

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

Volume 8 | Issue 3

2025

© 2025 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.

Algorithmic Offenders: When Corporations Commit Cybercrimes via AI

KHUSH DALBIR¹

ABSTRACT

The increasing reliance on Artificial Intelligence (AI) by corporations has revolutionised business operations, decision-making, and digital engagement. However, it has also introduced unprecedented avenues for corporate misconduct and cybercrime. From automated phishing tools and unethical data mining to algorithmic manipulation in financial markets and targeted disinformation, AI systems are now being weaponised—either deliberately or negligently—by corporations to gain unlawful advantages. This paper examines the phenomenon of AI-enabled corporate cybercrimes, identifies the regulatory and legal gaps in the Indian context, and explores the challenge of assigning liability when algorithms become the offenders. Through doctrinal analysis, comparative legal frameworks, and real-world case illustrations, the study proposes a robust policy and legal structure to prevent and regulate such emerging threats. It argues that India must move swiftly to enact algorithmic accountability legislation and empower regulatory agencies to audit, monitor, and penalise AI-based corporate misconduct.

I. INTRODUCTION

Artificial Intelligence has become a core component of modern corporate infrastructure, powering everything from financial forecasting and human resource management to customer engagement and cybersecurity. While the benefits of AI are undeniable, its increasing autonomy and opacity have led to new forms of corporate misconduct. Cybercrimes, which were once limited to human-driven actions, are now being orchestrated through algorithms—whether for profit maximisation, competitive sabotage, or regulatory evasion.

In India, the digital economy is growing at an exponential rate, with AI adoption by private companies becoming mainstream. However, the legal system remains inadequately equipped to deal with the implications of AI as an agent of crime. Unlike traditional corporate fraud, AI-driven misconduct often operates within legal grey zones, where the intent is programmed, the execution is autonomous, and accountability is diffused. This research aims to explore how corporations are exploiting AI tools to commit cyber offences, the challenges this poses for enforcement, and the reforms needed to establish effective regulatory oversight.

¹ Author is a LL.M (Cyber Law) Student at IILM University, Greater Noida, India.

II. DECODING AI-ENABLED CORPORATE CYBERCRIMES

Artificial Intelligence has transformed corporate operations by enabling data-driven decision-making, automating complex tasks, and personalising consumer engagement. However, in parallel with these advancements, a darker dimension of AI use has emerged—its exploitation in furthering corporate cybercrimes. These crimes differ from traditional cyber offences not only in scale and sophistication but also in the obscurity of agency and accountability. Corporates may intentionally or recklessly deploy AI tools that engage in deceptive, invasive, or fraudulent conduct, blurring the boundaries between operational efficiency and digital malfeasance.

Below are some key categories of AI-facilitated corporate cybercrimes:

Automated Phishing and Social Engineering: AI models trained on communication datasets can mimic human conversation and generate contextually accurate messages. Corporations can exploit these capabilities to launch automated phishing campaigns that impersonate stakeholders or regulators to obtain sensitive information. AI-driven chatbots may also be programmed to manipulate users into divulging passwords, financial details, or confidential business data.

Data Scraping and Predictive Profiling: AI systems can mine and process enormous amounts of publicly accessible and private user data—often through unsanctioned scraping of websites, social media platforms, or mobile applications. This data is then used for behavioural profiling, targeted manipulation, or discriminatory decision-making in employment, credit lending, and insurance sectors. Such practices may breach data protection laws and consumer rights.

Dark Pattern Marketing and Algorithmic Manipulation: Corporations have begun using AI to design 'dark patterns'—interface designs that trick users into making unintended choices, such as subscribing to services, sharing data, or clicking on misleading links. These manipulative practices rely on psychological profiling and optimisation algorithms to exploit user vulnerability, violating consumer protection norms.

High-Frequency Trading and Financial Market Manipulation: In financial sectors, AI is used for algorithmic trading that can influence market trends through rapid, high-volume transactions. Some corporations manipulate these algorithms to create false signals, engage in spoofing, or pump-and-dump schemes, distorting genuine investor behaviour and destabilising financial markets.

