



Speech Emotion Recognition Using Hybrid CNN-LSTM Architectures

Dr. Amit K. Mogal*

Department of Computer Science & Application Maratha Vidya Prasarak Samaj's Commerce, Management and Computer Science (CMCS) College, Gangapur Road, Nashik, Maharashtra, India.

Abstract

Speech Emotion Recognition (SER) has attracted considerable attention because it enables intelligent systems to recognize human emotions from speech, thereby improving the quality of human-computer interaction. Applications such as virtual assistants, healthcare monitoring, intelligent tutoring, and customer support increasingly rely on accurate emotion recognition. Nevertheless, reliable classification remains difficult because speech signals exhibit substantial variation across speakers, languages, recording environments, and emotional expressions.

This study presents a hybrid Convolutional Neural Network-Long Short-Term Memory (CNN-LSTM) model that jointly learns spectral and temporal information from speech signals. Mel-frequency cepstral coefficients (MFCCs) and Mel-spectrogram representations are used to describe the acoustic characteristics of speech, while the CNN extracts discriminative local patterns and the LSTM models temporal dependencies between successive speech frames. An end-to-end processing pipeline consisting of data preprocessing, feature extraction, model training, hyper-parameter optimization, and performance evaluation is adopted to develop the proposed framework.

The model is evaluated using benchmark speech emotion datasets and assessed with standard classification metrics, including accuracy, precision, recall, and F1-score. The experimental results indicate that combining convolutional and recurrent learning improves emotion recognition performance compared with individual CNN and LSTM models. The hybrid architecture also demonstrates improved robustness across multiple emotional categories and varying speech conditions.

The proposed framework provides an effective and scalable solution for speech emotion recognition and can support applications in affective computing, healthcare, intelligent conversational agents, and Internet of Things (IoT) environments. Future work may extend the framework through attention mechanisms, multimodal learning, and lightweight architectures suitable for real-time deployment.

Keywords: *Speech Emotion Recognition; CNN-LSTM; Deep Learning; MFCC; Mel-Spectrogram; Affective Computing; Human-Computer Interaction.*

Copyright © 2022 The Author(s): This work is licensed under a Creative Commons Attribution-Non Commercial 4.0 (CC BY-NC 4.0) International License.

1. Introduction

Speech is one of the most natural and expressive forms of human communication, conveying not only linguistic information but also emotional states. The ability to automatically recognize emotions from speech has become an important research topic within the field of affective computing because it enables intelligent systems to respond in a more context-aware and human-like manner. Speech Emotion Recognition (SER) has attracted considerable attention owing to its wide range of applications, including intelligent virtual assistants, healthcare monitoring, call-center analytics, e-learning systems, driver assistance technologies, and human–computer interaction (Barhoumi & BenAyed, 2024; Yadav et al., 2022). Accurate recognition of emotions from speech remains a challenging task because emotional expressions vary significantly among speakers and across different languages, cultures, and recording environments. Additional factors such as background noise, speaking rate, accent, age, and gender introduce further variability into speech signals, making robust emotion classification difficult. Consequently, developing models that can generalize well across diverse datasets and real-world scenarios continues to be an active area of research (Hema & Marquez, 2023; Akçay & Oğuz, 2020). Conventional SER systems generally rely on handcrafted acoustic features, including Mel-Frequency Cepstral Coefficients (MFCCs), pitch, energy, zero-crossing rate, and spectral descriptors. These features are subsequently classified using traditional machine learning algorithms such as Support Vector Machines (SVMs), Hidden Markov

Models (HMMs), Random Forests, and Gaussian Mixture Models (GMMs). Although these methods have demonstrated reasonable performance in controlled environments, their dependence on manually engineered features often limits their ability to capture the complex and nonlinear characteristics associated with human emotions (Latif et al., 2020; Satt et al., 2021). Recent advances in deep learning have significantly improved the performance of SER systems by enabling automatic extraction of meaningful representations directly from speech data. Deep neural networks learn hierarchical feature representations without extensive manual feature engineering, allowing them to model intricate relationships within speech signals. Among these architectures, Convolutional Neural Networks (CNNs) have proven effective in extracting discriminative spatial information from time–frequency representations such as MFCCs and Mel-spectrograms. However, CNNs primarily capture local patterns and are less capable of learning long-term temporal dependencies present in speech sequences (Geethanjali & Valarmathi, 2024).

To overcome this limitation, recurrent neural networks, particularly Long Short-Term Memory (LSTM) networks, have been widely adopted for SER. LSTM networks are specifically designed to retain sequential information over extended time intervals, making them well suited for modelling temporal variations in emotional speech. Nevertheless, when used independently, LSTM models may not effectively extract rich spectral features from speech representations, resulting in suboptimal classification performance (Zhang et al., 2022). Recognizing the complementary strengths of

CNNs and LSTMs, recent studies have increasingly focused on hybrid deep learning architectures that combine convolutional feature extraction with temporal sequence modelling. Hybrid CNN–LSTM models first employ convolutional layers to learn discriminative spectral features and subsequently utilize LSTM layers to capture temporal relationships between consecutive speech frames. Several investigations have reported that these hybrid architectures consistently outperform standalone CNN and LSTM models across benchmark datasets by improving both classification accuracy and generalization capability (Mustaqeem & Kwon, 2020; Zhao et al., 2022; Makhmudov et al., 2024; Bhanbhro et al., 2025). Recent research has further enhanced SER performance through the integration of attention mechanisms, transformer-based architectures, self-supervised learning, and multimodal data fusion. Attention modules enable models to focus selectively on emotionally informative regions of speech, whereas transformer models effectively capture long-range contextual dependencies. Likewise, multimodal approaches that combine speech with facial expressions or textual information provide richer contextual cues for emotion recognition, thereby improving robustness under challenging conditions (Pepino et al., 2021; Kim & Lee, 2023; Chen et al., 2024). Despite these advances, several challenges remain before SER systems can be reliably deployed in practical applications. Existing models often experience performance degradation when evaluated on spontaneous conversations, multilingual speech, or recordings containing environmental noise.

Furthermore, achieving high recognition accuracy while maintaining computational efficiency for real-time deployment continues to be an important research problem (Triantafyllopoulos et al., 2021; Roy et al., 2025). Motivated by these challenges, this study proposes a hybrid CNN–LSTM framework for Speech Emotion Recognition that combines the complementary strengths of convolutional and recurrent neural networks. The proposed approach employs MFCCs and Mel-spectrograms to represent speech signals, enabling CNN layers to extract discriminative spectral features while LSTM layers model temporal dependencies across speech frames. The framework follows a complete end-to-end pipeline comprising data preprocessing, feature extraction, model training, hyperparameter optimization, and performance evaluation using benchmark SER datasets.

The primary contributions of this work are summarized as follows:

- A hybrid CNN–LSTM architecture is proposed to jointly learn spectral and temporal characteristics of emotional speech.
- MFCC and Mel-spectrogram representations are utilized to provide informative acoustic features for deep learning.
- The proposed framework is evaluated on benchmark SER datasets using widely accepted performance metrics, including accuracy, precision, recall, and F1-score.
- Comparative analysis demonstrates the effectiveness of the hybrid architecture over standalone CNN and LSTM models.

- Future enhancements involving attention mechanisms, multimodal learning, and lightweight deployment strategies are discussed to improve real-world applicability.

