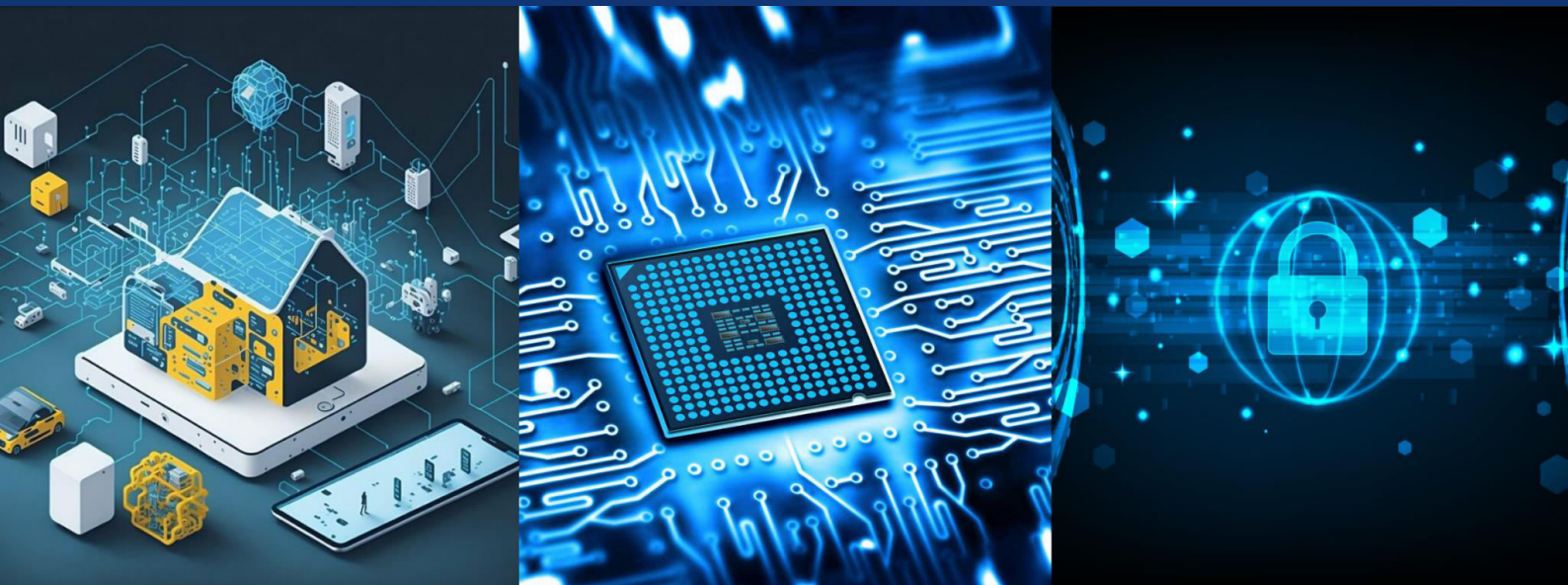
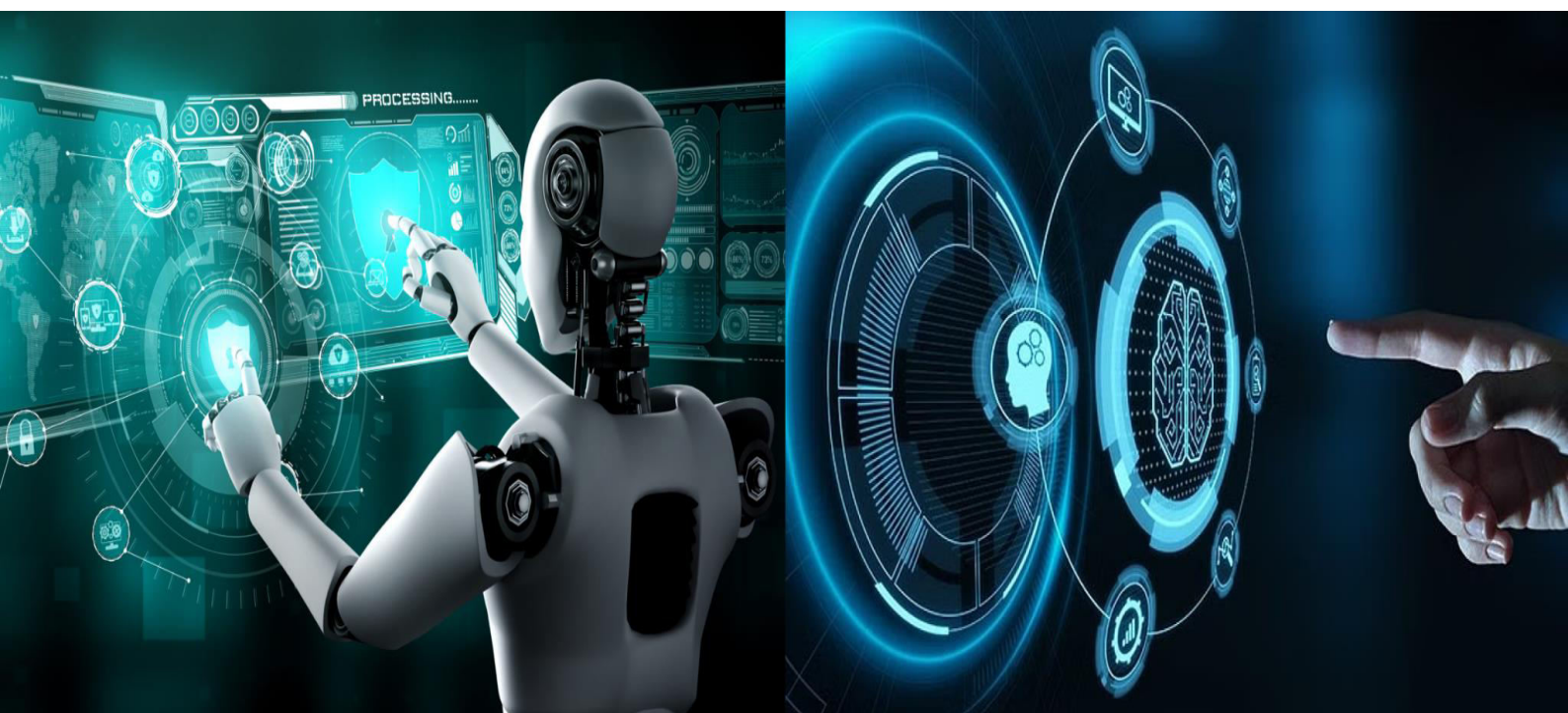




