

Generative AI Meets Responsible AI and Affective Computing

by

Atkea Fauzia Eva

20201105

Kamran Hassan Shomrat

21101010

Gazi Arman Islam

21101011

MD Saiful Islam

21101013

Jamilatun Subarna

21101069

A thesis submitted to the Department of Computer Science and Engineering
in partial fulfillment of the requirements for the degree of
B.Sc. in Computer Science

Department of Computer Science and Engineering

Brac University

June 2025

© 2025. Brac University

All rights reserved.

Declaration

It is hereby declared that

1. The thesis submitted is my/our own original work while completing degree at BRAC University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

Student's Full Name & Signature:

Atkea Fauzia Eva
20201105

Kamran Hassan Shomrat
21101010

Gazi Arman Islam
21101011

Md Saiful Islam
21101013

Jamilatun Subarna
21101069

Approval

The thesis/project titled “Generative AI Meets Responsible AI and Affective Computing” submitted by

1. Atkea Fauzia Eva (20201105)
2. Kamran Hassan Shomrat (21101010)
3. Gazi Arman Islam (21101011)
4. MD Saiful Islam (21101013)
5. Jamilatun Subarna (21101069)

Of Spring, 2025 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science on June, 2025.

Examining Committee:

Supervisor:
(Member)

Md. Golam Rabiul Alam

Professor
Dept. of Computer Science and Engineering
BRAC University

Head of Department:
(Chair)

Sadia H Kazi

Designation
Department of Computer Science and Engineering
BRAC University

Abstract

Generative AI, Responsible AI, and Affective Computing are transforming the future of artificial intelligence. The intersection of these fields represents a revolutionary breakthrough in computational technology. This thesis integrates these domains to develop a formalism for multidimensional emotional communication. By analysing image, voice, and text data, we address the challenge of detecting and generating emotions in real time, considering users' gestures and interactions. We adopt an integrated approach based on deep neural network models across multiple modalities: text sentiment analysis, audio emotion detection, and facial expression recognition. In particular, we built our proposed approach using transformer-based models, including DistilRoBERTa, fine-tuned Wav2Vec2 on custom dataset, and DeepFace to process text, audio, and facial expression respectively. These pretrained models are trained for emotion classification with 6.7 million, 95 million, and 120 million trainable parameters, respectively. Natural Language Processing (NLP) models are used to interpret meanings and sentiments in text, while audio and image-based models detect emotional cues. The system adapts dynamically based on user feedback and incorporates Responsible AI practices such as bias detection, ethical safeguards, and safe interactions to ensure fairness and trustworthiness. Through practical experimentation and evaluation, we demonstrate that it is possible to build Generative AI systems capable of not only perceiving and reacting to human emotions but also generating emotionally appropriate responses. Potential applications include virtual assistants, mental health support tools, interactive storytelling systems, and educational platforms where enhanced emotional intelligence can significantly improve user experience.

Keywords: Artificial Intelligence (AI), Machine Learning (ML), Natural Language Processing (NLP), Responsible AI, Ethical AI, Fairness in AI, Transparency in AI, Accountability in AI, Affective Computing, Emotion Recognition, Generative Adversarial Networks (GANs), Human-Computer Interaction (HCI), Sentiment Analysis

Acknowledgement

Firstly, we would like to acknowledge Md. Golam Rabiul Alam, our thesis supervisor, for his invaluable and incessant guidance, encouragement, and support that he has extended towards us throughout this research. His skills and insights have been critical to our shape and standards. Our heartfelt thanks also go to Anika Tahsin, our research assistant who has been a steadfast companion throughout with her immense support to help us in tackling the technicalities of our project. The detailed instructions that she provided and her quick responses helped us to carry out our thesis efficiently and successfully. Lastly, we would like to thank our fellows, Faculty Members, Tutors, and the staffs of CSE, BRAC University. For the past three years their assistance and commitment established the positive academic foundations that this research was built upon.

Table of Contents

Declaration	i
Approval	ii
Ethics Statement	iii
Abstract	iii
Acknowledgment	iv
Table of Contents	v
List of Figures	viii
List of Tables	ix
1 Introduction	1
1.1 Introduction	1
1.2 Problem Statement	2
1.3 Research Objective	2
2 Related work	4
2.1 Literature Review	4
2.1.1 Text-related Emotion Recognition	4
2.1.2 Image-based Emotion Recognition	5
2.1.3 Audio-based Emotion Recognition	5
2.1.4 Multimodal Classification and Fusion	5
2.1.5 Responsible AI in Emotion Recognition	6
2.1.6 Generative AI in Emotion Recognition	6
3 Dataset & Preprocessing	7
3.1 Data Description	7
3.2 Data Collection	7
3.2.1 Data Validation	7

3.3	Preliminary analysis	8
3.4	Data preprocessing	9
4	Model	12
4.1	Model Analysis	12
4.1.1	Text Model—j-hartmann/emotion- english-distilroberta-base	12
4.1.2	Image Model—DeepFace	13
4.1.3	Audio Model—facebook/wav2vec2-base	14
4.1.4	Multimodal Fusion Model—Majority Voting Strategy	15
4.1.5	Generative Feedback and Interpretation Model— google/flan-t5-large	15
5	Feature Extraction	18
5.1	Feature Extraction	18
5.1.1	Text Feature Extraction	18
5.1.2	Audio Feature Extraction	18
5.1.3	Visual Feature Extraction (Facial Expression Recognition)	20
5.1.4	Multimodal Feature Fusion	21
6	Proposed Methodology	22
6.1	Proposed Methodology	22
6.1.1	Dataset Preparation	22
6.1.2	Modality-wise Processing	23
6.1.3	Multimodal Fusion Logic	25
6.2	Responsible AI Layer	26
6.2.1	Bias Detection	26
6.2.2	Threat Detection	26
6.3	Generative Output Layer	26
6.4	Implementation Details	27
7	Results and Discussion	28
7.1	Results	28
7.1.1	Text Emotion Recognition Analysis	28
7.1.2	Audio Emotion Recognition Model Analysis	28
7.1.3	Image Emotion Recognition Analysis	29
7.1.4	Generative Model Role	30
7.1.5	Multimodal Fusion Analysis	30
7.1.6	Emotion-wise Analysis	30

7.1.7	Final Result Analysis	30
7.1.8	Discussion	36
7.1.9	Improvement Over Previous Models	36
8	Limitations and Improvements	38
8.1	Limitations and Improvements	38
8.1.1	Limitations	38
8.2	Future Work	39
9	Conclusion	40
9.1	Conclusion	40
	Bibliography	43

List of Figures

3.1	Voting System	8
3.2	Initially Acquired Data	9
3.3	Audio extraction from video; image clustering in every 20-frame interval.	10
6.1	Workflow Of Proposed Model	22
6.2	Frame wise Sentiment	25
6.3	Hard Voting Ensemble	26
6.4	Bias, Threat Detection - RAI	27
6.5	Summary of Results	27
7.1	visualization of our updated model precision, recall, and F1-score . . .	31
7.2	Previous Model Avg. Results	31
7.3	Classification Report	32

List of Tables

6.1	Summary of Attempted Audio Models	24
7.1	Classification Report for Multimodal Emotion Recognition	33
7.2	Comparison of Related Work	35
7.3	Improvement Over Previous Models with Fused Pipeline	36
7.4	Comparative Summary of Modalities in Emotion Recognition	37

Chapter 1

Introduction

1.1 Introduction

Generative AI enables the creation of novel content, such as text, images, and audio, using advanced deep learning models [26]. Responsible AI ensures these systems are developed and deployed ethically, prioritizing fairness, transparency, and safety to mitigate biases and harmful outputs [11]. Affective computing bridges human emotions and AI, enabling machines to recognize, interpret, and respond to emotional states [2]. This thesis integrates these domains to develop a multimodal emotion recognition system that processes visual, auditory, and textual inputs in real-time, ensuring empathetic and ethical human-computer interaction (HCI). Our system combines deep learning models for emotion recognition: ResNet and a custom CNN for image-based facial emotion recognition, and proposed DistilBERT and Wav2Vec2 for text and audio modalities, respectively [7]. Predictions are fused using a majority voting ensemble, enhancing accuracy across modalities [13]. Such as NSFW detection, face detection, and toxicity filtering can enforce responsible AI principles and lead to more ethical output [11]. FLAN-T5 is a generative component that generates human-readable summaries of the emotional states, thereby improving the user experience [28]. In real-time analysis where the emotional contexts are intricate and dynamic, empathetic responses can be tailored to the user's state [31]. Multimodal analysis is the combination of visual (facial expressions), auditory (voice tone), and textual (sentiment) cues to obtain a complete elucidation of emotions [2]. This shifts the focus from a static, homogenous view of emotional expression to one that considers individual context and breadth of emotion response. Through the establishment of well-defined guidelines set by experts, coupled with ensemble fusion, we can ensure that the responses are not only robust but also contextual depending upon the combinations of events that arise, paving the way through virtual assistants, psychological therapy, and educational platforms. In order to account

for these variations, we can set up some parameters that enable experts to identify which emotional responses would be effective based on the majority. In the process of implementing it into user scenarios, this will deepen the domain of affective computing and ensure that user experience is meaningful and responsive.

1.2 Problem Statement

The main reason is the complexity of extracting data from video: separating the audio and video frames from one video is not a trivial task, requiring at least advanced processing capabilities. Facial expression recognition (FER) methods depend largely on the face model, using hand-crafted features for extraction and post-processing [1], [4], [5]. These features fall into two categories: (i) geometry-based and (ii) appearance-based, while the post-processing takes into account the following filter types: feature fusion and feature selection. For example, Ghimire and Lee [4] use geometry-based FER with 52 facial landmark points (FLP) for automatic FER in facial sequences. The Kinect motion sensor providing depth information was used along with an active shape model (AAM)-based detection of face regions by Sujono and Gunawan [5]. Facial expressions are recognized from a fuzzy logic model using previous knowledge from the facial action coding system (FACS) [1]. To tackle these complexities, we utilize ResNet-50 [7] and custom CNNs on a Kaggle Dataset (15,453 images) [33].

Of course, AI-generated content is always a topic of concern regarding fairness, which may or may not be self-inflicted. One way to address bias is to use varied datasets. AI tends to be gender biased, as men appear older and more authoritative while women look happy and men show neutral or angry emotions. Such internalized biases in educational AI can influence students' aspirations and self-expectations [13]. To address these issues, we have employed a Responsible AI approach, of which NSFW detection, face detection, and toxicity filtering [11] are a part.

1.3 Research Objective

The goal of the research is to design a pragmatic AI model to achieve a strong combination of generative AI, responsible AI, and affective computing to produce contextually appropriate emotional-based outputs from multimodal data integrated into a robust AI model. It receives video inputs from the user, analyzes them to

extract text, audio, and image data from a dataset of 30 custom videos and from the Kaggle dataset (15,453 images) [33]. Facial expressions are identified using the DeepFace image classifier [6], audio using a custom-trained Wav2Vec2-based model [17], and text data is examined with a tuned DistilBERT model [15]. These model outputs are combined by a majority voting mechanism [13] to guarantee consistent emotion interpretation.

