

INTERNATIONAL JOURNAL OF LAW
MANAGEMENT & HUMANITIES
[ISSN 2581-5369]

Volume 9 | Issue 4

2026

© 2026 International Journal of Law Management & Humanities

Follow this and additional works at: <https://www.ijlmh.com/>

Under the aegis of VidhiAagaz – Inking Your Brain (<https://www.vidhiaagaz.com/>)

This article is brought to you for free and open access by the International Journal of Law Management & Humanities at VidhiAagaz. It has been accepted for inclusion in the International Journal of Law Management & Humanities after due review.

In case of any suggestions or complaints, kindly contact support@vidhiaagaz.com.

To submit your Manuscript for Publication in the International Journal of Law Management & Humanities, kindly email your Manuscript to submission@ijlmh.com.

Medical Liability in AI-Assisted Clinical Decision Making: Need for a New Legal Framework

SHANU SINGH CHOUHAN* AND DR. MANISH YADAV**

ABSTRACT

The progressive integration of artificial intelligence into clinical decision-making processes has generated a legal lacuna of considerable consequence, one that existing medical liability frameworks, designed for human actors operating within conventional physician-patient relationships, are structurally ill-equipped to address. This research paper examines the doctrinal inadequacies of current medical negligence law when applied to AI-assisted clinical decisions, analyses the liability attribution challenges arising from the triadic relationship between physician, patient, and algorithmic system, and evaluates comparative regulatory responses across the United States, European Union, Australia, and India. Drawing upon landmark judicial decisions, legislative frameworks, and academic literature, the paper argues that the resolution of medical liability in AI-assisted clinical decision-making requires not the incremental adaptation of existing tort principles but the deliberate construction of a new, purpose-built legal framework, one that assigns responsibility equitably, incentivises safety, preserves patient rights, and remains responsive to the pace of technological evolution.

Keywords: Artificial Intelligence, Medical Liability, Clinical Decision Making, Medical Negligence, Algorithm, Tort Law, Patient Rights, Regulatory Framework, India

I. INTRODUCTION

The physician-patient relationship, as it has been understood and governed by medical law for centuries, rests upon a foundational premise: that clinical decisions are made by identifiable, accountable human professionals who possess defined qualifications, owe legally cognisable duties of care, and bear personal responsibility for the consequences of their clinical judgements. This premise, which has informed medical negligence doctrine, professional ethics frameworks, and patient rights law across jurisdictions, is being fundamentally disrupted by the integration of artificial intelligence into clinical decision-making. Artificial intelligence systems are no longer experimental tools confined to research laboratories. They are operational clinical

* Author is a Research Scholar at National Law Institute University, Bhopal, Madhya Pradesh, India.

** Author is an Associate Professor at National Law Institute University, Bhopal, Madhya Pradesh, India.

instruments embedded in the daily practice of hospitals, diagnostic centres, and outpatient facilities worldwide. As of August 2024, the United States Food and Drug Administration had authorised approximately 950 AI-enabled medical devices, with 235 approvals in 2024 alone, the highest annual total in the agency's regulatory history, with radiology, cardiology, and neurology accounting for the majority of authorisations. AI systems are reading electrocardiograms, detecting pulmonary nodules, recommending cancer treatment protocols, triaging emergency patients, and generating clinical documentation, functions that were, until recently, the exclusive domain of licensed medical professionals.

This clinical reality has produced a legal paradox of considerable urgency. When an AI system contributes to a clinical error, whether by recommending an incorrect diagnosis, generating a biased treatment suggestion, or failing to flag a life-threatening abnormality, the existing architecture of medical liability law struggles to provide a coherent, just, and workable answer to the foundational question: who is responsible? The developer who built the algorithm? The hospital that deployed it? The clinician who relied upon it? The regulatory authority that approved it? Or some combination of all four, allocated in proportions that existing doctrine has not yet determined?

This paper argues that the answer to this question cannot be found within the existing framework of medical negligence law alone, and that the development of a new, AI-specific medical liability framework is both intellectually necessary and practically urgent.

