

FIND: An Internet wide Deepfake Detection Tool

Pradhyumna Jadhav¹, Reuben George² and Sunita Kulkarni³

¹⁻³MIT World Peace University, Dept of Electrical and Electronics Engineering, Pune, India
Email: pradhyuma.jadhav@gmail.com, reubengrg2624@gmail.com, sunita.kulkarni@mitwpu.edu.in

Abstract— The proliferation of deepfakes necessitates robust, multi-faceted detection mechanisms. This work introduces a synergistic framework that integrates two specialized pipelines to address the challenges of both video and image forgery detection [1]. For video deepfake detection (temporal analysis), a resource-efficient architecture combines a classical ResNet feature extractor with a Quantum Long Short Term Memory (QLSTM) network [1]. This quantum-enhanced model demonstrates up to a 99.16% reduction in recurrent layer parameters compared to classical LSTM, achieving performance superior to or comparable with its classical counterpart [1]. For image deepfake detection (frequency/spatial analysis), we adopt the WMamba module, which excels at exploiting subtle forgery artifacts by applying Dynamic Contour Convolution (DCCConv) to hierarchical Wavelet representations [1]. WMamba achieves state-of-the-art (SOTA) performance in cross-dataset evaluation [1]. This dual approach confirms that specialized, multi-modal strategies are critical for resilient, computationally viable deepfake detection systems.

Index Terms— Deepfake Detection, Quantum LSTM (QLSTM), Wavelet Transform, State Space Models.

I. INTRODUCTION

The advent of deepfakes — realistic yet synthetic media — has been bolstered by the improvements in the field of Artificial Intelligence and Machine Learning. These deepfakes can easily imitate people in images and videos, making it very difficult for the naked eyes to determine the fakes from the originals. While deepfakes can be used positively in many fields such as personalized advertisement targeting, movies and entertainment to depict various “stunts” as originally done by the actors themselves, and many other such avenues, the misuse creates equally drastic consequences. Deepfakes can be used to imitate people in power to cause chaos and spread misinformation, while loved ones can be imitated in ransom scams.

II. LITERATURE SURVEY

The current meta of deepfake detection technology is in a constant battle of tug of war: the generating models get more advanced, using the high speed of GPU developments to fuel their own agenda: better deepfakes that are not only difficult but sometimes impossible for the current detection models to detect — all for just a measly \$9.99 subscription, or sometimes free.

The limitations exist due to the field having rapidly moved past simple Convolutional Neural Networks (CNNs), and the challenges of video compression and real-world degradation of the data make it even more difficult to detect after multiple reshares and filters.

We will first discuss the historical and technological path of deepfake detection, identify why it is negatively affected by data laundering that happens due to online resharing, and justify the proposed next-generation hybrid solutions, specifically the role of 6 Qubit Quantum Long-Short Term Memory (QLSTM) in video detection temporally and the Wavelet Mamba (WMamba) model for high-resolution spectral image-based deepfake detection. These models directly make up for the compression vulnerability identified in classical approaches, encouraging us to utilize the FaceForensics++ (FF++) dataset. The FF++ dataset plays a vital role for training, given its inclusion of diverse manipulation techniques and controlled compression artifacts [1]. The initial successful models for deepfake detection relied on previously proven and established methods designed to focus on the predictable cues left behind by manual tampering. Then the first generation of mainstream deepfake detectors utilized powerful CNN architectures, often pre-trained on large-scale datasets. XceptionNet, originally trained on ImageNet, proved to be a strong baseline, demonstrating the ability to automatically learn fundamental manipulation patterns recovered from the data. XceptionNet could reach over 95% The use of deep learning techniques was a far cry from traditional detection methods that usually deployed steganalysis combined with Support Vector Machines (SVMs), which looked for the patterns that were based on statistically derived tells that manipulators left behind while faking any media. These traditional methods proved highly unpredictable and were adversely affected exposed to compression, as compression would normalise the aforementioned tells. [5] While CNNs offered greater robustness, researchers soon discovered that even the statistical patterns they discovered were tied to high-frequency components—subtle cues that reside in the spatial frequency domain of the image.

