

A Fine-Tuned Vision Transformer for Deepfake Image Detection

<i>Kanwerjit</i>	<i>GauravDeep</i>	<i>Sumandeep Kaur</i>
<i>Dept.ofComputerScience&Engineeri</i>	<i>ng Dept.ofComputerScience&Engineerin</i>	<i>g Dept.ofComputerScience&Engineeri</i>
<i>PunjabUniversity,Patiala,India</i>	<i>PunjabUniversity,Patiala,India</i>	<i>PunjabUniversity,Patiala,India</i>

Abstract

The proliferation of algorithmically synthesized visual content broadly labelled as deepfake poses a mounting threat to the integrity of digital information ecosystems. Despite rapid advances in generative modelling, robust automated detection remains challenging, as synthesis quality now routinely exceeds the threshold of reliable human inspection. This paper fine-tunes a Vision Transformer (ViT) classifier from the publicly released `dima806/deepfake_vs_real_image_detection` checkpoint via the Hugging Face Trainer API on a balanced corpus of 190,081 images drawn from Kaggle, comprising equal proportions of authentic photographs and AI-generated samples spanning GAN-based and latent diffusion architectures. The fine-tuned model achieves an overall classification accuracy of 99.20% and a macro-averaged F1-score of 0.9920 on a held-out evaluation set of 38,081 images, with symmetric per-class error rates (128 false positives; 175 false negatives). These results demonstrate that the Vision Transformer, whose globally unconstrained multi-head self-attention mechanism enables detection of the long-range spatial incoherence characteristic of synthetic imagery, constitutes a computationally tractable and high-performing architecture for synthetic image forensics, surpassing all surveyed CNN and frequency-domain baselines on the same benchmark.

Keywords—deepfake detection; Vision Transformer (ViT); synthetic image classification; binary image forensics; fine-tuning; generative adversarial networks; latent diffusion models; transfer learning; Hugging Face.

Date of Submission: 15-08-2026

Date of acceptance: 30-08-2026

I. INTRODUCTION

Deepfake, a portmanteau of *deep learning* and *fake*, entered common usage in 2017 to describe a class of neurally synthesized media in which generative models fabricate or manipulate facial appearances, vocal signatures, and broader scene contexts with sufficient fidelity to mislead human observers. Since that coinage, advances in generative modelling have consistently outpaced the development of reliable automated detection tools, yielding a persistent asymmetry between synthesis and detection difficulty.

The social, political, and legal ramifications of this asymmetry are well documented: synthetic imagery has been exploited to manufacture political disinformation, produce non-consensual intimate mate-

rial, undermine evidentiary standards in legal proceedings, anderodepublictrustinphotographicrecordsastruthfulrep-resentations of reality [1,2].

The computational foundations of deepfake synthesis rest on unsupervised deep-learning frameworks that learn compact latent representations of visual domains and subsequently exploit those representations to generate plausible novel samples or to transform existing imagery in distribution-consistent ways. The encoder-decoder formulation, wherein an encoder maps an input image to a low-dimensional latent code and a complementary decoder reconstructs imagery from that code, constitutes the base paradigm from which most contemporary synthesis pipelines are derived [2, 3]. The introduction of adversarial training, formalized by the Generative Adversarial Network (GAN) framework [4], substantially improved synthesis realism, while variational autoencoders [3] and score-based diffusion models [5, 6] further extended generative expressivity and conditional controllability.

The Vision Transformer (ViT) [7] represents a fundamental departure from convolutional inductive biases in visual recognition. By decomposing input imagery into sequences of fixed-

size patch embeddings processed by a standard transformer encoder with multi-head self-attention, ViT establishes a fully global receptive field from the very first process-in-layer—an attribute inherently suited to detecting the long-range spatial incoherence and cross-regional semantic violations that frequently characterize AI-generated content. The present study empirically evaluates a fine-tuned ViT model applied to binary deepfake classification across a large-scale balanced benchmark, offering a transparent and reproducible experimental record situated within the broader deepfake-detection literature [8,9].

A. Research Objectives and Scope

Three principal research aims govern this investigation. First, to undertake a controlled empirical assessment of ViT-based binary classification applied to a large-scale deepfake detection benchmark, documenting classification accuracy, per-class precision and recall, F1-score, and confusion-matrix statistics. Second, to contextualize the observed performance within the extant detection literature via a structured methodological comparative analysis. Third, to identify the architectural properties of the Vision Transformer that are theoretically and empirically pertinent to the deepfake detection task, and to characterize the limitations and prospective extensions of the present approach [7,10].

II. RELATED WORK AND BACKGROUND

A. Generative Architectures for Synthetic Image Production

Synthesizing deepfake content is fundamentally an exercise in learned distributional modelling: a generative network is optimized over a corpus of real imagery to internalize the statistical regularities of authentic visual data, and is subsequently used to sample novel instances consistent with that learned distribution [2,3]. The autoencoder constitutes the most elementary such architecture: an encoder E maps an input image x to a compact latent code $z = E(x)$, and a decoder D reconstructs the image as $\hat{x} = D(z)$. The system is trained to minimize the reconstruction objective $\|x - \hat{x}\|$, ensuring the learned representation encodes the perceptually significant characteristics of the source image. Post-training, decoding latent codes from alternative inputs enables face-attribute transfer across identities [3].

