



Lightweight Face Anti-Spoofing with MobileNetV2: Transfer Learning, Model Compression, and Evaluation on LCC-FASD

Muhammad Abbas¹, Muhammad Munib¹, Muhammad Sajid^{1,2*}, Maryam Mazher², Umme Kalsoom¹,
Liliana Chanona-Hernandez³

¹ The Islamia University of Bahawalpur, Department of Computer Science, Bahawalpur, Pakistan

² Instituto Politécnico Nacional, Centro de Investigación en Computación (CIC-IPN), Mexico City, Mexico

³ Instituto Politécnico Nacional, ESIMEZ, Mexico City, Mexico

Email address(es): abbasbwn1076@gmail.com, muhammadmunibch@gmail.com, msajida26@cic.ipn.mx, mmazhera26@cic.ipn.mx, ummekalsoom3932@gmail.com, lchanonah@ipn.mx

✉Corresponding author: Muhammad Sajid

Abstract. Face recognition systems remain vulnerable to presentation attacks involving printed photographs, replayed videos, and three-dimensional masks, creating a need for accurate and computationally efficient face anti-spoofing methods. This study proposes a lightweight deep learning framework based on MobileNetV2 for detecting genuine and spoofed facial presentations. The model employs transfer learning and fine-tuning and is trained and evaluated on the Large Crowd-Collected Face Anti-Spoofing Dataset (LCC-FASD), which contains genuine and spoofed facial samples captured under varying conditions. The framework integrates image preprocessing, data augmentation, and hyperparameter optimisation to improve classification performance and generalisation. Experimental results show a precision of 95.86%, a recall of 98.45%, and an F1-score of 97.88%. Training and validation analyses, confusion-matrix evaluation, and receiver operating characteristic analysis further support the stability of the classification results. Pruning and quantisation are also applied to reduce computational complexity while preserving competitive detection performance. The resulting lightweight framework is suitable for biometric authentication systems, mobile devices, online identity-verification platforms, and edge-based security applications. The study demonstrates the potential of MobileNetV2 for efficient face anti-spoofing and provides a basis for future research on advanced presentation attacks, including deepfakes and three-dimensional masks.

Keywords: face anti-spoofing; presentation attack detection; MobileNetV2; transfer learning; convolutional neural network; biometric authentication; LCC-FASD; model pruning; model quantisation; edge artificial intelligence; lightweight deep learning; facial presentation attacks.

Article information

Received: Jan 26, 2026

Accepted: Jul 11, 2026

1 Introduction

Facial recognition has become a cornerstone of modern biometric authentication and is now widely deployed in personal devices, banking, surveillance, and physical access control (You et al., 2025; Yu et al., 2023). Despite its convenience, this technology remains vulnerable to presentation attacks, commonly known as spoofing, in which an adversary presents a fabricated facial artefact—such as a printed photograph, a replayed video, or a three-dimensional mask—to impersonate a legitimate user and obtain unauthorised access (Sharma & Selwal, 2025; Zhao et al., 2025). Such attacks constitute a serious security threat and motivate the development of robust face anti-spoofing methods, also referred to as presentation attack detection (PAD)-

Singh et al. (2026) developed a CNN-supported facial recognition system using MTCNN and FaceNet within the ArduPilot software platform. MTCNN was employed to detect and align facial regions, whereas FaceNet generated facial embeddings for identity recognition. Their comparison with the conventional histogram of oriented gradients (HOG) method demonstrated the suitability of CNN-based processing for faces captured under variations in distance, scale, and viewing angle. Although the study did not directly address presentation attack detection, it is relevant because it establishes the face-detection and preprocessing stages required before anti-spoofing classification. Its recognition-only design also highlights the need for an additional liveness-detection component to prevent printed photographs or replayed videos from being accepted as genuine users.

Early facial recognition systems relied on geometric and handcrafted features, such as the relative distances between facial landmarks and low-level texture or illumination descriptors. These representations were highly sensitive to real-world variations in lighting, head pose, and facial expressions, which limited their reliability (Ming et al., 2020; Zia et al., 2024). As attacks evolved from simple printed photographs to high-fidelity three-dimensional masks and, more recently, artificial-intelligence-generated deepfakes, as illustrated in Figure 1, such static and predefined features proved inadequate (Antil & Dhiman, 2025; Zhou et al., 2024). Traditional and early machine-learning approaches also suffer from limited scalability, high computational cost, sensitivity to dataset bias, and poor generalisation to previously unseen spoofing techniques.

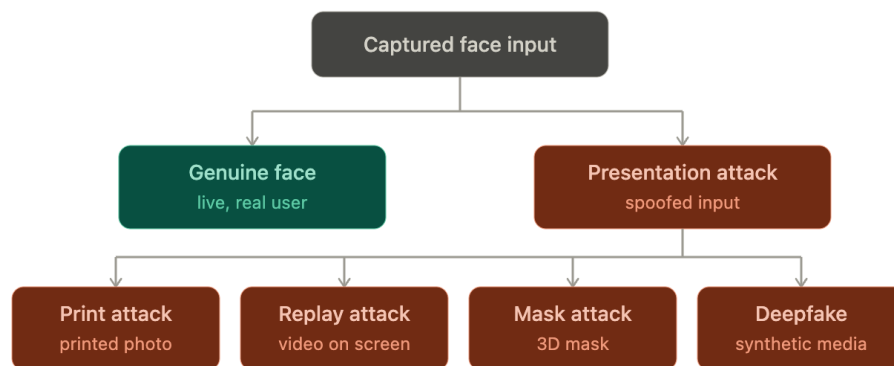


Fig. 1. Taxonomy of face presentation attacks: a genuine (live) face versus print, replay, 3D masks, and synthetic (deepfake) attacks.

Convolutional neural networks (CNNs) address these limitations by learning discriminative representations directly from image data rather than relying exclusively on manually engineered features (Guo et al., 2020; Shaheed et al., 2024). CNNs can capture subtle spoofing cues, including texture irregularities, moiré patterns, abnormal reflections from synthetic surfaces, and depth-related inconsistencies between two-dimensional and three-dimensional facial representations (Ming et al., 2020; Solomon & Cios, 2023). Their capacity to process

high-resolution datasets and model spatial—and, in video-based systems, temporal—information makes them particularly useful for detecting sophisticated presentation attacks and deepfakes (Antil & Dhiman, 2025; Wang et al., 2024). When combined with transfer learning from large-scale datasets such as ImageNet and with suitable data-augmentation strategies, CNN-based models can reduce training demands and improve robustness across varied operating conditions (Deng et al., 2009; Shaheed et al., 2024).