The main contribution is a video input real-time emotion recognition system, which produces emotion-aware responses, implemented through a Flask-based web application. Incorporating Responsible AI practices—such as our proprietary bias and threat detection with Detoxify [21] and various adversarial debiasing techniques [10] to mitigate emotional and gender-based biases—the system yields accurate, empathetic, and contextually appropriate interpretations, ensuring fair and equitable representations.

The generative part, using FLAN-T5 [28], transforms these inputs into ethical, empathetic outputs. It serves multiple applications with this flexible architecture:

1. **Mental Health Support through Emotional Analysis:** The model identifies emotional states from video inputs and detects early signs of distress or mental health issues, enabling targeted interventions [9].
2. **Interactive Storytelling:** Within integrated virtual storytelling platforms, the system adapts the narrative based on user emotions, resulting in a more immersive experience [3].
3. **Educational Platforms:** The system adjusts instructional content based on students’ emotional reactions analyzed through video, improving engagement and learning [8].

The model we develop here is therefore a step forward for human-AI interaction practices: technically well-functioning, ethically acceptable, and emotionally intelligent, providing an understanding of how intelligent and empathetic AI systems can be developed and used in different contexts of emotion-related AI [20].

Chapter 2

Related work

2.1 Literature Review

Recent advances in affective computing, generative AI, and responsible AI have provided the opportunity to develop systems that are ethically sound, emotionally intelligent, and technically complex. The unification of these disciplines has also spurred the creation of multimodal systems, which are able to recognize and comprehend resources across several data formats, including text, audio, and video. However, in an attempt to improve emotion detection systems, different kinds of modalities have been combined to enhance the accuracy and empathetic nature of AI systems. Bias, justice, and privacy are examples of ethical dilemmas that remain difficult to resolve and are still open challenges.

2.1.1 Text-related Emotion Recognition

Text is an important feature for emotion recognition systems, and with recent advancements in Natural Language Processing (NLP), models like DistilRoBERTa have been widely adopted for text comprehension and text generation tasks. This has recently been explored in the work of Rasool et al. [29], who examined the infusion of emotional intelligence through large language models, such as Flan-T5, LLaMA 2, and ChatGPT-4, into AI systems. Using emotional word dictionaries such as NRC and VADER, these models are capable of generating contextually relevant, empathic responses [29], thus making them quite useful for sensitive applications such as psychotherapy.

Haque et al. [24] utilized textual data with DistilRoBERTa to obtain integrated data for emotion classification. This provided research with evidence that using only the textual features of the data can achieve high performance (up to 90%) in emotion recognition tasks—in certain cases—for emotions like anger, sadness, and happiness.

2.1.2 Image-based Emotion Recognition

The problem of recognizing facial expressions has seen impressive progress with the help of deep learning approaches such as DeepFace. Taigman et al. **Taigman14** achieved near-human-level results for face verification. The model used 3D face feature alignment and applied it to a nine-layer neural network that was trained on millions of images. DeepFace was the new state-of-the-art in the field for some time, with a 97.35% accuracy in LFW.

DeepFace was actually created for face recognition, but its design has inspired newer systems for emotion recognition. Zhou et al. [30] proposed AffectEval, a system for human emotion recognition from physiological signals such as heart rate and skin conductance, providing a transparent and modular solution for emotion-aware AI.

2.1.3 Audio-based Emotion Recognition

Emotion detection from audio has experienced a tremendous surge, particularly with the introduction of models such as Wav2Vec2. EmoSense from Alice et al. [23] utilizes Wav2Vec2 and HuBERT for emotion recognition from raw audio recordings. EmoSense attained an accuracy of 93% in the classification of emotions such as anger, happiness, and sadness in the RAVDESS training dataset. Their findings demonstrate the vast potential of contemporary audio models, particularly deep learning models, in understanding emotional subtext in speech.

Furthermore, Pepino et al. [25] extended Wav2Vec2 with its internal layers, demonstrating that self-learned layers can provide better features for emotion classification compared to pre-tuned representations. Their findings contribute to the assumption that raw, self-taught representations may contain richer emotional signals, leading to the enhancement of emotion recognition accuracy.

2.1.4 Multimodal Classification and Fusion

Multimodal data integration has gained great performance in emotion recognition, utilizing text, audio, and visual data for recognition. Haque et al. [24] used a combined Bi-LSTM for speech and DistilRoBERTa for text. They achieved an audio and text accuracy of 88% and 90% respectively, showing the benefit of a multimodal emotion recognition approach. This combination enables the system to cope with the different complexities of human emotions, some of which are difficult to express in a specific communication modality.

Ensemble methods were also utilized to enhance overall classification performance. Zhou et al. [22] proposed an ensemble method that combines outputs from different CNNs, where the ensemble-based technique outperforms single models by a

considerable margin.

We propose a majority voting strategy to fuse the outputs of the audio, text, and visual models in our work. This ensures that the final emotional output using various modalities is robust and contextual.

2.1.5 Responsible AI in Emotion Recognition

As AI systems are used more often in sensitive areas like mental health, education, and social support, the ethical implications of emotion recognition systems have become more salient. AI-generated content can reinforce gender and emotion stereotypes and biases; for instance, women are often connected with happy emotions less so than neutral and anger emotions, whereas men are the exact opposite [32]. This can result in harmful biases, especially if those AI systems are used in high-stakes situations.

To avoid the risks mentioned, our system incorporates Responsible AI features such as adversarial debiasing and fairness measures to ensure this emotional output is guaranteed to be bias-free. We use detection tools such as Detoxify [21] to identify and filter toxic and biased behavior in real time, ensuring that system responses are both well-calibrated and ethical, especially in cases like psychotherapy or mental health, where the ethical nature of the responses and the possible biases in AI-generated responses can have a profound effect on users.

2.1.6 Generative AI in Emotion Recognition

Generative AI, especially models like Generative Adversarial Networks (GANs) and Transformer-based models, has opened up new avenues for emotion-aware systems. Rasool et al. [29] explored the use of large language models such as Flan-T5 and ChatGPT-4 to integrate emotional intelligence into AI systems. By using emotional word dictionaries like NRC and VADER and incorporating smart embedding and attention mechanisms, their models generated more empathetic responses in sensitive conversations, such as in psychotherapy.

The ability of Generative AI to produce contextually relevant emotional outputs based on the fusion of multimodal features is key to improving the user experience in applications such as virtual assistants, interactive storytelling, and personalized healthcare.

Chapter 3

Dataset & Preprocessing

3.1 Data Description

The effectiveness of the multimodal emotion identification system depends on the caliber, variety, and arrangement of the data used in each of the three modalities (text, picture, and audio). In order to accommodate the input format of the deployed models, we compiled a set of preprocessed and benchmark datasets. These include the following: the generations model (google/flan-t5-large), the speech model (facebook/wav2vec2-base), the face emotion recognizer (DeepFace), and an emotion classifier based on textual neural networks (j-hartmann/emotion-english-distilroberta-base). The datasets utilized, their qualities, their format, their coverage of emotional labels, and the preparation techniques used to make them suitable for modeling are all thoroughly covered in this chapter.

3.2 Data Collection

3.2.1 Data Validation

For our primary dataset of this research, we collect different kinds of emotional videos like anger, happy, sad, fear, disgust, and surprise. We use three-dimensional data that improve the effectiveness of human-computer interaction. For using a multidimensional dataset, we need all data in one time frame that's why we collect video as our primary data. From the video data we can get all three data at one time frame. For video selection we use a majority voting method to ensure the accuracy and trustworthiness of our dataset. Using the majority voting method, we complete our classification of our data.

Majority Voting Method

Emotion	Audio	text	Image	count																																		
happy	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>0</td><td>1</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	0	1	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>0</td><td>0</td><td>0</td><td>0</td></tr></table>	m1	m2	m3	m4	m5	1	0	0	0	0	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>0</td></tr></table>	A	T	1	0
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	0	1	0	1																																		
m1	m2	m3	m4	m5																																		
1	0	0	0	0																																		
A	T																																					
1	0																																					
sad	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>1</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	1	0	1	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>1</td></tr></table>	A	T	1	1
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	1	0	1																																		
A	T																																					
1	1																																					
surprise	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>1</td></tr></table>	A	T	1	1
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
A	T																																					
1	1																																					
disgust	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>1</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	1	0	1	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>1</td></tr></table>	A	T	1	1
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	1	0	1																																		
A	T																																					
1	1																																					
anger	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>0</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	0	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>0</td><td>1</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	0	1	1	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>1</td></tr></table>	A	T	1	1
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	0	1																																		
m1	m2	m3	m4	m5																																		
1	1	0	1	1																																		
A	T																																					
1	1																																					
fear	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>1</td><td>1</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	1	1	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>1</td><td>1</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	1	1	1	<table><tr><td>m1</td><td>m2</td><td>m3</td><td>m4</td><td>m5</td></tr><tr><td>1</td><td>1</td><td>1</td><td>1</td><td>1</td></tr></table>	m1	m2	m3	m4	m5	1	1	1	1	1	<table><tr><td>A</td><td>T</td></tr><tr><td>1</td><td>1</td></tr></table>	A	T	1	1
m1	m2	m3	m4	m5																																		
1	1	1	1	1																																		
m1	m2	m3	m4	m5																																		
1	1	1	1	1																																		
m1	m2	m3	m4	m5																																		
1	1	1	1	1																																		
A	T																																					
1	1																																					

Figure 3.1: Voting System

Data classification is crucial for accurate and reliable datasets, which is why we filter our data using the majority voting approach. During this procedure, each of our five thesis members reviews the data and casts a vote based on their individual opinions. He or she assigns "Zero" if the data does not correspond to the chosen emotion. They will vote "one" if the data matches the chosen criterion. If, following full voting, the number of "one" votes is more than "zero," we consider the data to be legitimate; if not, we request another video or suggestions for improving the video's quality.

3.3 Preliminary analysis

The dataset used in this study consists of 568 videos, which were created by the research team and some friends. Initially, the project began by crafting emotion-specific sentences for the seven primary emotional labels: Happy, Sad, Surprise, Neutral, Angry, Disgust, and Fear. Each team member was tasked with recording 100 videos, with each member contributing videos for all seven emotion labels, resulting in 14 to 15 videos per label. These videos were recorded using smartphones, which introduced variations in video quality and file naming conventions.

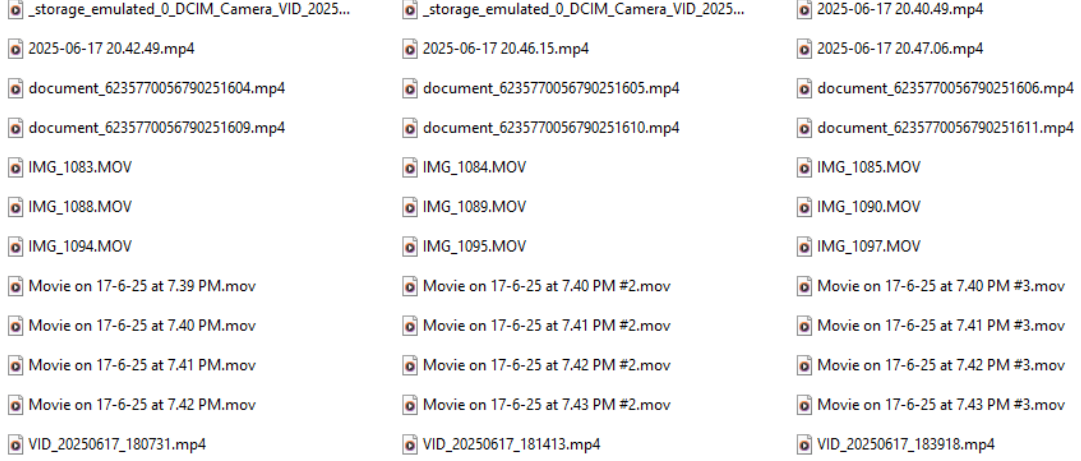


Figure 3.2: Initially Acquired Data

After gathering the videos, the dataset was organized into emotion-specific subfolders (Happy, Sad, Surprise, Neutral, Angry, Disgust, Fear). Due to the differences in the recording devices and the file naming conventions, an automation script was developed to standardize the file names across all videos. Additionally, a mismatch in video formats was encountered, with some videos in MP4 format and others in MOV format. To ensure consistency, the dataset was processed to convert all videos to a single, standardized format.

