

Research Article

Algorithms with a Bedside Manner: Regulating AI's Social and Legal Impact on U.S. Healthcare

¹Abdullah Mazharuddin Khaja, ²Iftikhar Bhatti, ³Tamoor Ali Sadiq, ⁴Michidmaa Arikhad

¹Department of Computer Science, Governors State University, University Park, Illinois.

²Heller School for Social Policy and Management, Brandeis University, Waltham, MA.

³Southern Business School, University of Southern Punjab.

⁴Department of Computer Science, American National University, Louisville, Kentucky.

Received Date: 11 August 2025

Revised Date: 29 August 2025

Accepted Date: 06 September 2025

Published Date: 11 September 2025

Abstract: Artificial intelligence (AI) is rapidly entering routine clinical practice in the United States, with adoption accelerating across imaging, triage, clinical decision support, and operational management. Reported benefits include faster and more accurate diagnosis, improved workflow efficiency, and the potential to reduce healthcare costs. However, widespread implementation also introduces critical risks, including safety failures, algorithmic opacity, embedded bias, privacy breaches, and unresolved questions of liability. The current regulatory environment remains fragmented. While the Food and Drug Administration (FDA) has begun adapting device approval pathways to encompass software and machine learning applications, these frameworks are still evolving. Privacy protections under existing health law provide a foundation for safeguarding patient data, yet they are increasingly strained by large-scale data aggregation and cross-system linkages. Liability doctrines in malpractice and product law address some harms but leave significant gaps when autonomous algorithmic logic shapes medical decisions. Intellectual property policies further complicate matters by influencing transparency, disclosure, and oversight. This paper synthesizes interdisciplinary literature on clinical performance, social impact, and legal governance of AI in healthcare. It proposes an analytic method for regulatory assessment that emphasizes four core principles: safety, equity, privacy, and accountability. A consolidated risk-response matrix and supporting figures are presented to assist policymakers and health system leaders in evaluating emerging tools. The discussion recommends lifecycle validation processes, mandated bias audits, strengthened data governance protocols, clarified liability standards, and structured education for clinicians and patients. The overarching aim is the development of ethically aligned and legally compliant AI that enhances, rather than undermines, the quality of bedside care.

Keywords: Artificial Intelligence, Clinical Governance, Healthcare Regulation, Patient Safety, Algorithmic Bias.

I. INTRODUCTION

Artificial Intelligence (AI) is making rapid inroads into healthcare, promising data-driven solutions in diagnostics, treatment recommendations, and patient monitoring. AI algorithms are already transforming diverse fields from financial risk assessment [1] to judicial administration [2]. In medicine, well-trained AI systems can sift vast health datasets to enhance decision-making and reduce errors, potentially transforming the healthcare system [3]. For example, AI-driven diagnostic tools have shown success in imaging analysis. As of 2024, the U.S. Food and Drug Administration (FDA) has authorized nearly 1,000 AI-enabled medical devices, a sharp rise from only six such approvals in 2015 [4]. This explosive growth underscores AI's perceived value in improving patient care. However, alongside these benefits come serious concerns: issues of patient safety, algorithmic biases, data privacy, and accountability have triggered calls to reexamine existing legal frameworks governing healthcare AI [5]. The concept of an "algorithm with a bedside manner" captures the challenge: AI tools must not only be technically proficient but also socially responsible and legally compliant in the sensitive context of patient care [6]. This paper explores AI's social and legal impacts on U.S. healthcare and how regulation can ensure its ethical, equitable, and lawful use.

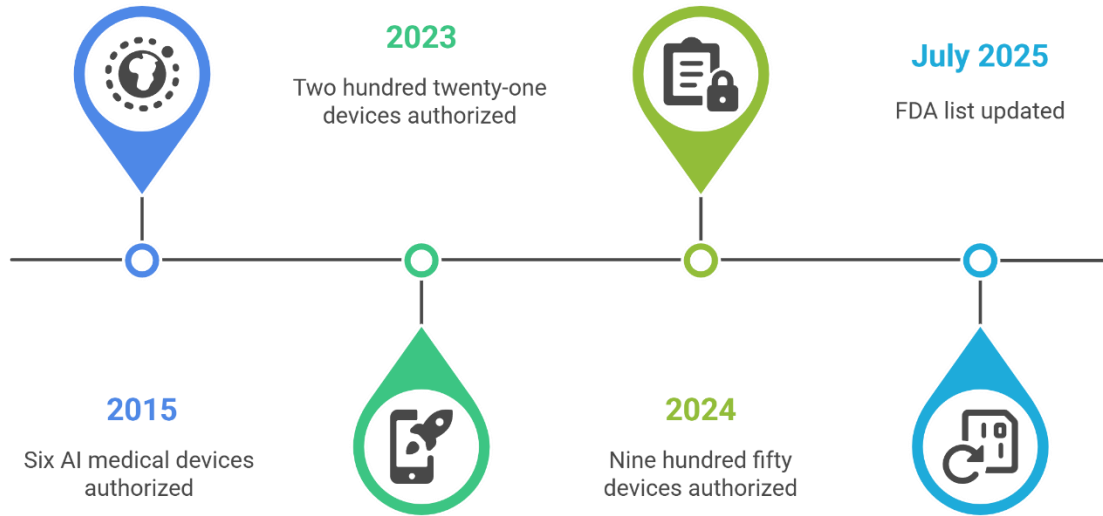
II. AI IN HEALTHCARE: SCOPE AND BENEFITS

AI applications in healthcare range from machine-learning algorithms that interpret medical images to robotic process automation in hospital workflows. In clinical diagnostics, AI can enhance the speed and accuracy of detecting conditions [7-8]. A notable milestone was in 2018, when the FDA approved IDx-DR – the first autonomous AI system for diabetic retinopathy screening that operates without specialist oversight. Such tools expand access to care by enabling early disease detection in primary care settings [9]. AI-driven software now assists in radiology, cardiology, oncology, pathology, and more. Indeed, radiology accounts for the majority of AI-enabled medical devices (around 750 of the 950 FDA-authorized AI devices by 2024), owing to the technology's strengths in image analysis and pattern recognition [10].



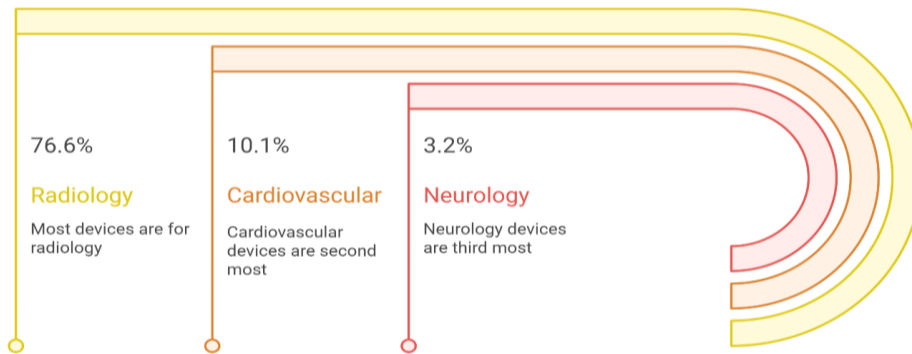
js

Figure 1: Growth in FDA-authorized AI medical devices.



Beyond diagnostics, AI-powered predictive models help identify at-risk patients (for example, those likely to be readmitted or develop complications), allowing for preventive interventions [11]. Hospitals are also employing AI to optimize scheduling, manage electronic health records, and even provide virtual nursing assistants. These innovations aim to improve outcomes and efficiency, aligning with the medical community’s “quadruple aim” of enhancing patient care, population health, provider work life, and cost-effectiveness [12]. Researchers have begun integrating AI into biomedical research as well – from drug discovery to genomics – anticipating that intelligent algorithms could uncover insights in complex biochemical data. In short, AI’s footprint in healthcare is expanding rapidly, bringing clear benefits such as faster diagnosis, personalized treatment, and streamlined operations [13]. At the same time, realizing these benefits at scale requires confronting the social and legal implications described in the following sections.