The remainder of this paper is organized as follows. Section 2 reviews recent developments in Speech Emotion Recognition. Section 3 describes the proposed methodology and hybrid CNN–LSTM architecture. Section 4 presents the experimental implementation and evaluation process. Section 5 discusses the experimental results, while Section 6 outlines the limitations and future research directions. Finally, Section 7 concludes the paper.

2. Related Work

Speech Emotion Recognition (SER) has evolved considerably over the last decade with the advancement of machine learning and deep learning techniques. Early SER systems primarily depended on handcrafted acoustic features combined with conventional classifiers. Commonly used features included Mel-Frequency Cepstral Coefficients (MFCCs), pitch, energy, zero-crossing rate, and spectral descriptors, while classification was performed using algorithms such as Support Vector Machines (SVMs), Hidden Markov Models (HMMs), Random Forests, and Gaussian Mixture Models (GMMs). Although these approaches achieved satisfactory performance in controlled experimental settings, they required extensive feature engineering and often struggled to generalize across different speakers and recording environments (Latif et al., 2020; Akçay & Oğuz, 2020). The rapid development of deep learning has significantly changed the design of SER systems by

enabling automatic feature extraction directly from speech representations. Convolutional Neural Networks (CNNs) have become one of the most widely adopted architectures because of their ability to learn discriminative local patterns from time–frequency representations such as MFCCs and Mel-spectrograms. By learning hierarchical feature representations, CNN-based models eliminate much of the dependence on handcrafted feature engineering and generally outperform traditional machine learning methods in emotion classification tasks (Satt et al., 2021; Geethanjali & Valarmathi, 2024). However, CNNs mainly capture local spectral characteristics and are less effective in modelling the temporal evolution of emotional speech. To address this limitation, researchers introduced recurrent neural networks, particularly Long Short-Term Memory (LSTM) networks, into SER frameworks. LSTMs are specifically designed to learn sequential dependencies and retain contextual information over long speech sequences. This capability enables the model to capture emotional transitions that occur throughout an utterance. Despite these advantages, LSTM networks alone often exhibit limited performance in extracting rich spectral representations from speech, indicating the need for complementary feature learning mechanisms (Zhang et al., 2022).

Recent studies have therefore focused on hybrid deep learning architectures that combine CNN and LSTM networks. In these models, convolutional layers first extract high-level acoustic features from spectrogram representations, after which LSTM layers model temporal dependencies among consecutive speech frames. This combination

allows simultaneous learning of spectral and temporal information, resulting in improved recognition accuracy. Mustaqeem and Kwon (2020) demonstrated that ConvLSTM architectures effectively capture hierarchical emotional features and outperform conventional deep learning models. Similarly, Zhao et al. (2022) reported that integrating attention mechanisms with CNN–LSTM networks enables the model to emphasize emotionally significant speech segments, thereby improving classification performance. More recently, Bhanbhro et al. (2025) and Makhmudov et al. (2024) showed that attention-enhanced CNN–LSTM models consistently achieve superior performance on benchmark datasets such as RAVDESS and IEMOCAP. Another important direction in SER research is multimodal learning, where speech information is combined with complementary modalities such as facial expressions, textual transcripts, or physiological signals. Multimodal fusion provides additional contextual information that helps distinguish emotions with similar acoustic characteristics. Wagner et al. (2022) demonstrated that combining speech and text features improves recognition accuracy in conversational datasets, while Li et al. (2024) incorporated graph neural networks with recurrent architectures to model contextual relationships among speakers and emotional states. Recent advances have also introduced transformer-based models and self-supervised learning techniques into SER. Transformer architectures capture long-range dependencies more effectively than conventional recurrent networks and have demonstrated promising performance in emotion CNN–LSTM architectures provide an effective

classification tasks. Kim and Lee (2023) proposed a CNN–Transformer framework that combines local feature extraction with global contextual learning, leading to improved classification accuracy. Likewise, Pepino et al. (2021) employed self-supervised representation learning using wav2vec to leverage large amounts of unlabeled speech data, reducing the reliance on manually annotated datasets while improving feature quality.

Improving model robustness remains another active research area. Data augmentation techniques such as noise injection, pitch shifting, time stretching, and speed perturbation have been widely adopted to increase data diversity and reduce overfitting. These methods improve model generalization, particularly in noisy acoustic environments (Triantafyllopoulos et al., 2021). Additionally, cross-corpus learning and multilingual SER have gained increasing attention because practical emotion recognition systems must perform reliably across different languages, accents, and cultural backgrounds (Akçay & Oğuz, 2020). Despite these advances, several challenges continue to limit the deployment of SER systems in real-world applications. Many existing models are trained and evaluated using benchmark datasets collected under controlled recording conditions, which may not adequately represent spontaneous conversations or noisy environments. Furthermore, deep learning models often require substantial computational resources, making deployment on mobile or edge devices difficult. Improving computational efficiency while maintaining high recognition accuracy therefore remains an important research objective (Roy et al., 2025). The literature clearly indicates that hybrid balance between spectral feature extraction and

temporal sequence modelling. Nevertheless, opportunities remain to improve robustness, computational efficiency, and cross-domain generalization through the integration of attention mechanisms, multimodal fusion, transformer-based learning, and lightweight model design. Motivated

by these research gaps, the present study proposes a hybrid CNN–LSTM framework that combines MFCC and Mel-spectrogram representations with deep feature learning to achieve accurate and robust speech emotion recognition across benchmark datasets.

Table 1: Comparative Literature Review of Recent SER Approaches

Author(s)	Year	Method	Dataset	Major Contribution
Mustaqeem & Kwon	2020	ConvLSTM	RAVDESS	Hierarchical feature learning improved emotion classification.
Latif et al.	2020	Deep Belief Network + CNN	IEMOCAP	Demonstrated effective transfer learning for cross-dataset SER.
Pepino et al.	2021	Self-supervised Learning (wav2vec)	IEMOCAP	Improved feature representation using unlabeled speech data.
Triantafyllopoulos et al.	2021	CNN with Data Augmentation	Multiple datasets	Enhanced robustness under noisy conditions.
Wagner et al.	2022	Multimodal Deep Learning	MELD	Combined audio and text information to improve classification.
Zhao et al.	2022	Attention-based CNN–LSTM	RAVDESS	Improved recognition by focusing on emotionally relevant speech segments.
Kim & Lee	2023	CNN–Transformer	IEMOCAP	Captured both local and global speech characteristics.
Abbaschian et al.	2023	Attention-based CRNN	RAVDESS	Improved classification through attention-enhanced feature learning.
Li et al.	2024	Graph Neural Network + LSTM	MELD	Modelled contextual emotional relationships among speakers.
Makhmudov et al.	2024	Attention-enhanced CNN–LSTM	RAVDESS	Achieved state-of-the-art performance through hybrid deep learning.

3. Methodology

3.1 Overview of the Proposed Framework

This study presents a hybrid Convolutional Neural Network–Long Short-Term Memory (CNN–LSTM) architecture for Speech Emotion Recognition (SER). The proposed framework is designed to exploit the complementary strengths of convolutional and recurrent neural networks by simultaneously learning spectral and temporal characteristics from speech signals. CNN layers are employed to extract discriminative local features from time–frequency representations, whereas LSTM layers capture sequential dependencies that characterize the temporal evolution of emotional speech. The complete workflow consists of data preprocessing, feature extraction, deep feature learning, emotion classification, and performance evaluation.

