

Identity-Aware Lightweight MobileNetV2 with Distillation and Optuna for Face Spoofing Detection

Alif Sahputra^{1,*}, Alva Hendi Muhammad¹

¹ Universitas AMIKOM Yogyakarta, Indonesia

* Corresponding author: Alif Sahputra, Universitas AMIKOM Yogyakarta, Indonesia

✉ alifsahputra@students.amikom.ac.id

Copyright: © 2026 by the authors

Received: 13 May 2026 | Revised: 20 May 2026 | Accepted: 11 June 2026 | Published: 9 July 2026

Abstract

Presentation attacks, commonly known as face spoofing, remain a major security challenge in facial recognition-based authentication systems because forged media such as printed photos and replayed videos can deceive biometric verification. Lightweight CNN models such as MobileNetV2 are suitable for practical implementation, but their limited representational capacity may affect their ability to capture subtle spoofing cues and generalize to unseen identities. Previous evaluations may also produce inflated performance estimates when images from the same identity appear across training and testing sets. This study evaluates Knowledge Distillation and Optuna-based hyperparameter tuning on MobileNetV2 for lightweight face anti-spoofing under an identity-aware evaluation protocol. The novelty lies in an identity-aware comparison between representation enhancement through EfficientNet-B0-based Knowledge Distillation and optimization-based improvement through Optuna. A total of 60,000 CelebA-Spoof images were divided using an 80:10:10 subject-disjoint split, and four scenarios were compared. The baseline MobileNetV2 achieved the best overall balance, with an accuracy of 0.9943, F1-score of 0.9958, and ACER of 0.0058. Meanwhile, Knowledge Distillation obtained the lowest APCER of 0.0035, indicating fewer spoof samples were incorrectly accepted as live under the identity-aware evaluation setting.

Keywords: face anti-spoofing; hyperparameter optimization; knowledge distillation; mobilenetv2; optuna.

To cite this article: Sahputra, A., & Muhammad, A. H. (2026). Identity-Aware Lightweight MobileNetV2 with Distillation and Optuna for Face Spoofing Detection. *Edumatic: Jurnal Pendidikan Informatika*, 10(2), 320–329. <https://doi.org/10.29408/edumatic.v10i2.34863>

INTRODUCTION

Facial recognition technology has been widely adopted as an authentication mechanism in digital systems, including access control, attendance management, financial services, and mobile-based applications (Mun & Lee, 2022). Despite its practicality, this technology remains vulnerable to presentation attacks or face spoofing, where forged media such as printed photos, replayed videos, or facial images displayed on screens are used to deceive the system. Xing et al. (2025) and Yu et al. (2023) emphasized that face spoofing remains a critical threat in facial authentication systems because failure to detect presentation attacks may lead to false acceptance. In real authentication scenarios, this condition may allow unauthorized users to gain access, manipulate attendance records, or misuse another person's identity. Therefore, face anti-spoofing should not be viewed only as a classification task, but also as a security requirement for biometric authentication systems.



The need for robust anti-spoofing becomes more challenging when facial authentication is deployed on mobile devices, attendance cameras, and edge-based systems. These environments encourage the use of compact models with lower computational complexity, although full deployment performance still requires device-level evaluation (Xiao et al., 2024). Cai et al. (2024) showed that high-capacity architectures such as Vision Transformer can provide strong feature representation for face anti-spoofing, but their computational requirements may limit practical deployment. In contrast, lightweight CNNs such as MobileNetV2 are more suitable for resource-constrained devices. MobileNetV2 has also been widely used as a lightweight transfer learning architecture in image classification tasks (Gulzar, 2023). However, its compact structure may limit its representational capacity, making it less reliable when facing subtle spoofing cues, different image qualities, complex attack variations, and unseen subject identities.

Yu et al. (2023) reviewed deep learning-based face anti-spoofing methods, including CNN-based models, domain generalization, frequency-based representation, and transformer-based approaches. Although these methods have improved detection performance, several limitations remain. Many studies still rely on high-capacity models, while lightweight CNNs may suffer from weaker representation and generalization capability (Cai et al., 2024). In addition, some evaluations may produce overly optimistic results when random data splitting allows images from the same identity to appear in both training and testing sets. Under this condition, the model may partially learn identity-specific facial patterns rather than spoof-related cues. Wang et al. (2023) emphasized that generalization is an important issue in face anti-spoofing because models must remain reliable when tested on unseen conditions or identities. Therefore, identity-aware evaluation is needed to reduce identity leakage and to assess whether a model can generalize to identities that were not seen during training.