Figure 2: Distribution of AI devices by clinical specialty



III. SOCIAL IMPACT ON PATIENTS AND SOCIETY

A) *Quality of Care and Patient Outcomes*

AI systems hold the potential to reduce human errors and variation in care. They can analyze symptoms, images, and lab results with a consistency that aids clinical decision-making. Many Americans are optimistic that AI can reduce medical mistakes. In one survey, 40% believed that using AI would decrease errors made by healthcare providers, compared to only 27% who feared it would increase errors [14-17]. In specific domains like medication dosing or image analysis, AI tools have already demonstrated accuracy matching or exceeding human experts [18]. For patients, this can mean quicker diagnoses (as with AI-based stroke detection or cancer screening) and more data-informed treatment choices. Furthermore, AI can improve access to care: for example, telehealth chatbots and symptom-checkers offer basic medical guidance in underserved or remote areas, and assistive technologies driven by AI help disabled patients navigate healthcare services [19-21]. By mining large datasets, AI can also identify public health trends or at-risk populations, enabling proactive interventions at the societal level.

Despite these advantages, the net effect of AI on patient outcomes remains a point of debate. Public confidence is guarded – only 38% of U.S. adults surveyed in 2022 thought that increased AI use would lead to better health outcomes, while 33% worried it would lead to worse outcomes [22]. Clinical evidence for AI's impact is still emerging, and outcomes depend on how well algorithms are validated and used. There have been high-profile disappointments (such as an AI system that failed to improve cancer treatment recommendations), reminding stakeholders that rigorous evidence is needed before trusting AI with life-critical decisions. To maximize the quality of care, healthcare AI systems should undergo extensive clinical evaluation and continuous monitoring in real-world use [23]. Regulators like the FDA have begun requiring robust validation for AI-based devices, including evidence that they improve diagnostic accuracy or patient management. Ultimately, AI's contribution to patient outcomes will be determined by how responsibly these tools are integrated into medical practice – complementing, not replacing, human providers [24-26].

B) Patient Trust and the Doctor–Patient Relationship

The human element in healthcare – empathy, communication, and trust – is essential to “bedside manner.” Introducing AI into care processes can affect how patients perceive their providers and treatments. A majority of patients express discomfort at the idea of doctors relying on AI for their care. In a Pew Research survey, 60% of Americans reported feeling uncomfortable if their physician primarily used AI for diagnosis or treatment decisions [27]. Many fear that an over-reliance on algorithms could depersonalize medical care. Indeed, 57% believed that AI would make the patient–provider relationship worse by reducing personal interaction and empathy [28-30]. Patients often value the reassurance and explanation that a human doctor provides – something a cold algorithm may lack. There is concern that if clinicians defer too much to AI recommendations, they might spend less time listening to patients or tailoring advice to individual values, thereby eroding trust.

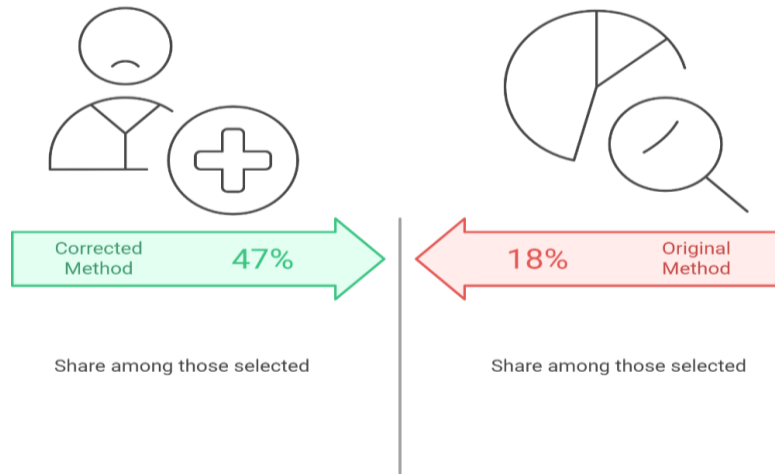
On the other hand, transparent and thoughtful use of AI could enhance trust if it leads to demonstrably better care. Patients might be more accepting of AI involvement when it is used as a tool by a caring clinician rather than as a replacement [31-32]. Clear communication is key: physicians need to explain how an AI concluded (in understandable terms) and why they are following or overriding its advice [33]. This aligns with emerging ethical guidance that patients should be informed when AI is involved in their care and given the choice to consent to its use [34-36]. When implemented with patient-centered design – for example, AI systems that provide understandable outputs or even empathy in patient-facing roles – algorithms could support the therapeutic relationship. Early experiments with AI chatbots for mental health counseling illustrate both potential and pitfalls. While these bots offer 24/7 support and some patients find it easier to open up to an uncritical machine, others find the interaction shallow or uncanny [37]. In sum, maintaining a good “bedside manner” in the age of AI will require healthcare providers to integrate algorithms in a way that augments human connection rather than diminishing it. Medical professionals may need training on how to effectively integrate AI insights into conversations with patients, ensuring that technology enhances the compassion and trust that define high-quality care [38-40].

C) Bias, Equity, and Fairness

AI in healthcare carries the risk of perpetuating or even amplifying biases and disparities in society. Algorithms learn from historical health data that may reflect unequal access or treatment [41]. If not carefully designed, an AI tool can exhibit racial, gender, or socioeconomic bias in its recommendations. A striking example is a widely used commercial algorithm for guiding high-risk care management, which was found to discriminate against Black patients [42]. The algorithm used healthcare spending as a proxy for health needs – a choice that, due to systemic inequalities, caused Black patients to appear “lower risk” than equally sick white patients. As a result, healthier white patients were being prioritized over sicker Black patients for extra care programs [43-44]. Researchers showed that fixing this bias would more than double the number of Black patients identified for enhanced care – from 18% of those selected to 47% [45]. This case highlights how embedded biases in training data can lead to inequitable healthcare decisions, even without any intent to discriminate.

Bias can enter AI systems through various pathways: unrepresentative training datasets, variables that serve as proxies for race or income, or prediction targets that reflect unequal treatment (such as using past healthcare costs, which are lower for historically underserved groups) [46]. If such biases go unchecked, AI could worsen healthcare disparities, giving privileged groups better access or tailored treatments while others get suboptimal care. This outcome would violate the ethical principle of equity in medicine and potentially run afoul of anti-discrimination laws [47]. For instance, if an AI systematically offers fewer pain medications to certain ethnic groups based on skewed data, it could raise liability under civil rights laws or healthcare quality regulations. Ensuring algorithmic fairness is thus a paramount social concern. Solutions under discussion include rigorous bias audits of healthcare algorithms and the requirement of diversity in training data [48]. Researchers Mullainathan and Obermeyer, who uncovered the bias in the risk prediction algorithm, advocate for routine auditing of algorithms “just as for medicine,” to prevent problems rather than only cure them after harm occurs [49-50].

Figure 3: Bias in a widely used care management algorithm

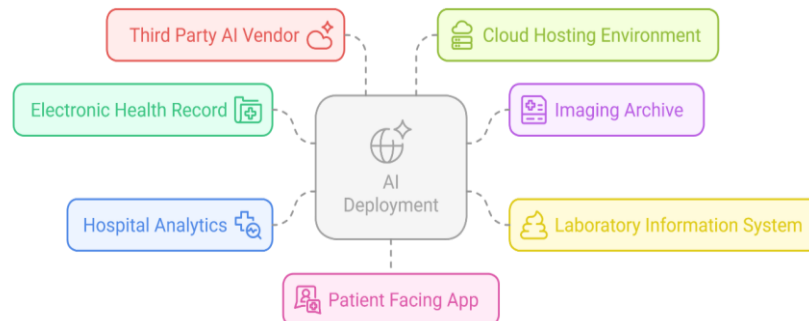


Additionally, greater transparency can allow external scrutiny. If companies disclose how their clinical AI models work (or at least their performance across different patient subgroups), biases can be identified and corrected sooner. Regulators may mandate such disclosures or bias mitigation steps as part of the approval process. In sum, achieving equitable AI in healthcare will require conscious efforts to address bias at every stage – from design and training to deployment – so that these technologies help close gaps in care rather than widen them [51]. This aligns with the vision that AI should “eliminate inequities rooted in historical and contemporary injustices” in healthcare [52], a goal increasingly emphasized by public health experts and professional bodies.

