

THE LIAR'S DIVIDEND AND DEEPFAKES: INTERNATIONAL AND INDIAN PERSPECTIVES

Dr. Janees Rafiq^{*}
Pranav Dalia^{**}

ABSTRACT

Deepfake technology, which uses artificial intelligence to generate realistic audio and video forgeries, has raised new challenges for democratic governance, law enforcement and media trust. One particularly insidious consequence is the liar's dividend, the ability of wrongdoers to deny genuine evidence by claiming it is a deepfake. This paper examines the liar's dividend in a comparative context, drawing on international and Indian case studies. It reviews the literature on deepfakes and misinformation, defines the liar's dividend, and analyses real incidents in the United States, Europe and India in which authentic evidence was dismissed as fabricated. It reports global and Indian data on deepfake proliferation, public trust, and forensic capacity, and compares three legal regimes: India's Information Technology Act and allied provisions, the European Union's Artificial Intelligence Act and Digital Services Act obligations, and United States legislative proposals including the DEEPFAKES Accountability Act. The paper discusses the institutional and democratic implications of this trend, highlighting the erosion of trust and the empowerment of malign actors, and concludes with recommendations spanning technical defence, legal reform, platform governance, media literacy and multi-stakeholder coordination, including India's Deepfakes Analysis Unit.

KEYWORDS: Deepfakes; Liar's Dividend; Misinformation; Digital Forensics; Artificial Intelligence Regulation

1. INTRODUCTION

Digital media forensics and AI-generated content are reshaping information integrity. In recent years, highly realistic deepfake videos and audio, created using artificial intelligence, have emerged globally, capable of depicting people saying or doing things they never did and enabling sophisticated disinformation campaigns.¹ Paradoxically, the advent of deepfakes has also created a new weapon for liars, a phenomenon legal scholars Robert Chesney and Danielle Citron term the liar's dividend: the benefit wrongdoers gain by invoking the possibility of a deepfake to cast doubt on real evidence, allowing them to escape accountability for their actions by denouncing authentic video and audio as fabricated.² The consequences are serious. If genuine evidence, a politician's recorded statement, for instance, can be discredited as fake, the very notion of objective truth is undermined, and a deeply skeptical public may come to disbelieve what their own eyes or ears tell them even when the information is

^{*} Assistant Professor, Department of Law, Model Institute of Engineering & Technology, Jammu.

^{**} LLB Hons., Department of Law, Model Institute of Engineering & Technology, Jammu.

¹ Robert Chesney & Danielle Citron, Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security, 107 Calif. L. Rev. 1753, 1758 (2019).

² Id. at 1785-86.

real, eroding the trust democratic government requires to function.³ Global surveys confirm a crisis of confidence in this regard: seventy per cent of Americans report concern that deepfakes can spread disinformation, and a Reuters Institute-affiliated poll found eighty-six per cent of Indians believe misinformation and deepfakes will affect future elections.⁴ The World Economic Forum has identified AI-generated misinformation among the foremost global risks of the current period.⁵

This paper explores the liar's dividend through case analysis and data. It first reviews the literature on deepfakes and misinformation; it then draws on international examples, including United States court cases and political controversy, and Indian incidents, to illustrate how real evidence has been dismissed as fake. It examines the current legal and policy frameworks in India, the European Union and the United States governing synthetic media, assesses the implications for governance and institutional trust, and closes with recommendations spanning detection technology, regulatory reform, media literacy and institutional coordination.

2. LITERATURE REVIEW

Scholarly attention to deepfakes has grown rapidly since the mid-2010s. Early concern centred on privacy and free expression, particularly non-consensual deepfake pornography, before shifting toward questions of politics and national security.⁶ Chesney and Citron's 2019 analysis remains the seminal contribution, warning that deepfakes could inflict unprecedented harm on privacy, democracy and national security, and coining the term liar's dividend to capture the perverse effect whereby educating the public about deepfakes actually makes liars more credible when they claim genuine evidence is fabricated.⁷ Subsequent scholarship has elaborated this concept empirically: a 2023 study by Murphy and Raines experimentally demonstrated that false claims of fake news can significantly boost public support for politicians following scandal, finding that people are indeed more willing to support politicians who cry wolf

³ Id. at 1786.