[1] The fact that ensemble architecture later integrated explicit frequency methods, such as Wavelet Transform alongside XceptionNet [6], suggested that relying purely on spatial CNN feature learning was not robust enough to preserve and model high-frequency cues against degradation. The fundamental weakness of these early models was their inherent dependence on low-level statistical artifacts that was highly subject to loss through common post-processing compression.[5] Due to the improving deepfakes, simple frame-by-frame analysis became redundant, as individually judging each frame was not enough to attach correlation between frames. This led to the focus shifting towards temporal inconsistencies in the video itself. Hence temporal artifacts were deemed to be critical and widespread adoption of hybrid models ensued. These architectures typically combine a CNN backbone (such as ResNet-50) for frame-level spatial feature extraction with a Recurrent Neural Network (RNN) component, often a Long Short-Term Memory (LSTM) or a Bidirectional LSTM (Bi-LSTM), to model sequential dynamics that emerge from the frame to frame sequencing in videos.[7] Through this hybridization, the models can detect both spatial and temporal discrepancies, which make it effective to identify manipulations between deepfakes and face-swaps and neural-textures in videos[7]. As a case in point, a hybrid structure that merged both ResNet-50 when dealing with visual frames and Bi-LSTM when dealing with audio streams had high accuracy rate of 94.6 percent on the DFDC test set [8]. Regardless of these high benchmark scores, the complexity and computational cost of these CNN-LSTM systems had become a major constraint of deploying them[8].

Also, the current classical temporal models tend to have challenges of providing total and complete modelling of the various temporal dynamics of different forged regions. Most of such methods are based on a fixed temporal convolution, simply put short term association of frames, which restricts their capacity to capture the long term temporal interactions required to make credible detection across multiple video clips. i—human—;Most of them are based on fixed temporal convolution, akin to saying short term relationship, which limits their capacity to capture the long-term temporal interaction required to achieve credible detection between two or more video clips. The fact that achieving the highest possible accuracy was not equal to success on the ground in the real world, partly because of such factors as bias in the dataset and high memory consumption, confirmed the eventual necessity of efficiency and generalization- oriented architectures, and not only raw performance figures.[8]

III. RESEARCH GAP

The primary limitation that can be identified in both classical and mid-era deepfake detection models is that they cannot extrapolate the performance to the laboratory to the data corrupted by the real-world platform during the sharing and reposting. The problem of generalization is the generalization issue that the quantum and spectral architectural improvements are aimed at reducing.

The unknown bad actor to the deepfake detectors is the violent after-processing of the social media and video sharing sites (or YouTube and Facebook) with high traffic. These websites proactively censor and modify the content uploaded to the Web (to avoid the infamous buffering of the 2G and 3G networks), which removes low-

level forensic evidence [15]. This is a harsh process since compression as such causes the loss of high-frequency components [1]. These high-frequency signals are important in the detection of the nuanced artifacts added by deepfake generation models that lead to the performance of deepfake detectors to degrade substantially when compressed.[1]. Such a systemic issue is a drastic change of domain. In response to this, scholars have been forced to come up with sophisticated structures that imitate the video sharing pipelines of social networks. To scale compression and resizing parameters to model platform-specific properties in large datasets, these emulators approximate these parameters to enable successful fine-tuning to close the gap between lab-based training and real-world tasks[15]. The extremely high attention given to the creation of tools to handle this loss of information proves compression laundering to be the most crucial factor that constrain the successful operation of the detection technology. Another complication that adds to the generalization crisis is the use of training materials that show old methods of forgery. Recent SOTA results can reveal a significant conflict of interest: leading training datasets, to their credit, include FF++, but they are based on outdated forgery methods that were widespread several years ago.[16].

