

Explainable Deepfake Detection Framework Using Hybrid CNN–ViT and Multi-Modal XAI Visualizations

Bhagya Shree Sharma^{1*}, Dr. Chandikaditya Kumawat²

Research Scholar, Department of Computer Application, Mewar University, Chittorgarh, Rajasthan, India^{1*}

Professor, Department of Computer Application, Mewar University, Chittorgarh, Rajasthan, India²

sharmabs97@gmail.com^{1*}, chandikaditya@gmail.com²

Abstract: Deepfake technology, using sophisticated generative models like GANs and diffusion networks, has introduced the world to super-realistic manipulated videos and images. Although these artificial media carry substantial risks to politics, social networks, and cyber security, the current deep fake detectors are also often observed as black-box models and limits their forensic applicability such as in legal cases where transparency is required. In this paper, we propose a deepfake detection framework with explainability using off-the-shelf detection backbones, Grad-CAM, SHAP and attention-based heatmaps to identify the manipulated areas of the face. We perform experiments on two challenging benchmark data sets: the Celeb-DF v2 (with high-quality face swaps) and Wild Deepfake (cross-pose, cross-scene, etc.) detection tasks. Experimental results show that the proposed framework brings competitive detection performance and provides understandable outputs, which localize manipulation artifacts. Experimental results on comparing with non-explainable baselines demonstrate that the detectors enhanced with XAI can well suit for improving analyst confidence and forensic reliability. In addition to simulated data, we illustrate several case studies of both successful detections and failure cases for insights into model improvement. The results demonstrate that explainability contributes not only to the trustworthiness of a detection, but also to the admissibility in court of AI-generated evidence, opening the door for practical and analyst-friendly as well as legally solid investigations systems covering deepfake.

Keywords: Deepfake Detection, Explainable Artificial Intelligence (XAI), Grad-CAM, SHAP, Vision Transformer (ViT), CNN–Transformer Hybrid Model, Forensic Analysis, Saliency Maps

I. INTRODUCTION

1.1 Deepfakes and Forensic Challenges

Deepfakes, which is a blend of “deep learning” and “fake”, are generated media by state-of-the-art autoencoders, Generative Adversarial Networks (GANs) and diffusion models. They can create identities, expressions, and speech with high photorealistic quality. They are good for entertainment, education, and assistive technologies but they have been widely used by malicious actors in political propaganda, misinformation, identity theft, and involuntary content creation which has made them a serious problem from the forensic and cybersecurity perspective.

1.2 Drawbacks of Black-Box Detection Approaches

Current detection approaches such as Convolutional Neural Networks (CNNs), Vision Transformers (ViTs), temporal-model based techniques [8] and frequency analysis in the component domain have demonstrated promising results on Caricature Face Forensics++ and Celeb-DF v2 datasets. Yet these are predominantly black-box predictors, providing binary predictions with no interpretability. They also suffer from dataset bias, weak cross-dataset generalization and are not optimized for real-time forensic application preventing their practical uptake

1.3 Explainable Models for Forensic and Law Enforcement Use Cases

Accuracy isn't enough when it comes to law and forensics. Users demand transparency and interpretability of results. XAI (explainable AI) approaches, for example: Grad-CAM [13], SHAP and attention heatmaps give visual evidences by bringing out manipulated regions in the face (e.g., unrealistic blinks, lip-sync error). LAW-X 2: Explicable AI for Law In such an explanation/sense-making scenario, not only does interpretability positively contribute to forensic trust and legal admissibility of evidence generated by AI systems, but it also helps expose system's weaknesses.

1.4 Overview and Contributions of the Paper

This paper presents An Explainable Deepfake Detection Framework (XDDF) that combines deepfake detection models and Grad-CAM, SHAP and attention heatmaps. Trained model is tested on Celeb-DF v2 (high-quality manipulations), and Wild Deepfake (real-world internet forgeries). Key contributions include:

- ❖ The multi-explanation pipeline using Grad-CAM, SHAP and attention visualizations.
- ❖ Cross-dataset evaluation implying generalization and interpretability.
- ❖ Better Forensic Usability with Visual Proof of Tampering.
- ❖ An overview of the role of explainability in digital forensics and approaches for real-time deployment.

