
AI-POWERED DECISION-MAKING IN CORPORATIONS: LEGAL LIABILITY, CORPORATE GOVERNANCE, AND THE QUESTION OF LEGAL PERSONHOOD

Ms. Rubi Chauhan, LL.M., Lovely Professional University, Phagwara

ABSTRACT

The rapid integration of artificial intelligence (AI) into corporate decision-making has transformed how companies operate, driving efficiency, scale, and strategic agility. Yet, this shift also presents unprecedented legal and governance challenges. This research paper examines the evolving intersection of AI deployment and corporate accountability, structured across three key dimensions: legal liability, corporate governance, and the theoretical question of AI personhood. Drawing on case law from the United States, European Union, and India, the paper analyses how courts and regulators attribute responsibility when AI systems cause harm, whether through biased hiring, financial mismanagement, or operational failures. It explores the duties of corporate boards under fiduciary and statutory frameworks, highlighting the increasing risk of liability for directors who fail to oversee algorithmic tools adequately.

The study further assesses arguments for and against granting AI systems legal personhood, ultimately rejecting this approach in favour of distributed accountability models that preserve human oversight. Sector-specific applications in finance, employment, healthcare, and logistics underscore the regulatory fragmentation and gaps in current legal regimes. Empirical data from corporate surveys and governance reports reveal a systemic shortfall in AI oversight, with boards often lacking the expertise or frameworks to manage AI-related risk.

The paper concludes with detailed policy recommendations, advocating for a risk-tiered governance structure, legally mandated AI audits, and enhanced boardroom literacy. By synthesizing comparative jurisprudence, empirical research, and normative theory, this study offers a foundational legal architecture for governing AI in the corporate context, balancing innovation with accountability in the algorithmic age.

Keywords: Artificial Intelligence, Corporate Governance, Corporate Liability, AI Personhood, Algorithmic Liability.

Introduction

The integration of Artificial Intelligence into corporate decision-making processes has reshaped how modern enterprises operate. AI tools have empowered corporations to process vast amounts of data and make swift data-driven decisions, from algorithmic trading to automated hiring. According to PwC's 2024 Global AI Study, over 85% of Fortune 500 companies reported using AI in at least one core business function which shows that AI is no longer experimental or optional but rather it's being actively used in daily business operations. Whereas 42% of these companies are using AI for very important strategic decisions made by top executives, such as risk assessment, compliance, and board-level strategy planning.¹

Despite the growing reliance on intelligent systems, legal systems worldwide have not yet evolved at the same pace. The central tension lies in reconciling the autonomy and opacity of AI systems with traditional legal doctrines of corporate liability, governance oversight, and personhood. When AI systems make harmful decisions, such as unlawfully rejecting job applicants, executing discriminatory credit decisions, or causing financial losses, key questions emerge:

- Who is liable when harm results from AI-driven corporate actions?
- How should corporate governance adapt to oversee algorithmic tools?
- Should AI systems be granted some form of legal personhood to streamline liability and regulation?

This paper examines these questions across three dimensions:

1. **Legal Liability:** How courts and regulators assign liability in cases where corporations use AI for core functions;
2. **Corporate Governance:** How boards of directors and corporate officers should manage algorithmic decision-making;
3. **Legal Personhood:** Whether advanced AI should be granted autonomous legal status

¹ PwC, *Global AI Study* (2024), <https://www.pwc.com/gx/en/issues/data-and-analytics/global-ai-study.html>.

or rights in corporate law.

To provide a well-rounded analysis, the paper draws on legal precedents from India, the United States, and the European Union, integrates empirical data from corporate surveys, and explores comparative governance models. It concludes with policy proposals for a balanced regulatory architecture that promotes innovation while safeguarding accountability.

Legal Liability in AI Driven Corporate Decision-Making

Traditional Legal Frameworks of Liability

Liability in corporate law is traditionally assigned to human agents whether individuals, corporate officers, or the corporation itself as a legal person. The concept of vicarious liability holds employers and corporations responsible for acts committed by employees or agents in the course of their duties.²

However, the advent of autonomous AI systems complicates this structure. When AI tools independently make decisions, especially those using machine learning (ML) models whose internal logic may be opaque even to developers (the so-called "black box" problem), it becomes difficult to:

- Identify the actor responsible for the harm;
- Trace the causal link between input, processing, and output;
- Apply doctrines such as mens rea or negligence.

Even though using AI creates complex challenges for legal responsibility, courts and regulators generally don't let the companies to escape the liability. Instead, liability is typically traced back to:

- The **developer or vendor** of the AI system- if defectively designed;
- The **deploying company**- as the party responsible for usage, supervision, and risk mitigation;

² Sir John W. Salmond, *Salmond on Jurisprudence* 346 (P.J. Fitzgerald ed., 12th ed. 1966).

- The **individuals or officers**- who failed to exercise proper oversight.

Case Law: United States and Europe

a. **Mobley v. Workday, Inc. (2024) – U.S. District Court, N.D. California**

The Northern District of California, in a landmark ruling, held that Workday's AI-driven recruitment tool violated federal anti-discrimination laws under Title VII and the Americans with Disabilities Act (ADA).³ The AI system which was used by a large employer, had screened applicants in a manner that disproportionately excluded Black candidates, older applicants, and individuals with mental health conditions.

