

AUDIO FEATURE EXTRACTION FOR PREDICTING VIRAL TIKTOK CONTENT USING THE CONVOLUTIONAL NEURAL NETWORK ALGORITHM

Edo Kurniawan^{1*}, Nur Aeni Widiastuti², Teguh Tamrin³

Universitas Islam Nahdlatul Ulama Jepara, Indonesia
pardameanestu@gmail.com^{1*}, Nuraeniwidiastuti@unisnu.ac.id, Teguh@unisnu.ac.id
**Corresponding author*

Received April 28, 2026; Revised May 29, 2026; Accepted May 30, 2026; Published May 31, 2026

ABSTRACT

The rapid growth of TikTok has increased interest in identifying factors that influence content virality, particularly audio elements that play a central role in trend formation and user engagement. This study aims to develop a model for predicting TikTok content virality based on audio characteristics using a Convolutional Neural Network (CNN). The proposed approach focuses exclusively on audio information to evaluate its independent contribution to virality prediction without incorporating visual, textual, or engagement-based features. The research employed a quantitative experimental design using audio extracted from publicly available TikTok videos categorized into viral and non-viral classes. Audio signals were preprocessed through normalization and duration standardization before being transformed into Mel-spectrogram representations. These spectrogram images were then used as input to a CNN model for automatic feature extraction and classification. Model performance was evaluated using precision, recall, and F1-score metrics. The experimental results demonstrate that the proposed CNN model effectively distinguished viral and non-viral TikTok content. Evaluation on the testing dataset produced a precision of 0.833, recall of 1.000, and F1-score of 0.909. In addition, Mel-spectrogram visualizations revealed distinct frequency-energy patterns between viral and non-viral audio samples, indicating that acoustic characteristics contain meaningful information associated with content virality. In conclusion, audio features can serve as reliable predictors of TikTok content virality, and the CNN-based framework successfully extracted discriminative acoustic patterns from Mel-spectrogram representations. This study contributes to the fields of audio analytics and social media intelligence by providing an audio-centered approach for early-stage virality prediction prior to content publication.

Keywords: TikTok, Content virality, Audio features, Mel-spectrogram, Convolutional neural network

INTRODUCTION

The rapid growth of social media platforms has fundamentally transformed the way information is created, distributed, and consumed. Among various social networking applications, TikTok has emerged as one of the most influential short-video platforms, enabling users to produce and share multimedia content with unprecedented speed and reach. Unlike conventional social media platforms that rely heavily on follower networks, TikTok employs a recommendation-driven ecosystem capable of exposing content to large audiences regardless of the creator's popularity. This mechanism significantly increases the likelihood of content becoming viral within a short period. According to Zeng et al. (2023), TikTok's success is closely associated with its algorithmic content recommendation system, user participation culture, and the integration of audiovisual elements that encourage continuous engagement and content replication. These characteristics have made TikTok an important subject in studies related to digital

communication, social media analytics, and content popularity prediction.

The increasing volume of content published on TikTok has motivated researchers to investigate factors that influence content virality. Viral content is commonly characterized by rapid growth in views, likes, comments, and shares shortly after publication. Previous studies have shown that predicting content popularity remains a challenging task because user engagement is influenced by multiple interconnected factors, including content characteristics, platform algorithms, and audience behavior. Recent developments in deep learning have provided new opportunities for addressing this challenge through automated feature extraction and pattern recognition. Convolutional Neural Networks (CNNs), in particular, have demonstrated strong capabilities in learning complex nonlinear relationships from high-dimensional data and have achieved excellent performance across various classification domains, including signal processing, image analysis, anomaly detection, and predictive analytics. Recent studies reported by Cai (2026) and Pal et al. (2026) demonstrate that CNN-based architectures consistently achieve high classification accuracy by automatically extracting discriminative features from complex datasets, outperforming many traditional machine learning approaches.