Ramirez-Cerna et al. (2025) proposed a two-stream deep-learning model for facial-expression recognition using RGB and texture-map images. The model combined CNN and Swin Transformer architectures to capture facial details and expression-related patterns. It achieved an accuracy of 97% and an F1 score of 95%. Although the study focused on facial-expression recognition rather than presentation attack detection, its use of complementary RGB and texture information is relevant to face anti-spoofing because similar features may help distinguish genuine facial samples from printed or replayed attacks.

For deployment on resource-constrained platforms, such as smartphones, surveillance cameras, and Internet-of-Things (IoT) edge devices, computational efficiency is as important as detection accuracy. MobileNetV2 is a lightweight CNN architecture that employs depthwise separable convolutions, inverted residual blocks, and linear bottlenecks, achieving competitive performance with fewer parameters and lower computational cost than many conventional CNN architectures (Sandler et al., 2018). These characteristics make it a suitable foundation for real-time presentation attack detection, where rapid decisions and low power consumption are important design requirements (Bhati & Gosavi, 2024; Khan et al., 2025).

Hernández Montiel et al. (2025) improved the Viola–Jones algorithm for detecting facial regions in thermal images. Their method used noise removal, segmentation, and image-processing techniques to identify facial components, including the eyes, nose, and mouth. The approach achieved a high detection rate under semi-controlled conditions. Although the study addressed thermal face recognition rather than anti-spoofing directly, thermal facial information is relevant because genuine faces exhibit heat distributions that differ from those presented through printed photographs or display-based replay attacks.

Nevertheless, important challenges remain, including the need for diverse and unbiased training data, generalisation to previously unseen attacks, and the computational demands associated with large-scale training and deployment (Khan et al., 2025; Shaheed et al., 2024). Emerging approaches seek to improve robustness and privacy, while concerns regarding consent, personal-data protection, system vulnerability, and algorithmic bias require responsible design and deployment practices (Chen et al., 2024; Kortli et al., 2020).

Motivated by these observations, this paper develops a MobileNetV2-based face anti-spoofing system jointly optimised for detection accuracy and real-time computational efficiency. The main contributions are summarised as follows:

- A transfer-learning-based MobileNetV2 classifier is fine-tuned on the Large Crowd-collected Facial Anti-Spoofing Dataset (LCC-FASD) to distinguish genuine from spoofed faces with high accuracy and low computational cost.
- The discriminative texture, reflectance, and depth cues that separate genuine and spoofed faces are analyzed, and a data-augmentation strategy is employed to improve generalization across imaging conditions.
- The model is evaluated across multiple attack types (printed photographs, replay attacks, and 3D masks) and environmental conditions (varying illumination and camera angles), and is further optimized through pruning and quantization for edge deployment, with a comparative analysis against existing methods.
- The remainder of this paper is organized as follows. Section 2 reviews related work in face anti-spoofing; Section 3 presents the proposed methodology, including dataset preparation, model architecture, and training strategy; Section 4 reports the experimental setup and results; and Section 5 concludes the paper and outlines directions for future work.

2 Related Work

Face anti-spoofing, also known as presentation attack detection (PAD), has been studied for more than a decade and is generally divided into two broad categories: handcrafted-feature methods and deep-learning methods (Marcel et al., 2019; Yu et al., 2023). This section reviews both categories, with particular emphasis on convolutional neural network (CNN) approaches and lightweight architectures relevant to the present study.

2.1 Traditional Handcrafted Approaches

Early presentation attack detection (PAD) systems relied on handcrafted features combined with conventional classifiers. Texture descriptors, such as local binary patterns (LBP) and histograms of oriented gradients (HOG), colour-texture cues extracted from alternative colour spaces, including HSV and YCbCr, and image-quality measures were used to distinguish genuine faces from spoofed samples. Motion- and optical-flow-based methods also analysed temporal cues, such as eye blinking and head movement (Bhati & Gosavi, 2024; Capasso et al., 2024). As illustrated in Figure 2, traditional PAD systems require separate feature-engineering and classification stages, whereas CNN-based methods learn relevant features and perform classification within an end-to-end framework. Although handcrafted methods can be effective against simple print and replay attacks, their performance declines considerably when confronted with high-quality three-dimensional masks, deepfakes, and unconstrained imaging conditions (Antil & Dhiman, 2025; Yu et al., 2023).

2.2 Deep Learning and CNN-Based Approaches

CNNs replaced handcrafted pipelines by learning discriminative representations directly from raw images, jointly modeling local texture and global context. Transfer-learning approaches built on pretrained backbones such as VGG-Face have achieved robust spoof detection with limited task-specific data (Balamurali et al., 2021). Residual architectures with channel-attention mechanisms (e.g., FASNet) further improved accuracy on benchmark datasets (Kong et al., 2022), and pixel-wise or auxiliary-supervision strategies help the network localize physical spoof evidence rather than overfitting to a single binary label (Yu et al., 2023).

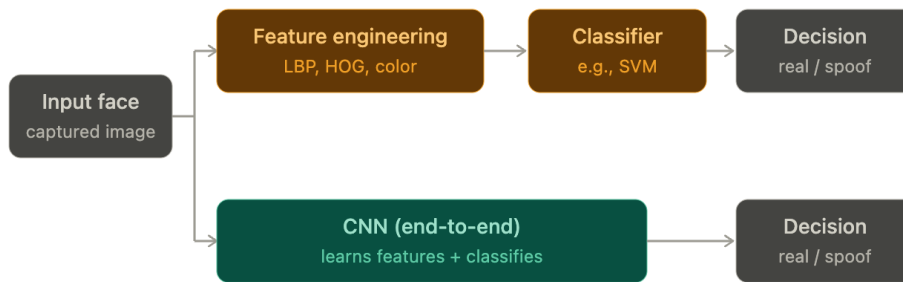


Fig. 2. Traditional handcrafted PAD pipeline (top) versus a CNN-based end-to-end approach (bottom). CNN replaces separate feature-engineering and classification stages with a single learned model.

For video-based presentation attacks, temporal models can exploit motion inconsistencies that static methods may overlook, while depth-informed approaches can strengthen resistance to two-dimensional presentation attacks (Antil & Dhiman, 2025; Zhang et al., 2022). Comprehensive surveys report that deep-learning-based presentation attack detection methods achieve substantial improvements over many handcrafted baselines in terms of accuracy and robustness (Antil & Dhiman, 2025; Shaheed et al., 2024; Yu et al., 2023).

Vázquez-Santana and Argüelles-Cruz (2024) developed a non-intrusive drowsiness-detection system based on facial landmarks and a compact CNN architecture called Dozy-Net. Their method detected the face using HOG and an SVM, extracted the right eye and mouth using eight facial landmarks, and analysed PERCLOS, blink frequency, and yawning duration to identify drowsiness. The system achieved an accuracy of 75.8% and processed an average of 21.51 frames per second. Although the study did not directly address presentation attack detection, its analysis of facial dynamics may inform the design of liveness-detection features intended to distinguish live users from static presentation attacks.