The primary issue with training algorithms on these older forgeries is that it hinders the creation of useful generalization abilities that are required to detect modern, state of the art deepfakes. To deliver the much-needed generalizable forensics, it is necessary to have a detector that is able to detect a broad spectrum of unseen deepfakes.[17]. Nevertheless, currently, the majority of available datasets consist of biased representation of forgery forms, frequently being specific to certain face swapping techniques.[16]. This requires the creation of universal models that no longer emphasize technique-specific artifacts, but instead various fundamental distortions, including Face Inconsistency Artifacts (FIA) and Up-Sampling Artifacts (USA), that cut across different types of face forgery. Another thing worth mentioning is that the processing of large datasets needed to achieve real-world robustness training, particularly of high-resolution content, puts a very serious strain on computing and memory resources[8]. Thus, a solution of next generation should be designed to be architecturally resilient to forensic needs as well as highly computationally efficient, and should demand specialized backbones that surpass the resource constraints of conventional CNNs and Vision Transformers. The choice of the quantum model is a deliberate step towards a strategic choice to concentrate on the time domain as it is supposed that any motion anomalies which may occur are less susceptible to video compression effects than to artifacts of stationary frames. Classical approaches tend to emphasize only spatial inconsistency and are unable to explicitly investigate fine-grained artifacts in time[21]. The framework of a typical QLSTM cell consists of six different VQC blocks to control the internal processes, i.e., circulation of information via the input, forget, output gates, and updating of the cell state[19]. The suggested 6-qubit architecture is an optimized VQC architecture that is intended to balance between computational complexity and expressibility with a higher feature dimension capacity to encode a sequence as compared to reduced number of qubits, as used in a prior 4 qubit architecture. [24]. The design is required to recover the low-amplitude temporal distortions (e.g. irregular blinking or inappropriate continuity of head pose) that compresses slowly, and eliminate them. Such quantum based architectures are typically developed using complex methods such as Differentiable Quantum Architecture Search (DiffQAS) [14], wherein the model actively identifies the most suitable circuit parameters to optimize the quantum benefit of the complex temporal sequence analysis needed in video forensics. WMamba concentrates on the wavelet analysis which is then connected to the efficient Mamba architecture to explicitly concentrate on the high frequency elements that the compression kills. It is a well-known fact that wavelet analysis can help improve detector generalization due to the fact that wavelets are efficient at capturing key contours as well as fine-grained and globally distributed artifacts that are otherwise not detected in the normal spatial domain[12]. This model is a manifestation of both classic forensic principles and contemporary deep learning efficiency, that have shown that SOTA detectors are slowly but steadily returning basic spectral analysis to the very essence of their architecture.

The WMamba architecture gets beyond a major limitation of the earlier wavelet architectures which did not make the most of the properties of wavelet data[12]. Using Mamba, WMamba finds long-range spatial relationships of linear complexity, thus allowing the extraction of fine-grained artifacts that are globally distributed and consistent even in small image patches to be robustly extracted by WMamba [12]. Because compression factors and eliminates the high-frequency signal content in which forgery artifacts are held [1], the ability of WMamba to effectively extract and process these spectral characteristics over the whole image offers a direct and robust mitigation against spectral smearing of implementation. The use of both WMamba and QLSTM is justified by the fact that the former specializes in videos, and is a spin-up of our earliest intention, which was to discover deepfakes on the internet, predominantly in the form of images, which is addressed by WMamba.

IV. PROPOSED SYSTEM

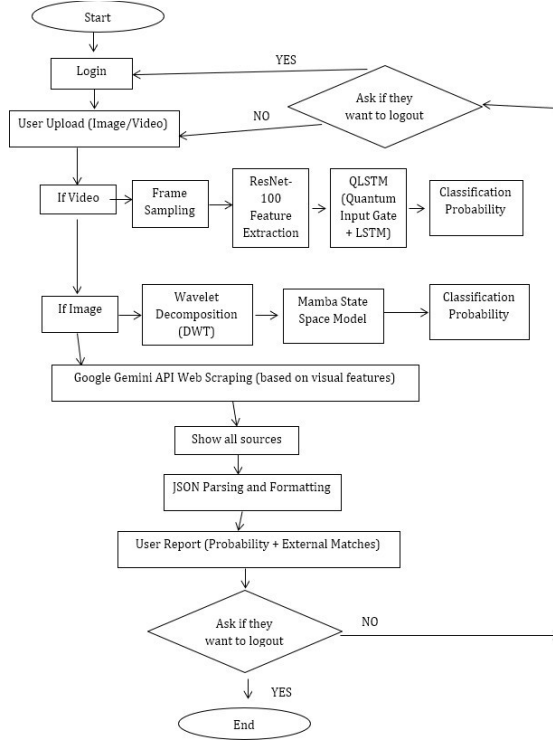


Fig. 1: System Architecture

This deepfake detection framework has two optimized pipelines, each designed for a specific type of data: video and images.