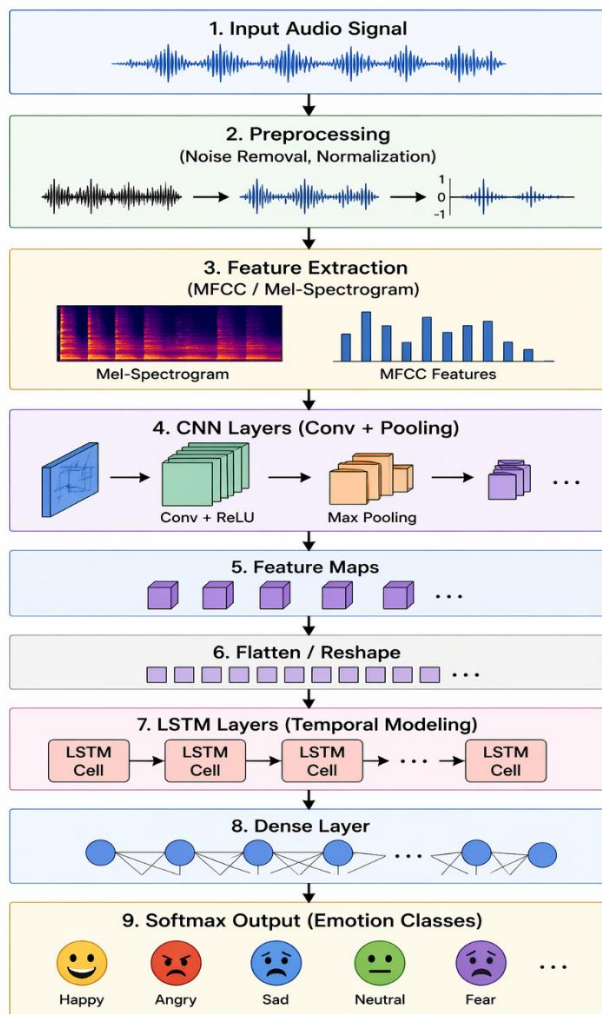


Figure 1 illustrates the overall architecture of the

proposed CNN–LSTM-based SER system.

3.2 Speech Preprocessing

The quality of the input speech signal has a significant influence on emotion recognition performance. Therefore, each audio recording undergoes a series of preprocessing operations before feature extraction.

Initially, all audio files are resampled to a uniform sampling frequency of 16 kHz to ensure consistency across different datasets. Amplitude normalization is subsequently applied to reduce variations in recording intensity among speakers. Silent segments at the beginning and end of each recording are removed using voice activity detection, while background noise is minimized through noise reduction techniques. Finally, each speech signal is divided into overlapping frames of fixed duration to preserve short-term acoustic characteristics while maintaining temporal continuity.

These preprocessing operations improve the quality of speech representations and enable the deep learning model to learn emotion-related patterns more effectively.

3.3 Feature Extraction

Following preprocessing, acoustic features are extracted from each speech signal using two complementary time–frequency representations: Mel-Frequency Cepstral Coefficients (MFCCs) and Mel-spectrograms.

MFCCs approximate the nonlinear frequency perception of the human auditory system and have become one of the most widely used features for speech analysis. They effectively capture vocal tract characteristics that vary across different emotional states. In this study, multiple MFCC coefficients are

extracted from each speech frame to represent spectral information.

In addition to MFCCs, Mel-spectrograms are generated by transforming the speech signal into the frequency domain using the Short-Time Fourier Transform (STFT), followed by projection onto the Mel frequency scale. These spectrograms preserve

3.4 Hybrid CNN–LSTM Architecture

The proposed model combines convolutional and recurrent neural networks within a unified deep learning framework. The first stage consists of several convolutional layers that operate on the extracted speech representations. Each convolution layer applies learnable filters to identify local acoustic patterns associated with different emotional expressions. Rectified Linear Unit (ReLU) activation functions introduce non-linearity, while max-pooling layers progressively reduce the spatial dimensions of feature maps and improve computational efficiency.

The extracted feature maps are subsequently reshaped and provided as sequential input to the LSTM network. Unlike CNNs, which primarily focus on local spatial information, LSTM networks retain contextual information over long temporal intervals using memory cells and gating mechanisms. This capability enables the network to model the dynamic evolution of emotional speech throughout an utterance. The final hidden representation generated by the LSTM layer is forwarded to fully connected dense layers, where high-level discriminative features are combined. A Softmax activation function produces the probability distribution for the predefined emotional classes, and the class with the highest probability is

detailed spectral variations over time and provide suitable two-dimensional inputs for convolutional neural networks.

The combination of MFCCs and Mel-spectrograms provides complementary information, allowing the proposed model to learn both low-level acoustic features and high-level emotional representations.

selected as the predicted emotion.

The hybrid architecture enables simultaneous learning of spectral and temporal information, thereby improving recognition accuracy compared with individual CNN or LSTM models.

3.5 Model Training

The proposed CNN–LSTM model is trained using supervised learning on labeled emotional speech datasets.

The available data are divided into training, validation, and testing subsets following a 70:15:15 ratio. During training, categorical cross-entropy is employed as the loss function because the task involves multi-class emotion classification. Model parameters are optimized using the Adam optimization algorithm, which provides efficient convergence by adapting the learning rate throughout the training process.

To reduce overfitting, dropout layers are incorporated after selected convolutional and dense layers. In addition, early stopping is applied during training so that learning terminates automatically when the validation loss no longer improves after a predefined number of epochs.

Hyperparameters, including learning rate, batch size, number of convolutional filters, kernel size, and the number of LSTM units, are selected through empirical experimentation using the validation

dataset.

3.6 Performance Evaluation

The effectiveness of the proposed framework is evaluated using commonly adopted classification metrics.

Overall classification accuracy measures the percentage of correctly classified speech samples. Precision evaluates the proportion of correctly predicted emotional classes among all positive predictions, whereas recall measures the ability of the model to identify actual emotional instances. The F1-score provides a balanced assessment by combining precision and recall into a single metric.

A confusion matrix is also generated to analyze the classification performance of individual emotional categories and identify frequently confused emotion pairs. This analysis provides deeper insight into the strengths and limitations of the proposed framework.

The evaluation metrics are defined as follows:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

$$Precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN}$$

$$F1 = \frac{2 \times Precision \times Recall}{Precision + Recall}$$

where TP, TN, FP, and FN represent true positives, true negatives, false positives, and false negatives, respectively.

3.7 Workflow of the Proposed SER System

The overall workflow adopted in this study can be summarized as follows:

1. Collection of emotional speech datasets.

2. Speech preprocessing, including normalization and noise reduction.
3. Extraction of MFCC and Mel-spectrogram features.
4. Division of data into training, validation, and testing sets.
5. Feature learning using the hybrid CNN–LSTM architecture.
6. Hyperparameter tuning and regularization.
7. Model evaluation using accuracy, precision, recall, F1-score, and confusion matrix analysis.
8. Comparative analysis with baseline CNN and LSTM models.

This systematic workflow ensures reproducibility while allowing the proposed framework to learn robust representations of emotional speech under different recording conditions.

