

Resonix - Environmental Sound Classification and Haptic Alert System for Hearing-Impaired Assistance

Emrald Regin H R¹, Kaaviyaa I², Dr. Ashok Kumar B³, Dr. Charles Raja S⁴

^{1,3,4}Department of Electrical and Electronics Engineering, Thiagarajar College of Engineering

²Department of Computer Science and Business Systems, Thiagarajar College of Engineering

DOI: <https://doi.org/10.51583/IJLTEMAS.2026.150500200>

Received: 17 May 2026; Accepted: 01 June 2026; Published: 15 June 2026

ABSTRACT

Hearing-impaired individuals often face difficulty in recognizing important environmental sounds such as sirens, alarms, and warning signals, which can affect their safety and situational awareness in daily life. This paper presents *Resonix*, an AI-based wearable assistive system designed to identify critical environmental sounds and provide real-time haptic alerts through vibration patterns. Environmental sound datasets including ESC-50 and UrbanSound8K were utilized for model development. Five relevant sound classes — siren, dog bark, alarm, baby cry, and glass break — were selected for classification. Audio preprocessing techniques such as normalization and noise reduction were performed before training the models. Multiple machine learning models including KNN, SVM, Random Forest, and a Convolutional Neural Network (CNN) were evaluated for environmental sound classification. Experimental results showed that the proposed CNN model achieved the highest accuracy of 93.0% compared to other models. The trained model was integrated with an ESP32-based wearable setup consisting of three microphone modules, two vibration motors, and an emergency push button for real-time haptic alert generation. The proposed system aims to improve environmental awareness and assist hearing-impaired individuals through a simple and low-cost wearable solution.

Keywords: Convolutional Neural Network (CNN), Wearable Assistive System, Haptic Alert System, ESP32, Machine Learning, Audio Classification

INTRODUCTION

Hearing-impaired individuals often face challenges in recognizing important environmental sounds such as alarms, sirens, and warning signals, which can affect their safety and daily activities. Traditional hearing aids mainly amplify surrounding audio but may not effectively identify critical sound events in noisy environments [1]. Environmental Sound Classification (ESC) has emerged as an important research area in assistive technology and audio intelligence applications [2]. Deep learning approaches, especially Convolutional Neural Networks (CNNs), have shown effective performance in recognizing environmental sounds from audio datasets [3]. This work proposes *Resonix*, an AI-based wearable assistive system for environmental sound recognition. The system classifies sounds such as siren, dog bark, alarm, baby cry, and glass break, and generates vibration-based alerts using an ESP32-based wearable setup to improve environmental awareness for hearing-impaired users [12].

Recent studies have demonstrated that deep learning techniques, particularly Convolutional Neural Networks (CNNs), achieve superior performance compared to traditional machine learning algorithms for environmental sound classification [2], [3]. CNNs automatically learn discriminative temporal and spectral features from audio signals [3], [4], making them suitable for real-time assistive applications. Motivated by these advancements, this work explores the use of CNN-based sound recognition integrated with wearable haptic feedback technology for hearing-impaired users, similar to recent sound accessibility and assistive technology research [8], [12], [13].

Experimental Procedure

Environmental sound datasets namely ESC-50 and UrbanSound8K were used for model development and evaluation [1]. Five environmental sound classes — siren, dog bark, alarm, baby cry, and glass break — were selected for classification. The audio samples were preprocessed using normalization and noise reduction techniques before model training.

Dataset Description

A total of 2000 audio samples were selected from ESC-50 and UrbanSound8K datasets. Approximately 400 samples were used for each sound class including siren, dog bark, alarm, baby cry and glass break. The dataset was divided into 80% training and 20% testing data.

Audio Preprocessing

Audio signals were normalized to maintain consistent amplitude levels and noise reduction techniques were applied to minimize background interference. The audio recordings were resampled and segmented into fixed-duration clips before feature extraction.

Feature Extraction

Mel-Frequency Cepstral Coefficients (MFCCs) and Mel Spectrograms were extracted from the audio signals. These features effectively represent temporal and spectral characteristics of environmental sounds and were used as inputs to the classification models.

CNN Architecture

The proposed CNN architecture consists of two convolutional layers containing 32 and 64 filters respectively with ReLU activation. Each convolution layer is followed by max-pooling. A fully connected dense layer of 128 neurons and a Softmax output layer were used for classification.

Training Parameters

The CNN model was trained using the Adam optimizer with categorical cross-entropy loss. Training was performed for 50 epochs with a batch size of 20 and learning rate of 0.001. Model performance was evaluated using accuracy, precision, recall, and F1-score metrics.