The court rejected the defense that the discriminatory outcomes were the result of an "autonomous" algorithm and held that: *"The deployment of artificial intelligence tools does not sever the chain of accountability. Employers have a duty to ensure that such tools comply with existing anti-discrimination laws."*⁴

By this decision it is established that AI is considered an instrument of the employer, not a legal actor in its own right. Rather, liability attaches to those who design, procure, and implement the technology and not the algorithm itself.

b. **Dyroff v. Ultimate Software Group, Inc. (2017)**

Although predating the full rise of AI governance frameworks, Dyroff offers insights into how courts approach algorithmic decision-making. The plaintiff alleged that the platform's recommendation algorithm led the victim to join a drug-related group, ultimately leading to death.⁵

While the court ultimately ruled in favour of the platform under Section 230⁶ of the Communications Decency Act, the judgment sparked discussions on the role of algorithmic agency. Critics warned that as algorithms gain autonomy and opacity, doctrines such as CDA immunity may increasingly shield entities from accountability for harm caused by machine-

³ *Mobley v. Workday, Inc.*, 3:23-cv-00770, (N.D. Cal. 2024).

⁴ *Ibid.*

⁵ *Dyroff v. Ultimate Software Group, Inc.*, 934 F.3d 1093 (9th Cir. 2019)

⁶ Communications Decency Act, 47 U.S.C. § 230 (2018).

based recommendations. This case illustrates the gap between the operation of AI systems and how existing laws apply to them.

c. Germany: Google Autocomplete Case (BGH VI ZR 269/12)

The German Federal Court of Justice held Google liable under §823⁷ of the German Civil Code for defamatory autocomplete suggestions generated by its search engine algorithm.⁸

The court emphasized that even though the suggestions were generated automatically by an AI tool, the platform's control over the system, and its failure to remove harmful content upon notification, sufficed to establish liability.

This case laid the groundwork for holding platforms and companies accountable for AI-driven content and decisions, reinforcing the idea that algorithmic outputs are attributable to the entity that deploys them.

Indian Perspective: Nascent Jurisprudence, Growing Risk

India's legal landscape on AI liability remains largely undeveloped, though there is growing interest in the subject. The Information Technology Act, 2000 and the Digital Personal Data Protection Act, 2023 provide some basis for regulatory enforcement, especially in cases involving data misuse, profiling, and discriminatory decision-making. However, courts are yet to deliver clear rulings on corporate AI use.

In the absence of case law, legal scholars have argued for a principle of constructive accountability, where:

- Liability is imposed on the company that benefits from the AI's use;
- Developers and vendors are held liable under product liability or negligent design frameworks;
- Corporate directors are held accountable for failures in oversight, training, or risk

⁷ Bürgerliches Gesetzbuch [BGB] [Civil Code], § 823, translation at https://www.gesetze-im-internet.de/englisch_bgb/ (Ger.).

⁸ Bundesgerichtshof [BGH] [Federal Court of Justice] VI ZR 269/12, Oct. 14, 2014 (Ger.)

disclosure.⁹

The SEBI (Listing Obligations and Disclosure Requirements) Regulations, 2015 may also impose liability if the deployment of an AI tool results in material misstatements or undisclosed risk to shareholders.¹⁰

Corporate Governance in the Age of AI

As AI systems progressively influence decision-making at strategic and operational levels, questions arise regarding how corporate governance frameworks should adapt to manage the associated risks. Corporate governance involves the structures, rules, and processes through which companies are directed and controlled. It is rooted in the core principles of transparency, accountability, and fiduciary duty.

The Board of directors are responsible for safeguarding shareholder interests and ensuring the long-term viability of the company. This includes oversight of technological risk, especially when that technology significantly affects regulatory compliance, market conduct, or stakeholder rights.

Fiduciary Duties and AI

Duty of Care and Skill

The duty of care requires directors to act diligently and prudently in overseeing corporate affairs. This includes:

- Evaluating major investments in AI tools;
- Understanding AI limitations, such as algorithmic bias, data drift, and lack of explainability;
- Conducting adequate due diligence before procurement or deployment.

⁹ Dr. Arvind Malhotra, "AI and the Indian Legal Vacuum," *NLUD J. of Tech. Law*, 2024

¹⁰ Securities and Exchange Board of India (Listing Obligations and Disclosure Requirements) Regulations, 2015, cl. 4 & 34.

In *In re Caremark Int'l Inc. Derivative Litigation*,¹¹ the Delaware Chancery Court emphasized that boards are liable when they fail to implement adequate reporting or information systems to detect and address risk.¹² This standard is particularly relevant to AI, where harm may not arise from deliberate wrongdoing but from system design flaws or training data errors.

A recent report by EY and Harvard Law School (2023) found that only 32% of boards had discussed AI-specific risks, and only 14% had developed formal oversight protocols.¹³ This gap suggests potential Caremark-type liability in the near future.

Duty of Loyalty and Fairness

The duty of loyalty prohibits directors from allowing personal interests to conflict with those of the corporation. In AI contexts, this duty may be implicated where:

- Directors have interests in AI vendor companies;
- Algorithms are used to manipulate markets or suppress whistleblowing

The Facebook–Cambridge Analytica scandal (though not AI-specific) illustrates the reputational and legal fallout from data-driven manipulation of stakeholder rights.

Practical Governance Mechanisms

Corporations are responding to these challenges by instituting various AI governance frameworks. Key components include:

AI Risk Committees

Leading firms like Google, JPMorgan Chase, and Reliance Industries have established board-level or executive-level AI oversight committees. These bodies are tasked with:

- Reviewing procurement of AI tools;

¹¹ *In re Caremark Int'l Inc. Derivative Litig.*, 698 A.2d 959 (Del. Ch. 1996)

¹² *Ibid.*

¹³ EY & Harvard Law School, *Board Oversight of AI: Survey Report* (2023), https://www.ey.com/en_us/board-matters/board-oversight-of-ai.

- Approving AI risk assessments;
- Ensuring compliance with regulatory standards (e.g., GDPR, DPDP Act, or the EU AI Act).

