regional interest, to assess how a transitioning legal system is beginning to tackle AI in healthcare and to identify unique challenges it faces (such as resource constraints and alignment with international standards). Here, materials reviewed include national strategic documents (e.g., the “Digital Uzbekistan 2030” strategy), draft legislation, and commentary by local legal experts.

The research also relies on secondary sources for context and critique. We reviewed a wide range of academic literature in law and ethics, reports by international organizations (OECD, WHO, Council of Europe, among others), and analyses by professional bodies and NGOs. These sources provide insight into normative debates and proposals – for example, calls for global governance of AI in health, or suggestions for new liability models (such as no-fault compensation schemes for AI-related injuries). All sources used are cited in accordance with academic standards. Given the evolving nature of this field, we prioritized recent publications (2019–2025) to capture the latest regulatory developments and ethical discussions.

Results

At the international level, a patchwork of soft-law guidelines and principles has emerged to address AI ethics in healthcare. While not legally binding, these frameworks carry significant normative weight and often inform national policies. A unifying theme is the insistence that AI be human-centric – i.e., that it serves human dignity, rights, and well-being, rather than undermining them. For example, the OECD Principles on Artificial Intelligence (2019), endorsed by dozens of countries, promote the innovative and trustworthy use of AI in a manner that respects human rights and democratic values. These principles enumerate values-based guidelines such as fairness, transparency, and accountability. Notably, the OECD calls for mechanisms to ensure human oversight of AI systems and to guard against misuse, whether intentional or inadvertent [3]. This implies that even at a high level, international consensus holds that AI should not be given free rein in sensitive fields like healthcare: human agency and the rule of law must frame its operation.

Another milestone is UNESCO’s Recommendation on the Ethics of Artificial Intelligence (2021), the first global normative instrument on AI ethics adopted by all 193 UNESCO member states. The Recommendation explicitly makes the “protection of human rights and dignity” the cornerstone of AI governance, urging that fundamental principles such as transparency, fairness,

and accountability be incorporated throughout an AI system's lifecycle . It emphasizes the importance of human oversight, reflecting a widespread concern that decisions affecting human lives – for instance, a critical diagnosis or a treatment allocation – should not be left to algorithms alone . Although UNESCO's text is not a treaty, its universal endorsement gives it moral authority. Member states, including Uzbekistan, have pledged to uphold its guidance, which includes assessing the impact of AI on human rights, ensuring diversity and inclusivity in AI development, and safeguarding data privacy.

The World Health Organization (WHO) has similarly taken a leading role by examining AI in the context of public health and medical care. In 2021, the WHO released a comprehensive guidance report on the Ethics and Governance of AI for Health, after extensive consultation with experts in ethics, law, and health policy . The WHO identified a set of six consensus principles for the ethical use of AI in health. These principles include: (1) Protecting human autonomy – e.g., AI should not override the clinician's judgment or the patient's informed choices; (2) Promoting human well-being and safety (the AI must be tested for safety and its benefits should outweigh risks); (3) Ensuring transparency and explainability; (4) Fostering responsibility and accountability; (5) Ensuring inclusiveness and equity (so that AI benefits all segments of society and does not discriminate); and (6) Promoting responsive and sustainable AI (meaning continuous evaluation and ability to be overridden or improved) [2]. In line with these principles, the WHO stresses that AI for health should be designed and used in ways that respect human dignity, fundamental rights and values, and that such systems must promote equity, fairness, inclusiveness, and accountability in healthcare . These broad ethical imperatives directly target the risks of misuse: for example, the call for inclusiveness and fairness addresses the danger of AI perpetuating health disparities, and the requirement of accountability ensures there are answerable parties if AI causes harm.

While endorsing AI's potential, the WHO also flags significant gaps and challenges in the current landscape. Notably, it observes a worrying lack of harmonization and coordination between countries in regulating AI for health, with few universally accepted standards beyond data privacy in some regions . This regulatory patchwork can allow unethical practices to slip through and makes it difficult for authorities to keep pace with rapid AI innovation . The WHO's report calls for capacity-building and international collaboration to develop a global governance framework for AI in healthcare, suggesting that without concerted action, disparate national rules may leave loopholes or conflict with each other . Indeed, WHO experts have even proposed exploring legally binding international rules on medical AI, potentially through instruments like an updated International Health Regulations, to ensure global alignment . In practical terms, the WHO recommends concrete safeguards, including rigorous validation of AI tools before deployment, mechanisms for algorithmic auditing to detect bias or errors, and even the creation of no-fault compensation systems for patients harmed by AI – a model that would compensate victims without needing to prove malpractice, thereby encouraging reporting of AI-related adverse events

