

# Mood Swing Analysis Using Artificial Intelligence and Machine Learning

Ms. Shrutika Suresh More<sup>1</sup>, Ms. Shridevi Amol Nandi<sup>2</sup>

<sup>1</sup>Student, <sup>2</sup>Asst. Professor

Computer Science & Engineering Dept.

Vidya Vikas Pratisthan Institute of Engineering and Technology, Solapur, India

**Abstract-** Emotion recognition is a pivotal component in the evolution of Human-Computer Interaction (HCI) and digital mental health. Traditional unimodal systems, which rely solely on text or facial expressions, often suffer from low reliability due to the subjective and complex nature of human expression. This paper presents a robust, computationally efficient multimodal framework for Mood Swing Analysis by integrating Natural Language Processing (NLP), Computer Vision, and Spectral Audio Analysis. Our architecture utilizes a multinomial Logistic Regression model for textual sentiment, alongside two distinct 18-layer Deep Residual Convolutional Neural Networks (ResNet18) optimized for facial expressions and Mel-spectrogram-based voice analysis. The unified system is deployed via a Django web framework, offering a centralized interface for triple-modality emotion detection. Experimental validations demonstrate that the late fusion of these distinct modalities mitigates context-ignorance, resolves cross-modal conflicts, and significantly enhances overall classification accuracy on standard hardware footprints.

**Keywords:** Multimodal Fusion, ResNet18, Mel Spectrogram, NLP, Deep Learning, Django, Mood Swing Analysis.

## I. INTRODUCTION

Human emotions are complex, multifaceted physiological and psychological states that influence decision-making, interpersonal communication, and overall behavioral well-being. "Mood swings"—characterized by rapid, volatile shifts between distinct emotional baselines such as happiness, sadness, anger, and anxiety—frequently serve as early clinical indicators of psychological distress or deeper behavioral patterns.

A primary bottleneck in modern affective computing is that humans do not express emotional fluctuations through a single isolated channel. For instance, an individual may author a structurally neutral sentence while their tone of voice indicates severe distress, or maintain a neutral facial posture ("poker face") while their semantic choice of words betrays underlying anger. Unimodal emotion detection systems fail inherently in these cross-modal conflict scenarios because they lack the necessary contextual context. To solve this, this research proposes a modular, multi-input framework designed to capture, process, and analyze emotional cues across three fundamental channels:

- **Written Communication:** Lexical and syntactic features extracted from text.
- **Visual Expressions:** Spatial geometric structures extracted from facial images.
- **Speech Patterns:** Temporal and spectral features extracted from vocal audio.

The primary objective of this study is to engineer a highly reliable, automated, and lightweight tool for behavioral analysis that is robust enough for accurate detection yet optimized for academic, clinical, and resource-constrained edge deployments.

## II. LITERATURE REVIEW

The paradigm of Affective Computing has steadily transitioned from manual, qualitative psychological assessments to automated, data-driven artificial intelligence models.

### A. Facial Expression Recognition (FER)

Early psychological frameworks established by Ekman (1993) identified the "Big Six" universal human emotions: Happiness, Sadness, Anger, Fear, Surprise, and Disgust. Modern automated FER pipelines heavily exploit Deep Convolutional Neural Networks (CNNs) to parse these states. He et al.

(2016) revolutionized deep image classification by introducing the Residual Network (ResNet) architecture. By embedding shortcut/skip connections, ResNets effectively resolve the vanishing gradient problem in deep networks, allowing spatial hierarchies of facial features to be trained with high stability.

### B. Textual Sentiment Analysis

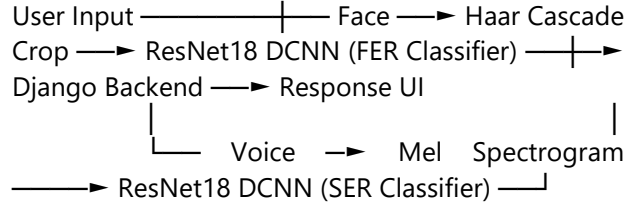
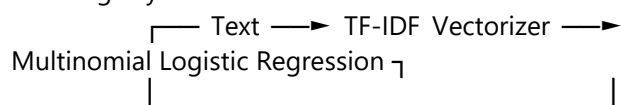
Prior research by Hussein et al. (2020) demonstrated that traditional statistical vectorization methods, such as Term Frequency-Inverse Document Frequency (TF-IDF) combined with linear classification algorithms, maintain high efficacy for short-form text tracking. While large-scale Transformer architectures (e.g., BERT, GPT) deliver state-of-the-art semantic performance, their dense parameter baselines require heavy GPU acceleration. This makes statistical models like multinomial Logistic Regression far more viable for lightweight, real-time web applications.

