

REGULATION OF SOCIAL MEDIA IMAGERY IN THE AGE OF ARTIFICIAL INTELLIGENCE: CHALLENGES AND FRAMEWORKS

AUTHOR – MOHD AKASH, RESEARCH SCHOLAR, FACULTY OF LEGAL STUDIES, MJP ROHILKHAND UNIVERSITY BAREILLY, UP (INDIA).

BEST CITATION – MOHD AKASH, REGULATION OF SOCIAL MEDIA IMAGERY IN THE AGE OF ARTIFICIAL INTELLIGENCE: CHALLENGES AND FRAMEWORKS, *INDIAN JOURNAL OF LEGAL REVIEW (IJLR)*, 5 (12) OF 2025, PG. 935-940, APIS – 3920 – 0001 & ISSN – 2583-2344

Abstract

Artificial intelligence (AI) has transformed the digital landscape, notably across social media platforms. These platforms have incorporated AI technology into their algorithms to improve user experience. However, the impact of artificial intelligence on social media is debatable. To sustain user confidence, social media platforms should ensure that their AI practices are transparent and ethical. The impact of AI on social media content is considerable and diverse, allowing for personalised content recommendations, automated content development, and real-time content analysis. However, there are also concerns regarding algorithmic bias and the possibility of job displacement. As AI technology advances, it is critical that ethical issues and social responsibility are prioritised in the development and application of AI in social media marketing. The impact of AI on social media posts is a multifaceted problem that has both positive and negative consequences. While AI algorithms can improve the user experience by offering personalised content, there are worries that they may also promote disinformation and create filter bubbles. With the growing development of AI-generated images, ranging from deepfakes to algorithmically augmented visuals, regulating such content presents significant legal, ethical, and technological issues. This article reviews the current frameworks and shortcomings in regulating AI-generated imagery on social media, investigates international and Indian legal methods, evaluates policy solutions, and recommends for a human-rights-based, technologically adaptive regulatory strategy. The goal is to promote accountability, transparency, and user protection against misinformation, privacy infringement, and the abuse of synthetic media.

Keywords: Artificial Intelligence, Social Media, Content, Impact, Machine Learning, Natural Language Processing, Algorithms, Data Analysis, Personalization, Automation.

Introduction

In recent years, Artificial Intelligence (AI) has developed as an effective tool for developing and managing social media content. AI algorithms can analyse enormous amounts of data to uncover patterns and insights that humans may overlook, resulting in more efficient and effective content

creation. However, the impact of artificial intelligence (AI) on social media material is still

being debated. Previous research has proven that generative artificial intelligence (AI) comprises systems that produce various forms of material, such as text, images, sounds, videos, and codes, using algorithms, models, and rules.² It works inside an unsupervised or partially supervised machine learning environment, using statistics and probability to create artificial artefacts. This study will look into the influence of artificial intelligence on social media content, as well as the implications for

content providers and users. Artificial intelligence (AI) is becoming more ubiquitous in all facets of our lives, including social media. Social media platforms like Facebook, Twitter, and Instagram have incorporated AI technology into their algorithms to improve user experience. AI algorithms can analyse user data and behaviour to deliver personalised information and adverts, improve search results, and control content. As a result, AI has a considerable influence on the content that users encounter on social media networks.

The impact of AI on social media material has far-reaching consequences for individuals, society, and democratic institutions. The spread of disinformation and the construction of filter bubbles can damage people's capacity to make educated decisions and participate in public debates. It is critical to understand how AI algorithms influence social media content and how this affects individuals and society. The purpose of this study is to look into the impact of AI on social media content, with a particular focus on the potential for AI to promote misinformation and filter bubbles. The incorporation of AI technology into social media platforms has resulted in a variety of possible benefits, but it has also raised worries about its impact on social media content.

One of the biggest concerns is the potential for AI systems to spread misinformation and create filter bubbles. Filter bubbles are the phenomena in which users are only exposed to content that confirms their own opinions, resulting in polarisation and a lack of diversity of viewpoints. Misinformation spreads swiftly on social media, and using AI algorithms to target people with personalised material may worsen the issue. Therefore, it is necessary to explore the impact of AI on social media content.

The rise of artificial intelligence, particularly generative AI, has had a substantial impact on the nature of content production and distribution on social media. AI-generated images abound on platforms such as Facebook, Instagram, TikTok, and Twitter, some of which

are innocuous and others fraudulent. These visuals have the potential to entertain, educate, deceive, or manipulate public opinion. Deepfakes, face-swapped pictures, and AI-enhanced propaganda blur the distinction between fact and fiction. Despite the increasing prominence of such imagery, legal and regulatory remedies are uneven and inadequate. This research fills this gap by providing a complete analysis of regulatory problems and solutions to AI-generated photography on social media.