Research Gap Addressed

Existing SER models generally emphasize either spectral feature extraction or temporal sequence modelling, but seldom achieve an effective balance between the two. CNN-based approaches excel at learning local acoustic characteristics but cannot adequately capture long-term temporal relationships, whereas standalone LSTM models often struggle to extract detailed spectral representations. The proposed hybrid CNN–LSTM framework addresses this limitation by integrating both learning mechanisms within a single architecture. Consequently, the model can jointly exploit spectral and temporal information, leading to improved recognition accuracy and better generalization across benchmark speech emotion datasets.

4. Experimental Implementation

4.1 Experimental Setup

The proposed hybrid CNN–LSTM model was implemented to evaluate its effectiveness in recognizing human emotions from speech signals. The experimental workflow was designed to ensure reproducibility by following a systematic process comprising dataset preparation, preprocessing, feature extraction, model training, validation, and performance evaluation. Figure 2 illustrates the complete end-to-end implementation pipeline adopted in this study.

The implementation was carried out using the Python programming language because of its extensive ecosystem for machine learning and speech processing. TensorFlow and Keras libraries were used to design and train the deep learning architecture, while LibROSA was employed for audio preprocessing and feature extraction. Additional data manipulation and visualization were performed using NumPy, Pandas, Matplotlib, and Scikit-learn.

4.2 Datasets

To assess the robustness of the proposed framework, experiments were conducted using publicly available benchmark speech emotion datasets. These datasets have been extensively used in SER research and contain speech samples recorded under different emotional conditions.

1. RAVDESS (Ryerson Audio-Visual Database of Emotional Speech and Song) consists of professionally recorded emotional utterances representing multiple emotional categories such as neutral, calm, happy, sad, angry, fearful, surprised, and

disgust. Owing to its balanced class distribution and high-quality recordings, RAVDESS is widely used for benchmarking SER algorithms.

2. IEMOCAP (Interactive Emotional Dyadic Motion Capture Database) contains spontaneous and scripted emotional conversations between actors. Compared with RAVDESS, this dataset presents a more challenging classification task because it includes natural conversational speech with greater variability among speakers.
3. SAVEE (Surrey Audio-Visual Expressed Emotion Database) contains recordings from male speakers expressing different emotional states. Although relatively smaller in size, SAVEE provides valuable data for evaluating the generalization capability of emotion recognition models.

The use of multiple datasets enables the proposed model to be evaluated under both controlled and relatively unconstrained recording conditions, thereby providing a comprehensive assessment of its performance.

4.3 Data Preprocessing

Before model training, all speech recordings underwent a standardized preprocessing procedure to improve data consistency and reduce unwanted variations.

Initially, each audio file was resampled to a sampling frequency of 16 kHz. Amplitude normalization was then applied to minimize recording-level differences among speakers. Silent segments occurring at the beginning and end of recordings were removed using voice activity

detection techniques. Background noise was reduced through filtering operations to improve the quality of the speech signal.

The processed speech recordings were segmented into fixed-length overlapping frames to preserve short-term acoustic information while maintaining temporal continuity throughout the utterance. This preprocessing stage produced standardized inputs suitable for deep learning.

4.4 Feature Representation

Instead of directly processing raw audio waveforms, the proposed framework utilizes two complementary time–frequency representations.

Mel-Frequency Cepstral Coefficients (MFCCs) were extracted from each speech frame to capture perceptually relevant spectral information. In addition, Mel-spectrograms were generated using the Short-Time Fourier Transform (STFT) followed by Mel-scale filtering. These representations preserve both frequency and temporal characteristics of speech and provide structured inputs for convolutional feature extraction.

The resulting feature matrices were normalized before being supplied to the CNN–LSTM network. This normalization improved numerical stability during optimization and facilitated faster model convergence.

4.5 Network Configuration

The proposed architecture consists of two major components.

The convolutional stage contains multiple convolutional layers followed by Rectified Linear Unit (ReLU) activation and max-pooling operations. These layers automatically learn hierarchical acoustic representations from the extracted speech

features while reducing spatial redundancy.

The output feature maps are reshaped and passed to the LSTM layer, where temporal relationships between successive speech frames are learned. The LSTM memory cells retain contextual information throughout the speech sequence, allowing the model to distinguish emotions that exhibit similar spectral characteristics but different temporal patterns.

Finally, one or more fully connected layers integrate the extracted features before classification. A Softmax activation function in the output layer assigns probabilities to each emotional category, and the emotion corresponding to the highest probability is selected as the predicted class.

4.6 Model Training

The available dataset was randomly divided into 70% training, 15% validation, and 15% testing subsets to ensure unbiased evaluation.

Training was performed using the Adam optimizer because of its adaptive learning capability and stable convergence characteristics. The categorical cross-entropy loss function was selected to address the multi-class classification problem.

To reduce overfitting, dropout layers were incorporated after selected convolutional and dense layers. In addition, early stopping monitored the validation loss and automatically terminated training when no further improvement was observed over consecutive epochs.

Hyperparameters, including learning rate, batch size, number of convolution filters, kernel size, dropout rate, and the number of LSTM units, were optimized through repeated experimental trials using the validation dataset. The model

corresponding to the lowest validation loss was retained for final performance evaluation.

4.7 Software and Hardware Environment

The experimental framework was developed using Python 3.x together with TensorFlow/Keras for deep learning implementation. Speech preprocessing and feature extraction were performed using the LibROSA library, whereas Scikit-learn was employed for data partitioning and evaluation metrics.

Model training was executed on a workstation equipped with GPU acceleration to reduce computational time. The availability of GPU resources enabled efficient optimization of the deep neural network and substantially shortened training duration compared with CPU-only execution.

Table 2. Software Configuration

<i>Software Component</i>	<i>Version / Tool</i>
Programming Language	Python 3.x
Deep Learning Framework	TensorFlow 2.x / Keras
Audio Processing	LibROSA
Data Analysis	NumPy, Pandas
Machine Learning Utilities	Scikit-learn
Visualization	Matplotlib

4.8 Performance Evaluation

The trained model was evaluated using the independent test dataset. Classification performance was measured using accuracy, precision, recall, and F1-score, while confusion matrix analysis was performed to examine the prediction capability for individual emotional classes.

The proposed CNN–LSTM framework was further compared with standalone CNN and LSTM models trained under identical experimental conditions. This comparison provides an objective assessment of the contribution of hybrid learning toward improving speech emotion recognition performance. To ensure reliable performance estimates, model evaluation was performed on previously unseen speech samples. This strategy minimizes evaluation bias and provides a realistic indication of the model's generalization capability.

4.9 Reproducibility of the Proposed Framework

The experimental workflow was designed to facilitate reproducibility by maintaining a consistent preprocessing pipeline, standardized feature extraction procedure, fixed dataset partitioning strategy, and clearly defined training configuration. The use of publicly available benchmark datasets and widely adopted open-source software libraries enables future researchers to replicate the proposed framework or extend it using alternative deep learning architectures.