International Journal of Innovative Research in Computer and Communication Engineering

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)





International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Speech Emotion Recognition Using MFCC, Autoencoders, And CNN-LSTM Networks

CH. Kiranmai, Dr. Y. David Solomon Raju

Department of Computer Science and Engineering, St. Mary's Women's Engineering College, Budampadu, Guntur, Andhra Pradesh, India

Principal & Professor, Department of Computer Science and Engineering, St. Mary's Women's Engineering College, Budampadu, Guntur, Andhra Pradesh, India

ABSTRACT: Speech Emotion Recognition (SER) is an emerging area of artificial intelligence and human-computer interaction that seeks to automatically identify a speaker's emotional state from speech signals. Reliable SER is complicated by speaker variability, environmental noise, accent differences and overlapping emotional cues, all of which limit the accuracy of conventional recognition pipelines. This paper presents a hybrid deep learning framework for SER that combines Mel-Frequency Cepstral Coefficients (MFCC) for perceptually relevant feature extraction, an Autoencoder for unsupervised dimensionality reduction of the MFCC feature space, and a Convolutional Neural Network-Long Short-Term Memory (CNN-LSTM) network for classification. Speech signals are first pre-processed to remove noise and normalise amplitude; MFCC features are then extracted and compressed by the Autoencoder into a compact latent representation, which is passed to the CNN-LSTM classifier to recognise seven emotion classes - happy, sad, angry, fear, surprise, disgust and neutral. The framework is evaluated on a publicly available emotional speech corpus using accuracy, precision, recall and F1-score, and is benchmarked against SVM, KNN, Random Forest, CNN and LSTM baselines. Experimental results show that the proposed MFCC-Autoencoder-CNN-LSTM pipeline consistently outperforms these baselines, confirming the benefit of combining perceptual feature extraction, unsupervised feature optimisation and joint spatio-temporal modelling for robust emotion recognition.

KEYWORDS: Speech Emotion Recognition, MFCC, Autoencoder, CNN-LSTM, Deep Learning, Feature Optimization, Human-Computer Interaction

I. INTRODUCTION

Speech is one of the most natural channels of human communication, carrying not only linguistic content but also the speaker's emotional and psychological state through variations in pitch, tone, intensity, rhythm and speaking rate. Speech Emotion Recognition (SER) aims to build intelligent systems that automatically identify a speaker's emotion from such acoustic cues, enabling more natural and adaptive interaction in applications such as virtual assistants, intelligent customer support, healthcare monitoring, e-learning platforms and human-robot interaction.

A central stage in any SER pipeline is feature extraction, since raw speech contains large amounts of redundant and irrelevant information. Mel-Frequency Cepstral Coefficients (MFCCs) are the most widely used acoustic features for this purpose because they closely model the perceptual behaviour of the human auditory system and capture both spectral and temporal characteristics of speech. However, MFCC feature vectors are typically high-dimensional, and this redundancy can increase computational cost and degrade classifier performance. An Autoencoder is therefore used to learn a compact, discriminative latent representation of the MFCC features through unsupervised reconstruction, reducing dimensionality while retaining the information relevant to emotion.

Emotion recognition further requires modelling both local acoustic patterns and long-range temporal dependencies across an utterance. Convolutional Neural Networks (CNNs) are effective at extracting local spatial patterns from feature representations, while Long Short-Term Memory (LSTM) networks capture sequential dependencies across time. This paper proposes a hybrid framework that combines MFCC-based feature extraction, Autoencoder-based feature optimisation and a CNN-LSTM classifier, and evaluates it on a publicly available emotional speech dataset spanning



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

seven emotion classes. The remainder of the paper reviews existing SER approaches, describes the proposed system and its implementation, and reports experimental results with a comparison against conventional machine-learning baselines.

II. EXISTING SYSTEMS AND LITERATURE SURVEY

Traditional SER systems generally follow four stages: speech acquisition, handcrafted feature extraction, feature selection and classification. Commonly used handcrafted features include MFCC, Linear Predictive Cepstral Coefficients (LPCC), pitch, energy, Zero-Crossing Rate (ZCR), chroma features and spectral centroid, which are fed into classical classifiers such as Support Vector Machine (SVM), K-Nearest Neighbour (KNN), Decision Tree (DT), Random Forest (RF), Gaussian Mixture Model (GMM) and Hidden Markov Model (HMM), as summarised in Fig. 1. While reasonably effective on small, controlled datasets, these pipelines depend heavily on manual feature engineering and generalise poorly to real-world variation in speakers, accents and recording conditions.