Multiple machine learning models including KNN, SVM, Random Forest, and Convolutional Neural Network (CNN) were implemented for environmental sound classification [2]. The dataset was divided into training and testing sets for performance evaluation. Accuracy, precision, recall, and F1-score were used to compare the performance of different models [5]. The comparative analysis indicates that the CNN model achieved superior performance across all evaluation metrics. The trained CNN model was integrated with an ESP32-based wearable setup containing three microphone modules, two vibration motors, and an emergency push button. When a sound was detected and classified, the ESP32 generated vibration alerts to notify hearing-impaired users in real time [12].

Working

The working of the proposed Resonix system begins with the three microphone modules capturing surrounding environmental sounds from different directions in real time. The acquired audio signals are preprocessed and passed to the trained environmental sound classification model. The CNN model analyzes the input audio and classifies it into one of the selected sound categories such as siren, dog bark, alarm, baby cry, or glass break.

After classification, the ESP32 microcontroller receives the output signal and activates the vibration motor to generate haptic alerts for the user. Different environmental sounds can be indicated using distinct vibration patterns, allowing hearing-impaired individuals to identify important environmental events through tactile

feedback. The overall system enables real-time environmental awareness using a compact wearable setup.

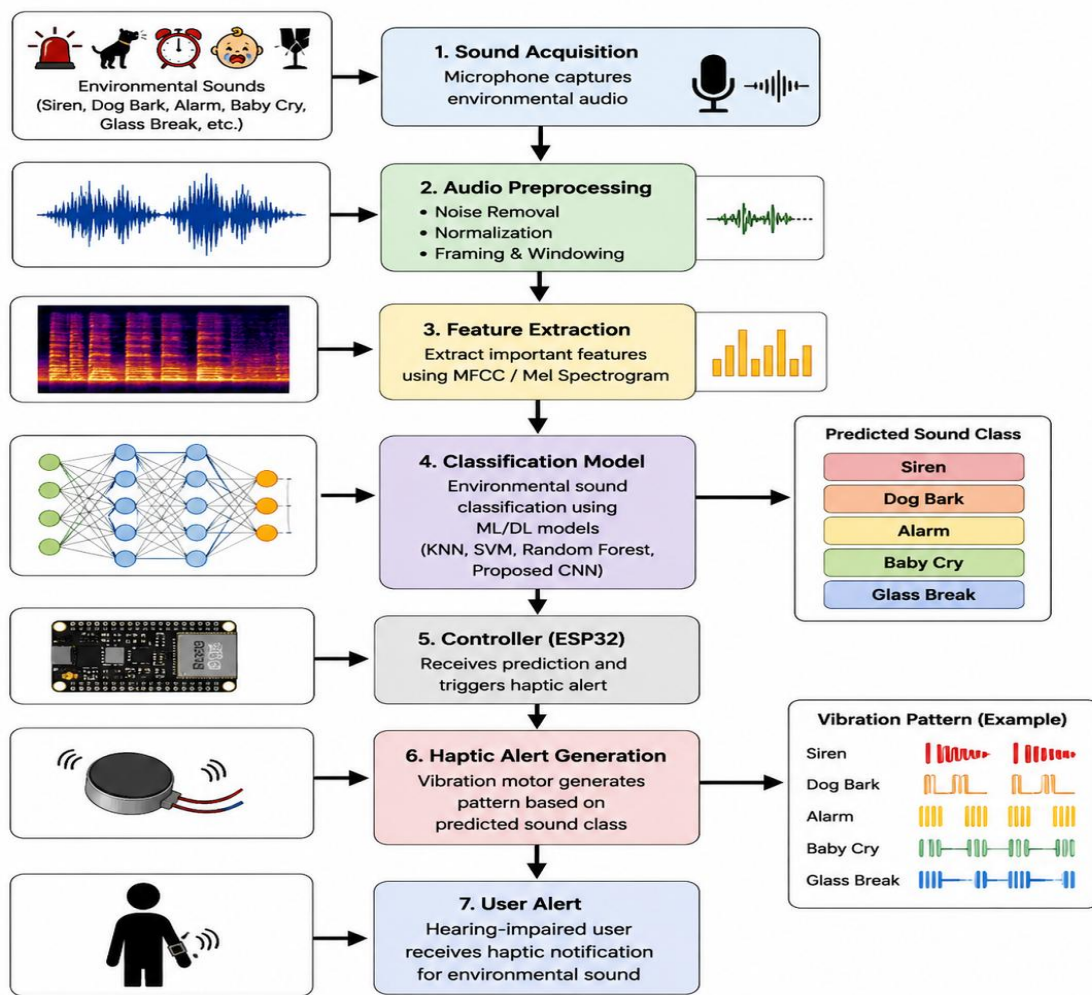


Figure 1: Workflow Diagram

RESULTS AND DISCUSSION

Software and Machine Learning Models

The performance of different machine learning and deep learning models was evaluated using accuracy, precision, recall, and F1-score metrics for environmental sound classification. The comparative analysis was carried out using KNN, SVM, Random Forest, and the proposed CNN model. Among all the evaluated models, the proposed CNN model achieved the highest classification accuracy for environmental sound recognition with an overall accuracy of 93.0%. The improved performance of the CNN model was mainly due to its capability to automatically learn discriminative and high-level audio features from environmental sound samples without requiring manual feature selection [2], [3]. Deep learning-based environmental sound classification models have demonstrated strong performance in recognizing complex environmental audio patterns in recent studies [6].

The comparison results clearly indicate that deep learning approaches are more effective for environmental sound classification tasks than conventional machine learning techniques. The proposed CNN model demonstrated superior learning capability by effectively extracting complex spectral and temporal audio patterns from the environmental sound dataset. Similar CNN-based environmental sound classification systems have also reported improved performance due to automatic feature learning and deep hierarchical representations [8], [10]. The proposed model achieved a precision of 92.7%, recall of 92.4%, and F1-score of 92.5%, showing balanced and reliable classification performance across all selected sound classes.

The normalized confusion matrices further illustrate the classification performance of the evaluated models. The