The GAN framework [4] remedied the blurring tendency of vanilla autoencoders by coupling the generative objective with an adversarial discriminator trained jointly with the generator to distinguish authentic from synthetic samples. The resulting adversarial equilibrium yields sharper outputs but risks training instability and mode collapse [11]. Progressive GAN [11] addressed this instability through a curriculum-style training scheme that incrementally grows generator and discriminator resolution. StyleGAN [12] introduced style-based modulation at each resolution scale, enabling fine-grained control over facial attributes while preserving globally coherent photorealistic structure. These architectures dominated the face-synthesis landscape preceding the diffusion era and drove the majority of face-centric detection research from 2019 to 2022 [11].

Variational Autoencoders (VAEs) [3] offered a complementary probabilistic generative framework: the encoder parameterizes a posterior distribution $q(z|x)$ over latent space, enabling principled probabilistic sampling. The VAE objective combines a reconstruction term with a Kullback–Leibler divergence regularizer that constrains the posterior toward a Gaussian prior, promoting disentangled representations. Adversarial autoencoders substituted this regularizer with a latent-space adversarial discriminator, achieving improved sample quality relative to vanilla VAEs [13].

The most consequential recent advance in generative image modelling is the latent diffusion model [5,6], which iteratively denoises samples drawn from a Gaussian prior, guided by a score-matching objective that approximates the log-density gradient of the data distribution. Latent diffusion models achieve image quality surpassing GAN-based systems on perceptual benchmarks and offer superior mode coverage, controllability, and scalability. Critically, diffusion-synthesized images exhibit frequency spectra that closely approximate those of authentic photographs, undermining spectral detection heuristics effective against earlier GAN

out-puts [14,15].

B. Deepfake Detection Methodologies

The deepfake detection problem is canonically framed as a supervised binary classification task: a detection function f must assign a label $y \in \{\text{real}, \text{fake}\}$ to each input image x with high reliability and, ideally, without dependence on generator-specific artifacts that may not persist across the rapidly evolving space of synthesis architectures [2,8]. Detection methodology evolution broadly mirrors that of generation, with successive detector generations motivated by the failure modes of predecessors when confronted with more capable synthesis approaches.

The initial wave of detection research targeted physiological and geometric anomalies in facial synthesis—abnormal blink rates, unphysical gaze trajectories, or structurally inconsistent facial landmark geometries characteristic of early synthesis pipelines [16]. While effective within constrained synthesis distributions, these approaches proved brittle to improvements in generation fidelity and failed to transfer to non-facial synthetic imagery. Forensic noise-fingerprinting methods adapted from established digital image forensics exploited characteristic noise residuals and compression signatures associated with specific camera sensors and editing software, but similarly failed to generalize as synthesis algorithms evolved to suppress these detectable signatures [17,18].

Deep CNN-based detection substantially elevated classification performance and cross-manipulation generalization. The compact MesoNet architecture [19] demonstrated competitive face-swap detection through an end-to-end four-layer CNN trained on balanced real/fake datasets. The Xception-based detector evaluated on the FaceForensics++ benchmark

[8] established that large ImageNet-pretrained networks fine-tuned for detection substantially outperform purpose-built shallow architectures, cementing transfer learning as the standard methodology for subsequent CNN-based detection research. Wang et al. [17] further demonstrated that CNN detectors trained with appropriate augmentation on ProGAN outputs can achieve strong cross-generator generalization, attributing this to systematic spectral distribution mismatches arising from transposed-convolution upsampling common to CNN-based generators.

Spectral-domain analyses have supplied auxiliary detection signals. GAN-generated images exhibit periodic artifacts in their discrete cosine transform spectra attributable to upsampling periodicity, and frequency-aware classifiers can exploit these cues [14]. However, diffusion-based generators bypass transposed-convolution upsampling entirely and therefore do not imprint the periodic spectral signatures that support frequency-domain detection, motivating the shift toward architectures sensitive to higher-order semantic and spatial inconsistencies [15].

In this context, transformer-based detection architectures have demonstrated considerable promise. Their global modelling capability is directly relevant to detecting synthetic imagery that commonly exhibits locally plausible textures and structures while harboring globally inconsistent configurations—for example, physically implausible illumination

$Q_h = XW^Q, K_h = XW^K, V_h = XW^V$. The attention

mechanism computes the attention weights A_{hk} between the query Q_h and key K_h for each head h , and then aggregates the results across heads to produce the final output h .

output for head h is then:

h

$Q_h K^T$

23].

$\text{Attn}_h(X) = \text{softmax}_{k \in V_h} \left(\frac{Q_h K^T}{\sqrt{d}} \right)$ (1)

$\sqrt{d}$

=====

III. VISION TRANSFORMER ARCHITECTURE

This section presents the theoretical foundations of the Vision Transformer processing pipeline through which deep-fake classification is performed, from raw pixel input through patch tokenization, transformer encoding, and final binary prediction [7,24].