Knowledge Distillation and hyperparameter optimization represent two different strategies for improving lightweight models. Guo et al. (2024) described Knowledge Distillation as a strategy for transferring knowledge from a stronger teacher model to a smaller student model. In this study, EfficientNet-B0 is used as the teacher model because of its stronger representational capacity, while MobileNetV2 is used as the lightweight student model. Kong et al. (2024) showed that Knowledge Distillation can support face anti-spoofing by transferring richer representational knowledge to the student model. Li et al. (2022) also demonstrated that Knowledge Distillation can be applied to face presentation attack detection to support knowledge transfer between teacher and student models. Meanwhile, Lai et al. (2024) demonstrated that Optuna-based hyperparameter optimization can improve model training by searching for better training configurations. However, it remains necessary to examine whether representation enhancement through Knowledge Distillation, optimization-based improvement through Optuna, or their combination provides a more favorable effect on lightweight CNN-based face anti-spoofing under the same evaluation setting.

Building on the limitations of previous studies, this research examines three aspects that remain insufficiently explored in lightweight face anti-spoofing. First, direct evidence comparing Knowledge Distillation and Optuna-based hyperparameter optimization for improving MobileNetV2 is still limited (Kong et al., 2024; Lai et al., 2024). Second, the ability of MobileNetV2 to recognize spoofing patterns from unseen identities has not been widely examined using an identity-aware evaluation setting. Third, evaluation results may become overly optimistic when subject identities overlap between training and testing data, making identity-aware evaluation more suitable for assessing generalization in face anti-spoofing (Wang et al., 2023; Zheng et al., 2024). The novelty of this study lies in an identity-aware comparative evaluation of representation enhancement through EfficientNet-B0-based Knowledge Distillation and optimization-based improvement through Optuna for MobileNetV2-based lightweight face anti-spoofing.

This study evaluates the effectiveness of Knowledge Distillation and Optuna-based hyperparameter optimization in improving MobileNetV2 for face anti-spoofing detection. Four scenarios were examined: baseline MobileNetV2, MobileNetV2 with Optuna, MobileNetV2 with Knowledge Distillation from an EfficientNet-B0 teacher initialized with ImageNet weights and fine-tuned on live-spoof data, and MobileNetV2 with both approaches. Using the CelebA-Spoof dataset with an identity-aware split strategy, this study contributes to unseen-identity generalization evaluation, a controlled comparison between representation enhancement and hyperparameter optimization, and the analysis of lightweight face anti-spoofing models for biometric authentication systems.

METHOD

This study was conducted as a computational experiment to evaluate the effect of Knowledge Distillation and Optuna-based hyperparameter tuning on MobileNetV2 for face anti-spoofing. Four controlled scenarios were examined: baseline MobileNetV2, MobileNetV2 with Optuna, MobileNetV2 with Knowledge Distillation, and MobileNetV2 with both Knowledge Distillation and Optuna. The experiment used CelebA-Spoof, a public face anti-spoofing dataset containing live and spoof face images with variations in identity, illumination, acquisition conditions, and attack media (Gong et al., 2024). A controlled subset of 60,000 images was selected using a fixed random seed while preserving live-spoof distribution and subject-level separation. The data were divided into training, validation, and testing sets using an 80:10:10 ratio, resulting in 48,000, 6,000, and 6,000 images, respectively. The split was performed at the subject level to ensure that the same identity did not appear across training, validation, and testing subsets. The split was performed based on subject identity to reduce identity leakage and support realistic generalization evaluation.