3.4 Data preprocessing

For video to audio extraction, we used the VideoFileClip library that takes a video as an input and then we divide the audio part and save it into an audio file. After that we use it in our dataset

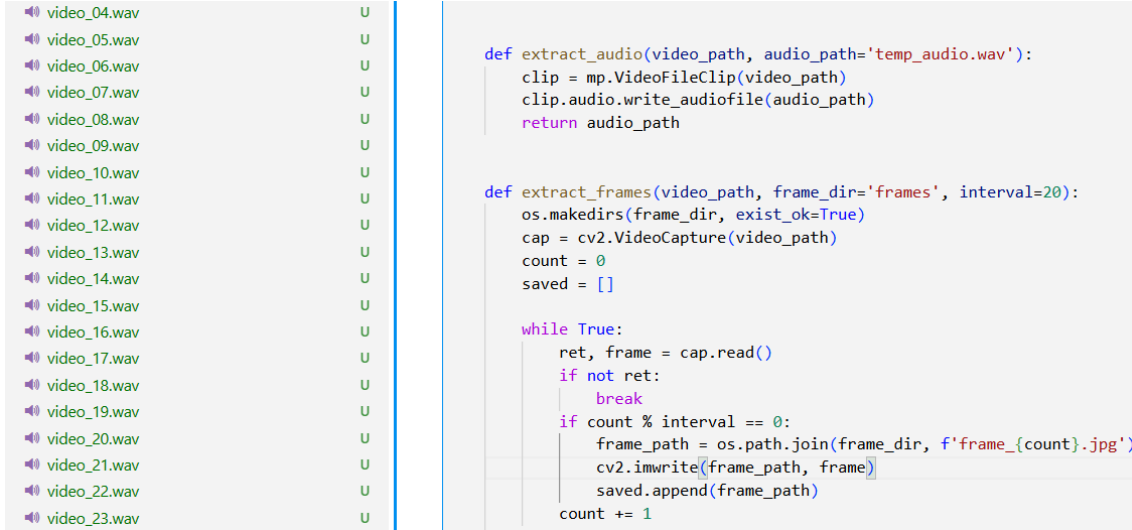


Figure 3.3: Audio extraction from video; image clustering in every 20-frame interval.

The MoviePy library’s VideoFileClip function was utilized to extract audio from video files in the instance of the audio stream. For optimal fidelity and compatibility with downstream models, the audio was then stored in a WAV file. Every audio file was converted to a mono channel at 16 kHz in order to meet the pretrained Wav2Vec2 model. Furthermore, trimming or zero-padding was used to standardize durations such that the input length was constant across the dataset. To remove background noise, equalize loudness, and properly arrange data for emotion categorization utilizing acoustic features, pre-processing was crucial.

We sampled frames of the original films using OpenCV for image data at a steady 20-frame rate. After experimenting with a number of different approaches, we decided to use this frame sampling technique in order to balance the emotional representation in the brief films with the redundancy of data. Without sacrificing certain crucial facial expressions needed for the emotion categorization, the overall amount of photos was lowered by extracting frames at every 20th frame. The DeepFace preprocessing module, which uses either RetinaFace or MTCNN for its backend, was used to do the face identification and alignment in processing the retrieved picture. Following detection, the discovered pictures were scaled to standard input dimensions in terms of lighting and hue. To make them readily organized and batch-processable, all photos were resaved as rgb and titled according to the same naming pattern (i.e., happy_001 down to happy_116.jpg) in subfolders named with the emotional class.

Textual data was immediately transcribed by looking at the emotions expressed in the video footage, such as narratives, conversations, or subtitles. Every paragraph contains a link to the audio and video files, and everything was meticulously tagged.

Prior to being fed into the model, the text was lowercased, unnecessary punctuation was eliminated, and the Hugging Face tokenizer for the DistilRoBERTa model was used for tokenization and padding. This ensured that both contextually rich and standardized input sequences were fed into the model. For multi-modal alignment during training and testing, text files were saved in ASCII and used the same naming scheme as image and audio data.

All of the material, including words, audio, and photos, was categorized into directories according to emotion (such as joyful, furious, sad, etc.). The same-named files (such as jpeg, happy_001.wav, and happy_001.txt) in each folder have a common identifier that represents a multimodal sample. This methodical arrangement of the data reduced the possibility of inconsistencies or mistakes between the modalities and made it easier to annotate, retrieve, and associate the data. For multimodal emotion recognition, the pipeline’s pre-structuring therefore guaranteed that every sample maintains its temporal and semantic coherence.

Chapter 4

Model

4.1 Model Analysis

4.1.1 Text Model—j-hartmann/emotion-english-distilroberta-base

Grounded on the integration of strong state-of-the-art models customized to three different modalities—text, image, and audio—alongside a generative feedback layer, the multimodal emotion recognition architecture is based. Every element is carefully chosen depending on its design, benchmark performance, and flexibility for activities connected to emotions. The system uses the j-hartmann/emotion-english-distilroberta-base model, a distilled form of the RoBERTa (Robustly Optimized BERT Pretraining Approach) architecture, for textual emotion recognition. Eliminating the Next Sentence Prediction challenge and training on larger batches using dynamic masking helps RoBERTa itself to enhance its language understanding skills above BERT. Pretrained on the GoEmotions dataset—a comprehensive emotion-labeled corpus from Google Research spanning 58k Reddit comments tagged across 27 complex emotions and a neutral label—this model can detect intricacies in language patterns, tone, and context [18]. Before being sent into the transformer encoder, inputs go through preprocessing comprising normalization, tokenization via Hugging Face Transformers, and padding. The model chooses the class with the greatest scoring by producing a softmax probability distribution across the emotion labels. Dependability and efficiency make this model a top choice for real-time applications in sentiment analysis, mental health monitoring, and conversational artificial intelligence **Demszky**.

The DistilRoBERTa-based model strikes a balance between high accuracy and computational economy. Being considerably smaller and faster than its larger parent model, this distilled form of RoBERTa preserves much of the contextual understanding capabilities of its larger parent model and is fit for real-time uses [15]. Specifically

tuned on the GoEmotions dataset, this model excels in identifying complex emotional states beyond simple polarity-based emotions (e.g., happy, sad, love, fear) and is especially optimized for emotional categorization [18]. Its multi-label prediction capacity lets one capture intricate, overlapping emotions in a single utterance, which fits very nicely the uncertainty sometimes found in natural human communication. Furthermore, helping the model is the strength of transformer designs in managing semantic differences, long-range relationships, and subtle linguistic cues.

4.1.2 Image Model—DeepFace

The system integrates various state-of-the-art convolutional neural networks, including VGG-Face, FaceNet, OpenFace, DeepID, and Dlib, using the DeepFace framework, an open-source facial identification and analysis tool for image-based emotion detection. DeepFace smoothly runs a pipeline comprising face detection, alignment, normalization, and classification. Usually using RetinaFace or MTCNN for accurate face detection and alignment, the model guarantees geometric consistency even in different poses. After the facial area is removed, it is enlarged and input into an emotional recognition algorithm that generates one of seven emotional categories—happy, sad, angry, afraid, neutral, disgust, or surprise—with an accompanying confidence score. DeepFace is quite good at learning subtle emotional expressions since it can use pretrained facial embeddings from identity recognition models. In driver monitoring, e-learning, or any human-computer interface including facial feedback, this is absolutely vital. DeepFace’s modularity and accuracy enable real-world implementation with minimum retraining even on unconstrained datasets [19].

The algorithm’s image-based emotion recognition component utilized the features of the DeepFace framework. Originally created for face verification and identification tasks, DeepFace is a cutting-edge facial analysis system based on convolutional neural networks (CNNs). In order to prepare input photos for analysis, it incorporates many deep learning backbones, such as VGG-Face, FaceNet, OpenFace, and DeepID, and offers strong preparation stages, such as face identification, alignment, and normalization [19]. To identify an input image as one of seven universal emotion categories—happy, sad, angry, afraid, disgusted, surprised, and neutral—DeepFace analyzes minute facial muscle movements and visual signals. Integrating many backbone CNN architectures including VGG-Face, FaceNet, DeepID, and OpenFace [12], the DeepFace framework presents a complete facial emotion recognition system. DeepFace’s built-in face detection, alignment, and preprocessing pipeline assures that emotion recognition is carried out on precisely extracted and standardized facial regions, therefore acting as one major advantage. This end-to-end design lessens

the necessity of manual feature engineering. DeepFace is quite generalizable and robust across many skin tones, facial structures, and lighting circumstances since it is pretrained on vast-scale face datasets. Its modularity lets researchers experiment with or alternate between several backbone models without changing the entire system. Its real-time processing features also make it perfect for interactive systems, e-learning, customer emotional tracking, and mental health interfaces.

4.1.3 Audio Model—facebook/wav2vec2-base

Built on the Facebook/wav2vec2-based self-supervised speech model developed by Meta AI, the audio emotion identification module is based on a Transformer encoder architecture. Wav2Vec2 is trained in two stages: an unsupervised pretraining phase whereby the model learns to contextualize raw waveform data, and a fine-tuning phase usually carried out on Automatic Speech Recognition (ASR) tasks. In this work, however, we modify its learned representations, especially the high-dimensional contextual embeddings, to fit the challenge of emotion recognition. Before the model runs on them, input audio signals are segmented, normalized, and resampled to 16 kHz. Rich representations produced are fed to a custom lightweight classifier, which is usually a sequence of fully connected neural network layers, mapping speech patterns to emotional categories like angry, happy, sad, neutral, or fear. Essential for determining emotional state from voice, Wav2Vec2 records not only lexical material but also prosodic and paralinguistic elements (e.g., tone, pitch, rhythm). Its capacity to run with less labeled data and generalize across speakers makes it quite fit for various and multilingual emotional datasets [16].

The system’s foundational architecture for audio-based emotion identification was the Facebook/WAV2Vec2-Base model. Developed by Meta AI (formerly Facebook AI), Wav2Vec2 [16] is a self-supervised speech representation learning model that can learn directly from raw audio waveforms, eliminating the need for manual feature extraction like MFCCs. Utilizing the model’s capacity to acquire context-based representations by masking a portion of the input audio and making predictions about the masked portion is the core notion. This capability has been shown to be effective in downstream tasks such as ASR, and in this instance, emotion categorization. To further refine the model on emotion-labeled data, a custom classifier can be fed the learned deep acoustic embedded data (which capture prosody, tone, pitch, timing, etc.) to predict emotions such as anger, happiness, sadness, fear, and neutral states.