Understanding AI-Generated Imagery on Social Media

AI-generated images are graphics created, modified, or transformed using algorithms with no—or little—human intervention. ChatGPT, developed by OpenAI, is a popular example of generative AI. OpenAI has produced various models, including GPT, GPT-2, GPT-3, GPT-4, and image pretraining iGPT, which are gaining popularity worldwide. Other software companies, including DeepMind, Google, SenseTime, and Ali, have created their own language models and invested in enormous datasets. Other common instances are:

Deepfakes: Synthetic media where a person's likeness is superimposed onto another's face or body using AI.

GANs (Generative Adversarial Networks): AI systems that generate hyper-realistic images from training data.

Image Enhancements: AI tools that modify facial expressions, skin tone, or background in photos.

Avatars & Virtual Influencers: Entirely artificial characters crafted to appear and behave like real humans.

These images have been used for harmless fun, political propaganda, satire, pornography, misinformation, and even cybercrime. Their prevalence on social media raises urgent concerns about misinformation, consent, privacy, defamation, and democratic integrity.

Technological Aspects of Generative AI

Generative AI refers to technologies that use models, typically deep learning neural networks, to generate information patterns that are comparable to the data on which they were trained. Examples include text, graphics, audio, and video. This type of AI differs dramatically from "narrow" or "applied" AI—systems that may excel at pattern recognition or data analysis but do not generate new material. The rapid evolution of generative AI has been fuelled by the creation of complex algorithms and models such as generative adversarial networks (GANs), variational autoencoders (VAEs), and transformer-based models like GPT-3. These technologies have permitted the development of very realistic and diversified content, including text, graphics, sound, and video.³ Previous works highlight the unsupervised or partially supervised nature of generative AI, which allows these systems to learn patterns and distributions from vast amounts of training data.¹⁷²⁰ This has led to the development of large language models like ChatGPT, which can generate humanlike text and engage in context-aware conversations.¹⁷²¹ Researchers have also explored the scalability and performance of generative AI systems, addressing challenges such as mode collapse and training stability.¹⁷²² Advances in computational resources and optimization techniques have further enabled the training of larger and more complex generative models.¹⁷²³

¹⁷²⁰ Mario Lucic, Karol Kurach, Marcin Michalski, Sylvain Gelly, and Olivier Bousquet, "Are Gans Created Equal? A Large-Scale Study," *Advances in Neural Information Processing Systems* 31 (2018), <https://proceedings.neurips.cc/paper/7350-are-gans-created-equal-a-large-scale-study>.

¹⁷²¹ Daniel Adiwardana, Minh-Thang Luong, David R. So, Jamie Hall, Noah Fiedel, Romal T. Hoppilan, Zi Yang, et al., "Towards a Human-like Open-Domain Chatbot," *arXiv* (February 27, 2020), <https://doi.org/10.48550/arXiv.2001.09977>.

¹⁷²² Lars Mescheder, Andreas Geiger, and Sebastian Nowozin, "Which Training Methods for Gans Do Actually Converge?" In *Proceedings of the 35th International Conference on Machine Learning* 80 (2018): 3481–90, <https://proceedings.mlr.press/v80/mescheder18a>.

¹⁷²³ Tero Karras, Samuli Laine, and Timo Aila, "A Style-Based Generator Architecture for Generative Adversarial Networks," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2019, 4401–10, http://openaccess.thecvf.com/content_CVPR_2019/html/Karras_A_Style-Based_Generator_Architecture_for_Generative_Adversarial_Networks_CVPR_2019_paper.html.

Legal Risks of Content Generated by Generative AI

The content produced by generative AI is inextricably related to the data it receives as input. This relationship between input data and output content generates a wide range of value judgements that can differ dramatically, posing new legal, regulatory, and ethical difficulties in the development and application of this technology. Among these challenges, data security emerges as a prominent concern, encompassing issues such as illegal data collection, algorithm abuse, difficulty distinguishing genuine content from false information, the proliferation of misleading or deceptive content, inadequate protection of personal privacy, and potential copyright infringement.¹⁷²⁴ Besides it, other important risks are:

Identification and Attribution: AI-generated images often look indistinguishable from real ones, making detection difficult. Determining the source or creator is challenging, especially with anonymized or decentralized uploads.

Consent and Privacy: Deepfakes and manipulated images are frequently created without the consent of the subject, violating privacy and personal dignity. Victims, often women, suffer reputational and emotional harm.

Misinformation and Public Harm: AI-generated imagery can create political disinformation, foment hate speech, or incite violence. Synthetic images used in fake news campaigns have already influenced elections and communal tensions.