Among the various components contained in TikTok videos, audio represents one of the most influential elements in shaping content dissemination and audience engagement. Trending music, sound effects, rhythm patterns, and audio signatures are frequently reused by content creators and often become central drivers of viral challenges and platform-wide trends. Despite its importance, audio remains relatively underexplored compared with visual features and engagement-based indicators in virality prediction studies. Recent research by Reis et al. (2026) demonstrated that acoustic signals transformed into Mel-spectrogram representations can be effectively analyzed using deep convolutional neural networks, achieving classification accuracy exceeding 94% in real-world audio recognition tasks. The study further highlights the ability of Mel-spectrogram-based CNN models to capture discriminative temporal and frequency-domain patterns that are difficult to identify through conventional feature engineering approaches. These findings suggest that audio characteristics contain rich predictive information and may serve as a valuable source for identifying potential viral content before extensive user interaction data becomes available.

Recent advances in deep learning have further strengthened the role of feature extraction in multimedia analytics. Convolutional Neural Networks have been successfully applied across a wide range of domains, including medical diagnosis, fault detection, signal classification, financial fraud detection, and image recognition. Studies conducted by Shi et al. (2026) and Alqassab et al. (2026) demonstrated that CNN-based architectures are capable of automatically extracting hierarchical features from complex spectral and image-based representations, resulting in high classification performance and improved robustness compared with conventional approaches. Similarly, Pal et al. (2026) reported that convolution-based architectures effectively capture local patterns and hidden correlations in high-dimensional data, making them particularly suitable for

classification problems involving complex nonlinear relationships. These findings indicate that CNN has become one of the most reliable deep learning techniques for extracting meaningful information from large-scale and heterogeneous datasets.

Despite the growing body of research on content popularity prediction, existing studies predominantly focus on user engagement metrics, visual content characteristics, textual information, or multimodal approaches that combine multiple data sources. While these approaches often achieve strong predictive performance, they provide limited understanding of the independent contribution of audio characteristics to content virality. Furthermore, many prediction models rely on post-publication engagement indicators that are only available after a video has been distributed to audiences. This limitation reduces their usefulness as early-stage decision-support tools for content creators. In contrast, recent research on acoustic signal analysis has demonstrated that time-frequency representations such as Mel-spectrograms contain discriminative information that can be effectively processed by CNN models for classification tasks. The successful application of Mel-spectrogram-based CNNs in acoustic fault detection and signal recognition suggests that similar approaches may be extended to analyze viral content patterns in social media environments.

Based on the identified research gap, this study proposes a Convolutional Neural Network-based framework for predicting TikTok content virality using audio characteristics as the primary source of information. Unlike previous studies that emphasize engagement metrics, metadata, visual features, or multimodal data integration, the proposed approach focuses exclusively on audio signals represented through Mel-spectrogram transformations. This design enables the model to learn frequency-domain and temporal audio patterns that may be associated with viral trends before substantial audience interaction occurs. Therefore, the novelty of this study lies in three main contributions: (1) investigating audio as an independent predictor of TikTok content virality, (2) utilizing Mel-spectrogram representations to capture discriminative acoustic features, and (3) implementing a CNN-based classification framework capable of automatically extracting complex audio patterns for viral content prediction. These contributions are expected to enrich the literature on audio analytics and social media intelligence while providing practical insights for content creators and digital marketers seeking to optimize audio selection strategies before content publication.

Beyond its academic significance, early prediction of content virality has substantial practical value for content creators, digital marketers, and businesses that rely on social media platforms to reach target audiences. The ability to identify potentially viral content before publication can support more effective content planning, optimize promotional strategies, and reduce the uncertainty associated with digital marketing campaigns. Recent advances in deep learning have demonstrated that predictive models can successfully extract hidden patterns from complex and high-dimensional data across various application domains, including financial anomaly detection, fault diagnosis, intention prediction, and healthcare analytics. Studies by Cai (2026), Pal et al. (2026) and Su et al. (2026) reported that convolution-based architectures are capable of

identifying subtle and nonlinear relationships that are often difficult to detect using traditional analytical approaches. These findings suggest that deep learning-driven prediction systems can serve as reliable decision-support tools in dynamic environments characterized by large volumes of heterogeneous data. Consequently, investigating audio-driven virality prediction on TikTok is not only scientifically relevant but also practically important in supporting data-informed content development strategies.

