

INTERNATIONAL JOURNAL FOR LEGAL RESEARCH AND ANALYSIS



Open Access, Refereed Journal Multi Disciplinary
Peer Reviewed

www.ijlra.com

DISCLAIMER

No part of this publication may be reproduced, stored, transmitted, or distributed in any form or by any means, whether electronic, mechanical, photocopying, recording, or otherwise, without prior written permission of the Managing Editor of the *International Journal for Legal Research & Analysis (IJLRA)*.

The views, opinions, interpretations, and conclusions expressed in the articles published in this journal are solely those of the respective authors. They do not necessarily reflect the views of the Editorial Board, Editors, Reviewers, Advisors, or the Publisher of IJLRA.

Although every reasonable effort has been made to ensure the accuracy, authenticity, and proper citation of the content published in this journal, neither the Editorial Board nor IJLRA shall be held liable or responsible, in any manner whatsoever, for any loss, damage, or consequence arising from the use, reliance upon, or interpretation of the information contained in this publication.

The content published herein is intended solely for academic and informational purposes and shall not be construed as legal advice or professional opinion.

**Copyright © International Journal for Legal Research & Analysis.
All rights reserved.**

ABOUT US

The *International Journal for Legal Research & Analysis (IJLRA)* (ISSN: 2582-6433) is a peer-reviewed, academic, online journal published on a monthly basis. The journal aims to provide a comprehensive and interactive platform for the publication of original and high-quality legal research.

IJLRA publishes Short Articles, Long Articles, Research Papers, Case Comments, Book Reviews, Essays, and interdisciplinary studies in the field of law and allied disciplines. The journal seeks to promote critical analysis and informed discourse on contemporary legal, social, and policy issues.

The primary objective of IJLRA is to enhance academic engagement and scholarly dialogue among law students, researchers, academicians, legal professionals, and members of the Bar and Bench. The journal endeavours to establish itself as a credible and widely cited academic publication through the publication of original, well-researched, and analytically sound contributions.

IJLRA welcomes submissions from all branches of law, provided the work is original, unpublished, and submitted in accordance with the prescribed submission guidelines. All manuscripts are subject to a rigorous peer-review process to ensure academic quality, originality, and relevance.

Through its publications, the *International Journal for Legal Research & Analysis* aspires to contribute meaningfully to legal scholarship and the development of law as an instrument of justice and social progress.

PUBLICATION ETHICS, COPYRIGHT & AUTHOR RESPONSIBILITY STATEMENT

The *International Journal for Legal Research and Analysis (IJLRA)* is committed to upholding the highest standards of publication ethics and academic integrity. All manuscripts submitted to the journal must be original, unpublished, and free from plagiarism, data fabrication, falsification, or any form of unethical research or publication practice. Authors are solely responsible for the accuracy, originality, legality, and ethical compliance of their work and must ensure that all sources are properly cited and that necessary permissions for any third-party copyrighted material have been duly obtained prior to submission. Copyright in all published articles vests with IJLRA, unless otherwise expressly stated, and authors grant the journal the irrevocable right to publish, reproduce, distribute, and archive their work in print and electronic formats. The views and opinions expressed in the articles are those of the authors alone and do not reflect the views of the Editors, Editorial Board, Reviewers, or Publisher. IJLRA shall not be liable for any loss, damage, claim, or legal consequence arising from the use, reliance upon, or interpretation of the content published. By submitting a manuscript, the author(s) agree to fully indemnify and hold harmless the journal, its Editor-in-Chief, Editors, Editorial Board, Reviewers, Advisors, Publisher, and Management against any claims, liabilities, or legal proceedings arising out of plagiarism, copyright infringement, defamation, breach of confidentiality, or violation of third-party rights. The journal reserves the absolute right to reject, withdraw, retract, or remove any manuscript or published article in case of ethical or legal violations, without incurring any liability.

FROM DEEPPAKES TO DIGITAL EXPLOITATION: ADDRESSING THE LEGAL VACUUM AROUND AI- GENERATED NON-CONSENSUAL CONTENT

AUTHORED BY - HARSHITA CHAUHAN¹

¹Law student, VSLLS, Vivekananda Institute of Professional Studies, Delhi

1. ABSTRACT

The rapid growth of artificial intelligence has made it easier than ever to create and share deepfakes and other forms of synthetic media in the digital world. While these technologies have useful applications, they are increasingly being misused to produce non-consensual and sexually explicit content, often targeting women and even minors. Recent incidents, including the circulation of inappropriate AI-generated images of celebrities, ordinary digital users with 99% of them being that of women and underage girls, and the misuse of AI tools like 'Grok' especially on social media platforms like 'X' and 'Instagram', show how quickly such content can spread online and how difficult it is to control once released due to lack of strict and clear legal provisions regarding them. These developments raise serious concerns about privacy, consent, identity, and the growing problem of gender-based digital exploitation as well as violence.

This paper aims to answer the question:- **“To what extent are existing legal frameworks adequate in addressing the misuse of AI Generated and synthetic media, particularly in protecting privacy, consent and prevention Gender based digital exploitation and violence?”**

It uses a doctrinal and comparative approach to examine the current legal position in India and compare it with emerging responses in other jurisdictions.

The analysis shows that existing laws are not fully equipped to deal with these challenges. Most legal provisions were not designed to address AI-generated harm and often fail to provide clear remedies or accountability. The absence of specific laws, combined with the speed and scale of these arising issues, creates serious gaps in protection, privacy and fundamental rights of users in digital spaces. There is a need for clearer and more focused legal frameworks that

recognise the importance of consent, address gendered harm, and assign responsibility more effectively in cases of AI misuse and Digital Exploitation.

Keywords: Deepfakes; Artificial Intelligence; Synthetic Media; Digital Exploitation; Consent; Privacy; Gender-Based Violence; Legal Regulation

2. INTRODUCTION

The rapid advancement of artificial intelligence has significantly reshaped the creation and circulation of digital content, particularly through the rise of deepfakes and other forms of synthetic media. These technologies are capable of generating highly realistic images, videos, and audio that can be difficult to distinguish from authentic material. While they offer innovative possibilities across various sectors, their misuse has emerged as a serious legal and social concern, especially in relation to privacy, consent, and identity.