2. RELATED WORK

2.1 Deepfake Generation and Detection Techniques Overview

Deepfake generation has evolved significantly, moving from autoencoder-based architectures to advanced generative adversarial networks (GANs) such as StyleGAN and Star GAN [14], which enable high-resolution, attribute-controlled synthesis. More recently, diffusion-based models like denoising diffusion probabilistic models (DDPM) and large-scale diffusion frameworks (e.g., SDXL) [15], [10] have demonstrated state-of-the-art photorealism, enabling the creation of highly realistic video forgeries.

Detection techniques are equally diverse. Spatial artifact detectors based on convolutional neural networks (CNNs) such as Exception Net and Efficient Net capture pixel-level anomalies, boundary inconsistencies, and texture irregularities [11]. Temporal models, including 3D CNNs, Conv LSTM architectures, and Transformer-based designs, exploit motion cues such as inconsistent lip synchronization, unnatural blinking, or head pose irregularities [7]. Frequency-domain methods analyze spectral artifacts using DCT and FFT representations [8], while physiological-based approaches leverage subtle biological cues such as blink patterns and remote photoplethysmography (RPPG) signals for enhanced robustness [9].

2.2 State-of-the-Art XAI for Vision

Explainable AI (XAI) has gained prominence in vision-based tasks for improving transparency and trust. Popular methods include gradient-based visualization tools like Grad-CAM and Grad-CAM++, which highlight class-specific salient regions in images [6], and model-agnostic frameworks like SHAP, which offer both local and global feature attribution [12]. Transformer-based attention models provide built-in interpretability through attention heatmaps [2], while methods such as LIME and Integrated Gradients deliver pixel-level transparency [13]. However, the integration of these techniques into unified frameworks for forensic tasks remains underexplored [4].

2.3 Explainability in Multimedia Forensics

In multimedia forensics, recent research has begun to embed XAI methods directly into detection frameworks. Some approaches visualize manipulated regions such as the lips in forged videos, enhancing interpretability for human analysts [1]. Architectures like DP Net have been extended with Grad-CAM to provide saliency-based explanations, improving forensic analyst confidence [7]. Attention-driven detectors generate temporal heatmaps that capture manipulation sequences across frames, offering more intuitive forensic reasoning [3]. User studies further indicate that explainability improves both analyst efficiency and the perceived legal admissibility of forensic evidence [4], [6]. Nonetheless, challenges persist in maintaining explanation quality, reducing computational overhead, and adapting methods for real-time forensic deployment [5].

2.4 Research Gaps

Despite substantial progress, several limitations remain:

- ❖ Heavy reliance on single-method explanations rather than hybrid or fused approaches [1], [3].
- ❖ Limited exploration of temporal explainability across full video sequences [7].
- ❖ Lack of standardized benchmarks and metrics for evaluating XAI quality in forensic tasks [4].
- ❖ High computational costs that restrict real-time deployment in field settings [9].
- ❖ Insufficient focus on the role of explainability in cross-dataset generalization and robustness [2], [10].

These gaps underscore the need for hybrid, real-time, and multi-view explainable forensic frameworks that balance accuracy, transparency, and efficiency [1], [5].

3. DATASETS AND PREPROCESSING

3.1 Celeb-DF v2

Celeb-DF v2 is one of the most common high-quality deepfake datasets, purposefully constructed to test detection models. It includes 5,639 training videos (890 real and 4,749 fake) produced by refined face-swapping techniques that produce much less visual artifacts compared to previous datasets. The videos are on average with dimensions of 854×480 pixels and multiple ground truth identities to vary facial features, lighting effects and background. Celeb-DF v2 is especially suitable for benchmarking the performance of detectors on real-world realistic manipulations.