Internal Audit and Impact Assessments

Firms increasingly use algorithmic impact assessments (AIA) to evaluate risks related to fairness, transparency, and data governance before deploying AI systems. These tools are inspired by GDPR's Data Protection Impact Assessments (DPIAs) and are mandated under Canada's Directive on Automated Decision-Making.¹⁴

"Responsible AI" Frameworks

A growing number of corporations are adopting voluntary "responsible AI" charters. For example:

- Microsoft has a six-principle AI framework (fairness, reliability, privacy, inclusiveness, transparency, and accountability);
- Tata Consultancy Services publishes an annual AI Ethics and Governance Report;
- In India, Infosys requires "algorithmic explainability and stakeholder review" for all AI projects impacting employment or customer engagement.

Despite these developments, a 2024 Deloitte study found that over 60% of Indian listed companies lack any formalized board policy on AI oversight.¹⁵

Jurisdictional Perspectives on Governance Standards

India: Evolving Landscape under SEBI and DPDP

- The SEBI LODR Regulations impose obligations on listed companies to disclose

¹⁴ Government of Canada, *Directive on Automated Decision-Making* (2021), <https://www.canada.ca/en/government/system/digital-government/digital-government-innovations/automated-decision-making.html>.

¹⁵ Deloitte India, *AI Governance in the Boardroom: 2024 Outlook* (2024), <https://www2.deloitte.com/in/en/pages/risk/articles/ai-governance-in-boardroom.html>.

material risks, which arguably include AI deployment in core business decisions.

- The Digital Personal Data Protection Act, 2023¹⁶ requires consent-based, fair, and purpose-limited data processing—implicating AI systems that rely on personal data for automated profiling or decision-making.
- However, India lacks a dedicated AI governance statute, and company law does not yet specify director duties in relation to algorithmic oversight.

Scholars have proposed amendments to the Companies Act, 2013 to:

- Add AI-related disclosures under financial statements;
- Require an “AI compliance report” for high-risk sectors;
- Establish director training mandates on AI ethics and law.¹⁷

European Union: Structured Regulation

- The EU AI Act (2025 Draft) introduces a risk-based classification of AI systems:
 - Prohibited: e.g., social scoring systems;
 - High-risk: e.g., credit scoring, biometric systems;
 - Limited risk: e.g., chatbots, spam filters.

High-risk AI systems must undergo conformity assessments, human oversight, and transparency obligations.

Boards of companies deploying AI in HR, finance, or compliance will be held responsible under the AI Act if they fail to implement proper governance safeguards.

United States: Private Sector-Led Standards

- The U.S. lacks a comprehensive AI law, but agencies like the FTC, SEC, and EEOC

¹⁷ Aditya Narula & Riya Mehta, Corporate Law and AI Risk: A Reform Agenda, 16 *NALSAR L. Rev.* 101 (2024).

have issued sectoral guidance:

- FTC: AI use must be “truthful, fair, and non-discriminatory”;
- EEOC: AI hiring tools must comply with Title VII;
- SEC: In 2023, issued guidance on AI-related market manipulation risks in trading firms.

Governance in U.S. corporations is often guided by private initiatives such as the NIST AI Risk Management Framework, which outlines governance structures, measurement tools, and stakeholder engagement strategies.

Empirical Data on Corporate AI Oversight

A 2023 joint study by Harvard Business Review and IBM surveyed 480 corporate board members across 8 countries:

- 87% believed AI posed “serious legal and reputational risk”;
- Only 19% had designated AI oversight at board level;
- 72% had experienced at least one “AI ethics crisis,” such as biased hiring or pricing tools;
- 41% were unsure whether their AI tools complied with existing data or employment laws.¹⁸

The study concluded that boardroom AI literacy is critically low, and that failure to govern AI may trigger the next wave of director liability suits.

¹⁸ IBM & Harvard Bus. Rev., *AI in the Boardroom: 2023 Survey of Directors* (2023), <https://hbr.org/2023/09/ai-in-the-boardroom>.

Recommendations for Strengthening AI Governance

Governance Measure	Description
Board AI Training Mandate	Annual training on AI, ethics, and risk compliance for all board members.
AI Governance Charter	Internal policy detailing roles, responsibilities, and escalation protocols.
Algorithmic Audit Trail	Maintain logs of model changes, training data, and human override events.
Stakeholder Consultation	Include employee, consumer, and regulator input in high-risk AI use cases.
AI Incident Reporting System	Anonymous whistleblowing system for AI misuse or bias.

Legal Personhood and the Autonomous Machine

What is Legal Personhood?

Legal personhood refers to the status of an entity being recognized by law as having rights and obligations, capable of owning property, entering contracts, and being sued or held liable. Traditionally, legal personhood is granted to:

- **Natural persons** (human beings);
- **Juridical or artificial persons** (corporations, trusts, NGOs).

The legal status of personhood is not synonymous with sentience or moral agency; it is a construct used to allocate responsibility and streamline legal operations. The legal fiction doctrine, notably applied in corporate law, allows corporations to act as legal persons despite lacking consciousness or physical form.

This raises the question: can AI systems, particularly those demonstrating high levels of autonomy and decision-making, be granted a similar status?

AI Autonomy and the Challenge to Legal Doctrine

AI systems, especially those using deep learning and reinforcement learning, can perform tasks previously reserved for human cognition, such as:

- Strategic investment decisions;
- Automated hiring and performance evaluations;
- Medical diagnostics and treatment recommendations.

Some of these systems can retrain themselves based on feedback loops (e.g., AlphaZero, GPT-like models), leading to emergent behaviors unforeseen by their creators.

In this context, traditional legal attribution, based on clear lines of intent, negligence, or corporate chain-of-command, becomes complex. If an AI system:

- Makes a discriminatory hiring decision;
- Manipulates financial data;
- Or causes injury via an automated vehicle

Who is liable? The programmer? The deploying company? The board of directors?