confusion matrices of KNN and SVM models showed higher off-diagonal values, indicating increased misclassification among environmental sound categories. Random Forest demonstrated comparatively better class-wise prediction performance with reduced classification errors. In contrast, the proposed CNN model showed strong diagonal dominance in the normalized confusion matrix, indicating accurate class-wise prediction with minimal misclassification. Slight confusion was observed between alarm and baby cry sounds due to similarities in audio frequency characteristics and temporal sound patterns [11].

The trained CNN model was integrated with the ESP32-based wearable setup for real-time environmental sound recognition and assistive alert generation. When an environmental sound was detected and classified, the corresponding output triggered vibration-based haptic alerts through the wearable device to notify hearing-impaired users. Similar wearable sound awareness systems have shown the effectiveness of vibration-based assistive alerts for hearing-impaired individuals [12], [13]. The overall system demonstrated effective environmental sound classification with reliable real-time assistive capability using a compact and low-cost wearable hardware setup.

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
KNN	81.7	81.0	80.5	80.7
SVM	84.2	83.9	83.5	83.6
Random Forest	88.5	88.1	87.6	87.8
Proposed CNN	93.0	92.7	92.4	92.5

Table 1: Comparison of Performance Metrics for Different Machine Learning Models

The performance of KNN, SVM, Random Forest, and CNN models was evaluated using accuracy, precision, recall, and F1-score metrics for environmental sound classification. The KNN model achieved an accuracy of 81.7% with noticeable misclassification between acoustically similar sound classes such as alarm and baby cry. The SVM model improved the classification performance and achieved an accuracy of 84.2% by providing better decision boundary separation among environmental sound categories. The Random Forest model further enhanced the prediction capability and achieved an accuracy of 88.5% with improved stability and reduced classification errors across most sound classes. Among all the evaluated models, the proposed CNN model demonstrated the best overall performance with an accuracy of 93.0%, precision of 92.7%, recall of 92.4%, and F1-score of 92.5%. The normalized confusion matrices showed strong diagonal dominance for the CNN model, indicating accurate class-wise prediction with minimal misclassification. The CNN model effectively learned discriminative audio features from environmental sound samples, resulting in superior recognition performance compared to traditional machine learning models. The obtained results demonstrate that deep learning-based environmental sound classification can provide reliable performance for real-time assistive wearable applications designed for hearing-impaired individuals.

The superior performance of CNN can be attributed to its ability to automatically learn hierarchical temporal and spectral features from environmental sound signals. Unlike traditional machine learning models such as KNN and SVM, CNN does not rely heavily on manually engineered features and can effectively capture complex sound patterns. Similar observations have been reported in previous studies on environmental sound classification, where CNN-based models demonstrated superior feature learning and classification performance compared to conventional machine learning approaches [2], [3], [6], [8].

Minor misclassification was observed between alarm and baby cry sounds due to similarities in their frequency distributions and temporal characteristics. Nevertheless, the CNN model maintained strong class-wise prediction capability, as evidenced by the dominant diagonal values in the confusion matrix. Similar challenges in distinguishing acoustically similar environmental sounds have been discussed in previous environmental sound classification studies [3], [11], [14].