Therefore, this study aims to develop and evaluate a Convolutional Neural Network (CNN) model for predicting TikTok content virality based exclusively on audio characteristics. Audio signals extracted from TikTok content are transformed into Mel-spectrogram representations to enable the automatic extraction of temporal and frequency-domain patterns through deep learning techniques. The selection of CNN is motivated by its proven effectiveness in learning discriminative features from complex representations, as demonstrated in recent studies involving acoustic signal analysis, spectral classification, image recognition, and pattern identification tasks. By focusing solely on audio information, this research seeks to provide a clearer understanding of the independent contribution of audio features to content virality while minimizing the influence of visual, textual, and engagement-based factors. Ultimately, the proposed framework is expected to contribute to the advancement of audio analytics, social media intelligence, and early-stage viral content prediction models for short-video platforms.

RESEARCH METHODOLOGY

Research Design and Data Preparation

This study employed a quantitative experimental approach to develop a TikTok content virality prediction model based on audio characteristics using a Convolutional Neural Network (CNN). The research workflow consisted of audio collection, audio preprocessing, Mel-spectrogram generation, CNN model training, and performance evaluation. The dataset comprised audio extracted from publicly available TikTok videos and categorized into viral and non-viral classes according to predefined engagement criteria. Each audio file was converted into WAV format, normalized, and standardized to a uniform duration. Subsequently, the audio signals were transformed into Mel-spectrogram representations using the Short-Time Fourier Transform (STFT), enabling the extraction of both temporal and frequency-domain information. The resulting Mel-spectrogram images served as input to the CNN model.

Convolutional Neural Network Architecture

The proposed CNN architecture consisted of convolutional layers, pooling layers, flatten layers, and fully connected layers. Convolutional layers were employed to automatically extract discriminative features from Mel-spectrogram images, while pooling layers reduced feature-map dimensionality and computational complexity. ReLU activation was utilized to enhance nonlinear learning capability, and dropout was applied to reduce overfitting.

a. Convolution Operation

The convolution layer extracts local features from the Mel-spectrogram input according to Equation (1).

$$y(i,j) = \sum_{m=0}^{M-1} \sum_{n=0}^{N-1} x(i+m,j+n) \cdot w(m,n) + b \dots (1)$$

where:

$x(i,j)$ = input value (Mel-spectrogram),
 $w(m,n)$ = convolution kernel weight,
 b = bias,
 $y(i,j)$ = output feature map,
 $M \times N$ = kernel size.

b. ReLU Activation Function

To improve nonlinear learning capability and accelerate convergence, the Rectified Linear Unit (ReLU) activation function was applied as expressed in Equation (2).

$$f(x) = \max(0, x) \dots (2)$$

c. Pooling Layer

Max-pooling was employed to reduce feature-map dimensions while preserving important information. The pooling operation is formulated in Equation (3).

$$y = \max(x_1, x_2, \dots, x_n) \dots (3)$$

d. Fully Connected Layer and Sigmoid Function

After feature extraction, the resulting feature maps were flattened and forwarded to fully connected layers for classification. For binary classification, the sigmoid activation function was used as shown in Equation (4).

$$\sigma(x) = \frac{1}{1+e^{-x}} \dots (4)$$

The output value ranges from 0 to 1 and represents the probability of belonging to the viral class.

Model Training

The preprocessed dataset was divided into training, validation, and testing sets. Model training was performed using the Adam optimizer and binary cross-entropy loss function. Validation performance was monitored throughout training to prevent overfitting, and early stopping was employed when no significant improvement was observed after several consecutive epochs.