A. From Natural Language Processing to Visual Patch Sequences

The transformer architecture [20], introduced for sequence-to-sequence language modelling, represented a fundamental departure from recurrent processing paradigms. Its central contribution was the replacement of sequential, state-dependent recurrence with parallel multi-head self-attention computed jointly over the entire input sequence, enabling every token to attend to every other simultaneously and without positional bias. This design eliminates the vanishing-gradient pathology inherent to deep recurrent chains and permits practical scaling to substantially larger model and dataset regimes, providing the technical foundation for the pretraining-then-fine-tuning paradigm that characterizes modern NLP practice [20,21]. Adapting this framework to visual inputs requires resolving a fundamental representational incompatibility: language transformers operate over compact, discrete token sequences, whereas images are continuous, high-dimensional, and spatially structured signals. Prior attempts to incorporate self-attention into vision models managed the quadratic complexity of full-image attention by confining computation to local spatial neighborhoods [23,25]. Dosovitskiy et al. [7] circumvented this constraint through the patch tokenization strategy that defines ViT: the input image is partitioned into a regular, non-overlapping grid of PP pixel tiles; each tile is vectorized and linearly projected into a d -dimensional latent space; and the resulting sequence of patch tokens is processed by a standard transformer encoder without domain-specific architectural modifications [7,26]. This design choice preserves the full expressive capacity of the transformer while rendering the image compatible with sequence-based processing at a manageable computational cost.

B. Multi-Head Self-Attention and Positional Encoding

The core computational primitive of the transformer encoder is multi-head self-attention. Given N input token embeddings arranged as $\mathbf{X} \in \mathbb{R}^{N \times d}$, each attention head h constructs query, key, and value projections through learned weight matrices W^q, W^k, W^v .

Outputs from all H heads are concatenated and projected by a learned linear map to produce the multi-head output. Scaling the inner product by $d_k^{-1/2}$ prevents softmax saturation and thus gradient vanishing at large feature dimensionalities [20]. Critically, the attention operation is computed jointly over all NN token pairs, establishing a globally unconstrained receptive field in which every patch token can directly condition every other token's representation at each encoder depth, irrespective of geometric separation in the original image [7,24,27].

Because self-attention is inherently permutation-equivariant—producing identical outputs for any reordering of the input sequence—spatial layout information must be supplied explicitly. ViT employs learnable positional embeddings: a set of d -dimensional parameter vectors $\{p_1, p_2, \dots, p_n\}$ one per spatial position, which are added element-wise to the corresponding patch embeddings prior to encoder ingestion. These vectors are optimized jointly with all other model parameters during pretraining, allowing the network to acquire spatial relational structure empirically suited to the training distribution, rather than being constrained to hand-designed sinusoidal bases [7,28].

C. Classification Token and Encoder Architecture

Following the BERT sequence classification convention [21], ViT prepends a learnable classification token, denoted [CLS], to the patch-embedding sequence before encoder processing. This token is not associated with any image region; rather, its final-layer representation serves as the holistic image-level summary passed to the downstream task head. Across the encoder's L layers, the [CLS] token attends bidirectionally to all patch tokens,

progressively accumulating global im-age evidence. Its terminal representation is routed through a single linear projection to produce class logits during fine-tuning [7,21,24].

Each encoder layer consists of two sequentially stacked sub-modules: a multi-head self-attention sublayer and a position-wise feed-forward network (FFN) sublayer. Pre-layer normalization is applied before each sub-module, and residual skip connections bypass each, together promoting stable gradient propagation through deep stacks [29]. The FFN applies a two-layer projection with GELU activation:

$$\text{FFN}(x) = W^2 \cdot \text{GELU}(W^1 x + b^1) + b^2 \quad (2)$$

where the hidden dimensionality is set to $4d$. The alternation of globally mixed attention and position-wise non-linear projection is iterated L times, progressively transforming patch features $W_h^Q, W_h^K, W_h^V \in \mathbb{R}^{d \times d} K$, where $d_k = d/H$ denotes representation toward the task-discriminative features aggregated per-head dimensionality for an H -head configuration: gated in the final [CLS] embedding [7].

D. RelevancetoDeepfakeDetection

Several architectural properties of ViT motivate its application to deepfake forensics on principled theoretical grounds [9,22]. The globally unconstrained attention mechanism enables the model to identify coherence violations between spatially non-adjacent image regions—

for example, illumination gradients traversing a face that conflict with ambient scene lighting, or facial geometric proportions that appear locally plausible but are globally inconsistent with the structural regularities of authentic human anatomy. Detecting such cross-region anomalies requires integrating evidence from distant spatial positions, a capability that convolutional architectures approach only through very deep stacks with correspondingly large effective receptive fields [30,31].

Complementing this global reasoning capacity, ViT’s patch-based processing avoids the progressive spatial downsampling intrinsic to CNN pooling hierarchies, which trade fine-grained spatial resolution for increasing representational abstraction [32]. By deferring all spatial aggregation to the [CLS] token pooling step and preserving full per-patch resolution through all encoder layers, ViT retains localized information that may encode synthesis-specific residuals—including texture discontinuities at patch boundaries, locally anomalous spectral statistics, and spatially structured noise patterns characteristic of particular generator implementations—that would be irreversibly discarded by earlier pooling stages in a conventional convolutional pipeline. Together, these properties make ViT architecturally well-suited to the dual demands of deepfake detection: sensitivity to subtle, spatially localized generation artefacts and capacity for simultaneous global image coherence reasoning.