During preprocessing, the training images were processed using random resized crop, horizontal flipping, low-intensity color jitter, tensor conversion, and ImageNet normalization. Meanwhile, the validation and test images were resized to 224×224 pixels, converted into tensors, and normalized using ImageNet mean and standard deviation values. MobileNetV2 was used as the student model, while EfficientNet-B0 was used as the teacher model. The teacher was initialized with ImageNet-pretrained weights, fine-tuned on the live-spoof training data, and the best checkpoint based on validation F1-score was frozen during the distillation process. Teacher logits were generated without gradient computation and used as soft-target guidance for MobileNetV2. The training objective combined class-weighted cross-entropy loss and distillation loss, following the teacher-student strategy in face anti-spoofing (Shahreza et al., 2024). The total loss was formulated as: $\mathcal{L}_{total} = (1 - \alpha)\mathcal{L}_{CE} + \alpha\mathcal{L}_{KD}$, where $\mathcal{L}_{CE}$ is the cross-entropy loss, $\mathcal{L}_{KD}$ is the KL-divergence-based distillation loss, and α controls the distillation contribution. This formulation allows the student model to learn from both ground-truth labels and teacher-generated soft targets without increasing the student architecture complexity.

Optuna was applied in S1 and S3 using validation F1-score as the optimization objective. The search space included learning rate, weight decay, dropout, optimizer type, alpha, and temperature. To ensure experimental fairness, all scenarios used the same identity-aware split, preprocessing pipeline, input size, batch size, random seed, early stopping strategy, and final test set. The test set was used only for final evaluation and was not involved in hyperparameter tuning, early stopping, or checkpoint selection. Model performance was evaluated using accuracy, precision, recall, F1-score, confusion matrix, APCER, BPCER, ACER, and ROC-AUC. In addition, ONNX-based feasibility was assessed by validating each exported model, counting parameters, measuring model size, and recording CPU-based dummy-input inference time.

Table 1. Final experimental configuration

Parameter	S0 Baseline	S1 Optuna	S2 KD	S3 KD & Optuna
Model	MobileNetV2	MobileNetV2	MobileNetV2	MobileNetV2
Teacher	-	-	EfficientNet-B0	EfficientNet-B0
Optimizer	AdamW	AdamW	AdamW	AdamW
Learning rate	1×10^{-4}	9.73×10^{-5}	1×10^{-4}	2.60×10^{-5}
Weight decay	0.01	6.62×10^{-4}	0.01	1.26×10^{-3}
Dropout	0.20	0.313	0.20	0.097
Alpha	-	-	0.6	0.788
Temperature	-	-	4.0	3.069
Optuna trials	-	10	-	10
Early stopping	5	2/5	5	2/5

RESULTS AND DISCUSSION

Results

The identity-aware CelebA-Spoof test set was used to evaluate four MobileNetV2-based face anti-spoofing scenarios: baseline MobileNetV2, MobileNetV2 with Optuna, MobileNetV2 with Knowledge Distillation, and MobileNetV2 with both Knowledge Distillation and Optuna. Accuracy, precision, recall, F1-score, and ROC-AUC were used to assess general classification performance, while APCER, BPCER, and ACER were used to examine anti-spoofing-specific errors. APCER is particularly important in biometric security because spoof samples classified as live may lead to false acceptance.

Table 2. General classification performance

Scenario	Accuracy	Precision	Recall	F1-score	ROC-AUC
S0	0.9943	0.9970	0.9945	0.9958	0.9998
S1	0.9940	0.9953	0.9958	0.9955	0.9998
S2	0.9940	0.9946	0.9965	0.9955	0.9994
S3	0.9913	0.9921	0.9950	0.9936	0.9991

Table 3. Anti spoofing oriented performance

Scenario	APCER	BPCER	ACER
S0	0.0055	0.0061	0.0058
S1	0.0042	0.0096	0.0069
S2	0.0035	0.0112	0.0073
S3	0.0050	0.0162	0.0106

Tables 2 and 3 show that S0 achieved the best overall balance, with the highest F1-score and the lowest ACER. This indicates that the baseline MobileNetV2 configuration was already highly competitive under the identity-aware evaluation setting. S2 achieved the highest recall and the lowest APCER, suggesting that Knowledge Distillation improved sensitivity toward spoof samples. However, this improvement was accompanied by a higher BPCER, meaning that more live samples were rejected as spoof. S3 produced the weakest overall result, indicating that combining Knowledge Distillation and Optuna did not necessarily improve the decision balance.