Regional human rights bodies have also weighed in. The Council of Europe, guardian of the European Convention on Human Rights (ECHR) and other bioethical treaties, has recognized that AI in healthcare implicates core rights. Under the auspices of the Oviedo Convention (1997) on human rights and biomedicine, the Council's Bioethics Committee (CDBIO) issued a 2023 report focusing on how AI affects the patient–doctor relationship and related rights . The report zeroes in on principles of *patient autonomy and consent*, *professional responsibility*, *privacy and control over personal health data*, and *equitable access to healthcare* . It warns that AI tools, if misused, could erode patient autonomy – for instance, if patients are not informed an AI is involved in their care or if they cannot obtain an explanation for an AI-driven decision. Likewise, the duty of care and professional standards expected of physicians may be jeopardized if doctors over-rely on AI recommendations without critical scrutiny, potentially conflicting with their ethical and legal obligations to patients . The Council of Europe emphasizes that privacy and data protection (as

protected by instruments like Convention 108+ and ECHR Article 8) must be reinforced in the AI era, since health data fuels many AI systems. Patients should retain a degree of *self-determination over their health data*, meaning robust informed consent and data governance are necessary when data is used to train or run AI . The principle of equitable access is also highlighted: AI should not just be a high-tech tool available in wealthy hospitals or to well-off patients; rather, its use should enhance, not diminish, fairness in healthcare access. These points from the Council of Europe echo globally recognized human rights and ethics standards, translating them to the AI context to guide lawmakers. Indeed, the Council of Europe is in the process of negotiating an international convention on AI and human rights, which – once adopted – could impose binding obligations on member states to regulate AI consistently with rights like privacy, non-discrimination, and due process. Such a convention would almost certainly cover medical AI as a domain requiring heightened safeguards .

Discussion

The **misuse of AI in healthcare** can manifest in various ways, often implicating fundamental ethical principles and patients' rights. This may identify: *bias and discrimination, privacy and data protection, transparency and explainability, and impact on patient autonomy and consent*. Each of these issues illustrates how an AI system, if not properly designed and governed, may run afoul of medical ethics and human rights, thereby requiring legal redress or preventive regulation.

A well-documented concern is that AI algorithms can perpetuate or even exacerbate biases present in training data, leading to unjust discrimination in healthcare delivery . Unlike a human doctor, who is bound by professional ethics and anti-discrimination laws, an AI system might inadvertently offer different quality of care to different demographic groups if those groups were not adequately represented or were misrepresented in the data used to create the model. One striking example comes from a study in the United States where a machine-learning algorithm used to allocate extra healthcare support to patients was found to systematically underestimate the risk scores of Black patients compared to equally sick white patients . The algorithm in question used healthcare spending as a proxy for health needs – assuming that patients with higher past expenditures were sicker – but because historically less money was spent on Black patients (due to systemic inequalities and access barriers), the AI concluded they were “healthier” than they truly were. As a result, many Black patients did not get flagged for additional care programs, effectively *denying them beneficial interventions* on a discriminatory basis . This is a paradigmatic case of AI misuse: a seemingly neutral tool producing a racially biased outcome, clashing with the ethical principle of justice and the legal right to equal treatment. From a human rights perspective, such outcomes could violate anti-discrimination laws and equality guarantees in constitutions or human rights treaties. They also raise questions of due process and transparency – the patients were unlikely to know an algorithm was triaging their care in this manner, let alone challenge its decision.

Bias is not limited to race. AI systems could exhibit gender bias (for example, under-diagnosing heart attacks in women because training data was skewed toward male patients), age bias (perhaps allocating fewer resources to elderly patients), or biases against people with rare conditions or disabilities. A related issue is that AI tools, such as image recognition in diagnostics, may perform worse for certain populations. Research has shown that some AI diagnostic systems for dermatology or radiology underperform on darker-skinned patients because the training images were predominantly of lighter-skinned individuals . In one study, algorithms tasked with identifying skin lesions had significantly higher error rates on images of patients with dark skin, raising the risk of misdiagnosis or delayed diagnosis for those patients . This undermines the ethical duty of non-maleficence (do no harm) and again implicates equality rights. Medical malpractice law could also come into play if a biased AI contributes to a misdiagnosis – the providers and developers could potentially face liability for resulting harm. Consequently,

addressing bias is a top priority in all ethical AI frameworks: measures such as diverse training data, bias testing/auditing of algorithms, and ongoing monitoring are recommended to ensure AI does not become a high-tech conduit for old prejudices .