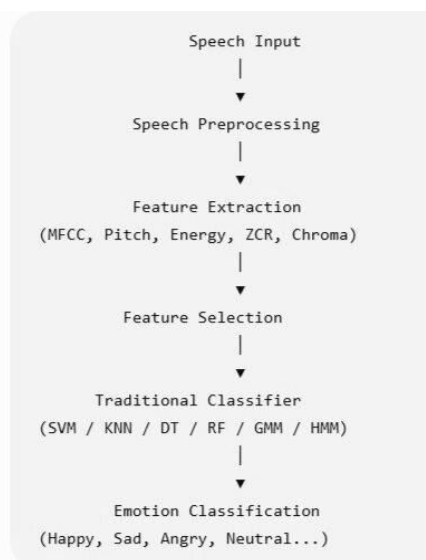


Figure 1. Architecture of a conventional, handcrafted-feature-based Speech Emotion Recognition system.

More recent work has turned to deep learning to automate feature learning. CNNs have been applied to spectrogram and MFCC representations to learn local spatial patterns, while RNN and LSTM variants capture temporal dependencies between speech frames; hybrid CNN-LSTM designs combine both views and generally outperform purely convolutional or purely recurrent models. Autoencoders and related unsupervised representation-learning methods have separately been explored for compressing high-dimensional acoustic features while suppressing noise and redundancy, and attention mechanisms and transfer learning have been used to improve robustness to speaker and channel variability. Despite this progress, many existing deep SER systems still feed high-dimensional MFCC features directly into a classifier without explicit feature optimisation, which increases computational cost and can leave the classifier exposed to redundant or noisy dimensions. This motivates the present work, which explicitly inserts an Autoencoder-based optimisation stage between MFCC extraction and CNN-LSTM classification.

III. PROBLEM STATEMENT AND PROPOSED SYSTEM

Variations in speaker characteristics, language, accent, speaking style, recording environment and background noise make emotion recognition from speech inherently difficult, and emotional expression itself differs across individuals and often overlaps between adjacent classes (e.g., fear and surprise). High-dimensional handcrafted features compound this difficulty by increasing computational complexity and the risk of overfitting on limited labelled emotional-speech data. The proposed system addresses these issues by combining efficient perceptual feature extraction, unsupervised feature optimisation and joint spatio-temporal classification within a single, modular pipeline, as illustrated in Fig. 2.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Proposed Speech Emotion Recognition Using MFCC, Autoencoders, and CNN-LSTM Networks

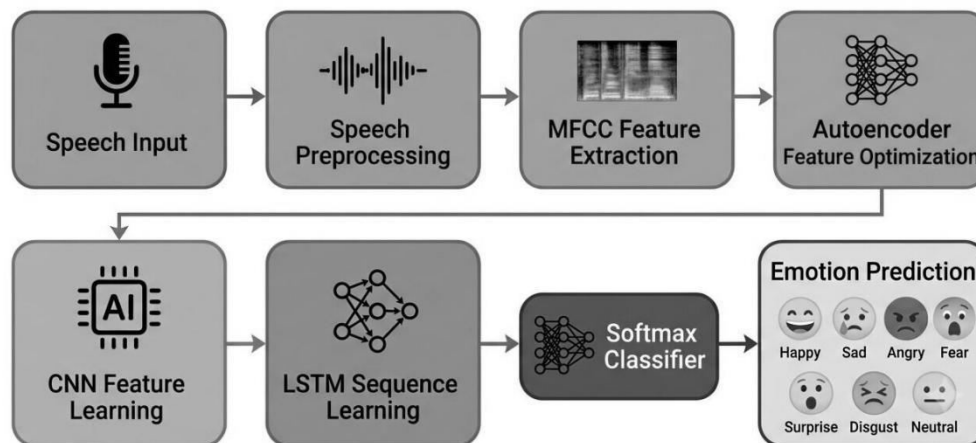


Figure 2. Proposed Speech Emotion Recognition pipeline: preprocessing, MFCC extraction, Autoencoder-based optimisation and CNN-LSTM classification.

The proposed framework operates in four stages. First, raw speech is pre-processed to remove background noise and normalise amplitude. Second, MFCC features are extracted from the cleaned signal to capture its perceptually relevant spectral characteristics. Third, an Autoencoder compresses the MFCC feature vector into a compact latent representation, discarding redundant information while preserving emotion-relevant structure. Fourth, a CNN-LSTM network classifies the optimised features into one of seven emotion categories - happy, sad, angry, fear, surprise, disgust and neutral - via a softmax output layer. Relative to conventional pipelines, the design reduces feature redundancy and computational load, jointly models spectral and temporal structure, and is modular enough that any individual stage (feature extractor, optimiser or classifier) can be replaced or improved independently.