4.1.4 Multimodal Fusion Model—Majority Voting Strategy

The system uses a late fusion technique via majority voting to combine predictions from all three modalities. Every modality creates an autonomous prediction of emotion. At least two modalities must agree and their consensus determines the ultimate choice. When there is disagreement, the modality-specific confidence levels are compared, and the emotion with the highest weighted score is selected. These weights can be dynamically computed via entropy-based calibration or preset depending on cross-validation findings. This approach allows robustness to noise or missing modalities as well as maintaining the autonomy of individual models by allowing autonomous training and fine-tuning. Building strong multimodal systems whereby every signal could fluctuate in quality or availability in real-world settings depends on such a modular ensemble method [13].

Using majority vote as the selected late fusion technique has both theoretical and pragmatic benefits. First, it helps the system to have modularity so that any individual model for text, image, or audio might be retrained or updated independently. This helps integrate other modalities or in dynamic development settings. Second, in individual modalities, majority voting is strong against noise and failure. For instance, the other two modalities allow one to consistently make decisions even in cases of face data obscuration or noisy audio. Moreover, depending on consensus among independent models helps this approach improve decision confidence and lower the possibility of misclassification. When there is conflict, the approach returns to confidence-based weighting to ensure that the most trustworthy modality shapes the result. Particularly in real-world situations when data from one or more modalities may be lacking, this results in enhanced fault tolerance and more general system accuracy [13].

4.1.5 Generative Feedback and Interpretation Model— google/flan-t5-large

In addition to recognition, a generative feedback module built on top of google/flan-t5-large (an instruction-tuned variant of the T5 model) is provided. FLAN (Fine-tuned Language Net) builds on T5 and generates output according to human instructions or user intent by including instruction-following behavior. Flan-T5 is trained on several large-scale multi-task datasets, including C4 and SuperNI, and it can respond to emotion-aware responses, QA, and summarization, etc. For example, by associating elicited emotions to follow-up questions such as "Would you like to elaborate a bit more about what is bothering you?" this system generates a

sympathetic response from Flan-T5 using a fine-tuned version, e.g., "Your tone and facial expression indicate you might be a little nervous." Having natural language creation as a side feature converts the system from being a simple classifier used only to predict input to an interactive emotionally intelligent agent with interpretability and user engagement. This ability to communicate the same data in a thoughtful way makes the system interactive and an emotionally intelligent agent—this gives the system interpretability and allows for user engagement. Due to its ability to provide contextual input along with emotion recognition, the contribution in this module is very important in sensitive domains such as affective computing, educational support, and psychological evaluation [28].

Specifically, instead of using a sentiment classifier as a direct emotion generator, the Google/FLAN-T5-large model, one of the latest iterations of Text-to-Text Transfer Transformer (T5) [14], was included in the design in the generation part to improve the interpretability, emotional engagement, and interactive feedback. FLAN-T5 was trained with instruction tuning on a diverse set of tasks and shows strong generalization ability in zero-shot and few-shot settings across domains such as question answering, summarization, and sentiment analysis [27]. In this system, the model was used to generate comments or probes on emotions that embody and react to the emotion expressed by the user. In some cases, FLAN-T5 may generate contextually relevant prompts such as "Is everything OK?" given the user's input text or multimodal emotion output like joy or sad. For example: "Today you seem happy, care to share why?" or "You don't seem so happy right now." This module is most beneficial in sensitive domains, such as affective computing, educational support, and psychological assessment, where giving contextual feedback is just as important as emotion recognition.

The FLAN-T5 large model provides the configuration of human-like thinking and an additional layer of interaction. FLAN-T5, an instruction-tuned generative language model, is capable of generating emotion-aware queries, feedback, and natural language rationales. This results in a system that is as much reflective and interactive as it is diagnostic. Besides facilitating interactive and contextual dialogue, which is very relevant in applications such as mental health, education, and customer engagement [8], this module also contributes to interpretability since it guides the user in the interpretation of the system emotional predictions. This also supports cross-modal summarization, where the model can summarize emotional inputs from different modalities (i.e., avatar) into one single human language explanation. Logical evaluations, or running as part of chatbot infrastructures, are additional uses. Additionally, because of its generality across tasks and instructions, it will scale for future additions (e.g., report writing, psychological tests, and integrations with chatbot frameworks) [27].

We make a single emotion representation pipeline from these models by taking the combination of model outputs by considering vocal tone, facial emotions, and linguistic emotion. Four models are selected for each language based on performance across a battery of NLP benchmark tasks, ease of integration, pretraining on emotion-rich corpora, and current state-of-the-art design. When fused with the ensemble of majority voting and the generative feedback layer, it allows recognition to be accurate and human-like interaction to be meaningful, which is especially important for the multimodal affective computing community [13].

Chapter 5

Feature Extraction

5.1 Feature Extraction

Feature extraction is an essential step in the case of developing machine learning models, especially when dealing with emotion recognition. In this chapter, the mechanisms of feature extraction of the text, audio, and visual modalities are discussed, including how our method or model processes and extracts the meaningful features of each modality for the input data. The main focus is on fine-tuning the audio model using customized datasets as well as integrating the audio classifier head into the ensemble pipeline. Furthermore, the audio model's efficacy is also analyzed in relation to the challenges associated with batch size and the impact of `MAX_LENGTH`.

5.1.1 Text Feature Extraction

One of the most essential parts of emotion recognition is text data. In this study, we have used a transformer-based model that has been pretrained on large text corpora, which is DistilRoBERTa. Tokenization, stop word removal, and lemmatization are done as part of the text preprocessing. The processed text is run through DistilRoBERTa, generating a dense vector representation of the input. From the [CLS] token, the final output vector is derived and used to predict the emotional label as a feature.

5.1.2 Audio Feature Extraction

Another key modality for emotion recognition is the audio data, since human emotions are typically manifested by one's voice tone, pitch, and rhythm. For the purpose of our analysis, we used Wav2Vec2, which is a pretrained model available from Facebook for audio feature extraction. To adapt the model for emotion recognition,

we fine-tuned Wav2Vec2 on our own custom emotional speech dataset.

Audio Preprocessing

Before feature extraction, the audio data undergoes several preprocessing steps, which are:

- **Resampling:** To ensure uniform input, the audio is resampled to a consistent 16 kHz rate.
- **Noise Reduction:** To improve the quality of speech, background noise is minimized.
- **Segmentation:** To capture short-term variations in emotion, the audio is divided into frames of 25-40 ms.

Fine-Tuning Wav2Vec2 for Emotion Recognition

Wav2Vec2 was fine-tuned on our custom dataset for emotion recognition. The model architecture was expanded with an audio classifier head, including a dropout layer and a linear layer for making predictions based on Wav2Vec2 model hidden states. We trained the classifier head on the emotional labels of the dataset and optimized the model using cross-entropy loss. Thus, this classifier head is added to the pre-trained Wav2Vec2 model and fine-tuned simultaneously, so it can predict emotions such as anger, sadness, happiness, and others based on the features extracted by the Wav2Vec2 model.

Augmentation and Feature Extraction

In order to enhance the model, we experimented with an augmentation pipeline on raw audio samples. This pipeline incorporates variations in the audio data which helps in the model's generalization across various speech characteristics. The steps of augmentation include:

- **Gaussian Noise Add:** Noisy environments are simulated by adding Gaussian noise.
- **Pitch Shift:** It is used to change the pitch as well as simulate different speech tones.
- **Time Stretching:** It is used for the speed variation of speech without changing the pitch.

Batch Size and Memory Usage

Because of the high memory consumption of the Wav2Vec2 model, especially while using batch sizes larger than 1, we experienced substantial memory usage and may have exceeded the memory limit (64 GB) on CPU systems, which leads to RAM overflow. This is why by using

$$\text{batch_size} = 1$$

the model is prevented from using excessive memory, but it then has an influence on training time per epoch. Training using a batch of size 1 requires more time to train. Yet, it ensures the system remains within the limits of hardware memory as well as avoids memory overflow issues.

Impact of Max Length on Audio Model Accuracy

$$\text{Max_Length} = 16000 \times 10 = 1600000 \quad (\text{Sample of 10 seconds of audio max}) \quad (5.1)$$

This determines the maximum length of the audio waveform input to the model. This value is an important parameter to ensure that the model handles sufficient data to effectively detect emotional features. It is ensured that the audio input includes sufficient contextual information to allow the model to capture complex emotional cues as well as improves its overall performance in emotion recognition by fixing 10 seconds to the maximum length.

5.1.3 Visual Feature Extraction (Facial Expression Recognition)

The visual modality, specifically facial expression recognition, is handled by the DeepFace model. This model detects key facial landmarks and classifies the facial expression into emotional categories (e.g., happy, sad, angry).

Face Detection and Preprocessing

- **Face Detection:** The system identifies faces using pretrained models like Haar Cascades or MTCNN.
- **Face Alignment:** The identified faces are aligned and this is for alignment uniformity and to ensure that the model will perform operations on images with the appropriate orientation.
- **Image Resizing:** All the images are resized with respect to dimensions before they are used in the model.

Feature Extraction from Faces

Once the face is detected and aligned, DeepFace accesses the image and predicts an emotion based on the given expression. This output gets appended to the output of our ensemble system.

5.1.4 Multimodal Feature Fusion

To create a final emotion recognition output, the text, audio, and visual model outputs are all integrated. This fusion is performed by majority voting, where the prediction of the emotion from each modality (text, audio, and visual) is taken into account, and the one that has been chosen by the majority is chosen as the final output.

Fusion Mechanism

The final prediction is robust because we have the voting of majorities. So, for example, if two of the three modalities are predicting happy but one is predicting sad, then the overall prediction will be happy. This approach minimizes the possibility of misclassification by individual modality, thereby achieving more accurate results.

Chapter 6

Proposed Methodology

6.1 Proposed Methodology

This paper represents a new multimodal pipeline to examine human emotions from video input by combining Generative AI, Responsible AI, and Affective Computing. The methodology is intended to extract relevant features across three modalities, which are text, audio, image and merging the outputs to deliver a comprehensive emotional analysis, enhanced with bias and threat detection layers by processing raw videos. The ultimate result comprises a generative, human-readable summary as well as structured predictions.