II. THE EXISTING FRAMEWORK OF MEDICAL LIABILITY: DOCTRINAL FOUNDATIONS AND STRUCTURAL ASSUMPTIONS

A. Medical Negligence: The Classical Doctrine

Medical liability in common law jurisdictions is primarily governed by the law of negligence, a doctrinal framework requiring the establishment of four elements: the existence of a duty of care owed by the defendant to the claimant; a breach of that duty through failure to meet the applicable standard of care; causation between the breach and the harm suffered; and resulting damage. In the medical context, the standard of care has been classically articulated through the *Bolam* test, established in *Bolam v. Friern Hospital Management Committee* [1957] 1 WLR 582, wherein McNair J held that a doctor is not negligent if he acts in accordance with a practice accepted as proper by a responsible body of medical men skilled in that particular art.¹ This standard, subsequently qualified by the requirement of logical defensibility established in

¹ *Bolam v. Friern Hospital Management Committee* [1957] 1 WLR 582.

Bolitho v. City and Hackney Health Authority [1998] AC 232, has remained the foundational benchmark of medical liability in England, India, and numerous common law jurisdictions for over six decades.

In the Indian context, the Supreme Court adopted the *Bolam* principle in *Jacob Mathew v. State of Punjab* (2005) 6 SCC 1, holding that a professional may be held liable for negligence only if she or he falls short of the standard of a reasonably competent professional in that field. The Court further distinguished between an error of judgement and negligence proper, establishing that not every adverse outcome in clinical practice constitutes actionable negligence, and that the standard applied must account for the inherent uncertainties of medical practice.² The *Jacob Mathew* decision has remained the governing precedent in India's medical negligence jurisprudence, supplemented by the consumer protection jurisdiction established by the Supreme Court in *Indian Medical Association v. V.P. Shantha* (1995) 6 SCC 651, which brought medical services within the ambit of the Consumer Protection Act and provided patients with an additional, more accessible avenue of redress.³

B. Structural Assumptions of Classical Medical Liability

The classical medical liability framework rests upon several structural assumptions that AI-assisted clinical decision-making directly and fundamentally challenges. First, it assumes that the decision-maker is a human professional, identifiable, licensed, and personally accountable. Second, it assumes that the decision-making process is, at least in principle, transparent and constructible, that the clinical reasoning leading to a particular decision can be articulated, examined, and evaluated against the standard of a reasonable professional. Third, it assumes a bilateral relationship, between physician and patient, in which responsibility flows along a single, clearly traceable chain of accountability. Fourth, it assumes that the tools and technologies employed by the clinician are instrumentalities of their professional judgement rather than autonomous contributors to the clinical decision.⁴

Each of these assumptions is disrupted when AI enters the clinical environment. AI systems are not human professionals and cannot hold licences, bear personal responsibility, or be subjected to disciplinary proceedings. The decision-making processes of deep learning systems, the most clinically prevalent AI architecture, are notoriously opaque, operating through millions of weighted parameters that resist meaningful human comprehension or reconstruction: the so-

² *Bolitho v. City and Hackney Health Authority* [1998] AC 232.

³ *Indian Medical Association v. V.P. Shantha*, (1995) 6 SCC 651.

⁴ Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, opened for signature Sept. 5, 2024, C.E.T.S. No. 225.

called "black box" problem that renders the transparency assumption of classical doctrine inapplicable. The physician-patient relationship acquires a triadic character when an AI system participates in clinical reasoning, introducing a developer, a deploying institution, and an algorithmic intermediary into a liability framework designed for two parties. And AI systems, particularly autonomous diagnostic tools, are not merely instruments of professional judgement but active contributors to, and in some cases primary generators of, the clinical recommendation upon which the patient's care is based.⁵

III. THE LIABILITY ATTRIBUTION PROBLEM IN AI-ASSISTED CLINICAL DECISIONS

A. The Triadic Relationship and the Diffusion of Responsibility

The introduction of AI into clinical decision-making creates what legal scholars have termed the "responsibility gap", a situation in which the diffusion of agency across multiple actors (developer, deployer, and supervising clinician) results in no single actor bearing unambiguous legal responsibility for an adverse clinical outcome. This diffusion operates across three distinct but interconnected dimensions.