AI-Driven Surveillance and Privacy Violations: Companies employ AI-powered surveillance systems—facial recognition, keystroke tracking, sentiment analysis—to monitor employees or customers, often without informed consent. These tools are used to enforce productivity, prevent leaks, or track consumer behaviour. When done covertly or excessively, they constitute cyber offences and contravene privacy laws.

Generative AI and Intellectual Property Theft: Corporates increasingly use generative AI models to replicate or modify content—text, code, designs—potentially infringing upon third-party intellectual property. When used to reverse-engineer patented technologies or plagiarise proprietary content, such use amounts to digital piracy and corporate espionage.

Exploitation of Vulnerable Populations through AI: Corporations may use AI tools to identify and exploit vulnerable demographics—children, the elderly, or marginalised groups—for predatory marketing or disinformation campaigns. These activities, though automated, have real-world consequences and may result in violations of consumer law, child protection statutes, and human rights.

In each of these instances, AI is not merely a passive tool but an active mechanism of offence—either because it was designed to behave that way or because the corporation failed to implement adequate oversight. The complexity of AI's role and its capacity to learn and adapt make detection and regulation of such corporate misconduct especially challenging. Therefore, decoding these behaviours is the first step in shaping appropriate legal, ethical, and technological safeguards.

III. LEGAL FRAMEWORK IN INDIA

India's current legal framework is ill-equipped to address the emerging threats posed by AI-enabled corporate cybercrimes. The country's regulatory infrastructure is largely reactive, sector-specific, and predicated on traditional models of human agency and intent—features that are fundamentally disrupted by the deployment of autonomous and opaque AI systems.

Several existing statutes partially address cyber offences and corporate misconduct, but none directly contemplate the legal complexities introduced by artificial intelligence:

Information Technology Act, 2000 (IT Act): This is India's primary legislation governing cyber activities. While provisions like Section 43A (compensation for failure to protect data), Section 66 (hacking), Section 66C (identity theft), Section 66E (violation of privacy), and Section 72A (unauthorised disclosure of personal information) offer some protection against digital misconduct, the IT Act does not define or regulate AI systems, nor does it recognise

the liability of autonomous decision-making by corporate-deployed algorithms. The Act is also silent on algorithmic manipulation, AI-generated content, or black-box opacity.²

Digital Personal Data Protection (DPDP) Act, 2023: Recently enacted, this law focuses on the processing of personal data by data fiduciaries and processors. It contains provisions around data minimisation, purpose limitation, and user consent. However, it does not impose mandatory auditability for AI tools, nor does it address algorithmic bias, discrimination, or the implications of AI-driven profiling. Moreover, the DPDP Act does not mandate impact assessments for high-risk automated systems, leaving a critical gap in accountability.³

Bharatiya Nyaya Sanhita, 2023 (BNS): Traditional offences like cheating (Section 316), criminal breach of trust (Section 315), and criminal conspiracy (Section 62) may be invoked in some cases of corporate cyber misconduct. However, these provisions are framed with human actors in mind and typically require demonstrable intent (*mens rea*) or knowledge. This becomes challenging when misconduct is executed through autonomous AI systems acting on trained algorithms or adaptive behaviour. The lack of direct human intervention in such acts often renders attribution under these sections legally ambiguous, thereby weakening the prosecutorial foundation in AI-related corporate offences.⁴

Securities and Exchange Board of India (SEBI) Regulations: In the financial sector, SEBI has frameworks to regulate algorithmic trading by market participants. However, these are primarily technical and administrative, with minimal emphasis on criminal liability in cases of algorithmic manipulation or fraud perpetrated through AI bots.⁵

Competition Act, 2002: The Act prohibits anti-competitive conduct such as cartelisation and abuse of dominant position. Yet, there is no guidance on algorithmic collusion, where companies use AI to coordinate pricing strategies without direct communication—a growing concern in global competition law discourse.