A recent and widely discussed example of this misuse is the circulation of non-consensual, sexually explicit AI-generated images of Taylor Swift across social media platforms. These images spread rapidly, with some reportedly receiving millions of views before being taken down¹. The incident demonstrated not only the speed at which such content can proliferate, but also the limitations of platform moderation systems in responding effectively. In response to the widespread circulation, even search functionalities were temporarily restricted, reflecting the scale and urgency of the issue.² Such developments have intensified global concerns regarding the absence of effective legal safeguards against AI-generated exploitation.³

The legal challenges posed by synthetic media are particularly complex because existing legal frameworks were not designed to address such technologically mediated harm. Traditional laws relating to obscenity, defamation, and identity theft often struggle to capture the unique nature of deepfakes, where the content may be entirely fabricated yet capable of causing real and lasting harm. This creates a gap between legal definitions and lived experiences, making enforcement difficult and often inadequate.

¹ Taylor Swift and the Dangers of Deepfake Pornography, National Sexual Violence Resource Center (Jan. 2024), https://www.nsvrc.org/blog_post/taylor-swift-and-dangers-deepfake-pornography/

² X Blocks Some Taylor Swift Searches as Deepfake Images Circulate, PBS News (Jan. 2024), <https://www.pbs.org/newshour/nation/x-blocks-some-taylor-swift-searches-as-deepfake-explicit-images-circulate>

³ Clare Duffy, Fake Explicit Images of Taylor Swift Show the Dark Side of AI, CNN (Jan. 2024), <https://edition.cnn.com/2024/01/26/tech/taylor-swift-ai-images/index.html>

In the Indian context, these concerns intersect with constitutional protections of privacy and dignity. The recognition of privacy as a fundamental right in Justice K.S. Puttaswamy v. Union of India ⁴ underscores the importance of individual control over personal information and identity. However, the application of this principle to AI-generated content remains uncertain, particularly in cases where the representation is artificial but the impact on the individual is real. This raises important questions about the scope and adaptability of existing legal doctrines.

Moreover, emerging evidence suggests that the misuse of deepfakes is not evenly distributed. A significant proportion of non-consensual synthetic content disproportionately targets women, often in the form of sexually explicit material. This has led to the recognition of such practices as a form of gender-based digital violence, where technology is used to harass, intimidate, and exploit. The scale, anonymity, and permanence of digital platforms further amplify this harm, making it difficult for victims to seek timely and effective remedies.

At the same time, the borderless nature of digital platforms complicates regulatory efforts. Content created in one jurisdiction can be shared globally within seconds, making it challenging for domestic legal systems to respond effectively. While some countries have begun to explore regulatory measures, the global legal response remains fragmented and reactive.

In this context, the misuse of AI-generated deepfakes and synthetic media raises a fundamental question regarding the adequacy of existing legal frameworks in addressing such harms, particularly in relation to privacy, consent, and the prevention of gender-based digital exploitation.

3. CONCEPTUALIZING SYNTHETIC MEDIA AND DIGITAL HARM

3.1. Defining Deepfakes and AI-Generated Content

Deepfakes and AI-generated content refer to media that is either manipulated or entirely created using artificial intelligence techniques, particularly deep learning models. These systems, often trained on large datasets, are capable of generating highly realistic images, videos, and audio that can convincingly imitate real individuals. Unlike traditional forms of digital editing,

⁴ Justice K.S. Puttaswamy v. Union of India, (2017) 10 SCC 1 (India)

deepfakes rely on neural networks to replicate facial expressions, voice patterns, and body movements with a high degree of accuracy.⁵

From a technical perspective, deepfakes are commonly produced using generative adversarial networks (GANs), where two neural networks work together to create increasingly realistic outputs.⁶ However, the legal understanding of deepfakes is still evolving. There is currently no universally accepted legal definition, and most jurisdictions attempt to regulate such content through existing laws on defamation, obscenity, identity theft, or data protection. This lack of precise legal categorization creates ambiguity, particularly when the content is fabricated but still causes real harm.

3.2. From Innovation to Exploitation

Artificial intelligence technologies were initially developed to enhance innovation across fields such as entertainment, education, and communication. However, the same tools have increasingly been misused for harmful purposes. The transition from innovation to exploitation has been driven largely by the growing accessibility of AI tools, many of which are now freely available or require minimal technical expertise to operate.

This accessibility has enabled individuals to create and distribute manipulated or synthetic content at an unprecedented scale. What once required advanced technical knowledge can now be achieved through user-friendly applications and online platforms. As a result, the production of non-consensual and exploitative content has become faster, cheaper, and more widespread.⁷ The scale of this misuse is particularly concerning. Reports indicate that a significant portion of deepfake content online is pornographic in nature, with the vast majority targeting women without their consent.⁸ This demonstrates how technological advancement, in the absence of adequate safeguards, can facilitate new forms of digital harm rather than purely beneficial outcomes.

3.3. Rethinking “Harm” in the Digital Age

The rise of synthetic media challenges traditional legal understandings of harm. In many cases,

⁵ Making deepfake tools doesn't have to be irresponsible, MIT Technology Review, <https://www.technologyreview.com/2019/12/12/131605/ethical-deepfake-tools-a-manifesto/>

⁶ Ian Goodfellow et al., Generative Adversarial Networks, arXiv (2014), <https://arxiv.org/abs/1406.2661>

⁷ The State of Deepfakes 2023, Sensity AI Report, <https://sensity.ai/reports/>

⁸ Deepfake Pornography Is Out of Control, WIRED (2019), <https://www.wired.com/story/deepfake-porn-is-out-of-control/>

deepfake content is entirely fabricated, yet it can cause serious damage to an individual's reputation, dignity, and psychological well-being.

This raises an important question: how should the law respond to harm that does not involve a “real” act but produces very real consequences?

Unlike physical or direct harm, digital harm is often intangible, but no less significant. Non-consensual deepfakes can lead to reputational damage, social stigma, emotional distress, and even professional consequences. The permanence and shareability of digital content further intensify this impact, as such material can be repeatedly circulated and remain accessible over long periods of time.⁹

Legal frameworks have traditionally focused on tangible harm or clearly identifiable wrongful acts. However, synthetic media blurs these boundaries by creating situations where the harm lies not in what actually occurred, but in what is perceived to have occurred. This disconnect exposes a critical gap in existing legal doctrines, particularly in areas concerning privacy, identity, and dignity.

The recognition of privacy as a fundamental right in *Justice K.S. Puttaswamy v. Union of India* emphasizes the importance of personal autonomy and control over one's identity. However, the application of these principles to AI-generated representations remains uncertain. As a result, there is a growing need to rethink legal definitions of harm in a way that accounts for the realities of the digital environment.

4. AI- ENABLED SEXUAL EXPLOITATION AND THE CRISIS OF CONSENT AND DIGITAL PRIVACY

AI-Enabled sexual exploitation has significantly amplified the grave crisis of consent and digital privacy in the modern world, powered by Gen AI tools that allow perpetrators to create, distribute, and manipulate inappropriate and intimate imagery without the permission of the victim. This is an alarming crisis because most of these imagery created are of women and even underage girls leaving women in a vulnerable state even in the digital world more than ever before. The lack of existing strict legal provisions of these issues allows more perpetrators to participate in these crimes without clear legal repercussions, making gender-based violence and exploitation inevitable.