⁴ George Washington University Program on Extremism, Post-Election Poll Study Reveals Deepening Distrust in Government and Information Sources (2024); Adobe, Misinformation and Harmful Deepfakes Will Affect Future Elections in India, reported in *The Indian Express* (May 22, 2024).

⁵ World Economic Forum, Deepfakes: How India Is Tackling Misinformation During Elections (Aug. 8, 2024).

⁶ M. Srikant, *Bharatiya Laws Against Deepfake Cybercrime: Opportunities and Challenges*, Vivekananda International Foundation (Apr. 28, 2025).

⁷ Robert Chesney & Danielle Citron, *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*, 107 *Calif. L. Rev.* 1753 (2019).

over fake news, confirming the political potency of the liar's dividend as a rhetorical strategy.⁸

The broader literature also documents a surging ecosystem of deepfake creation. One industry report found a tenfold global increase in detected deepfakes between 2022 and 2023, with the Asia-Pacific region recording a fifteen-hundred-and-thirty per cent increase over the same period.⁹ India is no exception: one analysis found the number of deepfake videos circulating online in 2023 was approximately five hundred and fifty per cent higher than in 2019.¹⁰ Alongside this raw proliferation, platforms have struggled to contain synthetic media; major platforms have instituted transparency policies requiring labels on AI-generated content, though enforcement remains uneven. The World Economic Forum has separately documented India's institutional response, noting that the government-backed Deepfakes Analysis Unit received hundreds of submissions through a public tip line during the 2024 general election, identifying AI alteration in numerous circulating political videos.¹¹ These developments compound existing institutional concern: the 2023 Edelman Trust Barometer recorded global trust in media at historic lows, and Indian-specific polling finds that eighty-one per cent of respondents consider online information altered and difficult to verify.¹² The literature converges on two related themes, that deepfakes are surging in both volume and technical capability, and that the liar's dividend leverages this surge to deepen institutional mistrust, while noting that legal and regulatory measures remain generally slower to develop than the underlying technology even in jurisdictions with active artificial intelligence legislation.¹³

1.1 Research Methodology

This study adopts a qualitative, multi-case approach supported by quantitative data. It systematically reviews recent academic, policy and journalistic literature on deepfakes and the liar's dividend, drawing on documented incidents, fact-checking reports and news investigation, supplemented by industry data including fraud-detection reports and government releases on forensic capacity. For the international analysis, illustrative cases from the United States and

⁸ Oliver Murphy & Anna Barbara Raines, Watch Out for False Claims of Deepfakes, and Actual Deepfakes, This Election Year, Brookings (Jan. 25, 2024).

⁹ Sumsb, Sumsb Research: Global Deepfake Incidents Surge Tenfold from 2022 to 2023 (Oct. 12, 2023); Global Initiative Against Transnational Organized Crime, Rogue Replicants: Criminal Exploitation of Deepfakes in South East Asia (June 17, 2024).

¹⁰ Global Initiative Against Transnational Organized Crime, Rogue Replicants: Criminal Exploitation of Deepfakes in South East Asia (June 17, 2024).

¹¹ World Economic Forum, Deepfakes: How India Is Tackling Misinformation During Elections (Aug. 8, 2024).

¹² Edelman, 2023 Edelman Trust Barometer: Global Report (2023); Adobe, *supra* note 4.

¹³ H. White et al., AI-Generated Deepfakes: What Does the Law Say?, Rouse (Feb. 29, 2024).

European Union were selected where authenticity disputes were central to the underlying controversy; for the Indian analysis, high-profile political and media incidents from 2020 to 2024 were examined, with case details triangulated across independent sources, including independent forensic analysis of the audio at issue in the Annamalai case discussed in Part 4.¹⁴ Legal frameworks were compared through direct review of legislative text, including India's Information Technology Act and the European Union's Artificial Intelligence Act, supplemented by credible legal commentary.¹⁵ Because the subject matter continues to evolve rapidly, the most recent available data, including 2024 industry and World Economic Forum studies, were incorporated throughout.