Consumer Protection Act, 2019: The Act covers unfair trade practices and misleading advertisements, and under the e-commerce rules, it imposes certain duties on digital platforms. However, there is no specific recognition of dark pattern marketing, algorithmic deception, or manipulative personalisation techniques driven by AI.

Additionally, India lacks:

² Information Technology Act, 2000 (India).

³ Digital Personal Data Protection Act, 2023 (India).

⁴ Government of India, *The Bharatiya Nyaya Sanhita, 2023*, Ministry of Law and Justice, Act No. 45 of 2023 (Published on December 25, 2023).

⁵ Securities and Exchange Board of India, “Algorithmic Trading Framework” (2021).

- A centralised statutory definition of AI and its classification by risk levels
- Mandates for AI audit trails, explainability, or transparency in high-risk sectors
- Procedural rules for evidentiary standards in cases involving autonomous AI-generated actions

Despite growing AI adoption across critical sectors like finance, healthcare, retail, and logistics, Indian law continues to treat AI tools as mere instruments rather than independent operational agents. This doctrinal gap allows corporations to evade liability by attributing blame to technological systems, third-party developers, or unintended consequences.

The situation is further complicated by the absence of sectoral regulators' jurisdiction over AI-based misconduct. While CERT-In is empowered to handle cybersecurity incidents, it lacks the mandate to investigate algorithmic bias, dark pattern marketing, or AI-generated disinformation.

In sum, India's legal framework is fragmented, outdated, and insufficiently anticipatory to address the nuanced legal questions posed by corporate cybercrimes involving artificial intelligence. There is an urgent need for legislative reform that integrates AI-specific responsibilities, liability regimes, and enforcement protocols into India's broader cyber and corporate legal architecture.

IV. CHALLENGES OF LEGAL ATTRIBUTION

One of the most formidable hurdles in addressing AI-enabled corporate cybercrimes is the issue of legal attribution. Traditional criminal and corporate liability frameworks are premised on clear lines of causality, intent, and human agency. However, AI technologies—especially those that learn and evolve over time—blur these lines and challenge existing doctrines of responsibility, making attribution of legal culpability highly complex.

Obscured Intent (Mens Rea): Establishing the mental element of a crime—intent or knowledge—is fundamental to criminal liability. In cases of AI-induced harm, especially those involving machine learning or unsupervised algorithms, intent becomes difficult to locate. The AI may “learn” harmful behaviour over time, or follow logical paths that were neither foreseen nor explicitly endorsed by the corporation. The lack of intentional or negligent human input in the algorithm's final action makes traditional frameworks of mens rea insufficient.

The ‘Black Box’ Problem: Many AI systems, particularly deep learning models, are inherently opaque. Their decision-making processes are not easily interpretable, even by their

own developers. This phenomenon—often referred to as the ‘black box’—poses significant difficulties in auditing the algorithm’s actions or producing legally admissible explanations of why a system acted in a certain way. Courts and regulators may struggle to establish causation or foreseeability when they cannot understand the underlying processes.

Fragmented Responsibility and the Chain of Custody: AI systems in corporate environments are rarely the result of a single actor’s work. Developers, data scientists, third-party vendors, IT teams, and corporate leadership all play roles in training, deploying, and overseeing these systems. Without a clear legal doctrine to assign degrees of responsibility, accountability is diluted. This makes enforcement and penalty application inconsistent, and opens up space for corporate actors to deflect liability by blaming system design, external consultants, or even the AI itself.

Proxy Liability and Human Shielding: Corporations may exploit the legal ambiguity of AI to use it as a proxy for decisions they wish to distance themselves from—such as discriminatory hiring, biased credit scoring, or manipulative marketing. By framing these outcomes as “algorithmic,” corporations attempt to sidestep human liability and shield decision-makers from scrutiny.

Lack of Regulatory Forensics and Technical Expertise: Enforcement agencies and regulators in India, such as CERT-In or state cyber cells, often lack the technical capacity to reverse-engineer AI outputs or conduct forensic analysis on advanced machine learning systems. This limits their ability to investigate cybercrimes involving AI, especially those orchestrated by well-resourced corporate entities with sophisticated technological infrastructure.