At the developer level, AI system developers design, train, validate, and commercially distribute clinical AI products. Their liability exposure under existing frameworks is primarily governed by product liability law, specifically the question of whether the AI system constitutes a defective product within the meaning of relevant product liability statutes. In India, product liability provisions under the Consumer Protection Act, 2019, specifically Chapter VI, Sections 82 to 87, impose liability on manufacturers for manufacturing defects, design defects, and inadequate warnings.⁶ However, AI systems present unique challenges for product liability doctrine: they are not static products but adaptive systems that learn and evolve post-deployment, potentially developing failure modes that did not exist at the time of manufacture and could not have been anticipated through conventional pre-market testing. The question of whether a post-deployment algorithmic failure constitutes a manufacturing defect, a design defect, or an adequacy-of-warning issue, and who bears the burden of demonstrating which, has not been definitively resolved in Indian or most comparative jurisdictions.⁷

⁵ Regulation 2024/1689, of the European Parliament and of the Council of 13 June 2024 Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act), 2024 O.J. (L 1689) 1.

⁶ W. Nicholson Price II, Sara Gerke & I. Glenn Cohen, *Liability for Use of Artificial Intelligence in Medicine*, in RESEARCH HANDBOOK ON HEALTH, AI AND THE LAW (Barry Solaiman & I. Glenn Cohen eds., 2024).

⁷ Nithesh Naik et al., *Legal and Ethical Consideration in Artificial Intelligence in Healthcare: Who Takes Responsibility?*, 9 FRONT. SURG. 862322 (2022).

At the deploying institution level, hospitals and healthcare organisations that integrate AI systems into clinical workflows potentially incur liability through several established doctrines: direct institutional negligence for failure to adequately validate the AI system for use in their specific patient population; negligent credentialing for deploying a system without adequate clinical governance protocols; and vicarious liability for the clinical errors of employed or contracted clinicians who rely upon AI recommendations. The institutional liability question is complicated by the reality that many AI systems are acquired through commercial licensing agreements that include liability limitation and indemnification clauses, contractual arrangements that may effectively transfer liability risk away from deep-pocketed developers to less resourced healthcare institutions, with potentially inequitable consequences for patient compensation.

At the clinician level, the supervising physician who relies upon an AI recommendation faces liability exposure under traditional medical negligence doctrine, but the nature and extent of this exposure is deeply uncertain. The central question is the degree to which a clinician is entitled to rely upon an AI recommendation without independent verification, and conversely, the degree to which the "automation bias", the well-documented tendency of human operators to over-defer to automated systems, constitutes a breach of the standard of care. As AI systems achieve diagnostic accuracy that rivals or exceeds human specialists in specific domains, the clinical and legal calculus of appropriate reliance becomes increasingly complex: a clinician who overrides a correct AI recommendation may be as legally vulnerable as one who defers to an incorrect one.⁸

B. The Black Box Problem and Causation

The opacity of deep learning systems creates a further doctrinal challenge at the level of causation, the requirement that a claimant demonstrate that the defendant's breach was a material cause of the harm suffered. In conventional medical negligence, causation is established through expert evidence reconstructing the clinical decision-making process and demonstrating that a different decision, in accordance with the applicable standard of care, would have produced a different outcome. Where the clinical error is generated or substantially influenced by an AI system whose internal reasoning cannot be meaningfully reconstructed or explained, this evidentiary exercise becomes practically impossible. The clinician cannot explain why the AI reached its recommendation; the developer can identify the algorithmic

⁸ Kavitha Palaniappan, Elaine Yan Ting Lin & Silke Vogel, *Global Regulatory Frameworks for the Use of Artificial Intelligence (AI) in the Healthcare Services Sector*, 12 HEALTHCARE (BASEL) 562 (2024).

output but cannot articulate the specific reasoning pathway that produced it; and the patient is left unable to establish the causal chain that legal doctrine requires.⁹