IV. PROPOSED SYSTEM

A. Hybrid Video Detection: ResNet-QLSTM

The QLSTM network replaces classical neural network components with Parameterized Quantum Circuits (PQCs) to achieve massive parameter reduction and faster learning.

Classical features (latent vectors I) extracted by ResNet are transformed into a quantum state via a Feature Map. We employ Angle Encoding, which is preferred for its simplicity with floating-point data. A single feature value $x \in \mathbb{R}$ is mapped to a rotation angle applied by a Pauli rotation gate R_k ($k \in \{x, y, z\}$):

$$R_k(x)|0\rangle = e^{-ix((\sigma k)/2)}|0\rangle$$

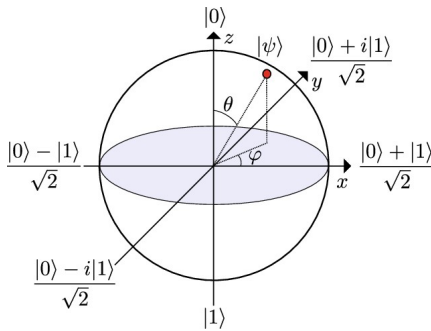


Fig. 2: Bloch Sphere

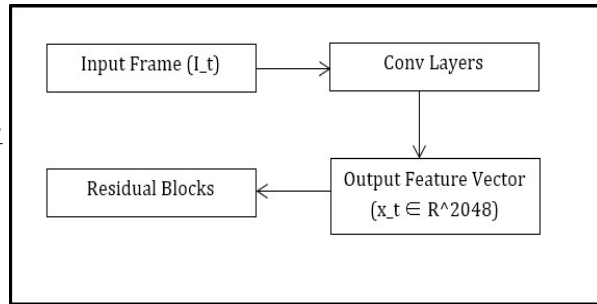


Fig. 3: Input Cycle

B. Input Representation

Given a video V consisting of T frames:

$$V = \{I_1, I_2, \dots, I_T\}, I_t \in \mathbb{R}^{H \times W \times C} \quad (1)$$

where I_t is the image frame at time t , and H , W , and C represent the height, width, and number of colour channels respectively.

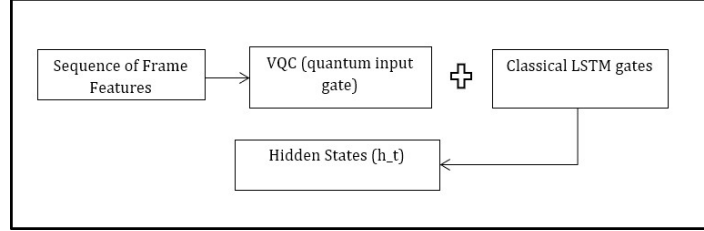


Fig. 4: Internal working of QLSTM Block

C. Feature Extraction

Each frame I_t is passed through a spatial feature extractor, typically a convolutional neural network (CNN)

$$x_t = f_{\text{CNN}}(I_t; \theta_{\text{CNN}}), \quad x_t \in \mathbb{R}^d \quad (2)$$

where f_{CNN} is a function parameterized by θ_{CNN} that produces a d -dimensional feature vector for each frame.

D. Temporal Modelling with QLSTM

The sequence of extracted features $X = \{x_1, x_2, \dots, x_T\}$ is input to a Quantum Long Short-Term Memory (QLSTM) network, which models temporal dependencies across video frames.