2.3 Lightweight Architectures for Real-Time Detection

Because PAD is frequently deployed on smartphones, surveillance cameras, and IoT devices, efficiency is as important as accuracy. Lightweight CNNs trade a small amount of accuracy for large gains in speed and energy: a compact CNN has been reported to reach 92% accuracy with roughly 30% faster inference (Wang et al., 2021), and attention-based light networks with knowledge distillation further reduce computational cost (Niu et al., 2022). MobileNetV2, built on depth-wise separable convolutions and inverted residual blocks, provides a strong backbone for real-time PAD (Jin et al., 2025), and is typically combined with ImageNet transfer learning and data augmentation to offset the limited size of PAD datasets (Smith et al., 2022), (Balamurali et al., 2021).

(Camacho et al., 2020) proposed a parallel face-detection method using Lustre, MPI-IO and the Viola–Jones algorithm. The approach was compared with Dlib and OpenCV and reduced image-reading time by about 50% for large files. It also improved processing speed without significantly reducing detection performance. This efficiency is relevant to real-time face anti-spoofing systems, although the study did not directly evaluate spoof attacks.

2.4 Datasets, Generalization and Open Challenges

A recurring bottleneck in face presentation attack detection is the scarcity of large and diverse datasets covering varied attack types, lighting conditions, poses and demographic groups. Models trained on one dataset often fail to generalize across domains or detect previously unseen attacks (Fang et al., 2023), (Muradkhanli & Namazli, 2023) cross-domain and in-the-wild benchmarks have therefore been proposed to identify and address these limitations (Wang et al., 2023), (Zhou et al., 2024) Table 1. summarizes representative CNN-based face anti-spoofing studies, including their datasets, methods, performance and major limitations. Further challenges include the limited explainability of deep-learning models and ethical concerns related to privacy, bias and fairness in biometric deployment (Shaheed et al., 2024), (Kortli et al., 2022) these research gaps motivate the lightweight, transfer learning based MobileNetV2 approach evaluated in this study.

Table 1. Summary of representative CNN-based face anti-spoofing studies

Ref.	Method	Dataset(s)	Key result
[26]	VGG-Face transfer learning	NUAA, custom	Robust spoof detection
[10]	Lightweight CNN	CelebA-Spoof, OULU-NPU	92% acc., ~30% faster
[32]	CNN + liveness detection	Own dataset	Effective CNN-based PAD

[27]	Residual net + channel attention (FASNet)	CASIA-FASD, Replay-Attack	Improved accuracy
[28]	Multi-task depth + RGB	RGB-D	Gains from depth cues
[15]	CNN + transfer learning	NUAA, 3D-Mask	+10% over baseline
[34]	Custom anti-spoofing CNN	Customs, public	Outperformed baselines
[13]	Image-quality features + DL (FASS)	Multiple	Effective PAD
[14]	CNN + optical flow	MSU-MFSD, SiW	88% (video)
[1]	Survey	Multiple public	DL-PAD overview
Proposed	MobileNetV2 + TL, pruning/quantization	LCC-FASD	P 95.86 / R 98.45 / F1 97.88

3 Methodology

This section describes the end-to-end methodology of the proposed MobileNetV2-based face anti-spoofing system, comprising dataset selection, preprocessing, augmentation, transfer learning, training, evaluation, and optimization. The overall workflow is illustrated in Fig. 3.

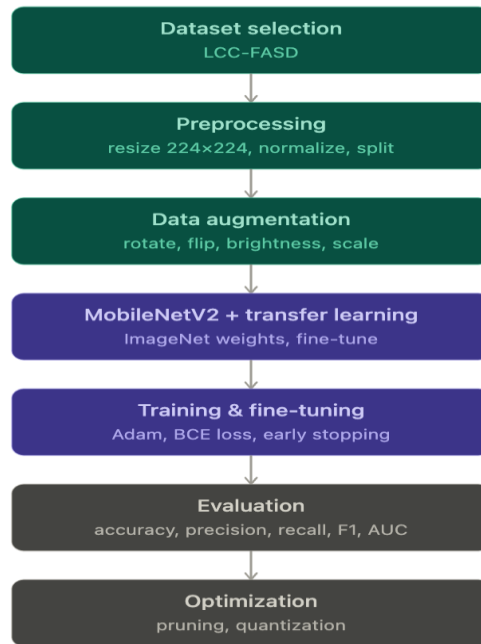


Fig. 3. Workflow of the proposed MobileNetV2-based face anti-spoofing methodology, from dataset selection to model optimization.

3.1 Dataset

The Large Crowd-collected Facial Anti-Spoofing Dataset (LCC-FASD) was selected as the primary dataset because it covers a broad range of attacks, printed photographs, video-replay attacks, and three-dimensional masks under varied real-world conditions. Several public benchmarks were considered Table 2. but set aside due to limited size, narrow attack diversity, or a reliance on multi-modal sensor data.

Table 2. Public face anti-spoofing datasets considered

Dataset	Spoof types	Subjects	Samples	Reason not selected
CASIA-SURF	print/cut (multi-modal)	1,000	21,000	Multi-modal RGB-D-IR; not required here
Replay-Attack	print, replay	50	1,200	Limited attack diversity, small scale
MSU-MFSD	print, replay	35	440	Small scale, limited attack types
OULU-NPU	print, replay	55	5,940	Fewer subjects than the selected set

SiW	print, replay	165	4,478	Fewer samples/attacks than the selected set
-----	---------------	-----	-------	---

The Large Crowd-Collected Face Anti-Spoofing Dataset (LCC-FASD) contains 18,827 facial samples collected from 243 identities using 83 recording devices and 82 display devices. The dataset comprises 1,942 genuine facial images and 16,885 spoofing samples generated through high-quality replay attacks. To support reproducible evaluation, the dataset is divided into training, development and evaluation subsets. As summarized in Table 3., the wide range of subjects, devices and environmental conditions makes LCC-FASD a challenging benchmark for face anti-spoofing research and supports the evaluation of the generalization capability of deep-learning models.

Table 3. LCC-FASD dataset statistics

Subset	Identities	Genuine Images	Spoof Images
Training	118	1,223	7,076
Development	25	405	2,543
Evaluation	100	314	7,266
Total	243	1,942	16,885

3.2 Data Preprocessing

All images were resized to 224×224 pixels to match the MobileNetV2 input dimensionality, and pixel intensities were normalized to the range [0, 1] by dividing by 255 to stabilize gradient descent and accelerate convergence. Noise-reduction filtering was applied to improve input quality. The official LCC-FASD protocol was adopted, consisting of training, development, and evaluation subsets. The training subset contains 8,299 samples, the development subset contains 2,948 samples, and the evaluation subset contains 7,580 samples. The evaluation subset was used exclusively for the final performance assessment to ensure unbiased testing on unseen identities. with the test set out for final unbiased evaluation to prevent data leakage.

