

Emotion Detection by Voice Analysis

Vaibhav Sharma¹, Vaibhav Sharma², Seema Mehla³ and Virender Kumar Mehla⁴

¹⁻³GLA University, Mathura, India

Email: {vaibhav.sharma_cs21, vaibhav.sharma2_cs21}@gla.ac.in, seemamehlakkr@gmail.com

⁴Sanskriti University, Mathura, India

Email: virender.mehla@gmail.com

Abstract— Emotion detection by voice, needed for human-computer interaction (HCI), is an emerging field that leverages advanced machine learning techniques and audio processing to identify human emotional states from vocal cues. This technology is capable of distinguishing various emotions by analyzing audio features like pitch, tone, intensity, and speech rate. Emotion-oriented systems can revolutionize industries such as healthcare, customer service, and interactions between people and computers by improving empathy and personalizing answers. This paper explores the design, development, and application of an emotion detection system using voice data, presenting technical challenges, solutions, and the societal implications of integrating emotion recognition into everyday technologies.

Index Terms—Emotion recognition, voice analysis, speech processing, machine learning, affective computing, Convolutional neural networks (CNN), Long Short-Term Memory (LSTM), Mel-frequency cepstral coefficients (MFCC)..

I. INTRODUCTION

Emotion detection by voice aims to bridge the gap between human emotional expression and machine interpretation. The science of Artificial Intelligence aims at making intelligent machines and by utilizing machine learning (ML) and signal processing techniques, the system processes audio signals to extract features indicative of emotional states. As speech is inherently emotional, it carries subtle cues that reveal the speaker's psychological state. This research focuses on applying these cues to develop systems that can detect and classify emotions in a variety of real-world applications, from healthcare to customer service.

The motivation for this paper is based on the increased demand for a system that can understand and respond to human emotions. In fields like mental health, emotion recognition can provide insights into mood fluctuations and stress levels. In customer service, real-time emotion detection can lead to more responsive and personalized customer interactions. Moreover, in virtual assistants and interactive gaming, integrating emotional intelligence can create more natural and human-like interactions with machines.

II. RELATED WORK

A. Existing Emotion Detection Systems

Emotion detection from voice signals has seen significant advancements, particularly in machine learning, signal processing, and audio feature extraction techniques. Several key algorithms and technologies have contributed to the development of these systems.

Speech Emotion Recognition (SER): Speech Emotion Recognition (SER) focuses on classifying emotions based on the features derived from speech signals. The primary goal of SER is to enable machines to recognize the emotional state of a speaker from their voice[11]. SER is becoming an important sub-discipline of human-machine interaction, that could bring benefits to different areas, such as minimizing fatigue driving, assisting some medical diagnosis, and collecting feedback from students in online education. [1]

Key datasets such as RAVDESS (Ryerson Audio-Visual Database of Emotional Speech and Song) and TESS (Toronto emotional speech set) are commonly used to provide emotional speech samples for training deep learning models. These datasets offer a wide variety of emotional expressions, helping systems generalize across diverse emotional cues. Commonly used algorithms in SER systems include:

- Support Vector Machines (SVMs) and Random Forests are traditional machine learning models frequently used for emotion classification, as they can handle high-dimensional feature spaces effectively. These methods often perform well when there is a clear boundary between different emotions, but they may not capture more complex patterns in speech.
- Neural Networks, particularly Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs), especially Long Short-Term Memory (LSTM) networks, are now widely adopted for processing sequential speech data. These models excel at recognizing temporal dependencies and emotional cues present in continuous speech, making them ideal for speech emotion recognition tasks where the emotional content evolves over time.

Deep Learning-Based Approaches: With recent advancements in deep learning, researchers have begun employing more sophisticated architectures, particularly CNNs and LSTMs, for both feature extraction and the recognition of temporal dependencies in emotional speech.

- The overall accuracy of the CNN baseline model is 73%. This indicates that in 73 out of 100 test samples, the emotion is correctly classified [2].