Difficulty in Establishing Corporate Criminal Liability: The current legal regime does not explicitly provide for corporate criminal liability in the context of AI. While the principle of vicarious liability may apply in civil contexts, applying it to criminal acts involving AI introduces uncertainties. Who within the corporation can be held responsible for the AI’s harmful actions—CEOs, CTOs, compliance officers, or the board? Without codified roles and thresholds, enforcement becomes subjective and arbitrary.

Jurisdictional and Cross-Border Complexities: AI systems are often hosted on cloud infrastructures, maintained by foreign vendors, or trained on datasets sourced globally. When such systems are involved in cross-border cybercrimes—such as data theft or disinformation—the questions of jurisdiction, cooperation between enforcement agencies, and evidentiary standards further complicate attribution.

The challenge of attribution in AI-enabled corporate cybercrime is not merely technical—it is a structural problem embedded in the assumptions of the current legal framework. To address this, India must not only modernise its laws but also rethink the foundations of liability, causation, and intent in the age of autonomous systems. Legal doctrines must evolve to incorporate concepts like algorithmic foreseeability, functional responsibility, and risk-based liability allocation, enabling a fair and enforceable regime of corporate accountability in the digital era.

V. GLOBAL COMPARISONS

As the legal and ethical implications of artificial intelligence unfold across the world, several jurisdictions have begun to develop regulatory frameworks and enforcement tools to address AI-enabled misconduct, including corporate cybercrimes. These global efforts provide valuable insights for India as it considers reforming its domestic regulatory architecture. While approaches vary in stringency, scope, and philosophical underpinnings, they all reflect an acknowledgment of AI's disruptive potential and the urgent need for accountability.

European Union – AI Act and AI Liability Directive: The European Union has been at the forefront of regulating artificial intelligence. The proposed AI Act (2021) classifies AI systems by risk levels—minimal, limited, high-risk, and unacceptable—and imposes obligations such as transparency, human oversight, and accuracy on high-risk AI systems.⁶ Of particular relevance is the AI Liability Directive,⁷ which aims to ease the burden of proof in cases where AI causes harm, enabling claimants to seek redress without needing to fully decipher the black-box nature of algorithms. This includes situations where corporations use AI systems that engage in discriminatory, deceptive, or manipulative conduct. The EU's approach combines ex-ante risk management with ex-post liability, offering a robust framework for regulating corporate AI use.

United States – Sectoral and Agency-Driven Oversight: In the United States, AI regulation is fragmented but proactive within specific agencies. The Federal Trade Commission (FTC) has issued guidance and taken enforcement actions against unfair or deceptive practices involving AI, including algorithmic bias and facial recognition misuse.⁸ The Securities and Exchange Commission (SEC) monitors algorithmic trading and has issued alerts on AI-

⁶ European Commission, "Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)," COM/2021/206 final.

⁷ European Commission, "Directive on Adapting Non-Contractual Civil Liability Rules to Artificial Intelligence (AI Liability Directive)," COM/2022/496 final.

⁸ Federal Trade Commission (US), "Aiming for Truth, Fairness, and Equity in Your Company's Use of AI," April 2021.

induced market manipulation.⁹ Although the U.S. lacks a comprehensive AI law, initiatives such as the Algorithmic Accountability Act and state-level bills (e.g., California Consumer Privacy Act) indicate a growing legislative focus. Importantly, U.S. regulators have emphasised corporate responsibility for algorithmic outcomes, including failures arising from flawed training data or lack of oversight.

China – Centralised and Security-Focused Regulation: China’s framework of AI regulation is grounded in societal stability and national security. The Cybersecurity Law (2017), Data Security Law (2021), and Personal Information Protection Law (2021) collectively form a strict regime over data usage, algorithmic processing, and cross-border transfers. In 2022, China introduced rules for algorithmic recommendation services, mandating registration, transparency, and user control features.¹⁰ The Cyberspace Administration of China (CAC) plays a pivotal role in enforcing compliance. Corporations using AI for manipulation or surveillance without adequate safeguards face administrative penalties and potential bans. China’s model is marked by strong state oversight and regulatory enforcement, with less emphasis on due process or civil liberties.