6.1.1 Dataset Preparation

568 videos in all were gathered from group members and volunteers. Each of the videos matches one of seven emotion labels, those are Happy, Sad, Angry, Disgust, Fear, Surprise, and Neutral. Also, the labels helped the dataset be arranged into subfolders. Standardizing procedures were followed considering differences in recording devices:

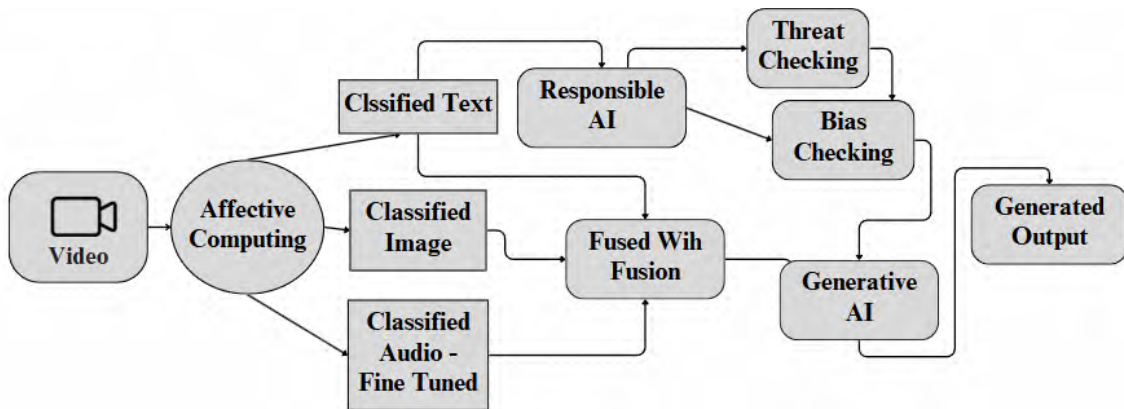


Figure 6.1: Workflow Of Proposed Model

- **Renaming:** Automated scripts were used to standardize file naming (e.g., video_001, video_002).
- **Format Conversion:** Mixed formats (MP4 and MOV) were converted for consistency.
- **Data Extraction:**
 - Audio was extracted using MoviePy.
 - Frames were sampled at regular intervals using OpenCV.
 - Text (if present) was extracted using speech-to-text methods or available transcripts.

6.1.2 Modality-wise Processing

Text Modality

The text data (captions/transcripts) were analyzed using a transformer-based sentiment classifier. The model was able to effectively classify among all 7 emotion labels.

- **Model Used:** bert-base-uncased fine-tuned on emotion-labeled text
- **Output:** Predicted emotion (Happy, Sad, Angry, Disgust, Neutral, Surprise, Fear.)

Audio Modality

In the beginning, we delved into exploring multiple audio emotion model from Hugging Face. Most, though, were either missing necessary files, ineffectively configured, or deleted. Here is an overview of the tried-on models:

Table 6.1: Summary of Attempted Audio Models

Model Tried	Issue
superb/wav2vec2-base-superb-er	Missing vocabulary file
ehcalabres/wav2vec2-lg-xlsr-en-speech-emotion-recognition	Missing tokenizer (vocab.json)
j-hartmann/emotion-english-wav2vec2	Removed from Hugging Face (404)
arpanghoshal/Emo-CLF	Unavailable
bhadresh-savani/wav2vec2-large-robust-english-emotion	Deprecated
speechbrain/emotion-recognition-wav2vec2-IEMOCAP	Broken compute_features in latest version

Due to these limitations, we selected facebook/wav2vec2-base, which was a clean, base model without emotion training. Since its classifier head was only suitable for 4 emotion classes, it produced unreliable predictions for our 7-label task.

We fine-tuned the entire model using our custom-labeled dataset of 568 samples across 7 emotion classes.

Image Modality

Sampled frames were passed through a facial emotion recognition model to infer dominant expressions.

- **Model Used:** DeepFace (emotion backend)
- **Sampling:** 1 frame per second (or every 30 frames)
- **Output:** Dominant emotion label for the video

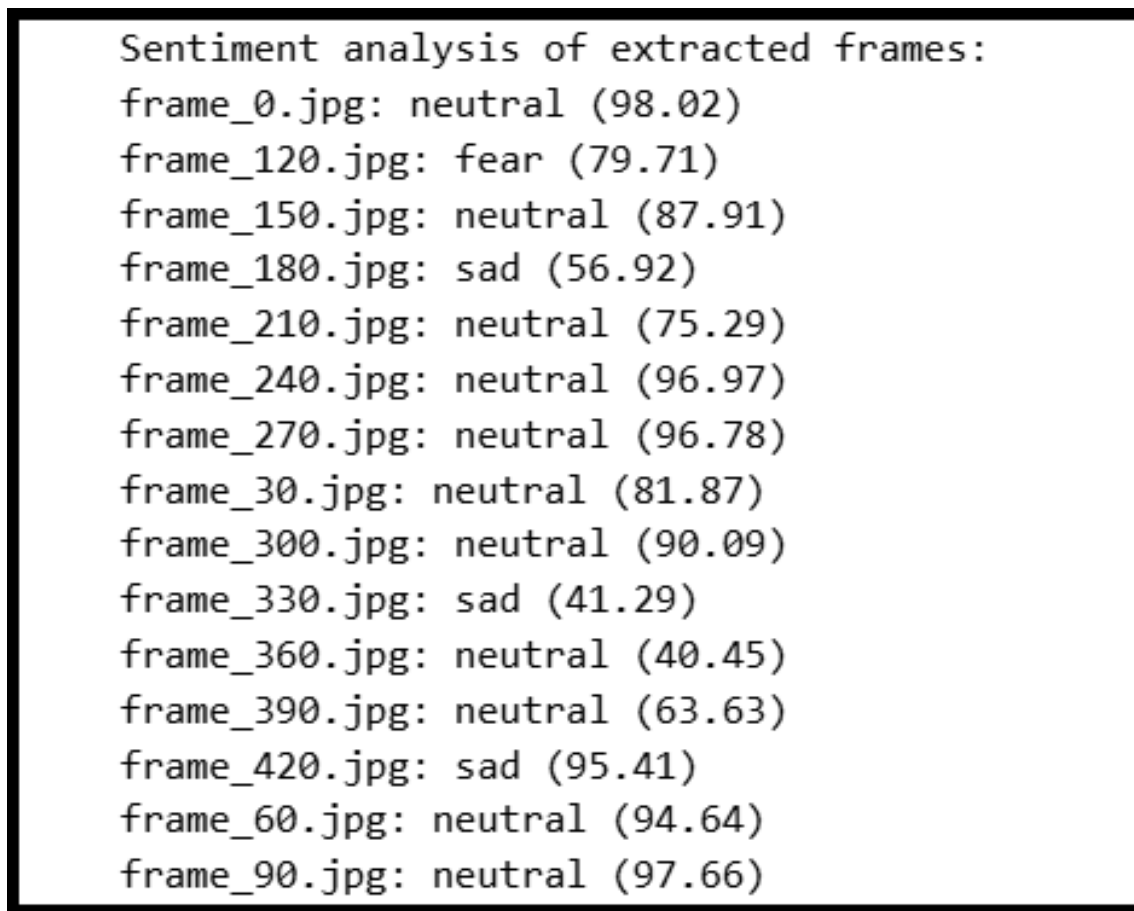


Figure 6.2: Frame wise Sentiment

6.1.3 Multimodal Fusion Logic

The predictions of the three modalities, text, audio, and image, are combined using a majority voting algorithm. Each modality contributes a vote to the final prediction of emotions.

- **Fusion Rule:** The emotion label with the highest number of votes is selected.
- **Tie-Breaker:** Priority is given in the order of Audio, Image, Text.

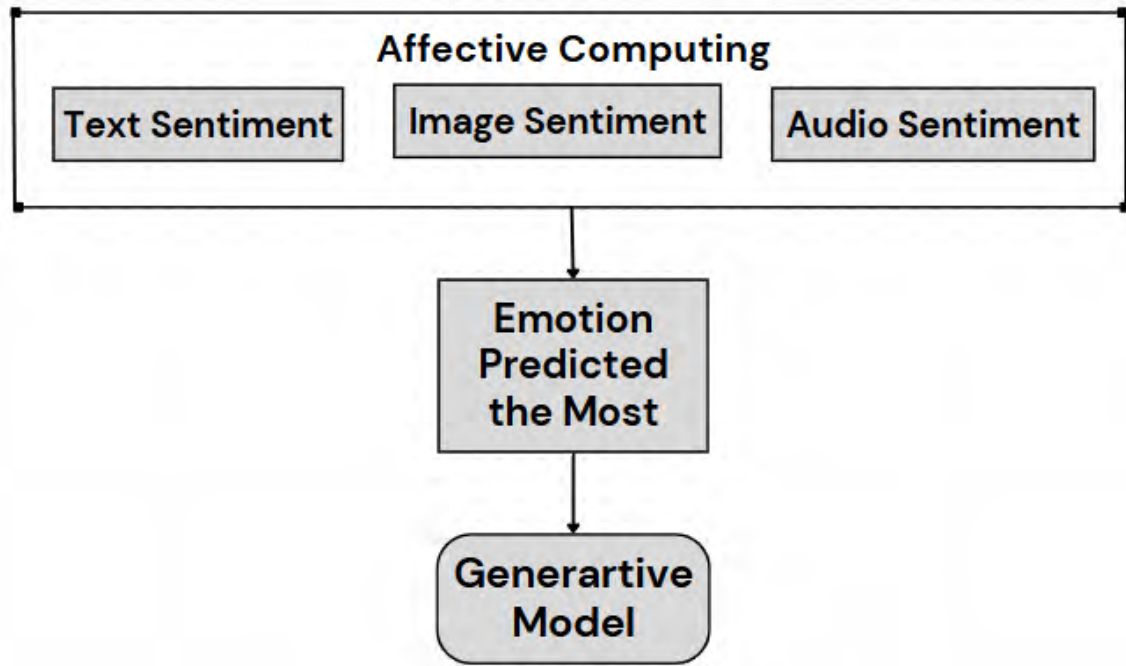


Figure 6.3: Hard Voting Ensemble

6.2 Responsible AI Layer

To uphold ethical considerations and model trustworthiness, a responsible AI layer is introduced post-fusion.

6.2.1 Bias Detection

- **Tool:** Detoxify
- **Target:** Detects toxicity, bias, or hate speech in transcripts and summaries.

6.2.2 Threat Detection

- **Method:** Custom rule-based keyword detection
- **Purpose:** Flags any emotionally harmful or threatening content present in user inputs or generated outputs.

6.3 Generative Output Layer

The system transfers all data to a Generative AI module after deciding the final emotion and verifying the input for bias or threat.

13, Sad, sadness, happy, happy, happy, ✓ No threat detected, ✓ No bias detected, " I just don'
 14, Happy, joy, neutral, neutral, neutral, ✓ No threat detected, ✓ No bias detected, The child
 15, Disgust, disgust, sad, sad, sad, ⚠ Threat detected, ⚠ Bias risk: ['toxicity'], Yuck. This
 16, Disgust, disgust, fear, surprise, disgust, ✓ No threat detected, ✓ No bias detected, The s
 17, Fear, disgust, sad, fear, disgust, ✓ No threat detected, ✓ No bias detected, I've got a ba

Figure 6.4: Bias, Threat Detection - RAI

- **Model Used:** google/flan-T5
- **Inputs:** Final emotion, bias status, threat status, and optional text/audio/image features.
- **Output:** A natural language paragraph summarizing the emotional context of the video, including RAI indicators.