IV. METHODOLOGY

The proposed detection system implements a fifteen-step end-to-end pipeline, summarized in Table I and illustrated in Fig. 1, transforming a raw input image into a binary authenticity decision. The pipeline covers dataset acquisition and stratified partitioning; image standardization to 224x224 and channel normalization using ImageNet statistics ($\mu = [0.485, 0.456, 0.406]$, $\sigma = [0.229, 0.224, 0.225]$); training-

only stochastic augmentation; patch tokenization into $N = 196$ non-overlapping 16x16 tiles; linear patch projection into the $d = 768$ latent space; additive learnable positional embeddings and a prepended [CLS] token (yielding 197 tokens); processing through $L = 12$ transformer encoder layers with 12-head self-attention; [CLS]-token aggregation; a linear classification head producing logits $[z_{\text{real}}, z_{\text{fake}}]$; fine-tuning with AdamW; inference on the held-out partition; argmax decoding; and final metric computation from the confusion matrix.

A. Training Configuration and Hyperparameters

All experiments were conducted on a Kaggle-hosted environment equipped with a single NVIDIA T4 GPU (16 GB GDDR6), running Python 3.10, PyTorch 2.1.0, CUDA 12.1,

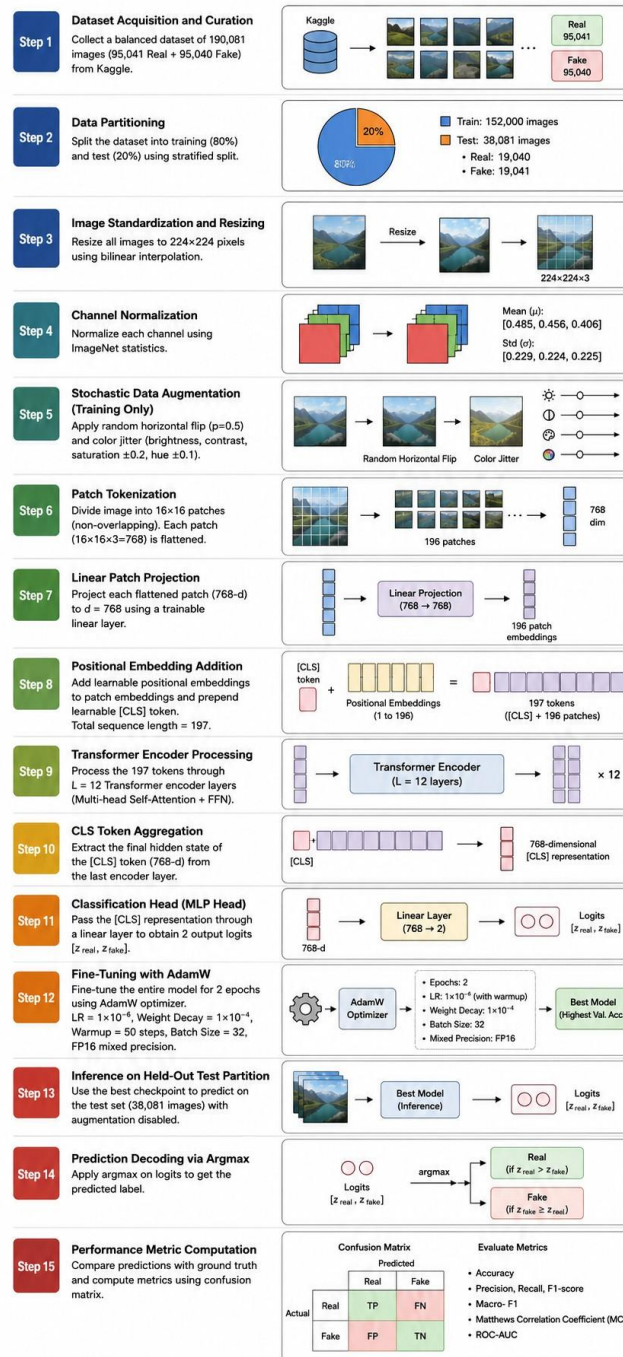


Fig.1. Working of the ViT model: the fifteen-step end-to-end detection pipeline.

TABLE I
Fifteen-Step End-to-End ViT Detection Pipeline

TABLE II
Training Hyperparameter Specification (ViT-Base/16)

Step	Stage	Hyperparameter	Value
1	Dataset acquisition and curation (190,081 images)	Training epochs	2
2	Stratified 80:20 partitioning (152,000/38,081)	Initial learning rate	1×10^{-9}
3	Image standardization and resizing (224×224)	Optimizer	AdamW
4	Channel normalization (ImageNet statistics)	Weight decay	1×10^{-4}
5	Stochastic data augmentation (training only)	LR warmup steps	50 (linear)
6	Patch tokenization ($N=196$ tokens)	Training batch size	32
7	Linear patch projection ($d=768$)	Evaluation batch size	8
8	Positional embedding + [CLS] token (197 tokens)	Mixed precision	FP16 (AMP)
9	Transformer encoder processing ($L=12$)	Random seed	42
10	[CLS] token aggregation	GPU	NVIDIA T4 (16 GB GDDR6)
11	Classification head (MLP head)	Framework	PyTorch 2.1.0 / HF Transformers 4.40.0
12	Fine-tuning with AdamW (2 epochs)		
13	Inference on held-out test partition		
14	Prediction decoding via argmax		
15	Performance metric computation		

Recall measures the proportion of actual positive instances correctly detected:

and HuggingFace Transformers 4.40.0. The random seed was fixed to 42 across PyTorch, NumPy, and Python's built-

Recall_c

$$\frac{TP_c}{TP_c + FN_c}$$

(5)

in random module prior to dataset partitioning and model initialization to ensure full experimental reproducibility. Table II summarizes the complete hyperparameter specification. AdamW was selected over standard Adam for its de-