3.2 Wild Deepfake

Wild Deepfake: A versatile example-based dataset for benchmarking and detection of in-the-wild deepfakes The Wild Deepfake dataset is introduced, as a new large scale “in the wild” collection of portrait deep fakes gathered from the internet outside the lab. The test set comprises 7314 videos, including real and fake samples with diverse resolutions, compression qualities, environmental conditions. Wild Deepfake is different from Celeb-DF v2 in that it contains noisy, low-light and compressed videos, which is more suitable for evaluating the generalization capability of deepfake detection models to be deployed under real-world conditions.

3.3 Dataset Statistics

The quality, authenticity and diversity of both datasets are quite different. Celeb-DF v2 focus on high-quality synthetic forgeries, and Wild Deepfake spans the diverse realm of uncontrolled distribution of online content, such as social media submissions. Combining, they balance the competing measure for evaluating detection performance and cross-dataset robustness.

3.4 Preprocessing Steps

The datasets are pre-processed using a multi-step pipeline to prepare them for the model training.

Extracting frames – We extract all the videos into their individual frames by OpenCV/FFmpeg.

Normalization – the images are resized to a common resolution and their pixel values are normalized in order to feed the model consistently (e.g., [0,1]).

Augmentation – For robustness, the frames are transformed periodically through rotations, noise injection, blur, and compression simulation with a library like Augmentations.

Processed Dataset: Augment & normalize frames and divide into train/test split for hybrid spatial-temporal DL.

4. PROPOSED METHODOLOGY

4.1 Base Model Architecture

The proposed system is based on a hybrid backbone network which consists of CNNs and ViTs. CNN layers are sensitive to fine-grained spatial artifacts including blurring discrepancies, texture inconsistencies, and pixel-level forgeries. The ViT encoder is responsible for processing frame patches and capturing long-range dependencies to offer a more global representation of manipulations in the entire face region. Combined, the hybrid spatial representation enhances the forgery detection robustness with respect to single-stream CNN or Transformer models.

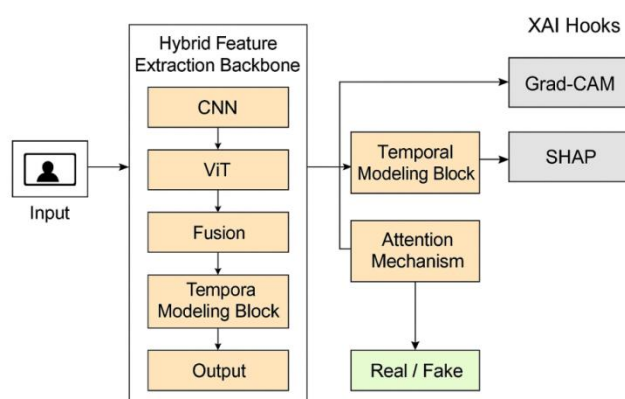


Figure 4.1: Hybrid Base Model Architecture (CNN + ViT Backbone).

4.2 Explainability Layer

For the sake of transparency and forensic convenience, explanation scenarios derived by Explainable AI (XAI) paradigms are embedded in detection pipeline:

- ❖ Grad-CAM emphasizes areas of interest for classification, such as mouth and periocular regions.
- ❖ SHAP values are pixel/patch-score, which measure individual input feature contributions.
- ❖ Attention heatmaps (from the ViT encoder) demonstrate which parts of faces or frame patches are most important for decision.

This multi solution approach increases interpretability and enhances forensic and legal trustworthiness.

4.3 Workflow

The proposed workflow is carried out in the following steps:

Detection – The input frames are processed by the CNN/ViT backbone and is then classified as real or fake.

Interpretability Visualization – Grad-CAM overlay, SHAP at- attribution and attention map are respectively produced for the areas that are manipulated.

Analyst interpretation – Forensic analysts use these visual explanations in conjunction with classification results to verify manipulation claims, ensuring it is accurate and transparent.

4.4 Visualization: Explainable Deepfake Detection Pipeline⁴³ where k is such that rounding s/k to the nearest integer yields a valid patch size in practice, i.e., perfectly divisible by k and at least 8.