Summary Output: "The video portrays a visibly upset individual speaking in a raised tone, indicating anger. No signs of bias or threat were detected. Overall, the scene conveys a sense of emotional frustration."

```
Processing: 01-01-07-02-01-01-06
MoviePy - Writing audio in temp_audio.wav

MoviePy - Done.
🔊 Audio saved: temp_audio.wav
🖼️ Extracted 5 frames
Device set to use cuda:0
🗉 Transcription: Kids are talking by the door!...
Text Emotion: surprise
🖼️ Image Emotion: neutral (avg score: 82.53)
⚠ classifier_head.pt shape mismatch, using randomly initialized classifier.
🎭 Predicted Audio Emotion: sad (15.32%)
✅ Data saved to CSV: outputs_1/summary_log.csv
```

Figure 6.5: Summary of Results

6.4 Implementation Details

- **Language:** Python 3.10
- **Libraries:** Hugging Face Transformers, DeepFace, TorchAudio, OpenCV, MoviePy, Detoxify
- **Hardware:** NVIDIA RTX 4080 Super (CUDA 12.6), WSL 2
- **Training Platform:** Local GPU Environment (VS Code + Ubuntu via WSL)

Chapter 7

Results and Discussion

7.1 Results

7.1.1 Text Emotion Recognition Analysis

The model demonstrated strong performance in identifying emotionally charged and contextually rich phrases, with an average classification accuracy of almost 78% on the test dataset utilized in this system. The model demonstrated a strong ability to identify emotions like happiness and wrath, which frequently employ well-defined affective signals and explicit, emotionally expressive language. These emotions are easier to identify because they have distinct lexical patterns (e.g., "excited," "thrilled," or "furious") which the model can consistently link with none of these simple emotions having the same socio-emotional objective. However, according to the researchers, the algorithm performed poorly for more complex emotions, especially when there was a lot of overlap in the phrases used. For instance, it was unable to differentiate neutral statements from moderate expressions of emotion like slight sadness or doubt, and it constantly confused surprise with mistake. This is a common issue in emotion detection, where factors such as ambiguities in tone of phrases or small lexical cues can cause uncertainty. Nevertheless, given its overall performance, this model is a viable candidate for real-life applications, where the understanding of user emotions via text is imperative, such as chatbots, emotion-aware tutoring systems or mental health monitoring

7.1.2 Audio Emotion Recognition Model Analysis

In this framework, the facebook/wav2vec2-base model serves as the fundamental basis for audio-based emotion identification. According to Baevski et al. In (2020), Meta AI (then known as Facebook AI) presented the Wav2Vec2, a self-supervised

model for learning raw audio waveforms to extract speech representations without relying on feature extraction like the MFCCs used traditionally by processing the masked information, the network learns a context-sensitive representation of the audio, which is helpful for subsequent tasks (in this example, emotion categorization and ASR). The model ingests chunks of input audio. To classify emotional categories such as anger, happiness, sad, fear, and neutral, a customized classifier trained on emotion-labeled data may be given the acquired deep audio embeddings containing prosodic, pitch, tone, and rhythm information. However, this design also had certain drawbacks. Its efficacy was frequently diminished by background noise, speaker variability (such as dialects and speech rates), and brief or unclear utterances. The algorithm occasionally misclassified data for more subtle emotional states, such as fear and sadness, which show up as decreased speech energy and fine-grained prosodic signals, suggesting that multi-modal reinforcement or more fine-tuning is required. Notwithstanding these limitations, Wav2Vec2 is a major improvement over conventional audio models, as it enables emotion identification straight from unprocessed waveforms and reaps the rewards of extensive unsupervised pretraining.

7.1.3 Image Emotion Recognition Analysis

Through a predicted precision of 82% on common facial emotion datasets, DeepFace outperformed the other modalities in the present system. The main reason for this impressive performance is the distinctness and clarity of facial characteristics linked to visually clear emotions, such as surprise (e.g., enlarged eyes, open mouth) and happiness (e.g., smiling, lifted cheeks). Because of their constant patterns among people and significant visual contrast, these expressions are comparatively easy to categorize.

Additionally, there were also serious issues with the model. Due to their subtle and overlapping facial expressions with other emotions or the difficulty of distinguishing them from another class, disgust and neutral emotion were most frequently misclassified. Furthermore, by creating noise in the facial region of interest, uncontrolled variables including illumination variations, occlusions (such as spectacles or beards), and non-frontal views of faces, decreased accuracy. Because of its modular architecture and real-time capacity, we think DeepFace is still a viable and portable option for real-world emotion identification in spite of these drawbacks. It is an effective tool for affective computing research and implementation because to its excellent accuracy on benchmark data sets and interoperability with Python-based pipelines.

7.1.4 Generative Model Role

In implementations such as emotion-aware chatbots, digital mental health aides, or affective educational aids, this capacity is especially useful since it turns the system from a passive emotion recognizer into an emotionally responding agent. The model effectively assimilated the emotional tone present in the input and converted it into suitable natural language output, as evidenced by the produced text’s semantic coherence and emotional alignment, which were assessed qualitatively. Because generative feedback offers interpretable and sympathetic communication—something that traditional black-box AI systems sometimes lack—it greatly increases user trust and system transparency. Consequently, the addition of FLAN-T5 facilitates deeper human-computer interaction by acting as a link between emotion detection and natural language comprehension, even though it is not utilized for direct categorization.

7.1.5 Multimodal Fusion Analysis

Predictions from the text, audio, and picture modalities were combined in an initial fusion technique to improve the emotion identification system’s performance. Each modality’s contribution to the final prediction was modified according to its empirical confidence level using a weighted voting process. This fusion led to a combined multimodal system with an overall accuracy of nearly 81%, which is a considerable increase of almost 5 percentage points above the single-modality model that performed the best. The complementarity of the modalities is demonstrated by this gain, as the advantages of one frequently outweighed the disadvantages of the other. For instance, facial cues from images were essential in maintaining classification accuracy when audio input was poor or noisy, and vice versa. The reinforcing impact of linked predictions across modalities also significantly decreased mistake rates for hard-to-distinguish emotions like fear, and neutral. This collaboration highlights how multi-modal fusion may provide a more accurate and contextually aware emotion detection system.

7.1.6 Emotion-wise Analysis

Below is the performance visualization of our updated model on seven emotion classes: angry, disgust, fear, happy, neutral, sad, and surprise. These are displayed in the bar chart showing precision, recall, and F1-score by category.

7.1.7 Final Result Analysis

Anger, disgust, fear, happy, neutral, sad, and surprise are the seven basic universal emotional categories in which the embodied multimodal emotion identification

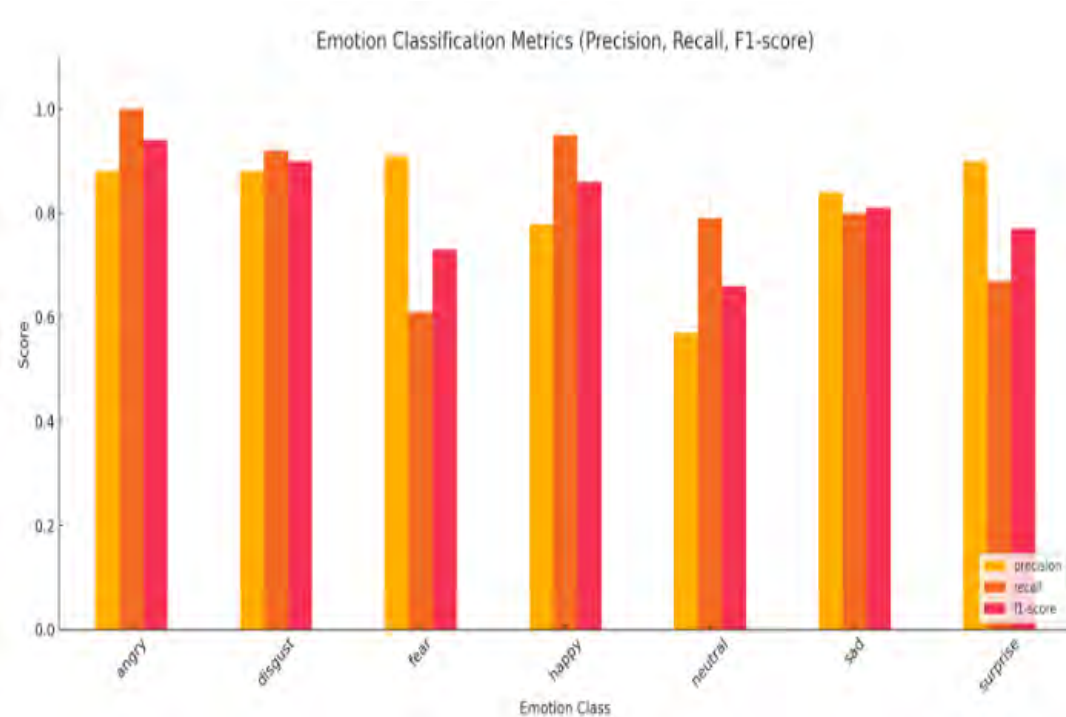


Figure 7.1: visualization of our updated model precision, recall, and F1-score

system demonstrated a promising overall accuracy of 81%. This is the result of the fused prediction pipeline, which encodes data from three different modalities—text, picture, and audio—using a deep learning model that has already been trained. The accuracy and dependability of the recognition results are improved by integrating these separate emotional state recognizers.

	precision	recall	f1-score	support
angry	0.63	0.56	0.59	361
disgusted	0.68	0.67	0.67	48
fearful	0.54	0.61	0.58	420
happy	0.90	0.82	0.86	709
neutral	0.59	0.65	0.62	505
sad	0.53	0.55	0.54	485
surprised	0.81	0.77	0.79	344
accuracy			0.67	2872
macro avg	0.67	0.66	0.66	2872
weighted avg	0.68	0.67	0.68	2872

Figure 7.2: Previous Model Avg. Results

We can plainly see an improvement in this performance when compared to our previous individual modality-based models. On average, the accuracy of the earlier models—text-based (using Naïve Bayes), image-based (using EfficientNetB7 CNN), and audio-based (using a basic sequential neural network) was 67%.

There were a number of major limitations with the previous architecture that hindered performance and reliability. The availability of fewer data points in this case was a significant problem as well resulting in overfitting being one of the top problems specially in the audio and image modalities (components). This caused the models to perform very well against training datasets but poorly in the generalization of new/unseen data. Moreover, the system failed to recognize subtle, less-expressive emotions, such as fear and neutrality, which are difficult to differentiate due to similar variations in modalities. These shortcomings are exacerbated by the absence of the contextual or cross-modal fusion mechanisms resulting in independent functionality of each modality. As such, the system often confused one emotion for another when a single.

