

# Detecting AI-Generated Speech: Challenges and Advances in Audio Forensics

Smar Ahmad Sohail \*, Aqsa Nawaz \*, Muhammad Zohaib \*, Iftikhar Rasheed \*, Umar Fayyaz \*,

\*Department of Information and Communication Engineering, Faculty of Engineering and Technology,  
The Islamia University of Bahawalpur, Punjab, Pakistan

**Abstract**—Recently, people and organizations have been having problems distinguishing real speech from AI-generated speech, making AI voices similar to natural-sounding voices due to voice-generation methods and deep learning developments. This extreme level of naturalness allows for synthesizing an individual target voice and regenerating someone's voice with very high accuracy. Audio forensics is cautious when enhancing speech intelligibility and conducting analysis to establish whether audio recordings are authentic and to determine if any manipulation has been performed on them.

**Index Terms**—Audio Forensics, Voice Processing, Deep-Fake Audio Detection, Voice Authentication, spectrogram, Mel-spectrum, MFCC, Chromagram, CNN.

## I. INTRODUCTION

Forensics refers to the science that uses scientific methods or expertise to investigate crimes or examine evidence that may be presented in a court of law [1]. Digital media forensics is a branch of forensic science that involves ensuring that the digital content is accurate and authentic [2]. Digital audio forensics involves a variety of tasks, including determining the environment or the recording device, identifying speakers from the audio, and assessing the audio content's integrity.

User authentication is the method used to verify an individual's claim to their identity, which determines whether they should be granted or denied access to a device or service [3]. Voice authenticity is a crucial aspect of ensuring secure communication and audio forensics. It involves confirming whether a voice is genuine or not. With advancements in Artificial Intelligence (AI), generating fake voices that sound natural has become exceedingly simple. AI-generated voices can be used in fraudulent activities, impersonation threats, and other malicious actions, making voice-based communication systems unreliable [4].

Recently, voice recognition has been successfully integrated with other biometric methods, such as body surface vibrations [5], touch gestures [6] and a combination of facial recognition, mouse movements, and keystrokes [7] in performing continuous authentication. Many systems, including banking applications, smart speakers, and voice-controlled devices, rely on voice commands. It becomes a severe security risk if these voices are not properly verified. If someone can generate the voice of a trusted user using the AI or play the recording of their voice, then that user would be compromised [8]. A strong

voice authenticity verification system can help us identify fake AI-generated voices.

If someone can generate the voice of a trusted user using the AI or play the recording of his voice [9], then that user would be authenticated, compromising the system and violating the user's privacy. The system also needs to discriminate between the AI-generated voice and recordings from the real voice. There are many techniques available to achieve this task. The first task is using the dataset with real and cloned voices (AI-generated) [8]. Extract the audio features from the audio, such as chromagram, spectrogram, Mel-Spectrum, and MFCC, and make a PNG of the feature for identification using the CNN, adjusting the network properly [10]. The AI-generated voice has some characteristic distortions in the spectrogram [11].

This review seeks to explore the methods and technologies used to verify the authenticity of voices, specifically in identifying AI-generated voices. It discusses different algorithms, techniques, and tools applied for authenticating the voices. The review also discusses the limitations and challenges of these systems, especially with the development of AI in generating highly realistic fake voices.

## II. OVERVIEW OF VOICE AUTHENTICATION

A Voice Authentication System consists of different important parts that work together to authenticate whether a voice is real or AI-generated [12]. These components include **pre-processing, feature extraction, voiceprint database, and matching logic**.

### A. Pre-processing

Since actual audio may contain some inappropriate data, the pre-processing phase handles all the preparation work needed to make the input capture audio suitable for implementation in a suggested Verification phase and to provide generalization for the system to become suitable for any input audio. This stage consists of three steps: audio capturing, rate conversion, and normalization [13].