AI in healthcare is fueled by data – often **sensitive personal health data**. This creates tension with privacy rights and data protection laws. The misuse of AI can occur if systems are developed or deployed without proper safeguards for patient data. For example, large datasets of electronic health records may be used to train an AI diagnostic model. If those records are shared with AI developers without patient consent or an appropriate legal basis, it may constitute a breach of privacy laws (such as HIPAA in the US or GDPR in Europe). A notorious case illustrating these concerns was the partnership between London’s Royal Free Hospital and Google DeepMind. In 2016, the hospital trust transferred 1.6 million patients’ records to Google DeepMind to develop an app for detecting kidney injury. The UK Information Commissioner’s Office later found this arrangement unlawful, because patients were not adequately informed and the data use exceeded the hospital’s authority . Essentially, patients who visited the hospital for routine treatment had no reasonable expectation their data would be siphoned to an AI project, even if the goal (improving diagnostics) was well-intentioned . The case underscores that data protection is not a mere technicality but a substantive right: any AI deployment must abide by privacy principles like informed consent, purpose limitation (using data only for the specific legitimate purpose collected), and data minimization. In Europe, these principles are enshrined in the GDPR and backed by potential fines; elsewhere, privacy commissioners and courts are increasingly ready to intervene when health data is misused.

Beyond legal compliance, privacy is linked to patient trust. If patients fear that their medical information will be fed into opaque AI systems or potentially misused (e.g., shared with tech companies or insurers without consent), they might be less candid with doctors or avoid seeking care, undermining public health outcomes. Moreover, AI models themselves can pose privacy risks through what is known as re-identification: if an AI is trained on “anonymized” health data, there’s a chance that the model’s outputs or patterns could indirectly reveal identities or sensitive facts about individuals, especially if combined with other data. Thus, an ethical deployment of AI requires strong technical safeguards (encryption, de-identification techniques, access controls) and legal safeguards (compliance with data protection laws, oversight by ethics boards or data protection officers). International standards like the OECD’s and UNESCO’s recommendations emphasize privacy and data governance as key components of trustworthy AI [2,3]. For instance, they call for algorithms to be developed in a manner that preserves confidentiality of personal health information and for transparent policies to inform patients when their data is used. In practice, some jurisdictions are creating secure health data spaces (like the European Health Data Space initiative) to enable AI innovation in research while protecting individual rights . The misuse scenario to avoid is one where an AI might technically excel at predicting illness, but in doing so it tramples privacy – such as an AI that scrapes data from fitness apps or social media without consent to infer health conditions, or uses genetic data in ways that could stigmatize individuals or groups. Such practices could violate not only privacy statutes but also broader notions of dignity and freedom from intrusion.

A frequent criticism of AI systems, especially those based on machine learning, is their “black box” nature – they may produce a result (e.g. a diagnosis or risk score) without offering a human-understandable rationale. In healthcare, lack of transparency or explainability is ethically problematic. Patients have a right to understand the basis of medical decisions affecting them, and healthcare providers have a duty to base care on understandable reasoning. When AI is misused in the sense of being applied without proper explainability, it can undermine informed consent and trust. For example, if an AI recommends a certain treatment but the doctor cannot explain to the patient why the AI chose that path (and the doctor herself doesn’t know, due to the AI’s complexity), how can the patient make an informed choice or feel confident in the

recommendation? This conflicts with the ethical principle of autonomy and legal doctrines requiring informed consent to treatment. It also potentially conflicts with due process rights in contexts like health insurance or public healthcare entitlements: if a government algorithm decides to deny someone a medical benefit, fairness may require an explanation or an opportunity to contest the decision. The Council of Europe's guidance and the EU's emerging regulations both stress the need for explainability in AI, especially in critical sectors like health . The EU AI Act, for instance, will require that users of high-risk AI systems are provided with information on how the system makes decisions and that such systems are subject to human oversight precisely to catch errors or unjust outcomes .