Wav2Vec2, a self-supervised model developed for speech recognition, has demonstrated its potential in extracting emotional information from raw audio by learning representations without the need for labeled data. The model's ability to process and understand speech at a high level makes it an appealing tool for detecting emotions. Furthermore, deep learning models like Transformers and Attention Mechanisms have shown promise in enhancing performance, allowing for more nuanced understanding of speech emotion dynamics. The ability of deep learning systems to automatically learn relevant functions from data reduces the need for handcrafted functional engineering and improves overall accuracy.

Multimodal Emotion Recognition: To improve the precision and reliability of emotion detection, modern systems increasingly integrate speech with additional data sources like text and facial expressions. This approach has given rise to advanced emotion recognition techniques. This approach recognizes that emotional expression can be conveyed through more than just voice, incorporating facial expressions, body language, and even contextual text. Platforms such as Beyond Verbal and Cogito adopt a multimodal framework, combining voice analysis with facial expressions and behavioral signals to achieve a deeper and more accurate assessment of emotional states [10].

Real-Time Emotion Detection: Real-time emotion detection systems, such as Affectiva and RealTime Emotion, focus on processing voice signals in real-time with low latency to deliver immediate sentiment analysis.

As IEMOCAP dataset contains 50 utterances in average per conversation, in the worst possible scenario it is justified to assume the time taken for retrieving an emotion for an arbitrary utterance, is the same as the time spent to retrieve emotion predictions for a conversation having 50 utterances on average. The results suggest approximately 0.7s of latency per utterance for passing through the complete pipeline [3]. These systems require efficient algorithms capable of performing complex analysis within milliseconds to ensure smooth user interaction. However, delivering low-latency processing remains a significant challenge, particularly in resource-constrained environments like mobile devices and edge computing platforms, where limited computational power restricts performance.

B. Challenges in Emotion Detection

While voice-based emotion recognition has seen remarkable progress, significant key challenges persist for researchers and engineers working in this field:

Noisy Environments: One of the major obstacles to accurate emotion detection is the presence of background noise, which can severely degrade the performance of emotion recognition systems, especially in crowded or dynamic environments. In settings like busy offices, streets, or public transportation, extraneous sounds such as traffic noise or conversations can overwhelm the emotional cues in speech, making it difficult for the system to isolate the emotional content. While techniques such as noise reduction and speech enhancement algorithms

have made progress, achieving perfect noise cancellation in all scenarios remains a difficult challenge. To enhance emotion detection reliability in real-world settings, future research should focus on developing more robust noise-filtering techniques and adaptive models that maintain accuracy in acoustically challenging environments.

Cultural Differences in Emotional Expression: Emotional expression can vary significantly across different cultures and societies, which poses a challenge for developing universally accurate emotion detection systems. What may be perceived as a sign of happiness in one culture could be interpreted differently in another, depending on the cultural norms surrounding emotional expression. Moreover, datasets used to train these models often lack sufficient cultural diversity, which limits the system's ability to generalize across various populations. As a result, systems may perform well in one cultural context but fail in another. To address this issue, the creation of more culturally diverse datasets and the development of culturally adaptive models are essential for improving the cross-cultural applicability of emotion detection technologies.

Variability in Speech Patterns: The human voice and associated speech patterns can be characterized by a number of attributes, the primary ones being pitch, loudness or sound pressure, timbre, and tone [4]. Timbre is the perceived sound quality of a musical note, sound or tone [5]. Tone is the use of pitch in language to distinguish lexical or grammatical meaning – that is, to distinguish or to inflect words [6]. The variations make it challenging for emotion detection models to accurately recognize emotions across diverse speakers. For example, some individuals may express anger in a high-pitched voice, while others may do so in a deeper tone. Similarly, speech patterns can change due to factors such as age, gender, or regional dialects. For emotion detection systems to be effective, they must be capable of handling these variances, meaning that models need to be sufficiently flexible to account for the wide range of emotional expressions and speech styles seen across different speakers.