Performance Evaluation

Model performance was evaluated using accuracy, precision, recall, and F1-score metrics. In addition, a confusion matrix was utilized to analyze classification errors between viral and non-viral classes.

a. Recall

Recall measures the model's ability to correctly identify all positive instances and is calculated using Equation (5).

$$\text{Recall} = \frac{TP}{TP+FN} \dots(5)$$

where:

TP (True Positive) = correctly classified viral content,

FN (False Negative) = viral content incorrectly classified as non-viral.

b. F1-Score

The F1-score represents the harmonic mean of precision and recall and is calculated using Equation (7).

$$\text{F1-score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \dots(7)$$

The F1-score provides a balanced measure of classification performance, particularly when class distributions are imbalanced.

RESULTS AND DISCUSSION

This section presents the empirical findings obtained from the training and testing processes of the proposed audio-based Convolutional Neural Network (CNN) model for TikTok content virality prediction. The discussion focuses on audio feature representation, model training performance, testing results, and overall classification performance.

Audio Preprocessing and Feature Representation

The collected TikTok audio samples were first subjected to preprocessing procedures, including amplitude normalization, duration standardization, and Mel-spectrogram transformation. This process converted raw audio signals into two-dimensional time-frequency representations suitable for CNN-based feature extraction. The resulting Mel-spectrograms revealed observable differences between viral and non-viral content. Viral audio samples generally exhibited stronger energy concentrations and more distinctive frequency distributions across specific frequency bands. In contrast, non-viral samples displayed relatively flatter and less dynamic spectral patterns. These differences provided discriminative acoustic features that enabled the CNN model to distinguish between the two classes.

CNN Training Results

The CNN model was trained using both training and validation datasets over

multiple epochs. During the initial stages of training, classification accuracy increased rapidly, indicating that the model effectively learned relevant audio patterns from the Mel-spectrogram inputs. As training progressed, both training and validation accuracy gradually stabilized, suggesting convergence of the learning process.

The training curves presented in Figure 1 demonstrate a continuous increase in accuracy accompanied by a consistent decrease in loss values. Similarly, the validation accuracy remained relatively stable throughout the training process, while validation loss gradually decreased and approached a low value. These results indicate that the model successfully learned meaningful feature representations without experiencing significant performance degradation.

Testing Results

The trained model was subsequently evaluated using testing data that were not involved during the training phase. The testing results demonstrate that the proposed CNN model was capable of distinguishing viral and non-viral TikTok content with a high level of confidence. Several testing samples were correctly classified as viral with prediction probabilities approaching 1.00, while non-viral samples were also correctly identified with similarly high confidence scores. These findings suggest that the learned audio representations captured meaningful acoustic characteristics associated with content virality and generalized effectively to unseen data.

Model Performance Evaluation

a. CNN Training Performance

The training performance illustrated in Figure 1 and Figure 2 shows that the CNN model learned audio feature patterns effectively throughout the optimization process. Accuracy improved substantially during the early epochs before reaching a stable plateau, while loss values consistently decreased as training progressed. The relatively small difference between the training and validation curves suggests that the model maintained stable learning behavior and did not exhibit substantial overfitting. Consequently, the model demonstrated good generalization capability when applied to previously unseen audio samples.