The EU's Flirtation with "Electronic Personhood"

In 2017, the European Parliament adopted a resolution on Civil Law Rules on Robotics recommending the creation of a specific legal status for "smart autonomous robots," potentially granting them electronic personhood.¹⁹ The rationale was to:

- Facilitate liability assignment when there is no clear human agent at fault;

¹⁹ European Parliament Resolution 2015/2103(INL) on Civil Law Rules on Robotics, 2017 O.J. (C 252) 239.

- Allow AI systems to hold insurance or funds to compensate victims;
- Create a formal category similar to corporations.

However, this proposal met with widespread criticism:

- Over 150 European AI experts rejected the idea, citing lack of sentience or moral agency.²⁰
- The EU AI Act (2025) ultimately rejected personhood and instead focused on human accountability frameworks.

This marked a normative and regulatory rejection of AI personhood in favour of maintaining human-based liability chains.

U.S. Position: Human-First Accountability

U.S. legal theory remains firmly anchored in human and corporate personhood. Courts have not recognized AI as legal persons, and liability is generally imposed via:

- **Product liability**- if the AI system is defective;
- **Vicarious liability**- if an AI is an agent of a corporation;
- **Negligence**- if human oversight was inadequate.

Notably, in *United States v. Athlone Indus., Inc.* (1987)²¹, the court emphasized that legal personhood for artificial entities depends on statutory intent.²² Since Congress has not extended personhood to AI, such status remains unavailable under U.S. law.

The National AI Initiative Act (2020) and Algorithmic Accountability Act (2019) recommend human-centered audit and control frameworks. The FTC has also reiterated that corporations

²⁰ Open Letter to the European Commission from AI Scholars (2018), <https://futureoflife.org/open-letter/open-letter-to-the-european-commission-on-the-eu-artificial-intelligence-strategy/>.

²¹ *United States v. Athlone Indus., Inc.*, 746 F.2d 977 (3d Cir. 1984).

²² *Ibid.*

remain liable for decisions made by their AI tools.²³

Indian Context: Legal Ambiguity

India does not currently recognize AI systems as legal persons. However, the Indian Supreme Court has shown a willingness to expand the definition of legal personhood in creative ways:

- *Shiromani Gurudwara Prabandhak Committee v. Som Nath Dass* (2000)²⁴: Recognized religious idols as legal persons.
- *Animal Welfare Board v. A. Nagaraja* (2014)²⁵: Referred to animals as entities with inherent dignity and certain rights.

This opens the theoretical door to legal fictions applied to AI. However, The Information Technology Act (2000), The Companies Act (2013), and The Digital Personal Data Protection Act (2023), all assign responsibility to natural or juridical persons, not machines. AI systems in India are thus treated as tools, not agents or persons.

Scholars like Ramanathan & Bhattacharya argue that Indian law may eventually adopt a trust-like model, where high-risk AI systems are managed under fiduciary duties imposed on a trustee (i.e., the deploying corporation).²⁶

Arguments For and Against AI Personhood

Argument	For	Against
Legal Clarity	Personhood allows courts to allocate blame in AI accidents or misconduct.	Blame is better assigned to humans who design, deploy, or supervise AI.

²³ Fed. Trade Comm'n, *Using AI? You're Still Responsible* (2023), <https://www.ftc.gov/business-guidance/blog/2023/04/using-ai-youre-still-responsible>.

²⁴ *Shiromani Gurudwara Prabandhak Committee v. Som Nath Dass*, (2000) 4 S.C.C. 146.

²⁵ *Animal Welfare Board of India v. A. Nagaraja*, (2014) 7 S.C.C. 547 (India).

²⁶ A. Ramanathan & S. Bhattacharya, *Artificial Agents and Legal Responsibility: Indian Frameworks*, 17 *NUJS L. Rev.* 88 (2024).

Argument	For	Against
Functional Equivalence	Corporations are non-sentient yet have personhood; why not AI?	Corporations are composed of humans and are accountable to shareholders and laws.
Insurance Models	AI personhood could enable standalone liability insurance (akin to car insurance).	Corporations can already insure AI tools without granting them personhood.
Ethical Concerns	Avoids scapegoating low-level employees for AI's autonomous decisions.	Risks diluting human accountability and creating "moral offloading" onto machines.
Precedent Creation	Could facilitate treaties or laws about AI rights or regulations.	Could create dangerous moral equivalence between humans and tools.

Alternative Legal Models

Recognizing the complications of AI personhood, scholars and policymakers are exploring alternative frameworks:

- a. **Agency Attribution Model:** Treats AI as an agent whose actions are legally attributed to its principal (e.g., the board or company). Inspired by Section 230 of the Restatement (Second) of Agency.
- b. **Distributed Accountability Model:** Treats AI as an agent whose actions are legally attributed to its principal (e.g., the board or company). Inspired by Section 230 of the Restatement (Second) of Agency.
- c. **Insurance-Funded Compensation Pools:** Similar to the Vaccine Injury Compensation Program, this model uses centralized or private insurance funds to compensate victims of AI errors without needing to determine fault precisely.

- d. Corporate Oversight Doctrine: A refined version of Caremark duties, requiring directors to implement and monitor AI compliance systems.

Philosophical Challenges

Legal personhood also touches on philosophical debates around:

- **Moral agency:** AI lacks intention, empathy, and moral reasoning.
- **Responsibility:** Personhood implies the capacity to act responsibly or be punished.
- **Consciousness:** AI cannot feel guilt, remorse, or understand the consequences of its actions.

Therefore, any attribution of personhood is a legal fiction, useful only insofar as it aids justice and administration. Most scholars agree that AI personhood risks becoming a liability shield for corporations.