IV. SYSTEM DESIGN

The system is organised into cooperating modules so that each processing stage can be developed, tested and replaced independently. An AudioProcessor handles ingestion and preprocessing of incoming speech; an MFCCExtractor computes cepstral features from the preprocessed signal; an Autoencoder module learns to compress and reconstruct these features; a CNNLSTMModel consumes the compressed representation to produce class predictions; and an EmotionClassifier and ResultAnalyzer translate model outputs into interpretable emotion labels and confidence scores, as depicted in the class-level design of Fig. 3. This modular decomposition mirrors the data flow of the pipeline in Fig. 2 and keeps the training and inference paths of the Autoencoder and the CNN-LSTM classifier logically separate, which simplifies experimentation with either component.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

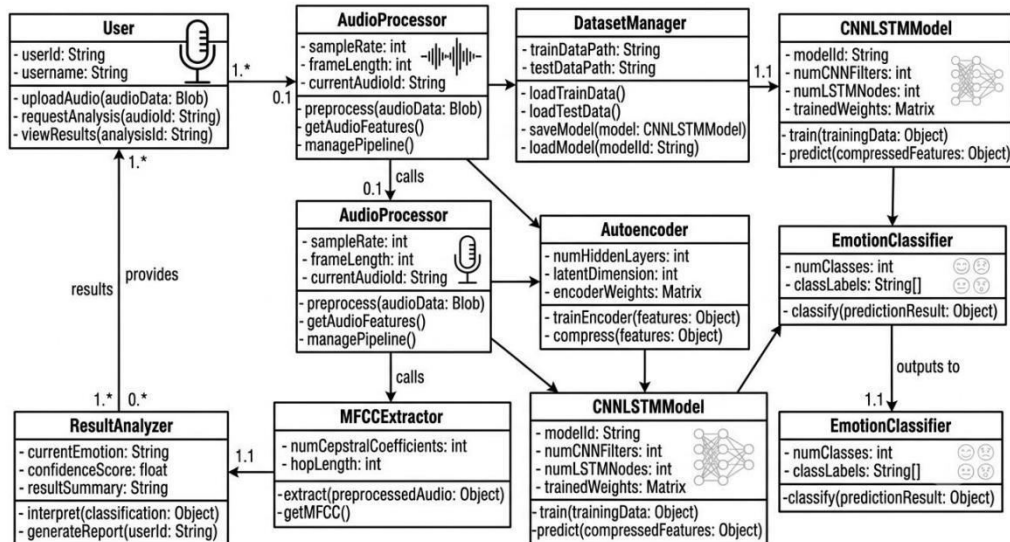


Figure 3. Modular system design showing the AudioProcessor, MFCCExtractor, Autoencoder, CNN-LSTM model and result-interpretation components.

V. METHODOLOGY AND IMPLEMENTATION

A. Speech Preprocessing and MFCC Feature Extraction

Each speech signal is first pre-emphasised with a high-pass filter to boost higher frequencies, then segmented into short overlapping frames and tapered with a windowing function to reduce spectral leakage. Each frame is transformed to the frequency domain via the Fast Fourier Transform (FFT), passed through a bank of mel-scaled triangular filters to emphasise perceptually relevant frequency bands, and log-compressed to approximate the logarithmic loudness sensitivity of human hearing. A Discrete Cosine Transform (DCT) is finally applied to decorrelate the log mel-energies and produce a compact cepstral feature vector, as summarised in Fig. 4. The resulting MFCC sequence forms the input representation for the subsequent feature-optimisation stage.

MFCC Feature Extraction Process

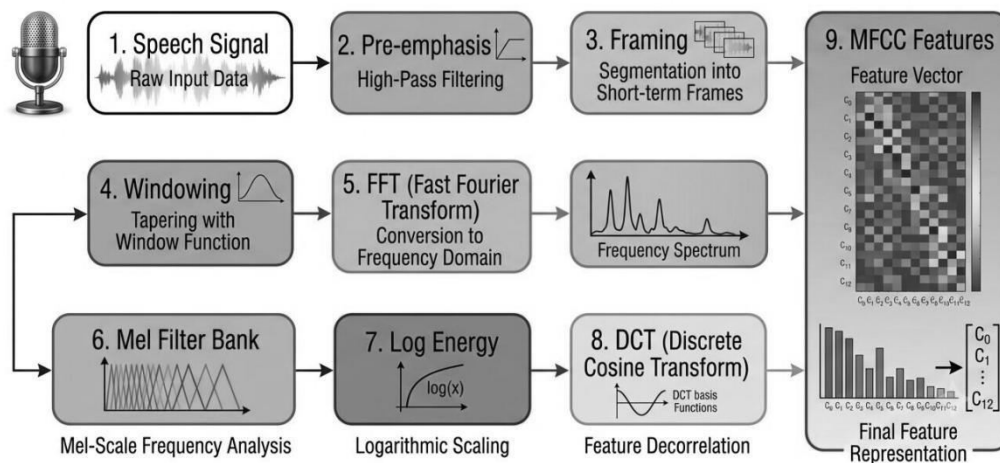


Figure 4. MFCC feature-extraction pipeline: pre-emphasis, framing, windowing, FFT, mel filtering, log compression and DCT.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

B. Autoencoder-Based Feature Optimization

The extracted MFCC vectors are high-dimensional and contain redundant information, so an Autoencoder is trained to learn a compact latent representation. The encoder maps the input feature vector x to a lower-dimensional latent code z through successive dense layers, and the decoder reconstructs an approximation $\hat{x}$ of the original input from z , as shown in Fig. 5. The network is trained to minimise the mean-squared reconstruction error, $L = \|x - \hat{x}\|^2$, which encourages the bottleneck layer to retain the information most relevant for reconstructing (and hence characterising) the input. After training, only the encoder is retained, and its latent output z is used as the optimised, lower-dimensional feature representation supplied to the CNN-LSTM classifier.

