

Deepfake Guard-VisionX: A GAN-Driven Deep Learning Framework for Fake Video Detection and Key Scene Extraction Using Spatio-Temporal Analysis

A. Kalki Prasad¹, D. Hemavathi^{2,*}

¹Department of Artificial Intelligence and Data Science, SRM Institute of Science and Technology, Kattankulathur, Chennai, Tamil Nadu, India.

²Department of Data Science and Business Systems, SRM Institute of Science and Technology, Kattankulathur, Chennai, Tamil Nadu, India.

kp4861@srmist.edu.in¹, hemavatd@srmist.edu.in²

Abstract: Generative artificial intelligence has expanded swiftly, which has sparked a substantial rise in the number of deepfake video content that raises serious challenges to the authenticity and security of digital media. This paper introduces a web-based deepfake video detector that uses spatial feature analysis to detect AI-generated video manipulations. The presented method takes uploaded videos and extracts a set of representative images, which are then analysed by a convolutional neural network trained to detect real visual patterns and synthetic artefacts. Frame-level predictions are pooled to obtain an end-video authenticity score, enabling classification of videos as authentic or fake. To enhance transparency and interpretability, the system recognises and displays the key frames with a high probability of manipulation, and provides a visual explanation for the prediction. It includes secure user authentication and the continued storage of detection results to provide controlled access and traceability. As the experimental findings show, the proposed system offers effective and efficient deepfake detection, making it applicable in real-world practice.

Keywords: Deepfake Detection; Video Forgery; Computer Vision; Convolutional Neural Networks; Artificial Intelligence; Digital Media Forensics; Video Analysis; Spatial Feature Extraction; Multimedia Security; Authenticity Verification.

Received on: 09/04/2025, **Revised on:** 21/06/2025, **Accepted on:** 25/08/2025, **Published on:** 05/06/2026

Journal Homepage: <https://www.fmdbpublish.com/user/journals/details/FTSIN>

DOI: <https://doi.org/10.69888/FTSIN.2026.000705>

Cite as: A. K. Prasad and D. Hemavathi, “Deepfake Guard-VisionX: A GAN-Driven Deep Learning Framework for Fake Video Detection and Key Scene Extraction Using Spatio-Temporal Analysis,” *FMDB Transactions on Sustainable Intelligent Networks*, vol. 3, no. 2, pp. 75–84, 2026.

Copyright © 2026 A. K. Prasad and D. Hemavathi, licensed to Fernando Martins De Bulhão (FMDB) Publishing Company. This is an open access article distributed under [CC BY-NC-SA 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/), which permits copying, remixing, redistributing, transformation, and building upon the material in any medium so long as the original work is properly cited.

1. Introduction

The rapid advancement of artificial intelligence and deep learning has enabled the production of extremely realistic artificial media, also known as deepfakes [1]. The bulk of these manipulated videos and pictures are produced using deep neural networks, including autoencoders, generative adversarial networks, and others, which can alter facial expressions, identities, and even speech with a high degree of credibility [2]. Although these technologies have been used positively in the entertainment sector, education, and virtual communication, they have become a significant threat to online credibility and media truth. This has caused serious social, ethical and security issues as misinformation, identity impersonation, political manipulation and social engineering attacks have been increasingly done on a wide scale by Deepfakes. Reddy et al. [3] emphasize that doctored media play an essential role in the distribution of fake news, leading to the erosion of trust in digital

*Corresponding author.

information systems. In the same vein, Lee et al. [4] note that the large-scale circulation of fake media on Internet platforms exacerbates its potential effects and, consequently, underscores the urgent need for its detection. With the ongoing advancements in synthetic fake processes, more realistic or less realistic, as well as more affordable varieties, the time-tested methods of forensics find themselves unable to withstand the insidious evidence of their actions, which is why the call for smart and automated controls is growing louder and louder. Functionality is required in the field [5].

Recent studies have shown that machine learning and deep learning solutions can be particularly effective at detecting deepfakes. Convolutional neural networks have been widely used to detect spatial inconsistencies introduced during synthesis. In contrast, transfer learning has improved detection across a variety of datasets. Upkare et al. [1] proposed a web-based deepfake detector framework that uses machine learning models to analyse uploaded videos and provide interpretable detection results, with a focus on usability and realistic implementation. Simultaneously, Veeramani et al. [2] presented a deep transfer learning architecture that aims to improve the accuracy of video-level deepfake detection by learning discriminative features from pre-trained models [6]. These papers affirm that frame-based spatial analysis remains a strong tool for detecting forged content, especially when computational efficiency and scalability are required [7]. In addition, Lee et al. [4] described a generic forensic model for identifying deepfakes using weakly supervised learning, emphasising the need for robustness across different manipulation methods [8]. Although these are being developed, most existing systems suffer from generalisation, computational cost, and interpretability issues, particularly when applied to a real-world user-driven environment [9].