⁹ The Psychological Effects of AI Clones and Deepfakes, Psychology Today (Feb, 2024), <https://www.psychologytoday.com/us/blog/urban-survival/202401/the-psychological-effects-of-ai-clones-and-deepfakes>

4.1. Non-Consensual Deepfakes and Nudification Technologies

The rapid evolution of artificial intelligence has enabled the emergence of non-consensual deepfakes and so-called “nudification” technologies, which allow users to generate sexually explicit images or videos of individuals without their knowledge or consent. These tools use machine learning models to manipulate or entirely fabricate visual content, often requiring nothing more than a single photograph. Unlike traditional image editing, such technologies are highly automated and widely accessible, allowing even individuals with no technical expertise to produce realistic explicit content.

Recent investigations reveal that these tools are not only widely available but also actively promoted across mainstream platforms. A report found that dozens of “nudify” applications capable of generating explicit deepfake images are still accessible through major app stores, raising serious concerns about platform regulation and oversight.¹⁰

Similarly, investigations have uncovered advertisements on social media platforms promoting tools that allow users to “upload a photo” and generate nude images, often without any meaningful age verification safeguards.¹¹

The accessibility and ease of use of these technologies have played a crucial role in their rapid spread. What previously required advanced technical skills can now be achieved in seconds, enabling the large-scale creation and distribution of non-consensual explicit content. Victims are often targeted using images taken from social media, making the harm both personal and widespread.¹²

A particularly alarming dimension of this issue is its growing impact on minors. Evidence suggests that these tools are increasingly being used to generate sexually explicit content involving children, often by manipulating otherwise innocent images. The Internet Watch Foundation has reported that AI-generated sexual imagery involving minors is rising rapidly, with cases involving the use of “nudification” tools to create exploitative material.¹³

In some instances, such fabricated content has been used for blackmail and coercion, including cases where perpetrators used AI-generated images to extort minors into sharing real explicit

¹⁰Apple, Google app stores still host dozens of AI ‘nudify’ apps, Indian Express (2026) <https://indianexpress.com/article/technology/apple-google-app-stores-still-host-dozens-of-ai-nudify-apps-report-claims-10498397/>

¹¹Meta platforms showed “nudify” deepfake ads, CBS News (2025) <https://www.cbsnews.com/news/meta-instagram-facebook-ads-nudify-deepfake-ai-tools-cbs-news-investigation/>

¹²How a ‘nudify’ site enabled AI-generated porn misuse, CNBC (2025) <https://www.cnn.com/2025/09/27/nudify-ai-generated-deepfake-fbi.html>

¹³Internet Watch Foundation, AI CSAM Report Update (2024) https://www.iwf.org.uk/media/opkpmx5q/iwf-ai-csam-report_update-public-jul24v11.pdf

material.¹⁴

Further data indicates that a significant proportion of reported sexual imagery involving young people now includes manipulated or AI-generated content. In the United Kingdom, approximately 19% of reports of nude or sexual imagery involving minors were found to involve digitally altered or AI-generated images, highlighting the scale of the issue.

This trend is particularly visible in school environments, where students have begun using AI tools to create explicit images of their peers. Research highlights that many educational institutions are unprepared to respond to such incidents, and existing legal frameworks often fail to clearly classify these acts as child sexual abuse material (CSAM), creating uncertainty in enforcement.¹⁵

Beyond minors, the broader pattern of harm is heavily gendered. Women and girls are disproportionately targeted by non-consensual deepfake content, often for purposes of sexual humiliation, harassment, or coercion. Reports indicate that the rise of such technologies is contributing to a chilling effect, discouraging women from participating freely in digital spaces due to fear of exploitation.¹⁶

From a legal standpoint, these developments expose significant gaps in existing frameworks. Traditional laws governing obscenity, defamation, and harassment were not designed to address AI-generated harm, particularly where the content is fabricated but the impact is real. The absence of clear legal recognition of non-consensual synthetic media, especially in cases involving minors, creates a regulatory vacuum that allows such practices to proliferate with limited accountability.

The scale, accessibility, and evolving sophistication of these technologies demand urgent legal attention. Without targeted regulation that explicitly addresses consent, identity, and digital exploitation, non-consensual deepfakes and nudification tools risk becoming normalized forms of abuse, particularly against women and children in the digital age.

4.2. Case Analysis: Taylor Swift Deepfake Controversy

The circulation of non-consensual AI-generated explicit images of Taylor Swift in January

¹⁴Internet Watch Foundation, AI nudification and child protection data (2025) <https://www.iwf.org.uk/news-media/news/ai-nudification-app-ban-and-on-device-protections-for-children-welcomed-following-iwf-campaign/>

¹⁵Addressing AI-generated CSAM in schools, Stanford HAI (2025) <https://hai.stanford.edu/policy/addressing-ai-generated-child-sexual-abuse-material-opportunities-for-educational-policy>

¹⁶AI deepfakes and impact on women online, The Guardian (2025) <https://www.theguardian.com/global-development/2025/nov/05/india-women-ai-deepfakes-internet-social-media-artificial-intelligence-nudify-extortion-abuse>

2024 marked a critical moment in the public and legal discourse surrounding deepfake exploitation. The images originated on online forums and quickly spread to mainstream platforms such as X (formerly Twitter), where a single post reportedly amassed over 45 million views within hours before it was removed. Despite platform policies prohibiting such content, the images continued to circulate widely, being repeatedly reposted across accounts.

One of the most striking aspects of the incident was the delay in effective content moderation. Reports indicate that it took approximately 17 to 18 hours for coordinated efforts, including those by Swift's team and platform intervention, to meaningfully curb the spread of the images.¹⁷ During this time, the content had already reached millions of users, demonstrating how quickly AI-generated exploitation can escalate beyond control. The inability to contain the spread in real time highlights a structural weakness in platform governance, where reactive moderation mechanisms fail to keep pace with viral dissemination.

In response to the widespread circulation, X took the unusual step of blocking search results for "Taylor Swift", preventing users from easily accessing or redistributing the content. While this measure reflected the seriousness of the situation, it also underscored the inadequacy of existing moderation systems, as such an extreme intervention would not be feasible in cases involving ordinary individuals.