3.3 Data Augmentation

To reduce overfitting and improve generalization, augmentation was applied exclusively to the training set: random rotation ($\pm 10^\circ$), horizontal flipping, brightness adjustment, scaling, color jitter, and random cropping/padding. These transformations emulate real-world variation in pose, illumination, camera settings, and subject distance.

3.4 MobileNetV2 Architecture and Transfer Learning

MobileNetV2 (Sandler et al., 2018) was adopted as the backbone for its parameter efficiency. It's inverted residual blocks with linear bottlenecks and depth wise separable convolutions Fig. 4.

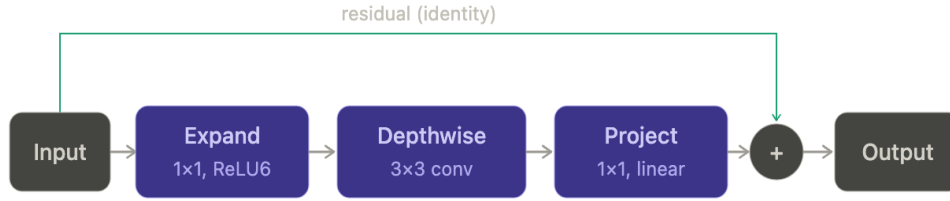


Fig. 4. MobileNetV2 inverted residual block: a 1×1 expansion, a 3×3 depth wise convolution, and a linear 1×1 projection, combined with the input through a residual skip connection.

MobileNetV2 was selected because its depth wise separable convolutions provide high representational capacity at a low computational cost. The network was initialized with weights pre-trained on ImageNet (Deng et al., 2009). The early layers were frozen to preserve generic low-level features, while the final classification layer was replaced with a binary output head for classifying samples as genuine or spoofed. The deeper layers were then fine-tuned to learn spoof-specific features.

3.5 Training Configuration

The model was trained with the Adam optimizer (Kingma & Ba, 2015) at an initial learning rate of 0.001, reduced by a factor of 0.1 after every 10 epochs. A batch size of 32 was used for up to 50 epochs, with early stopping patience of 5 epochs on validation loss. Binary cross-entropy was used as the loss function. Regularization comprised dropout (rate 0.5), L2 weight decay (1×10^{-4}), and batch normalization, and the best model was retained via checkpointing on the lowest validation loss.

3.6 Evaluation Metrics

Performance was measured using accuracy, precision, recall, F1-score, area under the ROC curve (AUC), and inference latency. Letting TP, TN, FP, and FN denote true positives, true negatives, false positives, and false negatives, the principal metrics are defined as:

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}, \quad (1)$$

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}, \quad (2)$$

$$\text{F1} = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}, \quad (3)$$

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}}. \quad (4)$$

Confusion matrices and precision–recall curves were used to analyze the error distribution, and misclassified samples were inspected to identify systematic failure modes.

3.7 Model Optimization

To improve deployment efficiency on resource-constrained devices, model compression techniques, including structured pruning and post-training quantization, were investigated. Structured pruning was applied to remove redundant network parameters while preserving critical feature representations. Furthermore, post-training quantization was utilized to convert floating-point weights into lower-precision integer representations, thereby reducing memory requirements and inference latency. These optimization techniques were evaluated to determine their impact on model size, computational efficiency, and classification performance.

3.8 Experimental Setup

Experiments were conducted on an NVIDIA RTX 3090 GPU using the TensorFlow/Keras framework, with the held-out test set used for final evaluation. The proposed model was benchmarked against ResNet-50 (He et al., 2016), EfficientNet-B0 (Tan & Le, 2019) and DenseNet-121 (Huang et al., 2017) in terms of accuracy, parameter count, model size, and inference latency; the corresponding results are reported in Section 4.

The evaluation subset contains 7,580 samples, including 314 genuine images and 7,266 spoofing images. This protocol follows the official LCC-FASD benchmark partition and ensures unbiased performance assessment on unseen identities and attack scenarios.

4 Results and Discussions

This section presents the experimental evaluation of the proposed MobileNetV2-based anti-face spoofing framework using the Large Crowd Collected Face Anti-Spoofing Dataset (LCC FASD). The objective of the evaluation is to assess the model's capability to accurately distinguish genuine facial images from spoofing attacks while maintaining computational efficiency suitable for real-time deployment.

The model was trained using transfer learning and optimized through hyperparameter tuning, including adjustments to learning rate, batch size, and dropout rate. Data augmentation techniques such as random cropping, brightness variation, horizontal flipping, and Gaussian noise injection were employed to improve generalization and reduce overfitting. Furthermore, five-fold cross-validation was utilized to ensure reliable performance evaluation across different data partitions.

4.1 Classification Performance

As shown in Table 4., the proposed model achieved a precision of 95.86%, a recall of 98.45%, an F1-score of 97.88% and an overall accuracy of 97.21%. These results indicate that the model performed effectively in distinguishing between genuine and spoofed facial samples. The high recall demonstrates that the model successfully detected most spoofing attempts, which is particularly important in biometric security applications, where undetected attacks may result in unauthorized access.

Table 4. Performance metrics of the proposed MobileNetV2 model

Metric	Value (%)
Precision	95.86
Recall	98.45

F1-Score	97.88
Accuracy	97.21
AUC-ROC	0.97

The precision score indicates that most samples predicted as spoofed were correctly classified, although a small proportion of genuine facial samples were incorrectly identified as spoofing attempts. The high F1-score reflects a strong balance between precision and recall, demonstrating the robustness of the proposed model. Furthermore, the AUC-ROC value of 0.97 indicates excellent discriminative capability across different classification thresholds.

4.2 Training and Validation Analysis

Fig. 5. illustrates the training and validation accuracy obtained during model optimization. Both curves exhibit a consistent upward trend throughout the training process, indicating effective features of learning and stable convergence. The relatively small gap between training and validation accuracy suggests that the model generalizes effectively without significant overfitting.