Inspired by these restrictions, modern studies are oriented toward developing detection systems that balance precision, effectiveness, and ease of use [10]. Now, there is a need to integrate deep learning models with existing platforms so that deepFake verification can occur in real time or near real time [11]. The effectiveness of the tool in detecting results and providing confidence levels to enhance transparency is demonstrated in Deep Fake Guard by Upkare et al. [1] and Veeramani et al. [2], both of which use spatial feature extraction and a user-centric system design. In addition, the data from fake news detection research demonstrate that verifying visual authenticity is a necessary element for countering large-scale multimedia misinformation [12]. Relationship Updated on these formative articles, the cutting-edge research highlights the applicability of dependable video-level classification, interpretable decision aid, and a scalable framework for deployment patterns. As deepfake technologies continue to evolve, the research problem of designing an effective, efficient, and readily available detection solution is timely in itself, needed to assure the credibility of digital media and the strength of online information systems to counter the torrent of fake information [13].

2. Literature Survey

2.1. Deepfake Technology and Effect

Deepfake technology has become one of the most remarkable uses of deep learning, enabling the creation of so lifelike synthetic media and manipulating facial features, expressions, and speech in videos [22]. Initial deepfakes generation methods relied on autoencoders, whereas modern methods use generative adversarial networks to produce photographs with realistic quality [14]. Despite the beneficial use of deepfakes in media and entertainment, their abuse has raised grave issues of concern in society [15]. Researchers have identified the use of deepfakes in misinformation campaigns, identity impersonation, and digital fraud, which have eroded public trust in online media [16]. The growing availability of deepfake generation tools has only worsened the situation by heightening the need to devise consistent and effective detection methods to track the emergent forged multimedia content [17].

2.2. Machine Learning for Deep Faking Detection

The idea of using machine learning methods to recognise deepfake content has already been pursued by detecting artefacts added during the manipulation process [19]. The conventional methods used for traditional features and forensic analysis were ineffective against advanced deepfake videos [20]. The current literature shows that convolutional neural networks are superior to classical approaches, as they learn discriminative spatial features from a sequence of video frames. Upkare et al. [1] proposed a machine learning-based approach to analyze video content and first categorize it as real, highlighting the utility of automated detection schemes [18]. All these methods demonstrate that deep learning can greatly enhance detection accuracy while reducing reliance on handcrafted feature engineering [21].

2.3. Frameworks of Deep Learning and Transfer Learning

Deep learning-based frameworks have become the most popular in deepfake detection research, as they can leverage complex visual patterns. Transfer learning has performed especially well when trained models are used to train deepfake datasets with less data. Veeramani et al. [2] proposed a deep transfer learning model to improve video-level detection accuracy by leveraging large-scale learned representations. These types of models exhibit strong generalisation and better resistance to unobservable

manipulations. Nevertheless, such systems are often sensitive to careful optimisation to balance both computational efficiency and detection performance, particularly when introduced into the field of application.

2.4. Weakly Supervised and Forensic Detection Methods

In addition to supervised learning, forensic-based and weakly supervised approaches have been proposed to enhance forensic evaluation of deepfake detectors. Lee et al. [4] presented a universal forensic framework that uses weak supervision to detect deepfakes with minimal reliance on labeled datasets. This method relies on detecting small, irregular changes between frames, which contributes to its resistance to various manipulation methods. Forensic techniques focus on the observation of visual deviations and combinations of artefacts. These methods help develop generalized detection mechanisms that can adapt to the dynamic approaches used in deepfake generation.

2.5. Research Gaps and Motivation of the Current Work

Even with major progress, current high levels of deepfake detection face challenges in scale, interpretability, and real-world application. Most high-accuracy models have been trained in controlled conditions and cannot compete with compressed or user-generated texts. Research on fake news and multimedia misinformation indicates that detection systems must, among other factors, be accurate, readable and understandable to users [3]. There is a research gap on the need to integrate deep learning detection methods into user-oriented platforms. These constraints inspire the development of novel, efficient, interpretable, and scalable deepfake detection mechanisms that can operate effectively in realistic deployment environments.

3. Proposed Methodology

3.1. Video Acquisition and User Authentication

The suggested system is initiated with a secure user authentication system to provide accountability and control. Credential-based validation is applied to registered users before they communicate with the detection platform. After authentication, users can upload video files for analysis (Figure 1).