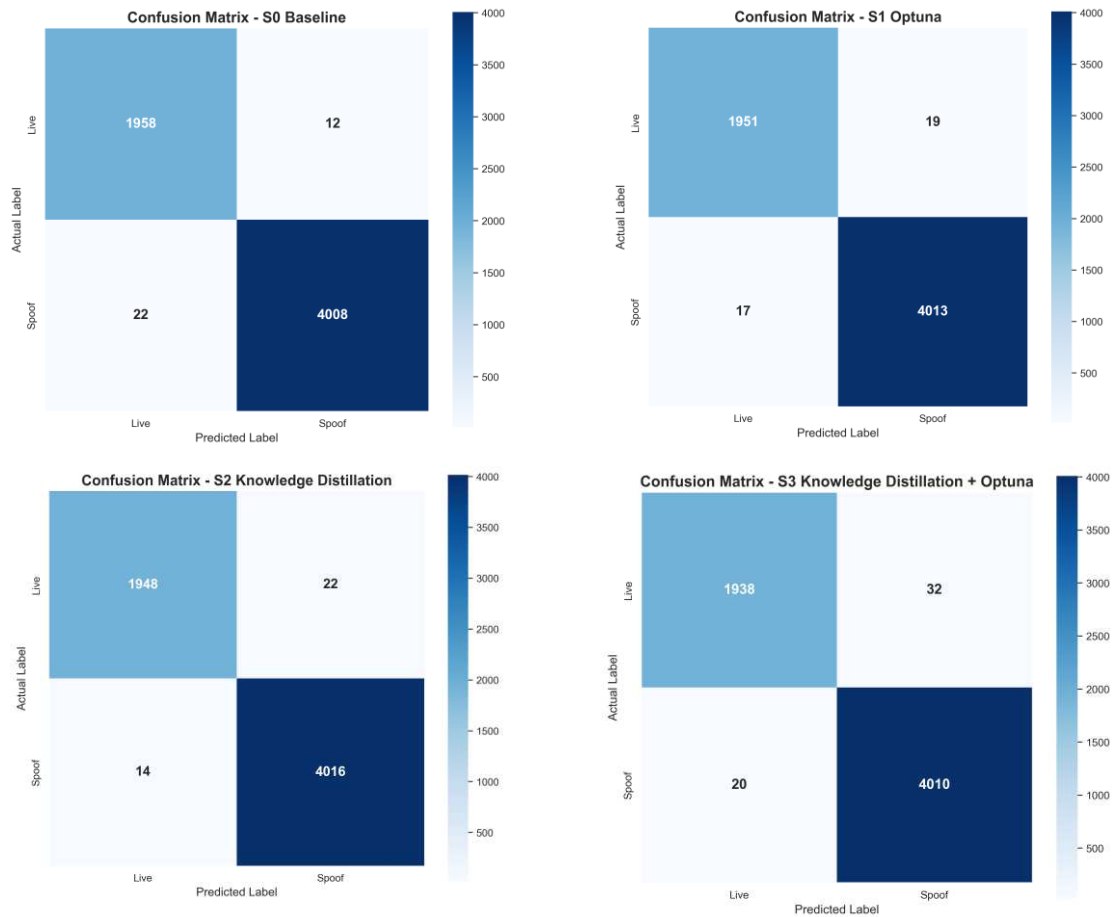


Figure 1. Confusion matrix results of four experimental scenarios

The confusion matrices in Figure 1 provide a clearer view of the error patterns. Confusion-matrix-based evaluation is useful because it shows the distribution of correct and incorrect classifications more explicitly than aggregate metrics alone (Zeng, 2025). S0 produced the most balanced distribution, with 12 false positives and 22 false negatives. S2 reduced false negatives to 14, confirming its lower APCER and stronger ability to reduce spoof acceptance errors. However, S2 also increased false positives to 22, showing a trade-off between spoof sensitivity and live-user acceptance. S3 produced the highest false positives, indicating less stable threshold-based behavior. These results suggest that Knowledge Distillation can improve spoof sensitivity, but it does not automatically improve all error types.