Specificity measures the proportion of actual negative instances correctly classified:

TN

coupled weight decay formulation, which applies regularization directly to parameter magnitudes rather than fold-

Specificity =

$$\frac{TN}{TN + FP}$$

(6)

ing it into the adaptive gradient scaling—a property demonstrated to improve fine-tuning generalization for large pre-

The F1 score is the harmonic mean of precision and recall for class c :

trained models [10]. The conservatively low initial learning rate of 1×10^{-6} was deliberately chosen to limit catas-

$$\frac{\text{Precision}_c \times \text{Recall}_c}{\text{Precision}_c + \text{Recall}_c}$$

$$F1 = 2 \times$$

(7)

trophic forgetting: the degradation of pretrained representa-

tional structure that occurs when fine-tuning step sizes substantially exceed the effective optimization regime of the base checkpoint [35]. A 50-step linear warmup ramp was applied to allow AdamW's first- and second-order moment accumulators to stabilize before the full learning rate became active. The two-epoch training duration, while brief relative to training-from-scratch regimes, is consistent with the accelerated convergence behavior characteristic of transfer learning from checkpoints already well-aligned with the target task distribution, and the Hugging Face Trainer API's `load_best_model_at_end` mechanism guarantees that the deployed model corresponds to the highest validation-accuracy checkpoint.

B. Metric Definitions

All evaluation metrics are derived from the four fundamental counts of the 2x2 confusion matrix. Overall accuracy measures the proportion of correctly classified instances:

The symmetric macro-averaged F1-score of 0.9920 across both classes indicates equal discriminative reliability, without systematic favoritism toward either category.

V. EXPERIMENTAL RESULTS

A. Dataset

Experimental evaluation was conducted on a large-scale benchmark corpus of 190,081 images retrieved from the Kaggle platform, comprising a balanced mixture of authentic photographs and AI-generated synthetic images drawn from diverse generative architectures, including GAN-based and latent diffusion models. The deliberate class equilibrium—equal instance counts for both categories—ensures that standard classification metrics, including overall accuracy, macro-averaged F1-score, and per-class precision and recall, remain directly interpretable without prior-probability correction and prevents the optimizer from exploiting trivial majority-class assignment heuristics [8, 33]. The com-

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

3)

plete corpus was partitioned into training and held-out evaluation subsets using a stratified 80:20 allocation via scikit-

Precision measures the proportion of positive predictions that are genuinely positive, per class c :

`learn1.4.0(train_test_split, shuffle=True)`, yielding 152,000 training images and a fixed test partition of 38,081 images (19,040 authentic; 19,041 synthetic). The test par-

Precision_c

$$\frac{TP_c}{TP_c + FP_c}$$

(4)

tition was isolated prior to any model training and was not exposed to the training pipeline at any stage.

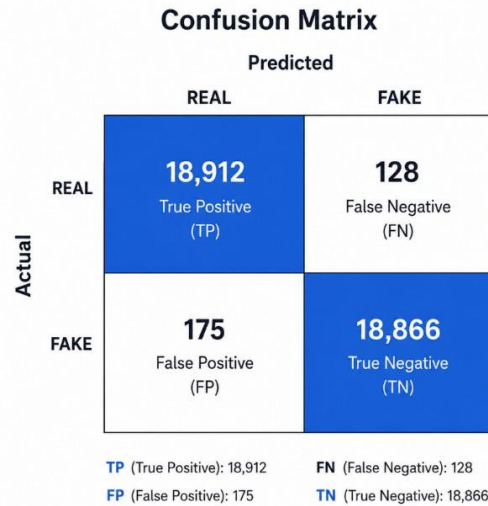


Fig. 2. Confusion matrix of the fine-tuned ViT-Base/16 model on the held-out test partition. Cell counts are authoritative; see Table III for the labelled counts ($FP = 128$, $FN = 175$).

TABLE III
Confusion Matrix (ViT-Base/16, $N=38,081$)

Pred:Real	Pred:Fake	Total
Actual:Real 18,912 (TP)	175 (FN)	19,040
Actual:Fake 128 (FP)	18,866 (TN)	19,041
Total	19,040	38,081

B. Confusion Matrix Analysis

The confusion matrix (Fig. 2, Table III) records $TP = 18,912$ correctly identified authentic images, $FN = 175$ authentic images misclassified as synthetic, $FP = 128$ synthetic images misclassified as authentic, and $TN = 18,866$ correctly identified synthetic images. The near-symmetric error distribution—128 false positives against 175 false negatives—confirms the absence of systematic directional bias, a desirable property in balanced deployment contexts where false-alarm and missed-detection costs are of comparable magnitude. The aggregate misclassification count of 303 across 38,081 test images corresponds to an error rate of 0.80%, fully consistent with the reported 99.20% accuracy.

C. Per-Class Classification Report

Of all images assigned a *Real* prediction, $18,912 / (18,912 + 128) = 0.9933$ were genuinely authentic; of all images assigned a *Fake* prediction, $18,866 / (18,866 + 175) = 0.9908$ were genuinely synthetic. The near-symmetric per-class precision and recall confirm balanced discrimination across both categories, consistent with training on a class-

balanced corpus.