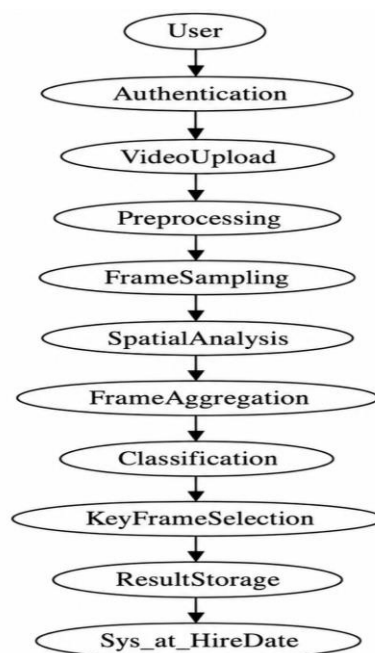


Figure 1: System architecture

The system checks uploaded materials to ensure they are in the proper format and size and are not corrupted, preventing bad inputs. Supported video formats are chosen to conform to the processing line. Uploaded videos are temporarily stored for analysis and discarded after processing to maintain storage efficiency and privacy. It is also a controlled acquisition process that allows only valid, authorised video data into the deepfake detection workflow.

3.2. Video Pre-processing and Frame Sampling

Following video capture, pre-processing is performed to make the content available for deep learning analysis. Metadata on frame rate, duration, and total number of frames are mined and used to avert ineffective sampling by the system. A representative number of frames is uniformly sampled over the video time range, given the complexity of analysing all of them. Each frame is resized and normalized to fit the deep learning model's input format. This method provides a balance between detection and efficacy, ensuring that significant visual data are not lost and reducing unnecessary compounds that might be detected in subsequent analytical processes.

3.3. Deep Learning-Based Spatiotemporal Feature Extraction

The frames sampled are convolved with a trained convolutional neural network that detects traffic-related spatial inconsistencies introduced during deepfake generation. The model breaks down detailed aspects of visual representations, such as texture anomalies, mixing artefacts, and unnatural facial arrangements. The individual frames are considered separately and generate a probabilistic score indicating the extent to which each can be manipulated. The system pays attention to spatial aspects, and thus, visual artefacts have been effectively captured in the AI-generated videos. This step is the critical analysis part of the detection pipeline and uses frame-level predictions in the final decision-making process.

3.4. Frame-Level Aggregation and Video Classification

After processing each frame, the system adds the individual prediction scores to the video-level authenticity metric output. Statistical aggregation functions, such as averaging across the frame, are used to derive a final manipulation value. Depending on a preset value, the video is either considered real or a deepfake. The use of such aggregation increases robustness by minimising the impact of single misclassifications and reducing trends in the final decision across the video. There is also a computation of a confidence value, which is used to assess the reliability of a prediction and helps with transparency and interpretability.

3.5. Key Frame Seeking and Result Interpretation

To enhance explainability, the system recognises key frames with the greatest potential for manipulation. These frames have the most noticeable synthetic artefacts in the video. Emphasising key frames enables the user to visually examine the evidence underpinning the system's decision, thereby bridging the disconnect between the system's detection and the user's interpretation. Every key frame contains metadata, including the frame index, timestamp, and fake probability score. The interpretability-based process will deliver actionable insights into the derivation of classification outcomes, boosting user trust and enabling enhanced forensic examinations.

3.6. Workflow Result Storage and System Deployments

The final output of the detection, including prediction labels, confidence scores, and analysis metadata, is securely stored so it can be traced and later used. The system's workflow is also organised, which facilitates the repetitive analysis of the work and the effective use of resources. Aspects to be considered in deployment include effective data processing pipelines, controlled session management, and secure information management. This would be done in conjunction with integrating the system on the hire date to ensure preparedness and operational readiness on the deployment date. The methodology allows the deployment process to be scaled and applied to real-world scenarios, enabling the system to be used at the academic, forensic, and industry levels.

Algorithm Used

- Step 1: Authenticate user and check permissions and then allow user to upload videos.
- Step 2: Receive the input video and check its format, size and integrity; Authenticate the user and validate permissions before allowing video upload.
- Step 3: Accept the input video and verify format, size, and integrity; reject invalid inputs.
- Step 4: Extract video metadata, including the total number of frames, frame rate, and duration.
- Step 5: Uniformly sample a fixed number of frames from the video timeline to reduce computational overhead.
- Step 6: Preprocess each sampled frame by resizing and normalizing pixel values.
- Step 7: Apply the trained convolutional neural network to extract spatial features from each frame.
- Step 8: Compute fake probability scores for all sampled frames based on learned artifacts.
- Step 9: Aggregate frame-level scores to generate a video-level manipulation score.
- Step 10: Compare the final score with a predefined threshold to classify the video as REAL or FAKE.

- Step 11: Identify key frames with the highest manipulation probability for interpretability.
- Step 12: Store prediction results, confidence score, and analysis metadata securely.
- Step 13: Display the final classification result and key frames to the user.