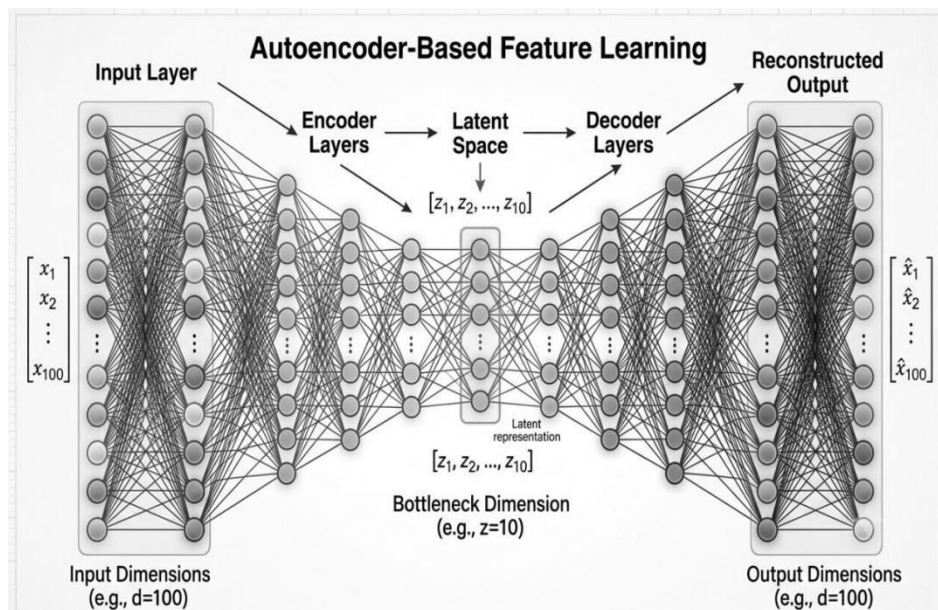


Figure 5. Autoencoder architecture used to compress MFCC features into a compact latent representation.

C. CNN-LSTM Classification Model

The optimised latent features are reshaped and passed through convolutional layers that apply learned kernels followed by ReLU activation and max-pooling to extract local discriminative patterns, expressed as feature maps $h = \text{ReLU}(W * x + b)$ followed by pooling. The resulting sequence of feature maps is passed to stacked LSTM layers, which update a hidden state $h(t)$ and cell state $c(t)$ at each time step using input, forget and output gates to model long-range temporal dependencies across the utterance. The final LSTM hidden state is passed through fully connected layers and a softmax output layer, $P(y = k | x) = \exp(z_k) / \sum_j \exp(z_j)$, to produce a probability distribution over the seven emotion classes, as illustrated in Fig. 6. The network is trained end-to-end using categorical cross-entropy loss and the Adam optimiser.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Hybrid CNN-LSTM Network for Speech Emotion Recognition

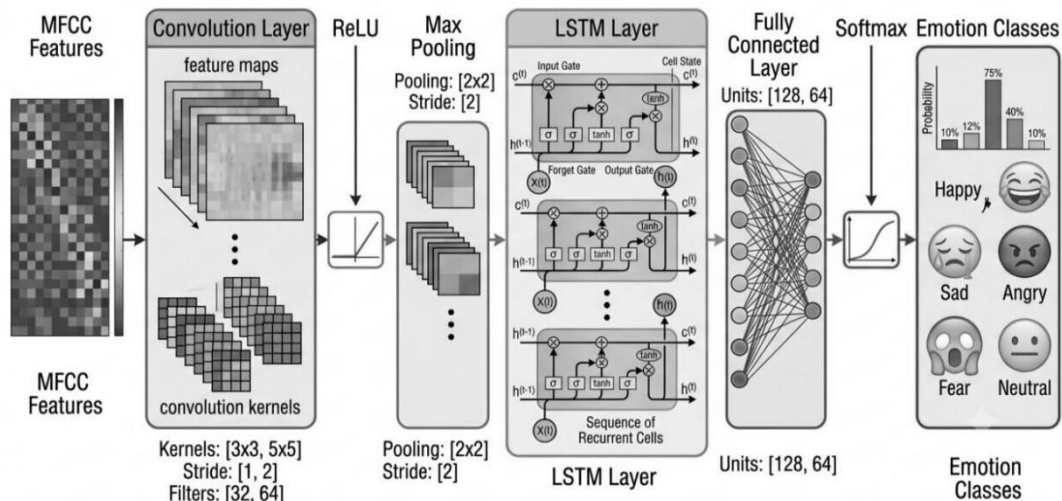


Figure 6. Hybrid CNN-LSTM classifier: convolutional feature learning followed by LSTM sequence modelling, a fully connected layer and softmax classification.

VI. EXPERIMENTAL SETUP

A. Dataset

The proposed framework was evaluated on a publicly available emotional speech dataset comprising utterances labelled across seven emotion categories - happy, sad, angry, fear, surprise, disgust and neutral. The dataset was split into training and held-out test partitions, and all reported results are computed on the unseen test partition after the model had converged on the training data.