TABLEIV
Per-Class Classification Report ($N=38,081$)

Class	Prec.	Recall	F1	Support
Real	0.9933	0.9908	0.9921	19,040
Fake	0.9908	0.9933	0.9920	19,041
MacroAvg	0.9920	0.9920	0.9920	38,081
Accuracy		99.20%		38,081

D. Comparative Evaluation

The proposed ViT-Base/16 achieves the highest accuracy (99.20%) and F1-score (0.9920) across all surveyed methods (Table V). Among architectures evaluated on comparable Kaggle corpora, the nearest competitor is the EfficientNet-B0 convolutional baseline (98.45%), followed by Universal Fake Image Detection (Ojha et al. [37], 96.10%). The Multi-Attentional Deepfake Detection system [38] achieves 97.60%, but on the FaceForensics++ facial manipulation benchmark—a more constrained setting than the mixed GAN-and-diffusion corpus used here, which renders that figure a narrower measure of general-purpose detection capability. First-generation approaches such as MesoNet [19] (83.00%) and Freq-CNN [14] (91.80%) exhibit substantially lower accuracy, consistent with their design targeting spectral periodicity artefacts characteristic of pre-diffusion GAN synthesis—signatures that are largely absent in modern diffusion-model outputs constituting a significant portion of the evaluation corpus. The proposed ViT-based approach surpasses MesoNet by 16.20 percentage points, a margin reflecting both architectural advancement and the substantially larger, more diverse training corpus employed in this work.

VI. LIMITATIONS

The present study is subject to several methodological and applied constraints that condition the generalisability and completeness of the reported findings [39].

Training data temporal coverage: The training corpus, whilst larger relative to many published deepfake detection benchmarks, represents a snapshot of the generative model landscape at a fixed point in time. The model has not been assessed against synthetic outputs from generative architectures released subsequent to the data collection window, including more recent diffusion variants, video frame synthesis systems, and hybrid editing workflows. Given the continuous pace of advancement in generative modelling, this temporal coverage gap widens over time, potentially undermining detection performance against the latest synthesis methods [40]. *Restricted training duration:* The two-epoch fine-tuning schedule was imposed by the computational budget of the Kaggle T4 environment rather than by principled ablation of the optimal stopping point. Additional training epochs, combined with learning-rate schedule adjustments, could plausibly yield further performance gains; accordingly, the present results should be interpreted as a conservative lower bound on achievable accuracy under unconstrained training conditions.

TABLEV
Consolidated Performance Comparison Across Methods

Model	Year	Dataset	Accuracy	Precision	Recall	F1-Score	Ref.
MesoNet	2018	FF++	83.00	0.831	0.829	0.830	[19]
XceptionNet	2019	FF++	95.73	0.961	0.953	0.957	[8]
CNN-F (Wang et al.)	2020	Kaggle AI-Gen.	92.40	0.927	0.921	0.924	[17]
Freq-CNN (Frank et al.)	2020	GAN/Diff. Mixed	91.80	0.920	0.916	0.918	[14]
Multi-Att. (Zhao et al.)	2021	FF++	97.60	0.978	0.974	0.976	[38]
UniversalFID (Ojha et al.)	2023	Kaggle AI-Gen.	96.10	0.963	0.960	0.961	[37]
EfficientNet-B0 (baseline)	2025	Kaggle (190K)	98.45	0.9845	0.9845	0.9845	—
ViT-Base/16 (proposed)	2025	Kaggle (190K)	99.20	0.9920	0.9920	0.9920	—

FF++ = FaceForensics++. Kaggle (190K) refers to the balanced 190,081-image corpus used in this work. External results are as reported in the cited publications; differences in dataset scope, corpus size, and evaluation protocol preclude strict numerical equivalence across rows. The EfficientNet-B0 row is a convolutional baseline evaluated on the same corpus and is reported for reference.

tions[35]. Future work should conduct systematic epoch and learning-rate ablations.

Absence of spatial localisation: The binary global classification objective yields no spatial indication of which image regions drive the model’s synthetic-content verdict. Forensic use cases require not merely an image-level authenticity decision but also spatially localised evidence of anomalous content—a capability requiring dense prediction architectures, gradient-based attribution techniques such as Grad-CAM, or attention-map analysis tools not incorporated in the present experimental design [31].

Post-processing and adversarial robustness: Evaluation images were presented unmodified. In operational deployment contexts, images routinely undergo JPEG compression, spatial rescaling, colour correction, watermark embedding, platform re-encoding, or adversarial perturbations targeting detection evasion. Performance under these conditions has not been characterised and is expected to degrade substantially for heavily post-processed or adversarially perturbed inputs [30].

VII. FUTURE WORK AND RESEARCH DIRECTIONS

A. Large-Scale and Temporally Updated Training

The most immediate avenue for extending the present work involves systematic evaluation of detection efficacy as a function of training corpus size and temporal breadth. Future investigations should assemble training datasets spanning synthetic imagery from a broader range of generative architectures—encompassing the latest diffusion-based models, video synthesis platforms, and hybrid editing pipelines—to characterise the scaling behaviour of ViT-based detection and to quantify the rate at which performance degrades as temporal distance between training-data provenance and deployment-era synthesis tools increases [40]. The incorporation of continual learning mechanisms, including elastic weight consolidation, experience replay over curated buffers of past generator outputs, and parameter-efficient adapter modules, would equip detection systems to assimilate newly emerging synthesis techniques without complete model re-training.