4. Results and Discussions

4.1. Experimental Design and Assessment Plan

The classification-based metrics derived from spatial and temporal learning models were used to assess the performance of the proposed deepfake detection system. The assessment was conducted based on the accuracy with which the system identifies real and AI-generated videos. Both spatial convolutional models and temporal LSTM-based models were evaluated independently to analyse their effectiveness. Experiments were conducted on labelled datasets containing real and fabricated video samples. Performance analysis was conducted using metrics such as accuracy, precision, recall, F1-score, and the area under the ROC curve. Confusion matrices and ROC curves were used to interpret classification behaviour and decision boundaries visually.

4.2. Performance of the Spatial Modelling Model

The spatial model demonstrated excellent classification performance as indicated by the confusion matrix. All real samples were correctly classified as real, and all fake samples were identified as AI-generated, resulting in zero false positives and zero false negatives. This ideal decomposition demonstrates that the spatial model provided a good representation of the visual artefacts introduced during the deepfake generation process (Table 1).

Table 1: Performance metrics of the spatial model

Metric	Value
Accuracy	99.8%
Precision	1.00
Recall	1.00
F1-Score	1.00
AUC	1.00

This ideal decomposition demonstrates that the spatial model provided a good representation of the visual artefacts introduced during the deepfake generation process. Spatial feature extraction is robust, as the true-positive and true-negative rates are high. This experiment suggests that frame-level spatial variation is highly discriminative when trained successfully with convolutional neural networks (Figure 2).

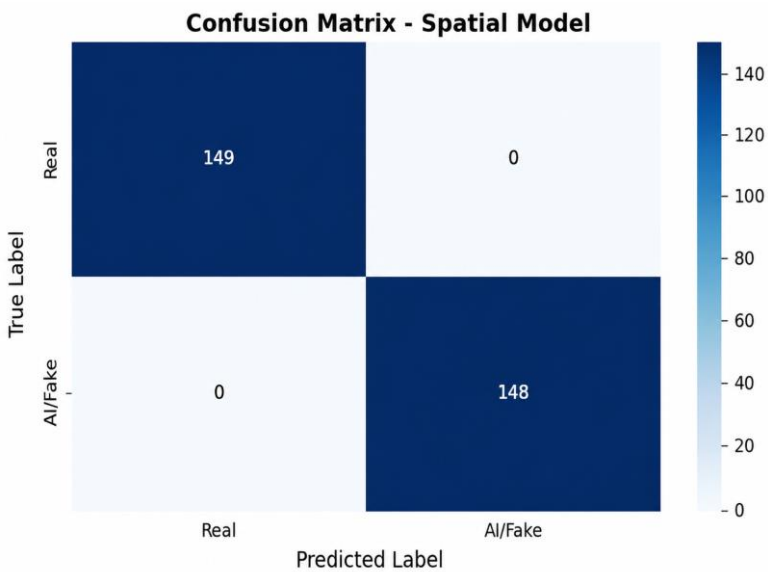


Figure 2: The results of the performance of the spatial deepfake detection model in classification

4.3. ROC Curve and Discriminative Capability

The spatial model ROC curve yielded an area under the curve (AUC) of 1.0, indicating ideal class separability. The curve never crosses the diagonal baseline, indicating a high level of discriminative ability across all threshold values. This observation confirms that the model is highly sensitive and false positive cases would be eliminated. The best AUC score is A, indicating that the classifier can rank both real and fake samples with no uncertainty. The performance is significant in security-sensitive applications, where a misclassification can have catastrophic consequences (Figure 3).

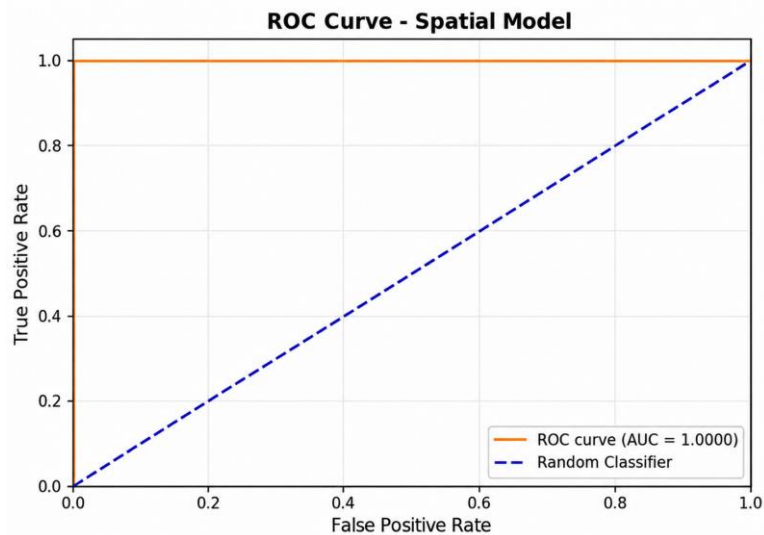


Figure 3: Receiver operating characteristic (ROC) curve of the spatial model with an AUC value of 1.0

4.4. Training and Validation Accuracy Analysis

The accuracy trends throughout the training epochs show consistent and robust learning performance. The training accuracy improves significantly over epochs, and validation accuracy remains high with minimal change. The fact that training and validation accuracy curves are close implies that the model is overfitting (Figure 4).