3. COMPARATIVE CASE ANALYSIS: THE UNITED STATES AND OTHER JURISDICTIONS

The United States has produced several notable instances of the liar's dividend in litigation. In January 2022, Capitol riot defendant Guy Reffitt became one of the first high-profile defendants to invoke deepfakes in a criminal trial: although video and audio evidence placed Reffitt at the Capitol on 6 January 2021, his defence counsel suggested to the jury that the evidence against him might be AI-manipulated, insinuating doubt notwithstanding the evidence's authenticity; Reffitt was ultimately convicted, but the episode illustrated how deepfake fears can be invoked strategically to cast doubt on genuine evidence.¹⁶ The significance of the Reffitt episode lies less in its outcome, since the jury was evidently unpersuaded, than in its timing: it arose within roughly a year of deepfake technology reaching genuine public prominence, demonstrating how quickly the liar's dividend strategy migrated from theoretical academic concern into actual courtroom practice once the underlying technology became a matter of common public awareness.

A more consequential instance arose in *Huang v. Tesla, Inc.*, a civil wrongful-death action arising from a fatal crash involving Tesla's Autopilot system, in which the plaintiffs sought to depose Elon Musk regarding a 2016 recorded statement asserting that Tesla's vehicles could, at that time, already drive more safely than a human being; Tesla's counsel argued that the recording, and other public statements attributed to Musk, might themselves be deepfakes given Musk's prominence as a frequent subject of synthetic media, and that he accordingly should not be compelled to testify regarding statements

¹⁴ P. Raina, *Indian Politician Blames AI for Alleged Leaked Audio: We Had It Tested*, *Rest of World* (June 27, 2023).

¹⁵ M. Srikant, *supra* note 6.

¹⁶ C.W. Martin & S. Etemad-Moghadam, *Deepfakes in Litigation*, Lutzker & Lutzker LLP (Mar. 12, 2024); *United States v. Reffitt*, No. 1:21-cr-00032 (D.D.C. 2022).

he could not verify as authentic.¹⁷ Santa Clara County Superior Court Judge Evette Pennypacker rejected this argument as deeply troubling, holding that accepting it would allow individuals, particularly public figures, to hide behind the mere possibility that their own recorded statements were deepfakes in order to avoid taking ownership of what they had actually said and done, and ordered Musk to submit to deposition.¹⁸ The Huang ruling is doctrinally significant beyond its immediate procedural context because it establishes, at least at the level of a single trial court, a workable evidentiary principle for addressing the liar's dividend directly: a bare assertion that evidence might theoretically be a deepfake, unsupported by any specific forensic indication that the particular recording in question displays signs of manipulation, does not suffice to relieve a party of accountability for authenticated prior conduct. This principle, that the liar's dividend defence requires some case-specific evidentiary foundation rather than mere reference to the general possibility of AI fabrication, offers a template other jurisdictions, including India, could usefully adopt as courts increasingly confront comparable disputes. Both cases exemplify American courts confronting the liar's dividend strategy directly: an attempt to discredit authentic evidence by invoking, without independent substantiation, the mere theoretical possibility of fabrication.

On the legislative front, American policymakers have focused predominantly on malicious deepfake application, including non-consensual pornography, though newer proposals address misinformation more directly; the DEEPPAKES Accountability Act and the NO FAKES Act, both introduced in Congress in 2023 and 2024 respectively, would criminalise certain deceptive uses of synthetic media and require disclosure of AI-altered content, though no federal statute presently targets the liar's dividend scenario specifically.¹⁹ In Europe, awareness of deepfake-related evidentiary denial has grown alongside policy development; Chesney and Citron's scholarship has been cited directly in European Union AI policy debate, and the Artificial Intelligence Act's transparency-marking obligations, examined further in Part 5, respond in part to this concern.²⁰ Italy's former Prime Minister Matteo Renzi was targeted by satirical deepfakes that, although initially presented as parody, generated broader concern about the blurring of synthetic and authentic political content.²¹ The Asia-Pacific region has experienced comparable pressure from a different direction: a Malaysian fraud ring's use of AI-generated voice cloning in telephone scams in early 2023 prompted regulatory warning regarding voice

¹⁷ Huang v. Tesla, Inc., No. 19CV346663 (Cal. Super. Ct. Santa Clara Cnty., tentative ruling Apr. 26, 2023).

