



Hybrid Deep Learning Framework for Robust and Computationally Efficient Image-Based Deepfake Detection

Priya Mishra¹, Dr. Deepika Pathak²

¹Research Scholar, ²Vice Chancellor

^{1,2}Faculty of Computer Sciences & Application, Dr. A. P. J. Abdul Kalam University, Indore, India

Abstract- The rapid advancement of artificial intelligence and generative models has significantly increased the creation and dissemination of deepfake images, posing severe threats to digital security, media authenticity and public trust. Conventional deepfake detection methods often suffer from limited generalisation capability, high computational complexity and reduced performance against advanced synthetic image generation techniques. This review paper presents a comprehensive analysis of hybrid deep learning frameworks for robust and computationally efficient image-based deepfake detection. The study critically examines existing architectures including Convolutional Neural Networks (CNNs), Vision Transformers (ViTs), attention-based models and explainable artificial intelligence (XAI)-enabled frameworks, and investigates the role of preprocessing and feature enhancement techniques such as normalisation, edge detection, artifact analysis and feature fusion in improving detection accuracy. Eleven studies are reviewed and organised into three themes covering deep learning architectures, preprocessing and feature enhancement, and benchmarking with societal impact. The reviewed evidence is consolidated into comparative tables that map each study to its focus, approach, principal finding and limitation, alongside a comparison of architecture families and a summary of preprocessing techniques. The review identifies six research gaps related to computational efficiency, cross-dataset generalisation, model interpretability, preprocessing strategy, benchmark inconsistency and lightweight hybrid design, and maps each to a corresponding future research direction. Directions focusing on lightweight hybrid architectures, explainable frameworks, benchmark-driven evaluation and multimodal detection strategies are presented to support the development of reliable and scalable deepfake detection systems.

Keywords- Deepfake detection; hybrid deep learning; convolutional neural networks; vision transformers; explainable AI; image forensics; deepfake images; computational efficiency; feature enhancement; artificial intelligence.

I. INTRODUCTION

The emergence of deep learning and generative artificial intelligence has transformed digital media creation by enabling the generation of highly realistic synthetic images, videos and audio content. Deepfake technology has gained considerable attention because of its capability to manipulate facial identities and visual appearance with remarkable realism. Deepfake images are commonly generated



using Generative Adversarial Networks (GANs), diffusion models and encoder-decoder architectures, making it increasingly difficult for humans to distinguish authentic from manipulated content.

Although deepfake technology has beneficial applications in entertainment, education, healthcare and virtual communication, its misuse has created serious concerns relating to misinformation, cybercrime, identity theft, political manipulation and digital fraud. The spread of synthetic media across social networks and digital communication channels has made reliable and efficient detection mechanisms an urgent requirement, and researchers have focused extensively on automated detection using deep learning.

Traditional machine learning approaches relied on handcrafted features and statistical image analysis, but these often failed against sophisticated manipulations produced by modern generative models. Deep learning techniques, particularly CNNs, have demonstrated superior performance in extracting complex visual patterns and identifying manipulation artifacts. Recent studies have introduced hybrid architectures combining CNNs, Vision Transformers, attention mechanisms and explainable AI frameworks to improve robustness and generalisation.

Despite these advances, existing detection models face several limitations including high computational overhead, poor cross-dataset adaptability, vulnerability to adversarial attacks and lack of interpretability. The continuous evolution of generation techniques creates new challenges for existing systems. There is therefore a growing need for hybrid frameworks that achieve high detection accuracy while maintaining computational efficiency and scalability.

1. Objectives and Contributions of this Review

This review provides a comprehensive analysis of hybrid deep learning approaches for image-based deepfake detection. It makes four contributions. First, it consolidates eleven recent studies into a structured thematic review supported by comparative tables rather than narrative description alone. Second, it compares the principal architecture families used in the field on the dimensions that matter for deployment, namely accuracy, generalisation, computational cost and interpretability. Third, it summarises the preprocessing and feature enhancement techniques reported across the corpus and states the purpose each serves. Fourth, it identifies six research gaps and maps each explicitly to a corresponding future research direction.