This problem is not merely theoretical. In *State v. Loomis*, 881 N.W.2d 749 (Wis. 2016), a case involving algorithmic criminal sentencing tools, the Wisconsin Supreme Court acknowledged the opacity of the COMPAS recidivism algorithm while upholding its use, in a decision that drew substantial academic criticism for sanctioning consequential algorithmic decision-making without requiring explainability. While not a medical case, *Loomis* illustrates the judicial discomfort with, and inadequate doctrinal tools for addressing, the black box problem that will inevitably confront medical liability courts in AI-assisted clinical negligence claims.¹⁰

IV. COMPARATIVE REGULATORY AND LEGAL RESPONSES

A. European Union: The AI Act and Medical Device Regulation

The European Union has developed the most comprehensive legal response to the challenge of AI liability, through the combined operation of the EU AI Act (2024), the Medical Devices Regulation (EU) 2017/745, and the proposed AI Liability Directive. The EU AI Act classifies AI systems used in patient diagnosis, monitoring, and treatment recommendation as high-risk AI systems under Annex III, subjecting them to mandatory conformity assessments, technical documentation requirements, algorithmic transparency obligations, and mandatory human oversight mechanisms before market deployment. The Act's Article 13 establishes a general transparency obligation requiring that high-risk AI systems be designed to allow deployers to interpret their outputs and use them appropriately, a direct regulatory response to the black box problem. The proposed EU AI Liability Directive, published in September 2022, introduces a rebuttable presumption of causation in AI-related harm cases, significantly reducing the evidentiary burden on claimants by presuming a causal link between the AI system's failure and the harm suffered, subject to the defendant's ability to rebut the presumption.¹¹ This legislative innovation represents a deliberate departure from the strict requirements of conventional causation doctrine and reflects the EU legislature's recognition that existing proof-of-causation standards are operationally incompatible with the opaque, multi-actor character of AI-related harm.

B. United States: Fragmented Regulation and Emerging Litigation

⁹ W. Nicholson Price, Sara Gerke & I. Glenn Cohen, *Potential Liability for Physicians Using Artificial Intelligence*, 322 JAMA 1765 (2019).

¹⁰ *State v. Loomis*, 881 N.W.2d 749 (Wis. 2016).

¹¹ Price II, Gerke & Cohen, *supra* note 6.

The United States regulatory response to AI medical liability remains fragmented and primarily sector-specific. The FDA governs AI-enabled medical devices under its Software as a Medical Device framework, applying existing medical device regulations supplemented by AI-specific guidance documents including the 2021 Action Plan for AI/ML-Based Software as a Medical Device. The FDA has introduced the concept of the Predetermined Change Control Plan, a mechanism allowing developers to prospectively define the parameters within which an AI system may learn and adapt post-approval without requiring a new regulatory submission, addressing the challenge of adaptive AI systems within a regulatory framework designed for static products. However, the United States lacks a unified federal AI liability statute, and medical liability remains primarily a matter of state tort law, producing significant jurisdictional inconsistency. Emerging AI medical liability litigation in the United States has begun to grapple with questions of standard of care in AI-assisted diagnosis, including whether a clinician's failure to override an erroneous AI recommendation constitutes actionable negligence, but appellate-level guidance on these questions remains limited.¹²

Scholarly commentary in the United States has increasingly advocated for the adoption of a strict liability standard for AI medical device developers, analogous to the strict liability doctrine established for defective products in *Greenman v. Yuba Power Products, Inc.*, 377 P.2d 897 (Cal. 1963), on the grounds that developers are best positioned to bear, distribute, and minimise the risks of AI system failures, and that negligence-based standards create inadequate incentives for safety investment given the evidentiary difficulties of proving developer fault in AI error cases.¹³

C. Australia: TGA Regulation and Evolving Standards

Australia regulates AI medical devices through the Therapeutic Goods Administration under the Therapeutic Goods Act, 1989, as amended in 2021 to specifically encompass AI-enabled medical software. The TGA applies a risk-based classification framework to AI medical devices, with higher-risk systems subject to more stringent pre-market assessment and post-market surveillance obligations. Australia's approach to AI liability more broadly is governed by its existing tort law framework supplemented by the Australian Consumer Law, which, like India's Consumer Protection Act, provides product liability remedies for defective goods, while the Voluntary AI Safety Standard (2024) and the AI Ethics Framework (2019) provide non-binding guidance on developer responsibilities. Australia is progressing toward mandatory regulation for high-risk AI applications, including healthcare, and ongoing governmental consultations