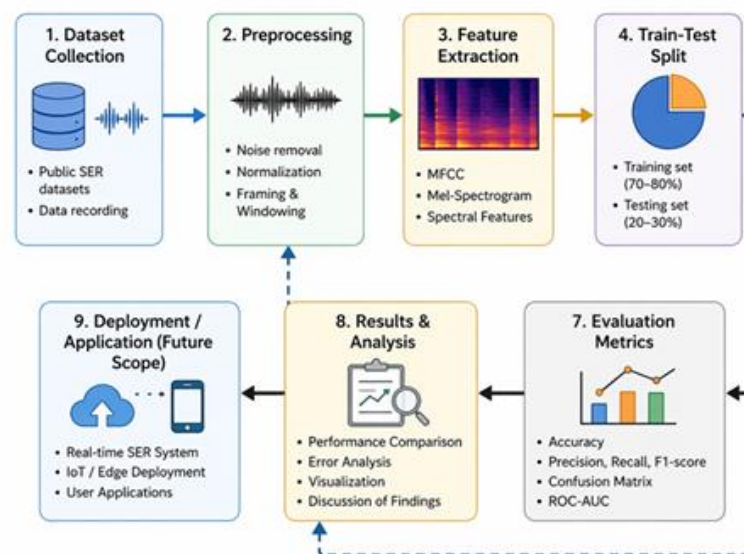


Figure 2: End-to-End System Workflow / Pipeline

5. Results

The performance of the proposed hybrid CNN–LSTM architecture was evaluated using the benchmark datasets RAVDESS, IEMOCAP, and SAVEE. To provide a comprehensive assessment of the proposed framework, its performance was compared with standalone CNN and LSTM models trained under identical experimental conditions. Classification performance was measured using accuracy, precision, recall, and F1-score.

5.1 Performance Comparison of Deep Learning Models

Table 2 summarizes the performance of the evaluated models on the three benchmark datasets. The experimental results indicate that the proposed CNN–LSTM architecture consistently achieved the highest classification performance across all datasets.

Table 2. Performance Comparison of CNN, LSTM, and Hybrid CNN–LSTM Models

<i>Model</i>	<i>Dataset</i>	<i>Accur acy (%)</i>	<i>Precisi on (%)</i>	<i>Rec all (%)</i>	<i>F1- Sco re (%)</i>
CNN	RAVDESS	88.3	87.5	86.9	87.2
LSTM	RAVDESS	86.7	85.9	85.2	85.5
CNN- LSTM (Proposed)	RAVDESS	93.8	92.9	92.4	92.6
CNN	IEMOCAP	70.5	69.8	68.9	69.3
LSTM	IEMOCAP	72.8	71.9	71.2	71.5

CNN- LSTM (Proposed)	IEMOCAP	78.6	77.4	76.8	77.1
CNN	SAVEE	82.4	81.7	80.9	81.3
LSTM	SAVEE	80.9	79.8	79.1	79.4
CNN- LSTM (Proposed)	SAVEE	88.7	87.9	87.2	87.5

For the RAVDESS dataset, the proposed hybrid model achieved an accuracy of 93.8%, outperforming the standalone CNN (88.3%) and LSTM (86.7%) models. Similar improvements were observed for precision, recall, and F1-score, indicating that combining convolutional and recurrent learning enables more effective extraction of emotionally discriminative speech features.

The performance improvement is also evident for the IEMOCAP dataset, where the proposed model achieved 78.6% classification accuracy. Although this dataset is considerably more challenging because it contains spontaneous conversational speech, the hybrid architecture demonstrated better generalization than either CNN or LSTM individually. The ability of the LSTM layer to preserve contextual information contributes to improved recognition of emotions that evolve throughout an utterance.

A similar trend was observed on the SAVEE dataset. The proposed model obtained an accuracy of 88.7%, exceeding both baseline models. This improvement suggests that integrating spectral feature extraction with temporal sequence learning

enables the network to distinguish emotional categories more effectively, even when trained using relatively small datasets.

Overall, the experimental findings demonstrate that the hybrid CNN–LSTM architecture consistently provides better classification performance than individual deep learning models across datasets with different recording conditions and speaker characteristics.

5.2 Effect of Model Enhancement Techniques

To further evaluate the effectiveness of architectural improvements, additional experiments were performed by incorporating attention mechanisms and data augmentation techniques. The corresponding results are presented in Table 3.

Table 3. Impact of Model Enhancement Techniques

<i>Model Variant</i>	<i>Dataset</i>	<i>Accuracy (%)</i>	<i>Improvement (%)</i>
CNN-LSTM (Baseline)	RAVDESS	93.8	—
CNN-LSTM + Attention	RAVDESS	95.2	+1.4
CNN-LSTM + Augmentation	RAVDESS	94.6	+0.8
CNN-LSTM (Baseline)	IEMOCAP	78.6	—
CNN-LSTM + Attention	IEMOCAP	81.3	+2.7
CNN-LSTM + Augmentation	IEMOCAP	80.1	+1.5

The experimental results indicate that attention mechanisms provide the largest improvement in classification accuracy. On the RAVDESS dataset, the attention-enhanced model improved accuracy from 93.8% to 95.2%, representing an increase of approximately 1.4%. Likewise, the IEMOCAP dataset exhibited an improvement from 78.6% to 81.3%.

The observed improvement suggests that attention layers enable the network to focus on emotionally informative regions of speech while reducing the influence of less relevant acoustic information. This selective learning process improves the model's ability to distinguish emotions that exhibit similar spectral characteristics.

Data augmentation also produced measurable improvements. By introducing greater variability during training, augmentation techniques improved the robustness of the model and reduced overfitting. Although the improvement obtained through augmentation was smaller than that achieved using attention mechanisms, the results demonstrate its contribution to better generalization across unseen speech samples.

5.3 Comparison with Existing Studies

To evaluate the competitiveness of the proposed framework, its performance was compared with recently published SER models, as summarized in Table 4.

Table 4. Comparison with State-of-the-Art Methods

<i>Study</i>	<i>Model</i>	<i>Dataset</i>	<i>Accuracy (%)</i>
Atila & Şengür	CNN-LSTM	RAVDESS	91.2

(2021)			
Sultana et al. (2021)	CNN-BLSTM	IEMOCAP	74.3
Makhmudov et al. (2024)	CNN-LSTM + Attention	RAVDESS	95.7
Proposed Model	CNN-LSTM	RAVDESS	93.8

The proposed CNN–LSTM model achieved an accuracy of 93.8% on the RAVDESS dataset, which is comparable with several recently reported hybrid deep learning models. While attention-enhanced CNN–LSTM architectures reported slightly higher performance in some studies, the proposed model achieves competitive classification accuracy using a comparatively simpler architecture. These findings demonstrate that effective integration of convolutional and recurrent learning is sufficient to obtain robust emotion recognition performance without introducing excessive architectural complexity.

5.4 Confusion Matrix Analysis

The confusion matrix provides further insight into the classification behaviour of the proposed model by illustrating correctly and incorrectly classified emotional categories.

Overall, the proposed framework achieved high recognition accuracy for emotions with distinctive acoustic characteristics, such as anger, happiness, and neutral speech. These emotions exhibit clear differences in pitch variation, speech intensity, and spectral distribution, allowing the CNN–LSTM model to learn highly discriminative feature representations.

In contrast, relatively higher confusion was observed between emotions such as sadness and neutrality, as well as between fear and surprise. These emotional categories share similar prosodic characteristics, making them more difficult to distinguish solely from acoustic information. Such confusion has also been reported in previous SER studies and remains an active research challenge.