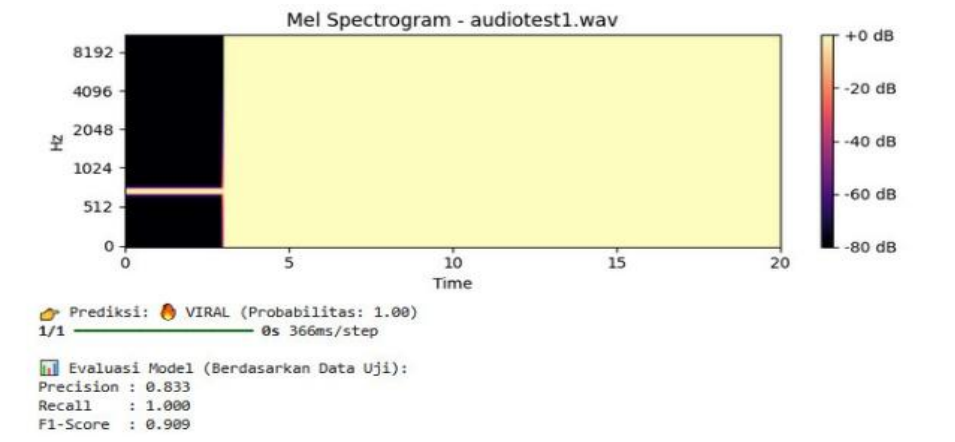


Figure 1. CNN Training Results

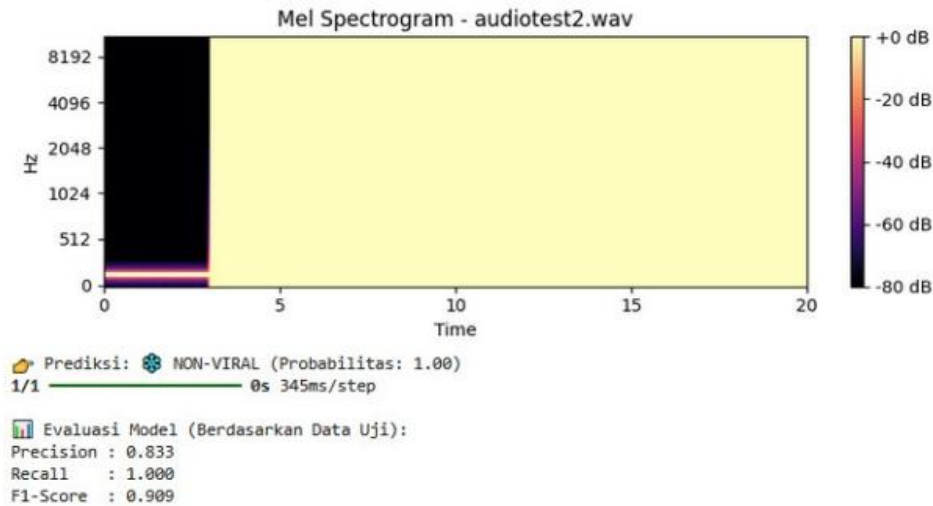


Figure 2. CNN Training Results (Continued)

b. CNN Testing Performance

The testing phase was conducted to evaluate the classification capability of the trained CNN model. Based on the evaluation results, the model achieved high performance in identifying viral content and maintaining balanced classification capability. The model obtained a Recall value of 1.000, indicating that all viral content samples in the testing dataset were successfully identified. Furthermore, the model achieved a Precision value of 0.833, suggesting that most samples predicted as viral were classified correctly. The resulting F1-score of 0.909 demonstrates a strong balance between precision and recall. Table 1 summarizes the classification performance of the proposed CNN model.

Table 1. Classification Performance Metrics

Metric	Value
Precision	0.833
Recall	1.000
F1-Score	0.909

The high recall value indicates excellent sensitivity in detecting viral content, while the precision score demonstrates that the majority of viral predictions were accurate. The F1-score further confirms that the model achieved a balanced classification performance between identifying viral content and minimizing incorrect predictions. To further analyze classification outcomes, a confusion matrix can be employed to evaluate the distribution of correctly and incorrectly classified samples. The confusion matrix framework used in this study is shown in Table 2.