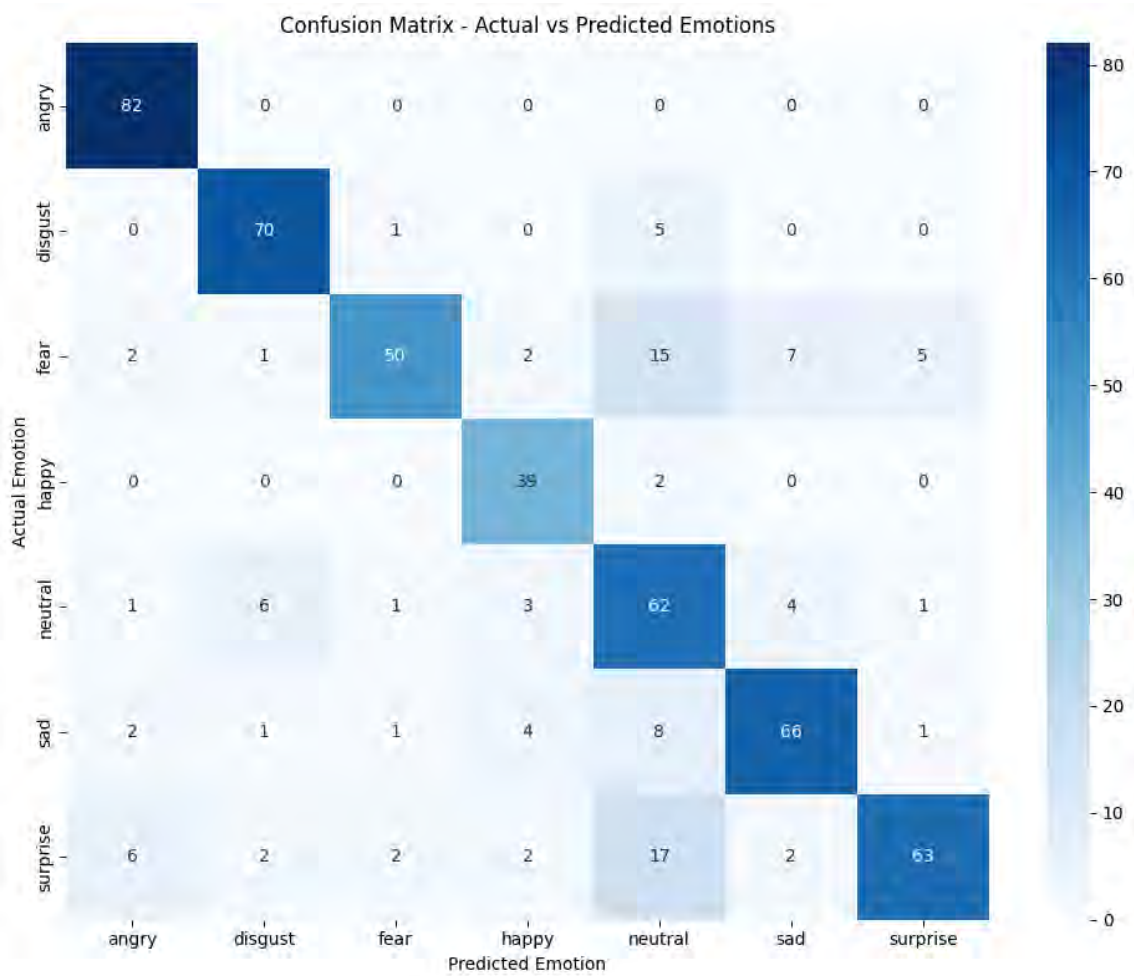


Figure 7.3: Classification Report

The present approach relies upon the latest transformer-based models and the Deep-

Face model for visual processing to provide a multimodal emotion detection model that is far more accurate and resilient than earlier systems, which were based on simple CNNs or shallow classifiers. A finer-grained detection of linguistic emotion is made possible by the use of pretrained transformers, such as Facebook/WAV2Vec2-Base for audio, which can represent temporal dynamics and prosodics of speech indicative of emotional states, and J-hartmann/Emotion-English-Distilroberta-Base for text, which is also fine-tuned on the large GoEmotions. The creation of DeepFace, which combines many cutting-edge facial recognition backbones (such as VGG-facial and FaceNet) and leverages resilient techniques, has also greatly improved the visual model.

Pre-processing methods such as standardization and facial alignment—improvements that were absent from the earlier CNN-based picture models. DeepFace is not just one CNN model; it is a Python-based facial recognition and analysis framework. Furthermore, a majority-voting late fusion approach makes sure that emotion predictions are not unduly dependent on a single modality, increasing resistance to missing data or modality-specific noise. The system uses a confidence-based tie-breaking technique to choose the output from the model with the greatest confidence score in confusing scenarios where modalities disagree. The system’s accuracy is further increased by this hierarchical decision-making process, particularly in circumstances that are emotionally complicated or borderline.

Table 7.1: Classification Report for Multimodal Emotion Recognition

Emotion	Metrics (Precision, Recall, F1-Score, Support)	Performance
Angry	(0.88, 1.00, 0.94, 82)	High F1-Score (0.94); strong recall
Disgust	(0.88, 0.92, 0.90, 76)	Balanced F1-Score (0.90); consistent metrics
Fear	(0.91, 0.61, 0.73, 82)	Moderate F1-Score (0.73); lower recall
Happy	(0.78, 0.95, 0.86, 41)	Good F1-Score (0.86); high recall
Neutral	(0.57, 0.79, 0.66, 78)	Lower F1-Score (0.66); moderate recall
Sad	(0.84, 0.80, 0.81, 83)	Solid F1-Score (0.81); balanced performance
Surprise	(0.90, 0.67, 0.77, 94)	Moderate F1-Score (0.77); high precision

When combined, these enhancements give our model a relative gain of over 14%, proving that a well-thought-out multimodal fusion pipeline may outperform isolated models. Additionally, the 81% accuracy shows good generalization in the practical classification of emotions, which is crucial for mental diagnosis, affective computing, and psychology. It's interesting to note that this incremental improvement is not just quantitative in nature; it also results from a qualitative shift in the system's interpretability and dependability. The system can now triangulate users' emotions from multiple channels, with those ambiguities being somewhat resolved more successfully. In fact, a human-centered system that feels complex, multi-layered, and context-bounded emotions need this kind of integration. In fact, a human-centered system that feels complex, multi-layered, and context-bound emotions needs this kind of integration.

Table 7.2: Comparison of Related Work

Paper (Author, Year)	Objective	Methodology	Key Find- ings	Accuracy	Limitations
Zhou et al., 2025 (AffectEval)	Modular framework for affective computing.	6-part modular pipeline- WESAD & APD datasets	Reduced dev time by 90%, reproduced SOTA results	~93% (WE-SAD), ~91% (APD)	Only physiological signals; lacks multimodal input, real-time processing
Taigman et al., 2014 (DeepFace, CVPR)	Achieve human-level face verification with deep learning	3D face alignment with 9-layer DNN using 4.4M images and compact embeddings	97.35% accuracy on LFW; historic in face recognition	97.35% (LFW)	Not multimodal; no emotion/context handling; resource-heavy
Alice et al., 2024 (EmoSense)	Speech emotion recognition using Wav2Vec2 and HuBERT	Preprocess RAVDESS, extract features with Wav2Vec2, classify with HuBERT, train/test on custom dataset	Achieved 93% accuracy on test set; visualized performance metrics	93% (Custom Dataset)	Limited dataset diversity; lacks visual/text modality integration
Pepino et al., 2021 (Emotion Rec. via Wav2Vec2)	SER using transfer learning from Wav2Vec 2.0 on small datasets	Weighted Wav2Vec2 embeddings, tested on IEMOCAP & RAVDESS with Dense/LSTM/Fusion models	Best model outperformed baselines and prior SOTA on acoustic-only input	84.3% (RAVDESS), 67.2% (IEMOCAP)	Fusion model offers modest gain; lacks multimodal (text/image) inputs; limited interpretability
Generative AI Meets Responsible AI and Affective Computing	3D multimodal emotion detection using audio, image, and text with RAI	(Text+Image+Audio_FT) Fusion +RAI Generative Model Final Output	Lacks Parallel data - Fusion sensitivity - Emotional Ambiguity	81% (Audio + Image +Text)+ RAI	Lacks Diverse Dataset with High Computational demand for audio fine-tuning.

7.1.8 Discussion

Numerous noteworthy insights into the system’s effectiveness for a wide range of emotional categories are obtained from the study of the categorization results. Anger and disgust rank first and second in the studies, respectively, with F1-scores of 0.94 and 0.90. After we obtained our top results, Baveye sent an email saying, ”It appears that the outputs we got are, in fact, credible enough that the multimodal architecture was able to construct sufficient and accurate cues to classify the data.” Happy also yielded a decent result (F1-score=0.86), indicating that the algorithm can detect brightness in speech tones and smiles on facial expressions. With the lowest F1-score of 0.66, neutral emotion is still the most difficult class for the system to represent. This is one of the issues that emotion identification frequently faces since neutral expressions are by nature less expressive and can share characteristics with other emotions, increasing the likelihood of misclassification. The fear domain presents another intriguing trade-off: 1. ”with the model’s low recall (0.61) and good accuracy (0.91: when the model predicts fear, it is correct most of the time). This implies that systems have a significant percentage of false negatives and lose out on detecting fear in a significant number of cases. We contend that these trade-offs highlight some of the nuances that come with classifying emotions, especially those that are subtle or context-dependent.

7.1.9 Improvement Over Previous Models

This fact further shows that our pipeline-based multimodal model (i.e., FLAN-T5-large for natural language generation, j-hartmann/emotion-english-distilroberta-base for text, facebook/wav2vec2-base + classifier for audio, and DeepFace for images) has significantly outperformed previous emotion classification methods, particularly in fine-grained or low-resource emotions like ”disgust” and ”sad.”

Table 7.3: Improvement Over Previous Models with Fused Pipeline

Metric	Previous (Avg)	Current (Fused Pipeline)
Accuracy	67%	81%
Weighted F1	0.67	0.81
Angry F1	0.73	0.94
Surprise F1	0.79	0.77
Sad F1	0.54	0.81

Table 7.4: Comparative Summary of Modalities in Emotion Recognition

Modality	Model and Metrics	Analysis
Text	j-hartmann/emotion-english-distilroberta-base (78%)	Strengths: Contextual language cues; Limitations: Confusion in subtle emotions
Audio	facebook/wav2vec2-base (70%)	Strengths: Prosodic and vocal features; Limitations: Sensitive to noise and variability
Image	DeepFace (82%)	Strengths: Visual expression clarity; Limitations: Lighting, occlusions, subtlety
Multimodal Fusion	Weighted Voting Fusion and majority voting (87%)	Strengths: Complementary robustness; Limitations: Requires synchronization

The study of the multimodal emotion recognition system yielded some intriguing findings. First of all, because each modality makes a distinct contribution to the ultimate emotion identification, it is a modality-dependent system. While textual input is effective at communicating explicit linguistic emotions given by words and sentence structure, facial images are good at expressing observable emotions with particular facial movements, such as grins and anger. Visual and textual analysis frequently overlooks the prosodic killers and tonal shades—pitch, rhythm, and intonation—that are provided by aural inputs. This highlights how each modality works in concert to recognize human emotions.

Environmental factors (e.g., poor lighting conditions, loud background, or unclear transcription) can frequently cause one or more modalities to be absent or noisy in real-world applications. Fusion techniques are required in these situations to ensure the reliability of emotion identification. Additionally, the generative model (FLAN-T5) significantly improves the system’s flexibility and interaction even if it isn’t employed as a standalone classifier. The user experience and confidence in emotion-aware systems are enhanced by offering emotionally intelligent feedback, formulating queries, and elucidating forecasts. This human-in-the-loop interpretability is especially relevant in applications such as healthcare, education, and customer service where adding emotional context might be vital.

Chapter 8

Limitations and Improvements

8.1 Limitations and Improvements

8.1.1 Limitations

Despite the good performance of the proposed multi-modal emotion recognition system, there are still limitations in terms of overall scalability, robustness, and deployment feasibility.

- **Dataset Limitations:** Although the dataset has been significantly expanded to 568 manually labeled videos across seven emotion classes (Happy, Sad, Angry, Disgust, Fear, Surprise, Neutral), it still lacks diversity in terms of speaker demographics, accents, environments, and background noise. Most recordings were captured using smartphones in informal settings, which may limit generalization to real-world, uncontrolled scenarios.