Misuse related to transparency can also occur when stakeholders deliberately or negligently hide the use of AI. If a hospital deploys an AI tool but presents its outputs to patients as if they were solely the doctor's judgment, patients are misled about their care. Ethically, patients should be informed when AI is involved in significant ways, a point echoed in guidelines like Singapore's AI in Healthcare Guidelines which call for patients (and providers) to know they are interacting with an AI or AI-assisted service . From a legal standpoint, jurisdictions like the EU, under GDPR Article 22, recognize a right not to be subject to purely automated decisions with legal or similarly significant effects without certain protections, including the right to an explanation of the decision. While medical treatment decisions are often not "fully automated" because a human clinician is in the loop, the spirit of such laws pushes towards transparency and the ability to contest algorithmic decisions. In short, using AI in a non-transparent way – whether by deploying inscrutable models or failing to disclose AI's role – is a misuse that can breach ethical obligations of candour and legal rights to information. To counteract this, several measures are advocated: using AI models that are inherently more interpretable for critical tasks, providing clear documentation and reason codes for AI outputs, training clinicians to interpret and communicate AI findings, and establishing procedures for second opinions or overrides when an AI's recommendation is questionable. The AMA (American Medical Association) has underscored that AI is to be viewed as "*augmented intelligence*" – a support tool rather than a replacement for human reasoning – precisely because doctors must be able to explain and stand by the decisions made . In practice, this means physicians should not blindly follow an AI but should use it as an aid, ensuring that the final decisions remain comprehensible and justified in medical terms. Failure to maintain this standard veers into unethical use of AI and could become fodder for malpractice claims or regulatory action if patients are harmed by unexplainable AI errors [10].

A critical legal dimension of AI in healthcare is determining who bears responsibility when something goes wrong. If an AI system's misuse or error leads to patient harm – for example, a misdiagnosis, inappropriate treatment, or denial of necessary care – the law must ascertain where the accountability lies. This is challenging because AI often involves multiple actors (developers, vendors, healthcare providers, and even the AI model itself which lacks legal personality) and can behave in ways not easily predictable by any one actor. The term "responsibility gap" is frequently used to describe the uncertainty over accountability for autonomous systems. Different legal theories and emerging norms are being explored to close this gap, ensuring that patients are not left without remedy and that there are incentives for safe deployment of AI.

Traditionally, healthcare injuries are addressed through medical malpractice law (holding physicians or hospitals liable for negligence) or through product liability law (holding manufacturers liable for defective products). In cases of AI misuse, both could potentially apply, but not without complications. For instance, if a doctor relies on an AI diagnostic tool that turns out to be wrong, and the patient is injured by a delayed or incorrect treatment, the patient might sue the doctor for malpractice. The question then becomes: was the doctor negligent to trust the AI? Malpractice is judged by the standard of a reasonably competent professional. As AI becomes common, a competent doctor might be expected to use AI as an aid (so not using AI could be negligence in some future scenarios), yet also expected to double-check AI outputs. If the AI was

obviously off-base and the doctor failed to notice, liability clearly remains with the doctor. However, if the AI's error was subtle or the doctor had no reasonable way to know the AI was flawed, should the doctor still be liable? Many clinicians worry about being held liable for decisions influenced by AI when the inner workings are beyond their understanding. Some have even called for safe harbor laws – changes to malpractice rules to protect doctors who use certified AI tools in good faith. In the US, this discussion is reviving debates on malpractice reform. As of now, no special immunity exists; thus, healthcare providers do remain on the hook if they mis-use AI (for example, blindly following an AI recommendation that a prudent doctor would question). This provides a strong incentive for clinicians to treat AI outputs as suggestions, not gospel.

A key factor is the level of human involvement. If an AI is *fully autonomous* – imagine a scenario in the future where an AI-driven robotic surgeon performs procedures start-to-finish – the law may lean towards treating the AI system (and by extension its creator or operator) akin to a human actor for liability purposes. Scholars have posited that in such cases, traditional negligence might not fit well, because there's no human conduct to evaluate against a standard of care. Instead, liability might shift more squarely to the developers under a product-liability or even a strict liability regime for AI-caused harms. Strict liability would mean the injured patient doesn't have to prove the manufacturer was negligent, only that the AI product caused harm while being used as intended. This is attractive for victim compensation but raises the stakes for AI developers, potentially chilling innovation. A middle-ground solution discussed in policy circles and by the WHO is the idea of compensation funds or insurance. The WHO has suggested exploring no-fault compensation schemes for AI injuries, similar to vaccine injury compensation programs, which could ensure patients are compensated swiftly without lengthy litigation, while cost is spread (via insurance or state funds) rather than falling solely on individual doctors or companies. This approach recognizes that with complex AI, pinning fault on one actor might be less important than making the patient whole and improving system safety.