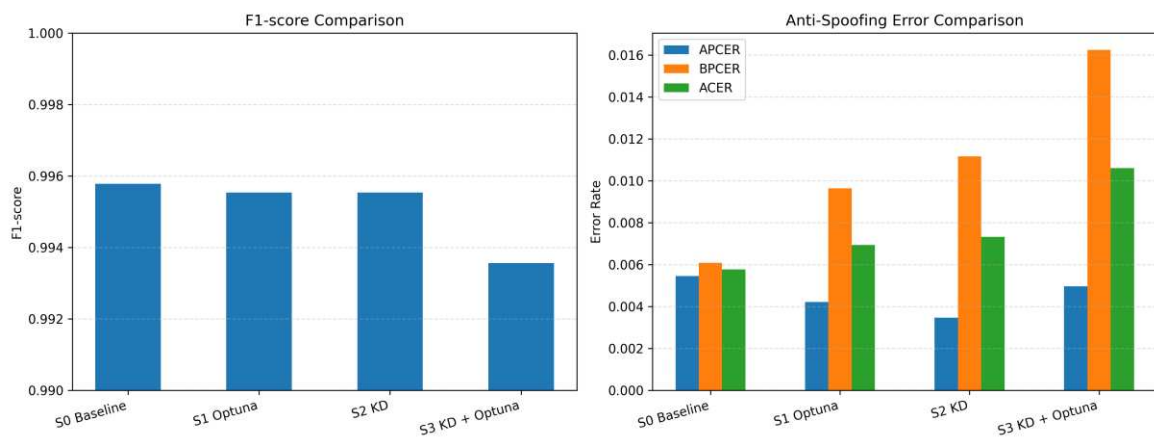


Figure 2. Comparison of f1-score and anti-spoofing error metrics across scenarios

Figure 2 summarizes the comparative trend across scenarios. The F1-score comparison confirms that S0 provided the strongest overall balance, while the anti-spoofing error comparison shows that S2 was more effective in reducing APCER. This supports the main empirical finding that the contribution of Knowledge Distillation was more specific to reducing spoof acceptance than improving all metrics uniformly. Optuna-based tuning did not provide a clear overall gain over the baseline, and its combination with Knowledge Distillation in S3 appeared less optimal.

The effect of class imbalance should also be considered. The selected subset contained more spoof samples than live samples, which may influence the model’s sensitivity toward the spoof class. Class-weighted cross-entropy was applied to reduce this effect, but the APCER and BPCER values show that each scenario still responded differently to the class distribution. S0 maintained the most balanced trade-off, while S2 became more sensitive to spoof samples.

Table 4. ONNX model complexity and execution test

Scenario	Parameters (M)	ONNX Size (MB)	Avg. CPU Time (ms)	Median Time (ms)
S0	2.23	8.46	3.47	3.17
S1	2.23	8.46	3.08	2.96
S2	2.23	8.46	2.98	2.73
S3	2.23	8.46	2.75	2.68

Table 4 shows that all exported ONNX models retained the same lightweight MobileNetV2 inference architecture, with 2.23 million parameters and a model size of 8.46 MB. The average CPU-based ONNX inference time using a dummy input ranged from 2.75 ms to 3.47 ms. Since Knowledge Distillation and Optuna affected only the training process, they did not increase the final inference complexity. These results should be interpreted as dummy-input ONNX execution readiness, not as a full real-world mobile benchmark.

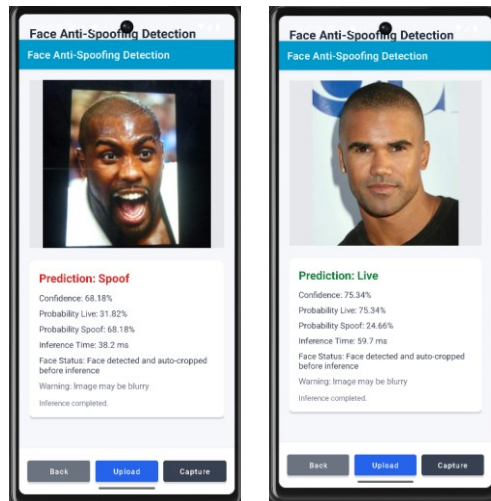


Figure 3. Mobile prototype interface for face anti-spoofing inference

The exported ONNX models were also implemented in a mobile prototype to demonstrate single-image inference. The prototype supports image input, face detection, preprocessing, prediction display, confidence score, and inference time. Figure 3 shows that the exported ONNX model could be integrated into an application environment and produce readable inference output.