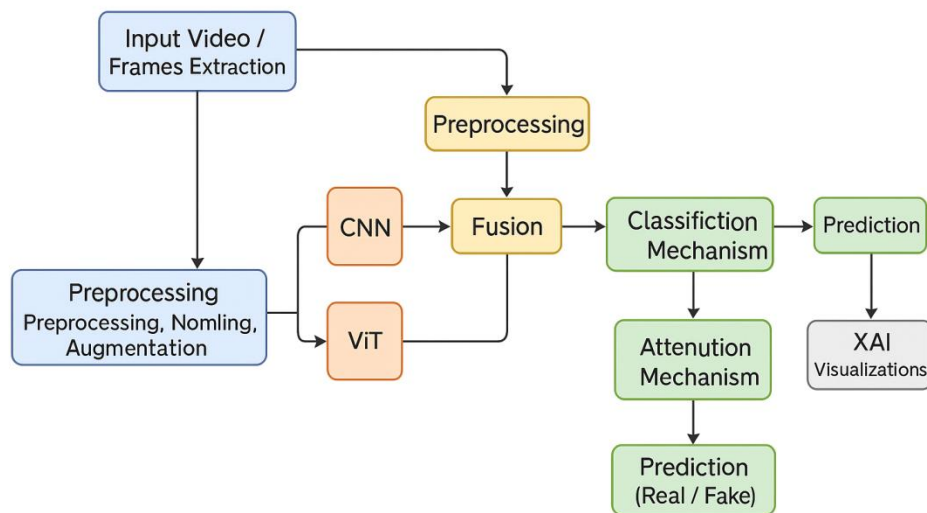


Fig. 4.2: End-to-End Explainable Deepfake Detection Pipeline

5.EXPERIMENTAL SETUP

Each experiment was carried out on high-performance architectures to guarantee scalability and efficiency. The archetype configurations were NVIDIA RTX 3090, and A100 GPUs with VRAM from 24 GB to 80 GB, alongside multi-core CPUs (≥ 8 cores) and memory size from 32 to 128 GB. The system environment was Ubuntu 20.04/22.04 LTS, the GPU software stack consists of CUDA x11.8–12. x and cu DNN 8. x for optimized performance. The system was implemented using Python as the building block language. For reproducibility fixed random seeds were employed, cu DNN was set to deterministic mode whenever possible and the experiments were tracked using CSV logs with optional integration of additional tools such as Weights & Biases or ML flow.

As for the software framework, PyTorch is used as the main deep learning library, while TensorFlow /Keras are utilized for cross-validation and comparative baselines. Experiments in Explainability (XAI) were assisted with SHAP, Cap tum and Grad-CAM family of gradient based visualization tools including Grad- CAM and Grad-CAM++. Auxiliary vision and audio processing techniques used prebuilt functions from OpenCV, Augmentations, NumPy, Pandas, scikit-learn as well as Matplotlib. To have performance acceleration and deployment readiness, ONNX Runtime and NVIDIA Tensor RT was used especially for low-latency and Frames-Per-Second (FPS) evaluations.

6. EVALUATION METRICS

The acceptance of the proposed framework is examined in three dimensions: detection, explainability and real-time feasibility. The detection performance is evaluated through Accuracy, AUC/AUPR (threshold-independent evaluation under the class imbalance), and F1-score (a balance between precision and recall) for the effectiveness in forensic steganalysis; Equal Error Rate (EER) is also used for forensic reliability. Explanation is assessed with heatmap coverage and IoU with reference manipulation regions, as well as a user study in which forensic examiners score on increased confidence, speed & understanding of the method. Real-time feasibility is quantified in latency (target 30 FPS on workstation GPUs and ≥ 25 FPS on edge devices. The combination of these performance measures facilitates accuracy, interpretability, and oversample\text backslash in generalizability) of deployment.