The incident carries significant implications for digital safety and gendered harm. If a globally influential figure such as Taylor Swift with substantial legal resources and public visibility could not prevent the widespread circulation of non-consensual deepfake content for nearly an entire day, it raises serious concerns about the vulnerability of ordinary women and minors. Unlike public figures, most victims lack the institutional support, legal resources, or visibility required to respond effectively, leaving them exposed to prolonged and often irreversible harm. From a legal perspective, the case reveals critical gaps in existing frameworks. Current laws addressing obscenity, defamation, and online harassment are largely reactive and fail to account for the speed, scale, and synthetic nature of AI-generated content. More importantly, the absence of a clear, consent-based legal framework for synthetic media allows such exploitation to exist in a regulatory grey area. The Taylor Swift controversy thus highlights the urgent need for legal reforms that explicitly address non-consensual AI-generated content, platform accountability, and the protection of digital identity.

¹⁷Vittoria Elliott, The Taylor Swift Deepfake Incident Shows How Broken the Internet Is, WIRED (Jan. 2024) <https://www.wired.com/story/taylor-swift-deepfake-images-ai/>

4.3. The Grok AI Controversy: Mass-Scale Exploitation

The controversy surrounding Grok, an AI system developed by xAI, illustrates a broader and more systemic dimension of AI-enabled exploitation. Unlike isolated deepfake incidents, Grok was reportedly used by users to generate explicit and non-consensual images of real individuals through simple prompts, enabling the rapid and large-scale production of exploitative content.¹⁸ This represents a significant escalation from individual misuse to automated, mass-scale harm. A key aspect of this controversy lies in the response of Elon Musk, who addressed concerns by emphasizing that the system generates outputs based on user prompts and that responsibility lies primarily with users rather than the platform or developers.¹⁹ This framing reflects a broader industry approach that treats AI systems as neutral intermediaries, despite their active role in generating harmful content.

However, such a position raises important legal questions regarding liability. If an AI system is designed and deployed in a manner that enables the creation of harmful or exploitative content, the extent to which developers and platforms can disclaim responsibility becomes highly contested. The Grok controversy highlights the inadequacy of existing legal frameworks, which are primarily structured around direct human actions and are ill-equipped to address harm generated through automated systems.

Furthermore, the scale at which such systems operate significantly amplifies the potential for harm. The ability to generate large volumes of explicit content in a short period not only increases the reach of exploitation but also makes traditional enforcement mechanisms ineffective. This shifts the legal focus from individual wrongdoing to systemic accountability, where the design and governance of AI systems must be critically examined.

The Grok incident therefore underscores a fundamental regulatory gap that while AI systems are capable of facilitating large-scale exploitation, existing legal doctrines have yet to clearly define the extent of responsibility borne by developers, platforms, and users. Addressing this gap is essential to ensuring that technological innovation does not come at the cost of basic rights to privacy, dignity, and consent.

¹⁸AI Tools Are Being Used to Generate Explicit Images at Scale, The Guardian (2024) <https://www.theguardian.com/technology/2024>

¹⁹Elon Musk Responds to Grok AI Concerns, Times of India (2024) <https://timesofindia.indiatimes.com/technology/social/elon-musk-defends-grok-image-generation-says-not-aware-of-any-naked-underage-images-generated/articleshow/126538422.cms>

5. Beyond Consent: Extreme and Emerging Harms

5.1. AI-Generated Child Sexual Abuse Material

One of the most alarming developments in the misuse of artificial intelligence is the rise of AI-generated child sexual abuse material (CSAM). Unlike traditional forms of CSAM, which involve real victims, AI-generated content can create entirely synthetic images of minors engaged in sexually explicit acts. While these images may not depict real children, they raise serious legal and ethical concerns due to their exploitative nature and potential to normalize harmful behavior.

Recent findings by the Internet Watch Foundation highlight the scale of this emerging threat. In its 2023–2024 reports, the organization identified thousands of AI-generated images depicting child sexual abuse, with a significant portion categorized as the most severe forms of abuse.²⁰ The report further noted that such content is becoming increasingly realistic, making it more difficult to distinguish from actual abuse material and complicating detection efforts. A particularly troubling aspect of this development is the role of AI in lowering the barriers to creating such content. Previously, the production of CSAM required direct abuse or access to illicit networks. However, AI now allows individuals to generate such material independently, without direct contact with victims. This raises complex ethical questions: even if no real child is involved in the creation of the image, does the act of generating and consuming such content still constitute harm?

From a legal perspective, jurisdictions differ in their treatment of synthetic CSAM. Some legal systems criminalize all forms of child sexual imagery regardless of whether a real child is involved, while others lack clear provisions addressing AI-generated material. This creates a regulatory gap, where perpetrators may exploit legal ambiguities to avoid liability.²¹ Furthermore, law enforcement agencies face practical challenges in distinguishing between real and synthetic content, which can hinder investigations and delay justice.

Beyond legality, the broader concern lies in normalization. The proliferation of AI-generated CSAM risks desensitizing individuals and reinforcing harmful attitudes toward minors. Reports have also indicated instances where such content is used as a precursor to coercion or grooming, including cases where perpetrators use fabricated images to blackmail minors into producing real explicit material.²² This demonstrates that even “synthetic” harm can have very real

²⁰Internet Watch Foundation, AI-Generated Child Sexual Abuse Imagery Report (2023–2024) <https://www.iwf.org.uk/resources/trends-and-data/ai-generated-csam/>

²¹AI-Generated Child Sexual Abuse Images Pose Legal Challenges, BBC News (2023) <https://www.bbc.com/news/technology-65910048>

²²Internet Watch Foundation, Annual Report & Emerging Threats (2024)

consequences.

5.2. Synthetic Identities and the Commodification of Women

In addition to explicit content generation, artificial intelligence has also enabled the creation of entirely synthetic identities that reflect and reinforce problematic social dynamics. One such example is the emergence of hyper-realistic, AI-generated female personas designed to cater to specific ideological or aesthetic preferences. The so-called “Jessica Foster” phenomenon refers to the creation and circulation of AI-generated representations of an idealized woman, often portrayed as submissive, hyper-feminine, and aligned with certain political or cultural narratives.²³

These synthetic identities are not merely fictional characters; they function as tools of influence and commodification. By presenting an artificial image of femininity that aligns with particular expectations, they contribute to the objectification of women and the reinforcement of harmful stereotypes. Unlike traditional media portrayals, AI-generated identities can be endlessly customized, replicated, and disseminated, amplifying their impact across digital platforms.

The implications of this phenomenon extend beyond individual harm to broader socio-political concerns. Synthetic identities can be used to shape public discourse, influence opinions, and promote specific ideological narratives. In some cases, they are deployed in online communities to engage users, build trust, and subtly reinforce gendered expectations. This blurs the line between authenticity and fabrication, raising concerns about manipulation and consent in digital interactions.