¹² *Id.*

¹³ *Greenman v. Yuba Power Products, Inc.*, 377 P.2d 897 (Cal. 1963).

signal a forthcoming legislative response to the liability gap that voluntary frameworks cannot adequately address.¹⁴

D. India: The Regulatory Vacuum and Its Consequences

India's legal framework governing medical liability in AI-assisted clinical decisions is characterised by a regulatory vacuum of considerable concern. No dedicated statute governs AI medical systems. Medical negligence continues to be governed by the *Jacob Mathew* standard under tort law and the Consumer Protection Act, 2019, neither of which addresses the specific complexities of AI-assisted clinical errors. The Digital Personal Data Protection Act, 2023 provides a partial data governance framework but does not address algorithmic accountability, bias, or clinical validation. The Medical Devices Rules, 2017, amended periodically by the Central Drugs Standard Control Organisation, govern medical devices but lack AI-specific provisions addressing adaptive learning systems, explainability requirements, or post-market surveillance obligations proportionate to the risk profile of clinical AI.¹⁵

The consequences of this regulatory vacuum are immediately and practically consequential. A patient harmed by an erroneous AI-assisted diagnosis in an Indian hospital currently has no clear statutory remedy specifically addressing AI-related clinical harm; no defined evidentiary framework for establishing causation in black box error cases; no mandatory reporting mechanism for AI clinical adverse events; and no regulatory authority specifically empowered to investigate, sanction, or remediate AI-related patient safety failures. The existing consumer forum and civil court structure, while theoretically available, lacks the technical capacity, evidentiary frameworks, and doctrinal tools necessary to adjudicate AI medical liability claims with the sophistication and consistency that such claims demand.¹⁶

V. TOWARDS A NEW LEGAL FRAMEWORK: PRINCIPLES AND ARCHITECTURE

A. Foundational Principles

The construction of a new legal framework for medical liability in AI-assisted clinical decision-making must be grounded in a coherent set of foundational principles that reflect both the specific character of AI-related harm and the broader imperatives of patient rights, clinical safety, and institutional accountability. Drawing from comparative analysis and academic literature, the following principles are proposed as the normative foundations of such a framework.

¹⁴ Viet-Thi Tran, Carolina Riveros & Philippe Ravaud, *Patients' Views of Wearable Devices and AI in Healthcare: Findings from the ComPaRe e-Cohort*, 2 NPJ DIGITAL MED. 53 (2019).

¹⁵ *Indian Medical Association v. V.P. Shantha*, *supra* note 3.

¹⁶ *Samira Kohli v. Dr. Prabha Manchanda*, (2008) 2 SCC 1.

The principle of graduated accountability requires that liability be distributed across the chain of actors contributing to an AI-assisted clinical decision, developers, deploying institutions, and supervising clinicians, in proportions calibrated to their respective roles, degrees of control, and capacities to prevent the harm in question. This principle rejects both the exclusive attribution of liability to the supervising clinician, which would unjustly burden practitioners who rely in good faith on certified AI systems, and the complete immunisation of developers from the clinical consequences of their product's failures.

The principle of algorithmic transparency requires that all AI systems deployed in clinical settings be designed and maintained to provide humanly comprehensible explanations of their recommendations, not as a post-hoc regulatory requirement but as a design-stage engineering obligation. Transparency is not merely an ethical preference; it is a legal prerequisite for the meaningful exercise of clinical supervision and the maintenance of informed consent.