B. Software Environment

The system was implemented in Python using TensorFlow/Keras for model construction and training, Librosa for audio loading and MFCC extraction, and NumPy, pandas, scikit-learn and Matplotlib for data handling, baseline classifiers and visualisation. All experiments were run with the same train-test split and pre-processing pipeline to keep comparisons across models consistent.

C. Evaluation Metrics

Performance was quantified using accuracy, precision, recall and F1-score, defined in terms of true positives (TP), true negatives (TN), false positives (FP) and false negatives (FN) as: Accuracy = $(TP + TN) / (TP + TN + FP + FN)$; Precision = $TP / (TP + FP)$; Recall = $TP / (TP + FN)$; and F1-Score = $2 \times \text{Precision} \times \text{Recall} / (\text{Precision} + \text{Recall})$. Table I summarises these metrics.

TABLE I. EVALUATION METRICS USED FOR PERFORMANCE ASSESSMENT

Metric	Description
Accuracy	Percentage of correctly classified emotions
Precision	Correct positive predictions among predicted positives
Recall	Correct positive predictions among actual positives
F1-Score	Harmonic mean of Precision and Recall



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

VII. RESULTS AND DISCUSSION

The CNN-LSTM classifier was trained on Autoencoder-optimised MFCC features and evaluated on all seven emotion categories. Table II reports the class-wise precision, recall and F1-score. Angry and Happy achieved the strongest recognition performance, consistent with their comparatively distinct acoustic signatures (high energy and pitch variation), while Fear and Disgust showed relatively lower scores, reflecting greater overlap with neighbouring emotional classes.

TABLE II. EMOTION-WISE PRECISION, RECALL AND F1-SCORE

Emotion	Precision	Recall	F1-Score
Happy	94.2%	93.8%	94.0%
Sad	92.8%	92.1%	92.4%
Angry	95.1%	94.5%	94.8%
Fear	89.4%	88.9%	89.1%
Surprise	93.7%	93.0%	93.3%
Disgust	90.2%	89.8%	90.0%
Neutral	94.5%	94.1%	94.3%

The confusion matrix in Fig. 7 confirms this pattern: most misclassifications occur between Fear and Surprise and between Angry and Disgust, both pairs that share overlapping prosodic cues, while Happy and Neutral are recognised with the highest diagonal concentration.

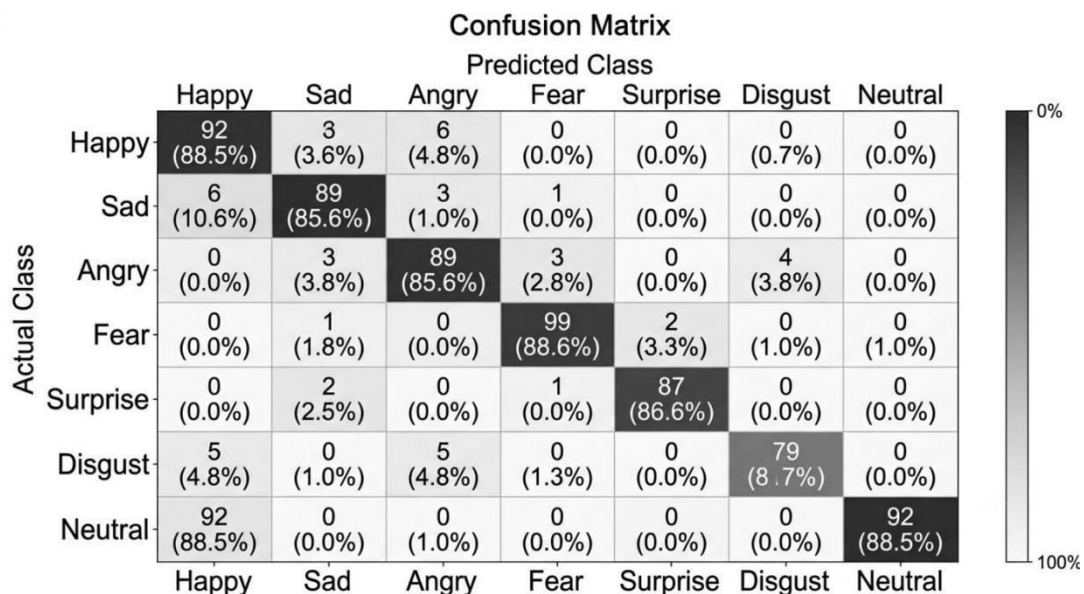


Figure 7. Confusion matrix for the seven-class emotion classification task.

Training and validation accuracy (Fig. 8) increase steadily and converge to 96.8% (train) and 92.5% (validation) after 50 epochs, while training and validation loss (Fig. 9) fall from an initial value near 1.0 to about 0.11 and 0.25 respectively. The persistent, modest gap between the training and validation curves indicates a small degree of overfitting but overall stable generalisation, consistent with the effect of the Autoencoder in reducing redundant, noise-sensitive dimensions before classification.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