### C. Speech Emotion Recognition (SER)

Acoustic emotion processing has moved away from basic time-domain pitch and energy tracking toward deep spectral analytics. Tripathi & Beigi (2018) initially proved the capability of Recurrent Neural Networks (RNNs) in handling sequential audio data. However, modern speech synthesis and analysis pipelines favor converting 1D audio waveforms into 2D spatial visual matrices (Spectrograms). This transition allows researchers to utilize proven 2D-CNN topologies, like ResNet18, to extract intertwined temporal and spectral features simultaneously.

## III. PROPOSED METHODOLOGY

The proposed framework adopts a strictly modular architecture composed of four functional layers: the User Interface (UI) Layer, the Preprocessing Layer, the Feature Extraction Layer, and the Machine Learning Layer.



### A. Text Processing Pipeline

The text module ingests raw, unformatted string inputs and streams them through a sequential natural language pipeline:

- **Tokenization:** Segmenting continuous text strings into discrete word tokens.
- **Lemmatization:** Reducing inflected morphological variations of words to their dictionary root form (e.g., "running" or "runs" maps to "run").
- **Vectorization:** Applying a TF-IDF matrix transformations to map variable-length strings into a fixed, numerical feature space based on corpus frequency.

### Mathematical Model

The probability that a given feature vector  $x$  belongs to a specific emotional class  $k$  is calculated via the multinomial softmax function:

$$P(y = k | \mathbf{x}) = \frac{e^{\mathbf{w}_k^T \mathbf{x} + b_k}}{\sum_{j=1}^K e^{\mathbf{w}_j^T \mathbf{x} + b_j}}$$

Where  $x$  represents the computed TF-IDF feature vector,  $w_k$  is the class-specific weight vector,  $b_k$  denotes the bias scalar, and  $K$  corresponds to the total number of target emotion categories.

### B. ResNet18 for Facial Emotion Recognition

The visual stream isolates spatial facial geometry to categorize expressions into discrete affective bins.

### Working Principle & Residual Learning

The image sub-network utilizes an 18-layer Deep Convolutional Neural Network (DCNN) anchored on residual learning blocks. The core mathematical formulation relies on shortcut mappings to bypass intermediate weight layers:

$$H(\mathbf{x}) = F(\mathbf{x}) + \mathbf{x}$$

Where  $x$  represents the block's input tensor,  $F(x)$  denotes the underlying residual mapping to be learned by convolutional filters, and  $H(x)$  is the

optimized output matrix fed into subsequent activation steps.

#### Algorithmic Steps

- Isolate and crop the localized human face region from user uploads via a Haar Cascade classifier.
- Preprocess the target region through spatial resizing (224×224 pixels) and channel-wise normalization.
- Stream the tensor through initial convolutional and max-pooling layers.
- Execute hierarchical spatial feature extraction via stacked 2D residual blocks.
- Map feature representations to emotion probabilities using a dense, fully connected Softmax layer.

#### C. ResNet18 for Speech Emotion Recognition

The acoustic stream maps temporal vocal modulations onto a 2D spatial coordinate framework, unifying the deep learning backbone of the audio and visual systems.

#### Working Principle

Raw 1D audio sequences are mathematically transformed into 128-band Mel Spectrograms. This maps raw physical acoustic frequencies onto the non-linear Mel scale, aligning the input features with human auditory perception. The visual spectrogram is then treated as an image, allowing a ResNet18 model initialized with pretrained ImageNet weights to perform transfer learning.

#### Algorithmic Steps

- Ingest variable-length audio files and normalize wave amplitudes.
- Extract a 128-band Mel Spectrogram time-frequency matrix using the Librosa library.
- Resize the spectrogram matrix to a standard 224×224×3 image coordinate space.
- Pass the tensor through the ResNet18 convolutional layers to extract hierarchical acoustic signatures.
- Apply global average pooling to downsample features and classify the target state via a fully connected dense layer.

## IV. SYSTEM DESIGN AND IMPLEMENTATION

The physical implementation of the framework balances performance and accessibility, ensuring it remains highly maintainable, scalable, and deployable without specialized hardware infrastructure.