The classical LSTM cell equations for a time step t are defined as:

$$f_t = \sigma(W_f x_t + U_f h_{t-1} + b_f) \quad (3)$$

$$i_t = \sigma(W_i x_t + U_i h_{t-1} + b_i) \quad (4)$$

$$o_t = \sigma(W_o x_t + U_o h_{t-1} + b_o) \quad (5)$$

$$\tilde{c}_t = \tanh(W_c x_t + U_c h_{t-1} + b_c) \quad (6)$$

$$c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t \quad (7)$$

$$h_t = o_t \odot \tanh(c_t) \quad (8)$$

In the QLSTM, part or all of these operations are replaced or enhanced by quantum circuits Q , allowing hybrid quantum-classical computation:

$$h_t = Q(x_t, h_{t-1}; \theta_Q) \quad (9)$$

where $Q(\cdot)$ is a parameterized quantum circuit acting on quantum states encoding the input feature x_t and the previous hidden state h_{t-1} , with parameters θ_Q trained via hybrid quantum-classical optimization.

The classical LSTM gates (forget, output, cell updates) operate as usual but instead of a classical input gate, a Variational Quantum Circuit (VQC) encodes frame features into quantum states. The VQC applies quantum rotations and entanglements on qubits, producing non-classical representations that are measured and fed as the input gate. Quantum simulation (using PennyLane) enables this gate to capture complex temporal dependencies and subtle frame-to-frame inconsistencies harder for classical models to represent.

E. Image Deepfake Detection: Wavelet + Mamba

For images, the system detects subtle visual artifacts by working in the frequency domain via wavelet decomposition: The input image I undergoes a 2D Discrete Wavelet Transform (DWT) and DWT outputs four sub-bands: low-frequency (LL) and three high-frequency (LH, HL, HH) capturing directional textures. These high-frequency sub-bands are combined into a multi-channel tensor highlighting fine textures and edges often altered in fakes.

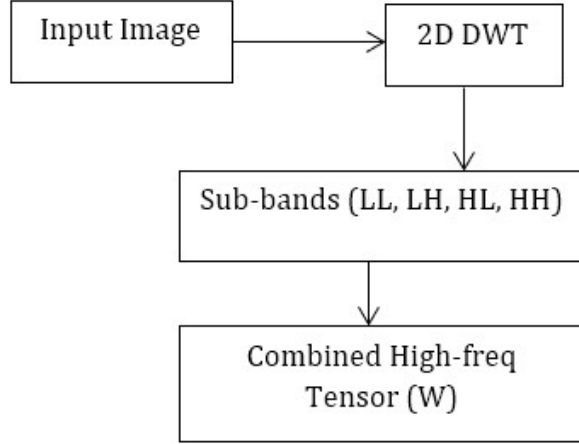


Fig. 5 Feature Extraction For WMamba

F. State-Space Modeling with MAMBA

This model captures long-range spatial-temporal correlations through a state-space representation, enabling efficient propagation of contextual information across frames.

State update:

$$s_{t+1} = As_t + Bu_t + w_t \quad (10)$$

Observation model:

$$y_t = Cs_t + Du_t + v_t \quad (11)$$

where s_t is the latent state vector at time t , u_t represents the input (e.g., spatial or frequency features), and y_t is the observed output (e.g., reconstructed features or predictions). Matrices A , B , C , and D define the system dynamics, while w_t and v_t are process and measurement noise terms, assumed to be zero-mean Gaussian.

The model parameters are learned to maximize temporal consistency and effectively discriminate between real and fake samples by minimizing a combined objective function consisting of reconstruction and classification losses.

The tensor W feeds into the Mamba visual state space model that processes W patch-wise and models a latent state vector s_t capturing long-range spatial relationships and nonlocal manipulation patterns. This helps the model to learn typical error distributions and subtle textural inconsistencies characteristic of deepfake images.

G. Classification Layer

The final hidden state h_T (from QLSTM) or the state s_T (from MAMBA) is fed into a classification layer to predict the authenticity of the input:

$$\hat{y} = \sigma(W_{\text{clf}}h_T + b_{\text{clf}}) \quad (12)$$

where $\hat{y}$ denotes the predicted probability of the input being fake or real, and σ represents the activation function—typically sigmoid or softmax, depending on the binary or multi-class setup.