D) Privacy and Data Security

AI’s hunger for data raises significant privacy concerns in healthcare. Machine learning models often require large datasets of patient information – medical records, lab results, genomic sequences, even real-time sensor data from wearables – to train and operate effectively [53]. This intersects with stringent privacy protections in the U.S., such as the Health Insurance Portability and Accountability Act (HIPAA), which safeguards personal health information [54]. The use of AI can strain these frameworks. For example, AI systems might combine data from multiple sources (hospital records, pharmacy data, fitness apps), pushing the boundaries of what current privacy laws cover. There is also the risk of re-identification: an AI given “de-identified” data might inadvertently learn patterns that enable patient identities to be inferred, defeating privacy safeguards [55]. Moreover, suppose healthcare AI is developed or hosted by third-party tech firms. In that case, questions arise over data sharing and ownership – who has the right to use patient data to train algorithms, and do patients meaningfully consent to such use? These issues underscore the tension between technological advancement and privacy [56]. Researchers note that robust AI development requires vast amounts of data, but this must be balanced against individuals’ rights over their sensitive health information [57-60].

Figure 4: Privacy and data flow risk surface for clinical AI



In practice, privacy lapses could erode public trust and lead to legal penalties. A survey found 37% of Americans worry that increased AI use would make the security of health records worse (only 22% thought it would improve security) [61]. High-profile incidents such as data breaches of AI health apps or hospitals sharing patient scans with AI startups without proper consent have drawn public ire and regulatory scrutiny. To address these concerns, regulation is focusing on ensuring AI

systems comply with privacy laws and cybersecurity standards. The U.S. Department of Health and Human Services has clarified that HIPAA applies to AI tools used by covered entities, meaning they must implement safeguards for any patient data processed by AI, and patients should have the right to access and control that data. However, gaps remain; for instance, if an AI company uses patient data to create an algorithm and then sells that model, it's unclear if HIPAA's protections fully extend to that scenario [62].

Lawmakers and scholars are examining whether new rules are needed specifically for AI, such as requiring explicit patient consent for AI analysis of their data, or giving patients a share in the benefits when their data contributes to a profitable AI tool. On the technology side, techniques like federated learning (where AI models learn from data without it leaving the healthcare provider's servers) and differential privacy are being explored to reconcile data needs with privacy. Regulators may encourage or mandate such approaches [63]. Cybersecurity is another facet: AI systems, especially if networked, could become targets for hackers. A corrupted clinical AI could be dangerous (imagine an attacker subtly manipulating a diagnostic AI to cause misdiagnoses). Thus, ensuring strong data encryption, access controls, and continuous security testing for AI in healthcare is critical and likely to be enforced through both HIPAA and FDA guidance. Overall, safeguarding patient privacy in the AI era will require updating legal interpretations and technical standards so that innovation does not come at the cost of confidentiality and trust [64].

E) Societal Perceptions and Acceptance

The successful integration of AI into healthcare will also depend on broader societal acceptance. Beyond individual patient trust, there is a collective question of whether the public feels comfortable with AI "having a seat in the clinic." Currently, caution is prevalent. About three-quarters of Americans (75%) are concerned that healthcare providers will adopt AI too rapidly before fully understanding the risks, rather than too slowly [65]. This suggests a public desire for a prudent, safety-first approach to the deployment of health AI. Education and transparency can play a role in shaping perceptions. When people understand how a particular AI improves care – for example, an algorithm that flags early signs of stroke that a physician might miss – they may be more willing to see it used.

On the other hand, media coverage of AI failures or controversies (such as an AI misdiagnosis leading to harm, or biases in a hospital's AI system) can quickly sour public opinion [66]. Thus, maintaining public confidence will require not only making AI safe and effective, but also demonstrating that safety and effectiveness are openly demonstrated. Health institutions could engage in public outreach, explaining in patient-friendly terms where and why they use AI. Some have suggested that hospitals create AI ethics committees or patient advisory panels to involve community voices in decisions about new AI tools [67]. Such participatory approaches can make society feel a sense of ownership over healthcare innovation. Additionally, societal values such as fairness, transparency, and accountability should be reflected in how AI is regulated, to reassure the public that these technologies won't undermine core principles [68]. In the next section, we examine how U.S. laws and regulations are evolving (or struggling to catch up) to govern AI in healthcare in line with these concerns.

IV. LEGAL AND REGULATORY LANDSCAPE

The U.S. legal system is beginning to address the challenges posed by AI in healthcare. Still, it currently resembles a patchwork of adapted laws and emerging guidelines rather than a comprehensive framework. No single federal law specifically regulates "AI in healthcare" as a distinct category. Instead, various existing laws and agencies cover pieces of the puzzle: the FDA oversees medical devices (increasingly including AI software), HIPAA and related statutes cover health data privacy, and tort law (malpractice and product liability) provides avenues for patients seeking redress from AI-related harm [69]. There are also intellectual property and trade secret considerations for AI algorithms, as well as evolving guidance from professional bodies. This section outlines the key aspects of the current regulatory landscape and identifies gaps that necessitate new approaches.

A) DA Regulation of AI/ML Medical Devices

The FDA is the primary regulator for medical technologies in the U.S. and has been actively adapting its policies to accommodate AI/ML (machine learning)- based medical devices. Under the FDA's definition, many AI algorithms used in diagnosis or treatment qualify as Software as a Medical Device (SaMD) and thus fall under its jurisdiction. The agency has cleared or approved hundreds of AI-powered tools in areas like radiology, cardiology, and endocrinology. As noted, FDA authorizations of AI medical devices have surged – a 37% increase from 2023 to 2024, with 950 total devices authorized by mid-2024 [70]. The FDA has demonstrated flexibility in utilizing pathways such as the 510(k) clearance, De Novo classification, and even Breakthrough Device designations to expedite the development of beneficial AI innovations. For example, the autonomous IDx-DR diabetic retinopathy tool was reviewed via the De Novo pathway. It granted Breakthrough Device designation, reflecting the FDA's willingness to fast-track novel AI that addresses serious conditions.

However, traditional medical device regulations were not fully built with adaptive, learning algorithms in mind. A significant challenge is how to regulate AI systems that can update themselves from new data (so-called “continually learning” algorithms). The FDA has recognized this and, in recent years, proposed a framework for regulating AI/ML medical software throughout its lifecycle, rather than a one-time static approval. This includes the concept of “Predetermined Change Control Plans,” where a manufacturer can get advance FDA clearance for the algorithm to evolve within specified limits [71]. Additionally, the FDA launched a Digital Health Center of Excellence and published guiding principles for Good Machine Learning Practice (GMLP), in partnership with other agencies, to ensure AI tools are developed in line with safety and effectiveness standards [72]. Key principles involve transparency about the AI’s intended use and limitations, rigorous validation (including re-training and testing if the model is updated), and monitoring of real-world performance. The FDA’s approach is still evolving, and critics note that regulatory oversight needs to keep pace with the rapid innovation cycle of AI. Nonetheless, the FDA is striving to balance innovation with patient safety [73] – a theme echoed in its statements that it does not want to stifle beneficial AI, but will act to prevent unsafe products. One regulatory gap is that many AI tools used solely for administrative or operational support (not for making medical decisions) fall outside the FDA’s purview, yet they could indirectly impact care quality. For now, the FDA focuses on AI that provides diagnoses, treatment recommendations, or other clinical decision support, especially if marketed as autonomous or as performing beyond human capabilities. In summary, FDA oversight is a critical pillar for ensuring healthcare AI systems are safe and effective, and it is gradually being reshaped to accommodate the unique nature of AI technology.