### B. Feature Extraction

This process gathers important features from a voice, which include things like the tone, pitch, and rhythm of the voice. Features must capture essential and distinguishing physical characteristics of a class of signals accurately. Essentially, they are a representation of these characteristics in the form of

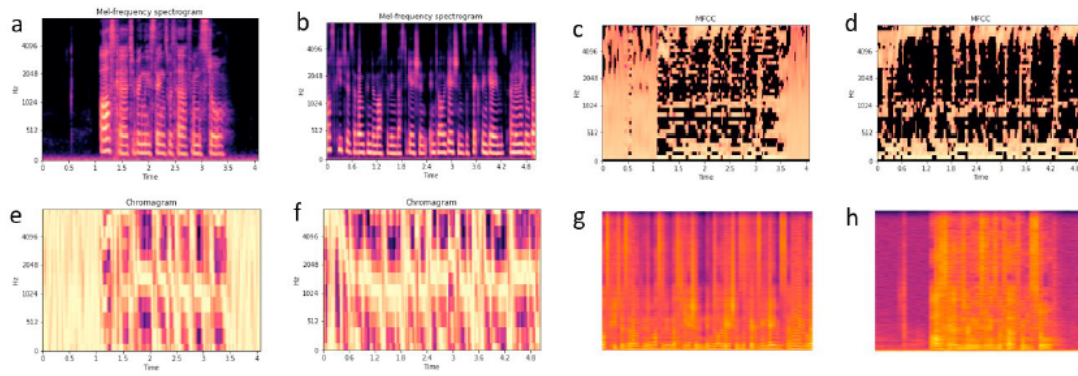


Fig. 1. (a) original Mel-freq audio, (b) cloned Mel-freq audio, (c) original MFCC audio, (d) cloned MFCC audio, (e) original chromagram audio, (f) cloned chromagram audio, (g) original spectrogram audio, (h) cloned spectrogram audio. [10]

values that can be manipulated mathematically [14]. These features are unique to each voice. Features must capture essential and distinguishing physical characteristics of a class of signals accurately. Essentially, they are a representation of these characteristics in the form of values that can be manipulated mathematically [15], [16].

Figure 1 shows that deep learning models consist of extracting audio features utilizing the Libros Python library to extract the Mel-Spectrograms, MFCC, Spectrogram, and Chromagram of every audio file and then converting them into PNG format images from every audio file. The images obtained were utilized to train selected deep learning models, including FG-LCNN [17], ResNet [18], VGG-16 [19], and a Custom Architecture [11]. [10]

### C. Datasets

While some work has shown the potential for detecting spoofing without prior knowledge or training data indicative of a specific attack [20], all previous work is based on some implicit prior knowledge, i.e. the nature of the spoofing attack and the targeted ASV system is known [21]. The Fake-or-Real (FoR) dataset includes more than 195,000 utterances from real humans and computer-generated speech. The dataset consists of 13,268 audio files, which can be used to train, validate, and test classifiers to detect synthetic speech and verify real speech. This dataset is divided into training, validation, and testing parts with 10,614, 2244, and 816 audio files, respectively. The second dataset used for evaluation of the proposed system and comparing it with other related works is ASVspoof 2019 dataset, organized in 2019, and is a collection of audio data designed for the Automatic Speaker Verification Spoofing (ASVspoof) challenge [13], [22].

### D. Decision Logic

Decision logic is the main component in the voice authenticity verification system. Once features are extracted from an uploaded voice sample, the system compares them with stored voiceprints in the database using algorithms CNNs, 1D CNN, RNN, and LSTM. Convolutional neural networks (CNN) serve as classifiers to ascertain whether the audio is

genuine or reproduced [23]. Implementing the 1D convolution operation allows the model to learn local patterns in the input sequence efficiently [13]. LSTM can exploit temporal dependencies across audio sequences, which leads it to detect very slight temporal patterns for spoofing attacks [24], [25].

## III. EXISTING SYSTEM

Discussed all the methods in [10] and [26] to identify the integrity of the audio file. It discusses how to distinguish between the original and deepfake voice. The current system either extracts audio features from and converts them into PNG and feeds them to the convolutional neural net, or it utilizes machine learning for classification purposes using the various extracted features. The audio extracted, such as MFCC, chromagram, and Mel-spectrogram, when combined with CNN, gives the best performance. [8] Existing systems can be trained using various datasets; some systems use ASVspoof [17] or VCTK speakers [11], achieving an accuracy of 100% with the VCTK speakers dataset [27]. Uses the system for voice recognition using linear Predictive Coding and a GMM model. The existing system focuses either on deepfake identification or can be used for voice recognition. The system for deepfake identification is useful in Digital Forensics. [8]

## IV. TECHNIQUES USED FOR VOICE AUTHENTICATION

The following techniques are used for Voice authentication:

### A. CNN

Convolutional Neural Networks (CNNs) are widely used in many applications due to their effectiveness in processing and analyzing sequences of data like time series [13]. CNNs are excellent at extracting spatial features from spectrograms, which are visual representations of audio signals. This is employed in the field of voice authentication [31].