2. Organisation of the Paper

Section 2 reviews the literature across three themes and consolidates it in comparative tables. Section 3 sets out the identified research gaps. Section 4 concludes, and Section 5 details future research directions.

II. LITERATURE REVIEW

1. Deep Learning Approaches for Image-Based Deepfake Detection

Recent advances in generative models have increased the sophistication of deepfake images, necessitating robust detection mechanisms. Deep learning approaches, particularly CNNs, have emerged as the dominant technique for identifying manipulated visual content. Potey et al. (2026) analysed the effectiveness of CNN architectures in detecting images generated using modern AI tools, and reported that CNN-based frameworks can identify forged facial artifacts and inconsistencies in image texture. Sahar et al. (2026) demonstrated the capability of InceptionResNetV2 in learning discriminative deepfake features, highlighting the importance of transfer learning and deep feature extraction for reliable media authentication.



Joshi et al. (2026) conducted a comparative study of multiple deep learning techniques and emphasised that hybrid architectures combining CNNs with attention mechanisms or transformer-based modules achieve improved classification accuracy relative to standalone models. Their study also identified generalisation across datasets as a critical unresolved challenge, a point that recurs throughout this corpus.

Chen et al. (2026) introduced the X2-DFD framework, integrating explainable artificial intelligence with deepfake detection to improve interpretability and adaptability. Their framework enables visualisation of manipulated regions, increasing model transparency and trustworthiness. This is the only study in the corpus to treat explainability as a first-class architectural concern rather than a post-hoc addition, which makes it the principal reference for the interpretability gap identified in Section 3.

Prabhu et al. (2026) reviewed recent advances in deepfake image detection and discussed the growing adoption of hybrid frameworks integrating CNNs, Vision Transformers and feature fusion mechanisms, noting that hybrid models outperform conventional CNNs on high-quality manipulations produced by diffusion and GAN-based techniques. These findings collectively indicate that the field is moving toward higher robustness, explainability and generalisation capability.

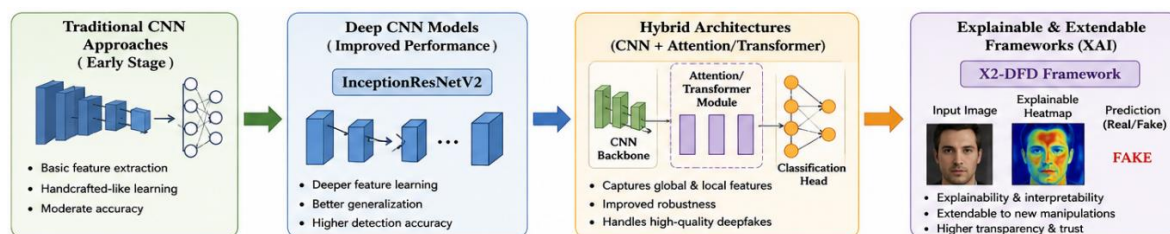


Figure 1: Evolution of deep learning architectures for image-based deepfake detection

Figure 1 illustrates the progression of deep learning approaches in image-based deepfake detection, from traditional CNN architectures to hybrid frameworks integrating Vision Transformers, attention mechanisms and explainable AI models. It highlights the comparative advantages of CNNs, InceptionResNetV2, hybrid CNN-transformer architectures and explainable frameworks such as X2-DFD in improving detection accuracy, feature learning and interpretability.

2. Preprocessing and Feature Enhancement Techniques

Preprocessing plays a crucial role in enhancing feature representation and improving detection performance. Existing studies indicate that image enhancement, edge detection, noise filtering and frequency-domain analysis contribute significantly to identifying subtle manipulation traces. Pawade and Mathur (2026) proposed a pattern analysis method based on edge detection and mutation weight scores, and demonstrated that edge irregularities and abnormal pixel transitions can expose synthetic alterations.