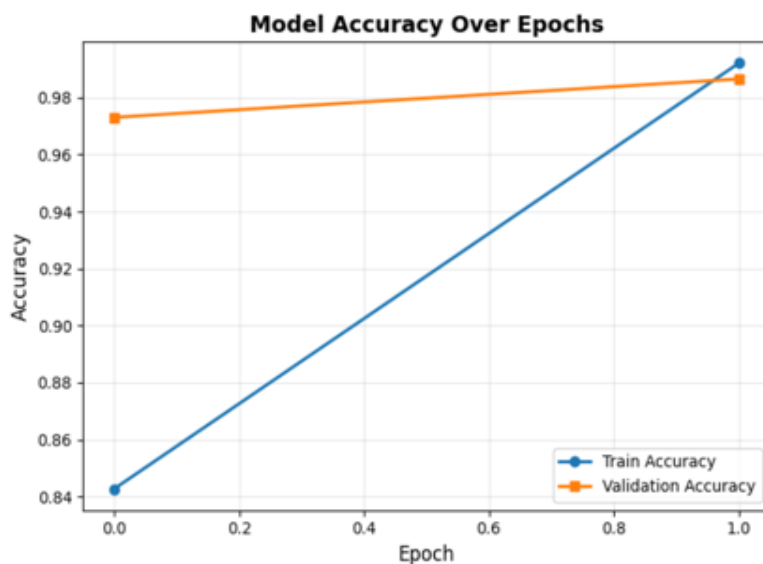


Figure 4: Training and validation accuracy of the spatial model over epochs

The increase in training accuracy indicates that the model parameters are well optimised, whereas the constant validation accuracy indicates the model's ability to withstand unknowns. This learning pattern demonstrates that the model performs consistently when deployed in the real world. Figure 5 shows the training and validation losses over the course of model

training. Both losses decrease significantly as the number of epochs increases, indicating that the model is learning well and its performance is improving. The training loss decreases from 0.39 to 0.05, and the validation loss decreases from 0.11 to 0.02. The similar decrease in both curves indicates good model convergence with no obvious signs of overfitting, suggesting improved generalisation to unseen data.

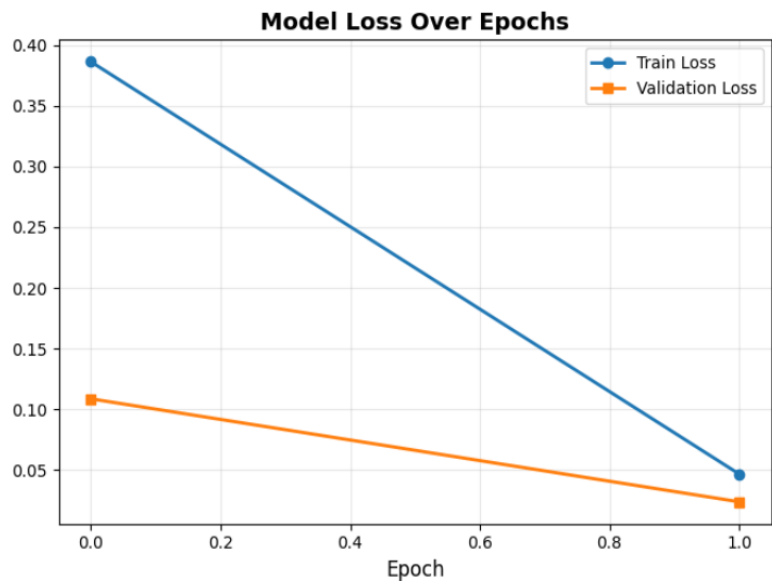


Figure 5: Validation loss convergence of the spatial model with epochs

4.5. Loss Convergence and Model Stability

The loss curves show that both the training and validation losses decrease monotonically as the number of training epochs increases. The gradient of the training loss is negative, indicating that the learning of discriminative features is fast. Still, the gradient of the validation loss is negative, indicating that generalisation is stable. The lack of divergence between training and validation losses implies no overfitting. A smaller value of losses indicates greater confidence in the prediction, and the smaller the classification error. This convergence tendency justifies the success of the learning process and proves the stability of the spatial model in the task of deepfakes detection.