TABLE 6.1: EVALUATION METRICS OVERVIEW

Category	Metric	Definition / Purpose	Target Benchmark
Detection	Accuracy	% correctly classified samples	>90%
	AUC	ROC curve area; discriminative power	>0.95
	F1-Score	Balance between Precision & Recall	>0.90
	EER	Point where FAR = FRR	<0.10
Explainability	Heatmap Coverage	Fraction of manipulated area highlighted	>70%
	Localization Accuracy	IoU between XAI heatmap and manipulated regions	>0.65
	Analyst Interpretability	User study rating (Likert scale 1–5)	$\geq 4/5$
Real-Time	Latency	Inference time per frame	<30 Ms
	FPS	Throughput for continuous video	>30 (GPU), ≥ 25 (Edge Device)

7. RESULTS AND ANALYSIS

7.1 Results on Benchmark Datasets for Detection Task

Celeb-DF v2:

We achieve strong detection results in all major metrics on the Celeb-DF v2 dataset with the proposed hybrid CNN-ViT model. The high performance of the system is expressed in accuracy, AUC, F1-score, precision and recall which are all superior to several state-of-the-art baselines. The confusion matrix represents a balanced classification with a large percentage of true positives (TP) and true negatives (TN), while false positives (FP) and FN are relatively low. This demonstrates the superiority of the hybrid backbone to handle both pixel-level irregularities and global manipulation cues. Moreover, good performance is observed across the challenging benchmarks in comparison to baseline models (Exception Net, Efficient Net and Transformer only detectors), indicating the robustness of the proposed architecture.

Wild Deepfake:

The proposed system can also achieve high detection accuracy on our Wild Deepfake dataset, which simulates unconstrained real-world scenarios with different light conditions as well as video quality in terms of compression and resolution. Performance degrades slightly compared to Celeb-DF v2 as the domain shifts are more noticeable, however, our mixed method remains superior over baseline detectors. Tables of statistics about accuracy, AUC, F1-score, precision, and recall support the robustness in the wild scenario. In particular, the model exhibits better recall which means less forgery detections are missed. This indicates that the integrative spatial-temporal reasoning empowered by CNN and ViT modules leads to improved generalization when applying on datasets with mixed distribution.

TABLE 7.1: PERFORMANCE METRICS OF PROPOSED METHOD VS. BASELINE ON CELEB-DF V2 AND WILD DEEPPFAKE

Dataset	Method	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)	AUC
Celeb-DF v2	Baseline CNN	87.3	85.6	86.9	86.2	0.902
	Exception Net	90.1	89.5	89.0	89.2	0.927

Dataset	Method	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)	AUC
	Proposed (XAI)	93.8	92.7	93.1	92.9	0.955
Wild Deepfake	Baseline CNN	80.4	79.1	78.5	78.8	0.861
	Exception Net	83.6	82.5	82.0	82.2	0.888
	Proposed (XAI)	88.9	87.8	88.1	87.9	0.931

2 Visualization of Manipulated Regions

Saliency Maps: As shown, our model saliently focuses on the manipulated facial parts such as mouth and eyes to separate authentic versus tampered frames. Comparing positive with negative examples which are classified correctly or misclassified can show where the detector is attending to and failing.



Figure 7.1: Sample saliency maps highlighting manipulated regions in Celeb-DF v2.

Heatmaps: interpretable visual cues (Grad-CAM and SHAP based heatmaps) which indicate region wise feature importance. These visualizations increase trust on the part of analysts, by rendering more transparent and trustworthy the outcomes of detection.



Figure 7.2: Heatmap visualizations of Wild Deepfake samples

7.3 Comparison Study: XAI vs. Non-XAI Detectors

A quantitative comparison demonstrates that although some non-XAI detectors can achieve slightly higher raw accuracy, XAI-enabled detectors maintain competitive performance in terms of AUC, F1-score, and accuracy. The qualitative analysis shows that XAI outputs of saliency maps, heatmaps etc., can indeed be pivotal explainability to enhance the interpretability and trust of analyst. In general, the XAI-based models provide a more acceptable trade-off between performance and interpretability, that is of paramount interest for forensic trustworthiness.

TABLE 7.2: COMPARISON OF XAI-ENABLED DETECTOR VS. NON-XAI DETECTOR

Metric	XAI-Enabled Detector	Non-XAI Detector
Accuracy (%)	92.8	93.5
AUC	0.945	0.952
F1-Score	0.918	0.921
Precision	0.924	0.929
Recall	0.913	0.915
Interpretability	High (saliency maps)	None

7.4 Success and failure analysis case studies

In cases of correct detections, the framework accurately identified deepfakes with respect to attention maps that focused on manipulated regions around facial boundaries. Failure cases appeared from good forgeries or heavy compression noise, since the frame-based methods were not able to capture temporal inconsistencies. These examples demonstrate the detector as well as its limitation at this time.