Potey et al. (2026) emphasised preprocessing strategies such as normalisation, image resizing and artifact enhancement for improving CNN performance on AI-generated images, and found that preprocessing directly affects feature extraction quality and model convergence. Chen et al. (2026) incorporated feature explainability within their framework, allowing models to focus on discriminative regions associated with tampering, which improves both accuracy and interpretability.

Gupta et al. (2026) proposed a dual-paradigm framework employing pre- and post-quantum-trained neural networks, combining advanced preprocessing with optimised architectures to achieve greater resilience against adversarial attacks. Xiong et al. (2026) introduced TalkingHeadBench, a benchmark



framework for multimodal deepfake analysis, highlighting the significance of standardised datasets and preprocessing pipelines in ensuring consistent evaluation.

Collectively these studies demonstrate that preprocessing and feature enhancement are essential for improving detection reliability, reducing false positives and strengthening robustness. It is notable, however, that preprocessing is reported as a supporting step in most of these studies rather than evaluated as an independent contributor, so its isolated effect on accuracy remains largely unquantified.

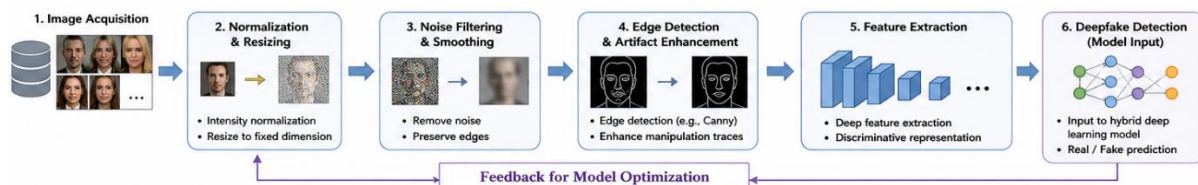


Figure 2: Preprocessing and feature enhancement pipeline for deepfake image detection

Figure 2 presents the preprocessing workflow used in deepfake detection systems, including image acquisition, normalisation, resizing, noise filtering, edge detection, artifact enhancement and feature extraction. It demonstrates how preprocessing improves feature representation quality and supports hybrid deep learning models in identifying synthetic manipulations and forged image patterns.

3. Challenges, Benchmarking and Societal Impact

Despite significant progress, several technical and societal challenges remain unresolved. A major issue is the rapid evolution of generative AI, which continuously produces more realistic and harder-to-detect images. Prabhu et al. (2026) noted that current frameworks often struggle to generalise across unseen datasets and emerging manipulation techniques, and emphasised the need for lightweight, computationally efficient hybrid models capable of real-time deployment.

Benchmarking and standardised evaluation have become important research concerns. Xiong et al. (2026) developed TalkingHeadBench to support comprehensive analysis across multimodal datasets, and highlighted inconsistencies in dataset diversity, evaluation protocols and performance metrics across existing studies. Joshi et al. (2026) similarly observed that many models achieve high accuracy on specific datasets but fail under cross-domain testing. Taken together, these two studies indicate that a substantial share of reported accuracy in the field may not transfer, which is a stronger claim than either makes alone.

Beyond technical concerns, deepfakes pose social, political and ethical threats. Hossain et al. (2026) conducted a bibliometric review of the psychological and societal impacts of deepfakes, revealing their influence on misinformation, public trust and digital manipulation. Livernoche et al. (2026) demonstrated the role of deepfakes in political misinformation during the 2025 Canadian election, providing empirical evidence of real-world deployment at scale and reinforcing the urgency of reliable detection.

Alnaqbi and Ikuesan (2026) reviewed audio deepfake detection techniques and suggested that future research should move toward multimodal frameworks integrating image, audio and video analysis, highlighting the growing importance of explainable AI, adversarial robustness and computational efficiency in next-generation systems.