Voice signals, also referred to as waveforms, are transformed into spectrograms of Mel Frequency Cepstral Coefficients (MFCCs). These spectrograms serve as visual representations of frequencies and amplitudes over time. Convolutional Neural Network (CNN) is utilized to analyze these spectrograms, facilitating the extraction of spatial features,

TABLE I  
SUMMARY OF VOICE AUTHENTICATION AND AI GENERATED VOICE DETECTION METHODS

| Ref. | Dataset                                                                                                                                         | Algorithm & Model                                            | Accuracy     | Purpose                                                                                                                                              | Strengths                                                                          | Challenges                                                              |
|------|-------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------|--------------|------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------|-------------------------------------------------------------------------|
| [27] | Data was collected from 30 respondents (20 men, 10 women).                                                                                      | Gaussian Mixture Model (GMM), Linear Predictive Coding (LPC) | 82%          | To test the model using voice recordings made in different environments (such as homes, labs, and classrooms)                                        | GMM and LPC provide a dependable and simple method for voice-based authentication. | Recordings without earphones showed lower accuracy due to higher noise. |
| [10] | The dataset from the Baidu Silicon Valley AI Lab. 10 original audios, 120 cloned, 4 altered audios                                              | CNN, VGG-16, FG-LCNN, ResNet                                 | 100%         | It compares deep learning models for identifying deepfake audio to help digital forensic investigations.                                             | Provides a comparative performance baseline between varying CNN models.            | Not all forensic requirements have been satisfied                       |
| [11] | The dataset from the Baidu Silicon Valley AI Lab. 124 cloned + 4 morphed samples                                                                | CNN                                                          | 100%         | To protect biometric systems and voice-driven interfaces (like Alexa and Siri) from AI-generated fake audio attacks and detect cloned speech voices. | Identifies characteristic artifacts cloned speech spectrograms.                    | Evolving cloning technology is major challenge.                         |
| [17] | The ASVspoof 2019, Train: 2,580 bonafide, 22,800 spoofed                                                                                        | Light CNN using Log Power Spectrum (LPS)                     | —            | To identify Synthetic speech attacks* (VC and TTS) that compromise speaker verification systems.                                                     | The CNN is light and has a compact design.                                         | High variability of spoofing attacks by TTS and VC systems.             |
| [28] | The ASVspoof 2019, Size 80k audio samples including Bonafide and spoof audio samples                                                            | EfficientCNN and RES-EfficientCNN (Residual version)         | 97.61%       | Create effective and flexible models for detecting fake speech.                                                                                      | Very low parameter count (less than 50,000) and high accuracy                      | It is challenging to generalize to unknown attack vectors.              |
| [29] | 62 minutes of real speech from 8 celebrities, Fake speech generated using Retrieval-based Voice Conversion (RVC)                                | XGBoost, RF, SVM                                             | 99.3%, 98.9% | Identify AI-generated speech in real time in order to stop abuses such as impersonation and privacy violations.                                      | High precision (up to 99.3)<br>Real-time detection (inference time of 0.004 ms)    | Real-time applications require lightweight models.                      |
| [30] | The ASVspoof 2017                                                                                                                               | CNN+RNN Hybrid                                               | —            | Identify replay attacks (non-professional spoofing) in speaker verification systems (banking, call centers, etc.).                                   | Deep learning (LCNN) performed better than classical techniques.                   | CNN+RNN needed smaller input dimensions.                                |
| [26] | The data set from ASVspoof2015, ASVspoof2017, ASVspoof2019, Fake-or-Real (FoR) Datasets contains over 195,000 real and synthetic audio samples. | SVM, Random Forest, XGBoost, MLP.                            | 98%          | To identify deepfake audio through machine learning and deep learning methods.                                                                       | High accuracy with different datasets.                                             | Training time and computational resources needed for deep models.       |

including pitch, voice timbre, and other distinctive attributes of an individual's voice. CNNs are particularly effective at identifying patterns, such as voiceprints, within images, there by enhancing the authentication process. ‘

### B. RNNs

A type of neural network called a recurrent neural network (RNN) uses the output from the preceding step as the input for the current step. All of the inputs and outputs of conventional neural networks are independent of one another [32]. RNN was created, and it used a Hidden Layer to tackle this problem [32]. It does the same task on all inputs or hidden layers to produce the output, using the same settings for each input. Unlike other neural networks, this lowers the parameter complexity [33]. Recurrent Neural Network design is depicted in Figure 2 [33]. This is used in voice authentication.