B) Liability and Accountability in AI-Driven Care

One of the thorniest legal issues is determining who is liable when an AI system causes harm to a patient. Suppose an AI recommendation leads to a misdiagnosis or inappropriate treatment – if the patient is injured, can they sue the doctor, the hospital, the AI software manufacturer, or all of the above? As of now, there is little case law directly on point, because truly autonomous AI decisions in medicine are still emerging and have not frequently been tested in court [36]. Thus, we must extrapolate from general principles of medical malpractice and product liability law [36]. Under current doctrine, physicians are expected to meet the “standard of care,” typically defined as what a reasonably competent peer would do in similar circumstances. If a physician relies on an AI and a mistake occurs, the question becomes whether relying on that AI was consistent with the standard of care. Early commentary suggests that doctors cannot blindly blame the algorithm – they retain a duty to critically evaluate AI recommendations [35]. For example, if an AI misreads an X-ray and a physician, without reviewing the image themselves, conveys a wrong diagnosis, a court would likely find the physician negligent for not acting as a careful professional (assuming peers would have caught the error). This implies that, at least in the near term, clinicians will be held responsible for AI’s actions as if they were their own. Indeed, courts have historically been reluctant to excuse physicians based on external tools or guidelines when those lead to errors [35].

Hospitals and health systems can also face liability. If a hospital implements an AI system that is flawed or fails to properly train its staff in using it, the hospital may be directly liable for negligence. Additionally, hospitals may incur vicarious liability for the acts of their employees using AI. E.g., if a nurse uses an AI tool incorrectly and harms a patient, the hospital could be on the hook as the employer. There is also the notion of “negligent credentialing”. If a hospital adopts a faulty AI software without due diligence, it could be seen as failing to ensure quality care, analogous to credentialing a dangerous physician. What about the developers of the AI – the companies that design and sell the algorithms? They could potentially face product liability claims if their software is deemed a defective product that caused injury [35]. Product liability (under strict liability or negligence theories) usually applies to manufacturers of medical devices, so an AI diagnostic program could be viewed as a product. Suppose it can be proven that the algorithm had an inherent defect (e.g., a flaw in its training that rendered it unsafe for a certain patient group) and that defect directly led to harm.

In that case, the software company might be liable. However, there are legal uncertainties here: software has not traditionally been treated as a “product” in all jurisdictions, and companies often include disclaimers. Moreover, if the FDA approves the AI, some product liability claims could be preempted by federal law, depending on the circumstances. As it stands, liability frameworks are inadequate and still evolving, which creates a risk of both gaps (no one accountable) and overlaps (multiple parties being sued) [35]. Scholars argue that this uncertainty itself can hinder innovation – physicians may be hesitant to use AI if they fear liability. Developers might be deterred if they face an excessively high litigation risk. To address these challenges, various policy solutions are being discussed. Some propose updating the standard of care: for instance, if an AI becomes widely adopted and proven to improve outcomes, the standard of care might evolve to require its use in certain situations. In contrast, if using an AI without human oversight is known to be risky, the standard might prohibit fully autonomous use. Another idea is safe harbor laws – for example, a law could state that if a physician followed a validated AI recommendation in good faith, they would have some protection from malpractice liability.

Alternatively, an insurance or indemnification model could be developed: AI manufacturers might provide liability insurance along with their product, or healthcare institutions might carry special coverage for AI-related incidents [35]. More radically, some suggest a no-fault compensation system for injuries caused by AI (similar to vaccine injury funds), which would compensate patients without needing to prove negligence, while encouraging reporting of errors. Any such changes would require legislative action or new case law precedents. In the meantime, to reduce liability exposure, hospitals are establishing thorough validation and training for AI tools and keeping humans in the loop for critical decisions. Rethinking liability in the era of AI is essential. Ensuring accountability while not unduly punishing clinicians who use AI responsibly is a delicate balance that regulators and courts will have to strike in the coming years [74].

C) Intellectual Property and Transparency

AI algorithms in healthcare also raise Intellectual Property (IP) issues that can have legal and ethical ramifications. Companies developing advanced medical AI often seek to protect their algorithms and the data behind them. Two common strategies are patents and trade secrets. Patenting AI algorithms (or their software implementations) is possible, but can be tricky if the invention is deemed an abstract algorithm without sufficient technical application. Nonetheless, some firms have obtained patents on specific AI-based medical techniques or devices. Patents provide exclusivity but require public disclosure of how the algorithm works, at least to some extent, in the patent filing. On the other hand, many AI developers rely on trade secret protection – they do not reveal the inner workings of the model, treating it as proprietary secret sauce [1]. For example, a company might keep secret the features or variables its clinical AI analyzes to make predictions, arguing that this knowledge is commercially sensitive. From a business perspective, trade secrets can be effective (no time limit like patents, and no need to disclose). However, from a regulatory and ethical standpoint, trade secrets in healthcare AI are problematic [1].

Lack of transparency can make it hard for clinicians, patients, or regulators to understand or trust an AI's decisions. It also complicates accountability – if something goes wrong, one cannot easily scrutinize a “black box. There is an ongoing debate about whether companies should be required to disclose more about their AI, at least to regulators or in liability litigation. Scholars like Raza (2024) have critically examined the trend of using trade secrets as a substitute for formal IP in AI, noting that it can hinder oversight and patient safety in contexts like healthcare [75]. In some cases, even the training data is kept secret, which means biases or gaps in the data cannot be detected externally. From a legal perspective, regulators could compel a certain degree of algorithmic transparency. The FDA, for instance, could require documentation of how an AI makes decisions and evidence that it has been tested for biases – some of this information might be kept confidential during the FDA review. Still, it creates at least a regulatory check. If an AI's logic is too opaque, it might fail to meet the FDA's safety and effectiveness requirements. Additionally, suppose an algorithm's recommendation process is so obscure that it cannot be explained to a user. In that case, it may conflict with emerging principles of “explainable AI” in medicine, as advocated by bodies such as the American Medical Association [38].

On the IP law front, there have been discussions about whether algorithms used in critical areas should enjoy trade secret protection or whether mandatory disclosure (or even making certain algorithms open source) is in the public interest. No consensus yet, but the tension is clear: how to ensure innovation incentives for AI developers while demanding transparency for patient safety. Intellectual property law also intersects with questions of data ownership – for example, if a hospital and an AI company collaborate, who owns the resulting model or discoveries? And if an AI invents something novel (say, identifies a new drug molecule or treatment protocol), can it be an “inventor” under patent law? Current U.S. law says human inventors must make inventions, but cases are testing AI-generated invention scenarios [76]. While this is a bit tangential to direct patient care, it illustrates that the legal system is grappling with AI's role as a creator or decision-maker in its own right [7]. For our focus, the key takeaway is that IP protections for AI algorithms should not undermine the transparency and accountability needed in healthcare. Regulators and possibly legislation may need to set boundaries – for instance, requiring that critical medical AI algorithms be subject to third-party audits even if their details are secret, or perhaps limiting enforcement of trade secrets when public health is at stake. The balance between encouraging innovation (through IP rights or secrecy) and protecting patients (through transparency and oversight) is an evolving frontier in AI and health law [1].

D) Evolving Standards and Guidelines

In addition to formal laws, a number of standards, guidelines, and ethical frameworks are shaping the regulation of AI in healthcare. Professional organizations have begun issuing guidance to fill gaps faster than legislation can. The American Medical Association (AMA), for example, adopted principles for augmented intelligence in medicine in 2018–2019, emphasizing that AI tools should be evidence-based, transparent, and designed to enhance physician decision-making while respecting patient rights [38]. The AMA's policies call for physician involvement in AI development, equitable access to AI, and advocating for liability frameworks that protect patients and providers appropriately [38]. The AMA Journal of Ethics and other medical journals have published extensive discussions on the ethics of AI, emphasizing issues such as patient consent for AI use, avoiding bias, and ensuring that AI decisions can be explained within the clinical context [38]. These ethical

guidelines, while not law, often influence institutional policies. For instance, a hospital might require that any AI system used undergo an ethics review or meet certain transparency criteria, based on recommendations from medical associations [77].