Table 5. Android emulator-based prototype inference test

Scenario	Correct/Total	Accuracy	Avg. Inference Time (ms)
S0	10/10	100%	35.67
S1	9/10	90%	41.19
S2	9/10	90%	37.55
S3	8/10	80%	35.74

Table 5 presents an Android Studio emulator-based prototype test using 10 selected samples for each scenario. S0 correctly classified all samples, while S1 and S2 each produced one error; S1 misclassified one spoof sample as live, whereas S2 misclassified one live sample as spoof. S3 produced two spoof-to-live errors, indicating that the emulator-based prototype test still showed different error tendencies across scenarios. Since this test used limited selected samples in an emulator environment, it should be interpreted only as an implementation feasibility check, while the main quantitative benchmark remains the full identity-aware CelebA-Spoof test set. In addition, because each scenario was trained once using a fixed random seed, the results should be interpreted as controlled empirical trends rather than statistically significant improvements. Future work should include repeated runs, confidence intervals, cross-dataset evaluation, real-device testing, and broader mobile-captured spoof samples to assess model stability and deployment robustness.

Discussion

The findings indicate that Knowledge Distillation and Optuna-based hyperparameter tuning influenced MobileNetV2 in different ways. The baseline MobileNetV2 achieved the most balanced overall performance, as shown by the highest F1-score and the lowest ACER, while the Knowledge Distillation scenario produced the lowest APCER. This means that Knowledge Distillation should not be interpreted as universally superior across all metrics, but as a strategy that provides a more security-oriented benefit by reducing spoof samples incorrectly accepted as live. Therefore, this study clarifies that representation enhancement and hyperparameter optimization do not necessarily produce the same type of improvement in lightweight face anti-spoofing.

The lower APCER under Knowledge Distillation can be explained through soft-target learning, representational transfer, and decision boundary smoothing. Unlike hard-label training, distillation provides softened class-distribution information from the EfficientNet-B0 teacher, allowing MobileNetV2 to learn more nuanced relationships between live and spoof samples. This guidance may help the student model become more sensitive to subtle spoofing cues, such as texture distortion, printed-photo artifacts, screen reflection, and illumination inconsistencies. This is relevant because face anti-spoofing performance can be affected by variations in image appearance, acquisition conditions, and visual quality (Xu et al., 2023). This interpretation is also consistent with previous studies showing that feature enhancement and stronger representation are important for improving face anti-spoofing performance (Li et al., 2025; Qi et al., 2025). However, the lower APCER in S2 was accompanied by higher BPCER, indicating a trade-off between stronger spoof sensitivity and live-user acceptance.

In contrast, Optuna-based hyperparameter tuning did not provide a clear overall improvement over the baseline. Lai et al. (2024) showed that Optuna can improve model training by searching for better configurations, such as learning rate, weight decay, dropout, and optimizer settings. However, hyperparameter tuning mainly adjusts training dynamics and does not directly expand the representational capacity of MobileNetV2. This may explain why S1 remained competitive but did not surpass the baseline in F1-score or ACER. The weaker performance of S3 further suggests that combining Knowledge Distillation and Optuna does

not automatically produce an additive effect, especially when alpha, temperature, learning rate, and dropout make the model more sensitive to the validation objective.

Compared with previous face anti-spoofing studies, this research offers a different perspective by directly comparing representation enhancement and training-configuration optimization under the same lightweight architecture. Previous studies have shown the usefulness of Knowledge Distillation for improving face anti-spoofing generalization (Kong et al., 2024) and the relevance of hyperparameter optimization for improving model performance (Lai et al., 2024). However, the present study shows that Knowledge Distillation mainly improves spoof sensitivity, while Optuna does not necessarily improve the final decision balance. This finding suggests that improving lightweight CNN performance depends not only on tuning training parameters, but also on the quality of representation learned by the student model.

The identity-aware evaluation protocol strengthens the methodological contribution of this study. Random splitting may allow the same subject identity to appear in both training and testing subsets, causing identity leakage and overly optimistic performance estimates. By separating identities across training, validation, and testing sets, this study provides a stricter assessment of whether the model learns spoof-related cues rather than identity-specific facial patterns. Ham et al. (2025) emphasized that generalization remains a central issue in face anti-spoofing, especially under unseen identities, domain shifts, and different acquisition conditions. This concern is also consistent with Wang et al. (2023), who highlighted the importance of evaluating face anti-spoofing models under unseen conditions. Thus, the identity-aware protocol makes the evaluation closer to real biometric authentication scenarios.