Sectoral Applications of AI and Their Legal Implications

Financial Sector: Algorithmic Trading, Credit Scoring and Fraud Detection

The financial services industry is among the earliest adopters of artificial intelligence, harnessing its capabilities to optimize speed, precision, and pattern recognition in high-stakes environments. AI is now embedded across a range of operations including high-frequency algorithmic trading (HFT), automated credit scoring systems, and sophisticated fraud detection protocols. These tools analyse massive datasets at real-time speeds, enabling institutions to react to market shifts, customer behaviour, or fraudulent anomalies far more swiftly than human analysts.

However, the deployment of AI in finance has raised significant legal and regulatory challenges. In the context of HFT, for instance, AI systems may exacerbate volatility or inadvertently engage in market manipulation if not carefully calibrated. Similarly, automated credit scoring models, which learn from historical lending data, risk perpetuating or amplifying biases that disproportionately disadvantage minority or low-income borrowers. The opacity of AI “black-box” models further complicates transparency, making it difficult for regulators or courts to determine intent, causation, or fault when things go wrong. These issues not only pose

systemic risks to the integrity of financial markets but also raise questions of corporate liability and fiduciary oversight.

The 2012 Knight Capital Group incident remains a cautionary example of such risks. In this case, a malfunctioning trading algorithm executed erroneous trades at a massive scale, causing losses of approximately \$460 million within 45 minutes and disrupting broader market stability. Although the U.S. Securities and Exchange Commission (SEC) did not apply any AI-specific laws, since none existed at the time, it held Knight Capital in violation of Rule 15c3-5, which mandates the implementation of appropriate risk management controls for firms with direct access to markets.²⁷ This event galvanized financial regulators to place greater emphasis on pre-deployment testing, algorithmic audits, and fail-safe mechanisms such as "kill switches" capable of instantly halting trading activity.

In response to such incidents, regulatory bodies have developed frameworks to enhance accountability. In the United States, the SEC now requires financial institutions to document algorithmic logic, conduct rigorous testing, and implement real-time monitoring protocols.²⁸ The European Union's Markets in Financial Instruments Directive II (MiFID II), in conjunction with the proposed EU AI Act, classifies algorithmic trading systems as "high-risk," mandating human oversight, auditability, and traceability of trading decisions.²⁹ Similarly, the Securities and Exchange Board of India (SEBI) issued a circular in 2019 requiring pre-approval for the deployment of algorithmic strategies and mandating circuit breakers to contain market volatility.³⁰

From a legal standpoint, firms remain strictly liable for the consequences of AI-driven actions undertaken by their systems. Directors and officers may also be personally liable under fiduciary duty doctrines, such as the Caremark standard articulated by the Delaware Chancery Court, if they fail to ensure the implementation of adequate compliance and oversight systems. In India, the SEBI (Listing Obligations and Disclosure Requirements) Regulations impose

²⁷ *Knight Capital Group Inc.*, Exchange Act Release No. 70694, 107 SEC Docket 1315 (Oct. 16, 2013), <https://www.sec.gov/litigation/admin/2013/34-70694.pdf>.

²⁸ *Consumer Fin. Prot. Bureau, CFPB Report on Credit Scoring and Fair Lending* (2022), <https://www.consumerfinance.gov/data-research/research-reports/>.

²⁹ Directive 2014/65/EU of the European Parliament and of the Council of 15 May 2014 on Markets in Financial Instruments, 2014 O.J. (L 173) 349.

³⁰ Securities & Exchange Board of India, *Framework for Algorithmic Trading by Stock Brokers*, Circular No. SEBI/HO/MIRSD/DOP/CIR/P/2019/68 (May 21, 2019), https://www.sebi.gov.in/legal/circulars/may-2019/framework-for-algorithmic-trading-by-stock-brokers_42975.html.

analogous duties on directors of listed companies. Moreover, the use of AI in credit decisioning implicates anti-discrimination statutes, including fair lending laws in the U.S., and may invite scrutiny from equality commissions or consumer protection bodies in other jurisdictions.

Employment Sector: AI in Hiring, Appraisal, and Workplace Automation

Artificial intelligence has significantly transformed the human resources (HR) domain, reshaping the way companies attract, evaluate, and manage employees. AI-driven recruitment platforms are now used to scan résumés, assess personality traits through video interviews, and even rank candidates based on predicted performance outcomes. Performance appraisal systems likewise employ machine learning to generate employee evaluations, predict attrition, and guide promotion decisions.

While these innovations promise efficiency and consistency, they have also exposed organizations to new legal vulnerabilities. AI systems trained on biased historical data can replicate or reinforce discriminatory patterns, especially in relation to gender, race, age, or disability. This was starkly illustrated by the 2018 controversy surrounding Amazon's internal AI recruitment tool. Initially developed to streamline candidate selection, the system was found to disproportionately favour male applicants, a result of biased input data derived from male-dominated hiring records.³¹ Amazon ultimately discontinued the tool, but the case highlighted how algorithmic systems, if left unchecked, can result in indirect discrimination that contravenes labour and civil rights laws.

In jurisdictions such as the United States, discriminatory outcomes generated by AI recruitment tools may breach Title VII of the Civil Rights Act, exposing employers to potential liability even in the absence of explicit intent. Although no formal litigation arose in the Amazon case, it raised the spectre of class action suits based on systemic bias. In response to such concerns, the New York City Council enacted Local Law 144 in 2023, which mandates that companies conduct independent bias audits of automated employment decision tools and disclose their use to job applicants.³²

The European Union has likewise moved toward a risk-based regulatory framework under its

³¹ Jeffrey Dastin, *Amazon Scraps Secret AI Recruiting Tool That Showed Bias Against Women*, Reuters (Oct. 10, 2018), <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G>.

³² Civil Rights Act of 1964, tit. VII, 42 U.S.C. §§ 2000e–2000e-17 (2024).