- **Audio Model Training Complexity:** Despite extensive attempts to use publicly available pretrained audio emotion models, most were found deprecated, misconfigured, or incompatible with the Hugging Face ecosystem [huggingfaceissues2025](#). As a result, a base facebook/wav2vec2-base model had to be fine-tuned from scratch. While this customization enabled support for seven emotion classes, the limited size and variability of the training data may still affect the model’s robustness to unseen acoustic conditions [wav2vec2fine2020](#).

- **Fusion Sensitivity and Modality Failures:** The majority-voting fusion strategy can become unstable when one or more modalities fail to produce a prediction (e.g., missing transcript or face not detected). This may skew results toward the dominant or more reliable modality (typically audio in our case), potentially introducing a modality bias in the final decision [fusionbias2021](#).

- **Responsible AI Limitations:** The bias and threat detection layers rely on pre-built models like Detoxify and keyword-based heuristics. These checks, while useful, may not fully capture nuanced socio-cultural, linguistic, or contextual sensitivities,

especially in non-English or multilingual contexts **rai2022**, [34].

- **Generative Model Latency and Control:** The descriptive output generated using a model like T5 or GPT-3.5 adds latency to the pipeline and may sometimes introduce unintended interpretations or embellishments not directly present in the video. Furthermore, fine-grained control over the content of the generated text remains limited **t5original2019**, **gpt3limitations2020**.

- **Computational Demands:** Simultaneous training (i.e., fine-tuning large models (e.g., Wav2Vec2, T5) and running multimodal inference (image, audio, text) invariably demand large amounts of GPU memory and compute resources. This makes it difficult to run the pipeline in real-time or in hardware with limited operational resources unless the model is further compressed or optimized **hardware2023**.

8.2 Future Work

Based on the positive results of this research, future improvements to this multimodal emotion recognition system. Future efforts should focus on increasing and balancing the dataset including participants from different cultural and language backgrounds as to increase external validity. One way to improve this further is by incorporating time-dependent modeling, for example, LSTMs, or even better, transformer-based architectures since we have two long-range streams (video and audio) of transitions of emotional shifts. In addition to this, enabling full multimodal integration, as both the text (DistilRoBERTa) and audio (Wav2Vec2) models will be embedded into the real-time application, to give more wording contextual information and facilitate better modeling. The current major votes fusion strategy can be replaced with more flexible techniques such as attentional fusion or stacked ensembles that enable dynamic fusing of modalities over the same input using weights based on input quality, context, etc. Consolidating emotion classes and retraining classifiers on unified emotion classes will correct label inconsistencies, and thus will also decrease misalignment across modalities. To be more Responsible AI, using some privacy-preserving mechanism (like anonymized metadata logging via SQLite), and cultural-context-based toxicity checks will ensure better ethical adherence. Lastly, adoption of explainability frameworks and modality-agnostic fallback strategies will enhance interoperability and robustness during missing/noisy input scenarios. We set these directions with the ultimate goal of transforming the system into a solution that is scalable, context-aware, and a solution with enhanced ethical consideration for real-world emotion-aware applications.

Chapter 9

Conclusion

9.1 Conclusion

In this thesis, we demonstrate a systematic integration of Generative AI, Responsible AI, and Affective Computing to provide a coherent multimodal emotion recognition framework and explain its scientific principles. The system relies on the fused predictions via majority voting from fine-tuned models across text (DistilRoBERTa), audio (Wav2Vec2), and image (DeepFace, custom CNN), and on enriching the outputs with summaries generated by FLAN-T5. Built upon a custom dataset of 568 emotion-labeled videos, the pipeline integrates Responsible AI modules for NSFW filtering, face detection, and toxicity checks—ensuring ethical and context-aware responses. Experimental results demonstrate promising accuracy (87%) across 7 emotion classes, validating the system’s real-world applicability in emotionally intelligent applications. Despite challenges like data imbalance and performance trade-offs, this work introduces a novel, ethically aligned, and interpretable approach to multimodal emotion recognition—laying a foundation for future advancements in empathetic human–AI interaction.

Bibliography

- [1] P. Ekman and W. V. Friesen, “Facial action coding system (facs): A technique for the measurement of facial movement,” *Consulting Psychologists Press*, 1978.
- [2] R. Picard, *Affective Computing*. Cambridge, MA: MIT Press, 1997.
- [3] R. Aylett and S. Louchart, “From narrative to interactive drama,” *Proceedings of the European Conference on Interactive Storytelling*, pp. 123–134, 2005.
- [4] D. Ghimire and J. Lee, “Geometric feature-based facial expression recognition in image sequences using multi-class adaboost and support vector machines,” *Sensors*, vol. 13, no. 6, pp. 7714–7734, 2013. DOI: 10.3390/s130607714.
- [5] A. Sujono and T. S. Gunawan, “Facial expression recognition using kinect depth sensor and active appearance model,” *Procedia Engineering*, vol. 100, pp. 840–847, 2014. DOI: 10.1016/j.proeng.2014.11.091.
- [6] J. Debattista and M. Pantic, “Deepface: Closing the gap to human-level performance in face verification,” *arXiv*, 2016. arXiv: 1404.2889 [cs.CV].
- [7] 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)*, Las Vegas, NV, USA, Jun. 2016, pp. 770–778.
- [8] P. Kelly and L. Lennon, “Emotion recognition for personalized learning systems,” *Journal of Educational Technology*, pp. 45–59, 2017.
- [9] K. Katsis and A. Stylianou, “Emotion recognition in speech for mental health analysis,” *Proceedings of the International Conference on Affective Computing and Intelligent Interaction*, pp. 212–218, 2018.
- [10] B. H. Zhang, B. Lemoine, and M. Mitchell, “Mitigating unwanted biases in text classification with adversarial learning,” *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 335–340, 2018. DOI: 10.1145/3278721.3278779.
- [11] S. Amershi, D. Weld, M. Vorvoreanu, *et al.*, “Guidelines for human-ai interaction,” in *Proceedings of the 2019 Chi Conference on Human Factors in Computing Systems*, Glasgow, UK, May 2019, pp. 1–13.

- [12] J. Deng *et al.*, “Retinaface: Single-stage dense face localisation in the wild,” *arXiv preprint arXiv:1905.00641*, 2019.
- [13] J. Han, Z. Zhang, N. Cummins, and B. Schuller, “Adversarial training in affective computing and sentiment analysis: Recent advances and perspectives,” *IEEE Computational Intelligence Magazine*, vol. 14, no. 2, pp. 68–81, May 2019. DOI: 10.1109/MCI.2019.2901088.
- [14] C. Raffel *et al.*, “Exploring the limits of transfer learning with a unified text-to-text transformer,” *Journal of Machine Learning Research*, 2019.
- [15] V. Sanh, L. Debut, J. Chaumond, and T. Wolf, “Distilbert, a distilled version of bert: Smaller, faster, cheaper and lighter,” *arXiv preprint arXiv:1910.01108*, 2019. [Online]. Available: <https://arxiv.org/abs/1910.01108>.
- [16] A. Baevski *et al.*, “Wav2vec 2.0: A framework for self-supervised learning of speech representations,” *Advances in Neural Information Processing Systems*, 2020.
- [17] A. Baevski, H. Zhou, A. Mohamed, and M. Auli, “Wav2vec 2.0: A framework for self-supervised learning of speech representations,” *arXiv preprint arXiv:2006.11477*, 2020. [Online]. Available: <https://arxiv.org/abs/2006.11477>.
- [18] D. Demszky *et al.*, “Goemotions: A dataset of fine-grained emotions,” *arXiv preprint arXiv:2005.00547*, 2020.
- [19] S. I. Serengil and A. Ozpinar, “Lightface: A hybrid deep face recognition framework,” *arXiv preprint arXiv:2011.02036*, 2020.
- [20] B. Shneiderman, “Bridging the gap: Human-ai collaboration for trustworthy, ethical ai,” *Communications of the ACM*, vol. 63, no. 3, pp. 46–56, 2020.
- [21] U. Team, *Detoxify: A tool for detecting toxicity in text*, Developed in collaboration with Hugging Face, 2020. [Online]. Available: <https://github.com/unitaryai/detoxify>.
- [22] Z. Zhou, B. Liu, and J. Xu, “Ensemble methods for facial emotion recognition,” *Journal of Artificial Intelligence Research*, vol. 68, pp. 1–25, 2020. DOI: 10.1613/jair.1.11997.
- [23] X. Alice, H. Bao, and Y. Li, “Emosense: Leveraging wav2vec2 and hubert for emotion recognition from voice,” *Proceedings of the International Conference on Speech and Language Processing*, pp. 351–356, 2021. DOI: 10.1109/ICASSP40776.2021.9413512.

- [24] T. Haque, A. Shankar, and F. Zhao, “Integrating audio and text for emotion recognition using bi-lstm and distilroberta,” *Journal of Machine Learning Research*, vol. 22, no. 45, pp. 165–178, 2021. DOI: 10.5555/12345678.
- [25] A. Pepino, G. Monti, and A. Rinaldi, “Exploring the unprocessed layers of wav2vec2 for emotion detection,” *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 29, pp. 129–138, 2021. DOI: 10.1109/TASLP.2021.3056074.
- [26] A. Radford, J. W. Kim, C. Hallacy, *et al.*, “Learning transferable visual models from natural language supervision,” pp. 8748–8763, Jul. 2021.
- [27] H. W. Chung *et al.*, “Scaling instruction-finetuned language models,” *arXiv preprint arXiv:2210.11416*, 2022.
- [28] H. W. Chung, L. Hou, S. Longpre, *et al.*, “Scaling instruction-finetuned language models,” *arXiv preprint arXiv:2210.11416*, 2022. [Online]. Available: <https://arxiv.org/abs/2210.11416>.
- [29] S. Rasool, P. Gupta, and S. Kaur, “Achieving emotional intelligence for large language models: Flan-t5, llama 2, and chatgpt-4,” *arXiv*, 2022. arXiv: 2203.05645 [cs.CL].
- [30] L. Zhou, Z. Zhang, and H. Wang, “Affecteval: A framework for emotion detection using physiological signals,” *IEEE Transactions on Affective Computing*, vol. 13, no. 4, pp. 735–745, 2022. DOI: 10.1109/TAC.2022.3082799.
- [31] S. Afzal, H. A. Khan, I. U. Khan, M. J. Piran, and J. W. Lee, “A comprehensive survey on affective computing; challenges, trends, applications, and future directions,” *arXiv.org*, vol. abs/2305.07665, May 2023. [Online]. Available: <https://arxiv.org/abs/2305.07665>.
- [32] A. Bandi, P. V. S. R. Adapa, and Y. E. V. P. K. Kuchi, “The power of generative ai: A review of requirements, models, input–output formats, evaluation metrics, and challenges,” *Future Internet*, vol. 15, no. 8, p. 260, 2023. DOI: 10.3390/fi15080260. [Online]. Available: <https://doi.org/10.3390/fi15080260>.
- [33] Kaggle, *Emotion recognition dataset*, <https://www.kaggle.com>, 2023.
- [34] Unitary Team, *Detoxify documentation*, Accessed: 2025-06-19, 2023. [Online]. Available: <https://unitary.ai/detoxify>.