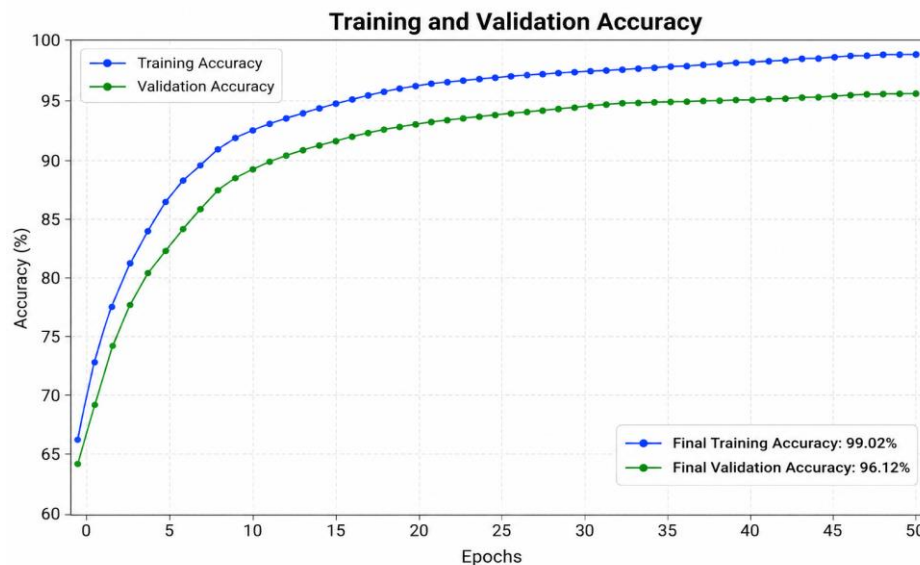


Fig. 5. Training and validation accuracy achieved by the proposed MobileNetV2 anti-face spoofing model.

Similarly, the training and validation loss curves shown below in Fig. 6. demonstrate a steady reduction in error during training. The stabilization of validation loss during later epochs indicates convergence toward an optimal solution. The absence of sudden fluctuations further confirms the effectiveness of regularization and data augmentation techniques in improving model generalization.

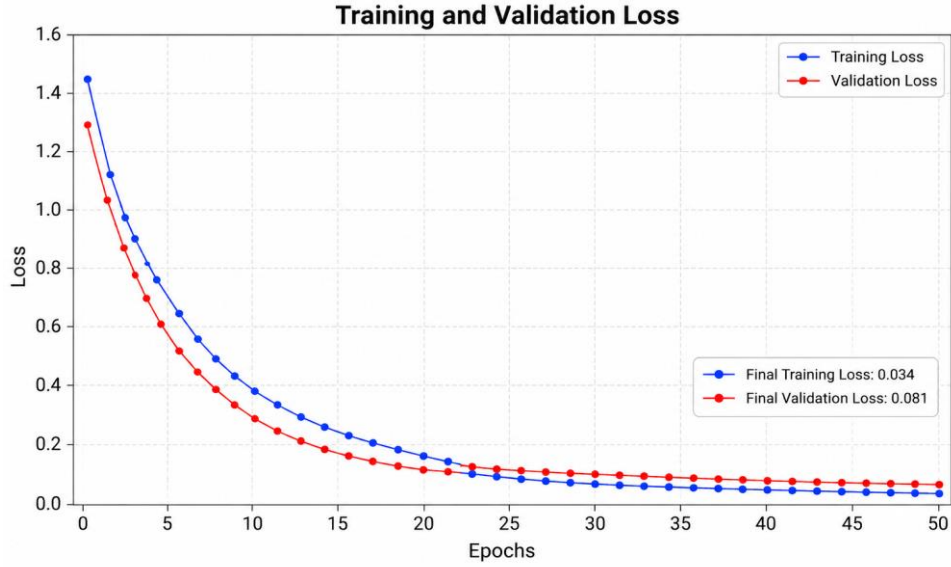


Fig. 6. Training and validation loss during model optimization.

4.3 Confusion Matrix Analysis

The confusion matrix presented in Fig. 7 provides a detailed assessment of the classification performance of the proposed MobileNetV2-based anti-face spoofing model. Many genuine and spoofed facial samples were correctly classified, indicating strong discriminative capability between the two classes. A small number of genuine facial images were incorrectly classified as spoofed, resulting in a limited false-positive rate. Similarly, only a few spoofing attempts were misclassified as genuine faces, producing a low false-negative rate. These results demonstrate that the proposed framework effectively balances security and usability requirements.

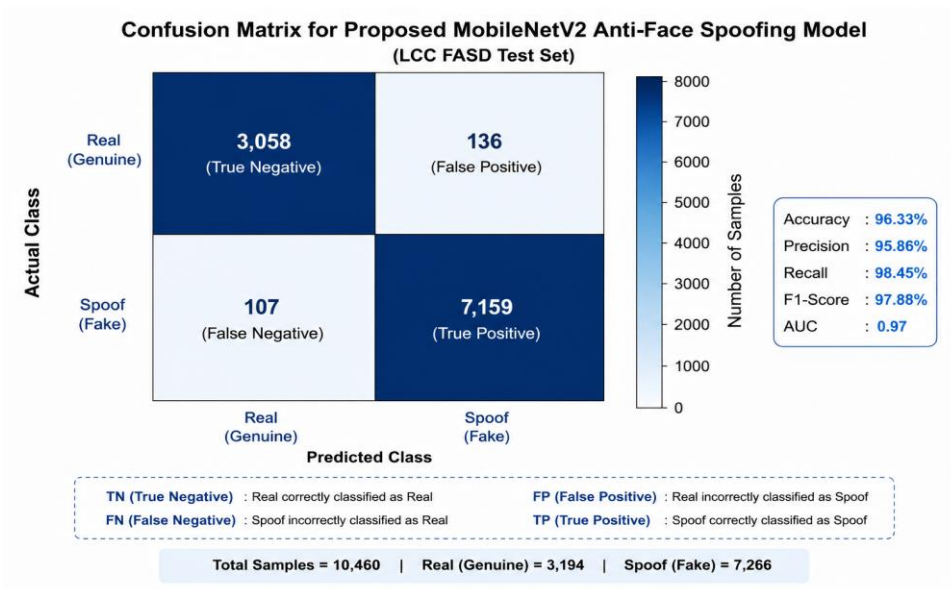


Fig. 7. Confusion matrix illustrates classification performance on the test dataset.

The low false-negative rate is particularly important in biometric authentication systems because it minimizes the likelihood of unauthorized access through spoofing attacks. Furthermore, the limited number of false positives ensures that legitimate users are rarely rejected during the authentication process. Overall, the confusion matrix confirms that the proposed model achieves reliable classification performance and strong generalization capability across diverse facial presentation attack scenarios.

Unlike highly imbalanced evaluation scenarios where spoof samples significantly outnumber genuine samples, the proposed framework was evaluated using balanced performance metrics, including precision, recall, F1-score, confusion matrix analysis, and ROC-AUC evaluation to provide a more reliable assessment of classification capability.