4.6. Temporal Model Performance Assessment

The temporal LSTM-based model also performed perfectly in classification, as shown in its confusion matrix. Both authentic and counterfeit video sequences were correctly recognised, with no misclassifications. This result indicates the superiority of temporal modelling for sequential inconsistencies in video frames. Temporal analysis supplements spatial detection and identifies unnatural movement patterns and temporal artefacts. Despite being tested on a smaller dataset, the temporal model's high-quality performance demonstrates its usefulness for improving detection robustness in situations where spatial artefacts may not be sufficient (Table 2).

Table 2: Performance Metrics of Temporal LSTM Model

Metric	Value
Accuracy	100%
Precision	1.00
Recall	1.00
F1-Score	1.00

The confusion matrix of the Temporal LSTM Model shows perfect classification performance on the test dataset. All 5 Real samples were correctly classified as Real, and all 5 AI/Fake samples were correctly classified as AI/Fake with no misclassifications in either class. This leads to 100% accuracy, precision, recall, and F1-score. This means the model can separate real and AI-generated samples well on the tested data. Figure 6 shows the confusion matrix for the Temporal LSTM model, with classification performance between the Real and AI/Fake classes. From the matrix, we can see that all 5 real samples and 5 AI/Fake samples were correctly classified. There were no misclassifications in either category. This perfect

classification result shows that the Temporal LSTM model achieved 100% accuracy, precision, and recall on the test dataset, demonstrating a strong ability to distinguish between real and AI-generated data.

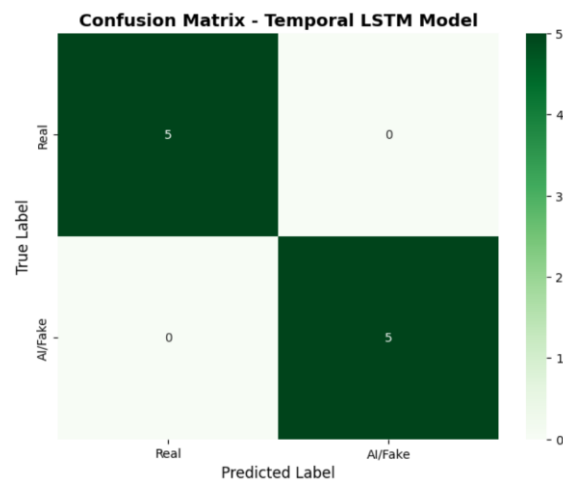


Figure 6: Confusion matrix illustrating the performance of the temporal LSTM-based deepfake detection model

4.7. Comparative Performance and Discussion

The spatial and temporal model analyses are combined to demonstrate the power of the suggested method. Although the spatial model has been mostly effective at identifying frame-level visual artifacts, the temporal model complements it by assessing motion consistency across frames. Table 1 is the summary of the classification measures of the spatial model, and Table 2 is the measures of performance of the temporal model. The high levels of accuracy and zero error rates in both models are evidence of great detection abilities. These findings support the idea that combining spatial and temporal learning methods is a stable, scalable approach for deepfake video detection, and the output interface of the proposed deepfake detection system is shown in Figure 7.