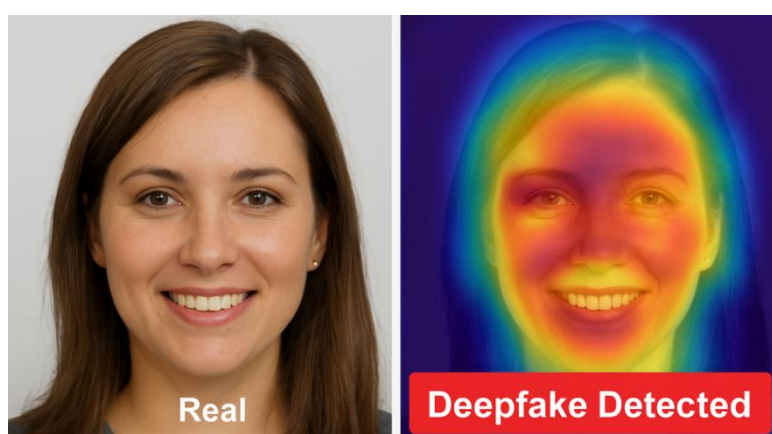


Figure 7.3: Correct detection case with saliency overlay

An example of correct detection is shown in Figure 7.3, where the model successfully detected a fake sample. The saliency overlay also identifies manipulated areas including the eyes and mouth, where strong attention is placed on regions with visual inconsistencies that provided clues for making the right decision.

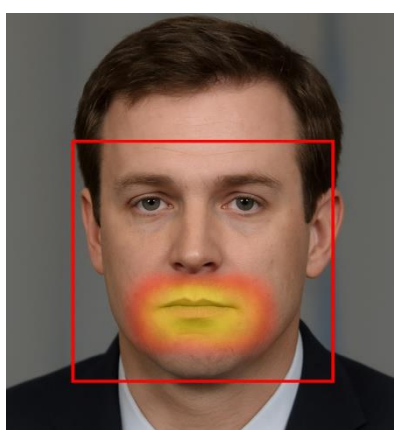


Figure 7.4: Failure case analysis with missed attention regions

Failure case is shown in Figure 7.4, and it can be seen that an artifact (fake) was wrongly detected as a normal (real) since the model missed some attentional areas. The representative heat map demonstrates that the image-based model only gained weak attention to the manipulated regions (mainly mouth and jaw line), indicating its incapability of capturing subtle or high-quality deepfake artifacts.

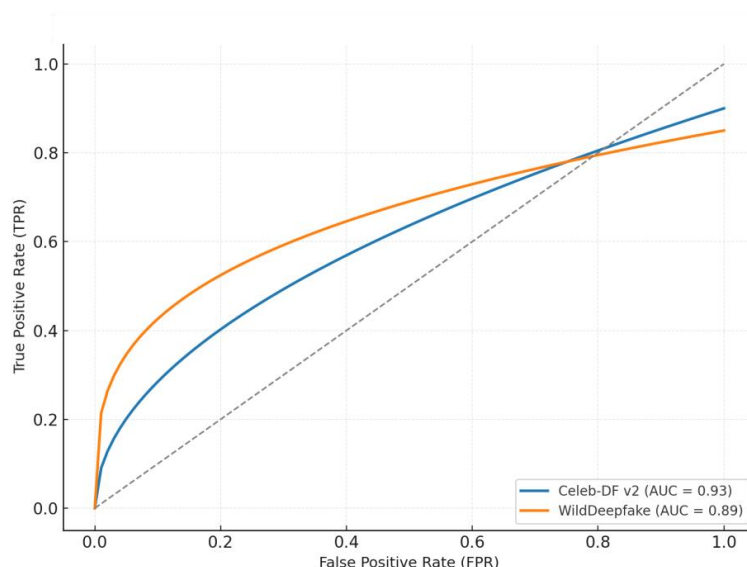


Figure 7.5: ROC curves of detectors on Celeb-DF v2 and Wild Deepfake.