The principle of mandatory clinical validation requires that all AI systems deployed in clinical settings demonstrate safety and efficacy through rigorous, demographically representative clinical trials before deployment, with particular attention to the performance of AI systems across the diverse patient populations of the jurisdiction in question, including the genetic, linguistic, and socioeconomic diversity of India's 1.4 billion people.¹⁷

The principle of causation facilitation requires that the evidentiary burden on patients harmed by AI clinical errors be modified to reflect the structural opacity of AI decision-making, through rebuttable presumptions of causation, disclosure obligations on developers and deploying institutions, and judicially managed technical evidence frameworks that enable meaningful liability adjudication without requiring patients to demonstrate the internal workings of complex algorithmic systems.

The principle of patient-centred compensation requires that liability frameworks ensure meaningful, accessible, and timely compensation for patients harmed by AI clinical errors, through a combination of direct liability claims, no-fault compensation mechanisms, and mandatory insurance requirements, without making compensation contingent on the resolution of complex multi-party liability disputes that may take years to adjudicate.

B. Proposed Architecture of a New Framework for India

Building upon these principles, a new legal framework for medical liability in AI-assisted clinical decision-making in India should comprise the following structural elements.

¹⁷ Seizing the Opportunities of Safe, Secure and Trustworthy Artificial Intelligence Systems for Sustainable Development, G.A. Res. 78/265 (Mar. 21, 2024).

A Dedicated AI Healthcare Liability Act should establish the primary legislative architecture, creating a purpose-built liability regime for AI-assisted clinical harm that operates alongside, rather than replacing, existing medical negligence and consumer protection law. This Act should define AI medical systems with precision; establish a mandatory developer liability provision imposing strict liability for manufacturing and design defects in AI clinical systems; create an institutional liability standard for deploying healthcare organisations that fail to implement adequate AI governance, validation, and oversight protocols; and codify a modified clinician liability standard that distinguishes between culpable automation bias and reasonable reliance on certified AI systems.

A Mandatory AI Clinical Adverse Event Registry should be established under the authority of an empowered regulatory body, analogous to the FDA's MAUDE database, requiring all healthcare institutions to report AI-related clinical adverse events, near-misses, and performance anomalies within defined timelines. Registry data should be publicly accessible in anonymised form, enabling regulators, researchers, and policymakers to identify systemic AI safety risks and develop evidence-based regulatory responses.¹⁸

A Modified Causation Standard should be legislatively enacted, introducing a rebuttable presumption of causation in cases where a claimant demonstrates that an AI system contributed to a clinical decision that deviated from the standard of care and resulted in demonstrable harm, shifting the evidentiary burden to the defendant developer or deploying institution to disprove the causal connection. This modification reflects both the structural opacity of AI decision-making and the principle that defendants with superior access to algorithmic information should bear the primary burden of causation evidence.

A No-Fault Compensation Mechanism should be established for a defined category of AI clinical adverse events, analogous to vaccine injury compensation programmes, providing swift, accessible compensation to patients without requiring protracted litigation, and funded through mandatory levies on AI medical system developers and deploying institutions. This mechanism would operate in parallel with the fault-based liability regime, ensuring that patient compensation is not held hostage to the resolution of complex multi-party responsibility disputes.

A Technical Expert Panel System should be established within the judicial framework, comprising AI engineers, clinical informaticists, medical professionals, and legal scholars, empowered to assist courts in evaluating technical evidence in AI medical liability proceedings,

¹⁸ WORLD HEALTH ORGANIZATION, *ETHICS AND GOVERNANCE OF ARTIFICIAL INTELLIGENCE FOR HEALTH* (2021).

assessing the adequacy of clinical validation data, and providing judicially recognised expert determinations on questions of algorithmic design, failure analysis, and causation that fall beyond the competence of conventional medical expert witnesses.

VI. THE QUESTION OF INFORMED CONSENT IN AI-ASSISTED CARE

A dimension of medical liability that requires specific attention in the AI context is the doctrine of informed consent, the patient's right to receive adequate information about proposed treatment, including its risks and alternatives, as a precondition of valid consent to clinical intervention. The Supreme Court of India in *Samira Kohli v. Dr. Prabha Manchanda* (2008) 2 SCC 1 established that a patient's consent to medical treatment must be real and informed, based on adequate disclosure of risks and alternatives, and that breach of the duty of informed disclosure constitutes an independent basis of medical liability.¹⁹ In the AI context, the informed consent doctrine raises urgent questions that existing jurisprudence has not addressed: Does a patient have the right to know that AI contributed to their diagnosis or treatment recommendation? What specific information about the AI system's capabilities, limitations, validation data, and error rates must be disclosed? Does a patient have the right to refuse AI-assisted clinical decision-making without prejudice to their access to care?