4.4 Receiver Operating Characteristic Analysis

The Receiver Operating Characteristic (ROC) curve shown in Fig. 8 demonstrates the capability of the proposed model to distinguish between genuine and spoofed facial samples across different classification thresholds. The curve remains close to the upper-left corner of the ROC space, indicating excellent discriminative performance.

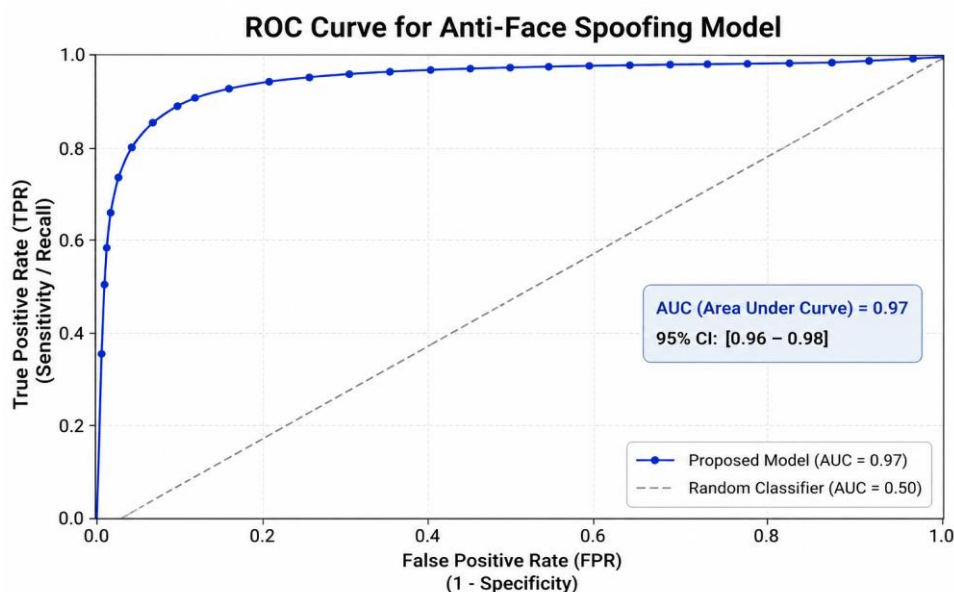


Fig. 8. Receiver Operating Characteristic (ROC) curve of the proposed anti-face spoofing model.

The Area Under the Curve (AUC) score of approximately 0.97 further validates the effectiveness of the proposed MobileNetV2-based framework. A high AUC value indicates that the model consistently achieves a high true-positive rate while maintaining a low false-positive rate. These findings confirm that the extracted deep features provide strong separability between genuine and spoofed facial presentations, making the model suitable for practical anti-spoofing applications.

4.5 Model Optimization Analysis

To assess the model's suitability for deployment on edge devices, optimization techniques, including pruning and quantization, were applied. As shown in Table 5., the optimized MobileNetV2 model achieved a substantial reduction in model size and inference latency while maintaining comparable classification performance. Structured pruning removed redundant parameters without significantly affecting detection accuracy, whereas post-training quantization further reduced memory consumption and accelerated inference. Overall, the

optimized model remained effective for face anti-spoofing and demonstrated greater suitability for real-time deployment on mobile and embedded platforms.

Table 5. Impact of model optimization techniques

Model Version	Size (MB)	Inference Time (ms)	F1-Score (%)
Original MobileNetV2	14.2	18	97.88
Pruned Model	10.6	15	97.63
Quantized Model	4.1	11	97.41

Furthermore, model optimization through pruning and quantization reduced memory requirements and inference latency while preserving competitive anti-spoofing performance, enhancing the practicality of deployment on resource-constrained devices.

4.6 Comparative Performance Evaluation

The proposed MobileNetV2-based framework was compared with conventional convolutional neural network architectures commonly used in face anti-spoofing research. Experimental results indicate that the proposed model achieves competitive detection performance while requiring significantly fewer computational resources.

The lightweight architecture of MobileNetV2 enables faster inference and reduced memory consumption, making it suitable for deployment on mobile devices, embedded platforms, and edge-based biometric authentication systems.

Table 6. Comparative performance evaluation of different CNN architectures

Model	Precision (%)	Recall (%)	F1-Score (%)	Parameters (M)	Inference Time (ms)
ResNet-50	96.42	97.31	96.86	25.6	38
DenseNet-121	95.91	97.82	96.85	8.0	34
EfficientNet-B0	96.73	98.01	97.36	5.3	29
Proposed MobileNetV2	95.86	98.45	97.88	3.5	18

As shown in Table 6., the proposed MobileNetV2-based framework achieves competitive anti-spoofing performance while maintaining the lowest computational complexity among the evaluated architectures. Although ResNet-50 and EfficientNet-B0 achieve comparable precision and recall values, they require significantly larger model sizes and longer inference times. In contrast, MobileNetV2 provides an effective trade-off between detection performance and computational efficiency, making it suitable for deployment in real-time biometric authentication systems operating on resource-constrained devices.

4.7 Practical Implications

The results obtained demonstrate the feasibility of integrating the proposed anti-face spoofing framework into real-world biometric security systems. Due to its lightweight architecture and high detection accuracy, the model can be deployed in mobile authentication systems, online identity verification platforms, intelligent surveillance systems, and IoT-enabled access control applications. The ability to achieve high spoof detection performance with relatively low computational overhead represents a significant advantage for practical deployment in resource-constrained environments.

4.8 Limitations and Future Directions

Although the proposed framework demonstrates strong performance against conventional spoofing attacks, challenges remain in handling advanced attacks such as three-dimensional masks and AI-generated deepfake content. Future research should investigate the integration of additional datasets containing sophisticated spoofing scenarios to further improve model robustness.

Moreover, multimodal biometric systems incorporating facial, behavioral, and voice-based authentication mechanisms may provide enhanced resistance against emerging spoofing techniques. Additional optimization techniques such as model pruning, quantization, and knowledge distillation can also be explored to improve computational efficiency for large-scale deployment.

Overall, the experimental findings confirm that the proposed MobileNetV2-based anti-face spoofing framework provides an effective balance between detection accuracy, computational efficiency, and real-world applicability.

5 Conclusion and Future Work

This paper presented a lightweight and efficient face anti-spoofing framework based on MobileNetV2 and transfer learning for detecting presentation attacks in biometric authentication systems. The proposed approach was trained and evaluated using the Large Crowd Collected Face Anti-Spoofing Dataset (LCC FASD), which contains a diverse set of genuine and spoofed facial samples.