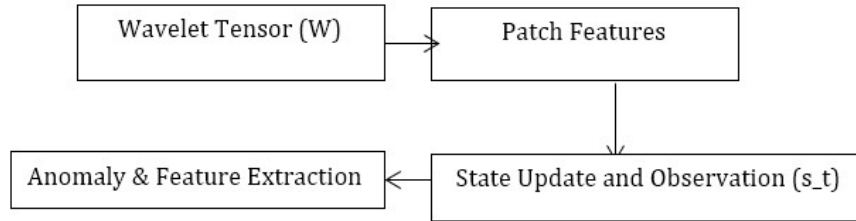


Fig. 6: Working of the Mamba Block

VI. RESULTS

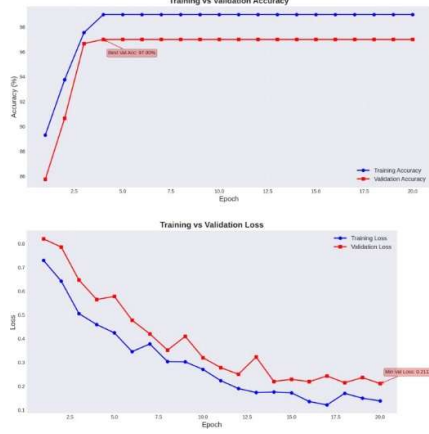


Fig. 7: Plot of Accuracy and Loss of ResNET 101+QLSTM

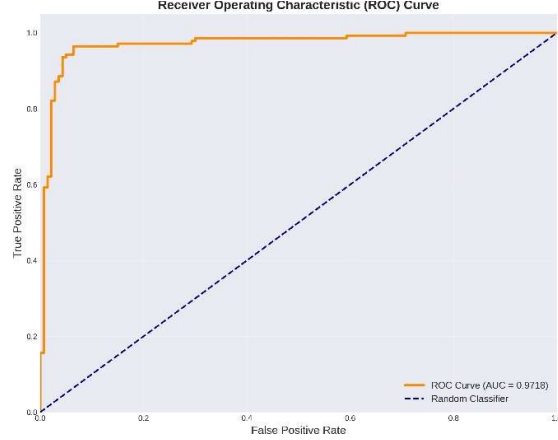


Fig. 8: ROC Curve of ResNET 101+QLSTM

The ResNet101+QLSTM model demonstrates fast convergence and stable learning across epochs. The validation accuracy closely follows training accuracy, indicating reduced overfitting and strong generalization capability. The ROC curve highlights a near-ideal trade-off between sensitivity and specificity, confirming that the QLSTM effectively models temporal dependencies in frame sequences. This result suggests that quantum-inspired recurrent mechanisms can capture fine-grained temporal variations in video forgeries more efficiently than conventional LSTM networks.

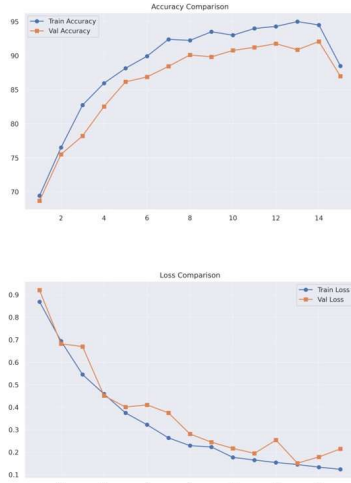


Fig. 9: Plot of Accuracy and Loss of MAMBA + HWFEB

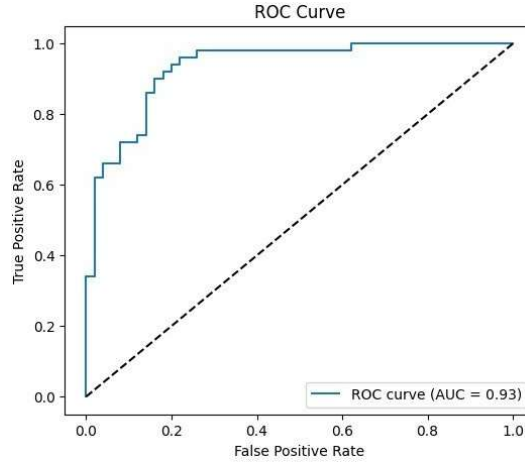


Fig. 10: ROC Curve of MAMBA+HWFEB

The MAMBA+HWFEB model exhibits steady accuracy improvements and a consistent decline in loss, proving the advantage of combining state-space modelling with wavelet-based frequency enhancement. The ROC curve further confirms its strong separability between authentic and forged samples. The integration of HWFEB allows the model to extract high-frequency texture distortions and subtle manipulation traces that are often missed in spatial-only methods, making this combination highly effective for static image forgery detection.