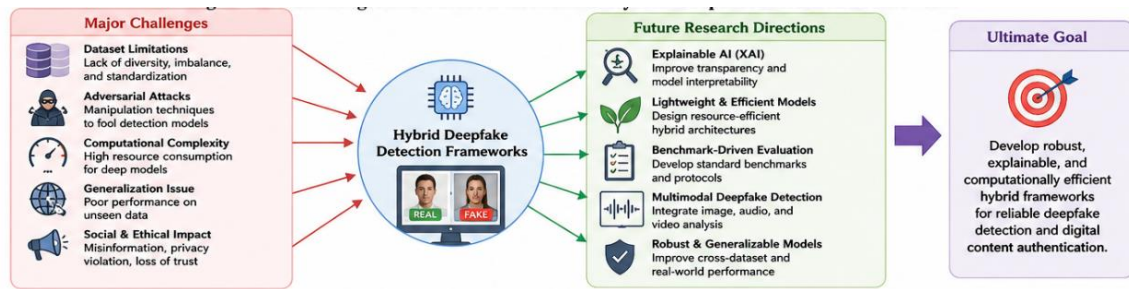


Figure 3: Challenges and future directions in hybrid deepfake detection frameworks

Figure 3 depicts the major challenges and emerging research directions in deepfake detection, including dataset generalisation, adversarial attacks, computational complexity, multimodal deepfakes and misinformation threats. It also highlights future trends including explainable AI, lightweight hybrid architectures, benchmark-driven evaluation and multimodal detection frameworks.

4. Consolidated Summary of the Reviewed Literature

Table 1 organises the reviewed studies by theme, and Table 2 compares each study on research focus, approach, principal finding and the limitation it leaves unresolved. The limitation column reflects the assessment of this review rather than statements made by the original authors.

Table 1: Thematic organisation of the reviewed literature

Theme	Analytical focus	Studies	Count
T1. Deep learning architectures for detection	CNN, transfer learning, hybrid CNN-transformer and XAI-enabled detection models	Potey et al. (2026); Sahar et al. (2026); Joshi et al. (2026); Chen et al. (2026); Prabhu et al. (2026)	5
T2. Preprocessing and feature enhancement	Edge detection, normalisation, artifact enhancement, adversarial-resilient preprocessing and benchmark pipelines	Pawade & Mathur (2026); Gupta et al. (2026); Xiong et al. (2026)	3
T3. Benchmarking, societal impact and multimodal detection	Standardised evaluation, societal and political impact, and audio-visual extension	Hossain et al. (2026); Livernoche et al. (2026); Alnaqbi & Ikuesan (2026)	3

Table 2: Comparative summary of the reviewed studies

Study	Research focus	Approach	Principal reported finding	Limitation for robust efficient detection
Potey et al. (2026)	Detection of images from modern AI tools	CNN architectures with preprocessing	CNNs identify forged facial artifacts and texture inconsistencies	Single-architecture focus; efficiency not assessed
Sahar et al. (2026)	Reliable media authentication	InceptionResNetV2 transfer learning	Transfer learning yields discriminative deepfake features	Heavy backbone; unsuited to constrained deployment
Joshi et al. (2026)	Comparison of detection techniques	Comparative evaluation of deep models	Hybrid CNN-attention models outperform standalone models	Reports cross-dataset failure without proposing a remedy
Chen et al. (2026)	Explainable and extendable detection	X2-DFD framework with XAI integration	Visualising manipulated regions improves transparency and trust	Explainability gains not weighed against computational cost