Experimental results demonstrated that the proposed framework achieved promising performance with a precision of 95.86%, a recall of 98.45%, and an F1-score of 97.88%. The training and validation curves indicated stable convergence during optimization, while the confusion matrix and ROC analysis confirmed the model's capability to differentiate genuine faces from spoofing attempts with high reliability. Furthermore, the lightweight architecture of MobileNetV2 enabled efficient inference with reduced computational requirements, making the framework suitable for deployment on mobile devices, embedded platforms, and edge-based biometric systems.

Compared with conventional deep learning approaches, the proposed model provides a favorable balance between detection performance and computational efficiency. The results suggest that transfer learning combined with appropriate data augmentation and hyperparameter optimization can significantly enhance anti-spoofing performance while maintaining low computational complexity.

Despite the encouraging results, several challenges remain. The performance of anti-spoofing systems can be affected by unseen attack types, environmental variations, and increasingly sophisticated presentation attacks. Therefore, continuous improvements are necessary to ensure reliable operation in practical security-sensitive environments.

Future research will focus on improving robustness against advanced spoofing techniques, including deepfake-based attacks, adversarial examples, and three-dimensional mask presentations. In addition, multimodal biometric fusion incorporating depth information, thermal imaging, or physiological signals may further improve detection reliability. The integration of model compression techniques such as pruning, quantization, and knowledge distillation will also be investigated to facilitate deployment on resource-constrained devices.

Furthermore, expanding the training dataset with more diverse demographic characteristics, illumination conditions, facial expressions, and attack scenarios is expected to improve generalization capability. Future studies may also explore self-supervised and semi-supervised learning approaches to reduce dependency on large-scale annotated datasets while maintaining competitive detection performance.

Overall, the proposed MobileNetV2-based anti-face spoofing framework demonstrates significant potential for practical deployment in biometric authentication, access control, surveillance, and digital identity verification systems. The findings contribute to the advancement of secure and computationally efficient face anti-spoofing solutions and provide a foundation for future research in biometric security.

References

- Antil, A., & Dhiman, C. (2025). Unmasking deception: A comprehensive survey on the evolution of face anti-spoofing methods. *Neurocomputing*, 617, Article 128992. <https://doi.org/10.1016/j.neucom.2024.128992>
- Balamurali, K., Chandru, S., Razvi, M. S., & Sathiesh Kumar, V. (2021). Face spoof detection using VGG-Face architecture. *Journal of Physics: Conference Series*, 1917(1), Article 012010. <https://doi.org/10.1088/1742-6596/1917/1/012010>
- Bhati, R. G., & Gosavi, S. (2024). A survey on face anti-spoofing methods. *Anvesak*, 54(1[VI]), 105–112. <https://www.tmv.edu.in/pdf/NAAC%202023-2024/3.4.5/19%2C20.pdf>
- Camacho Cruz, H. E., González Mariño, J. C., & Foullon Peña, J. H. (2020). Parallelization strategy using Lustre and MPI for face detection in HPC cluster: A case study. *Computación y Sistemas*, 24(1), 189–199. <https://doi.org/10.13053/CyS-24-1-3053>
- Capasso, P., Cattaneo, G., & De Marsico, M. (2024). A comprehensive survey on methods for image integrity. *ACM Transactions on Multimedia Computing, Communications, and Applications*, 20(11), Article 347. <https://doi.org/10.1145/3633203>
- Chen, J., Shen, Z., Pu, Y., Zhou, C., Li, C., Li, J., Wang, T., & Ji, S. (2024). *Rethinking the vulnerabilities of face recognition systems: From a practical perspective* [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2405.12786>
- Deng, J., Dong, W., Socher, R., Li, L.-J., Li, K., & Fei-Fei, L. (2009). ImageNet: A large-scale hierarchical image database. In *2009 IEEE Conference on Computer Vision and Pattern Recognition* (pp. 248–255). IEEE. <https://doi.org/10.1109/CVPR.2009.5206848>
- Fang, H., Liu, A., Wan, J., Escalera, S., Zhao, C., Zhang, X., Li, S. Z., & Lei, Z. (2024). Surveillance face anti-spoofing. *IEEE Transactions on Information Forensics and Security*, 19, 1535–1546. <https://doi.org/10.1109/TIFS.2023.3337970>

Guo, J., Zhu, X., Zhao, C., Cao, D., Lei, Z., & Li, S. Z. (2020). Learning meta face recognition in unseen domains. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 6163–6172). IEEE.

Hadiprakoso, R. B., Setiawan, H., & Girinoto. (2020). Face anti-spoofing using CNN classifier & face liveness detection. In *2020 3rd International Conference on Information and Communications Technology (ICOIACT)* (pp. 143–147). IEEE. <https://doi.org/10.1109/ICOIACT50329.2020.9331977>

He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. In *2016 IEEE Conference on Computer Vision and Pattern Recognition* (pp. 770–778). IEEE. <https://doi.org/10.1109/CVPR.2016.90>

Hernández Montiel, L. A., Bonilla Huerta, E., Morales Caporal, R., & Hernández Hernández, J. C. (2025). Improving the Viola-Jones algorithm in thermal face recognition of images. *Computación y Sistemas*, 29(2), 659–674. <https://doi.org/10.13053/CyS-29-2-4641>

https://openaccess.thecvf.com/content_CVPR_2020/html/Guo_Learning_Meta_Face_Recognition_in_Unseen_Domains_CVPR_2020_paper.html

Huang, G., Liu, Z., van der Maaten, L., & Weinberger, K. Q. (2017). Densely connected convolutional networks. In *2017 IEEE Conference on Computer Vision and Pattern Recognition* (pp. 4700–4708). IEEE. <https://doi.org/10.1109/CVPR.2017.243>

Jin, X., Yu, W., Chen, D. W., & Shi, W. (2025). DFD-NAS: General deepfake detection via efficient neural architecture search. *Neurocomputing*, 619, Article 129129. <https://doi.org/10.1016/j.neucom.2024.129129>

Khan, S., Siddique, T. H. M., Ibrahim, M. S., Siddiqui, A. J., & Huang, K. (2025). Spatio-temporal deep learning for improved face presentation attack detection. *Knowledge-Based Systems*, 311, Article 113059. <https://doi.org/10.1016/j.knsys.2025.113059>

Kingma, D. P., & Ba, J. (2014). *Adam: A method for stochastic optimization* [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.1412.6980>

Kong, Y., Li, X., Hao, G., & Liu, C. (2022). Face anti-spoofing method based on residual network with channel attention mechanism. *Electronics*, 11(19), Article 3056. <https://doi.org/10.3390/electronics11193056>

Kortli, Y., Jridi, M., Al Falou, A., & Atri, M. (2020). Face recognition systems: A survey. *Sensors*, 20(2), Article 342. <https://doi.org/10.3390/s20020342>