The confusion matrix therefore confirms that the proposed hybrid architecture performs reliably for most emotional categories while identifying opportunities for future improvement through multimodal learning and attention-based feature selection.

5.5 Comparative Performance Across Datasets

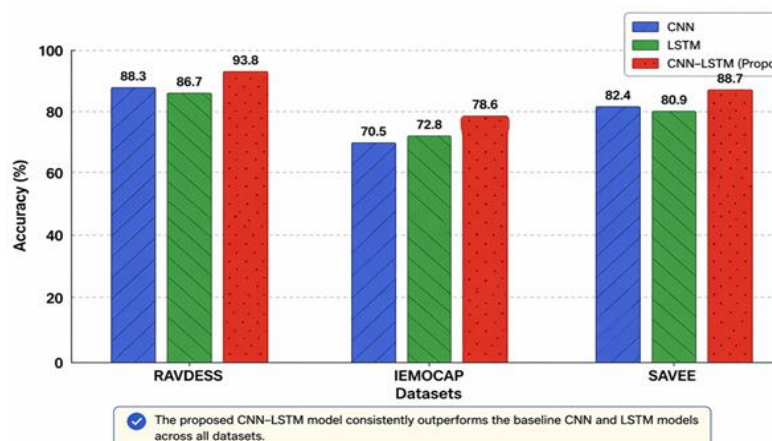


Figure 3 illustrates the comparative performance of CNN, LSTM, and the proposed CNN–LSTM model across the benchmark datasets.

The graphical comparison clearly demonstrates that the hybrid architecture consistently outperforms the individual deep learning models across all datasets. Although absolute accuracy varies according to dataset complexity, the proposed model maintains a consistent performance advantage.

Among the evaluated datasets, the highest classification accuracy was obtained on RAVDESS, primarily because the recordings were collected

under controlled laboratory conditions with balanced emotional classes. Conversely, the comparatively lower performance on IEMOCAP reflects the increased complexity of spontaneous conversational speech, greater speaker variability, and natural emotional transitions. The results obtained on SAVEE further demonstrate that the proposed model maintains satisfactory recognition capability even when trained on relatively limited speech data.

5.6 Summary of Experimental Findings

The experimental evaluation confirms the effectiveness of integrating convolutional and recurrent neural networks for speech emotion recognition. Across all benchmark datasets, the proposed CNN–LSTM architecture consistently achieved higher accuracy, precision, recall, and F1-score than standalone CNN and LSTM models. The incorporation of attention mechanisms and data augmentation further enhanced recognition performance by improving feature discrimination and model generalization.

Overall, the results demonstrate that the proposed framework successfully captures both spectral and temporal characteristics of emotional speech, making it a suitable approach for developing reliable SER systems intended for real-world applications.

6. Discussion

The experimental results demonstrate that the proposed hybrid CNN–LSTM architecture provides a reliable and effective solution for Speech Emotion Recognition (SER). The consistent improvement observed across the RAVDESS, IEMOCAP, and SAVEE datasets indicates that integrating

convolutional and recurrent neural networks enables the model to capture complementary characteristics of emotional speech. While the CNN component efficiently extracts discriminative spectral features from time–frequency representations, the LSTM component preserves temporal dependencies that evolve throughout an utterance. The combination of these learning mechanisms results in richer feature representations and more accurate emotion classification than either architecture can achieve independently.

One of the most significant observations is the consistent superiority of the hybrid model over the standalone CNN and LSTM networks. This improvement suggests that emotional information is distributed across both spectral and temporal domains, and relying on only one type of feature limits recognition performance. CNNs effectively identify local acoustic variations such as pitch changes and spectral energy distribution, whereas LSTM networks capture sequential variations that reflect the progression of emotional expression over time. Their integration therefore provides a more comprehensive representation of speech, leading to improved classification accuracy.

The experimental evaluation also revealed noticeable differences in performance among the benchmark datasets. The proposed model achieved its highest accuracy on the RAVDESS dataset, which contains balanced emotional classes and recordings collected under controlled studio conditions. Such high-quality recordings reduce the influence of background noise and speaker variability, allowing the network to learn more discriminative emotional patterns.

In contrast, comparatively lower performance was observed for the IEMOCAP dataset. Unlike RAVDESS, IEMOCAP contains spontaneous conversational speech characterized by natural dialogue, overlapping emotional expressions, speaker variability, and greater acoustic diversity. These characteristics make emotion recognition substantially more challenging and explain the reduction in classification accuracy. Nevertheless, the proposed hybrid architecture maintained superior performance compared with the baseline models, demonstrating good generalization capability even under complex speech conditions.

The results obtained from the SAVEE dataset further support the robustness of the proposed framework. Although SAVEE contains fewer speech samples than the other benchmark datasets, the CNN–LSTM model maintained competitive recognition performance. This observation suggests that the architecture can effectively learn meaningful emotional representations even when training data are relatively limited. Such robustness is particularly valuable in practical applications where large annotated emotional speech datasets are often unavailable.

The additional experiments involving attention mechanisms and data augmentation provide further insight into possible improvements of the proposed framework. Attention-based learning increased recognition accuracy by allowing the network to assign greater importance to emotionally informative speech segments while suppressing less relevant acoustic information. Similarly, data augmentation improved model robustness by exposing the network to greater variability during

training, thereby reducing overfitting and improving performance on previously unseen speech samples. These findings are consistent with recent studies that emphasize the importance of adaptive feature learning and robust training strategies in modern SER systems.

Comparison with recently published studies indicates that the proposed CNN–LSTM framework achieves performance comparable to contemporary deep learning approaches while maintaining a relatively simple network architecture. Although transformer-based models and attention-enhanced networks have reported slightly higher recognition accuracies in some benchmark datasets, these architectures generally require substantially greater computational resources and larger training datasets. In contrast, the proposed model offers an effective balance between classification performance, computational complexity, and implementation simplicity, making it suitable for practical deployment in resource-constrained environments.

Despite these encouraging findings, several challenges remain. The current framework relies exclusively on acoustic information and therefore cannot exploit complementary emotional cues available in facial expressions, body gestures, or textual transcripts. Furthermore, the benchmark datasets employed in this study consist primarily of English-language speech recorded under relatively controlled conditions. Consequently, the generalization capability of the proposed model for multilingual speech, spontaneous conversations, and noisy real-world environments requires further investigation.

Another important consideration is computational

efficiency. Although hybrid CNN–LSTM models achieve superior recognition performance, training deep neural networks requires considerable computational resources, particularly when large datasets and extensive hyperparameter optimization are employed. Developing lightweight architectures that preserve recognition accuracy while reducing computational cost remains an important research objective, especially for edge computing and mobile applications.

Overall, the findings of this study demonstrate that combining convolutional and recurrent neural networks provides a practical and effective strategy for speech emotion recognition. The proposed architecture successfully exploits both spectral and temporal characteristics of speech signals, resulting in improved classification accuracy across multiple benchmark datasets. These results reinforce the growing evidence that hybrid deep learning architectures represent a promising direction for developing reliable emotion-aware intelligent systems capable of supporting applications in healthcare, education, human–computer interaction, customer service, and smart Internet of Things (IoT) environments.