Furthermore, various multi-stakeholder groups and government advisory bodies have weighed in. In 2022, the White House Office of Science and Technology Policy released a “Blueprint for an AI Bill of Rights,” which, although not specific to healthcare, outlines rights such as the right to safe and effective systems and the right to privacy, both of which are particularly relevant in medical AI contexts. There is also ongoing work by the National Institute of Standards and Technology (NIST) on an AI Risk Management Framework, which healthcare organizations could use to assess and mitigate risks of AI tools. At the federal agency level, beyond the FDA, entities like the Federal Trade Commission (FTC) have warned against unfair or deceptive practices involving AI (for example, if a direct-to-consumer AI health app makes misleading claims, the FTC could intervene). The Office for Civil Rights (which enforces HIPAA) has guidance on AI and data sharing. No unified policy exists yet, but these pieces are gradually forming a regulatory ecosystem. We see an analogous situation internationally: the European Union is finalizing an AI Act that will likely classify medical AI as “high-risk” requiring strict oversight, a contrast to the U.S.’s more sectoral approach. U.S. regulators are certainly watching these global developments.

Table 1 summarizes some of the key challenges posed by healthcare AI and the regulatory responses either in place or under consideration:

Table 1. Key Challenges of AI in Healthcare and Regulatory Responses

Challenge	Description & Impact	Regulatory/Policy Response
Patient Safety & Efficacy	Risk of inaccurate diagnoses or recommendations harming patients. AI may not work as expected across all cases.	FDA pre-market review for AI medical devices; post-market surveillance of AI performance; requirement of clinical trials and evidence for safety/effectiveness [33].
Bias & Health Disparities	Algorithms may exhibit racial, gender, or socioeconomic biases, leading to inequitable care (e.g., fewer services for minorities) [34].	Proposed bias audits and validation on diverse data before deployment, as well as potential FDA guidance on bias testing, and professional guidelines urging equity in AI design [38].
Transparency (“Black Box”)	Many AI models are opaque, making it hard for clinicians/patients to understand decisions. Lack of explainability can undermine trust.	Emerging requirement for algorithmic transparency or explainability in high-risk AI. FDA and AMA encourage interpretable models [38]. Possible mandates for documentation of AI decision logic to regulators.
Privacy & Data Security	AI requires vast amounts of patient data, which raises the risk of privacy violations or data breaches. Sensitive health data could be misused or leaked.	Enforcement of HIPAA on AI data use; HHS/OCR guidance on de-identification with AI; exploring updates to privacy laws for AI context [4]. Emphasis on cybersecurity standards for AI systems.
Liability & Accountability	Unclear who is responsible if AI causes harm, creating legal uncertainty for providers and developers [36], [35]. This can slow adoption or leave patients without recourse.	No new laws have been enacted yet; we are relying on malpractice and product liability doctrines. Policy proposals include safe harbor laws for clinicians following AI, or no-fault compensation schemes [35]. Ongoing legal scholarship to adapt liability frameworks.
Integration & Training	Challenges in integrating AI into clinical workflows and ensuring staff are adequately trained to use AI outputs appropriately.	Soft regulatory approaches: FDA and professional orgs provide best practices for human-AI teaming. Hospitals are implementing internal policies to train and ensure the proper use of AI tools. Possibly part of accreditation standards in the future.
Ethical Use & Oversight	Ensuring AI is used ethically (e.g., with patient consent, fairness, respect for autonomy) and with proper human oversight to prevent abuse or errors.	Development of ethics guidelines (AMA, WHO, etc.) for AI in healthcare. Institutional review boards (IRBs) are looking at AI in research. Some states are considering laws on AI transparency in healthcare decisions. Federal “AI Bill of Rights” principles advocating safe and ethical AI use.

As Table 1 indicates, many of the responses are still in formative stages – guidelines rather than hard rules. The regulatory system is struggling to keep pace with the rapid advancement of technology. Notably, there is a push for more interdisciplinary collaboration, with legal experts, technologists, clinicians, and ethicists working together to craft rules that are both practical and robust [77]. Study emphasizes close collaboration among stakeholders and periodic reevaluation of AI laws to keep them effective as technology evolves [1]. This kind of adaptive, cooperative approach will be essential to govern AI in a way that protects patients and society without stifling beneficial innovation.

V. ENSURING ETHICAL, EQUITABLE, AND COMPLIANT AI IN HEALTHCARE

Given the impacts and gaps discussed, how can we move forward to ensure that AI is used in healthcare in an ethical, equitable, and legally compliant manner? Several key strategies emerge:

Figure 5: Human in the loop risk ladder for clinical AI use



1. **Strengthening Regulatory Oversight:** Regulators must continue refining their frameworks to address the unique challenges posed by AI. The FDA’s ongoing efforts to update its approach for AI/ML-enabled devices are crucial. This might include requiring algorithm developers to submit not just performance data but also impact assessments covering bias, privacy, and explainability before approval. Regulators could institute conditional approvals for AI tools, monitoring them with real-world evidence collection and requiring regular updates. Data protection regulators (like HHS for HIPAA) should clarify rules for AI datasets, possibly mandating techniques like anonymization or consent management tailored to AI uses. A potential idea is a certification or “FDA seal” for AI systems that meet not only safety but also transparency and bias standards – giving providers and patients confidence that a tool is trustworthy. At the same time, regulators should stay technology-neutral enough not to hinder beneficial advances; sandboxes or pilot programs can help test AI under supervision. Agencies might also increase coordination – for example, FDA, FTC, and OCR (Office for Civil Rights) could jointly issue guidance on acceptable practices for AI developers in health, covering safety, truthfulness in marketing, and privacy compliance in one package.
2. **Legal Clarification of Liability and Standards of Care:** To alleviate uncertainty, professional boards and possibly legislatures can articulate how the standard of care applies to AI. Medical specialty boards could issue statements like “Using a validated AI for X condition is acceptable as aiding diagnosis, but does not replace clinical judgment.” Courts, when cases do arise, will set precedents – a likely early stance is that clinicians are expected to know the limits of AI tools and will be judged on overall care quality, not just blindly following or ignoring AI. Meanwhile, developers of AI should anticipate product liability claims and proactively ensure quality and safety to mitigate that risk. Going forward, a balanced liability environment might include requiring AI companies to carry insurance or indemnify users for certain failures, which effectively internalizes the risk. Policymakers could explore creating a legal safe harbor when providers use certified AI tools and follow recommended usage guidelines – this would encourage adoption of vetted systems. Conversely, if a provider chose to use a non-approved AI tool and it caused harm, that should clearly fall outside the safe harbor and likely be deemed negligent. Such measures would channel AI use towards well-regulated products and practices, benefiting patients [78].
3. **Emphasizing Ethics, Equity, and Human Oversight:** Ethical use of AI should be embedded as a norm in healthcare institutions. Hospitals can develop AI ethics committees or integrate AI considerations into existing ethics review processes. These bodies can review proposed AI deployments for potential biases or ethical pitfalls (much like an IRB would for research). To ensure equity, developers must prioritize diversity in training data and test algorithms on various subpopulations, reporting performance results by group [34]. Regulators or payers could mandate this reporting. There should also be plans for continual auditing: for instance, a hospital using an AI scheduling system might audit whether it’s inadvertently giving fewer appointments to certain demographics. If issues are found, the AI must be adjusted (or taken out of service) – ethical practice demands responsiveness. Human oversight remains a critical

safeguard; even as AI gets more autonomous, healthcare should retain a human-in-the-loop for final decisions, at least until we have absolute confidence in specific AI actions. This oversight helps catch AI errors and also provides accountability, as a named professional supervises the outcome. The level of oversight may vary: an AI reading a routine X-ray might operate mostly autonomously with spot checks, whereas a tumor board should likely double-check an AI's recommendations for cancer treatment plans. Professional guidelines increasingly echo that AI should assist, not replace, clinical judgment, and that clinicians should not rely on AI output that they do not understand or cannot validate.