TABLE I: PERFORMANCE METRICS OF EXISTING AND PROPOSED MODELS

Model	Accuracy	Precision	Recall	F1-Score	AUC-ROC
ResNet 101+LSTM	88	90	88	88	97
MAMBA + ResNet	60	59.5	59.5	50	57.1
MAMBA + HWFEB	87	83.64	92	87.62	93
ResNet 101+QLSTM	96.5	96.43	96.43	96.43	97.18

The experiments show a clear difference between traditional baselines and the proposed hybrid models. Among video-based approaches, the ResNet50-QLSTM pipeline achieved the best performance, reaching 96.5 percent accuracy and consistently higher precision, recall, and F1 scores than the traditional ResNet50-LSTM baseline. This demonstrates the effectiveness of quantum-inspired recurrent modeling, which reduces the LSTM parameter load by more than 99 percent while maintaining predictive quality. On the image side, the Mamba plus HWFEB architecture outperformed the Mamba plus ResNet alternative, showing that wavelet-driven frequency features help distinguish forged artifacts more clearly. The ROC curves for both proposed models indicate stable generalization and strong separability. Overall, the results confirm that specific temporal and frequency-domain pipelines perform better than standard deepfake detectors while still being efficient enough for edge-oriented deployment.

VII. LIMITATIONS

The project has several obvious limitations. It will focus only on searching publicly accessible online platforms, such as social media, video sharing and streaming sites, and major forums, for potential deepfake content. This means it will only cover freely available data. Detection accuracy relies heavily on the availability and quality of publicly posted images and videos, which can vary widely and affect reliability. The way this issue is further intensified can be explained via the loss of high frequency data in lossy compression (affecting WMAMBA model) and the facial features as well as bodily appendages of the subject in the said content being blurred or softened due to the same lossy compression (affecting the QLSTM adversely). The system will not, or rather cannot, analyze private, encrypted, or restricted-access content to respect privacy and access regulations. Additionally, the project excludes efforts to remove or request takedowns of detected deepfake content, as legal and enforcement actions are beyond its scope. It will also not address non-visual synthetic media, like text-only deepfakes or text impersonations, focusing primarily on visual content. Automated data collection may be limited by anti-scraping measures on certain websites, which could result in incomplete coverage. Ultimately, the project aims to provide accurate detection and user notifications, but it will not offer remediation, legal support, or content removal services.

VIII. FUTURE SCOPE

Future expansions could broaden detection capabilities by integrating methods that combine, video, and physiological signals for better defense against these deepfakes. Adding real-time detection and monitoring features through browser extensions or social media integrations could help prevent the spread of deepfakes more proactively. Updating and retraining the model continuously is critical to keep up with evolving deepfake techniques and attacks. Optimizing the system for use on resource-limited devices, like smartphones and edge platforms, will increase accessibility and usability. Research into interpretable and explainable AI models for better forensic validation, as well as platform-agnostic designs that work across various digital ecosystems, will also be critical. The project may explore new quantum-inspired temporal models and scalable systems to improve detection efficiency and long-term adaptability. This aims to empower users to protect their digital identities amid growing synthetic media threats.

IX. CONCLUSION

The designed system shows that combining advanced temporal modeling techniques with a modular detection pipeline produces effective deepfake identification across various inputs. Using Quantum LSTM and State-Space Models greatly enhances the temporal analysis needed to spot subtle frame-wise manipulation. The flexible workflow, which accommodates multiple input modes and model complexity levels, allows users to adjust detection depth and speed based on their needs. While addressing key issues of scalability, reliability, and ease of use, this framework represents an important step toward next-generation AI-powered media verification tools. The project highlights how blending practical methods with cutting-edge AI research can create robust systems to reduce the increasing risks of artificial media fabrication.