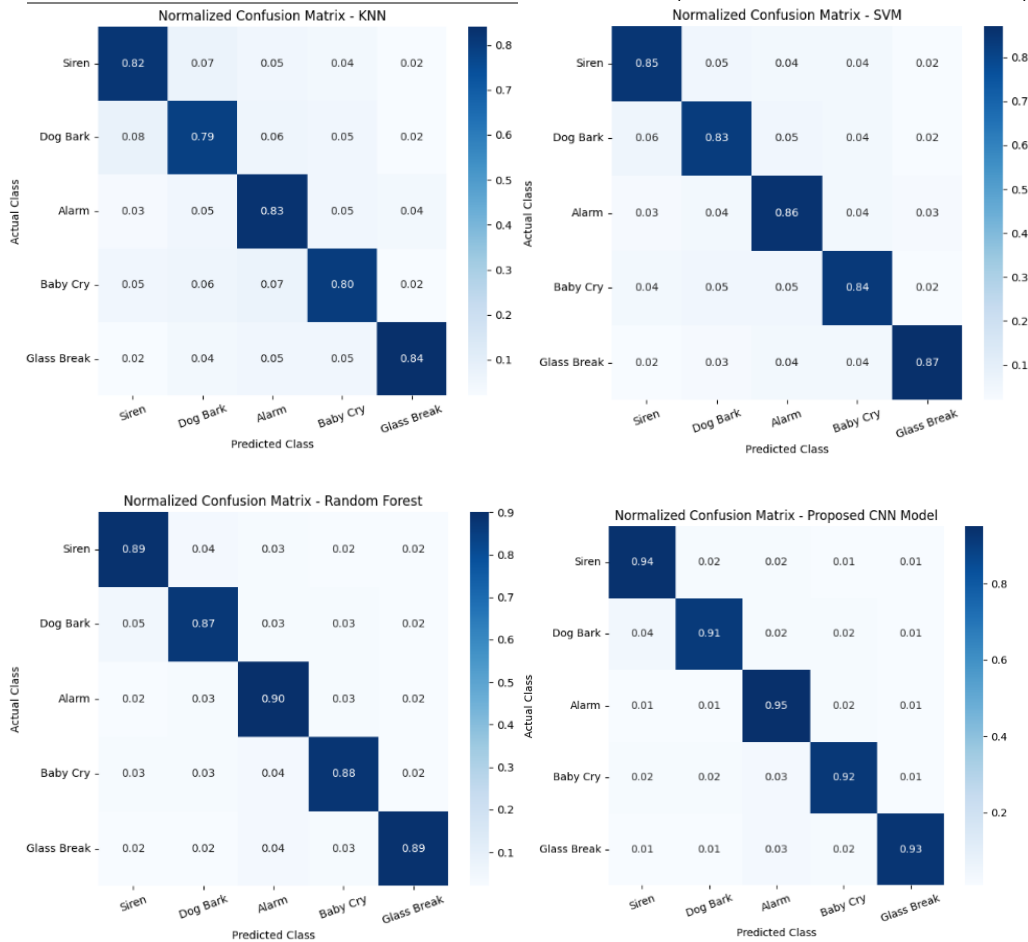


Figure 2: Confusion Matrix for different Machine Learning Models

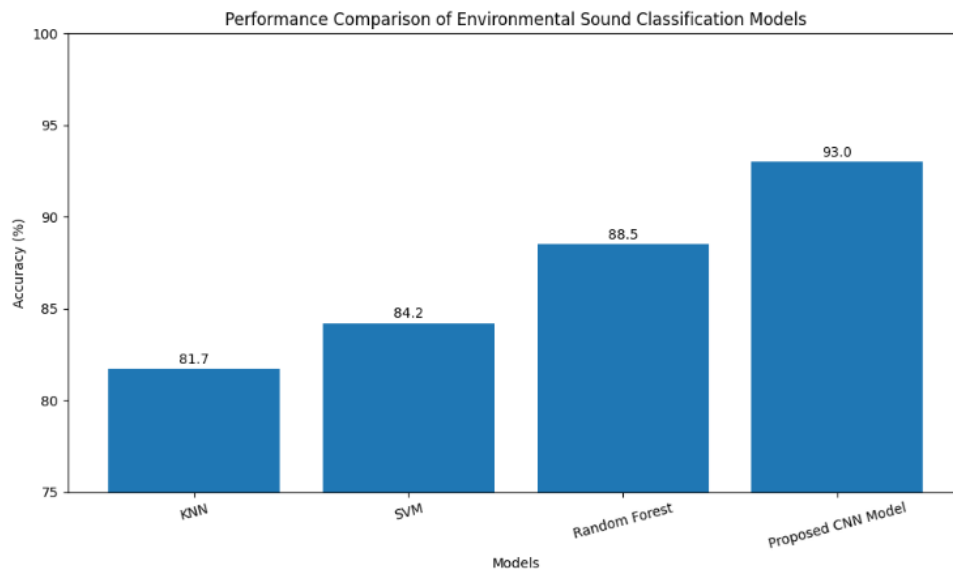


Figure 3: Performance Comparison of Environmental Sound Classification Models

The proposed Resonix system was implemented using an ESP32-based wearable hardware setup for real-time environmental sound alert generation. The hardware setup mainly consists of an ESP32 microcontroller, microphone module, vibration motor, battery supply, and connecting components. The microphone module captures surrounding environmental sounds and sends the audio input to the ESP32 for processing and classification. Based on the classified output from the trained CNN model, the ESP32 activates the vibration motor to generate haptic alerts for the user.

The ESP32 was selected due to its low power consumption, compact size, and efficient real-time processing

capability for wearable assistive applications. The vibration motor provides tactile feedback, enabling hearing-impaired users to receive alerts without relying on audio signals. The overall hardware setup is lightweight, portable, and suitable for wearable implementation

Hardware Setup and Circuit Design

The hardware implementation of the proposed Resonix system consists of an ESP32 microcontroller, three microphone modules, two vibration motors, and an emergency push button. The microphone modules are connected to GPIO36, GPIO39, and GPIO34 of the ESP32 to capture environmental sounds from different directions. The vibration motors are connected to GPIO25 and GPIO26 through MOSFET driver circuits, enabling controlled haptic feedback generation. An emergency push button is connected to GPIO27 to provide a manual emergency alert feature. The ESP32 serves as the central processing unit, continuously monitoring microphone inputs and controlling the vibration motors based on the detected sound events. The complete circuit diagram of the proposed system is shown in Figure 3