This is used in voice authentication.

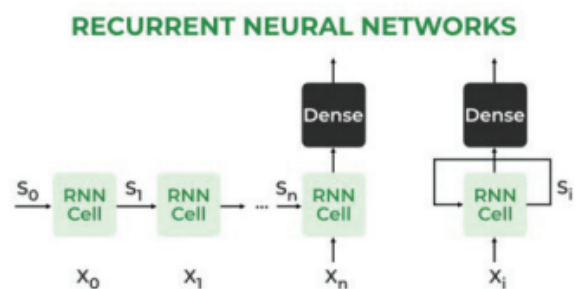


Fig. 2. Architecture of RNN algorithm

RNNs examine temporal dependencies, such as how your tone or pitch changes over time, using sequential audio features (MFCCs, spectrogram slices, etc.). They record patterns

over time, such as rhythm, intonation, and speaking speed, which are crucial for speaker verification.

### C. LSTM(Long Short-Term Memory Network)

LSTMs are an advanced version of RNNs. They are designed to manage long-term dependencies and avoid issues with vanishing gradients. This is used in voice authentication, as follows.

LSTMs are typically used to capture the temporal evolution of audio signals. They learn how a person's voice changes over longer periods, capturing speaker-dependent trends in prosody and speech rhythm. LSTMs can monitor whether the speech sequence matches the registered speaker and whether it matches the expected phrase in text-dependent voice verification, which requires the user to say a specific phrase.

## V. APPLICATIONS

Voice authenticity verification system has a wide range of applications across most fields. These systems play a very important role in enhancing the security of voice-based systems. Here are some key applications:

### A. Banking and Financial Services

Used for secure authentication of mobile banking and phone-based transactions. It helps prevent fraud by confirming that the voice belongs to the authorized account holder [34], [35].

### B. Smart Devices

Make sure that only the smart device's owner can issue commands. Prevents misuse or illegal access to personal information [34].

### C. Secure Communications

By verifying the identity of calls and voices, this feature protects confidential business communications. Stops phishing and impersonation attempts [36].

### D. Law Enforcement and Forensics

Helps verify voice recordings used as evidence in court cases or investigations. In criminal investigations, it detects deepfake AI-generated audio [37].

## VI. THREATS AND CHALLENGES

Voice authentication systems have many different threats and challenges due to the development of technology and attack methods. These are the major threats and challenges:

### A. Spoofing Attacks (Voice Impersonation)

Attackers use recorded or AI-generated voices to bypass the authentication system. In a recent well-known episode, malicious actors used AI to impersonate a German chief executive's voice on the phone, obtaining a fraudulent money transfer of 220,000 dollars [21], [38], [39].

### B. Replay Attacks

Attackers trick the system with pre-recorded audio snippets. An authorized user's pre-recorded audio message was used to enter a secure banking system in a security breach [40].

### C. Environmental Noise and Poor Audio Quality

Background noise, echoes, or low-quality microphones can lower the performance of the system [41]. In public places, voice authentication systems in mobile phones fail to verify users accurately.

### D. Privacy Concerns and Data Security

Voiceprint storage and processing can result in privacy threats if the database is compromised or leaked. It can cause serious privacy threats to the users.

## VII. LIMITATION

Getting greater accuracy is a challenging task. The presence of background noises leads to lower accuracy. [8] Most of the systems are used for deep fake voice detection and can be extended for the use of voice recognition. [8] People with strong accents can experience challenges with voice authentication systems, resulting in exclusion.

## VIII. CONCLUSION

Voice authenticity verification systems now play an important countermeasure to the misuse of AI-created voices, playing a critical role in securing voice-based technologies. The review has shown the dramatic methodological and algorithmic advances, e.g., CNNs, RNNs, and LSTMs, towards distinguishing real and AI-created voices. However, with the remarkable improvements, there are still some issues that involve susceptibility to external noise, the growing complexity of voice generation technology, as well as computational issues with privacy and efficiency.