Different jurisdictions are handling AI liability in healthcare in distinct ways, which we detail in the comparative section. However, some emerging common themes include: (a) the use of regulatory oversight to assure AI quality before it reaches patients (e.g., requiring clinical trials or evidence of safety for AI tools, similar to drug approvals), (b) the push for transparency and documentation (so that when things go wrong, it's easier to trace why and how – for instance, logging the AI's decision process and the human's involvement), and (c) potential updates to legal definitions (such as expanding “medical malpractice” to explicitly cover algorithm use or refining what counts as a “defective” algorithm in product law). There are also calls for establishing clear standards of care for AI usage. For example, if most competent oncologists use a particular AI tool to guide chemotherapy decisions in 2030, then not using it might be a breach of standard care – but also, using it without following the proper validation protocol could be a breach. Professional bodies and accrediting organizations will likely develop guidelines (they already have started) on the proper use of AI, which courts may later treat as authoritative on what a careful practitioner would do.

In sum, while the law has not fully settled on a single model for AI liability in healthcare, the clear trajectory is that there will always be a human or entity accountable for patient harm, even if the immediate cause was an AI's decision. The maxim “AI may be autonomous, but it's not independent of human responsibility” is guiding regulators. Whether through adapting old laws (malpractice, product liability) or crafting new ones (strict liability for AI, insurance requirements, no-fault funds), legal systems are striving to ensure that patients injured by AI do not fall through the cracks and that those developing and using AI have proper incentives to prioritize safety and ethics. Negligent or reckless misuse of AI – such as deploying an unvetted algorithm or ignoring an AI's known limitations – will likely be judged harshly under these emerging standards, just as a breach of duty in more traditional medical practice would be.

Jurisdictions around the world are beginning to develop legal and regulatory responses to the ethical challenges of AI in healthcare, each reflecting different legal traditions and policy priorities. We will make an attempt to shortly explore how the Singapore is addressing the risks of AI misuse in medical decision-making and moving next to the Uzbekistan's experience.

Singapore offers an interesting example of a jurisdiction that actively seeks to harness AI in healthcare while implementing a governance framework that blends regulation, guidelines, and self-regulation. As a smaller common-law jurisdiction with a strong tech focus, Singapore's approach is often cited as agile and innovation-friendly yet conscious of ethical duties. Rather than heavy-handed legislation, Singapore has developed detailed guidelines and sectoral regulations to manage AI use in the medical field, supported by its existing laws on data and healthcare.

In 2021, Singapore's Ministry of Health (MOH), together with the Health Sciences Authority (HSA) and Integrated Health Information Systems (the national health IT agency), issued the "Artificial Intelligence in Healthcare Guidelines" (AIHGle) . This comprehensive guideline is a key pillar of Singapore's strategy. It serves as a resource for both developers of healthcare AI and healthcare providers implementing AI. The AIHGle builds upon principles of Singapore's broader Model AI Governance Framework (which addresses AI ethics in all sectors) and adapts them to healthcare specifics . The guidelines enumerate five key principles for AI in healthcare: fairness, responsibility, transparency, explainability, and patient-centricity .

- *Fairness* means AI systems should avoid creating unjustified disparate impacts among different patient groups; the guideline urges careful design to prevent discrimination (echoing the concerns about bias discussed earlier) .
- *Responsibility* places accountability on both developers and implementers (like hospitals) for the outcomes of AI. They are expected to be answerable for how the AI is designed, tested, and used .
- *Transparency* requires that end-users (not only clinicians but also patients where relevant) are made aware when AI is being used and have access to information about the AI's role .
- *Explainability* goes further by saying AI decisions should be understandable and reproducible to a level that meets user expectations – essentially, clinicians should be able to get an explanation from the AI that they can interpret and relay to patients .
- *Patient-centricity* underscores that AI deployment must prioritize patient safety and well-being, not just hospital efficiency or cost-saving .

These principles are not just abstract: the AIHGle provides practical recommendations, such as conducting bias testing, ensuring informed patient consent when AI is directly involved in care decisions, and maintaining human oversight. It also deals with delineating roles – recognizing that sometimes the developer and the implementer might be the same entity (for instance, a tech-savvy hospital building its own AI) or different (a vendor selling to a hospital). The guidelines suggest having clear agreements (SLAs) that define each party's responsibilities, especially in maintenance and accountability for errors . This is meant to preempt finger-pointing after a failure: if, say, a prediction was wrong due to outdated software, was it the vendor's job to update it or the hospital IT's job to install updates? Clarifying such issues in advance is seen as good governance.

Singapore's HSA acts similarly to the US FDA in regulating medical devices, including AI software. They have issued specific regulatory guidelines for AI as medical devices, taking a life-cycle approach . This means an AI tool for, say, diagnosing ECG readings would need to be registered, meet safety/effectiveness criteria, and be monitored after deployment. Singapore has been proactive in updating definitions so that adaptive algorithms are covered. They also use

“regulatory sandbox” concepts – allowing some innovations to be trialed in controlled environments like test beds in hospitals with regulatory guidance rather than full licensing, to gather evidence and refine rules.