Table 2. Confusion Matrix Structure

Actual Class	Predicted Viral	Predicted Non-Viral
Viral	TP	FN
Non-Viral	FP	TN

Where TP denotes True Positive, TN denotes True Negative, FP denotes False Positive, and FN denotes False Negative classifications. The confusion matrix provides a detailed assessment of model predictions and facilitates further analysis of classification errors. Overall, the evaluation results demonstrate that the proposed audio-based CNN model provides reliable performance for TikTok virality prediction. The obtained metrics indicate that audio characteristics contain valuable information for distinguishing viral and non-viral content, supporting the effectiveness of Mel-spectrogram-based CNN classification in social media content analysis.

The evaluation results demonstrate that the proposed CNN-based model achieved strong classification performance in distinguishing viral and non-viral TikTok content based solely on audio characteristics. The model obtained a Precision value of 0.833, Recall of 1.000, and F1-score of 0.909, indicating that the extracted audio features contained sufficient information to support reliable virality prediction. The perfect recall score suggests that all viral samples in the testing dataset were successfully identified, while the high F1-score reflects a balanced relationship between detection capability and classification precision. Similar findings have been reported in CNN-based predictive systems across various domains. Cai (2026) demonstrated that convolutional neural networks effectively identify hidden patterns associated with financial anomalies, while Paul & Das (2023) reported that deep learning models outperform conventional approaches in forecasting complex temporal trends. Likewise, Fadli et al. (2021) showed that CNN architectures can achieve high classification performance by automatically extracting relevant features from structured input representations.

The effectiveness of the proposed approach is closely related to the use of Mel-spectrogram representations as input to the CNN model. By transforming audio signals into time-frequency images, Mel-spectrograms preserve both temporal and spectral information that can be efficiently processed through convolutional operations (Khalil & Bakar, 2023). Reis et al. (2026) demonstrated that acoustic signal representations combined with deep learning significantly improve classification performance in fault detection applications, highlighting the ability of CNN models to identify subtle spectral variations. Similar observations were reported by (Shi et al., 2026) who successfully applied spectral convolutional neural networks to analyze complex frequency-domain patterns. Furthermore, Naufal et al. (2021) showed that CNN-based classification consistently outperformed traditional machine learning approaches by leveraging hierarchical feature extraction mechanisms. These findings support the selection of Mel-spectrogram representations as an effective medium for capturing audio characteristics associated with content virality.

The findings also reinforce the growing recognition that audio constitutes an important component of user engagement within digital media ecosystems. Unlike engagement metrics such as views, likes, and shares, audio characteristics exist before a video is published and therefore provide an opportunity for early-stage virality prediction. The recommendation-driven environment of short-video platforms encourages the repeated use of recognizable sounds, music segments, and audio patterns, allowing

specific audio characteristics to influence content dissemination. Zeng et al. (2023) emphasized the importance of multimedia interaction and intelligent digital systems in supporting user engagement, while Reis et al. (2026) highlighted the discriminative capability of acoustic patterns in classification tasks. Consequently, the present findings suggest that audio should not merely be viewed as a supplementary content element but rather as an intrinsic feature capable of influencing viral diffusion dynamics (Samuel et al., 2026).

The strong performance achieved by the proposed model is consistent with a broader trend in recent artificial intelligence research, where convolution-based architectures have demonstrated remarkable capability in extracting meaningful representations from complex and high-dimensional data. Daka & Bhattacharyya (2026) reported that convolutional neural network frameworks successfully captured discriminative visual patterns for classification tasks, while Vijayabalan, (2026) demonstrated the effectiveness of CNN-based techniques in identifying subtle structural characteristics within neuroimaging data. Similar outcomes have been observed in intelligent sensing and engineering applications. Wu et al. (2027) utilized deep learning architectures to process complex signals generated by advanced electronic skin systems, whereas Hua et al. (2027) showed that modern neural architectures can accurately identify intricate patterns in highly dynamic combustion environments. Although these studies were conducted in different application domains, they collectively highlight the ability of deep neural networks to learn informative representations from complex data structures, supporting the suitability of CNNs for audio-based virality prediction.