In order to deduce emotion, three test cases have been examined, corresponding to the three emotional states: normal emotional state, angry emotional state, and panicked emotional state. For carrying out the analysis, four vocal parameters have been taken into consideration: pitch, SPL, timbre, and time gaps between consecutive words of speech [7].

Imbalanced Datasets: A common issue in emotion detection is the presence of imbalanced datasets, where certain emotions, such as happiness or sadness, are over-represented, while others, such as surprise or fear, may be underrepresented. This imbalance leads to biased models that are more accurate at predicting emotions with higher frequency in the dataset, but less effective at recognizing emotions that are less common. To mitigate this issue, data augmentation techniques, such as generating synthetic data or applying transformations to existing data, can help balance the distribution of emotions in training datasets. Additionally, class balancing methods, such as cost-sensitive learning or oversampling underrepresented classes, can help address this bias, allowing the model to perform more equally across various emotions.

Model Limitations: Despite the advancements in machine and deep learning, many emotion detection models suffer from overfitting to the specific datasets on which they were trained. This makes it difficult for the models to generalize to new, unseen environments, or to detect subtle emotional expressions. Additionally, detecting subtle or overlapping emotions, such as frustration versus anger or sadness versus depression, remains a significant challenge. Models may struggle to differentiate between emotions that are context-dependent or that share similar acoustic features.

Real-Time Processing: Real-time processing requirements pose substantial implementation challenges for emotion recognition systems, especially when operating on resource-constrained platforms like mobile devices or embedded system. Real-time emotion detection requires low-latency processing, meaning the system must be capable of making quick, accurate predictions without introducing noticeable delays in interaction. Achieving this level of efficiency is especially difficult when dealing with large, complex deep learning models that require substantial computational power.

III. METHODOLOGY

A. Data Collection

The emotion detection model was trained using two publicly available and well-regarded speech datasets, each containing a wide variety of emotional speech samples. These datasets offer diverse emotional expressions and serve as a strong foundation for training and evaluating emotion recognition systems.

Ryerson Audio-Visual Database of Emotional Speech and Song (RAVDESS): The RAVDESS dataset comprises 2,452 audio clips (1,440 speech, 1,012 song) recorded by 24 professional voice actors (12 male, 12 female). The dataset includes a broad range of emotions such as happiness, sadness, anger, fear, surprise, calm, and disgust.

Each sample in RAVDESS is accompanied by annotations evaluating the emotional validity and intensity of the performance[9]. This emotional intensity evaluation ensures that the emotions expressed in the dataset are genuine and recognizable, making it a reliable resource for emotion detection model training. Additionally, the dataset also covers both speech and song, offering a richer range of vocal expressions. The varied nature of this dataset helps in improving the model's generalization across different emotional contexts.

Toronto Emotional Speech Set (TESS): The TESS dataset contains 2,800 audio recordings featuring 200 target words embedded in carrier phrases, produced by two female actors. The recordings encompass seven emotional states: anger, disgust, fear, surprise, happiness, sadness, and neutral. This dataset was designed to focus on neutral-accented emotional speech, which makes it particularly useful for studying emotion detection across a range of linguistic and cultural backgrounds. By using target words embedded within neutral carrier phrases, TESS provides a controlled way to examine emotional expression, helping ensure that the emotional content is conveyed through voice alone rather than contextual cues. This makes TESS a valuable resource for emotion classification tasks requiring clear emotional signals.

Both RAVDESS and TESS offer a diverse spectrum of emotional speech, making them ideal datasets for training an emotion classification model capable of recognizing a broad array of emotional states.

B. Preprocessing

The raw audio recordings were processed through multiple preprocessing stages to optimize them for feature extraction and model training. This pipeline serves two key purposes: enhancing speech signal quality by reducing noise and artifacts, and converting the audio into an appropriate representation for deep learning architectures.