AI Act. Employment-related AI tools are classified as “high-risk” and are therefore subject to stringent requirements including human oversight, transparency disclosures, and record-keeping obligations. In contrast, India currently lacks sector-specific AI regulation. Nevertheless, Article 16 of the Indian Constitution guarantees equality of opportunity in public employment, and the Equal Remuneration Act prohibits gender-based discrimination in pay and recruitment. These provisions may be interpreted to extend liability to employers who deploy discriminatory AI systems without adequate safeguards.

From a governance perspective, companies that fail to audit their AI systems or apply fairness-correcting algorithms risk breaching both regulatory expectations and their own corporate HR obligations. Vicarious liability may attach not only for overtly biased outcomes but also for failures to detect or mitigate algorithmic discrimination, particularly where the company has delegated critical employment decisions to automated systems without meaningful human review.

Healthcare Sector: AI in Diagnosis, Risk Assessment, and Resource Allocation

The integration of artificial intelligence into healthcare has profoundly reshaped clinical diagnostics, treatment planning, and administrative efficiency. AI tools are now used to predict disease trajectories, recommend treatment regimens, automate medical billing, and allocate critical resources such as ICU beds or donor organs. In theory, these tools can enhance accuracy, reduce costs, and address clinician shortages. However, their use also presents complex legal and ethical dilemmas, especially when deployed in high-stakes scenarios involving life and death.

The case of IBM Watson for Oncology is illustrative of both the promise and peril of healthcare AI. Marketed as an advanced clinical decision-support system, Watson was intended to assist oncologists in identifying optimal treatment options for cancer patients. However, due to poorly curated training data and overreliance on expert-generated rules, the system occasionally produced unsafe or ineffective recommendations.³³ While IBM maintained that Watson was intended to supplement, not replace, clinical judgment, partner institutions such as MD Anderson Cancer Center faced reputational fallout and potential liability exposure. Legal

³³ Casey Ross & Ike Swetlitz, *IBM's Watson Supercomputer Recommended 'Unsafe and Incorrect' Cancer Treatments, Internal Documents Show*, STAT (July 25, 2018), <https://www.statnews.com/2018/07/25/ibm-watson-recommended-unsafe-incorrect-treatments/>.

responsibility was difficult to assign, due in part to the opacity of the algorithms and the fragmentation of accountability between software vendors and healthcare providers.

From a tort law perspective, healthcare providers remain directly liable for medical malpractice, even when AI systems are involved. Courts have generally taken the position that AI may augment, but not supplant, clinical decision-making; thus, reliance on flawed AI outputs does not absolve physicians of negligence. Furthermore, the principle of informed consent mandates that patients be clearly informed when AI technologies are used in diagnosis or treatment planning, particularly when such tools play a decisive role.

In regulatory terms, many jurisdictions now treat advanced healthcare AI systems as Software as a Medical Device (SaMD), subjecting them to oversight under frameworks such as the U.S. Food and Drug Administration (FDA) regime or the European Union Medical Device Regulation (EU MDR). These frameworks require evidence of safety, efficacy, and transparency, as well as ongoing post-market surveillance.

In India, while no dedicated AI regulation exists for the healthcare sector, indirect coverage is provided through the Clinical Establishments (Registration and Regulation) Act and the Telemedicine Practice Guidelines issued in 2020. Hospitals and diagnostic centers using AI are expected to validate the performance of such systems, maintain proper records, and ensure human oversight. A failure to do so may attract liability under tort law or the Consumer Protection Act. The Supreme Court of India, in *V. Kishan Rao v. Nikhil Super Specialty Hospital* (2010),³⁴ established that once a patient alleges medical negligence, the burden of proof shifts to the hospital to demonstrate due care.³⁵ This precedent could be extended to include AI-enabled negligence, placing an evidentiary burden on institutions that fail to audit or explain algorithmic decisions.

Logistics and Autonomous Transport

The deployment of artificial intelligence in logistics and transportation is rapidly transforming global supply chains, warehousing operations, and vehicular mobility. AI is now integral to route optimization, predictive inventory management, warehouse automation (as exemplified by Amazon Robotics), and the operation of autonomous vehicles and drones. These innovations

³⁴ *V. Kishan Rao v. Nikhil Super Specialty Hosp.*, (2010) 5 SCC 513 (India).

³⁵ *Ibid.*

promise significant gains in efficiency, cost reduction, and safety. However, they also introduce unprecedented legal complexities surrounding attribution of liability, regulatory compliance, and ethical accountability in scenarios involving machine-led decision-making.

A particularly consequential incident in this domain was the 2018 fatal crash involving an autonomous vehicle operated by Uber in Tempe, Arizona. During a test drive, the vehicle struck and killed a pedestrian despite the presence of a human safety driver, who was later charged with negligent homicide for failing to intervene in time. Uber itself avoided criminal prosecution, largely because prosecutors found no evidence of criminal intent attributable to the company.³⁶ Nevertheless, the incident provoked widespread public and legal scrutiny, leading to civil litigation, regulatory investigations, and reputational damage for the firm.

This case illustrates a central challenge in the legal governance of AI in transport: assigning liability in a multi-stakeholder ecosystem. Autonomous vehicles rely on a complex interplay of components, including AI-driven perception modules, software control systems, hardware sensors, and mapping infrastructure, often developed by different entities. When a malfunction occurs, courts must determine whether the fault lies with the vehicle manufacturer, the software developer, the fleet operator, or a combination thereof. Legal doctrines such as negligence, product liability, and strict liability are all potentially applicable, but their adaptation to autonomous systems remains a work in progress. Notably, the traditional negligence standard, based on human behaviour, may be ill-suited for systems that function without direct human control.