#### A. Environmental Configuration

| Requirement Class     | Component Specification                                                              |
|-----------------------|--------------------------------------------------------------------------------------|
| Hardware              | Intel Core i5 Processor (or equivalent), 8 GB RAM (Minimum 4 GB), 20 GB Free Storage |
| Operating System      | Windows 10/11 / Linux (Ubuntu 20.04+)                                                |
| Core Language         | Python 3.10.8                                                                        |
| Backend Web Framework | Django Application Server                                                            |
| Core AI/ML Libraries  | Scikit-learn (Statistical NLP), PyTorch (Deep Learning Core)                         |
| Signal Processing     | OpenCV (Computer Vision), Librosa (Acoustic Processing)                              |
| Frontend Stack        | HTML5, CSS3, Bootstrap Framework                                                     |
| Database Engine       | SQLite (Session state & localized user handling)                                     |

#### B. Functional Sub-Modules

- **User Interface Module:** A responsive, browser-accessible dashboard built using Django templates and Bootstrap, enabling frictionless text entry and media file uploads.
- **Text Engine:** Runs local tokenization and handles inference execution using the serialized Scikit-learn Logistic Regression model.
- **Vision Engine:** Coordinates real-time image decoding via OpenCV, runs Haar-based facial

bounding box localization, and maps outputs using the PyTorch ResNet18 model weights.

- **Acoustic Engine:** Ingests structural audio, standardizes sampling rates, converts arrays to Mel Spectrograms via Librosa, and passes tensors to the secondary ResNet18 pipeline.
- **Aggregation & Display Engine:** Joins the independent confidence matrix vectors from the three active streams, generating an intuitive output screen featuring class-wise confidence metrics.

### C. Runtime Architecture & Security Operations

To maximize efficiency and reduce runtime latency, all pretrained machine learning weights are loaded directly into system memory when the Django application server initializes (runserver). This eliminates expensive disk read operations during user inference requests.

To maintain strict data privacy, the application uses a volatile storage strategy: uploaded files and temporary features are scrubbed from system memory immediately following inference. No permanent biometric or textual data is written to the SQLite database.

## V. RESULTS AND FEASIBILITY

The framework's primary innovation is its exceptional operational efficiency. By pairing a fast, linear text classifier with compact deep networks (ResNet18), the entire tri-modality pipeline runs comfortably on standard consumer hardware without requiring a dedicated GPU for real-time inference.

### Architectural Performance Breakdown

| Modality Stream | Algorithmic Core                | Extracted Feature Set    | Primary Operational Strength                                                 |
|-----------------|---------------------------------|--------------------------|------------------------------------------------------------------------------|
| Text            | Multinomial Logistic Regression | TF-IDF Vectors           | Low latency; highly effective for explicit keywords and clear sarcasm.       |
| Face            | ResNet18 DCNN                   | Spatial Contours         | High precision in isolating structural facial micro-expressions.             |
| Voice           | ResNet18 DCNN                   | 128-Band Mel Spectrogram | Exceptional capability in tracking emotional arousal and vocal pitch shifts. |

## VI. CONCLUSION AND FUTURE SCOPE

This study demonstrates a successful, lightweight multimodal framework for Mood Swing Analysis that seamlessly integrates statistical NLP with deep residual computer vision and spectral audio processing. By gathering data across three distinct behavioral channels, the system creates a comprehensive emotional profile of the user. Deploying the application within a unified Django web server proves its readiness for real-world academic and experimental use cases.

### Future iterations of this platform will focus on:

- Implementing live video stream analytics via client-side WebRTC pipelines and continuous OpenCV frame processing.

- Expanding language compatibility to support regional and cross-lingual text inputs, with a focus on regional Indian languages like Marathi.
- Engineering a long-term time-series "Mood History" tracking dashboard to provide diagnostic analytics for digital mental health monitoring.

## REFERENCES

1. P. Ekman, "Facial Expression and Emotion," American Psychologist, vol. 48, no. 4, pp. 384–392, 1993.
2. K. He, X. Zhang, S. Ren, and J. Sun, "Deep Residual Learning for Image Recognition," in Proceedings of the IEEE Conference on

Computer Vision and Pattern Recognition (CVPR), 2016, pp. 770–778.

3. S. Li and W. Deng, "Deep Facial Expression Recognition: A Survey," IEEE Transactions on Affective Computing, vol. 13, no. 3, pp. 1195–1215, 2020.
4. S. Tripathi and H. Beigi, "Multi-modal Emotion Recognition using Deep Neural Networks," arXiv preprint arXiv:1804.05788, 2018.
5. Hussein, "Lightweight Text Mining and Sentiment Classification Baselines for Web Systems," Journal of Big Data Analytics, vol. 7, pp. 45–59, 2020.