Noise Filtering: Background noise can significantly degrade the quality of audio signals, especially in real-world environments with a lot of extraneous sounds. To minimize the impact of noise, techniques such as spectral gating and voice activity detection (VAD) were applied. Spectral gating helps in isolating the speech signal from unwanted noise by filtering out frequencies that are not part of the vocal signal. Voice activity detection is used to identify the sections of the audio signal where speech is present, discarding periods of silence or background noise. These techniques ensure that only relevant speech data is passed to the model, improving the accuracy of emotion detection.

Feature Extraction: Essential features from the audio data were extracted using the LibROSA library, which is a popular Python package for audio processing. Key features extracted include:

- Mel-frequency cepstral coefficients (MFCCs), which capture the timbral texture of speech and are commonly used in speech processing.
- Chroma features, which represent the harmonic and melodic aspects of audio signals, useful for distinguishing emotional tonal qualities.
- Pitch, which can reveal emotional states through variations in voice frequency, such as a higher pitch when expressing excitement or fear.

These features were selected due to their relevance in emotion recognition, as they can effectively capture both the acoustic properties of speech and the emotional nuances embedded in the vocal performance.

Normalization: After feature extraction, the extracted features were normalized to ensure that they were on a consistent scale. This step is crucial because variations in speech loudness, speed, and other factors could lead to inconsistent data. By normalizing the features, the model can focus on the inherent emotional signals in the speech rather than being affected by variations in speech characteristics that are not related to emotion. This helps in improving the robustness and performance of the emotion classification model.

C. Model Architecture

The emotion classification model for this project employed a Deep Learning-based approach, utilizing both Convolutional Neural Networks (CNNs) and Long Short-Term Memory networks (LSTMs). These architectures were selected due to their ability to effectively analyze sequential speech data and capture the intricate patterns inherent in emotional speech.

CNN: CNNs were used for feature extraction from the Mel spectrograms of the audio data. A Mel spectrogram is a visual representation of the audio signal where the intensity of the frequencies is displayed over time. CNNs are well-suited for extracting spatial patterns from these spectrograms, helping the model to learn distinctive features of the emotional speech signal. The ability of CNNs to detect local patterns in spectrograms allows them to recognize key acoustic features related to emotions, such as tone, pitch, and rhythm.

LSTM: LSTMs were employed to capture temporal dependencies in the speech signal. Unlike conventional neural networks, Long Short-Term Memory (LSTM) architectures are specifically engineered to process

sequential information, making them particularly suitable for analyzing speech signals due to their inherent temporal dependencies[12]. Emotions in speech often emerge over time through variations in tone, pitch, and pacing, and LSTMs are capable of retaining long-term dependencies in the data. This allows the model to recognize emotional states that develop gradually, such as sadness or anger, based on the temporal evolution of the speech signal.

This combination allows the model to effectively process sequential voice data and detect emotional states.

D. Some Common Mistakes

Data Imbalance: Some emotions (e.g., sadness or calm) may have fewer samples, leading to biased models and poor performance on underrepresented emotions.

Noise in Audio: Inadequate noise reduction can harm model accuracy. Effective preprocessing techniques like spectral subtraction are necessary.

The extracted features include frequency, wavelength, period, speed of travel, intensity, loudness, echo, pressure, amplitude, and resonance, among others [8].

Overfitting: Deep learning models may memorize training data, leading to poor generalization. Techniques like dropout or early stopping help prevent overfitting.

Insufficient Feature Extraction: Missing crucial features (e.g., MFCC, pitch) can hinder emotion detection. Comprehensive feature extraction is key for performance.

Ignoring Multilingual and Cultural Differences: Emotional expression varies across languages and cultures. Ignoring these differences can lead to inaccurate detection in multilingual settings.

Improper Evaluation Metrics: Dependence on classification accuracy alone often yields deceptive performance estimates, particularly when evaluating models on imbalanced datasets where class distributions are skewed. It's important to also report precision, recall, and F1 score.