Singapore – Balanced Innovation and Risk Management: Singapore’s Model AI Governance Framework (first released in 2019 and updated in 2020) promotes responsible AI by outlining principles such as explainability, fairness, and human-centricity.¹¹ While non-binding, the framework has influenced sectoral guidelines and is supported by tools such as AI Verify, a self-assessment mechanism for AI deployment. The Monetary Authority of Singapore (MAS) also has risk management guidelines for the use of AI in financial services. Singapore’s approach exemplifies a cooperative, industry-led model that balances innovation with ethical use.

United Kingdom – Adaptive Regulation and Regulatory Coordination: The UK’s strategy, detailed in its 2022 policy paper "Establishing a Pro-Innovation Approach to Regulating AI," focuses on principles rather than prescriptive rules. It encourages individual regulators (e.g., ICO, FCA) to issue AI-specific guidance within their sectors. The UK's approach is adaptive and seeks to avoid overregulation while emphasising accountability,

⁹ Securities and Exchange Commission (US), "Office of Compliance Inspections and Examinations: Observations on Cybersecurity and Resilience Practices" (January 2020).

¹⁰ Cyberspace Administration of China, "Regulations on the Management of Algorithmic Recommendation Services," 2022.

¹¹ Singapore Personal Data Protection Commission, "Model AI Governance Framework," Version 2.0, 2020.

transparency, and explainability. Pilot projects on AI in the judicial and tax systems illustrate the UK's willingness to test regulatory sandboxes and promote digital due process.¹²

Key Takeaways for India:

Risk-Based Classification: India can adopt a tiered risk classification model like the EU to determine obligations on AI deployment by corporations.

Strict Corporate Liability: Enforcement actions like those by the FTC and SEC demonstrate how sectoral regulators can hold corporations accountable for AI misuse.

Algorithmic Transparency and Audits: Mandating explainability and auditability, as seen in Singapore and the EU, can enhance public trust and legal accountability.

Dedicated AI Regulator or Coordination Mechanism: India can consider establishing a central AI authority or empower existing regulators with AI-specific mandates, akin to the UK's model.

Mandatory Registration and Disclosure of AI Tools: Inspired by China's approach, corporations could be required to register high-risk AI systems and disclose their functionality.

While each jurisdiction reflects different cultural and governance priorities, they all converge on the need to treat AI as a legal subject of regulation, not just a tool. For India, these global experiences offer a roadmap to build a nuanced, enforceable, and innovation-friendly legal framework for AI governance in the corporate and cybercrime context.

VI. ETHICAL AND GOVERNANCE CONCERNS

Beyond legal liability, the use of Artificial Intelligence by corporations to facilitate cybercrimes raises profound ethical and governance issues. As AI systems increasingly influence decision-making in areas like finance, marketing, surveillance, and employment, questions arise about fairness, transparency, discrimination, and exploitation. These concerns are magnified when corporate incentives—often driven by profit maximisation—lead to the deployment of AI in ways that harm users, workers, or society at large.

Algorithmic Bias and Discrimination: One of the most documented ethical issues in AI systems is the risk of bias in decision-making. AI algorithms trained on historical or unbalanced data can reinforce societal prejudices in hiring, credit scoring, law enforcement, and healthcare. When corporations deploy such biased systems, they contribute to systemic

¹² United Kingdom Department for Science, Innovation and Technology, "Establishing a Pro-Innovation Approach to Regulating AI," Policy Paper (March 2023).

discrimination while insulating themselves from accountability by attributing outcomes to “machine decisions.” This undermines equality and due process.¹³

Lack of Explainability and Informed Consent: AI systems used by corporations often operate without transparency, leaving users unaware of how decisions affecting them are made. This violates the principle of informed consent, particularly when consumers are subject to algorithmic pricing, personalised content delivery, or automated rejection in job applications or loan processing. Ethical AI governance requires explainable AI (XAI), where decisions can be understood, challenged, and audited.¹⁴