¹⁸ Id.

¹⁹ M. Srikant, *supra* note 6.

²⁰ H. White et al., *supra* note 13.

²¹ From Fake News to False Elections, EuropeNow (Aug. 2, 2020).

deepfakes specifically, illustrating that the underlying technology's misuse extends well beyond the evidentiary-denial scenario this paper's title identifies.²² Industry monitoring recorded a seventeen-hundred-and-forty per cent surge in deepfake detections across North America in 2023 alongside the fifteen-hundred-and-thirty per cent Asia-Pacific increase noted in Part 2, confirming that the underlying technological trend is genuinely global rather than concentrated in any single region.²³ Taken together, these cases confirm that the liar's dividend is not a theoretical construct but an active litigation and political strategy already tested before courts and electorates in multiple jurisdictions, a context that motivates the Indian-specific analysis in the following Part.

4. THE INDIAN EXPERIENCE: DEEPFAKES AND DENIAL IN PRACTICE

India's deepfake landscape has grown rapidly amid intense electoral competition. A 2024 Adobe survey found India ranked highest globally among the countries surveyed for the share of respondents, eighty-one per cent, who believe online content is being digitally altered, and government and civil society bodies have begun documenting instances both of synthetic media used for manipulation and, more significantly for this paper's argument, of the liar's dividend deployed to deny authentic material.²⁴ This ranking is itself analytically significant: it suggests that Indian public anxiety regarding deepfake authenticity now exceeds that recorded in the United States and European Union despite India's comparatively nascent regulatory response, examined in Part 5, a mismatch between public concern and institutional preparedness that the case analysis below illustrates in concrete terms.

The clearest Indian example arose in June 2023, when Tamil Nadu's Bharatiya Janata Party state president K. Annamalai released two viral audio clips, one purporting to show an opposition leader offering a bribe and another of separately questionable content, embarrassing the ruling Dravida Munnetra Kazhagam government; when evidence subsequently indicated that at least the first clip was genuine, Annamalai dismissed both as fabricated using artificial intelligence.²⁵ Independent forensic analysis conducted by journalists at Rest of World subsequently confirmed that the first clip was authentic, featuring identifiable voice and contextual detail consistent with the officials concerned, while the second recording appeared to have been AI-modified; the press characterised the episode as India's first high-profile instance of the liar's dividend, with generative AI functioning as a strategic cover for plausible

²² Global Initiative Against Transnational Organized Crime, *supra* note 9.

²³ Sumsb, *supra* note 9.

²⁴ Adobe, *supra* note 4.

²⁵ P. Raina, *supra* note 14.

deniability even where independent verification ultimately established the primary clip's authenticity.²⁶ The Annamalai episode is instructive precisely because of this mixed outcome: rather than illustrating a pure liar's dividend scenario in which entirely authentic evidence was falsely denied, it demonstrates a more sophisticated and arguably more dangerous variant, in which a genuine recording and a fabricated one were bundled together and denied collectively, allowing the underlying denial to borrow credibility from the fact that at least one component of the dispute genuinely did involve synthetic media. This bundling strategy, blending real manipulation with false claims of manipulation, may prove considerably harder for courts, fact-checkers and the public to unravel than the simpler binary scenario, an authentic recording falsely denied outright, that the Huang v. Tesla and Reffitt cases examined in Part 3 each present.

Other Indian incidents involve genuine deepfakes rather than false denial of authentic material, but illustrate the same underlying public sensitivity. In April 2024, Bollywood actor Ranveer Singh discovered a manipulated video of himself purportedly endorsing a political party; the video was in fact fabricated, and Singh's public response, a viral warning urging audiences to beware of deepfakes, together with a formal police complaint, illustrates a markedly different institutional posture from cases involving genuine evidence subsequently denied.²⁷ A comparable incident involving actor Aamir Khan, in which a video appearing to show him endorsing a political party was similarly disavowed as fabricated, reinforces the same pattern.²⁸ Neither case involved a liar's dividend denial of genuinely authentic material, but both demonstrate the pervasiveness of deepfake creation targeting public figures and the corresponding public alertness the phenomenon has generated. During the 2024 general election, India's Deepfakes Analysis Unit, established through coordination between the Ministry of Electronics and Information Technology and election authorities, processed hundreds of public submissions via a WhatsApp tip line, encompassing both unsophisticated manual edits and genuine AI-generated deepfakes, including lip-synced clips in which a candidate's speech was spliced onto unrelated footage in a manner that generated genuine public confusion.²⁹ In several contested instances, this environment produced false claims that authentic press conferences and speeches were themselves AI-altered, requiring professional fact-checking intervention to resolve; as one contemporaneous analysis of the Annamalai episode observed, generative AI carries extraordinary potential to tarnish

²⁶ Id.