E. Abbreviations and Acronyms

AI: Artificial Intelligence

CNN: Convolutional Neural Network

MFCC: Mel-Frequency Cepstral Coefficients

LSTM: Long Short-Term Memory

RAVDESS: Ryerson Audio-Visual Database of Emotional Speech and Song

TESS: Toronto Emotional Speech Set

VAD: Voice Activity Detection

HCI: Human-Computer Interaction

F1 Score: The harmonic mean of precision and recall, a performance metric for classification models represents the harmonic mean of precision and recall, serving as a balanced evaluation metric for classification models that accounts for both false positives and false negatives.

API: Application Programming Interface

SVM: Support Vector Machines

PCD: Principal Component Decomposition

F. Units

Time: Seconds (s)

Frequency: Hertz (Hz)

Audio Duration: Seconds (s)

Mel-frequency scale: Unitless (used for audio feature extraction)

Accuracy: Percentage (%)

Loss: Unitless (commonly used in deep learning models)

F1 Score: Unitless (range between 0 and 1)

Precision/Recall: Unitless (percentage or ratio between 0 and 1)

Intensity: Decibels (dB)

G. Equations

MFCC Calculation: The Mel-Frequency Cepstral Coefficients (MFCCs) are calculated through a series of steps, including Fourier Transform, Mel scale filter banks, and Discrete Cosine Transform (DCT):

$$\text{MFCC} = n = 1 \sum N \log |X(\text{fn})|. W(\text{fn}). \quad (1)$$

Where:

$X(f_n)$ is the Fourier-transformed signal,

$W(f_n)$ is the Mel filter bank,

N is the number of frequency bands.

Cross-Entropy Loss (commonly used in classification problems):

$$\text{Loss} = -i \sum y_i \log(y') \quad (2)$$

Where:

y_i is the actual class label (0 or 1 for binary classification or one-hot encoded for multiclass),

$y_i^{\wedge}$ is the predicted probability for class i .

Convolutional Neural Network (CNN) Output: The output of a CNN layer is determined by the following equation for each neuron:

$$\text{Output} = f(W \cdot X + b) \quad (3)$$

Where:

W is the weight matrix,

X is the input,

b is the bias,

f is the activation function (e.g., ReLU, Sigmoid).

Accuracy Calculation:

$$\text{Accuracy} = \frac{\text{Number of Predictions}}{\text{Number of Correct Predictions}} \times 100 \quad (4)$$

TABLE I. PERFORMANCE COMPARISON OF EMOTION DETECTION MODELS

<i>Model</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
MLP (Multilayer Perceptron)	86%	94%	95%	94%
Gaussian Naive Bayes	51%	88%	66%	56%
SVM	85%	93%	94%	93%
Decision Tree	67%	73%	70%	69%
Random Forest	82%	87%	88%	87%
CNN-LSTM	88%	90%	89%	90%

IV. RESULTS

The emotion recognition system was comprehensively assessed using multiple evaluation metrics: classification accuracy, precision (positive predictive value), recall (sensitivity), and the F1-score (harmonic mean of precision and recall). The results from our deep learning-based model (CNN-LSTM) were compared against different models to demonstrate the effectiveness of the proposed approach. The evaluation was based on two key speech datasets: RAVDESS and TESS.

The CNN-LSTM model demonstrated strong performance across all evaluation metrics. The accuracy of the model reached 88%, with precision, recall, and F1-score values of 90%, 89%, and 90%, respectively. This highlights the superior performance of CNN-LSTM in accurately classifying emotions based on speech data.

Similarly, the Random Forest model demonstrated strong performance across all evaluation metrics, achieving 82% accuracy with closely balanced precision (87%) and recall (88%), resulting in an F1-score of 87% that confirms robust classification consistency.

Random Forest, as an ensemble method, benefits from multiple decision trees to reduce variance, offering balanced performance across various emotional categories.