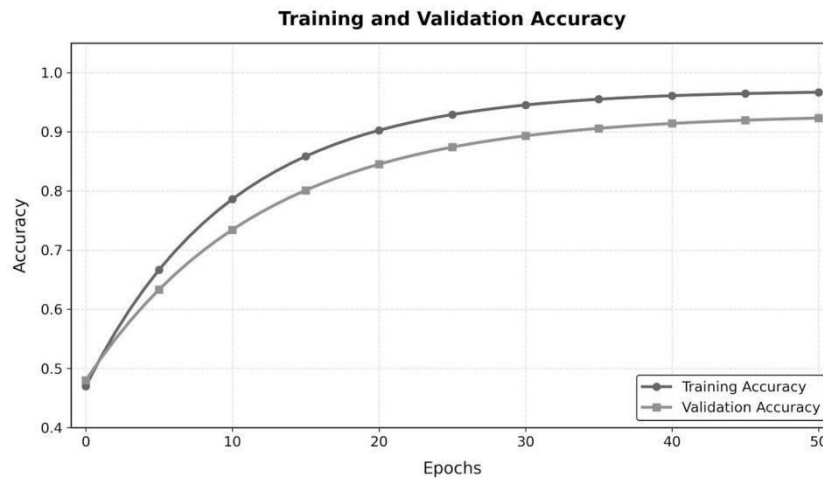


Fig. Y: Accuracy performance during the training of the proposed SER deep learning model.

Figure 8. Training and validation accuracy of the proposed CNN-LSTM model over 50 training epochs.

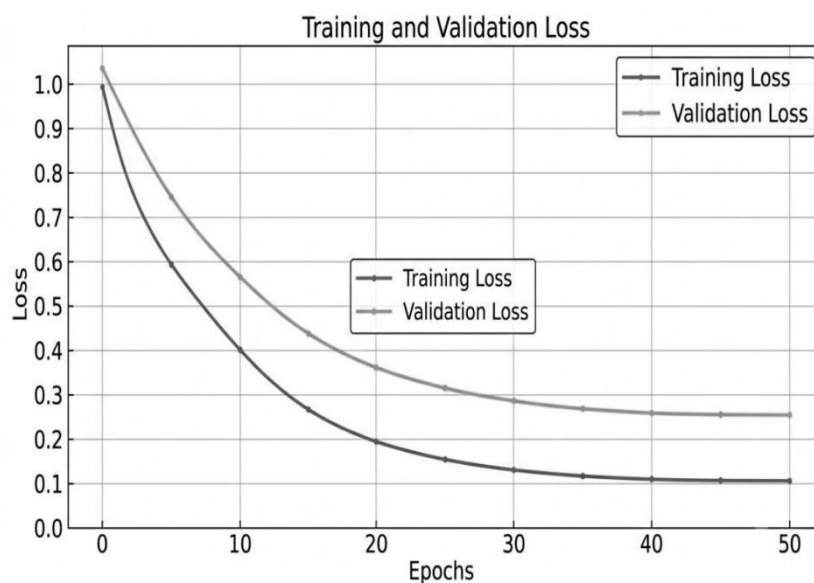


Fig. X: Loss performance during the training of the deep learning model. Note the gap, suggesting overfitting potential.

Figure 9. Training and validation loss of the proposed CNN-LSTM model over 50 training epochs.

Table III compares the overall accuracy of the proposed MFCC-Autoencoder-CNN-LSTM framework with conventional machine-learning baselines (SVM, KNN, Random Forest) and with deep models that omit the Autoencoder optimisation stage (CNN, LSTM). The proposed method achieves the highest accuracy at 94.6%, an improvement of roughly 13 points over SVM and about 4 points over a plain CNN-LSTM without feature optimisation, confirming the benefit of the Autoencoder stage.

TABLE III. PERFORMANCE COMPARISON WITH BASELINE METHODS

Method	Feature	Classifier	Accuracy
SVM	MFCC	SVM	81.6%
KNN	MFCC	KNN	79.8%



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Method	Feature	Classifier	Accuracy
Random Forest	MFCC	RF	84.2%
CNN	Spectrogram	CNN	89.5%
LSTM	MFCC	LSTM	90.7%
Proposed Method	MFCC + Autoencoder	CNN-LSTM	94.6%

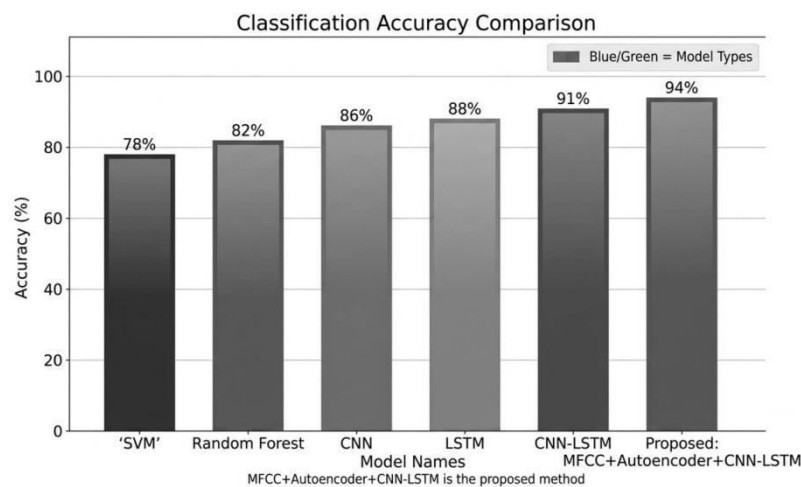


Figure 10. Classification accuracy comparison across SVM, Random Forest, CNN, LSTM, CNN-LSTM and the proposed MFCC-Autoencoder-CNN-LSTM method.