The present findings are also aligned with recent studies investigating prediction and classification problems in industrial, behavioral, and healthcare contexts. Szrama (2026) demonstrated that deep learning frameworks can effectively model complex temporal relationships for predictive maintenance and remaining useful life estimation. Similarly, Pal et al. (2026) reported that hybrid CNN-based architectures provide robust fault diagnosis performance by automatically extracting relevant features from large-scale industrial processes. In the context of human behavior analysis, Szrama (2026) showed that convolution-based models successfully captured latent patterns associated with human intention prediction, while Alqassab et al., (2026) achieved reliable disease identification performance using advanced CNN frameworks. These studies consistently indicate that CNN models are capable of identifying subtle and nonlinear relationships that may not be apparent through conventional analytical techniques. Consequently, the ability of the proposed model to distinguish viral and non-viral TikTok content further confirms the adaptability of CNN architectures across diverse predictive and classification tasks.

The primary novelty of this study lies in its exclusive use of audio characteristics as an independent predictor of TikTok content virality. Most previous studies have focused on metadata, user engagement indicators, visual content, multimodal information, or domain-specific classification problems. Studies by Daka & Bhattacharyya (2026) and Shi et al. (2026) primarily investigated the effectiveness of

convolution-based architectures in solving classification, prediction, diagnostic, and recognition problems across various domains. In contrast, the present study specifically examines whether audio information alone can provide sufficient predictive capability for identifying viral content on TikTok. The obtained Precision (0.833), Recall (1.000), and F1-score (0.909) demonstrate that meaningful virality-related information is embedded within audio signals and can be successfully extracted through Mel-spectrogram representations and CNN-based learning. Therefore, this research contributes to both social media analytics and audio intelligence by establishing an audio-centered framework for early-stage virality prediction without relying on post-publication engagement metrics or multimodal data integration.

CONCLUSION

This study developed an audio-based Convolutional Neural Network (CNN) framework to predict TikTok content virality using Mel-spectrogram representations. The experimental results demonstrated that audio characteristics contain meaningful information for distinguishing viral and non-viral content, as evidenced by the model's Precision of 0.833, Recall of 1.000, and F1-score of 0.909. The Mel-spectrogram analysis further revealed observable differences in frequency-energy distributions between viral and non-viral audio samples, enabling the CNN model to effectively learn discriminative acoustic patterns. These findings confirm that audio can serve as an independent predictor of content virality without relying on visual features, textual information, or post-publication engagement metrics. The primary contribution of this research lies in establishing an audio-centered virality prediction framework that supports early-stage content evaluation prior to publication. Future studies are recommended to employ larger and more diverse datasets, incorporate additional audio features, and compare CNN performance with other deep learning architectures to further improve prediction accuracy and model generalizability.

REFERENCES

- Alqassab, A. I. M., Luque-Nieto, M.-Á., & Mohammed, M. A. (2026). Identification of multiple ocular diseases using a hybrid quantum convolutional neural network with fundus images. *Scientific Reports*, 16(1), 1–14. <https://doi.org/10.1038/s41598-026-38063-z>
- Cai, M. (2026). Detecting Financial Fraud and Anomalies in Corporate Transactions with Convolutional Neural Networks. *International Journal of Computational Intelligence Systems*, 19(1), 1–27. <https://doi.org/10.1007/s44196-026-01275-2>
- Daka, C., & Bhattacharyya, S. (2026). A novel quantum convolutional neural network framework for quantum-enhanced classification of pixelated colour images. *Scientific Reports*, 16(1), 1–23. <https://doi.org/10.1038/s41598-026-45140-w>
- Fadli, A., Ramadhani, Y., & Aliim, M. S. (2021). Purwarupa Sistem Deteksi COVID-19 Berbasis Website Menggunakan Algoritma CNN. *Jurnal RESTI*, 5(5), 876–883. <https://doi.org/10.29207/resti.v5i5.3450>