The SVM model exhibited superior performance, attaining 85% classification accuracy with balanced precision (93%) and recall (94%) metrics, culminating in an F1-score of 93% that demonstrates exceptional predictive consistency.

The Multi layer Perceptron (MLP) classifier achieved competitive performance, attaining 86% accuracy with well-balanced precision (94%) and recall (95%) scores, resulting in an 94% F1-score that indicates consistent classification capability across all evaluation metrics. MLP, a fully connected neural network, captures non-linear patterns in the data, though it generally falls short compared to more specialized models like CNN-LSTM, which handle temporal data more effectively.

The Gaussian Naive Bayes model delivered modest but meaningful results (51% accuracy, 88% F1-score), showing its typical high recall (66%) despite lower precision (56%). Although Naive Bayes is based on strong assumptions about the independence of features, it still offers reasonable performance for emotion detection. The Decision Tree model demonstrated reasonable performance, attaining 67% classification accuracy with well-balanced precision (73%) and recall (70%) metrics, culminating in an 69% F1-score that indicates consistent predictive capability across all evaluation dimensions. While Decision Trees are interpretable and simple, they tend to overfit, which can reduce their generalization ability in complex tasks like emotion detection. Overall, the CNN-LSTM, MLP models emerges as the most powerful models for emotion detection, with superior results in all evaluated metrics, followed by the ensemble models (SVM, Random Forest) that also show promising performance. These findings suggest that CNN-LSTM, MLP models offer an effective approach for building emotion detection systems capable of handling a wide range of emotional states, improving human-computer interaction applications in diverse industries such as healthcare, customer service, and security.

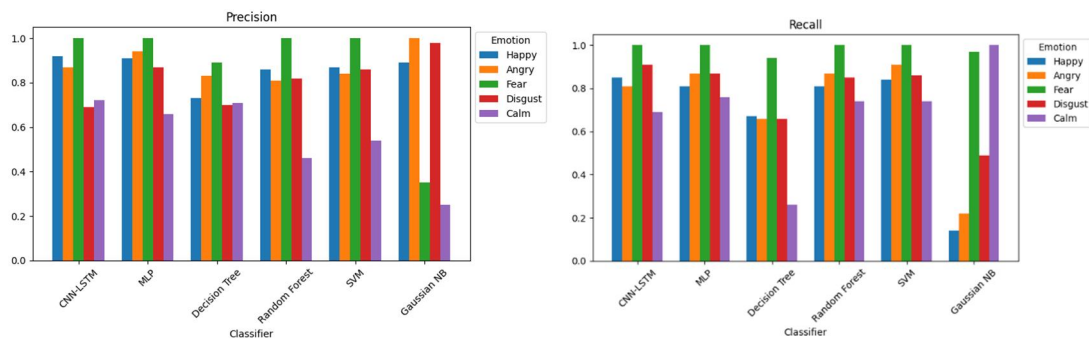


Figure 1. Precision of Emotion Detection Across Different Emotion Categories Figure 2. Recall of Emotion Detection Across Different Emotion Categories

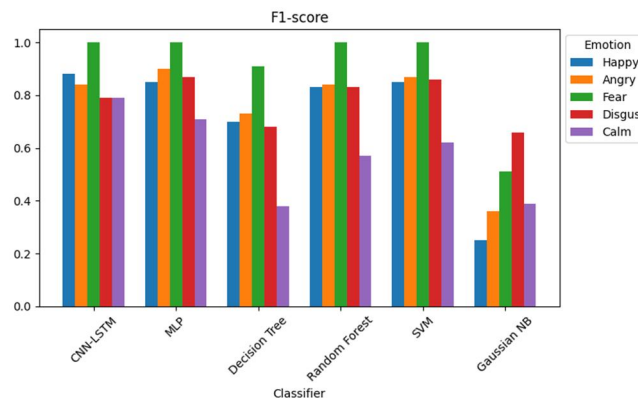


Figure 3. F1-Score of Emotion Detection Across Different Emotion Categories

V. CONCLUSION

The "Emotion Detection by Voice" paper demonstrated the use of deep learning and audio processing to identify emotions from speech. By utilizing the RAVDESS and TESS datasets, the model successfully recognized emotions like happiness, sadness, and anger. Techniques like noise filtering, feature extraction, and CNN-LSTM architectures proved effective in analyzing emotional expression.