REFERENCES

- [1] M. Rössler, D. Cozzolino, L. Verdoliva, C. Riess, J. Thies and M. Nießner, "FaceForensics++: Learning to Detect Manipulated Facial Images," in *Proc. IEEE/CVF Int. Conf. Computer Vision (ICCV)*, 2019, pp. 46–55.
- [2] N. U. Gal, M. A. Chaudhry, J. Younas and A. Rafiq, "Deepfake Detection Using XceptionNet," *Int. J. Sci. Res. Sci. Technol.*, vol. 10, no. 3, pp. 158–165, 2025.

- [3] M. A. Rashid, M. M. Rahman and T. Sharmin, "Hybrid Deepfake Image Detection Based on Haar Wavelet Transform and Xception Model," *Image Vis. Comput.*, vol. 130, Art. 104891, 2025.
- [4] S. Tipper, H. F. Atlam and H. S. Lallie, "An Investigation into the Utilisation of CNN with LSTM for Video Deepfake Detection," *Applied Sciences*, vol. 14, no. 21, Art. 9754, 2024.
- [5] A. Puri, C. Pancholi, B. Guo, R. Ren, T. Kang and Y. Jung, "Deepfake Detection Using CNN-LSTM and Multimodal Analysis: A Hybrid AI Approach," in *AMCIS 2025 TREOs*, p. 208.
- [6] T. Kim et al., "Beyond Spatial Frequency: Pixel-wise Temporal Frequency-based Deepfake Video Detection," in *Proc. IEEE/CVF Int. Conf. Computer Vision (ICCV)*, 2025, pp. 11198–11207.
- [7] A. Montibeller, D. Shullani, D. Baracchi, A. Piva and G. Boato, "Bridging the Gap: A Framework for Real-World Video Deepfake Detection via Social Network Compression Emulation," in *Proc. ACM DeepFake Forensics Workshop (DFF)*, 2025.
- [8] Z. Yan et al., "DF40: Toward Next-Generation Deepfake Detection," in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2024.
- [9] L. Ma et al., "From Specificity to Generality: Revisiting Generalizable Artifacts in Detecting Face Deepfakes," *arXiv:2504.04827 [cs.CV]*, 2025.
- [10] S.-Y. Chen, S. Yoo and Y.-L. Fang, "Quantum Long Short-Term Memory," *arXiv:2009.01783 [quant-ph]*, 2020.
- [11] S.-Y. Chen et al., "Differentiable Quantum Architecture Search in Quantum-Enhanced Neural Network Parameter Generation," *arXiv:2505.09653 [quant-ph]*, 2025.
- [12] S. Peng, T. Zhang, M. Zhu and Y. Pang, "WMamba: Wavelet-based Mamba for Face Forgery Detection," in *Proc. ACM Int. Conf. Multimedia (MM)*, 2025, Art. 518.
- [13] A. Pandey, B. Rudra and R. K. Krishnan, "Framework for Quantum-Based Deepfake Video Detection (Without Audio)," *Intelligent Systems*, Jun. 2025, doi: 10.1155/int/3990069.
- [14] "Quantum Computing," *Wikipedia*. Accessed: Feb. 2025. [Online]. Available: https://en.wikipedia.org/wiki/Quantum_computing
- [15] S. Bravyi and J. M. Gambetta, "Improved Classical Simulation of Quantum Circuits Dominated by Clifford Gates," *Phys. Rev. Lett.*, vol. 116, Art. 250501, 2016. doi: 10.1103/PhysRevLett.116.250501.
- [16] S. Aaronson et al., "Where Quantum Complexity Helps Classical Complexity: A Survey," *arXiv:2312.14075 [quant-ph]*, 2023. Available: <https://arxiv.org/html/2312.14075v3>
- [17] S. Dash et al., "A Systematic Review of Quantum Machine Learning for Digital Health," *NPJ Digit. Med.*, 2024. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12048600/>
- [18] SRI International, "Solving 'Unsolvable' Math Challenges with Quantum-Inspired Computers," Accessed: Feb. 2025. [Online]. Available: <https://www.sri.com/press/story/solving-unsolvable-math-challenges-with-quantum-inspired-computers/>