In the regulatory sphere, jurisdictions have begun to formulate frameworks to address these challenges. In the United States, the Department of Transportation (DOT) released its “AV 4.0” framework, which emphasizes transparency, data sharing, and corporate accountability in the testing and deployment of autonomous vehicles. However, the federal approach remains largely voluntary and fragmented, with substantial discretion left to state authorities. In contrast, China has proposed more centralized draft regulations assigning primary legal liability to the human operator or deploying company rather than the AI system itself. India currently lacks a comprehensive statutory framework for autonomous transport. However, the Motor Vehicles (Amendment) Act, 2019 increased penalties for road accidents and imposed enhanced liability

³⁶ Daisuke Wakabayashi, *Uber's Self-Driving Car Didn't Know Pedestrians Could Jaywalk*, N.Y. Times (Mar. 19, 2019), <https://www.nytimes.com/2019/03/19/technology/uber-self-driving-car-arizona.html>.

on corporations operating commercial fleets. As autonomous vehicle trials begin in Indian metropolitan areas, the absence of a clear legal regime may leave both victims and companies uncertain about their rights and obligations.

Legal exposure for firms operating in the autonomous logistics space may arise from various sources, including tort claims, product defect litigation, insurance disputes, and regulatory fines. The uncertainty surrounding causation and foreseeability in AI-related harm makes litigation unpredictable and may increase pressure for the establishment of sector-specific liability standards or insurance pools. Additionally, the deployment of drones and unmanned aerial vehicles raises further concerns under airspace regulation, privacy law, and public nuisance doctrine.

To mitigate these risks, companies are increasingly investing in governance mechanisms such as real-time fail-safes, redundancy systems, black-box recorders, and continuous validation of algorithmic performance. Nonetheless, the law remains underdeveloped in this area, and jurisprudence will likely evolve in parallel with technological maturity and societal acceptance of autonomous mobility.

Cross-Sectoral Themes and Comparative Legal Risks Analysis

The foregoing sectoral analysis reveals a number of recurring legal and regulatory themes that transcend individual industries. While the nature of AI applications varies across finance, employment, healthcare, and logistics, common concerns arise in relation to transparency, accountability, bias mitigation, and liability attribution.

In the financial sector, regulatory focus is centered on market stability, algorithmic transparency, and fiduciary responsibility. Legal risks often take the form of sanctions from oversight bodies such as the SEC or SEBI, as well as shareholder derivative actions based on failures of board-level oversight. In employment, the dominant concerns involve algorithmic discrimination, violations of civil rights statutes, and corporate exposure to class-action litigation. Here, governance mechanisms tend to emphasize fairness audits, transparency disclosures, and maintaining human review in AI-assisted hiring processes.

Healthcare, by contrast, is governed by strong normative imperatives around patient safety, medical ethics, and the sanctity of informed consent. Legal liability in this domain is typically

framed in terms of malpractice and product safety, with a growing emphasis on classifying AI systems as medical devices subject to regulatory approval. Finally, in the logistics and transport sector, the core legal challenge lies in attributing responsibility across complex socio-technical systems, with regulators grappling to adapt traditional doctrines to machine autonomy.

Across these sectors, empirical studies confirm a growing gap between AI deployment and governance maturity. According to McKinsey's Global AI Survey (2024), only 30% of firms conduct regular audits of AI outputs for fairness, explainability, or regulatory compliance.³⁷ The Stanford AI Index Report (2025) similarly noted that enforcement actions related to AI-based decision-making tripled between 2021 and 2024, underscoring a shift from aspirational ethics to active legal oversight.³⁸ In India, a 2023 NASSCOM study found that more than 75% of companies lack formal AI governance policies, leaving them vulnerable to compliance failures and reputational damage.³⁹

To navigate this emerging legal terrain, firms must invest in cross-functional AI governance architectures that combine legal compliance, ethical principles, technical safeguards, and board-level accountability. Sector-specific guidelines will remain essential, but as AI continues to blur boundaries between domains, regulators and courts alike will be called upon to develop coherent frameworks capable of managing its multifaceted risks.

Corporate Governance and AI: The Board's Emerging Role

Board-Level Responsibilities in the AI Era

In the digital age, corporate boards bear a pivotal role in ensuring that artificial intelligence (AI) deployment is not only technologically sound but also ethically grounded and legally compliant. Directors are bound by fiduciary duties to act in good faith, exercise due care, and prioritize the interests of the company and its stakeholders. In the context of AI, these duties translate into an obligation to:

³⁷ McKinsey & Co., *The State of AI in 2024: Generative AI's Breakout Year* (May 2024), <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-in-2024-generative-ais-breakout-year>.

³⁸ Stanford Institute for Human-Centered Artificial Intelligence (HAI), *Artificial Intelligence Index Report 2025* (Mar. 2025), <https://aiindex.stanford.edu/report/>.

³⁹ NASSCOM, *State of AI Adoption in India 2023: Enterprises and the Road Ahead* (2023), <https://www.nasscom.in/knowledge-center/publications/state-ai-adoption-india-2023>.

- Ensure that AI systems are deployed in accordance with legal norms and ethical principles.
- Align AI initiatives with the company's strategic objectives and risk appetite.
- Supervise risk management frameworks that address AI-related risks such as bias, discrimination, privacy breaches, and operational failures.
- Establish clear lines of accountability for AI-driven decisions and ensure remediation mechanisms are in place.

Neglecting these responsibilities could expose directors to liability under various jurisdictions. In the United States and India, breach of the duty of care can give rise to derivative litigation. In the United Kingdom, directors may face liability under Sections 172 to 174 of the Companies Act 2006 for failing to act with due diligence and care. In cases of gross negligence or regulatory violations, corporate criminal liability may also ensue.

The Caremark Standard and AI Governance (U.S.)

A foundational precedent in board oversight is the Delaware Chancery Court's ruling in *In re Caremark International Inc. Derivative Litigation* (1996), which established that directors must implement and monitor systems to ensure legal compliance and identify operational risks. This doctrine is especially relevant in the era of AI, where opaque algorithmic processes can introduce hidden liabilities.