4. **Education and Training:** To ensure compliance and effective use, clinicians, administrators, and patients alike require education about AI. Medical schools and continuing education programs are beginning to include training on how AI algorithms work, their limitations, and how to interpret their output. A physician who understands an AI tool is better equipped to use it correctly (and defend that use if legally challenged). Training also extends to technical staff – for example, IT departments must be trained in maintaining AI systems and protecting the data involved. On the patient side, improving general AI literacy can help patients provide informed consent and have realistic expectations for their care. Healthcare providers might develop patient pamphlets or web portals explaining the role of AI in care (e.g., “Your care may include AI analysis of your tests; here’s what that means for you”). Engaging patients and advocacy groups in discussions about AI policies can further align technology with patient values.
5. **Interdisciplinary Collaboration and Continuous Improvement:** Ensuring ethical AI is a moving target. It requires ongoing dialogue between AI engineers, healthcare professionals, ethicists, and legal experts. Multi-disciplinary panels could routinely review emerging AI technologies and advise regulators on needed actions. The study highlights the need for regular evaluation of the efficacy of AI legal frameworks, suggesting that public confidence can be built through awareness and inclusive policymaking [79]. Feedback loops should be established; for instance, any adverse events or near-misses involving AI should be collected in a national database (perhaps an extension of the FDA’s medical device reporting system) so that patterns can be identified and disseminated. This culture of learning from mistakes is vital given AI’s novelty. If a specific algorithm is found to perform poorly for a subgroup, that finding should be shared widely so that others can avoid similar issues.
6. **Legislative Initiatives:** In the longer term, dedicated legislation may be beneficial. While a comprehensive “AI in Healthcare Act” does not yet exist, lawmakers have shown interest in algorithmic accountability in other sectors. A tailored law could, for example, require that any AI used in clinical care meet certain baseline requirements (transparency, bias testing, and liability coverage) and maybe establish an oversight body or empower the FDA further. It could also clarify how patients can seek compensation if harmed by AI and even set rules for reimbursements (will Medicare/insurance pay for AI-enabled services? Under what conditions?). Some experts suggest a hybrid regulatory model: leverage existing institutions like the FDA for technical evaluation, but have new legal provisions for areas like liability and data rights where current laws are insufficient. Any new regulation, however, must be carefully crafted to avoid over-regulation that could stifle innovation or lock in specific technologies. Sandboxing new laws – perhaps through pilot programs at federal and state levels – could refine approaches before scaling them nationally [80].

VI. CONCLUSION

Artificial intelligence has undeniably arrived at the patient’s bedside, offering transformative possibilities for U.S. healthcare. From interpreting radiology scans to predicting clinical deterioration, AI systems are poised to become standard tools in the medical arsenal. The promise is great – improved diagnostic accuracy, personalized treatments, efficient hospital operations – but so are the social and legal responsibilities that accompany these algorithms. Ensuring that AI in healthcare has a “good bedside manner” is not just a metaphor for patient comfort; it signifies the broader mandate that technology must respect human values, rights, and the rule of law in medicine. Unregulated or careless use of AI could lead to erosion of trust, hidden biases exacerbating disparities, violations of privacy, and ambiguous accountability when something goes wrong. Conversely, with thoughtful governance, AI can be harnessed to enhance healthcare equitably and ethically, reinforcing the strengths of human caregivers rather than undercutting them.

The United States is still in the early days of adapting its healthcare laws and systems for the AI era. The current approach is a patchwork: the FDA is adapting its oversight of medical devices, existing privacy and liability laws are being tested, and organizations are voluntarily issuing their own ethical guidelines. This incremental progress has addressed some immediate issues, yet gaps remain. Moving forward, a more coherent framework will likely emerge, blending updated regulations with industry standards and professional norms. Key priorities will be validating AI tools rigorously, monitoring them continuously, guarding against bias and privacy breaches, and clarifying legal accountability. Equally important is maintaining a human-centric focus: technology in healthcare must ultimately serve patient well-being and autonomy. This means doctors and patients should be well-informed and in control of how AI is applied in each case.

In crafting regulations, the U.S. can draw lessons from other domains and countries. Still, solutions must be tailored to the unique context of American healthcare and adhere to established legal principles. Stakeholder engagement is vital – patients, providers, technologists, ethicists, and policymakers all need a voice in how AI is deployed on the frontlines of care. When regulation is done right, it does not stifle innovation; rather, it creates a trusted environment in which innovation can flourish responsibly. As one legal scholar put it, the goal is to encourage “disruptive innovation” in medicine while ensuring safety and public trust. That balance can be achieved through adaptive policies that evolve with the technology.

In conclusion, algorithms with a bedside manner should be more than a catchy phrase – it should be our collective aim for AI in healthcare: systems that are technically proficient, socially sensitive, and legally accountable. By proactively addressing the social impacts and fortifying the legal/regulatory guardrails now, we can welcome AI as a genuine partner in healing, rather than a source of new problems. The path ahead will require vigilance, flexibility, and collaboration across disciplines. If we succeed, the coming years will see AI delivering on its potential in healthcare – diagnosing diseases earlier, managing chronic conditions better, and extending high-quality care to more people – all while upholding the ethical standards and legal protections that are the bedrock of medicine in a democratic society.