²⁷ A. Gautam, Ranveer Singh Files Police Case After Deepfake Video Goes Viral, NDTV (Apr. 22, 2024).

²⁸ Id.

²⁹ P. Raina, *supra* note 14.

reputation on one hand, while proving potentially even more useful, on the other, in letting powerful individuals evade accountability for conduct genuinely their own.³⁰ India has accordingly witnessed both the creation of deepfakes and the deployment of the liar's dividend against authentic material within a comparatively compressed period, a pattern this paper's legal analysis in the following Part assesses against the country's existing statutory framework.

5. LEGAL AND POLICY FRAMEWORKS: INDIA, THE EUROPEAN UNION AND THE UNITED STATES

India currently has no statute criminalising deepfakes as such, relying instead on general-purpose provisions applied to deepfake-related harm. Under the Information Technology Act, 2000, Sections 66C and 66D address identity theft and cheating by impersonation respectively, capable of application to unauthorised audio or visual impersonation, while Section 66E prohibits the capture or transmission of a person's private image without consent, and Sections 67A and 67B penalise obscene or pornographic synthetic content specifically.³¹ Defamation provisions under the Bharatiya Nyaya Sanhita, including Section 356 addressing harm to reputation, are additionally invoked where deepfake content causes reputational injury, and the Digital Personal Data Protection Act, 2023 governs the underlying biometric or voice data whose unauthorised use to generate a deepfake may itself constitute a distinct statutory violation.³² These provisions remain general-purpose rather than deepfake-specific, however, and a 2025 policy analysis concludes that India presently lacks a dedicated legal framework addressing deepfake-related cybercrime, relying instead on this patchwork of catch-all provision.³³ The structural weakness this patchwork approach produces is worth stating explicitly: none of these provisions was drafted with the liar's dividend scenario in view, since each addresses either the creation of deepfake content itself, through the impersonation, privacy and obscenity provisions, or the underlying data misuse, through the DPDP Act, rather than the distinct harm of falsely denying that authentic content is authentic. A prosecutor or litigant confronting an Annamalai-style denial accordingly has no direct statutory hook comparable to the specific offence of, say, wiretapping under the Telegraph Act, and must instead attempt to fit the denial itself within general defamation or fraud categories never designed to address it.

The Government has responded partially through a 2024 advisory issued under the Intermediary Guidelines Rules, instructing platforms to label

³⁰ Id.

³¹ M. Srikant, *supra* note 6.

³² Id.

³³ Id.

synthetic media and cooperate with fact-checkers in demarcating misinformation, imposing a due-diligence obligation requiring intermediaries to flag potentially AI-generated misinformation for verification, though the advisory's disclosure requirement remains considerably lighter than binding statutory obligation.³⁴ The Supreme Court's recognition of digital privacy in Justice K.S. Puttaswamy v. Union of India indirectly shapes how deepfake-related evidence may eventually be treated in Indian courts, though no Indian court has yet ruled specifically on deepfake admissibility.³⁵

Table 1: Comparative Legal Measures Addressing Deepfakes

Aspect	India	United States	European Union
Dedicated deepfake statute	None; general IT Act and BNS provisions applied	No federal statute; some state laws address AI sexual content	AI Act (2024) mandates transparency marking for AI-generated content
Data privacy framework	DPDP Act, 2023 covers personal and biometric data	State privacy laws vary; no federal statute covering voice likeness specifically	GDPR covers personal data; AI Act extends to certain non-personal content
Platform regulation	Intermediary Guidelines Rules, 2021; 2024 government advisory on labelling	Largely voluntary platform self-regulation; no binding disclosure requirement	Digital Services Act requires content moderation and removal reporting
Penalties for non-compliance	IT Act penalties for identity fraud and defamation; advisory compliance presently voluntary	Fines proposed under pending federal bills; targeted state penalties	DSA permits fines of up to six per cent of global turnover for breach
Pending legislative reform	No dedicated AI or deepfake statute under active debate as of 2025	DEEPFAKES Accountability Act and NO FAKES Act pending	AI Act fully in force; Digital Services Act operative since 2024