In biometric security systems, the trade-off between APCER and BPCER plays a crucial role in determining overall performance. Reducing APCER helps prevent spoof attacks from being misclassified as genuine users, thereby strengthening security. However, higher BPCER values may lead to the rejection of legitimate users, which can negatively affect usability. The results obtained in S2 indicate that Knowledge Distillation contributes to more stringent spoof detection. Even so, improvements in security should be considered alongside their impact on user acceptance. Therefore, systems with stringent security requirements may emphasize lower APCER values, while attendance and general access-control applications may prioritize a more balanced ACER to maintain both security and usability.

The efficiency analysis shows that all scenarios retained the same MobileNetV2 inference architecture. Knowledge Distillation and Optuna changed the training process, but they did not increase the number of parameters or ONNX model size. This indicates that the observed performance differences were caused by training strategy rather than by a larger inference model. In practical terms, this supports the feasibility of improving lightweight CNN decision behavior without increasing deployment complexity. This implication is important because lightweight face anti-spoofing models are often needed for real-time or resource-constrained authentication environments (Zawar & Chakkarwar, 2023). Nevertheless, the ONNX execution test and Android emulator prototype should be interpreted only as implementation feasibility evidence, not full real-world deployment validation. Since each scenario was trained once using a fixed random seed, the findings should be viewed as controlled empirical trends rather than statistically significant improvements. Future work should include repeated runs, confidence intervals, paired statistical testing, cross-dataset evaluation, real-device benchmarking, and broader mobile-captured live and spoof samples.

CONCLUSION

This study concludes that representation enhancement through Knowledge Distillation and optimization-based improvement through Optuna affect MobileNetV2-based lightweight face anti-spoofing in different ways under an identity-aware evaluation protocol. The baseline

MobileNetV2 achieved the most balanced overall performance, while Knowledge Distillation provided a more security-oriented benefit by reducing spoof samples incorrectly accepted as live, indicating that richer teacher-guided representation is more useful for improving spoof sensitivity than hyperparameter tuning alone. Theoretically, the findings support the view that lightweight CNNs require stronger representational guidance to capture subtle spoofing cues, while methodologically, the use of identity-aware splitting provides a stricter evaluation by reducing identity leakage and testing unseen-subject generalization. Practically, the ONNX execution test and Android emulator prototype show implementation feasibility, but not full deployment readiness. This study is limited by its controlled single-run setting, the use of CelebA-Spoof as the main benchmark, and prototype-level emulator testing; therefore, future work should include repeated runs, confidence intervals, cross-domain evaluation, broader mobile-captured spoof samples, real-device benchmarking, and further investigation of face anti-spoofing fine-tuned teacher models.