VII. REFERENCES

- [1] H. Rafi, H. Rafiq, and M. Farhan, “Agmatine alleviates brain oxidative stress induced by sodium azide,” 2023.
- [2] S. I. Haider and M. Ali, “Mitigating the challenges of open and distance learning education system through use of information technology: a case study of Allama Iqbal Open University, Islamabad, Pakistan,” *Pakistan Journal of Distance and Online Learning*, vol. 5, no. 2, pp. 175–190, 2019.
- [3] A. Raza, “AI and privacy – navigating a world of constant surveillance,” *Euro Vantage Journal of Artificial Intelligence*, vol. 1, no. 2, pp. 74–80, 2024.
- [4] S. Khan, S. Mehmood, and S. I. Haider, “Child abuse in automobile workshops in Islamabad, Pakistan,” *Pakistan Journal of Criminology*, vol. 12, no. 1, pp. 61–74, 2020.
- [5] T. Ghulam, H. Rafi, A. Khan, K. Gul, and M. Z. Yusuf, “Impact of SARS-CoV-2 treatment on development of sensorineural hearing loss,” *Proceedings of the Pakistan Academy of Sciences: B. Life and Environmental Sciences*, vol. 58, suppl., pp. 45–54, 2021.
- [6] D. K. Mizouri et al., “Use of large language model chatbots in health care: privacy, cybersecurity, and ethical risks,” *Journal of Medical Internet Research*, vol. 25, p. e47617, 2023.
- [7] H. Rafiq, M. Farhan, H. Rafi, S. Rehman, M. Arshad, and S. Shakeel, “Inhibition of drug-induced Parkinsonism by chronic supplementation of quercetin in haloperidol-treated Wistars,” *Pakistan Journal of Pharmaceutical Sciences*, vol. 35, pp. 1655–1662, 2022.
- [8] A. Raza, “Trade secrets as a substitute for AI protection: a critical investigation into different dimensions of trade secrets,” 2024.
- [9] G. Palmeri, “Divine disguises on the crossroads of Khotan: The iconographies from Dandan Oilik,” *Journal of Asian Civilizations*, vol. 44, no. 2, pp. 67–107, 2021.
- [10] A. Ali, S. I. Haider, and M. Ali, “Role of identities in the Indo-Pak relations: a study in constructivism,” *Global Regional Review*, vol. 2, no. 1, pp. 305–319, 2017.
- [11] A. Raza, “Equality before law and equal protection of law: contextualising its evolution in Pakistan,” *Pakistan Law Journal*, 2023.
- [12] L. Notebaert, R. Harris, C. MacLeod, M. Crane, and R. S. Bucks, “The role of acute stress recovery in emotional resilience,” *PeerJ*, vol. 12, p. e17911, 2024.
- [13] A. Raza and N. Bashir, “Artificial intelligence as a creator and inventor: legal challenges and protections in copyright, patent, and trademark law,” Dec. 2023.
- [14] M. Farhan, H. Rafi, and H. Rafiq, “Dapoxetine treatment leads to attenuation of chronic unpredictable stress-induced behavioral deficits in rat model of depression,” *Journal of Pharmacy and Nutrition Sciences*, vol. 5, no. 4, pp. 222–228, 2015.
- [15] Z. Ahmed, S. Khan, S. Saeed, and S. I. Haider, “An overview of educational policies of Pakistan (1947–2020),” *Psychology and Education Journal*, vol. 58, no. 1, pp. 4459–4463, 2021.
- [16] A. Raza, “Credit, code, and consequence: how AI is reshaping risk assessment and financial equity,” *Euro Vantage Journal of Artificial Intelligence*, vol. 2, no. 2, pp. 79–86, 2025.
- [17] S. I. Haider and F. M. Burfat, “Improving self-esteem, assertiveness and communication skills of adolescents through life skills-based education,” *Journal of Social Sciences and Humanities (JSSH)*, vol. 26, no. 2, 2018.
- [18] H. Rafi and M. Farhan, “Dapoxetine: an innovative approach in therapeutic management in animal model of depression,” *Pakistan Journal of Pharmaceutical Sciences*, vol. 2, no. 1, pp. 15–22, 2015.
- [19] S. Zuberi, H. Rafi, A. Hussain, and S. Hashmi, “Upregulation of Nrf2 in myocardial infarction and ischemia-reperfusion injury of the heart,” *PLOS One*, vol. 19, no. 3, p. e0299503, 2024.
- [20] A. Raza, B. Munir, G. Ali, M. A. Othi, and R. A. Hussain, “Balancing privacy and technological advancement in AI: a comprehensive analysis of the US perspective,” *International Journal of Contemporary Issues in Social Sciences*, vol. 3, no. 3, pp. 3732–3738, 2024.
- [21] S. I. Haider and A. Waqar, “Projection of CPEC in print media of Pakistan from 2014–2019,” *Global Strategic and Security Studies Review*, vol. 1, pp. 45–64, 2019.
- [22] H. Rafi, H. Rafiq, R. Khan, F. Ahmad, J. Anis, and M. Farhan, “Neuroethological study of AlCl₃ and chronic forced swim stress induced memory and cognitive deficits in albino rats,” *The Journal of Neurobehavioral Sciences*, vol. 6, no. 2, pp. 149–158, 2019.
- [23] I. Bhatti, H. Rafi, and S. Rasool, “Use of ICT technologies for the assistance of disabled migrants in the USA,” *Revista Española de Documentación Científica*, vol. 18, no. 1, pp. 66–99, 2024.
- [24] F. Rasool and S. I. Haider, “Exploitation and low wages of labor migrants in Gulf countries,” *Global Management Sciences Review*, vol. 1, pp. 32–39, 2020.
- [25] M. Farhan, H. Rafiq, H. Rafi, S. Rehman, and M. Arshad, “Quercetin impact against psychological disturbances induced by fat-rich diet,” *Pakistan Journal of Pharmaceutical Sciences*, vol. 35, no. 5, 2022.
- [26] U.S. Department of Health and Human Services, “Administrative simplification: proposed modifications to the HIPAA Security Rule to strengthen cybersecurity,” Dec. 2024. [Online]. Available: <https://www.hhs.gov>.
- [27] H. Rafi, H. Rafiq, I. Hanif, R. Rizwan, and M. Farhan, “Chronic agmatine treatment modulates behavioral deficits induced by chronic unpredictable stress in Wistar rats,” *Journal of Pharmaceutical and Biological Sciences*, vol. 6, no. 3, pp. 80–85, 2018.