In order to overcome such limitations, subsequent research must focus on designing light, responsive models that excel across varied settings and against future threat vectors. Developing such systems to be robust to noise and variation in accent without compromising their real-time applicability for scalability also remains essential. Additionally, the incorporation of multi-modal biometrics along with federated learning has the potential to yield additional security levels without exposing the user's privacy to any concerns.

As AI develops, so will the technology used to exploit it for malicious use. Therefore, sustained innovation, interdisciplinarity, and the establishment of global standards for voice authentication systems are called for to ensure the security and integrity of voice interactions.

## REFERENCES

- [1] M. A. Qamhan, H. Altaheri, A. H. Meftah, G. Muhammad, and Y. A. Alotaibi, "Digital audio forensics: Microphone and environment classification using deep learning," *IEEE Access*, vol. 9, pp. 62 719–62 733, 2021.
- [2] M. Zakariah, M. K. Khan, and H. Malik, "Digital multimedia audio forensics: past, present and future," *Multimedia tools and applications*, vol. 77, pp. 1009–1040, 2018.

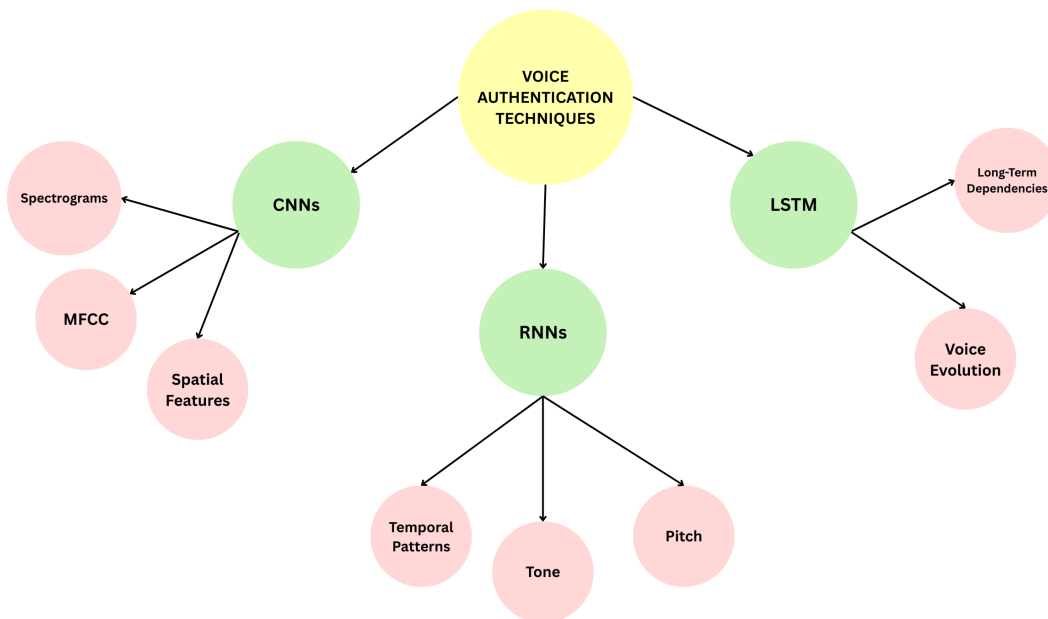


Fig. 3. Summary of Voice Authentication Techniques

- [3] Z. Meng, M. U. B. Altaf, and B.-H. F. Juang, "Active voice authentication," *Digital Signal Processing*, vol. 101, p. 102672, 2020.
- [4] Q. Wang, X. Lin, M. Zhou, Y. Chen, C. Wang, Q. Li, and X. Luo, "Voicepop: A pop noise based anti-spoofing system for voice authentication on smartphones," in *IEEE INFOCOM 2019-IEEE Conference on Computer Communications*. IEEE, 2019, pp. 2062–2070.
- [5] H. Feng, K. Fawaz, and K. G. Shin, "Continuous authentication for voice assistants," in *Proceedings of the 23rd Annual International Conference on Mobile Computing and Networking*, 2017, pp. 343–355.
- [6] G. Peng, G. Zhou, D. T. Nguyen, X. Qi, Q. Yang, and S. Wang, "Continuous authentication with touch behavioral biometrics and voice on wearable glasses," *IEEE transactions on human-machine systems*, vol. 47, no. 3, pp. 404–416, 2016.
- [7] G. Fenu, M. Marras, and L. Boratto, "A multi-biometric system for continuous student authentication in e-learning platforms," *Pattern Recognition Letters*, vol. 113, pp. 83–92, 2018.
- [8] A. Deore, O. Badhakh, S. Agase, S. Dalvi, and P. Futane, "Voice auth: Audio verification with deep-learning," in *2024 8th International Conference on Computing, Communication, Control and Automation (ICCUBEA)*. IEEE, 2024, pp. 1–6.
- [9] Z.-F. Wang, G. Wei, and Q.-H. He, "Channel pattern noise based playback attack detection algorithm for speaker recognition," in *2011 International conference on machine learning and cybernetics*, vol. 4. IEEE, 2011, pp. 1708–1713.
- [10] M. Mcuba, A. Singh, R. A. Ikuesan, and H. Venter, "The effect of deep learning methods on deepfake audio detection for digital investigation," *Procedia Computer Science*, vol. 219, pp. 211–219, 2023.