Aside from device regulation, Singapore leverages existing laws in the healthcare sector. For instance, the Private Hospitals and Medical Clinics Act and its successor, the Healthcare Services Act, set standards for healthcare providers (e.g., on maintaining medical records securely – which applies to data used in AI). Professional boards like the Singapore Medical Council have an Ethical Code that applies to use of new technologies: doctors must ensure they have sufficient understanding and that patient care quality isn’t compromised. In fact, a 2020 amendment to Singapore’s Civil Law (reflecting the Montgomery ruling on informed consent) sets how doctors must inform patients – arguably, if AI is materially part of the decision process, not disclosing it could risk falling short of informed consent standards.

Singapore hasn’t yet had cases of AI malpractice, but under its law, similar principles to other common law jurisdictions apply: if a provider uses AI recklessly or an institution deploys a flawed AI causing harm, negligence law would apply. However, Singapore’s smaller size and strong government oversight may mean many issues get resolved through regulatory action or internal review rather than litigation. Interestingly, Singapore often prefers prospective risk management over courtroom battles. For example, a public consultation in 2023 by Singapore’s Bioethics Advisory Committee explicitly looked at ethical, legal, and social implications (ELSI) of AI in healthcare, indicating that regulations or advisories may be updated based on stakeholder feedback. This participatory approach aims to surface potential problems (like liability ambiguities or consent issues) and address them through policy guidance before they manifest in crises.

Singapore markets itself as a hub for health tech innovation. To avoid stifling AI, the government provides sandboxes (via the IMDA – Infocomm Media Development Authority – for general AI governance and perhaps in health through MOH) and encourages self-regulation within a provided framework. They also invest in research on AI ethics and have formed partnerships (for example, with AI researchers and companies to develop AI for medical imaging, with oversight committees ensuring ethical use). The Model AI Governance Framework (which AIHGle builds on) is even offered internationally as a blueprint, reflecting Singapore’s desire to lead in AI best practices. But critically, all this happens without a specific “AI Act” like the EU’s. Instead, it’s interwoven into domain-specific rules and guidelines that collectively ensure AI is used responsibly.

For countries like Uzbekistan, Singapore’s approach demonstrates a balanced strategy: not rushing to legislate AI in isolation, but updating and reinforcing existing health laws, issuing detailed guidance to shape ethical use, and involving industry and experts in governance. The advantage is flexibility and keeping pace with innovation, while building an accountability culture. The potential downside could be that guidelines are not legally binding – they rely on professional compliance and reputation incentives. However, in a tightly regulated health system like Singapore’s, soft laws often effectively guide behavior, and they can be codified later if needed.

Moving to Uzbekistan’s emerging approach, Uzbekistan, as a civil law country in Central Asia with a transitioning economy, is at a nascent stage in dealing with AI in healthcare. The nation has ambitious digital development goals and recognizes AI as a driver of innovation. However, specific laws or regulations on AI, especially in the medical domain, are still in development. This section outlines the current state of play – including existing legal norms that could apply to AI misuse in healthcare, recent initiatives and draft legislation, and the challenges and prospects ahead.

The “**Digital Uzbekistan-2030**” Strategy, approved by presidential decree in 2020, lays the groundwork for integrating advanced digital technologies across sectors including healthcare . While not an AI-specific law, this strategy mandates investments in ICT infrastructure, capacity building, and the adoption of modern technologies in public services (e-health included). It signals political will at the highest level to embrace AI: for example, in public statements, the government has highlighted the need to implement AI solutions in diagnostics, telemedicine, and management of healthcare resources as part of improving service delivery and efficiency. The strategy and subsequent government resolutions call for pilot projects and training programs on AI. One tangible outcome has been the introduction of AI elements in certain healthcare projects – for instance, using AI software for analyzing medical images in a few hospitals as experimental programs (often in collaboration with foreign tech partners). However, until recently, these moves were happening without a dedicated regulatory framework, essentially under general laws on healthcare and technology.