- [28] Z. Zafar, I. Sarwar, and S. I. Haider, "Socio-economic and political causes of child labor: the case of Pakistan," *Global Political Review*, vol. 1, no. 1, pp. 32–43, 2016.
- [29] A. Raza, "Navigating the intersection of artificial intelligence and law in healthcare: complications and corrections," 2024.
- [30] M. Farhan, H. Rafi, and H. Rafiq, "Behavioral evidence of neuropsychopharmacological effect of imipramine in animal model of unpredictable stress-induced depression," *International Journal of Biology and Biotechnology*, vol. 15, no. 22, pp. 213–221, 2018.
- [31] A. Raza, "Trade secrets as a substitute for AI protection: a critical investigation into different dimensions of trade secrets," 2024.
- [32] J. N. Allen, "AI chatbots and HIPAA compliance: responsibilities for developers and health providers," *Harvard Law Review Blog*, 2023.
- [33] M. F. Muhammad, H. Rafiq, and H. Rafi, "Prevalence of depression in animal model of high-fat diet-induced obesity," 2015.
- [34] H. Rafi, H. Rafiq, and M. Farhan, "Pharmacological profile of agmatine: an in-depth overview," *Neuropeptides*, vol. 105, p. 102429, 2024.
- [35] S. I. Haider and N. K. Mahsud, "Family, peer group and adaptation of delinquent behavior," *Dialogue (Pakistan)*, vol. 5, no. 4, 2010.
- [36] H. Rafi, F. Ahmad, J. Anis, R. Khan, H. Rafiq, and M. Farhan, "Comparative effectiveness of agmatine and choline treatment in rats with cognitive impairment induced by AICl₃ and forced swim stress," *Current Clinical Pharmacology*, vol. 15, no. 3, pp. 251–264, 2020.
- [37] A. Raza, "Credit, code, and consequence: how AI is reshaping risk assessment and financial equity," *Euro Vantage Journal of Artificial Intelligence*, vol. 2, no. 2, pp. 79–86, 2025.
- [38] D. D. Farhud and S. Zokaie, "Ethical issues of artificial intelligence in medicine and health care," *Iranian Journal of Public Health*, vol. 50, no. 11, pp. i–vi, 2021.
- [39] M. Zubair, S. I. Haider, and F. Khattak, "The implementation challenges to women protection laws in Pakistan," *Global Regional Review*, vol. 3, no. 1, pp. 253–264, 2018.
- [40] H. Rafiq, M. Farhan, H. Rafi, S. Rehman, and M. Arshad, "Quercetin impact against psychological disturbances induced by fat-rich diet," *Pakistan Journal of Pharmaceutical Sciences*, vol. 35, no. 5, 2022.
- [41] S. Khan, S. I. Haider, and R. Bakhsh, "Socio-economic and cultural determinants of maternal and neonatal mortality in Pakistan," *Global Regional Review*, vol. 1, pp. 1–7, 2020.
- [42] H. Rafi, H. Rafiq, and M. Farhan, "Antagonisation of monoamine reuptake transporters by agmatine improves anxiolytic and locomotive behaviors commensurate with fluoxetine and methylphenidate," *Beni-Suef University Journal of Basic and Applied Sciences*, vol. 10, no. 1, p. 26, 2021.
- [43] A. Raza, "The application of artificial intelligence in credit risk evaluation: obstacles and opportunities in the path to financial justice," *Center for Management Science Research*, vol. 3, no. 2, pp. 240–251, 2025.
- [44] L. Gordon, M. Loeb, and W. Lucyshyn, "The economics of information security investment," *Information Systems Frontiers*, vol. 17, no. 1, pp. 51–70, 2015.
- [45] H. Rafi, H. Rafiq, and M. Farhan, "Inhibition of NMDA receptors by agmatine is followed by GABA/glutamate balance in benzodiazepine withdrawal syndrome," *Beni-Suef University Journal of Basic and Applied Sciences*, vol. 10, no. 1, p. 43, 2021.
- [46] A. Raza and N. Bashir, "Artificial intelligence as a creator and inventor: legal challenges and protections in copyright, patent, and trademark law," Dec. 2023.
- [47] I. H. Syed, W. A. Awan, and U. B. Syeda, "Caregiver burden among parents of hearing impaired and intellectually disabled children in Pakistan," *Iranian Journal of Public Health*, vol. 49, no. 2, pp. 249–256, Feb. 2020.
- [48] A. Raza, "Equality before law and equal protection of law: contextualising its evolution in Pakistan," *Pakistan Law Journal*, 2023.
- [49] S. Zuberi, H. Rafi, A. Hussain, and S. Hashmi, "Role of Nrf2 in myocardial infarction and ischemia-reperfusion injury," *Physiology*, vol. 38, suppl. 1, p. 5734743, 2023.
- [50] M. Farhan, H. Rafiq, H. Rafi, F. Siddiqui, R. Khan, and J. Anis, "Study of mental illness in rat model of sodium azide induced oxidative stress," *Journal of Pharmacy and Nutrition Sciences*, vol. 9, no. 4, pp. 213–221, 2019.
- [51] A. Raza, B. Munir, G. Ali, M. A. Othi, and R. A. Hussain, "Balancing privacy and technological advancement in AI: a comprehensive analysis of the US perspective," *International Journal of Contemporary Issues in Social Sciences*, vol. 3, no. 3, pp. 3732–3738, 2024.
- [52] S. I. Haider, B. A. Shah, and N. Jehan, "Socio-economic impact of emigration on the family members left behind: A case study of district Rawalpindi," *Global Regional Review*, vol. 2, no. 1, pp. 241–252, 2017.
- [53] H. Rafi, H. Rafiq, and M. Farhan, "Agmatine improves oxidative stress profiles in rat brain tissues induced by sodium azide," *Current Chemical Biology*, vol. 18, no. 3, pp. 129–143, 2024.
- [54] S. I. Haider and N. K. Mashud, "Knowledge, attitude, and practices of violence: A study of university students in Pakistan," *Journal of Sociology and Social Work*, vol. 2, no. 1, pp. 123–145, 2014.
- [55] A. M. Desai et al., "Generative AI in health systems: adoption and challenges," *Journal of the American Medical Informatics Association*, vol. 31, no. 2, pp. 133–150, 2025.
- [56] H. Rafi, H. Rafiq, R. Khan, F. Ahmad, J. Anis, and M. Farhan, "Neuroethological study of AICl₃ and chronic forced swim stress induced memory and cognitive deficits in albino rats," *The Journal of Neurobehavioral Sciences*, vol. 6, no. 2, pp. 149–158, 2019.
- [57] D. J. Cullen and D. J. Nicholson, "Artificial intelligence and privacy in health care," in *Research Handbook on Health, AI and the Law*, 1st ed., Cheltenham: Edward Elgar, 2023, ch. 8.
- [58] M. Farhan, H. Rafiq, and H. Rafi, "Prevalence of depression in animal model of high-fat diet induced obesity," *Journal of Pharmacy and Nutrition Sciences*, vol. 5, no. 3, pp. 208–215, 2015.
- [59] F. Shafqat, S. I. Haider, A. R. Rao, and S. Waqar, "Depression, anxiety, and stress in rural and urban population of Islamabad," *The Rehabilitation Journal*, vol. 2, no. 1, pp. 44–48, 2018.
- [60] M. Zubair, S. I. Haider, and F. Khattak, "The Implementation Challenges to Women Protection Laws in Pakistan," *Global Regional Review*, vol. 3, no. 1, pp. 253–264, 2018.
- [61] A. Raza, "Credit, code, and consequence: how AI is reshaping risk assessment and financial equity," *Euro Vantage Journal of Artificial Intelligence*, vol. 2, no. 2, pp. 79–86, 2025.
- [62] S. I. Haider and A. Waqar, "Projection of CPEC in print media of Pakistan from 2014–2019," *Global Strategic and Security Studies Review*, vol. 1, pp. 45–64, 2019.
- [63] H. Rafi, H. Rafiq, and M. Farhan, "Pharmacological profile of agmatine: an in-depth overview," *Neuropeptides*, vol. 105, p. 102429, 2024.
- [64] M. Farhan, H. Rafi, and H. Rafiq, "Dapoxetine treatment leads to attenuation of chronic unpredictable stress-induced behavioral deficits in a rat model of depression," *Journal of Pharmacy and Nutrition Sciences*, vol. 5, no. 4, pp. 222–228, 2015.
- [65] L. Notebaert, R. Harris, C. MacLeod, M. Crane, and R. S. Bucks, "The role of acute stress recovery in emotional resilience," *PeerJ*, vol. 12, p. e17911, 2024.
- [66] M. Farhan, H. Rafiq, H. Rafi, R. Ali, and S. Jahan, "Neuroprotective role of quercetin against neurotoxicity induced by lead acetate in male rats," unpublished.

- [67] S. Khan and S. I. Haider, "Women's education and empowerment in Islamabad, Pakistan," *Global Economics Review*, vol. 5, no. 1, pp. 50–62, 2020.
- [68] A. Raza, "AI and privacy – navigating a world of constant surveillance," *Euro Vantage Journal of Artificial Intelligence*, vol. 1, no. 2, pp. 74–80, 2024.
- [69] D. D. Farhud and S. Zokaie, "Ethical issues of artificial intelligence in medicine and health care," *Iranian Journal of Public Health*, vol. 50, no. 11, pp. i–vi, 2021.
- [70] S. I. Haider and N. K. Mahsud, "Family, Peer Group and Adaptation of Delinquent Behavior," *Dialogue (Pakistan)*, vol. 5, no. 4, 2010.
- [71] A. Raza and N. Bashir, "Artificial intelligence as a creator and inventor: legal challenges and protections in copyright, patent, and trademark law," Dec. 2023.
- [72] Z. Ahmed, S. Khan, S. Saeed, and S. I. Haider, "An overview of educational policies of Pakistan (1947–2020)," *Psychology and Education Journal*, vol. 58, no. 1, pp. 4459–4463, 2021.
- [73] A. Ali, S. I. Haider, and M. Ali, "Role of identities in the Indo-Pak relations: a study in constructivism," *Global Regional Review*, vol. 2, no. 1, pp. 305–319, 2017.
- [74] A. Raza, "Navigating the intersection of artificial intelligence and law in healthcare: complications and corrections," 2024.
- [75] D. K. Mizouri et al., "Use of large language model chatbots in health care: privacy, cybersecurity, and ethical risks," *Journal of Medical Internet Research*, vol. 25, p. e47617, 2023.
- [76] H. Rafi, F. Ahmad, J. Anis, R. Khan, H. Rafiq, and M. Farhan, "Comparative effectiveness of agmatine and choline treatment in rats with cognitive impairment induced by AICl₃ and forced swim stress," *Current Clinical Pharmacology*, vol. 15, no. 3, pp. 251–264, 2020.
- [77] J. N. Allen, "AI chatbots and HIPAA compliance: responsibilities for developers and health providers," *Harvard Law Review Blog*, 2023.
- [78] M. Farhan, H. Rafiq, and H. Rafi, "Behavioral evidence of neuropsychopharmacological effect of imipramine in an animal model of unpredictable stress-induced depression," *International Journal of Biology and Biotechnology*, vol. 15, no. 22, pp. 213–221, 2018.
- [79] S. Khan, S. Mehmood, and S. I. Haider, "Child abuse in automobile workshops in Islamabad, Pakistan," *Pakistan Journal of Criminology*, vol. 12, no. 1, pp. 61–74, 2020.
- [80] A. M. Desai et al., "Generative AI in health systems: adoption and challenges," *Journal of the American Medical Informatics Association*, vol. 31, no. 2, pp. 133–150, 2025.