From a legal standpoint, the regulation of synthetic identities remains largely underdeveloped. Existing laws on impersonation or identity theft are not easily applicable where the individual being represented does not exist. At the same time, the harm caused is indirect but significant, contributing to a culture that normalizes the commodification and objectification of women.

Moreover, this phenomenon reflects a deeper shift in how identity itself is understood in the digital age. When artificial representations can be created and circulated at scale, the distinction between real and synthetic begins to blur, challenging traditional legal concepts of personhood, dignity, and representation.

<https://www.iwf.org.uk/annual-report/>

²³AI-Generated Jessica Foster: The Rise of Digital Political Personas, mtsoln (2026) <https://mtsoln.com/blog/ai-news-727/thousands-have-swooned-over-this-maga-dream-girl-shes-made-with-ai-5984>

5.3. AI as a Tool for Gender-Based Violence (GBV)

The misuse of artificial intelligence must also be understood within the broader framework of gender-based violence (GBV). International organizations, including UN Women, have recognized that digital technologies are increasingly being used to perpetrate and amplify violence against women and girls in online spaces. This includes harassment, stalking, non-consensual image sharing, and AI-generated sexual exploitation.

According to UN reports, technology-facilitated gender-based violence is one of the fastest-growing forms of violence globally, with digital tools enabling new methods of abuse that are scalable, anonymous, and difficult to regulate.²⁴ AI systems, particularly those capable of generating realistic images and videos, significantly intensify this problem by allowing perpetrators to create and disseminate harmful content with minimal effort.

The gendered nature of AI misuse is particularly evident in the disproportionate targeting of women and girls. Deepfake pornography, nudification technologies, and synthetic exploitation are overwhelmingly directed at female subjects, reflecting existing societal inequalities. The United Nations has emphasized that such forms of digital abuse are not isolated incidents but part of a continuum of violence that extends from offline to online environments.²⁵

From a legal and policy perspective, this framing is crucial. Viewing AI misuse through the lens of GBV shifts the focus from isolated technological harms to systemic patterns of discrimination and violence. It highlights the need for comprehensive legal frameworks that not only regulate technology but also address underlying gender inequalities.

Moreover, the integration of AI into everyday digital interactions risks normalizing such forms of abuse if left unregulated. Without clear legal standards and accountability mechanisms, AI technologies may continue to reinforce existing power imbalances, further marginalizing vulnerable groups. Recognizing AI-enabled exploitation as a form of gender-based violence is therefore essential to developing effective legal and policy responses.

²⁴UN Women, Online and ICT-Facilitated Violence Against Women and Girls
<https://www.unwomen.org/en/what-we-do/ending-violence-against-women/faqs/online-and-ict-facilitated-violence>

²⁵United Nations, Technology-Facilitated Gender-Based Violence Reports
<https://www.un.org/en/observances/ending-violence-against-women-day>

6. Case Study Analysis: The Collien Fernandes Controversy (Germany)

Facts and Allegations

The case involving Collien Fernandes represents one of the most significant real-world examples of AI-enabled sexual exploitation exposing legal gaps. In 2026, Fernandes publicly accused her ex-husband, Christian Ulmen, of engaging in a prolonged campaign of digital abuse, including the creation and distribution of AI-generated pornographic deepfakes depicting her without consent.²⁶

According to her allegations, Ulmen created fake online profiles in her name and used them to share explicit content and interact with others as if he were her. These actions reportedly continued over several years, involving the circulation of manipulated images and videos as well as impersonation across digital platforms.

Fernandes described the experience as deeply violating, referring to it as “virtual rape,” a term that has since become central to the legal and public discourse surrounding the case. She further pursued legal action not only in Germany but also in Spain, citing stronger legal protections against digital gender-based violence.

It is important to note that these allegations have been denied by Ulmen, whose legal representatives have challenged the claims and initiated action against media reporting. However, regardless of the outcome, the case itself has triggered a broader legal and societal debate.²⁷

Concept of “Virtual Sexual Violence”

The Fernandes case has brought the concept of “virtual sexual violence” into mainstream legal discussion. This refers to forms of sexual harm carried out through digital means, particularly through AI-generated or manipulated content that violates a person’s autonomy, dignity, and identity.

Traditionally, sexual offences in law are grounded in physical acts. However, this case challenges that assumption by demonstrating that harm need not be physical to be deeply invasive. The creation and circulation of explicit deepfakes effectively simulate sexual violation, causing psychological trauma, reputational harm, and loss of control over one’s own

²⁶Collien Fernandes Deepfake Porn Allegations Spark Debate on Digital Violence, The Guardian (Mar. 30, 2026), <https://www.theguardian.com/world/2026/mar/30/collien-fernandes-deepfake-porn-allegations-digital-violence-against-women>

²⁷Collien Fernandes Case Raises Questions on Deepfake Laws, The Times (2026), <https://www.thetimes.com/world/europe/article/collien-fernandes-deepfake-germany-christian-ulmen-8k8wsstw6>

identity.

The idea of “virtual rape,” as articulated by Fernandes, reflects this shift. Even though the acts were digitally constructed, the experience of violation was real and severe. This challenges existing legal frameworks, which often fail to recognize non-physical forms of sexual violence with the same seriousness as physical offences.²⁸

Moreover, the case illustrates how such technologies can be weaponized within intimate relationships, transforming AI tools into instruments of coercion and abuse. This expands the scope of gender-based violence beyond public or anonymous actors to include intimate partner digital abuse, a dimension that existing laws are largely unprepared to address.

Public Outrage and Legal Reform

The case triggered widespread public outrage across Germany, evolving into a national movement demanding stronger protections against digital sexual violence. Thousands of people participated in protests, particularly in Berlin, calling for recognition of AI-enabled abuse as a serious legal offence.

More than 10,000 demonstrators reportedly gathered in solidarity with victims of digital exploitation, highlighting the broader societal resonance of the issue. The case also prompted a coalition of prominent women and activists to issue formal demands for legislative reform, emphasizing the urgent need to address gaps in the law.²⁹

At the governmental level, the controversy led to concrete steps toward legal reform. Germany’s Justice Minister announced plans to introduce legislation specifically targeting the creation and distribution of non-consensual deepfake pornography, with proposed penalties including imprisonment. The proposed reforms also aim to improve victim protection, enhance investigative powers, and increase platform accountability.³⁰

The case has also influenced broader European discussions, with calls for EU-level regulation of AI-based “nudification” tools and synthetic sexual content.