As of now, Uzbekistan does not have a comprehensive AI law, but several existing laws would govern aspects of AI use in healthcare:

- The *Law on the Protection of Citizens’ Health* (2019) provides patients’ rights to quality care and informed consent, and obligations on healthcare providers to follow standards of diagnosis and treatment. This law could implicitly cover AI: if a doctor uses AI in treating a patient, the general duty to provide care with all necessary knowledge and skill would apply. If AI misuse leads to harm, theoretically this could be a breach of that law, enforceable via malpractice suits or disciplinary actions. Notably, Uzbek law (like many post-Soviet systems) does not have as robust a tort litigation culture as the US, but patients can complain to health authorities and providers can be sanctioned or sued under civil codes for substandard care.
- The *Law on Personal Data* (2019, amended 2021) is crucial if health data is used for AI. It classifies health data as sensitive personal data. Processing such data (e.g., collecting it for an AI training set or sharing it with an AI developer) requires consent of the individual or another legal basis, and data security measures. Uzbekistan has been tightening personal data regulations recently (with data localization requirements and significant fines for violations). Misuse of patient data in AI development (like sharing medical records without permission) could violate this law. For example, if a hospital gave a tech company patient X-ray archives to train an algorithm without patient consent or anonymization, that might breach personal data protection rules. Enforcement is something Uzbekistan is ramping up as it aligns more with international privacy norms.

Recognizing the need for an AI-specific framework, Uzbekistan’s parliament took up its first AI bill in 2023-2024, which passed a first reading in May 2025 . While the full text isn’t public as of writing, summaries indicate this draft law addresses ethical principles, data protection, and liability for AI broadly (not just in health) . Notably, it mandates that AI systems be developed and used in an ethical manner to prevent harm and avoid bias or discrimination . This general requirement would directly cover healthcare AI, meaning, for instance, that a hospital or developer could be violating the law if they deploy an AI that is known to be biased against a certain group. The bill also introduces guidelines for data usage, which likely reinforce health data privacy obligations, and stipulates accountability for AI developers and users, including penalties for misuse . This could form the legal basis for holding an AI provider responsible if, say, their system generates illegal content or decisions. Importantly, the bill calls for a regulatory body to oversee AI and for promotion of AI in sectors like healthcare under responsible use conditions . The creation of an oversight agency could mean future certification or auditing of AI algorithms in sensitive fields [8].

Another development is Uzbekistan’s work on a National AI Strategy (. Aligning with OECD AI policy guidance, the strategy likely emphasizes capacity building, education, and international

cooperation on AI. For healthcare, this might translate to investing in AI tools for telemedicine in rural areas, or decision support in under-staffed clinics, but accompanied by training healthcare professionals in AI literacy.

In conclusion, Uzbekistan is in the foundational stage of addressing AI misuse in healthcare. The awareness at the top levels of government is there – they cite the importance of “ethical use of AI” and preventing harm and bias. But turning that into operational law and practice will require sustained effort. The country’s legal system will have to evolve to handle questions like liability if an AI error injures a patient (currently untested, but the principles of fault in civil law would apply), or accountability if a public hospital’s AI policy inadvertently denies someone life-saving treatment (raising constitutional rights issues perhaps). International experience, as discussed throughout this article, provides a roadmap. For instance, Uzbekistan’s lawmakers might choose to incorporate provisions similar to the EU (mandating human oversight in high-risk AI, or requiring bias audits), and to strengthen patient rights (explicitly giving patients the right to an explanation for algorithmic decisions in healthcare, for example). The challenges of resources and expertise are non-trivial, but with ongoing digital reforms and support from development partners, Uzbekistan is laying the legal groundwork to ensure that as AI enters its healthcare sector, it does so in a manner that upholds human rights, medical ethics, and public trust.

Conclusion

Artificial intelligence is reshaping decision-making in healthcare, bringing both immense opportunities and sobering risks. This comprehensive review has illustrated that ethical misuse of AI in healthcare is a real and pressing concern, one that straddles the domains of law, medicine, and ethics. From biased algorithms that inadvertently discriminate against vulnerable patient groups, to opaque “black box” systems that challenge informed consent and trust, to complex questions of liability when automated decisions lead to harm – the potential pitfalls are numerous. These are not merely theoretical problems of the future; they are emerging now in various forms across the globe, as evidenced by the cases and examples discussed.

A clear finding is that no single legal instrument or ethical guideline will suffice to address these multifaceted challenges. Instead, a layered approach is needed – combining international principles, national legislation, industry standards, and professional ethics. International frameworks provide a valuable foundation by articulating universal values of human rights, safety, and accountability that AI in healthcare must uphold. National experiences, from the EU’s rigorous regulatory regime to the US’s adaptive use of existing laws and Singapore’s innovation-friendly guidelines, offer diverse models that countries like Uzbekistan can learn from. The comparative analysis underscores that different legal cultures produce different tools: some more preventive and rule-based, others more principles-based and adaptive. There is merit in each approach, and importantly, they are not mutually exclusive. Indeed, the trend appears to be towards convergence – even jurisdictions initially hesitant to regulate are increasingly recognizing the need for at least baseline rules, and those with strict rules are learning to allow flexibility for innovation.