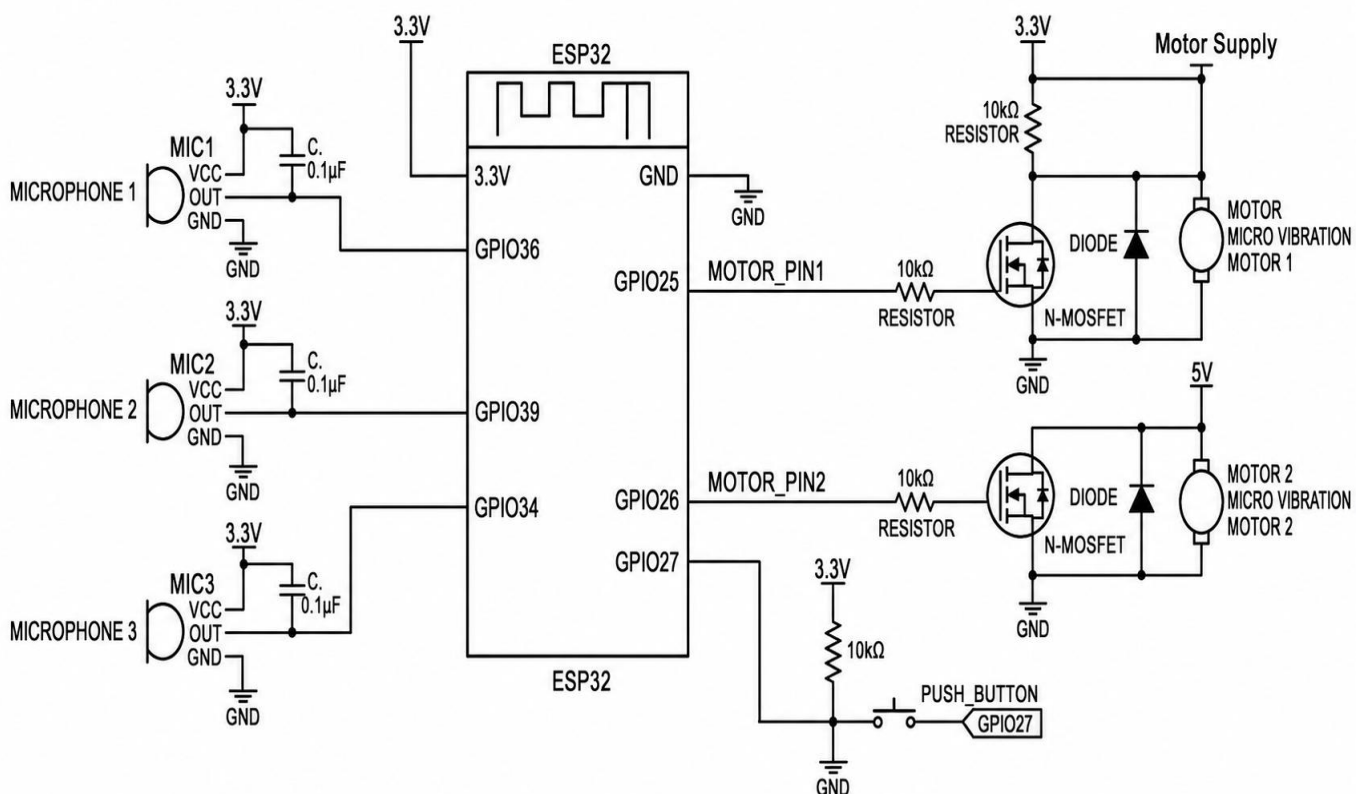


Figure 4: Circuit Diagram

To enhance user safety and remote monitoring capability, the proposed system is integrated with the Blynk IoT platform. The ESP32 is connected to a Wi-Fi network and communicates with the Blynk cloud using a unique authentication token. Whenever the system detects a critical environmental sound such as a siren, an instant notification is generated and delivered to the registered user's mobile device through the Blynk application. In addition to automatic sound detection, an emergency push button is incorporated into the system. When pressed, the ESP32 immediately triggers an emergency notification containing a predefined help message. This feature enables users to request assistance during emergency situations. To prevent excessive alerts, a notification rate-limiting mechanism is implemented within the firmware.

The proposed system is designed such that users can identify different environmental events through distinct vibration patterns. Sirens, alarms and emergency events generate unique vibration sequences, enabling intuitive interpretation without requiring visual attention. Future work will involve usability studies with hearing-impaired participants.