However most significantly, the Fernandes controversy exposed a fundamental weakness in existing legal systems: while the distribution of explicit material may be punishable, the creation of synthetic sexual content itself often remains outside clear criminal liability. This

²⁸Why the Collien Fernandes Case Matters for Digital Violence, Campact (2026), <https://www.campact.de/blog/2026/03/deepfakes-warum-der-fall-collien-fernandes-uns-alle-angeht/>

²⁹Thousands Rally in Berlin Against Deepfake Pornography, South China Morning Post (2026), <https://www.scmp.com/news/world/europe/article/3347509/thousands-rally-berlin-against-online-sexual-violence-pornographic-deepfakes>

³⁰Germany Plans Legal Reform to Tackle Deepfake Pornography, Heise Online (2026), <https://www.heise.de/en/news/Protection-against-digital-violence-draft-in-the-pipeline-11219979.html>

gap allows harmful conduct to occur without immediate legal consequences, reinforcing the need for proactive, rather than reactive, regulation.

7. The Question of Liability: Who Bears Responsibility?

7.1. User, Developer, or Platform?

One of the most complex legal challenges in the context of AI-generated harm is determining who should be held responsible. Unlike traditional legal situations where liability can be clearly attributed to a single actor, AI systems involve multiple participants. These include users who generate content, developers who design the systems, and platforms that enable its distribution. At a basic level, users who create and share non-consensual deepfake content can be held liable under existing laws relating to obscenity, defamation, or harassment. However, placing responsibility solely on users does not fully reflect how such harm occurs. AI developers create tools that are capable of producing harmful outputs, and platforms provide the infrastructure through which this content spreads rapidly.

This has led to different approaches to liability. One approach focuses on user responsibility, treating AI systems as neutral tools. Another approach supports a shared model of liability, where developers and platforms may also be held accountable, especially when adequate safeguards are not in place. The issue becomes even more complicated when considering whether AI developers should be treated as intermediaries under existing laws. Some argue that they simply provide tools, while others believe that their role in designing systems that generate content makes them more directly involved in the harm. This lack of clarity reflects a larger gap in how the law currently understands AI systems.³¹

7.2. Intermediary Responsibility and Platform Governance

Social media platforms play a central role in the spread of AI-generated content. They act as intermediaries by hosting and distributing user-generated material, including deepfakes. Traditionally, many legal systems have granted these platforms protection from liability, provided they act responsibly when notified of harmful content.

In India, Section 79 of the Information Technology Act, 2000 provides safe harbour protection to intermediaries. This protection applies as long as platforms act as neutral facilitators and remove unlawful content when required. However, recent amendments to the IT Rules have

³¹Are AI Developers Intermediaries Under Synthetic Media Rules?, Medianama (2025) <https://www.medianama.com/2025/11/223-ai-developers-intermediaries-meity-synthetic-information-rules/>

placed greater responsibility on platforms, particularly in relation to AI-generated content. Platforms are now expected to take proactive steps such as labelling synthetic media, ensuring transparency, and responding quickly to harmful content.³²

These changes also reflect the urgency of the issue. The time frame for removing harmful content has been significantly reduced in certain situations, recognizing how quickly deepfake material can spread. If platforms fail to comply with these obligations, they risk losing their safe harbour protection and may be held directly liable.³³

A similar debate exists in the United States under Section 230 of the Communications Decency Act, which provides broad immunity to platforms for user-generated content. While this provision has supported the growth of digital platforms, it is increasingly being questioned in the context of AI. Platforms are no longer just hosting content. In many cases, they are actively shaping or even generating it through integrated AI systems. This makes it harder to treat them as passive intermediaries

.

7.3. The Limits of Existing Liability Doctrines

Existing legal doctrines struggle to keep up with the realities of AI-generated harm. Traditional frameworks are based on clear distinctions between creators, publishers, and intermediaries. AI systems blur these roles by producing content through a combination of user input and automated processes.

One major limitation lies in the concept of safe harbour protection. This principle was developed at a time when platforms were seen as passive hosts. Today, with the integration of AI tools, platforms may play a more active role in enabling harmful content. This challenges the idea that they should be shielded from liability in all cases.

Another issue is the difficulty of establishing causation. Harm caused by AI-generated content often involves multiple actors and processes. A user provides input, the system generates content, and the platform distributes it. This makes it difficult to assign responsibility to any single party, which weakens the effectiveness of traditional legal approaches.

There is also the problem of speed and scale. AI-generated content can be created and shared within minutes, while legal remedies take much longer to apply. By the time action is taken, the damage is often already done.

³²AI-Generated Content and Combating Deepfakes: India's New Rules, Nishith Desai Associates (2026) <https://www.nishithdesai.com/research-and-articles/hotline/technology-law-analysis/ai-generated-content-and-combating-deepfakes-what-indias-new-rules-mean-for-global-platforms-15532>

³³India IT Rules and AI Liability Implications, Vaish Associates (2026) <https://www.vaishlaw.com/india-it-rules-before-and-after-amendment-ai-generated-co>

Because of these challenges, there is growing support for new approaches to liability. Some scholars suggest moving towards risk-based or design-based models, where responsibility is linked to how systems are built and whether harm could have been prevented. This would require a shift in how the law evaluates responsibility, focusing not just on individual actions but also on the role of technology in enabling harm.

At present, the law has not fully adapted to these changes. The result is a system where responsibility is often unclear and victims are left without effective remedies. Addressing this gap is essential if legal frameworks are to remain relevant in the age of artificial intelligence.

8. The Existing Legal Framework: India

8.1. Statutory Provisions in India

India does not yet have a specific law dealing directly with AI-generated deepfakes or synthetic sexual content. Instead, such cases are addressed through existing provisions under the Information Technology Act, 2000 and the Indian Penal Code, 1860.

Under the IT Act, Section 66E penalizes violation of privacy, while Sections 67 and 67A deal with the publication of obscene and sexually explicit content online. Similarly, provisions under the IPC such as Section 354C (voyeurism), Section 499 (defamation), and Section 509 (insult to modesty) may be applied depending on the nature of harm.

Recently, the government has taken steps to address deepfakes through advisories issued by the Ministry of Electronics and Information Technology, directing platforms to remove such content quickly and ensure user accountability.³⁴ However, the current framework remains indirect and largely reactive.

8.2. Constitutional Dimensions of Privacy and Dignity

The issue of AI-generated exploitation directly engages constitutional rights. In Justice K.S. Puttaswamy v. Union of India, the Supreme Court recognized the right to privacy as a fundamental right under Article 21.

This includes the right to control one's identity and personal data. Deepfake content violates this by using a person's image without consent, often in harmful or sexualized contexts. It also affects dignity, which is an essential part of the right to life.