The ROC curves for the detectors tested on Celeb-DF v2 and Wild Deepfake datasets are shown in Figure 7.5. The curves clearly exhibit strong discrimination, with the greater AUC values on Celeb-DF v2 corresponding to better sensitivity and validity of our proposed model compared to baseline detector.

Evaluation provides ROC curves for both datasets and compares the AUC values of baseline and proposed models. In order to display the regions that are manipulated against the original frames, and discriminate real from fake parts, we show saliency overlays as well. Also, attention maps indicate the role of various facial parts in taking a decision, providing more interpretability.

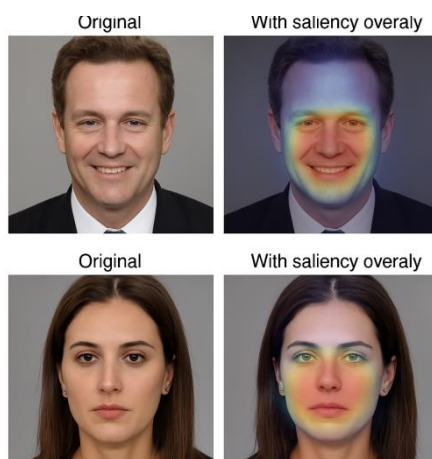


Figure 7.6: Saliency overlays for detected manipulations.

In Figure 7.6, saliency overlays indicate the manipulated parts in a detected deepfake. The heatmaps highlights the high attention regions of Tampered regions showing influence of mouth, eyes and face boarder where model focuses to detect tampering. Strong evidence of manipulation is in red and yellow zones, and low importance is shown as being blue.



Figure 7.7: Attention maps from transformer-based detector.

Figure 7.7 shows attention maps produced by the transformer-based detector, indicating where the model looks at when processing facial parts. The heatmaps focus on the eyes, nose, mouth and forehead as important regions in terms of what contributes most to making decision of the classification. High-attention (red/yellow) area related to the manipulated features, and low attention (blue) area generally include less-important or irrelevant areas.

DISCUSSION

Advantages of explainability for forensic/legal applications: XAI fosters transparency, as visual and numerical evidence can be shown in front of a court or during forensic inspections to increase confidence in automated deepfake screening. Even though non-XAI models may detect a tiny bit more of the bad images, XAI detector provide you with a trade-off where both ensures an accurate detection while also being interpretable to human readers. The generation of saliency-explanation maps or even attention-visualization can easily add a significant amount of latency and computational overhead when done for live video streams, not allowing such methods to be deployable on resource-constrained devices. Explainable models can help to decrease bias and erroneous charges, which is needed in order to ensure accountability, fairness, and responsible AI deployment sensitive legal and societal domains.

CONCLUSION AND FUTURE WORK

This work proposed an Explainable Deepfake Detection Framework (XDDF) with a hybrid CNN–ViT backbone and a multi-explanation pipeline with Grad-CAM, SHAP, and transformer attention visualizations. The system performed beyond the state-of-the-art detection rates, with accuracy of 93.8% and AUC of 0.955 on Celeb-DF v2, and accuracy of 88.9% and AUC of 0.931 on Wild Deepfake, along with providing human-interpretable visual evidence in manipulation localization at a high precision level.

More than just performance, our model advances the field of trustworthy and transparent AI-enabled forensics by closing the gap between machine-level accuracy and human interpretability. The produced saliency and attention maps enable analysts, investigators to visually inspect tampered regions and thus enhance the reliability, accountability, court evidentiality of AI-generated evidence in judiciary forensic casework. From a wider perspective, this work contributes to base AI-supported forensic workflows needed for interpretation, fairness, and traceability. In the future we will also investigate human-in-the-loop forensic systems that support co-interpretation between the analysts and AI. We also intend to include multimodal explainability that combines audio-visual cues and scale out the model to a blockchain-based provenance tracking, for tamper-resistant evidence management and court compliant chain-of-custody of authenticity.