Study	Research focus	Approach	Principal reported finding	Limitation for robust efficient detection
Prabhu et al. (2026)	Advances in deepfake image detection	Review of hybrid CNN-ViT and fusion frameworks	Hybrid models handle diffusion and GAN manipulations better	Synthesis only; no empirical validation
Pawade & Mathur (2026)	Pattern analysis of forged images	Edge detection with mutation weight score	Edge irregularities and pixel transitions expose alterations	Handcrafted signal; may not survive higher-fidelity synthesis
Gupta et al. (2026)	Adversarially resilient detection	Dual-paradigm pre and post quantum-trained networks	Combined preprocessing and optimised architecture improve resilience	Computational demands of the dual paradigm not addressed
Xiong et al. (2026)	Multimodal benchmarking	TalkingHeadBench benchmark and analysis	Exposes inconsistency in datasets, protocols and metrics	Benchmark scope limited to talking-head content
Hossain et al. (2026)	Societal and psychological impact	Bibliometric review	Deepfakes influence misinformation and public trust	Non-technical; offers no detection guidance
Livernoche et al. (2026)	Political misinformation	Empirical study of the 2025 Canadian election	Documents real-world prevalence and platform dynamics	Descriptive; detection performance not evaluated
Alnaqbi & Ikuesan (2026)	Audio deepfake detection	Systematic review for digital investigation	Recommends multimodal integration of image, audio and video	Audio-focused; image-domain transferability untested

Three observations follow from Table 2. First, the corpus is architecture-rich but efficiency-poor: no reviewed study reports computational cost, inference latency or parameter count alongside accuracy, even though efficiency is a stated concern in several. Second, cross-dataset generalisation is identified as a failure point by Joshi et al. (2026) and Xiong et al. (2026) but is not addressed by any proposed method in the corpus. Third, the societal studies establish urgency and scale but are entirely separate from the technical work, so the field's threat evidence and its detection research do not currently inform each other.

5. Comparison of Architecture Families

Table 3 compares the principal architecture families discussed across the corpus on the four dimensions that determine practical deployability. The comparison is qualitative and reflects the relative positions reported in the reviewed studies rather than measured benchmark values, which the corpus does not report on a common basis.

Table 3: Comparison of deep learning architecture families for deepfake detection

Architecture family	Detection accuracy	Cross-dataset generalisation	Computational cost	Interpretability
Conventional CNN	Effective on known manipulation types	Weak on unseen generators	Moderate	Low, black-box
Deep transfer learning (e.g. InceptionResNetV2)	Strong feature discrimination	Improved over plain CNN	High, large backbone	Low



Architecture family	Detection accuracy	Cross-dataset generalisation	Computational cost	Interpretability
Vision Transformer	Strong on high-quality synthesis	Better global context modelling	High, data-hungry	Low to moderate via attention maps
Hybrid CNN with attention or transformer	Highest reported in the corpus	Improved but still dataset-sensitive	High unless optimised	Moderate via attention weighting
XAI-enabled framework (e.g. X2-DFD)	Comparable to hybrid models	Extendable to new manipulation types	Additional explanation overhead	High, region-level visualisation

The table makes the central design tension explicit: the families that deliver the best accuracy and interpretability are also the most computationally demanding, which is precisely why lightweight hybrid design is the principal open problem identified in Section 3.

6. Summary of Reported Preprocessing Techniques

Table 4 summarises the preprocessing and feature enhancement techniques reported across the corpus, together with the purpose each serves in the detection pipeline.

Table 4: Preprocessing and feature enhancement techniques reported in the reviewed literature

Technique	Purpose in the detection pipeline	Reported in
Normalisation and resizing	Standardising input scale and intensity to improve convergence	Potey et al. (2026)
Noise filtering	Suppressing acquisition noise so manipulation traces remain distinguishable	Potey et al. (2026); Gupta et al. (2026)
Edge detection	Exposing boundary irregularities introduced by facial compositing	Pawade & Mathur (2026)
Artifact enhancement	Amplifying generator-specific residual patterns before feature extraction	Potey et al. (2026)
Mutation weight scoring	Quantifying abnormal pixel transitions as a manipulation indicator	Pawade & Mathur (2026)
Feature fusion	Combining complementary representations from multiple extractors	Prabhu et al. (2026)
Explainability-guided feature focus	Directing the model toward discriminative tampered regions	Chen et al. (2026)
Standardised benchmark pipelines	Ensuring preprocessing consistency for fair cross-study comparison	Xiong et al. (2026)

III. RESEARCH GAP

Although significant progress has been achieved in image-based deepfake detection, six research gaps remain in current detection frameworks. They are stated below and mapped to corresponding future research directions in Table 5.