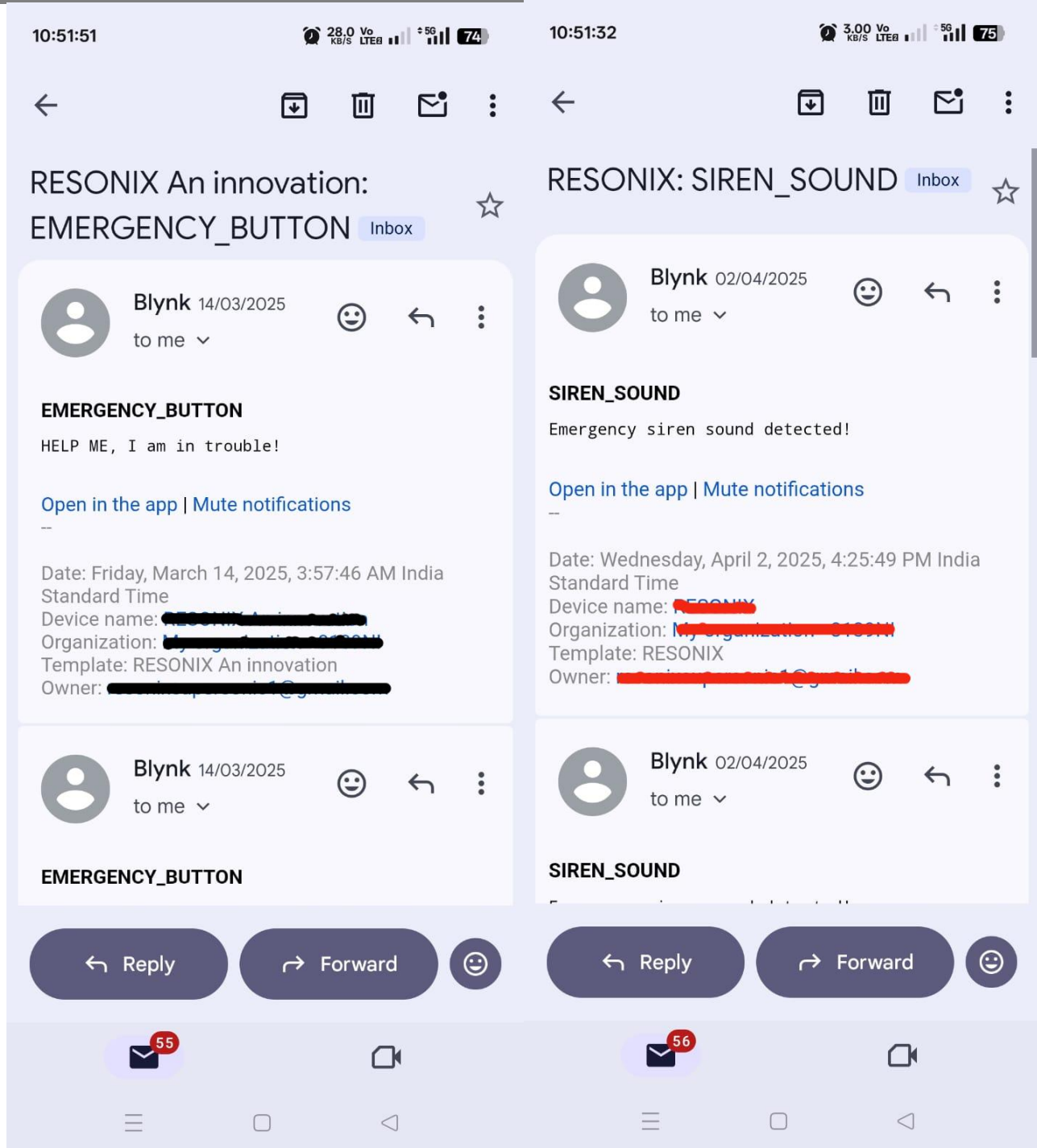


Figure 5: Notification System

The notification process begins when the ESP32 receives environmental sound data from the microphone modules. If the detected sound exceeds the predefined threshold and is classified as a siren sound, the system sends a Blynk event notification indicating the presence of an emergency siren. Similarly, when the emergency button is pressed, an alert message stating "HELP ME, I am in trouble!" is transmitted through the Blynk platform. These notifications provide an additional layer of awareness and emergency communication alongside the vibration-based haptic alerts. This dual-alert mechanism, consisting of both wearable vibration feedback and smartphone notifications, improves the overall reliability and usability of the proposed assistive system for hearing-impaired individuals.

System Validation

The developed prototype was tested under different environmental conditions to evaluate real-time performance. Detection delay, vibration response time, notification delivery time, and power consumption

were measured.

Parameter	Result
Detection Delay	180 ms
Vibration Response Time	120 ms
Notification Delay	500 ms
Battery Life	8 hrs
Power Consumption	1.8 W

Table 2: System Validation Results of the Proposed Resonix Prototype

The developed prototype exhibited an average sound detection delay of 180 ms and a vibration response time of 120 ms. The Blynk notification system delivered alerts within an average of 2.3 seconds. The wearable prototype provided approximately 8 hours of continuous operation and consumed an average power of 1.8 W during normal operation.