For Uzbekistan, situated at the beginning of its AI regulatory journey, the way forward should involve contextualizing global best practices to local realities. The analysis of Uzbekistan’s current steps reveals both the challenges and opportunities in crafting a framework for AI in its healthcare system. The existence of political will, as shown by the draft AI law focusing on ethical AI use, is a strong starting point. The task now is implementation: ensuring that laws are effectively enforced, that healthcare providers are trained in both using and questioning AI, and that patients’ rights are safeguarded. Uzbekistan’s legal system may need to adapt concepts like informed consent and professional liability to explicitly cover AI-related scenarios, to remove any doubt that, for example, a patient has the right to know about AI’s involvement or that a developer can be held to account for a faulty algorithm. Coordination between the health sector regulator

and whatever AI oversight body is created will be key, so that domain-specific expertise informs AI governance in health.

The promise of AI in healthcare – improved diagnostics, personalized treatment, streamlined operations – is too great to ignore. But as this study has shown, realizing that promise responsibly requires embedding AI within the rule of law and ethical medical practice. We must avoid the false dichotomy that it's a choice between innovation and safety; rather, smart regulation and vigilant ethics can guide innovation towards safe and equitable outcomes. Indeed, many of the legal and ethical measures discussed not only prevent harm but can improve AI performance (for instance, eliminating bias makes AI more accurate for more people; transparency can increase trust and therefore adoption of beneficial tools). In that sense, ethical AI is not an impediment to progress but a pathway to sustainable progress.

Finally, the analysis reinforces a crucial point: *the centrality of the human element*. Medicine is, at its core, a human endeavor about care and trust. AI is a powerful tool, but it remains a tool to be wielded by humans in service of human welfare. Laws and ethics surrounding AI in healthcare, therefore, seek not to reject the new and automated, but to ensure it operates in harmony with human values and judgment. When an AI system misfires, it is ultimately human responsibility – whether of a doctor, a developer, or an institution – that must answer for it. By clearly delineating those responsibilities and upholding patients' rights, we ensure that technology enhances rather than undermines the healthcare mission.

In conclusion, the misuse of AI in healthcare decision-making is a governance challenge that the international community and individual nations must meet with urgency and care. The evolving tapestry of laws, from international guidelines to national statutes, reflects an ongoing commitment to steer AI towards ethical use. As Uzbekistan and others refine their approaches, continued international dialogue and cooperation will be invaluable – sharing experiences of what works and what doesn't, learning from mistakes (preferably minor ones) before major mishaps occur. The stakes – human lives and dignity – could not be higher. But with informed, well-crafted legal frameworks and a steadfast ethical compass, we can harness AI's benefits while safeguarding against its risks, ensuring that the future of healthcare is both technologically advanced and profoundly humane.

REFERENCES:

1. Recommendation on the Ethics of Artificial Intelligence.
2. Ethics and Governance of Artificial Intelligence for Health, World Health Organization, 2021.
3. OECD Principles on AI, OECD Legal Instrument OECD/LEGAL/0449, adopted May 2019 (updated 2024).
4. Report on the Application of AI in Healthcare and its impact on the Doctor-Patient Relationship, Steering Committee for Human Rights in Biomedicine and Health, CDBIO, April 2023.
5. Artificial Intelligence in Healthcare Guidelines (AIHGle), MOH/HSA/IHIS, October 2021. Singapore Ministry of Health.
6. Health Sciences Authority (HSA) Regulatory Guidelines on AI Medical Devices, 2022.
7. Decree of the President of the Republic of Uzbekistan No. DP-6079, "On approval of the Strategy 'Digital Uzbekistan-2030' and measures for its effective implementation", 5 October 2020.
8. Kucharov, K., "Regulating AI in Uzbekistan: Balancing Innovation and Risk," Uzbekistan Law Blog, 2024 .

9. Ibaroudene, L., et al., “Shaping the future of AI in healthcare through ethics and governance,” *Humanit. Soc. Sci. Commun.* 10, 75 (2023)
10. Wong, J., “AI in healthcare: legal and ethical considerations in this new frontier,” *International Bar Association*, 2023 .
11. Payne, D., “Who pays when AI steers your doctor wrong?”, *Politico*, 24 March 2024 .
12. Hern, A., “Royal Free breached UK data law in 1.6m patient deal with Google’s DeepMind,” *The Guardian*, 3 July 2017 .