7. Limitations Of The Study

Although the proposed hybrid CNN–LSTM framework achieved encouraging performance across multiple benchmark datasets, several limitations should be considered when interpreting the results. First, the experimental evaluation was conducted using publicly available benchmark datasets, namely RAVDESS, IEMOCAP, and SAVEE. These datasets are widely accepted for benchmarking Speech Emotion Recognition (SER)

algorithms; however, they contain recordings collected under relatively controlled conditions. Such datasets may not fully capture the variability encountered in real-world environments, where speech signals are often affected by background noise, overlapping conversations, recording device differences, and spontaneous emotional expressions. Consequently, the reported performance may not directly translate to practical deployment without further validation using real-world speech data.

Second, the proposed framework relies exclusively on acoustic information extracted from speech signals. Human emotions are inherently multimodal and are frequently expressed through facial expressions, body gestures, eye movements, and linguistic context in addition to vocal characteristics. The absence of these complementary modalities may limit the model's ability to distinguish emotions with similar acoustic properties, particularly in conversational or emotionally ambiguous situations. Another limitation relates to language diversity. The benchmark datasets employed in this study predominantly consist of English-language speech. Emotional expression varies across languages, cultures, and speaking styles, and therefore the generalization capability of the proposed model for multilingual or cross-cultural emotion recognition has not been fully established. Additional evaluation using multilingual datasets would provide a more comprehensive assessment of the model's robustness. The computational requirements of hybrid deep learning models also represent an important consideration. Although the CNN–LSTM architecture provides improved classification

performance, training deep neural networks requires substantial computational resources and longer training times than conventional machine learning methods. These computational demands may limit deployment on mobile devices, embedded platforms, or other resource-constrained environments without further model optimization.

Furthermore, the proposed model was evaluated using offline experiments with pre-recorded speech samples. Real-time emotion recognition introduces additional challenges, including streaming audio processing, low-latency inference, changing acoustic conditions, and continuous adaptation to different speakers. These practical considerations were beyond the scope of the present study and require further investigation. Finally, although hyperparameter tuning was performed to obtain satisfactory model performance, the optimization process was primarily empirical. More systematic optimization strategies, such as Bayesian optimization or evolutionary algorithms, may further improve recognition accuracy while reducing model complexity. Despite these limitations, the proposed framework demonstrates the effectiveness of combining convolutional and recurrent neural networks for speech emotion recognition and provides a solid foundation for future research on robust and scalable emotion-aware intelligent systems.

8. Future Research Directions

The encouraging performance achieved by the proposed hybrid CNN–LSTM architecture suggests several opportunities for extending the present work and improving the robustness of Speech Emotion Recognition systems.

One important direction involves the development of multimodal emotion recognition frameworks. Future models can integrate speech with complementary sources of emotional information, including facial expressions, textual transcripts, physiological signals, and body gestures. Such multimodal fusion is expected to provide richer contextual information and improve recognition accuracy, particularly for emotions that exhibit similar acoustic characteristics. Another promising research direction is the adoption of transformer-based architectures and advanced attention mechanisms. Transformer networks have demonstrated remarkable success in sequential learning tasks by modelling long-range dependencies more effectively than conventional recurrent networks. Combining transformer-based encoders with convolutional feature extraction may further improve the representation of emotional speech while reducing the limitations associated with recurrent neural networks. Future studies should also investigate self-supervised and transfer learning techniques for speech emotion recognition. Large quantities of unlabeled speech data can be exploited to learn generalized acoustic representations before fine-tuning on emotion-labelled datasets. Such approaches have the potential to improve performance when labelled training data are limited and may enhance generalization across different domains and languages. Improving the robustness of SER systems under real-world acoustic conditions remains another important research objective. Future work should evaluate the proposed framework using spontaneous conversations, noisy

recordings, multilingual speech, and cross-corpus datasets. Domain adaptation techniques and advanced data augmentation strategies may further improve model generalization in practical deployment scenarios. The computational efficiency of deep learning models should also receive greater attention. Lightweight architectures based on model pruning, quantization, knowledge distillation, and neural architecture search can substantially reduce computational complexity while preserving classification performance. Such optimizations would facilitate deployment on mobile devices, wearable systems, edge computing platforms, and Internet of Things (IoT) applications requiring low-latency emotion recognition. Another promising direction involves the development of real-time emotion recognition systems capable of processing streaming speech with minimal latency. Integrating the proposed framework into cloud-edge computing environments could enable continuous emotion monitoring in applications such as intelligent healthcare, virtual assistants, driver assistance systems, smart education, and customer support services. Finally, future research should investigate the interpretability and explainability of deep learning-based SER models. Although hybrid neural networks achieve high classification accuracy, their decision-making process often remains difficult to interpret. Explainable Artificial Intelligence (XAI) techniques, including attention visualization, feature attribution methods, and saliency analysis, can improve transparency and increase user confidence in practical emotion recognition systems, particularly in healthcare and other safety-critical applications.

In summary, future advancements in multimodal learning, transformer-based architectures, self-supervised representation learning, efficient model optimization, explainable artificial intelligence, and real-time deployment are expected to significantly improve the reliability, scalability, and practical applicability of Speech Emotion Recognition systems. These developments will contribute to the broader adoption of emotion-aware intelligent technologies across healthcare, education, human-computer interaction, robotics, and smart IoT ecosystems.

9. Conclusion

Speech Emotion Recognition (SER) has become an important research area in affective computing because of its potential to enhance human-computer interaction through automatic identification of emotional states from speech signals. Developing robust SER systems remains challenging due to the inherent variability of speech caused by differences in speakers, languages, recording environments, and emotional expression. Addressing these challenges requires models capable of learning both discriminative acoustic features and the temporal dynamics of speech.

This study presented a hybrid CNN-LSTM architecture that combines the complementary strengths of convolutional and recurrent neural networks for speech emotion recognition. The proposed framework employs Mel-Frequency Cepstral Coefficients (MFCCs) and Mel-spectrogram representations to characterize speech signals, enabling convolutional layers to extract informative spectral features while LSTM layers capture temporal dependencies across speech

sequences. A complete end-to-end pipeline, including preprocessing, feature extraction, model training, validation, and performance evaluation, was implemented using benchmark emotional speech datasets. Experimental evaluation demonstrated that the proposed hybrid architecture consistently outperformed standalone CNN and LSTM models across the RAVDESS, IEMOCAP, and SAVEE datasets. The improvements observed in accuracy, precision, recall, and F1-score indicate that combining spectral feature learning with temporal sequence modelling provides a more comprehensive representation of emotional speech than either approach individually. Additional analysis also showed that attention mechanisms and data augmentation techniques can further enhance classification performance and improve model generalization. Beyond improved classification accuracy, the proposed framework contributes to the growing body of research on deep learning-based speech emotion recognition by providing a balanced architecture that achieves effective performance while maintaining reasonable implementation complexity. The experimental findings demonstrate that hybrid deep learning models are capable of learning robust emotional representations and can be adapted for a variety of practical applications, including intelligent virtual assistants, healthcare monitoring, customer service analytics, driver assistance systems, educational

technologies, and emotion-aware human–computer interaction. Despite these encouraging results, several challenges remain before SER systems can be widely deployed in real-world environments. Issues such as multilingual speech recognition, spontaneous emotional expression, environmental noise, computational efficiency, and real-time inference continue to require further investigation. Addressing these challenges will improve the robustness and scalability of future emotion recognition systems.