Platform Liability: Social media platforms play a dual role—enablers of innovation and gatekeepers of content. However, regulating imagery without stifling creativity or free speech is a legal and ethical tightrope.

¹⁷²⁴ Oren Bar-Gill, "Algorithmic Price Discrimination: When Demand Is a Function of Both Preferences and (Mis)Perceptions," *SSRN Scholarly Paper* (Rochester, NY, May 29, 2018), <https://papers.ssrn.com/abstract=3184533>.

Cross-border Enforcement: Images often transcend jurisdictions, making regulatory enforcement complex. What is legal in one country may be illegal in another.

Existing Legal Frameworks

Global Overview

United States: Given ChatGPT's tremendous influence and the significance of retaining its international leadership position in the field of AI, the United States should set an example by implementing a reasonably open regulatory framework for AI regulation. This strategy aims to develop AI models' legitimacy while also ensuring that they adhere to legal and ethical requirements. On April 11, 2023, the National Telecommunications and Information Administration, part of the US Department of Commerce, requested public feedback on proposed accountability measures for AI systems. The First Amendment generally protects expression, but several jurisdictions, including as California and Texas, have implemented legislation against nonconsensual deepfakes, particularly in elections and pornography.

European Union: The EU has been in the forefront of AI regulation, introducing the AI Act, which seeks to establish a complete legal framework for AI. This risk-based approach classifies AI systems according to their potential dangers and imposes stringent criteria on high-risk applications. The EU's GDPR also plays an important role in regulating AI, notably in terms of data security and privacy. The GDPR's principles of data minimisation, purpose limitation, and explainability have important consequences for the creation and deployment of generative AI models, which frequently rely on large volumes of data and are difficult to comprehend. Platforms are required by the Digital Services Act (DSA) and the AI Act to promote transparency in algorithmic content while also mitigating risk.

China: China's "New Generation Artificial Intelligence Development Plan" outlines its lofty

intentions to lead the world in AI by 2030. With an emphasis on security, transparency, and controllability, the nation has also published a number of ethical principles and governance rules for AI. Launched in 2019, the "Beijing AI Principles" highlight the significance of human-centered AI and the necessity of global cooperation in AI governance. However, there have been worries expressed over the possible abuse of AI technologies, especially when it comes to social control and surveillance. Strict regulations requiring watermarks on synthetic media and holding platforms responsible for fraudulent content have been put in place by the Cyberspace Administration.

Indian Context

AI-generated content is not specifically covered by any laws in India. But current laws only apply in part:

Sections 66E, 67, and 67A of the Information Technology Act of 2000 address electronic communication, obscene content, and privacy violations.

IPC Sections 499–500: Synthetic images that damage reputation are subject to defamation laws. IT Rules 2021: Platforms are required by intermediary norms to remove hazardous content within specified deadlines.

Despite this, there remains a significant legal vacuum in India since AI-generated pictures and deepfakes are not explicitly recognised.

Case Studies

Deepfake of Indian Actress: A Bollywood actress was shown in a compromising scenario in an AI-generated film that went viral on Instagram. Despite being a hoax, the video went viral and caused emotional distress and considerable outrage. Platform and legal shortcomings were exposed by the lack of quick take down tools.

Indian Election Fake Photos for 2024: AI-generated campaign images depicting rival politicians in unfavourable situations were used on a number of political pages. Some were parodies, but others veered into deceptive and

slandorous material. It was difficult for the Election Commission to react immediately.

Technological Responses

Additionally, technology can support regulation by using the following methods:

Watermarking: Tracking the origin of AI-generated photos by embedding digital signatures. Adobe, Microsoft, and other companies support the Content Authenticity Initiative (CAI), which establishes guidelines for recognising fake media.

Tools for Reverse Image Search and Detection: Websites such as Google and Truopic assist in identifying photos that have been altered. Adversarial AI, however, frequently evades detection methods, requiring ongoing development.

Policy Proposals

A multifaceted approach is required to control AI-generated photos on social media:

Reforms to the Law:

Create an AI Imagery Regulation Act in India that addresses responsibility, consent, and deepfakes.

Define and make harmful AI-generated images illegal, such as impersonation, political manipulation, and non-consensual pornography.

Platform Requirements:

Require AI-generated photos to be labelled or watermarked. Demand that platforms put strong verification and removal procedures in place for content that has been detected.

Implement new digital regulations or the IT Act's obligations for algorithmic transparency.

International Cooperation: Cross-border collaboration, enforcement, and standard-setting can be governed by international agreements or frameworks, such as a UN Charter on AI Media Ethics.