These questions are not merely theoretical. Research consistently demonstrates that patients have differentiated attitudes toward AI involvement in their care, accepting AI assistance in administrative functions while expressing significant reservations about autonomous AI clinical decision-making, and that the right to meaningful consent to AI-assisted care is both ethically grounded and legally cognisable. A new legal framework must therefore incorporate specific informed consent provisions for AI-assisted clinical decisions, defining the minimum disclosure obligations of healthcare providers, establishing the patient's right to human-only clinical assessment upon request, and ensuring that AI consent is integrated into, rather than obscured within, existing consent documentation processes.

VII. INTERNATIONAL HUMAN RIGHTS DIMENSIONS

The legal framework governing medical liability in AI-assisted clinical decisions must be situated within the broader architecture of international human rights law. The right to health, recognised under Article 12 of the International Covenant on Economic, Social and Cultural Rights and interpreted through General Comment No. 14 of the UN Committee on Economic, Social and Cultural Rights, imposes on state parties an obligation to ensure that health facilities,

¹⁹ *Samira Kohli v. Dr. Prabha Manchanda*, *supra* note 16.

goods, and services are available, accessible, acceptable, and of adequate quality. AI clinical systems that are biased, inadequately validated, or deployed without adequate oversight violate this obligation, not merely as instruments of individual clinical error but as systemic threats to the equitable quality of healthcare to which all persons are entitled. The UN General Assembly Resolution on AI safety (2024) and the Council of Europe's Framework Convention on AI and Human Rights, the first internationally binding AI treaty, opened for signature in September 2024, each reinforce the principle that AI systems must be designed and deployed in conformity with human rights obligations, and that states bear accountability for ensuring rights-protective AI governance in domains as consequential as healthcare.²⁰

VIII. CONCLUSION

The integration of artificial intelligence into clinical decision-making has generated a liability crisis that the existing framework of medical negligence law, designed for a different technological era and a different structure of clinical relationships, cannot adequately resolve. The classical doctrines of duty, breach, causation, and damage; the *Bolam/Jacob Mathew* standard of professional accountability; the bilateral framework of physician-patient liability; and the conventional tools of evidentiary proof, each faces structural inadequacy when applied to the opaque, multi-actor, algorithmically mediated character of AI-assisted clinical harm.

The comparative analysis of regulatory responses in the European Union, United States, and Australia demonstrates both the urgency of this challenge and the diversity of approaches available for its resolution, with the EU's binding, rights-based, presumption-of-causation framework representing the most structurally comprehensive response to date. For India, a nation of 1.4 billion people, a rapidly expanding healthcare AI sector, a documented regulatory vacuum in AI medical governance, and a patient population bearing disproportionate exposure to the risks of inadequately governed AI clinical systems, the development of a new, purpose-built legal framework for medical liability in AI-assisted clinical decision-making is not a future policy aspiration. It is a present constitutional obligation, a patient safety imperative, and a legal justice demand that the existing framework is structurally incapable of satisfying.

The framework proposed in this paper, grounded in the principles of graduated accountability, algorithmic transparency, mandatory clinical validation, causation facilitation, and patient-centred compensation, provides a principled architectural foundation for legislative action. Its implementation would require coordinated engagement across legislative, judicial, regulatory, and professional institutional domains, but the imperative of that engagement is beyond

²⁰ Council of Europe Framework Convention on Artificial Intelligence, *supra* note 4.

scholarly doubt. As artificial intelligence assumes an increasingly central role in the clinical decisions that determine patient outcomes, the law's failure to keep pace with technology is not merely an academic inconvenience. It is a systemic threat to the rights, safety, and dignity of every patient who enters a healthcare system in which algorithmic systems are making decisions that the law has not yet learned to govern.