By contrast, the European Union's Artificial Intelligence Act, Regulation (EU) 2024/1689, addresses synthetic content directly: Article 50 requires providers of systems generating deepfake images, audio or video to ensure

³⁴ Id.

³⁵ Justice K.S. Puttaswamy (Retd.) v. Union of India, (2017) 10 S.C.C. 1 (India).

outputs are marked in a machine-readable format detectable as artificially generated, and requires deployers of such systems to disclose the synthetic nature of the resulting content to those exposed to it.³⁶ The Digital Services Act, Regulation (EU) 2022/2065, further obliges very large online platforms to remove unlawful misinformation with reasonable promptness and to maintain transparency reporting regarding content moderation.³⁷ The United Kingdom's Online Safety Act 2023 imposes a comparable obligation on platforms to address harmful user-generated content, including manipulative AI media, while the United States retains no comprehensive federal statute specifically addressing deepfakes, relying instead on pending proposals including the DEEPFAKES Accountability Act, examined in Part 3, alongside state-level statutes, several of which, including in Texas and California, criminalise non-consensual pornographic deepfakes specifically without extending to the liar's dividend scenario this paper examines.

India's framework, in sum, possesses the criminal-law building blocks, fraud, defamation and privacy provisions, necessary to prosecute deepfake-related harm, but its judiciary and enforcement agencies presently lack explicit statutory guidance on AI-specific evidentiary questions, and no Indian law currently addresses the liar's dividend scenario, the false denial of authentic evidence, directly; institutions, principally courts and independent fact-checkers, must instead resolve authenticity disputes case by case using available digital forensic technique, a gap this paper's recommendations in Part 7 address directly.

6. INSTITUTIONAL AND DEMOCRATIC IMPLICATIONS

The liar's dividend exacerbates a pre-existing crisis of institutional trust. As Chesney and Citron warn, once deepfakes become sufficiently widespread, the public may struggle to believe what their own eyes or ears report even where the information in question is genuinely authentic, threatening the trust democratic governance requires to function.³⁸ Empirical survey evidence corroborates this concern directly: a post-election poll found a plurality of Americans trust neither the federal government nor national news media to accurately report fact, with seventy per cent of respondents specifically identifying the emergence of deepfakes as having made it harder to know what to believe online.³⁹ Indian sentiment mirrors this pattern, with over eighty per cent of respondents reporting difficulty verifying online content and only a

³⁶ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), art. 50, 2024 O.J. (L 1689).

³⁷ Regulation (EU) 2022/2065 (Digital Services Act), 2022 O.J. (L 277) 1.

³⁸ Robert Chesney & Danielle Citron, *supra* note 1, at 1786.

³⁹ George Washington University Program on Extremism, *supra* note 4.

small minority expressing confidence in most information encountered online.⁴⁰ Media literacy scholarship cautions that motivated reasoning will likely worsen this dynamic over time: once individuals suspect that even genuine footage might be synthetic, they become inclined to selectively doubt evidence that incriminates figures they support while readily accepting misinformation that exonerates them, a pattern that deepens political polarisation by allowing each side to construct a self-reinforcing account of what counts as credible evidence.⁴¹