Knowledge of Digital:

It is essential to educate the public about identifying synthetic media. Campaigns for media literacy ought to be included in platform user manuals and school curricula.

Development of Ethical AI.

The principles of responsible innovation, such as openness, equity, and user consent, must be followed by AI engineers. Certifications for ethical AI might be implemented.

Balancing Regulation with Free Expression

Overregulation can stifle legitimate innovation, satire, or political criticism, even though dangerous synthetic media must be restricted. Therefore, a human rights-based strategy is crucial. This comprises:

- observing international agreements and Article 19 of the Indian Constitution regarding freedom of expression.
- ensuring that the moderation of content is transparent, reasonable, and appealable.

avoiding outright prohibitions on AI tools or imagery until there is obvious and observable harm.

The Role of Judiciary and Civil Society

The Indian judiciary has the authority to gradually interpret current legislation to address new risks posed by AI imagery. Reforms to laws and policies can be pushed for through Public Interest Litigations (PILs) to support responsible use of AI and argue for safeguards, journalists, technologists, ethicists, and civil society organisations should collaborate.

Conclusion

Human ingenuity is demonstrated by generative AI's capacity to produce material, make judgements, and carry out jobs that have historically needed human intelligence. But this advancement also raises important questions about permission, privacy, intellectual property rights, and the nature of human creativity. The legal system must handle basic issues of authorship, ownership, and misuse potential as

these technologies develop. Our legal systems, moral principles, and social standards must develop together with generative AI technologies.

This is a dedication to maintaining human dignity, equity, and justice in the face of this new technology era, not just a legal requirement. It is difficult, but we must make an effort to strike a balance between protecting human interests and fostering innovation. AI's future reflects human ingenuity as well as our humanity, morals, and ideals.

The emergence of AI-generated photography poses a significant challenge to democratic nations and legal systems. Since social media is the main platform for these images, regulation must be swift, complex, and cooperative. Despite implementing IT regulations, India is still lagging behind in acknowledging the distinct danger posed by artificial visuals.

To guarantee that the promise of AI does not compromise truth, dignity, or trust in public discourse, a specific legislative framework supported by technology, ethics, user awareness, and international cooperation is essential.

References

1. The Information Technology Act, 2000.
2. IT (Intermediary Guidelines and Digital Media Ethics Code) Rules, 2021.
3. Digital Services Act, European Union (2022).
4. California Assembly Bill No. 602 – Deepfake Law.
5. "The Malicious Use of AI: Forecasting, Prevention, and Mitigation" – Future of Humanity Institute, Oxford.
6. Deeptech Report on Deepfakes (2019).
7. Supreme Court of India Judgments on Privacy – Justice K.S. Puttaswamy v. Union of India.
8. Adobe Content Authenticity Initiative – <https://contentauthenticity.org>
9. "Ethics of AI in Media" – UNESCO (2023).

Endnotes

2 Oleksandr Striuk, Yuriy Kondratenko, Ievgen Sidenko, and Alla Vorobyova, "Generative Adversarial Neural Network for Creating Photorealistic Images," in 2020 IEEE 2nd International Conference on Advanced Trends in

Information Theory (ATIT) (Kyiv, Ukraine: IEEE, 2020), 368–71,

<https://ieeexplore.ieee.org/abstract/document/9349326/>; Rajkumar Palaniappan, Kenneth Sundaraj, and Sebastian Sundaraj, "Artificial Intelligence Techniques Used in Respiratory Sound Analysis—A Systematic Review,"

Biomedizinische Technik/Biomedical Engineering 59, no. 1 (January 1, 2014), <https://doi.org/10.1515/bmt-20130074>; Kesavan Venkatesh, Samantha M. Santomartino, Jeremias Sulam, and Paul H. Yi, "Code and Data Sharing Practices in the Radiology Artificial Intelligence Literature: A Meta-Research Study," Radiology: Artificial Intelligence 4, no. 5 (September 1, 2022): e220081, <https://doi.org/10.1148/ryai.220081>; Mohammad Javed Ali and Ali Djalilian, "Readership Awareness Series—Paper 4: Chatbots and ChatGPT— Ethical Considerations in Scientific Publications," Seminars in Ophthalmology 38, no. 5 (July 4, 2023): 403–4, <https://doi.org/10.1080/08820538.2023.2193444>.

3 Carl Vondrick, Hamed Pirsiavash, and Antonio Torralba, "Generating Videos with Scene Dynamics," Advances in Neural Information Processing Systems 29 (2016), <https://proceedings>

[.neurips.cc/paper/2016/hash/04025959b191f8f9de3f924f0940515f-Abstract.html](https://neurips.cc/paper/2016/hash/04025959b191f8f9de3f924f0940515f-Abstract.html).