Surveillance Capitalism and Privacy Erosion: Many corporations employ AI systems for data harvesting, behavioural tracking, and targeted advertising. These tools often collect and process user data without explicit consent or clear limitations, amounting to a form of surveillance capitalism. The ethical dilemma lies in using user vulnerability for monetisation, especially when this data is repurposed for predictive policing, political profiling, or discriminatory filtering.¹⁵

Manipulation and Psychological Targeting: AI models can exploit cognitive biases to manipulate users through personalised nudges, dark patterns, or emotionally resonant content. This is especially dangerous when used to encourage compulsive behaviour, misinformation spread, or consumer purchases based on fear or urgency. The ethical problem is not just misuse, but the deliberate design of AI to exploit human psychology.¹⁶

Governance Gaps Within Corporations: Despite the growing use of AI, few Indian companies have instituted internal governance frameworks to oversee the ethical deployment of such systems. There is often no designated body to audit algorithms, assess risks, or ensure compliance with ethical standards. The absence of AI Ethics Committees, Chief AI Ethics Officers, or responsible innovation protocols allows corporate actors to prioritise profit over public interest without checks or consequences.

Inequitable Impact and Digital Marginalisation: AI systems may disproportionately harm marginalised communities who are less likely to be represented in training datasets or more likely to be misclassified. When AI systems are deployed in areas like welfare distribution,

¹³ Sandra Wachter, Brent Mittelstadt, and Luciano Floridi, “Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation” 7 *International Data Privacy Law* 76 (2017).

¹⁴ Jenna Burrell, “How the Machine ‘Thinks’: Understanding Opacity in Machine Learning Algorithms” *Big Data & Society* 1 (2016).

¹⁵ Nasscom-DSCI, “AI and Cybersecurity: Realising the Promise while Managing the Risks,” Industry Report (2021).

¹⁶ Mozilla Foundation, “The Rise of AI in Corporate Governance,” White Paper (2022).

policing, or employment screening, the consequences of inaccuracy are more severe for vulnerable populations, exacerbating inequality.¹⁷

Automated Harm Without Accountability: The ethical conundrum deepens when autonomous AI systems cause harm but no individual can be held directly accountable. This dilution of human responsibility creates a governance vacuum. Ethical frameworks must incorporate the principle of “meaningful human oversight,” ensuring that decisions with significant consequences are not left to machines alone.¹⁸

In light of these challenges, India must promote a culture of ethical AI use through a combination of legal mandates and voluntary governance. Corporations should be encouraged—if not required—to adopt internal ethical audits, publish AI impact assessments, and train employees on responsible AI use. At the policy level, a national framework on AI ethics, aligned with constitutional values and international best practices, would help establish a moral compass for AI deployment in corporate India.

Ultimately, ethical governance is not an abstract ideal but a necessary pillar of trust in digital transformation. Without it, AI risks becoming a powerful instrument of exploitation rather than empowerment.

VII. RECOMMENDATIONS

Given the multifaceted challenges posed by AI-enabled corporate cybercrimes—spanning legal gaps, attribution dilemmas, ethical failures, and weak institutional oversight—India must adopt a multi-pronged strategy to mitigate risks and establish accountability. The following recommendations aim to provide a comprehensive policy and legal roadmap:

Enact a Comprehensive Algorithmic Accountability Law: India needs a dedicated legislation that clearly defines algorithmic misconduct, imposes liability on corporations for harms caused by AI systems, and mandates transparency obligations. This law should cover both civil and criminal liabilities, introduce a classification system for high-risk AI, and provide remedies for victims of AI-induced harm.

Amend Existing Legal Frameworks (IT Act and DPDP Act): The Information Technology Act, 2000 and the Digital Personal Data Protection Act, 2023 should be amended to address AI-specific challenges. Key additions should include:

- Mandatory impact assessments for high-risk AI systems

¹⁷ NITI Aayog, “Responsible AI for All,” Discussion Paper, June 2021.

¹⁸ Council of Europe, “Recommendation CM/Rec(2020)1 on the Human Rights Impacts of Algorithmic Systems,” 2020.