These dynamics carry a specific institutional cost for courts and investigative bodies, which face a growing evidentiary burden: digital evidence increasingly requires supplementary proof of authenticity, including metadata analysis, multiple independent detection methods, and corroborating testimony, before it can be relied upon with confidence. India's forensic laboratories, already reportedly carrying substantial case backlogs, face mounting strain as this evidentiary burden grows, underscoring the practical urgency of the capacity-building recommendations examined in Part 7. This strain compounds an existing structural vulnerability specific to the Indian context: because courts, election authorities and fact-checking bodies in India presently rely on essentially the same forensic infrastructure to investigate both genuine deepfake creation and liar's dividend denial claims, a surge in either category of dispute directly diminishes the system's capacity to resolve the other promptly, creating a shared bottleneck that a jurisdiction with more developed, differentiated forensic capacity might avoid. Public agencies including election commissions and independent fact-checking organisations must correspondingly build and sustain their own institutional credibility, since a political culture in which any scandal-plagued figure can plausibly claim genuine evidence is an AI fabrication risks corroding public confidence not merely in the specific evidence at issue but in democratic accountability mechanisms more generally. Judge Pennypacker's warning in *Huang v. Tesla*, that accepting a deepfake defence without independent substantiation would set a precedent permitting individuals to avoid ownership of their own recorded conduct, captures precisely the institutional stake this paper's analysis identifies: if that principle were to spread from isolated litigation into mainstream political practice, meaningful accountability for recorded public conduct could erode substantially.

7. RECOMMENDATIONS

Addressing the liar's dividend requires a coordinated, multi-pronged strategy rather than any single legal or technical fix. On the technical front, deepfake detection tools, while presently imperfect, can flag likely manipulation for

⁴⁰ Adobe, *supra* note 4.

⁴¹ Oliver Murphy & Anna Barbara Raines, *supra* note 8.

further human review, and continued research investment, alongside international content-provenance standards such as the Coalition for Content Provenance and Authenticity framework, should be pursued to strengthen detection capacity; India's stated plan to link its forensic laboratories into a unified digital forensics network, together with continued investment through the National Forensic Sciences University, would meaningfully reduce the evidentiary backlog identified in Part 6 and build the specialised capacity authenticity disputes increasingly require. On the legal front, India should consider legislation specifically addressing malicious AI-generated content, including criminalising the wilful misuse of AI systems for public deception and mandating transparent labelling of AI-generated political content, drawing on the European Union's Article 50 watermarking obligation and India's own 2024 advisory as starting templates; a statutory duty of candour requiring disclosure of digital alteration, alongside a requirement that parties introducing disputed digital evidence in litigation consent to independent forensic validation, would directly target the liar's dividend scenario that no current Indian provision addresses, and civil defamation remedies should be examined for their capacity to compensate victims of false "it's a deepfake" denials specifically, distinct from the underlying defamation the original evidence may itself document.

Platforms bear an independent obligation to moderate deepfake content proactively, extending existing AI-content labelling policy toward permanent, tamper-resistant watermarking of AI-generated political advertising specifically, and adopting notice-and-takedown procedures for viral deepfakes comparable to the Digital Services Act's requirements; the Indian government's existing directive that intermediaries appoint grievance officers for synthetic-media complaints under the Intermediary Guidelines represents a foundation this recommendation would extend rather than a mechanism requiring wholesale replacement. Sustained public media literacy investment remains equally indispensable: national digital literacy campaigns, integration of media literacy curricula within schools and universities, and adequately funded, independent fact-checking organisations empowered to certify or debunk contested content publicly would collectively help citizens distinguish credible authentication from baseless denial, while publicising successful forensic authentication of contested evidence would help rebuild the institutional trust Part 6 shows has already eroded substantially. Finally, institutional coordination across government, industry and civil society should be formalised and expanded beyond its present ad hoc form: India's Deepfakes Analysis Unit offers a workable public-private partnership template that could evolve into a permanent centre of excellence on AI-generated misinformation, complemented by judicial training programmes preparing judges to assess deepfake-related evidentiary disputes and by international cooperation mechanisms addressing the transnational character many deepfake campaigns already exhibit.

Implementing these measures in combination, rather than pursuing any single intervention in isolation, offers the most realistic prospect of narrowing both the prevalence of convincing deepfakes and the ease with which AI-generated doubt can be weaponised against authentic evidence.

8. CONCLUSION

The liar's dividend, the denial of genuine evidence through the mere invocation of deepfake technology, represents a distinct information threat from the direct harms synthetic media itself causes, since it undermines accountability by rendering truth itself perpetually contestable. This paper's analysis demonstrates that the phenomenon is already materialising globally, and that India is not immune to it: from a state political leader's audio scandal to American civil litigation invoking the possibility that a chief executive's own recorded statements might be artificially generated, powerful actors already possess clear incentive to invoke this defence strategically. Without effective countermeasures, any recorded speech, video or image risks dismissal on purely technical grounds, eroding the factual foundation on which public accountability and democratic discourse both depend.