- [11] H. Malik, "Fighting ai with ai: fake speech detection using deep learning," in *2019 AES INTERNATIONAL CONFERENCE ON AUDIO FORENSICS (June 2019)*, 2019.
- [12] Z. Khanjani, G. Watson, and V. P. Janeja, "Audio deepfakes: A survey," *Frontiers in Big Data*, vol. 5, p. 1001063, 2023.
- [13] N. D. AL-Shakarchy, Z. N. Abdullah, Z. M. Alameen, and Z. A. Harjan, "Audio verification in forensic investigation using light deep neural network," *International Journal of Information Technology*, vol. 16, no. 5, pp. 2813–2821, 2024.
- [14] S. J. Joysingh, P. Vijayalakshmi, and T. Nagarajan, "Significance of chirp mfcc as a feature in speech and audio applications," *Computer Speech & Language*, vol. 89, p. 101713, 2025.
- [15] B. Peng, K. I.-K. Wang, and W. H. Abdulla, "Robust classification of urban sounds in noisy environments: A novel approach using spwvd-mfcc and dual-stream classifier," *Acoustics Australia*, pp. 1–16, 2025.
- [16] P. L. De Leon, M. Pucher, J. Yamagishi, I. Hernaez, and I. Saratzaga, "Evaluation of speaker verification security and detection of hmm-based synthetic speech," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 20, no. 8, pp. 2280–2290, 2012.
- [17] Z. Wu, R. K. Das, J. Yang, and H. Li, "Light convolutional neural network with feature genuinization for detection of synthetic speech attacks," *arXiv preprint arXiv:2009.09637*, 2020.
- [18] K. Chugh, P. Gupta, A. Dhall, and R. Subramanian, "Not made for each other-audio-visual dissonance-based deepfake detection and localization," in *Proceedings of the 28th ACM international conference on multimedia*, 2020, pp. 439–447.
- [19] R. Reimao and V. Tzerpos, "For: A dataset for synthetic speech detection," in *2019 International Conference on Speech Technology and Human-Computer Dialogue (SpED)*. IEEE, 2019, pp. 1–10.
- [20] Z. Wu, E. S. Chng, and H. Li, "Detecting converted speech and natural speech for anti-spoofing attack in speaker recognition," in *13th Annual Conference of the International Speech Communication Association (Interspeech 2012)*, 2012.
- [21] N. W. Evans, T. Kinnunen, and J. Yamagishi, "Spoofing and countermeasures for automatic speaker verification," in *INTERSPEECH 2013, 14th Annual Conference of the International Speech Communication Association*, 2013.
- [22] Z. Wu, J. Yamagishi, T. Kinnunen, C. Haniłçi, M. Sahidullah, A. Sizov, N. Evans, M. Todisco, and H. Delgado, "Asvspoof: the automatic speaker verification spoofing and countermeasures challenge," *IEEE Journal of Selected Topics in Signal Processing*, vol. 11, no. 4, pp. 588–604, 2017.
- [23] K. A. Al-Karawi, "Convolutional neural network-based detection of audio replay attacks in speaker verification systems," *International Journal of Speech Technology*, pp. 1–10, 2025.
- [24] K. Shahzad, S. Farhan, Y. U. Haq, R. Sana, and M. S. Pathan, "En-