- Recognition of automated decision-making as a basis for liability
- Provisions for algorithmic transparency, auditability, and non-discrimination
- Explicit obligations for human oversight of critical AI functions

Establish a Central Regulatory Body for AI Oversight: A specialised AI regulatory authority or an empowered division within existing bodies like MeitY should be created. This body should be responsible for:

- Accrediting AI developers and platforms
- Conducting audits and inspections
- Investigating AI-related violations
- Issuing compliance guidelines and penalties for corporate offenders

Introduce Corporate AI Governance Requirements: Large corporations and AI-using entities should be mandated to implement internal governance mechanisms, including:

- Appointment of Chief AI Ethics Officers
- Creation of AI Ethics Committees
- Routine ethical risk assessments
- Publication of AI impact reports and risk registers

Mandate Algorithmic Explainability and Audit Trails: Corporations should be required to maintain detailed logs of AI training data, decision processes, and override mechanisms. These records must be made available to regulators and courts during investigations or disputes.

Enhance Technical Capacity of Law Enforcement and Judiciary: Invest in training cybercrime units, forensic labs, and judicial officers in AI technologies, algorithmic forensics, and automated system analysis. Introduce AI-specific modules in police and judicial academies to bridge the current knowledge gap.

Implement a Regulatory Sandbox for Ethical AI: Encourage innovation through controlled experimentation. A regulatory sandbox model would allow corporations to develop and test AI systems under close scrutiny, ensuring compliance with ethical and legal norms before mass deployment.

Strengthen International Cooperation on AI-Related Cybercrime: Cybercrime facilitated by AI often crosses national boundaries. India must enter into bilateral and multilateral

cooperation treaties for cross-border investigation, evidence sharing, and enforcement related to AI misuse.

Launch Public Awareness and Whistleblower Protection Initiatives: Establish public education campaigns on AI risks and user rights. Also, provide legal protection and incentives for whistleblowers who report AI-related corporate misconduct or unethical practices.

Promote Ethical AI Certifications and Industry Standards: Encourage industry bodies and standard-setting organisations to develop voluntary codes of conduct, certification schemes, and benchmarking tools for responsible AI deployment in the corporate sector.

VIII. CONCLUSION

The integration of Artificial Intelligence into corporate operations has ushered in a new era of efficiency and innovation. However, this technological advancement has also created fertile ground for sophisticated and often concealed forms of cybercrime. As this paper has illustrated, AI is increasingly being exploited by corporations—intentionally or negligently—as a mechanism to perpetrate cyber offenses ranging from algorithmic manipulation and data exploitation to surveillance abuses and financial fraud.

The Indian legal and regulatory landscape, while evolving, remains ill-equipped to handle the unique challenges posed by autonomous systems. The opacity of AI decision-making, the difficulty in assigning intent or accountability, and the absence of enforceable governance structures have created significant blind spots in corporate liability. When AI becomes the silent accomplice in corporate misconduct, traditional legal doctrines fall short in ensuring justice, transparency, and public protection.

Comparative analysis of global regulatory regimes reveals that a proactive, risk-based, and transparency-driven framework is essential for addressing the threats posed by corporate misuse of AI. Jurisdictions such as the European Union, the United States, and China have made tangible progress by implementing AI-specific regulations, algorithmic accountability mandates, and sectoral oversight mechanisms. India must draw from these examples to craft a bespoke regulatory model that addresses its specific legal, cultural, and technological contexts.

Ethical governance must go hand in hand with legal reform. Corporations cannot be allowed to externalise risk while internalising profits through opaque algorithms. They must be held to standards of fairness, transparency, and responsibility, particularly when deploying technologies that directly affect human rights, market integrity, and societal trust.

Ultimately, the challenge of AI-enabled corporate cybercrime is not merely technological—it is constitutional, institutional, and moral. The way forward must combine innovation with accountability, freedom with oversight, and automation with justice. Only through a comprehensive legal and ethical framework can India safeguard its digital economy while ensuring that AI serves the public good rather than subverting it.