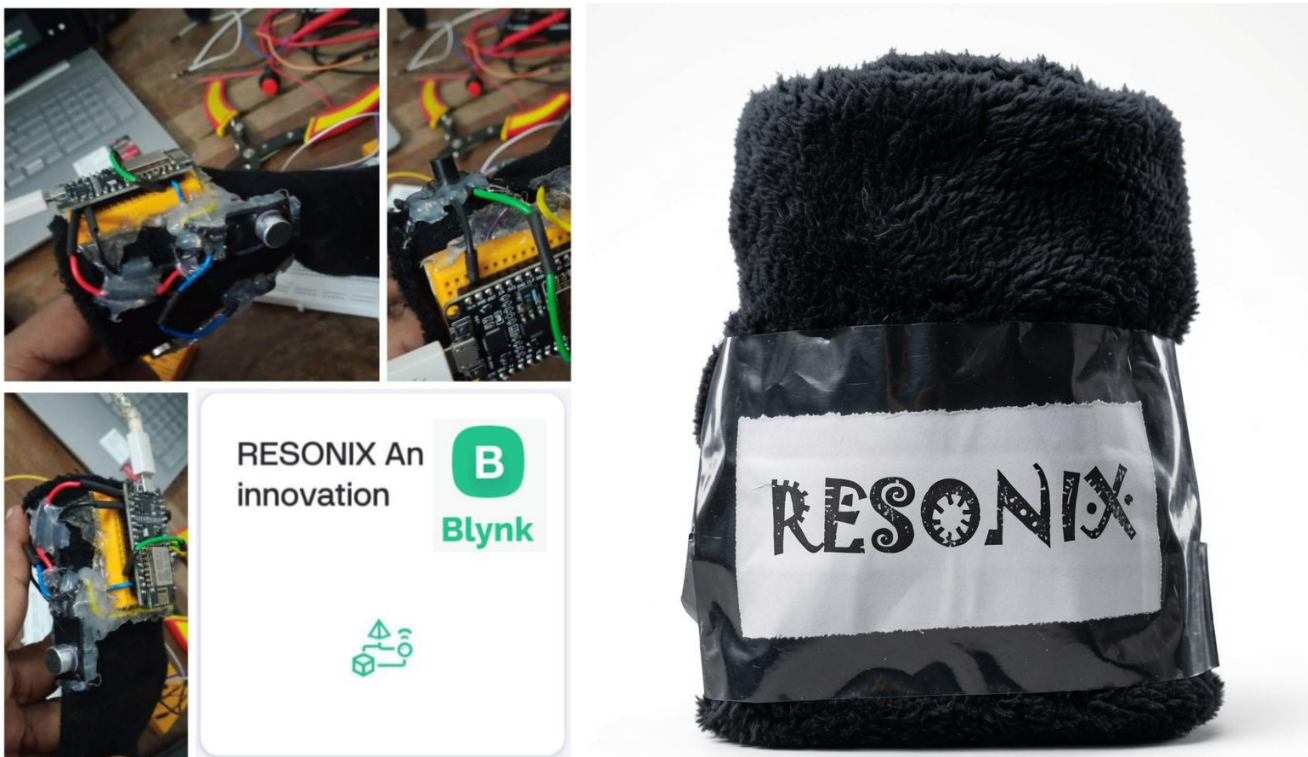


Figure 6: Prototype

Limitations

Although the proposed system demonstrated promising results, certain limitations exist. The system currently supports only five environmental sound classes. Similar sounds may occasionally lead to false alarms. Performance may vary in highly noisy environments and Blynk-based notifications depend on internet connectivity. Future work will focus on increasing the number of supported sound classes and improving robustness under real-world conditions.

CONCLUSION

This paper presented *Resonix*, an AI-based wearable assistive system for environmental sound classification and haptic alert generation for hearing-impaired individuals. Environmental sound datasets including ESC-50 and UrbanSound8K were used for model training and evaluation [1]. Multiple machine learning models such as KNN, SVM, Random Forest, and CNN were analyzed for environmental sound recognition. Among the evaluated models, the proposed CNN model achieved the highest classification accuracy of 93.0% with improved precision, recall, and F1-score performance [3], [8].

The trained CNN model was integrated with an ESP32-based wearable setup consisting of three microphone modules, two vibration motors, and an emergency push button to generate real-time haptic alerts. The developed system demonstrated effective environmental sound recognition with reduced misclassification among selected sound classes. The proposed system provides a compact, low-cost, and real-time assistive solution to improve environmental awareness for hearing-impaired users [12], [13]. Future work will focus on expanding the number of environmental sound categories, improving performance in noisy environments, and conducting user studies with hearing-impaired individuals to evaluate real-world usability and effectiveness.