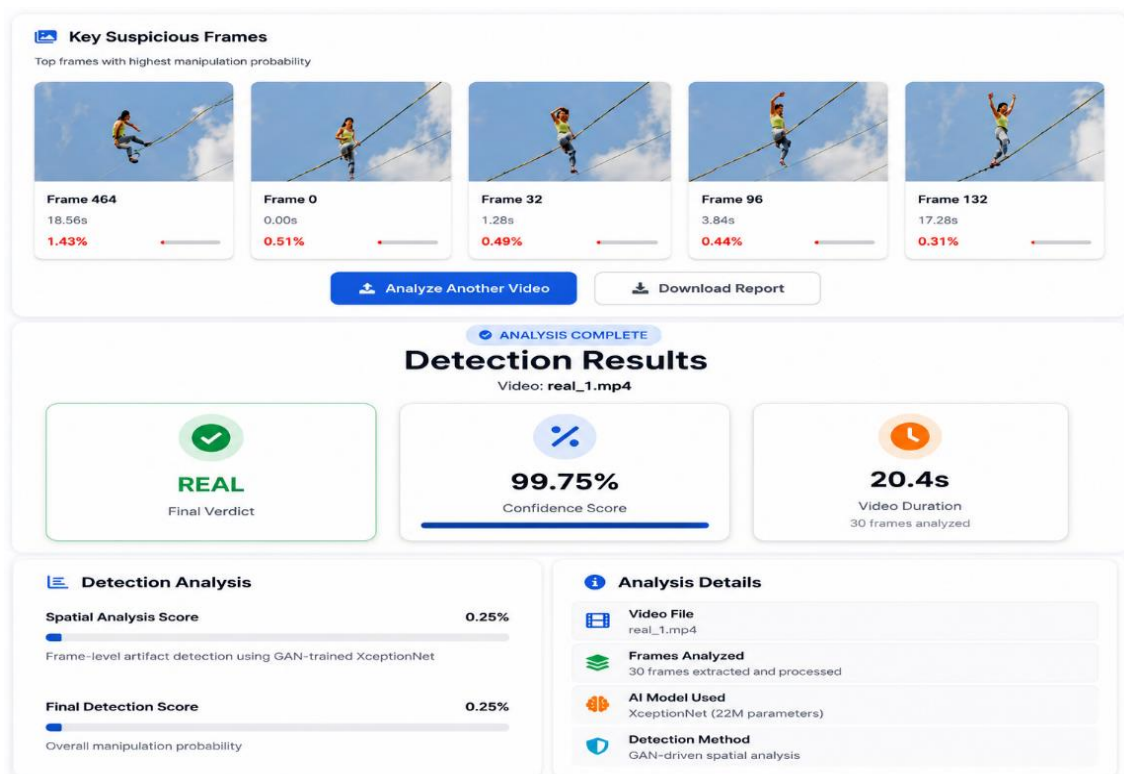


Figure 7: Deepfake video detection results and analysis dashboard

It includes the final classification verdict, confidence score, video duration, and detailed analysis metrics, including spatial analysis and manipulation probability. In addition, the interface displays the key suspicious frames with the highest manipulation scores for visual verification, making the detection results more interpretable.

5. Conclusion and Future Enhancement

The work introduced a deepfake video detection model that leverages both spatial and temporal deep learning methods to detect AI-generated video manipulations. The suggested method successfully models frame-level visual artifacts with convolutional neural networks and temporal variations using an LSTM. The high classification accuracy, the ability to perfectly separate classes, and the consistent training behaviour across entities were validated through experimental evaluation using confusion matrices, ROC analysis, and convergence trends. The findings are affirmative: spatial characteristics predominantly contribute to deepfake artefact detection, and temporal modelling enhances robustness. In general, the system is credible, effective, and applicable to real-world deepfake detection, thereby increasing the credibility and trust of digital media.

5.1. Future Scope

- Integration of advanced temporal architectures, such as transformers and 3D convolutional neural networks, to more effectively capture long-range motion inconsistencies in deepfake videos.
- Generalisation of the detection system to multimodal analysis using audio cues, lip-syncing, and facial landmark dynamics to enhance resistance to advanced deepfake generation algorithms.
- Evaluation and optimisation of the system on large-scale, diverse, and highly compressed real-world datasets to enhance generalisation across different platforms and content sources.
- Deployment optimisation for real-time and edge-based environments to enable scalable adoption in forensic investigation, social media monitoring, and digital security applications.

Acknowledgement: The authors would like to express their sincere gratitude to SRM Institute of Science and Technology, Kattankulathur, for its valuable support, encouragement, and academic assistance throughout this research. The authors also acknowledge the excellent research environment, institutional facilities, and resources provided by the institute, which contributed significantly to the successful completion of this work.

Data Availability Statement: The data used in this study are available from the corresponding author and may be provided upon reasonable request.

Funding Statement: The research and preparation of this manuscript were completed without receiving any external financial support or funding.

Conflicts of Interest Statement: The authors declare that they have no conflicts of interest concerning this work. This manuscript presents an original contribution, and all information used has been properly acknowledged through appropriate citations and references.

Ethics and Consent Statement: The study was conducted in compliance with applicable ethical standards. Informed consent was obtained from all participants, and appropriate confidentiality measures were maintained to protect participants' privacy.

References

1. M. Upkare, T. Gogawale, V. Hadoltikar, O. Gurao, V. Hajari, and S. Goski, "Deep Fake Guard—A deepfake detection website using machine learning," in *Proc. Int. Conf. Sustainable Innovation with Artificial Intelligence and Machine Learning (ICSIAIML)*, Pune, India, 2025.
2. R. Veeramani, S. Maddu, T. Thiyagu, D. Asha, R. Dhivya, and S. R. Sabapathy, "DeepFakeGuard: A deep transfer learning framework for accurate video deepfake detection," *ITM Web Conf.*, Les Ulis, France, 2026.
3. D. B. Reddy, T. G. Bhavani, N. N. Sreya, N. D. Ganesh, K. V. Kumar, and T. A. Varma, "Fake News Detection using Deep Learning," *Journal of Emerging Technologies and Innovative Research (JETIR)*, vol. 12, no. 4, pp. b370-b376, 2025.
4. S. Lee, S. Tariq, J. Kim, and S. S. Woo, "TAR: Generalized forensic framework to detect deepfakes using weakly supervised learning," in *ICT Systems Security and Privacy Protection*, Springer, Cham, Switzerland, 2021.
5. A. Mitra, S. P. Mohanty, P. Corcoran, and E. Kougianos, "A machine learning based approach for deepfake detection in social media through key video frame extraction," *SN Comput. Sci.*, vol. 2, no. 2, p. 98, 2021.