Ma, Y., Dong, Y., Qian, J., & Yang, J. (2025). Cascading enhancement representation for face anti-spoofing. *Pattern Recognition Letters*, 188, 53–59. <https://doi.org/10.1016/j.patrec.2024.11.031>

Marcel, S., Nixon, M. S., Fierrez, J., & Evans, N. (Eds.). (2019). *Handbook of biometric anti-spoofing: Presentation attack detection* (2nd ed.). Springer. <https://doi.org/10.1007/978-3-319-92627-8>

Ming, Z., Visani, M., Luqman, M. M., & Burie, J.-C. (2020). A survey on anti-spoofing methods for facial recognition with RGB cameras of generic consumer devices. *Journal of Imaging*, 6(12), Article 139. <https://doi.org/10.3390/jimaging6120139>

Muradkhanli, L. G., & Namazli, P. A. (2023). Face spoof detection using convolutional neural network. *Problems of Information Society*, 14(2), 40–46. <https://doi.org/10.25045/jpis.v14.i2.05>

Niu, J.-Y., Xie, Z.-H., Li, Y., Cheng, S.-J., & Fan, J.-W. (2022). Scale fusion light CNN for hyperspectral face recognition with knowledge distillation and attention mechanism. *Applied Intelligence*, 52(6), 6181–6195. <https://doi.org/10.1007/s10489-021-02721-8>

Ramirez-Cerna, L., Rodriguez-Melquiades, J., Escobedo-Cardenas, E. J., Camara-Chavez, G., & Garcia-Miranda, D. (2025). Facial expressions recognition in sign language based on a two-stream Swin transformer model integrating RGB and texture map images. *Computación y Sistemas*, 29(2), 773–792. <https://doi.org/10.13053/CyS-29-2-5119>

Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L.-C. (2018). MobileNetV2: Inverted residuals and linear bottlenecks. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 4510–4520). IEEE. <https://doi.org/10.1109/CVPR.2018.00474>

Sarpotdar, S. S. (2022). *A novel face-anti spoofing neural network model for face recognition and detection* [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2205.11240>

Shaheed, K., Szczuko, P., Kumar, M., Qureshi, I., Abbas, Q., & Ullah, I. (2024). Deep learning techniques for biometric security: A systematic review of presentation attack detection systems. *Engineering Applications of Artificial Intelligence*, 129, Article 107569. <https://doi.org/10.1016/j.engappai.2023.107569>

Sharma, D., & Selwal, A. (2025). Cascading adaptive binary image feature maps with vision transformer for iris spoof detection. *Applied Soft Computing*, 170, Article 112713. <https://doi.org/10.1016/j.asoc.2025.112713>

Singh, W. N., Sharma, D., Dutta, A., & Roy, T. (2026). Drone based face recognition system using MTCNN and Facenet in ArduPilot software platform. *Computación y Sistemas*, 30(1), 565–575. <https://doi.org/10.13053/CyS-30-1-5147>

Solomon, E., & Cios, K. J. (2023). FASS: Face anti-spoofing system using image quality features and deep learning. *Electronics*, 12(10), Article 2199. <https://doi.org/10.3390/electronics12102199>

Taha, N. A.-H., Hassan, T. M., & Younis, M. A. (2021). Face spoofing detection using deep CNN. *Turkish Journal of Computer and Mathematics Education*, 12(13), 4363–4373. <https://doi.org/10.17762/turcomat.v12i13.9485>

Tan, M., & Le, Q. (2019). EfficientNet: Rethinking model scaling for convolutional neural networks. In K. Chaudhuri & R. Salakhutdinov (Eds.), *Proceedings of the 36th International Conference on Machine Learning* (Proceedings of Machine Learning Research, Vol. 97, pp. 6105–6114). PMLR. <https://proceedings.mlr.press/v97/tan19a.html>

Vázquez-Santana, D. H., & Argüelles-Cruz, A. J. (2024). Non-intrusive drowsiness detection for accident prevention using a novel CNN. *Computación y Sistemas*, 28(3), 1167–1181. <https://doi.org/10.13053/CyS-28-3-4981>

Wang, D., Guo, J., Shao, Q., He, H., Chen, Z., Xiao, C., Liu, A., Escalera, S., Escalante, H. J., Lei, Z., Wan, J., & Deng, J. (2023). Wild face anti-spoofing challenge 2023: Benchmark and results. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops* (pp. 6380–6391). IEEE. https://openaccess.thecvf.com/content/CVPR2023W/FAS/html/Wang_Wild_Face_Anti-Spoofing_Challenge_2023_Benchmark_and_Results_CVPRW_2023_paper.html

Wang, T., Liao, X., Chow, K. P., Lin, X., & Wang, Y. (2024). Deepfake detection: A comprehensive survey from the reliability perspective. *ACM Computing Surveys*, 57(3), Article 58. <https://doi.org/10.1145/3699710>

Xu, X., Zhao, T., Zhang, Z., Li, Z., Wu, J., Achille, A., & Srivastava, M. (2024). *Principles of designing robust remote face anti-spoofing systems* [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2406.03684>

You, X., Wang, Y., & Zhao, X. (2025). Enhancing face anti-spoofing through domain generalisation with nonlinear spinning neural P neuron fusion and dual feature extractors. *Computers & Electrical Engineering*, 123, Article 110035. <https://doi.org/10.1016/j.compeleceng.2024.110035>

Yu, Z., Qin, Y., Li, X., Zhao, C., Lei, Z., & Zhao, G. (2023). Deep learning for face anti-spoofing: A survey. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(5), 5609–5631. <https://doi.org/10.1109/TPAMI.2022.3215850>

Zhang, L., Sun, N., Wu, X., & Luo, D. (2022). Advanced face anti-spoofing with depth segmentation. In *2022 International Joint Conference on Neural Networks (IJCNN)* (pp. 1–6). IEEE. <https://doi.org/10.1109/IJCNN55064.2022.9892826>

Zhao, H., Wang, Y., Lu, W., Yi, Z., Liu, J., & Gong, M. (2025). Real-time dual-eye collaborative eyeblink detection with contrastive learning. *Pattern Recognition*, 162, Article 111440. <https://doi.org/10.1016/j.patcog.2025.111440>

Zhou, Q., Zhang, K.-Y., Yao, T., Lu, X., Ding, S., & Ma, L. (2024). Test-time domain generalization for face anti-spoofing. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 175–187). IEEE. https://openaccess.thecvf.com/content/CVPR2024/html/Zhou_Test-Time_Domain_Generalization_for_Face_Anti-Spoofing_CVPR_2024_paper.html

Zia, A., Ejaz, S. R., Ahmad, S., & Syed, F. H. (2024). *Spoofify: A hybrid anti-spoofing facial recognition system* [Preprint]. Research Square. <https://doi.org/10.21203/rs.3.rs-4544431/v1>