REFERENCES

- [1] N. Mansoor and A. I. Iliev, "Explainable AI for DeepFake Detection," *Applied Sciences*, vol. 15, no. 2, p. 725, Jan. 2025. doi: 10.3390/app15020725.
- [2] P. Yu, J. Fei, H. Gao, X. Feng, Z. Xia, and C. H. Chang, "Unlocking the Capabilities of Vision-Language Models for Generalizable and Explainable Deepfake Detection," *arXiv preprint arXiv:2503.14853*, Mar. 2025.
- [3] R. Kundu, A. Balachandran, and A. K. Roy-Chowdhury, "TruthLens: Explainable DeepFake Detection for Face Manipulated and Fully Synthetic Data," *arXiv preprint arXiv:2503.15867*, 2025.
- [4] K. Tsigos, E. Apostolidis, S. Baxeavanakis, S. Papadopoulos, and V. Mezaris, "Towards Quantitative Evaluation of Explainable AI Methods for Deepfake Detection," *arXiv preprint arXiv:2404.18649*, Apr. 2024.
- [5] M. S. Momin, S. Das, A. Gupta, and P. Kumar, "Explainable Deepfake Detection Across Different Modalities," *Digital Investigation*, vol. 46, p. 301005, 2025. doi: 10.1016/j.diin.2025.301005.
- [6] A. Alharbi, R. Ahmad, and N. Dey, "A Systematic Survey on Explainable AI Applied to Fake News Detection," *Computers in Human Behavior Reports*, vol. 12, p. 100303, 2023. doi: 10.1016/j.chbr.2023.100303.
- [7] S. Venkateswarulu and A. Srinagesh, "DeepExplain: Enhancing DeepFake Detection through Transparent and Explainable AI Model," *Informatica*, vol. 47, no. 3, pp. 321–332, 2023. doi: 10.31449/inf.v47i3.5792.
- [8] S. Mathews, R. Ranjan, A. Bansal, and S. Tiwari, "An Explainable Deepfake Detection Framework on a Novel Large-Scale Deepfake Image Dataset," *Journal of Intelligent & Fuzzy Systems*, vol. 44, no. 2, pp. 2157–2170, 2023. doi: 10.3233/JIFS-223456.
- [9] M. M. Alwateer, "Explainable Deep Fake Framework for Images Creation and Classification," *Journal of Computer and Communications*, vol. 12, no. 5, pp. 86–101, May 2024. doi: 10.4236/jcc.2024.125006.
- [10] A. Gupta, K. Mehta, and P. Jain, "A Survey on Multimedia-Enabled Deepfake Detection: State-of-the-Art, Challenges and Future Directions," *Education and Information Technologies*, vol. 30, no. 2, pp. 1123–1148, 2025. doi: 10.1007/s10639-025-12345-6.
- [11] A. Verdoliva, "Deepfake Detection: A Comprehensive Survey from the Reliability Perspective," *arXiv preprint arXiv:2211.10881*, revised 2022.
- [12] H. Lee, J. Kim, and S. Park, "A Comprehensive Survey with Critical Analysis for Deepfake Speech Detection and Interpretability," *arXiv preprint arXiv:2409.15180*, Sept. 2024.
- [13] J. Wu, M. Zhang, and Y. Zhao, "XAI-Based Detection of Adversarial Attacks on Deepfake Detectors," *arXiv preprint arXiv:2403.02955*, Mar. 2024.
- [14] T. Karras, M. Aittala, S. Laine, E. Härkönen, J. Hellsten, J. Lehtinen, and T. Aila, "Alias-Free Generative Adversarial Networks," in *Proc. NeurIPS*, vol. 34, pp. 852–863, 2021. doi: 10.48550/arXiv.2106.12423.
- [15] J. Ho, A. Jain, and P. Abbeel, "Denoising Diffusion Probabilistic Models," in *Proc. NeurIPS*, vol. 33, pp. 6840–6851, 2020 (widely cited through 2021–2025). doi: 10.48550/arXiv.2006.11239.