- Hua, M., Zhang, G., Zhang, F., Wen, C., Li, J., & Xu, H. (2027). Cellpose-SAM: Transformer-based segmentation for cellular instability analysis of spherical hydrogen-air premixed flames. *Fuel*, 427, 1–11. <https://doi.org/10.1016/j.fuel.2026.139669>
- Khalil, R. A. B., & Bakar, A. A. B. U. (2023). A Comparative Study of Deep Learning Algorithms in Univariate and Multivariate Forecasting of the Malaysian Stock Market. *Sains Malaysiana*, 52(3), 993–1009. <https://doi.org/10.17576/jsm-2023-5203-22>
- Naufal, M. F., Shania, S., Millenia, J., Axel, S., Soebroto, J. T., Rizka Febrina, P., & Mercifia, M. (2021). Analisis Perbandingan Algoritma Klasifikasi MLP dan CNN pada Dataset American Sign Language. *Jurnal RESTI*, 5(3), 489–495. <https://doi.org/10.29207/resti.v5i3.3009>
- Pal, P. K., Ray, S., Behera, N., Chowdhury, S., Hens, A., & Lahiri, S. K. (2026). A Compact CNN-Transformer Model for Robust Fault Diagnosis in Large-Scale Chemical Processes. *International Journal of Prognostics and Health Management*, 17(1), 1–21. <https://doi.org/10.36001/ijphm.2026.v17i1.4744>
- Paul, M. K., & Das, P. (2023). A Comparative Study of Deep Learning Algorithms for Forecasting Indian Stock Market Trends. *International Journal of Advanced Computer Science and Applications*, 14(10), 932–941. <https://doi.org/10.14569/IJACSA.2023.0141098>
- Reis, J., Garcia, R., & Varejão, F. (2026). Enhancing Wagon Fault Detection Using Acoustic Signals and Deep Transfer Learning: From Scratch to Full Fine-Tuning. *International Journal of Prognostics and Health Management*, 17(1), 1–16. <https://doi.org/10.36001/ijphm.2026.v17i1.4707>
- Samuel, J., Murugan, T. K., Govindaraj, L., Balaji, M., SenthilKumar, V., & Sundararajan, S. (2026). Adversarial robust EEG-based brain–computer interfaces using a hierarchical convolutional neural network. *Scientific Reports*, 16(1), 1–18. <https://doi.org/10.1038/s41598-025-34024-0>
- Shi, Y., Liu, T., Yang, S., Di, Y., Li, Y., Yu, W., Huang, Y., Chen, D., & Cui, K. (2026). Diagnosis of meibomian gland dysfunction based on spectral convolutional neural network chip. *Photonix*, 7(1), 1–18. <https://doi.org/10.1186/s43074-026-00246-2>
- Szrama, S. (2026). Turbofan Engine Remaining Useful Life Prediction based on Physics-aware Hybrid Framework and Fatigue Cycles. *International Journal of Prognostics and Health Management*, 17(1), 1–16. <https://doi.org/10.36001/ijphm.2026.v17i1.4733>
- Vijayabalan, D. (2026). An Innovative Algorithm-Assisted Neuroimaging Technique for Calculating Brain Age. *Sains Malaysiana*, 55(3), 589–900. <https://doi.org/10.17576/jsm-2026-5503-18>
- Wu, S., Li, P., Li, K., Guo, Q., & Li, Y. (2027). Environmentally friendly high-performance electronic skin for intelligent human–machine interaction systems with feedback capabilities. *Journal of Materials Science and Technology*, 276, 122–132. <https://doi.org/10.1016/j.jmst.2026.03.077>
- Zeng, J., Chen, W., Chen, W., Wang, Y., & Li, X. (2023). Optimization of Pre-Hospital

First Aid Management Strategies for Patients with Infectious Diseases in Huizhou City Using Deep Learning Algorithm. *Acta Clinica Croatica*, 62(1), 131–140. <https://doi.org/10.20471/acc.2023.62.01.16>