However, despite this strong constitutional foundation, there is limited judicial guidance on how these principles apply specifically to AI-generated content.

³⁴MeitY Advisory on Deepfakes and Synthetic Media, Press Information Bureau (2023–2024) <https://pib.gov.in/PressReleasePage.aspx?PRID=1987652>

8.3. Doctrinal Gaps and Enforcement Challenges

Significant gaps remain in India's legal approach. There is no dedicated legislation addressing deepfakes or AI-generated sexual exploitation, forcing courts to rely on outdated provisions. Enforcement is another major challenge. Such content spreads rapidly and often across borders, making it difficult to trace perpetrators and take timely action. Technical challenges in detecting and proving AI-generated content further complicate the process. There are ongoing discussions around the proposed Digital India Act, which is expected to address emerging technologies like AI.³⁵ However, until such reforms are enacted, the legal framework remains fragmented and insufficient.

9. Comparative Legal Perspectives: Emerging Global Responses

9.1. Denmark's Copyright-Based Approach to Identity Protection

Denmark has emerged as a global frontrunner in addressing AI-generated deepfakes through an innovative legal approach grounded in copyright law. In 2025, the Danish government proposed amendments to its Copyright Act that would grant individuals legal rights over their own image, voice, and physical likeness. This proposal, expected to take effect around 2026, is widely considered the first of its kind in Europe.³⁶

Under this framework, any unauthorized use of a person's likeness through AI-generated content would be treated as a violation of copyright. Individuals would have the right to demand the removal of such content and seek compensation for misuse.³⁷ Importantly, the law also introduces obligations for platforms, which may face penalties if they fail to act upon such requests.³⁸

This approach represents a significant shift in legal thinking. Instead of relying solely on privacy or defamation laws, Denmark treats identity itself as a form of intellectual property. The idea is simple but powerful: an individual should have ownership over their own face, voice, and digital representation.

³⁵Digital India Act to Replace IT Act: Government Plans AI Regulation, The Hindu (2024) <https://www.thehindu.com/sci-tech/technology/digital-india-act-it-act-replacement-ai-regulation/article67634567.ece>

³⁶Denmark's Copyright Amendment Against Deepfakes, Denneweyer IP Blog (2026) <https://www.denneweyer.com/ip-blog/news/a-new-sense-of-self-denmarks-copyright-amendment-against-deepfakes/>

³⁷Denmark to Tackle Deepfakes by Giving People Copyright Over Their Features, The Guardian (2025) <https://www.theguardian.com/technology/2025/jun/27/deepfakes-denmark-copyright-law-artificial-intelligence>

³⁸Denmark Deepfake Law: Copyright Protection for Identity, LiveLaw (2025) <https://www.livelaw.in/articles/regulation-of-deepfakes-and-denmark-law-on-copyright-analysis-297718>

At the same time, the law attempts to balance freedom of expression by allowing exceptions for parody and satire. However, this raises further questions about how to distinguish legitimate expression from harmful misuse. While Denmark's model is proactive and rights-based, its practical implementation will depend heavily on enforcement and interpretation.

9.2. European and International Regulatory Trends

Across Europe and globally, governments are increasingly recognizing the risks posed by deepfakes and AI-generated content. However, most regulatory approaches differ significantly from Denmark's model.

At the European Union level, frameworks such as the AI Act and the Digital Services Act focus primarily on transparency and platform accountability rather than direct prohibition. These laws require AI-generated content to be labelled and impose due diligence obligations on platforms to manage harmful content.³⁹ However, they do not always provide clear mechanisms for victims to remove deepfake content once it has spread online.

This reflects a broader trend toward preventive regulation. Instead of waiting for harm to occur, these frameworks aim to reduce risks through measures such as watermarking, disclosure requirements, and algorithmic accountability. While these steps are important, critics argue that they are insufficient in cases involving non-consensual sexual content, where immediate removal and victim protection are crucial.

In contrast, some jurisdictions continue to rely on reactive legal models, addressing harm only after it occurs. This includes the use of existing laws on harassment, defamation, or obscenity, which often fail to capture the unique nature of AI-generated harm.

The growing divergence between preventive and reactive approaches highlights a key tension in global regulation. Preventive frameworks focus on systemic risk and technological design, while reactive models emphasize individual liability and post-harm remedies. Neither approach is fully sufficient on its own.

Denmark's model stands out because it attempts to bridge this gap. By granting individuals enforceable rights over their identity, it combines elements of both prevention and remedy. As deepfake technology continues to evolve, this hybrid approach may influence future international regulation.

³⁹Deepfake Regulation and EU Legal Frameworks, Overview of EU AI Act & DSA
<https://en.wikipedia.org/wiki/Deepfake>

10. Findings: The Structural Failure of Existing Legal Systems

The analysis across jurisdictions shows that the problem is not just a lack of laws, but a deeper structural failure in how legal systems respond to AI-generated harm. Existing laws are scattered, reactive, and not built for technologies that can create and spread harmful content within minutes.

One major issue is the inadequacy of current laws. Most provisions under criminal and cyber law were created before generative AI became widespread. As a result, they are applied indirectly to deepfake cases, which often leads to weak or inconsistent outcomes.⁴⁰

Another key gap is the absence of consent-focused regulation. In many cases, the law looks at whether content is obscene or defamatory, but misses the core issue, which is the use of a person's identity without permission. Without clearly recognizing consent as central, many harmful acts fall into legal grey areas.⁴¹

Finally, there is a growing gap between technology and law. AI tools are advancing quickly, becoming more accessible and more realistic, while legal reforms take time. This delay means that by the time action is taken, the harm has already spread and is difficult to reverse.

Overall, the issue is not just about regulating technology, but about rethinking how the law understands harm, consent, and protection in the digital age.

11. Recommendations: Reclaiming Control in the Age of Synthetic Reality

The law cannot afford to remain silent while technology is used to strip individuals, especially women and girls, of their dignity, identity, and sense of safety. What is at stake is not just regulation, but the basic right to exist without fear of being digitally violated.

There is an urgent need for the enactment of dedicated deepfake legislation that directly recognizes non-consensual AI-generated content as a serious offence. The law must clearly state that creating or circulating such material without consent is not a grey area. It is harm. It is violation. And it must carry strict and enforceable consequences.

⁴⁰The Challenge of Regulating Deepfakes and AI Content, World Economic Forum (2024) <https://www.weforum.org/agenda/2024/02/deepfakes-ai-regulation-challenges/>

⁴¹Regulating Deepfakes: Legal and Ethical Challenges, Brookings Institution (2023) <https://www.brookings.edu/articles/deepfakes-and-the-law/>

At the same time, legal systems must move toward recognizing identity and likeness as protected rights. A person's face, voice, and digital presence are not public property. They are deeply personal. The unauthorized use of these elements must be treated as a direct violation of autonomy, not just an indirect form of harm.