Applied to AI governance, the Caremark standard implies that directors must:

- Ensure that internal reporting systems are capable of identifying ethical and legal risks associated with AI use.
- Monitor ongoing AI operations for signs of misuse, bias, or regulatory breach.
- Act decisively when presented with red flags relating to algorithmic harm.

The scope of this duty was reaffirmed in *Marchand v. Barnhill* (2019), where the Delaware Supreme Court held that directors of Blue Bell Creameries failed to monitor a listeria outbreak, thereby breaching their oversight duties. The court emphasized that boards must pay special

attention to "mission-critical" risks—a category increasingly applicable to AI in sectors like healthcare, finance, and logistics. Companies heavily reliant on AI must therefore consider establishing AI-specific risk committees, adopt robust reporting mechanisms, and ensure access to technical expertise at the board level.

Regulatory Expectations: EU and US

In Europe, the AI Act (2025) imposes stringent governance obligations on companies deploying high-risk AI systems. Key provisions include:

- Mandated internal control frameworks to ensure human oversight.
- Obligations to conduct impact assessments before and during AI deployment.
- Requirements for appointing AI Compliance Officers and reporting serious incidents to regulators.

Failure to comply can attract fines up to €35 million or 7% of global turnover, making board-level accountability a legal and financial imperative.

In the United States, while a federal AI statute remains in progress, the Securities and Exchange Commission (SEC) issued a rulemaking in 2023 that requires public companies to disclose:

- AI-related risk exposure.
- Governance structures overseeing AI deployment.
- Material incidents involving algorithmic errors, bias, or cyber threats.

These disclosures place explicit responsibilities on audit committees and board members to maintain visibility over AI activities and respond proactively to governance gaps.

India: Corporate Law and AI Oversight

Although India lacks AI-specific corporate governance mandates, existing legal frameworks impose indirect obligations. Section 166 of the Companies Act, 2013, compels directors to act with due care, skill, and diligence. Schedule IV emphasizes the role of independent directors in overseeing risk management, while SEBI's (LODR) Regulations, 2015 mandate listed

companies to adopt enterprise risk frameworks—which, by implication, encompass AI-related risks.

Judicial precedent further clarifies the board's accountability. In *Sunil Bharti Mittal v. CBI* (2015), the Supreme Court held that directors could be held liable for corporate acts if such acts were conducted with their knowledge and under their control. Applied to AI, this suggests directors could face liability if they knowingly allowed AI systems with foreseeable risks to operate unchecked.

Key AI Governance Tools for Boards

To navigate these emerging expectations, boards are deploying specialized tools, including:

- **AI Ethics Charter:** A declaration of principles on fairness, accountability, and data privacy.
- **AI Risk Committee:** A dedicated subcommittee to evaluate algorithmic risks and oversee compliance efforts.
- **Algorithmic Impact Assessment (AIA):** A structured evaluation of the legal, ethical, and social impact of AI tools, often mandated by regulatory bodies like the EU.
- **Explainability Reports:** Documentation that elucidates how AI models reach decisions, particularly crucial in regulated sectors.
- **Incident Reporting Mechanisms:** Internal systems that enable employees to flag AI-related errors or misconduct.
- **Director Training Programs:** Educational initiatives to build AI literacy and governance capacity at the board level.

Comparative Overview: Governance Expectations

Jurisdiction	Board Duties	AI Regulation	Enforcement Risk
U.S.	Caremark duties; SEC disclosures	Algorithmic Accountability Act (proposed); sectoral laws	Moderate to high (shareholder suits, SEC action)
EU	Human oversight; conformity and incident reporting	EU AI Act; GDPR	High (AI Act fines, DPA investigations)
India	Sec. 166 Co. Act; SEBI LODR; consumer law	Draft Digital India Act; DPDPA	Emerging (judicial + regulatory)

These scenarios underscore the growing necessity for boards to proactively engage with AI governance. As AI becomes integral to business operations, directors must adopt a digitally prudent approach that ensures compliance, manages risk, and aligns innovation with stakeholder interests.

Conclusion

The integration of artificial intelligence into corporate ecosystems has outpaced the evolution of legal and governance frameworks. This study has demonstrated that AI systems, while delivering efficiency and scale, introduce unprecedented challenges to established doctrines of liability, corporate oversight, and legal personhood. Courts in various jurisdictions have consistently emphasized that accountability must rest with human actors and institutions, regardless of the autonomy or complexity of the AI tools involved.

Corporate directors face a growing fiduciary burden to proactively govern AI systems, implement internal controls, and ensure compliance with sectoral regulations. The adoption of algorithmic audit trails, ethics charters, and AI-specific risk committees is no longer optional but essential for managing the legal, reputational, and operational risks associated with automated decision-making.

Efforts to conceptualize AI as a legal person have been largely rejected in favour of distributed accountability models that retain human agency at the core of corporate legal responsibility. Granting AI systems personhood, while theoretically appealing to some, risks eroding accountability and shifting liability away from those in power.

Across sectors, finance, employment, healthcare, logistics, the regulatory landscape remains fragmented. Yet a common pattern emerges: inadequate governance structures, low board-level AI literacy, and a reactive approach to oversight contribute significantly to legal exposure and stakeholder harm. Jurisdictions such as the European Union have begun to codify AI-specific obligations through the AI Act, whereas India and the United States are progressing through piecemeal or sector-led initiatives.

Ultimately, the future of AI governance in corporations depends on embedding ethical safeguards, mandatory legal compliance systems, and informed boardroom leadership. This paper calls for a structured regulatory architecture that combines sectoral best practices, cross-border legal harmonization, and corporate accountability, ensuring that innovation is not achieved at the expense of legality, fairness, or human dignity.