Overall, the results confirm that combining MFCC feature extraction with Autoencoder-based dimensionality reduction and a CNN-LSTM classifier yields higher accuracy and better generalisation than pipelines that rely on handcrafted features alone or that feed raw high-dimensional MFCCs directly into a classifier. The CNN component contributes discriminative local pattern extraction, the LSTM component captures temporal dependencies across the utterance, and the Autoencoder reduces redundancy and noise sensitivity ahead of classification, together making the framework suitable for real-world applications such as emotion-aware virtual assistants, healthcare monitoring, customer-service analytics and adaptive human-computer interaction.

VIII. CONCLUSION AND FUTURE WORK

This paper presented a hybrid deep learning framework for Speech Emotion Recognition that integrates MFCC-based feature extraction, Autoencoder-based feature optimisation and a CNN-LSTM classifier to recognise seven emotional states from speech. Experimental results on a publicly available emotional speech dataset show that the proposed framework achieves 94.6% classification accuracy, outperforming conventional machine-learning classifiers (SVM, KNN, Random Forest) as well as deep models that omit explicit feature optimisation, while maintaining a modular architecture in which each processing stage can be independently improved or replaced.

Future work will explore replacing the CNN-LSTM backbone with Transformer-based architectures to better capture long-range dependencies via self-attention; extending the framework to multilingual and cross-accent emotional speech corpora to improve generalisation; applying model-compression techniques such as pruning, quantisation and knowledge distillation to enable real-time deployment on mobile, wearable and embedded platforms; and incorporating multimodal cues such as facial expression, physiological signals and gesture alongside speech to build more robust, real-world emotion-aware systems. Explainable AI techniques will also be investigated to improve the interpretability of the model's classification decisions.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

REFERENCES

- [1] S. Latif, R. Rana, S. Younis, J. Qadir, and J. Epps, "Transfer learning for improving speech emotion classification accuracy," IEEE Access, vol. 9, pp. 18185-18195, 2021.
- [2] H. M. Fayek, M. Lech, and L. Cavedon, "Evaluating deep learning architectures for speech emotion recognition," Neural Netw., vol. 92, pp. 60-68, 2021.
- [3] Z. Zhang, B. Schuller, and S. Wengner, "Deep learning-based speech emotion recognition: A survey," IEEE Trans. Affective Comput., vol. 13, no. 2, pp. 617-639, 2022.
- [4] Y. Kim and H. Kim, "Speech emotion recognition using CNN-LSTM hybrid networks," Expert Syst. Appl., vol. 186, pp. 115-128, 2022.
- [5] J. Deng, X. Xu, and B. Schuller, "Deep neural networks for speech emotion recognition," Inf. Fusion, vol. 82, pp. 145-160, 2022.
- [6] S. Hochreiter and J. Schmidhuber, "Long short-term memory," Neural Comput., vol. 9, no. 8, pp. 1735-1780, 1997.
- [7] D. P. Kingma and M. Welling, "Auto-encoding variational Bayes," in Proc. Int. Conf. Learning Representations (ICLR), 2014.
- [8] T. N. Sainath and B. Li, "Deep convolutional neural networks for speech recognition," IEEE Signal Process. Mag., vol. 29, no. 6, pp. 82-97, 2021.
- [9] B. Logan, "Mel frequency cepstral coefficients for speech recognition," in Proc. Int. Symp. Signal Processing (ISSP), 2020, pp. 1-5.
- [10] F. Eyben, M. Wollmer, and B. Schuller, "OpenSMILE: Audio feature extraction toolkit," in Proc. ACM Multimedia, 2021, pp. 1459-1462.
- [11] I. Goodfellow, Y. Bengio, and A. Courville, Deep Learning. Cambridge, MA, USA: MIT Press, 2016.
- [12] R. B. Patel and S. Sharma, "Speech emotion recognition using MFCC and deep learning," IEEE Access, vol. 10, pp. 45122-45135, 2022.
- [13] K. Han, D. Yu, and I. Tashev, "Speech emotion recognition using deep neural networks," in Proc. IEEE Int. Conf. Acoustics, Speech and Signal Processing (ICASSP), 2021, pp. 5800-5804.
- [14] M. Neumann and N. Vu, "Improving speech emotion recognition with attention mechanisms," in Proc. INTERSPEECH, 2022, pp. 1003-1007.
- [15] Y. Huang and J. Wang, "A hybrid CNN-LSTM framework for speech emotion recognition," Knowledge-Based Syst., vol. 241, pp. 108-120, 2023.
- [16] P. Verma and R. Singh, "Deep learning techniques for speech emotion recognition: A comparative study," Appl. Soft Comput., vol. 126, 2023.
- [17] A. Vaswani et al., "Attention is all you need," in Proc. Advances in Neural Information Processing Systems (NeurIPS), 2017.
- [18] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," in Proc. Int. Conf. Learning Representations (ICLR), 2015.
- [19] TensorFlow Developers, "TensorFlow: An end-to-end open source machine learning platform." [Online]. Available: <https://www.tensorflow.org>
- [20] Librosa Development Team, "Librosa: Python library for audio and music analysis." [Online]. Available: <https://librosa.org>



INTERNATIONAL
STANDARD
SERIAL
NUMBER
INDIA



INTERNATIONAL JOURNAL OF INNOVATIVE RESEARCH

IN COMPUTER & COMMUNICATION ENGINEERING

 9940 572 462  6381 907 438  ijircce@gmail.com



www.ijircce.com

Scan to save the contact details