REFERENCES

- Cai, R., Yu, Z., Kong, C., Li, H., Chen, C., Hu, Y., & Kot, A. C. (2024). S-Adapter: Generalizing Vision Transformer for Face Anti-Spoofing With Statistical Tokens. *IEEE Transactions on Information Forensics and Security*, 19, 8385–8397. <https://doi.org/10.1109/TIFS.2024.3420699>
- Gong, L. Y., Li, X. J., & Chong, P. H. J. (2024). Facial Anti-Spoofing Using “Clue Maps.” *Sensors*, 24(23). <https://doi.org/10.3390/s24237635>
- Gulzar, Y. (2023). Fruit Image Classification Model Based on MobileNetV2 with Deep Transfer Learning Technique. *Sustainability (Switzerland)*, 15(3). <https://doi.org/10.3390/su15031906>
- Guo, Z., Wang, D., He, Q., & Zhang, P. (2024). Leveraging logit uncertainty for better knowledge distillation. *Scientific Reports*, 14(1). <https://doi.org/10.1038/s41598-024-82647-6>
- Ham, H., Saptawijaya, A., & Arymurthy, A. M. (2025). Domain Generalization and Domain Adaptation Approaches for Face Anti-Spoofing: A Survey. *IEEE Access*, 13, 149390–149408. <https://doi.org/10.1109/ACCESS.2025.3592034>
- Kong, Z., Zhang, W., Wang, T., Zhang, K., Li, Y., Tang, X., & Luo, W. (2024). Dual Teacher Knowledge Distillation with Domain Alignment for Face Anti-Spoofing. *IEEE Transactions on Circuits and Systems for Video Technology*, 34(12), 13177–13189. <https://doi.org/10.1109/TCSVT.2024.3451294>
- Lai, L. H., Lin, Y. L., Liu, Y. H., Lai, J. P., Yang, W. C., Hou, H. P., & Pai, P. F. (2024). The Use of Machine Learning Models with Optuna in Disease Prediction. *Electronics (Switzerland)*, 13(23). <https://doi.org/10.3390/electronics13234775>
- Li, Y., Sun, W., Li, Z., & Guo, X. (2025). Face Anti-Spoofing Based on Adaptive Channel Enhancement and Intra-Class Constraint. *Journal of Imaging*, 11(4). <https://doi.org/10.3390/jimaging11040116>
- Li, Z., Cai, R., Li, H., Lam, K. Y., Hu, Y., & Kot, A. C. (2022). One-Class Knowledge Distillation for Face Presentation Attack Detection. *IEEE Transactions on Information Forensics and Security*, 17, 2137–2150. <https://doi.org/10.1109/TIFS.2022.3178240>
- Mun, H. J., & Lee, M. H. (2022). Design for Visitor Authentication Based on Face Recognition Technology Using CCTV. *IEEE Access*, 10, 124604–124618. <https://doi.org/10.1109/ACCESS.2022.3223374>
- Qi, H., Han, R., Duan, K., Shi, Y., Qi, X., Gao, C., & Ren, L. (2025). A high-performance adaptive fusion network for face anti-spoofing detection. *Scientific Reports*, 15(1). <https://doi.org/10.1038/s41598-025-21461-0>
- Shahreza, H. O., George, A., & Marcel, S. (2024). Knowledge Distillation for Face Recognition

- Using Synthetic Data With Dynamic Latent Sampling. *IEEE Access*, 12, 187800–187812. <https://doi.org/10.1109/ACCESS.2024.3505621>
- Wang, W., Liu, P., Zheng, H., Ying, R., & Wen, F. (2023). Domain Generalization for Face Anti-Spoofing via Negative Data Augmentation. *IEEE Transactions on Information Forensics and Security*, 18, 2333–2344. <https://doi.org/10.1109/TIFS.2023.3266138>
- Xiao, J., Wang, W., Zhang, L., & Liu, H. (2024). A MobileFaceNet-Based Face Anti-Spoofing Algorithm for Low-Quality Images. *Electronics (Switzerland)*, 13(14). <https://doi.org/10.3390/electronics13142801>
- Xing, H., Tan, S. Y., Qamar, F., & Jiao, Y. (2025). Face Anti-Spoofing Based on Deep Learning: A Comprehensive Survey. In *Applied Sciences (Switzerland)* (Vol. 15, Number 12). Multidisciplinary Digital Publishing Institute (MDPI). <https://doi.org/10.3390/app15126891>
- Xu, M., Yoon, S., Fuentes, A., & Park, D. S. (2023). A Comprehensive Survey of Image Augmentation Techniques for Deep Learning. *Pattern Recognition*, 137(137), 109347. <https://doi.org/10.1016/j.patcog.2023.109347>
- 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>
- Zawar, R., & Chakkarwar, V. (2023). Real-Time Face Liveness Detection and Face Anti-spoofing Using Deep Learning. *International Journal of Intelligent Systems and Applications in Engineering*, 11, 626–636. https://doi.org/10.2991/978-94-6463-196-8_47
- Zeng, G. (2025). Invariance Properties and Evaluation Metrics Derived from the Confusion Matrix in Multiclass Classification. *Mathematics*, 13(16). <https://doi.org/10.3390/math13162609>
- Zheng, T., Li, B., Wu, S., Wan, B., Mu, G., Liu, S., Ding, S., & Wang, J. (2024). MFAE: Masked Frequency Autoencoders for Domain Generalization Face Anti-Spoofing. *IEEE Transactions on Information Forensics and Security*, 19, 4058–4069. <https://doi.org/10.1109/TIFS.2024.3371266>