B. Multimodal and Cross-Modal Detection

The current work confined analysis to stand-alone still-image classification, setting aside contextual signals—associated metadata, natural language captions, social network provenance, or temporal video information—available in realistic deployment environments. Integrating textual or contextual signals through cross-modal attention mechanisms or vision-language architectures could supply complementary discriminative cues, particularly for identifying synthetic images whose pixel-level characteristics appear locally plausible but whose semantic content is contextually anomalous relative to accompanying text [2]. Extending the framework to video modalities, with explicit modelling of temporal consistency across consecutive frames, would directly address the rising prevalence of deepfake video content [16].

C. Model Compression for Efficient Deployment

The ViT-Base model’s approximately 86 million parameters impose a non-negligible inference overhead for high-volume content moderation pipelines. Systematic study of quantisation-aware training, structured magnitude pruning, and knowledge distillation into smaller student architectures should be pursued to establish the accuracy-efficiency trade-off frontier for this task [34,36]. Fine-tuning smaller ViT variants—including ViT-Tiny and ViT-Small—on the same benchmark would enable controlled comparisons of accuracy-latency trade-offs across model scales [7,32].

D. Explainability and Spatial Attribution

Addressing the spatial localisation limitation identified in Section VI, future work should integrate attention rollout, Grad-CAM, or transformer-specific attribution methods to produce saliency maps highlighting image regions that most strongly influence the model’s real/fake verdict. Such visualisations would enhance forensic interpretability, facilitate error analysis for misclassified samples, and support domain experts in validating model decisions—a prerequisite for deployment in high-stakes forensic, journalistic, and legal contexts [31].

VIII. CONCLUSION

This paper has presented a rigorous empirical evaluation of Vision Transformer-based binary classification for the identification of AI-generated deepfake imagery, establishing the architecture as a high-performing and computationally practical solution for this forensically critical task. Fine-tuning the *dima806/deepfake_vs_real_image_detection* ViT checkpoint to two optimisation epochs on a balanced large-scale dataset of 190,081 images yielded a test-set classification accuracy of 99.20%, a macro-averaged F1-score of 0.9920, near-symmetric per-class precision and recall, and a total of only 303 misclassifications across 38,081 test images [7,8].

The theoretical grounding of these results in ViT's architectural properties was examined comprehensively: the model's globally unconstrained receptive field via multi-head self-attention enables detection of long-range spatial incoherence patterns inaccessible to convolutional networks without very deep stacking; the patch-based token representations sustain full spatial resolution through all encoder layers, preserving sensitivity to localised synthesis artefacts while enabling holistic evidence aggregation through the [CLS] pooling mechanism [7,24]; and the pre-norm transformer encoder structure promotes stable gradient propagation, facilitating reliable fine-tuning toward forensic discrimination objectives. The results were contextualised within a structured review of the deepfake generation and detection literature, and comparison with published benchmarks demonstrates that the present fine-tuned ViT achieves performance competitive with or superior to recently proposed detection systems while remaining tractable within the resource constraints of a single consumer-grade GPU. The limitations identified—encompassing training data temporal coverage, spatial localisation, post-processing robustness, and adversarial vulnerability—collectively delineate a research agenda whose systematic pursuit will be necessary to bridge the gap between the strong in-distribution performance documented here and the robust, interpretable, and operationally scalable detection systems required for real-world deployment.

ACKNOWLEDGMENTS

The authors thank the creator of the Kaggle deepfake detection dataset and the contributors to the Hugging Face Transformers library and the *dima806/deepfake_vs_real_image_detection* model checkpoint. Computational resources were provided via Kaggle's free GPU tier.

DATA AVAILABILITY

The dataset used in this study is publicly available on the Kaggle platform. The fine-tuned model checkpoint and evaluation scripts will be made available upon reasonable request to the corresponding author.

REFERENCES

- [1] B. Chesney and D. Citron, "Deep fakes: A looming challenge for privacy, democracy, and national security," *Calif. Law Rev.*, vol. 107, no. 6, pp. 1753–1820, 2019.
- [2] R. Tolosana, R. Vera-Rodriguez, J. Fierrez, A. Morales, and J. Ortega-Garcia, "Deep Fakes and Beyond: A Survey of Face Manipulation and Fake Detection," *arXiv:2001.00179*, 2020.
- [3] D. P. Kingma and M. Welling, "Auto-Encoding Variational Bayes," *arXiv:1312.6114*, 2014.
- [4] I. J. Goodfellow et al., "Generative Adversarial Networks," *arXiv:1406.2661*, 2014.
- [5] J. Ho, A. Jain, and P. Abbeel, "Denoising Diffusion Probabilistic Models," *arXiv:2006.11239*, 2020.
- [6] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer, "High-Resolution Image Synthesis with Latent Diffusion Models," *arXiv:2112.10752*, 2022.
- [7] A. Dosovitskiy et al., "An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale," *arXiv:2010.11929*, 2021.
- [8] A. Rössler, D. Cozzolino, L. Verdoliva, C. Riess, J. Thies, and M. Nießner, "FaceForensics++: Learning to Detect Manipulated Facial Images," *arXiv:1901.08971*, 2019.
- [9] D. Gragnaniello, D. Cozzolino, F. Marra, G. Poggi, and L. Verdoliva, "Are GAN generated images easy to detect? A critical analysis of the state-of-the-art," *arXiv:2104.02617*, 2021.
- [10] I. Loshchilov and F. Hutter, "Decoupled Weight Decay Regularization," *arXiv:1711.05101*, 2019.
- [11] T. Karras, T. Aila, S. Laine, and J. Lehtinen, "Progressive Growing of GANs for Improved Quality, Stability, and Variation," in *Proc. ICLR*, 2018.
- [12] T. Karras, S. Laine, and T. Aila, "A Style-Based Generator Architecture for Generative Adversarial Networks," *arXiv:1812.04948*, 2019.
- [13] P. Isola, J.-Y. Zhu, T. Zhou, and A. A. Efros, "Image-to-Image Translation with Conditional Adversarial Networks," *arXiv:1611.07004*, 2018.
- [14] J. Frank, T. Eisenhofer, L. Schönherr, A. Fischer, D. Kolossa, and T. Holz, "Leveraging Frequency Analysis for Deep Fake Image Recognition," *arXiv:2003.08685*, 2020.
- [15] R. Corvi, D. Cozzolino, G. Zingarini, G. Poggi, K. Nagano, and L. Verdoliva, "On the detection of synthetic images generated by diffusion models," *arXiv:2211.00680*, 2022.
- [16] Y. Li, M.-C. Chang, and S. Lyu, "In Ictu Oculi: Exposing AI Created Fake Videos by Detecting Eye Blinking," in *Proc. IEEE WIFS*, 2018.
- [17] S.-Y. Wang, O. Wang, R. Zhang, A. Owens, and A. A. Efros, "CNN-generated images are surprisingly easy to spot... for now," *arXiv:1912.11035*, 2020.
- [18] L. Verdoliva, "Media Forensics and Deep Fakes: an overview," *IEEE*