First, many existing models focus on improving detection accuracy while neglecting computational efficiency. Advanced architectures such as deep CNNs and transformer-based models require



substantial computational resources, making them unsuitable for real-time and resource-constrained applications.

Second, most current frameworks suffer from poor cross-dataset generalisation. Models trained on specific datasets frequently fail to maintain performance on unseen or real-world manipulated images produced by different tools and synthesis techniques.

Third, existing methods lack effective explainability and interpretability mechanisms. Most systems function as black-box models, limiting user trust and reducing applicability in sensitive domains such as digital forensics, legal investigation and cybersecurity.

Fourth, preprocessing and feature enhancement strategies remain inadequate. Many studies overlook the importance of artifact enhancement, edge analysis and frequency-domain representations that can improve feature extraction and detection reliability.

Fifth, benchmark inconsistency remains a major challenge. Existing studies use different datasets, evaluation metrics and experimental settings, making fair comparison and reproducibility difficult.

Finally, limited research has addressed lightweight hybrid architectures capable of balancing detection accuracy, robustness, explainability and computational efficiency simultaneously. This creates the need for optimised hybrid frameworks integrating preprocessing, feature fusion, attention mechanisms and explainable AI techniques for reliable image-based deepfake detection.

Table 5: Research gaps mapped to corresponding future research directions

No.	Identified research gap	Corresponding future research direction
G1	Accuracy prioritised over computational efficiency	Lightweight hybrid models reporting latency and parameter count alongside accuracy (Section 5.1)
G2	Poor cross-dataset generalisation	Cross-dataset training and evaluation across diverse benchmarks (Section 5.3)
G3	Limited explainability and interpretability	XAI-integrated detection frameworks with region-level explanation (Section 5.2)
G4	Inadequate preprocessing and feature enhancement	Systematic evaluation of artifact, edge and frequency-domain enhancement as isolated contributors (Sections 5.1 and 5.3)
G5	Benchmark and evaluation inconsistency	Standardised benchmark protocols, datasets and metrics (Section 5.6)
G6	Absence of balanced lightweight hybrid architectures	Optimised hybrids integrating preprocessing, fusion, attention and XAI, with adversarial robustness (Sections 5.1, 5.4 and 5.5)

IV. CONCLUSION

Deepfake technology has emerged as one of the most challenging threats in the digital era because of its capability to generate highly realistic synthetic images that deceive both humans and automated systems. This review presented a comprehensive analysis of hybrid deep learning frameworks for robust and computationally efficient image-based deepfake detection, examining CNN-based architectures,



Vision Transformers, hybrid models, explainable AI frameworks and preprocessing techniques used for feature enhancement and manipulation detection.

The review found that hybrid frameworks combining multiple feature extraction and learning strategies achieve superior performance relative to standalone models, and that preprocessing methods including normalisation, edge detection, artifact enhancement and feature fusion contribute materially to accuracy and robustness. Comparing architecture families on a common set of dimensions also made the central design tension explicit: the families delivering the best accuracy and interpretability are the most computationally demanding.

Three further observations emerged from the comparative synthesis. No reviewed study reports computational cost alongside accuracy, so claims about efficiency in this field are currently unverifiable. Cross-dataset failure is diagnosed by two studies but addressed by none. And the societal evidence establishing the scale of the threat is entirely separate from the technical detection literature.

Despite notable progress, several challenges remain unresolved, including computational complexity, cross-domain generalisation, lack of interpretability and vulnerability to adversarial manipulation. Future research should therefore focus on designing lightweight, scalable, explainable and benchmark-driven hybrid systems capable of addressing real-world deepfake detection challenges.