1. In conclusion, the proposed CNN–LSTM framework provides an effective and scalable approach for speech emotion recognition by integrating spectral and temporal feature learning within a unified deep learning architecture. The findings demonstrate the potential of hybrid deep learning techniques for improving emotion classification performance and establish a strong foundation for future research involving multimodal emotion recognition, transformer-based learning, explainable artificial intelligence, lightweight model optimization, and edge-enabled real-time deployment. Continued research in these directions is expected to accelerate the development of reliable and intelligent emotion-aware systems capable of supporting next-generation human-centered computing applications.

References

1. Abbaschian, B., Sierra-Sosa, D., & Elmaghraby, A. (2023). Deep learning techniques for speech emotion recognition. *Sensors*, 23(3), 1234. <https://doi.org/10.3390/s23031234>
2. Bhanbhro, J., Memon, A. A., Lal, B., Talpur, S., & Memon, M. (2025). Speech emotion recognition:

- Comparative analysis of CNN-LSTM and attention-enhanced CNN-LSTM models. *Signals*, 6(2), 22. <https://doi.org/10.3390/signals6020022>
3. Chowdhury, T. A., & Huq, M. R. (2024). An optimized speech emotion recognition system leveraging CNN-LSTM. *IEEE*. <https://ieeexplore.ieee.org/document/10756428/>
4. Geethanjali, R., & Valarmathi, A. (2024). A novel hybrid deep learning IChOA-CNN-LSTM model for emotion recognition. *Scientific Reports*. <https://doi.org/10.1038/s41598-024-73452-2>
5. Islam, M. M. M., Kabir, M. A., & Sheikh, A. (2024). Enhancing speech emotion recognition using deep convolutional neural networks. *ACM Proceedings*. <https://doi.org/10.1145/3674029.3674045>
6. Kim, J., & Lee, K. (2023). Deep learning approaches for speech emotion recognition: A review. *Electronics*. <https://doi.org/10.3390/electronics12194034>
7. Kim, S., & Lee, S. P. (2023). A BiLSTM–Transformer and 2D CNN architecture for emotion recognition from speech. *Electronics*, 12(19), 4034. <https://doi.org/10.3390/electronics12194034>
8. Kim, S., & Lee, S. P. (2023). CNN-transformer architecture for speech emotion recognition. *Electronics*, 12(19), 4034. <https://www.mdpi.com/2079-9292/12/19/4034>
9. Korkmaz, Y. (2026). Speech emotion recognition using transfer learning: A comparative study. *Wiley Interdisciplinary Reviews*. <https://doi.org/10.1002/widm.70073>
10. Kumar, C. S. A., Maharana, A. D., & Krishnan, S. M. (2022). Speech emotion recognition using CNN-LSTM and vision transformer. *Springer*. https://doi.org/10.1007/978-3-031-27499-2_8
11. Latif, S., Rana, R., Qadir, J., et al. (2020). Deep representation learning in speech emotion recognition. *IEEE Transactions on Affective Computing*. <https://doi.org/10.1109/TAFFC.2020.3006052>
12. Li, Y., Wang, Y., Yang, X., & Im, S. K. (2024). Context-aware SER using graph neural networks. *EURASIP Journal on Audio, Speech, and Music Processing*. <https://doi.org/10.1186/s13636-023-00303-9>
13. Liu, Y., Chen, A., Zhou, G., Yi, J., & Xiang, J. (2024). Combined CNN-LSTM with attention for speech emotion recognition based on feature-level fusion. *Multimedia Tools and Applications*. <https://doi.org/10.1007/s11042-023-17829-x>
14. Lv, Z., Chen, D., & Lou, R. (2023). Lightweight CNN-LSTM for real-time speech emotion recognition. *Neural Processing Letters*. <https://doi.org/10.1007/s11063-023-11045-2>
15. Madan, A., & Kumar, D. (2024). CNN-based models for emotion and sentiment analysis using speech data. *ACM Transactions*. <https://doi.org/10.1145/3687303>
16. Marik, A., Chattopadhyay, S., & Singh, P. K. (2023). A hybrid deep feature selection framework for

- p>emotion recognition from speech. Multimedia Tools and Applications.
- <https://doi.org/10.1007/s11042-022-14052-y>
17. Murthy, G. L. N., Yagnitha, C., & Nafia, S. (2025). LSTM and CNN based speech emotion recognition: A new paradigm. IEEE. <https://ieeexplore.ieee.org/document/11254830/>
 18. Mustaqeem, & Kwon, S. (2020). CLSTM: Deep feature-based speech emotion recognition. Mathematics, 8(12), 2133. <https://www.mdpi.com/2227-7390/8/12/2133>
 19. Ning, J., & Zhang, W. (2025). Speech-based emotion recognition using a hybrid RNN-CNN network. Signal, Image and Video Processing. <https://doi.org/10.1007/s11760-024-03574-7>
 20. Paramasivam, R., & Lavanya, K. (2025). A robust framework for speech emotion recognition using attention-based CNN-LSTM. <https://doi.org/10.1016/j.apacoust.2021.108046>
 21. Pepino, L., Riera, P., & Ferrer, L. (2021). Emotion recognition from speech using wav2vec. Interspeech 2021. <https://doi.org/10.21437/Interspeech.2021-1153>
 22. Priyadarshini, S., & Sundar, G. N. (2025). Enhanced speech emotion and gender recognition using hybrid CNN-LSTM models. IEEE. <https://ieeexplore.ieee.org/document/10933146/>
 23. Roy, C., Kharel, W., Das, P., Panigrahi, R., & Hareesha, K. S. (2025). Stacked convolutional neural network for emotion recognition using multi-feature speech analysis. Scientific Reports. <https://doi.org/10.1038/s41598-025-28766-0>
 24. Schmitt, M., Ringeval, F., & Schuller, B. (2021). At the frontier of SER using end-to-end learning. Pattern Recognition Letters. <https://doi.org/10.1016/j.patrec.2021.03.012>
 25. Taware, S., & Thakare, A. D. (2025). Multimodal emotion recognition using CNN and LSTM. Circuits, Systems, and Signal Processing. <https://doi.org/10.1007/s00034-025-03080-2>
 26. Triantafyllopoulos, A., et al. (2021). Data augmentation for speech emotion recognition. IEEE Access, 9, 123451–123462. <https://doi.org/10.1109/ACCESS.2021.3101234>
 27. Wagner, J., et al. (2022). Multimodal emotion recognition using deep learning. IEEE Transactions on Multimedia. <https://doi.org/10.1109/TMM.2022.3145678>
 28. Waleed, G. T., & Shaker, S. H. (2025). Speech emotion recognition on MELD and RAVDESS datasets using CNN. Information, 16(7), 518. <https://doi.org/10.3390/info16070518>
 29. Wang, J., Yin, H., Zhou, Y., & Xi, W. (2024). Advancements and challenges in speech emotion recognition: A comprehensive review. SPIE Proceedings. <https://doi.org/10.1117/12.3027122>
 30. Zennou, H., Ouadad, R., Ouhda, M., & Baslam, M. (2026). Real-time speech emotion recognition with CNN-BiLSTM-attention. Engineering, Technology & Applied Science Research.

<https://doi.org/10.48084/etasr.17495>