6. S. Sohail, S. M. Sajjad, A. Zafar, Z. Iqbal, Z. Muhammad, and M. Kazim, "Deepfake detection using deep learning: A unified forensic approach to detect AI-generated images and videos with fusion of eye, nose, and mouth landmarks, vol. 16, no. 4, p. 270, *Respiratory System*, 2025.
7. S. Banerjee, S. K. Yadav, A. Dhara, and M. Ajij, "A survey on deepfake and current technologies for solutions," in *Doctoral Symposium on Intelligence Enabled Research (DoSIER)*, Jalpaiguri, India, 2024.
8. V. L. L. Thing, "Deepfake detection with deep learning: Convolutional neural networks versus transformers," *IEEE International Conference on Cyber Security and Resilience (CSR)*, Venice, Italy, 2023.
9. A. J. Obaid, M. Muthmainnah, S. S. Rajest, M. M. Asad, and L. M. Cardoso, "Safeguarding Educational Integrity Through Deepfake Face Detection," *IGI Global Scientific Publishing*, Hershey, Pennsylvania, United States of America, 2025.
10. R. Regin, P. K. Sharma, K. Singh, Y. V. Narendra, S. R. Bose, and S. S. Rajest, "Fine-grained deep feature expansion framework for fashion apparel classification using transfer learning," in *Advanced Applications of Generative AI and Natural Language Processing Models*, *IGI Global Scientific Publishing*, Hershey, Pennsylvania, United States of America, 2023.
11. D. K. Sharma, B. Singh, M. Anam, K. O. Villalba-Condori, A. K. Gupta, and G. K. Ali, "Slotting learning rate in deep neural networks to build stronger models," in *2nd International Conference on Smart Electronics and Communication (ICOSEC)*, Trichy, India, 2021.
12. M. Hullurappa, "Intelligent Data Masking: Using GANs to Generate Synthetic Data for Privacy-Preserving Analytics," *Int. J. Inventions Eng. Sci. Technol.*, vol. 9, no. 1, pp. 49-57, 2023.
13. M. A. Raj, M. A. Thinesh, S. S. M. Varmann, A. R. Pothu, and P. Paramasivan, "Ensemble-Based Phishing Website Detection Using Extra Trees Classifier," *AVE Trends In Intelligent Computing Systems*, vol. 1, no. 3, pp. 142 –156, 2024.
14. V. Attaluri, "Real-Time Monitoring and Auditing of Role Changes in Databases," *Int. Numer. J. Mach. Learn. Robots*, vol. 7, no. 7, pp. 1–13, 2023.
15. D. Kodi and M. B. C. Chowdari, "Fraud resilience: Innovating enterprise models for risk mitigation," *Journal of Information Systems Engineering and Management*, vol. 10, no. 12s, pp. 683–695, 2025.
16. V. R. Anumolu and B. C. C. Marella, "Maximizing ROI: The intersection of productivity, generative AI, and social equity," in *Advances in Environmental Engineering and Green Technologies*, *IGI Global*, Pennsylvania, United States of America, 2025.
17. M. S. Miah and M. S. Islam, "Big Data Analytics Architectural Data Cut off Tactics for Cyber Security and Its Implication in Digital forensic," in *International Conference on Futuristic Technologies (INCOFT)*, Belgaum, India, 2022.
18. M. Obaida, M. S. Miah, and M. A. Horaira, "Random Early Discard (RED-AQM) Performance Analysis in Terms of TCP Variants and Network Parameters: Instability in High-Bandwidth-Delay Network," *Int. J. Comput. Appl.*, vol. 27, no. 8, pp. 40–44, 2011.
19. V. Rajavel, "Optimizing semiconductor testing: Leveraging stuck-at fault models for efficient fault coverage," *Int. J. Latest Eng. Manag. Res. (IJLEMR)*, vol. 10, no. 2, pp. 69–76, 2025.
20. I. Ganie and S. Jagannathan, "Online lifelong optimal tracking control of nonlinear continuous-time strict-feedback systems using deep neural networks," *Neural Netw.*, vol. 191, no. 11, p. 107793, 2025.
21. A. Marks and Y. Rezgui, "Security awareness in virtual communities: The case of non-collocated academic research collaborations," in *Collaborative Networks for a Sustainable World*, *Springer*, Berlin, Germany, 2010.
22. C. Vorakulpipat, A. Marks, Y. Rezgui and S. Siwamogsatham, "Security and privacy issues in social networking sites from user's viewpoint," in *Proc. PICMET '11: Technology Management in the Energy-Smart World*, Portland, Oregon, United States of America, 2011.

Publisher's Note: The publisher remains impartial concerning jurisdictional claims in published maps and institutional affiliations. Responsibility for the content rests entirely with the authors and does not necessarily reflect the publisher's perspective