Future Work

Lightweight Hybrid Deep Learning Models

Future research should develop lightweight hybrid architectures that reduce computational complexity while maintaining high detection accuracy for real-time applications. Given that no study in this corpus reports inference cost, future work should report latency, parameter count and memory footprint alongside accuracy as standard practice.

Explainable and Interpretable Detection Frameworks

Integrating explainable artificial intelligence techniques into detection systems can improve model transparency, user trust and forensic reliability, particularly where detection outputs may be used as evidence.

Cross-Dataset Generalisation

Future models should be trained and evaluated across diverse benchmark datasets to improve robustness against unseen and evolving generation techniques, with cross-domain performance reported as a primary result rather than an ancillary observation.

Multimodal Deepfake Detection

Combining image, audio and video analysis within unified hybrid frameworks can improve detection capability against sophisticated multimodal deepfakes.

Adversarially Robust Detection Systems

Future studies should explore adversarial defence mechanisms to enhance the resilience of detection systems against manipulation-aware attacks.

Benchmark-Driven Evaluation Frameworks

Standardised benchmark protocols, datasets and evaluation metrics should be developed to ensure fair comparison, reproducibility and scalability of future deepfake detection research.



REFERENCES

1. Potey, M. A., Khan, Y. S., Gupta, A., Gupta, A., Ranjan, A., & Wani, A. (2026, March). Analysis of CNN for deepfake detection on images generated by latest AI tools. In 2026 3rd International Conference on Emerging Trends in Engineering and Medical Sciences (ICETEMS) (pp. 1-6). IEEE.
2. Sahar, N. M., Rozi, M. F. S. M., Suriani, N. S., Sari, S., Ismail, S., Jamal, A. A., & Alessa, A. M. (2026). Advances in deepfake detection: Leveraging InceptionResNetV2 for reliable video authentication. Vol. 1, 90-105.
3. Joshi, S., Besoya, R., Kumar, P., & Bidhuri, K. (2026, April). Dissecting deepfakes: A comparative study of deep learning detection techniques. In AIP Conference Proceedings (Vol. 3387, No. 1, p. 030001). AIP Publishing LLC.
4. Hossain, S., Sudarsan, D., Zaffar, H., & Ahmad, N. (2026). Social and psychological impact of deepfakes: A comprehensive bibliometric review. Global Knowledge, Memory and Communication, 1-20.
5. Chen, Y., Yan, Z., Cheng, G., Zhao, K., Lyu, S., & Wu, B. (2026). X2-DFD: A framework for explainable and extendable deepfake detection. Advances in Neural Information Processing Systems, 38, 83792-83839.
6. Prabhu, O., Naik, S. S., Prabhu, S., & Chadaga, K. (2026). Countering synthetic realities: Advances, challenges, and future directions in deepfake image detection. Systems and Soft Computing, 200476.
7. Pawade, V. S., & Mathur, S. (2026). Novel pattern analysis of deepfake images using image edge detection and mutation weight score. Journal of Forensic Science and Medicine, 12(1), 54-60.
8. Xiong, X., Patel, P., Fan, Q., Wadhwa, A., Selvam, S., Guo, X., & Sengupta, R. (2026). TalkingHeadBench: A multi-modal benchmark and analysis of talking-head deepfake detection. In Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (pp. 4139-4149).
9. Gupta, S., Hariprasad, Y., Iyengar, S. S., Gurappa, S., & Mohanty, P. (2026). Enhancing digital security: A novel dual-paradigm approach for robust deepfake detection using pre and post quantum-trained neural networks. Digital Threats: Research and Practice, 7(1), 1-15.
10. Livernoche, V., Musulan, A., Yang, Z., Godbout, J. F., & Rabbany, R. (2026, April). Deepfakes in the 2025 Canadian election: Prevalence, partisanship, and platform dynamics. In Proceedings of the ACM Web Conference 2026 (pp. 8717-8720).
11. Alnaqbi, M., & Ikuesan, R. A. (2026). A systematic review of audio deepfake detection techniques for digital investigation. Discover Computing, 29(1), 202.