Addressing this threat requires simultaneous progress across several fronts examined in this paper: strengthened forensic capacity and legal clarity to establish genuine evidence beyond reasonable doubt, transparency obligations of the kind the European Union's Artificial Intelligence Act already imposes to raise the practical cost of deceptive denial, more assertive platform-level content verification and labelling, and a citizenry equipped through sustained media literacy investment to seek corroboration before accepting or rejecting contested digital evidence. As Chesney and Citron themselves observed, successfully educating the public about the existence of deepfakes paradoxically strengthens the liar's dividend unless that education is deliberately paired with the practical tools needed to authenticate genuine evidence.⁴² Combining technical detection capacity, calibrated legal reform and sustained public literacy offers the most credible path toward narrowing the liar's dividend and preserving public trust in truthful evidence as deepfake technology continues to advance.

BIBLIOGRAPHY

Journal Articles and Reports

1. Robert Chesney & Danielle Citron, *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*, 107 *Calif. L. Rev.* 1753 (2019).

⁴² Robert Chesney & Danielle Citron, *supra* note 1, at 1795.

2. Oliver Murphy & Anna Barbara Raines, Watch Out for False Claims of Deepfakes, and Actual Deepfakes, This Election Year, Brookings (Jan. 25, 2024).
3. Sumsb, Sumsb Research: Global Deepfake Incidents Surge Tenfold from 2022 to 2023 (Oct. 12, 2023).
4. Global Initiative Against Transnational Organized Crime, Rogue Replicants: Criminal Exploitation of Deepfakes in South East Asia (June 17, 2024).
5. World Economic Forum, Deepfakes: How India Is Tackling Misinformation During Elections (Aug. 8, 2024).
6. Edelman, 2023 Edelman Trust Barometer: Global Report (2023).
7. M. Srikant, Bharatiya Laws Against Deepfake Cybercrime: Opportunities and Challenges, Vivekananda International Foundation (Apr. 28, 2025).
8. H. White et al., AI-Generated Deepfakes: What Does the Law Say?, Rouse (Feb. 29, 2024).
9. C.W. Martin & S. Etemad-Moghadam, Deepfakes in Litigation, Lutzker & Lutzker LLP (Mar. 12, 2024).
10. P. Raina, Indian Politician Blames AI for Alleged Leaked Audio: We Had It Tested, Rest of World (June 27, 2023).
11. George Washington University Program on Extremism, Post-Election Poll Study Reveals Deepening Distrust in Government and Information Sources (Press Release, 2024).
12. Adobe, Misinformation and Harmful Deepfakes Will Affect Future Elections in India (2024 Study), reported in The Indian Express (May 22, 2024).

Cases

1. Justice K.S. Puttaswamy (Retd.) v. Union of India, (2017) 10 S.C.C. 1 (India).
2. United States v. Reffitt, No. 1:21-cr-00032 (D.D.C. 2022).
3. Huang v. Tesla, Inc., No. 19CV346663 (Cal. Super. Ct. Santa Clara Cnty., tentative ruling Apr. 26, 2023).

Statutes and Regulations

1. The Information Technology Act, No. 21 of 2000, §§ 66C, 66D, 66E, 67A, 67B (India).
2. The Bharatiya Nyaya Sanhita, No. 45 of 2023, § 356 (India).
3. The Digital Personal Data Protection Act, No. 22 of 2023 (India).
4. Information Technology (Intermediary Guidelines and Digital Media Ethics Code) Rules, 2021 (India).
5. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), art. 50, 2024 O.J. (L 1689).
6. Regulation (EU) 2022/2065 (Digital Services Act), 2022 O.J. (L 277) 1.
7. Online Safety Act 2023, c. 50 (UK).
8. DEEPFAKES Accountability Act, H.R. 5586, 118th Cong. (2023) (US) (proposed).
9. NO FAKES Act, S. 4875, 118th Cong. (2024) (US) (proposed).