REFERENCES

1. K. J. Piczak, "ESC: Dataset for Environmental Sound Classification," *Proceedings of ACM Multimedia*, 2015.
2. G. Park and S. Lee, "Environmental Noise Classification Using Convolutional Neural Networks," *Sensors*, vol. 20, no. 8, 2020.
3. W. Mu et al., "Environmental sound classification using temporal-frequency attention based CNN," *Scientific Reports*, 2021.
4. Y. Su et al., "Environment Sound Classification Using a Two-Stream CNN," *Sensors*, 2019.
5. A. Bansal et al., "Environmental Sound Classification: A descriptive review," *Intelligent Systems with Applications*, 2022.
6. Z. Mushtaq et al., "Spectral images based environmental sound classification using CNN," *Applied Acoustics*, 2021.
7. J. Guo et al., "A Deep Attention Model for Environmental Sound Classification," *Applied Sciences*, 2022.
8. R. Akter et al., "A hybrid CNN-LSTM model for environmental sound classification," *Digital Signal Processing*, 2025.
9. S. Abdoli et al., "End-to-End Environmental Sound Classification using 1D CNN," *arXiv*, 2019.
10. J. Sharma et al., "Environment Sound Classification using Multiple Feature Channels and Attention-based DCNN," *arXiv*, 2019.
11. B. Zhu et al., "Environmental Sound Classification Based on Multi-temporal Resolution CNN," *arXiv*, 2018.
12. D. Jain et al., "SoundWatch: Exploring Smartwatch-based Deep Learning for Sound Accessibility," *ASSETS*, 2020.
13. D. Jain et al., "Deep Learning for Sound Accessibility on Smartwatches," *Communications of the ACM*, 2022.
14. A. Ashurov et al., "Environmental Sound Classification Based on Transfer Learning," *Electronics*, 2022.
15. Z. Zhang et al., "Deep CNN with Mixup for Environmental Sound Classification," *arXiv*, 2018.