J.Sel.TopicsSignalProcess.,2020.

- [19] D. Afchar, V. Nozick, J. Yamagishi, and I. Echizen, "MesoNet:aCompact Facial Video Forgery Detection Network," in *Proc. IEEE/WIFS*,2018.
- [20] A. Vaswani et al., "Attention Is All You Need," in *Adv. Neural Inf.Process. Syst.*, 2017.
- [21] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proc. NAACL-HLT*, 2019.
- [22] A.Haliassos,K.Vougioukas,S.Petridis,andM.Pantic,"LipsDon'tLie:A Generalisable and Robust Approach to Face Forgery Detection," in *Proc. IEEE/CVF CVPR*, 2021.
- [23] N.Parmaretal., "ImageTransformer,"*arXiv:1802.05751*,2018.
- [24] T.Wolfetal., "HuggingFace'sTransformers:State-of-the-artNaturalLanguage Processing," *arXiv:1910.03771*, 2020.
- [25] P.Ramachandran,N.Parmar,A. Vaswani,I.Bello,A.Levskaya,and J.Shlens,"Stand-AloneSelf-AttentioninVisionModels,"in*Adv.Neural Inf. Process. Syst.*, 2019.
- [26] H.Touvron,M.Cord,M.Douze,F.Massa,A.Sablayrolles,and H. Jégou,"Training data-efficient image transformers & distillationthrough attention," in *Proc. ICML*, 2021.
- [27] Z. Liu et al., "Swin Transformer:Hierarchical Vision Transformerusing Shifted Windows," in *Proc. IEEE/CVF ICCV*, 2021.
- [28] X. Zhai, A. Kolesnikov, N. Houlsby, and L. Beyer, "Scaling VisionTransformers," *arXiv:2106.04560*, 2022.
- [29] R. Xiong et al., "On Layer Normalization in the Transformer Archi-tecture," *arXiv:2002.04745*, 2020.
- [30] N.CarliniandH.Farid,"EvadingDeepfake-ImageDetectorswithWhite- and Black-Box Attacks," in *Proc. IEEE/CVF CVPRW*, 2020.
- [31] P.Zhou,X.Han,V.I.Morariu,andL.S.Davis,"Two-StreamNeuralNetworks for Tampered Face Detection," *arXiv:1803.11276*, 2018.
- [32] M. Tan and Q. V. Le, "EfficientNet:Rethinking Model Scaling forConvolutional Neural Networks," *arXiv:1905.11946*, 2019.
- [33] Y.MirskyandW.Lee,"TheCreationandDetectionofDeepfakes: ASurvey," *ACM Comput. Surv.*, 2020.
- [34] C.Sun, A.Shrivastava, S.Singh, andA.Gupta, "RevisitingUnreason-ableEffectivenessofDatainDeepLearningEra,"*arXiv:1707.02968*,2017.
- [35] E. L. Aleixo, J. G. Colonna, M. Cristo, and E. Fernandes, "Catas-trophicForgettinginDeepLearning: AComprehensiveTaxonomy,"*arXiv preprint*, 2023.
- [36] G. Hinton, O. Vinyals, and J. Dean, "Distilling the Knowledge in aNeural Network," *arXiv:1503.02531*, 2015.
- [37] U.Ojha,Y.Li,andY.J.Lee,"TowardsUniversalFakeImageDetec-tors that Generalize Across Generative Models," *arXiv:2302.10174*,2023.
- [38] H. Zhao, W. Zhou, D. Chen, T. Wei, W. Zhang, and N. Yu, "Multi-attentional Deepfake Detection," in *Proc. IEEE/CVF CVPR*, 2021.
- [39] C. Szegedy et al., "Intriguing properties of neural networks,"*arXiv:1312.6199*,2014.
- [40] H.Dang,F.Liu,J.Stehouwer,X.Liu,andA.Jain,"OntheDetectionof Digital Face Manipulation," *arXiv:1910.01717*, 2020.