- hancing voice spoofing detection: A hybrid approach with vggish-lstm model for improved security in automatic speaker verification systems,” *IEEE Access*, 2025.
- [25] T. Kinnunen, M. Sahidullah, H. Delgado, M. Todisco, N. Evans, J. Yamagishi, and K. A. Lee, “The asvspoof 2017 challenge: Assessing the limits of replay spoofing attack detection,” 2017.
  - [26] A. Hamza, A. R. R. Javed, F. Iqbal, N. Kryvinska, A. S. Almadhor, Z. Jalil, and R. Borghol, “Deepfake audio detection via mfcc features using machine learning,” *IEEE Access*, vol. 10, pp. 134 018–134 028, 2022.
  - [27] N. J. Perdana, D. E. Herwindiati, and N. H. Sarmin, “Voice recognition system for user authentication using gaussian mixture model,” in *2022 IEEE International Conference on Artificial Intelligence in Engineering and Technology (IICAET)*. IEEE, 2022, pp. 1–5.
  - [28] N. Subramani and D. Rao, “Learning efficient representations for fake speech detection,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 34, no. 04, 2020, pp. 5859–5866.
  - [29] J. J. Bird and A. Lotfi, “Real-time detection of ai-generated speech for deepfake voice conversion,” *arXiv preprint arXiv:2308.12734*, 2023.
  - [30] G. Lavrentyeva, S. Novoselov, E. Malykh, A. Kozlov, O. Kudashev, and V. Shchemelinin, “Audio replay attack detection with deep learning frameworks,” in *Interspeech*, 2017, pp. 82–86.
  - [31] T. Liu, D. Yan, R. Wang, N. Yan, and G. Chen, “Identification of fake stereo audio using svm and cnn,” *Information*, vol. 12, no. 7, p. 263, 2021.
  - [32] K. D. Bhargavi, V. Deepa, and I. Harshitha, “Audio classification through mfcc features using rnn algorithm,” in *2024 4th Asian Conference on Innovation in Technology (ASIANCON)*. IEEE, 2024, pp. 1–5.
  - [33] A. Aishwary, “Introduction to recurrent neural network,” *Retrieved from geeksforgeeks: <https://www.geeksforgeeks.org/introduction-to-recurrent-neural-network>*, 2023.
  - [34] F. Afandi and R. Sarno, “Android application for advanced security system based on voice recognition, biometric authentication, and internet of things,” in *2020 International Conference on Smart Technology and Applications (ICoSTA)*. IEEE, 2020, pp. 1–6.
  - [35] K. A. Kamiński, A. P. Dobrowolski, Z. Piotrowski, and P. Ścibiorek, “Enhancing web application security: Advanced biometric voice verification for two-factor authentication,” *Electronics*, vol. 12, no. 18, p. 3791, 2023.
  - [36] S. Safavi, H. Gan, I. Mporas, and R. Sotudeh, “Fraud detection in voice-based identity authentication applications and services,” in *2016 IEEE 16th international conference on data mining workshops (ICDMW)*. IEEE, 2016, pp. 1074–1081.
  - [37] M. Britz, *Computer forensics and cyber crime: An introduction*, 2/e. Pearson Education India, 2009.
  - [38] A. Pianese, D. Cozzolino, G. Poggi, and L. Verdoliva, “Training-free deepfake voice recognition by leveraging large-scale pre-trained models,” in *Proceedings of the 2024 ACM Workshop on Information Hiding and Multimedia Security*, 2024, pp. 289–294.
  - [39] Z. Wu, T. Kinnunen, N. Evans, J. Yamagishi, C. Hanilçi, M. Sahidullah, and A. Sizov, “Asvspoof 2015: the first automatic speaker verification spoofing and countermeasures challenge,” in *INTERSPEECH 2015, Automatic Speaker Verification Spoofing and Countermeasures Challenge, colocated with INTERSPEECH 2015*. ISCA, 2015, pp. 2037–2041.
  - [40] Z. Wu, S. Gao, E. S. Cling, and H. Li, “A study on replay attack and anti-spoofing for text-dependent speaker verification,” in *Signal and Information Processing Association Annual Summit and Conference (APSIPA), 2014 Asia-Pacific*. IEEE, 2014, pp. 1–5.
  - [41] M. A. Qamhan, H. Altaheri, A. H. Meftah, G. Muhammad, and Y. A. Alotaibi, “Digital audio forensics: microphone and environment classification using deep learning,” *Ieee Access*, vol. 9, pp. 62 719–62 733, 2021.