Future improvements could involve integrating more diverse datasets and exploring additional voice features, such as prosody and speech rate, to enhance emotion detection across languages and contexts. The technology holds great potential in areas like mental health, customer service, and human-computer interaction.

ACKNOWLEDGMENT

We sincerely thank the creators of the RAVDESS and TESS datasets for providing the foundation for our emotion detection model. We also appreciate the guidance of our research mentors and the computational resources provided by cloud platforms. Special thanks to our colleagues for their valuable feedback and support.

REFERENCES

- [1] Samaneh Madanian, Talen Chen, Olayinka Adeleye, John Michael Templeton, Christian Poellabauer, Dave Parry, Sandra L. Schneider, "Speech emotion recognition using machine learning — A systematic review", 2019, International Journal of Speech Technology, pp. 1-18.
- [2] Ashish Singh, Sohib, Nagula Prabhath, Ravi Kumar Pandey, Abhishek Yadav, Chanan Singh, "Emotion Analysis from Voice Signals: A Machine Learning Approach", Journal of Machine Learning, 2018, pp. 12-28.
- [3] Chamishka, S., Madhavi, I., Nawaratne, R., Alahakoon, D., De Silva, D., Chilamkurti, N., & Nanayakkara, V., "A voice-based real-time emotion detection technique using recurrent neural network empowered feature modelling", Futuristic Trends and Innovations in Multimedia Systems Using Big Data, IoT and Cloud Technologies (FTIMS), 2022, pp. 1-9.
- [4] Trevor R Agus, Clara Suied, Simon J Thorpe, Daniel Pressnitzer. "Characteristics of human voice processing", IEEE International Symposium on Circuits and Systems (ISCAS), 2010, pp.509-512.
- [5] Richard Lyon & Shihab Shamma. "Auditory Representation of Timbre and Pitch.", Harold L. Hawkins & Teresa A. McMullen. Auditory Computation. Springer, 1996, pp. 221–23.
- [6] R.L. Trask, "A Dictionary of Phonetics and Phonology", Routledge, 2004, pp. 280-292.
- [7] Poorna Banerjee Dasgupta, "Detection and Analysis of Human Emotions through Voice and Speech Processing", 2017, Journal of Voice Research, pp. 32-48.
- [8] Y. Ü. Sonmez and A. Varol, "New Trends in Speech Emotion Recognition", 7th International Symposium on Digital Forensics and Security (ISDFS), Barcelos, Portugal, 2019, pp. 1-7.
- [9] Aditi Agrawal, Anagha Jain, Balpreet Kaur, Saumya Jangid, Karuna Kadian, Vimal Dwived, "Comparative Analysis of Speech Emotion Recognition Models and Technique", Journal of Ambient Intelligence and Humanized Computing, 2023, pp. 1-15.
- [10] Jiachen Luo, Huy Phan, Joshua Reiss, "Cross-Modal Fusion Techniques for Utterance-Level Emotion Recognition from Text and Speech", ICASSP IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), 2023, pp. 3-8.
- [11] Sudharshanam Buddha, Rohan Sawant, Sumit Singh Ingle, Rushikesh Suryawanshi, Geeta Atkar, "Emotion Sense: - Real-time Speech Emotion Recognition for live calls", Second International Conference on Advances in Information Technology (ICAIT), 2024, pp. 5-7.
- [12] Harshawardhan S. Kumbhar, Sheetal U. Bhandari, "Speech Emotion Recognition using MFCC features and LSTM network", 5th International Conference On Computing, Communication, Control And Automation (ICCUBE), 2019, pp. 1-10.