There is also a pressing need for gender-sensitive legal reform. The reality cannot be ignored. Women and girls are being disproportionately targeted, silenced, and humiliated through AI tools. Laws that treat these harms as gender-neutral fail to capture their true impact. Recognizing AI-enabled abuse as a form of gender-based violence is not optional. It is necessary for justice.

Equally important is holding platforms accountable. Technology companies cannot continue to benefit from engagement while distancing themselves from the harm their systems enable. If platforms have the power to amplify content, they also have the responsibility to prevent its abuse. Faster takedowns, stronger safeguards, and real consequences for inaction must become the norm, not the exception.

Finally, this is not a problem that can be solved within borders. International cooperation is essential. Without it, perpetrators will continue to exploit legal gaps across jurisdictions. There must be collective global effort to establish standards that protect individuals, no matter where they are.

This is not just about adapting the law to technology. It is about deciding what kind of digital world we are willing to accept. One where identity can be stolen, altered, and violated without consequence, or one where dignity, consent, and safety are non-negotiable.

12. Conclusion

The rise of AI-generated content has fundamentally changed how identity, consent, and harm are understood in the digital age. What was once limited to physical or direct forms of abuse can now be replicated, altered, and spread at scale through synthetic media.

This paper has shown that existing legal systems are struggling to keep pace. While laws on privacy, obscenity, and online harm provide some level of protection, they are not sufficient to address the unique challenges posed by AI. The absence of consent-based regulation, the lack

of gender-sensitive frameworks, and the slow pace of legal reform have created a gap where harm continues to grow.

The urgency of the issue cannot be overstated. As AI technologies become more accessible and more advanced, the potential for misuse will only increase. Without timely and effective legal intervention, individuals will continue to lose control over their own identity and digital presence.

At the same time, it is important to maintain a balance. AI is not inherently harmful and has the potential to drive innovation and progress. The goal of regulation should not be to restrict technology entirely, but to ensure that its development and use respect fundamental rights such as privacy, dignity, and consent.

Ultimately, the challenge is not just technological or legal, but deeply human. Reclaiming control in the age of synthetic reality means ensuring that individuals remain at the center of the digital world, with their rights protected and their identity respected.

References / Bibliography

I. Judicial Decisions

Justice K.S. Puttaswamy v. Union of India, (2017) 10 SCC 1.

II. Statutes and Legal Frameworks

Information Technology Act, 2000 (India).

Indian Penal Code, 1860 (India).

Section 230, Communications Decency Act, 47 U.S.C. § 230 (United States).

III. Government Reports, Policy Documents, and Institutional Sources

Internet Watch Foundation, AI-Generated Child Sexual Abuse Material Trends Report (2023–2024),

<https://www.iwf.org.uk/resources/trends-and-data/ai-generated-csam/>

Internet Watch Foundation, Annual Report (2024),

<https://www.iwf.org.uk/annual-report/>

UN Women, Online and ICT-Facilitated Violence Against Women and Girls,

<https://www.unwomen.org/en/what-we-do/ending-violence-against-women/faqs/online-and-ict-facilitated-violence>

Press Information Bureau, Government of India, MeitY Advisory on Deepfakes and Synthetic Media (2023–2024),

<https://pib.gov.in/PressReleasePage.aspx?PRID=1987652>

IV. News Articles and Media Reports

Clare Duffy, Fake Explicit Images of Taylor Swift Show the Dark Side of AI, CNN (2024),

<https://edition.cnn.com/2024/01/26/tech/taylor-swift-ai-images/index.html>

Vittoria Elliott, The Taylor Swift Deepfake Incident Shows How Broken the Internet Is, WIRED (2024),

<https://www.wired.com/story/taylor-swift-deepfake-images-ai/>

BBC News, X Blocks Searches for Taylor Swift After Explicit AI Images Spread (2024),

<https://www.bbc.com/news/technology-68070367>

The Guardian, Denmark to Tackle Deepfakes by Giving People Copyright Over Their Features (2025),

<https://www.theguardian.com/technology/2025/jun/27/deepfakes-denmark-copyright-law-artificial-intelligence>

The Guardian, Collien Fernandes Deepfake Porn Allegations Spark Debate on Digital Violence (2026),

<https://www.theguardian.com/world/2026/mar/30/collien-fernandes-deepfake-porn-allegations-digital-violence-against-women>

Reuters, German Deepfake Porn Case Sparks Protests and Pressure for Law Reform (2026),

<https://www.reuters.com/business/media-telecom/german-deepfake-porn-case-sparks-protests-pressure-change-law-2026-03-26/>

South China Morning Post, Thousands Rally in Berlin Against Deepfake Pornography (2026),

<https://www.scmp.com/news/world/europe/article/3347509/thousands-rally-berlin-against-online-sexual-violence-pornographic-deepfakes>

The Week, German Deepfake Scandal and Legal Gaps (2026),

<https://theweek.com/law/german-deepfake-scandal-ai-pornography-protest>

V. Academic Articles and Research Papers

Brookings Institution, Regulating Deepfakes: Legal and Ethical Challenges (2023),

<https://www.brookings.edu/articles/deepfakes-and-the-law/>

World Economic Forum, The Challenge of Regulating Deepfakes and AI Content (2024),
<https://www.weforum.org/agenda/2024/02/deepfakes-ai-regulation-challenges/>
Henderson et al., Where's the Liability in Harmful AI Speech? (2023),
<https://arxiv.org/abs/2308.04635>

VI. Legal Commentary and Industry Analysis

Nishith Desai Associates, AI-Generated Content and Combating Deepfakes: What India's New Rules Mean (2026),
<https://www.nishithdesai.com/research-and-articles/hotline/technology-law-analysis/ai-generated-content-and-combating-deepfakes-what-indias-new-rules-mean-for-global-platforms-15532>
LiveLaw, AI-Generated Content and Deepfake Regulation in India (2026),
<https://www.livelaw.in/articles/ai-generated-content-deepfakes-524064>
Vaish Associates, India IT Rules and AI Liability Implications (2026),
<https://www.vaishlaw.com/india-it-rules-before-and-after-amendment-ai-generated-content/>
Dennemeyer IP Blog, Denmark's Copyright Amendment Against Deepfakes (2026),
<https://www.dennemeyer.com/ip-blog/news/a-new-sense-of-self-denmarks-copyright-amendment-against-deepfakes/>