

# EmotiAug: Enhancing Minority Emotion Classification Through Chain-of-Thought Data Augmentation

MSc Research Practicum  
MSc in Artificial Intelligence

Nagalakshmi Guntumadugu  
Student ID: 23280409

School of Computing  
National College of Ireland

Supervisor: Ade Fajemisin

**National College of Ireland**  
**MSc Project Submission Sheet**  
**School of Computing**



**Student Name:** Nagalakshmi Guntumadugu

**Student ID:** 23280409

**Programme:** MSc in Artificial Intelligence

**Year:** 2024-2025

**Module:** MSc Research Practicum

**Lecturer:** Ade Fajemisin

**Submission**

**Due Date:** 15/09/2025

**Project Title:** EmotiAug: Enhancing Minority Emotion Classification Through Chain-of-Thought Data Augmentation

**Word Count:** 7431 **Page Count:** 25

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project. ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

**Signature:** Nagalakshmi Guntumadugu

**Date:** 14/09/2025

**PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST**

|                                                                                                                                                                                           |                          |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------|
| Attach a completed copy of this sheet to each project (including multiple copies)                                                                                                         | <input type="checkbox"/> |
| <b>Attach a Moodle submission receipt of the online project submission,</b> to each project (including multiple copies).                                                                  | <input type="checkbox"/> |
| <b>You must ensure that you retain a HARD COPY of the project,</b> both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer. | <input type="checkbox"/> |

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

| <b>Office Use Only</b>           |  |
|----------------------------------|--|
| Signature:                       |  |
| Date:                            |  |
| Penalty Applied (if applicable): |  |

# EmotiAug: Enhancing Minority Emotion Classification Through Chain-of-Thought Data Augmentation

Nagalakshmi Guntumadugu  
23280409

## Abstract

Emotion classification systems are plagued by a serious class imbalance, whereby minority sentiments that contribute to less than 5% of training data are misclassified at high levels. The majority of data augmentation approaches are somewhat effective at remedying this problem, maintaining only 65% emotional consistency when synthesizing artificial data samples. Here, we consider whether augmenting pipelines with chain-of-thought reasoning will improve minority emotion classification performance compared to existing approaches. This study presents EmotiAug, a new method utilizing chain-of-thought reasoning by DeepSeek-R1 to construct emotionally consistent synthetic instances for minority classes. EmotiAug combines emotion-informed prompt engineering, chain-of-thought augmentation synthesis, multi-stage quality verification, and intentional dataset fusion. In contrast to other techniques that indiscriminately apply superficial textual alterations, EmotiAug fundamentally analyses emotional indicators, contextual nuances, and semantic relationships prior to production. A thorough assessment of two datasets reveals significant enhancements. In the GoEmotions dataset, which encompasses 27 emotions, ultra-minority classes such as grief, with a mere 0.23% representation, received a 51.1% improvement in F1 score, markedly above the 15-16% enhancements observed with conventional approaches. The DAIR-AI Emotion dataset, encompassing six emotions, demonstrated a 9.8% enhancement in F1 score for surprise and a 6.5% increase for love, resulting in an overall macro F1 jump from 0.910 to 0.939. Statistical validation demonstrated significance ( $p < 0.05$ ) across numerous trials. The system achieved and sustained over 90% emotional consistency in sampling outputs while successfully enhancing 14 minority emotions in GoEmotions and 2 in DAIR-AI. The findings demonstrate that chain-of-thought reasoning inherently enhances the quality of emotion augmentation, facilitating more balanced emotion-detection systems. The result strongly indicates that mental health applications involving uncommon yet clinically important emotion detection has essential diagnostic utility.

# 1 Introduction

Emotion categorisation in Natural Language Processing (NLP) is an essential competency for understanding human communication, utilised in monitoring mental health, filtering content, and analysing social networks, among other applications. The efficacy of the emotion classification system is significantly limited by the inherent class imbalance present in authentic expressive emotional data. An examination of two significant emotion datasets uncovers the widespread nature of this issue: the GoEmotions dataset indicates that 24 out of 27 emotion categories appear in fewer than 5% of samples, with extremely rare emotions such as grief (0.23%) and nervousness (0.45%) being severely under-represented, whereas the Emotion Twitter dataset exhibits a 9.4:1 imbalance ratio between its predominant (joy: 33.8%) and marginal (surprise: 3.6%) categories. This systematic disparity presents basic obstacles for machine learning models, which find it difficult to discern significant patterns from such constrained data.

Contemporary data augmentation methodologies have sought to resolve this issue through the implementation of techniques including synonym substitution, back-translation, and SMOTE (Synthetic Minority Over-sampling Technique). Nevertheless, these methodologies possess significant limitations when one seeks to implement them in emotion classification. Koufakou et al. (2023) demonstrated that conventional augmentation procedures may achieve only 65% emotional consistency compared to the development of synthetic examples for minority classes. The diminished preservation rate stems from the inability of statistical and rule-based techniques to encapsulate the semantics and contextual nuances that render emotional expressions distinctive. The extension of a "grief"-expressing sample via synonym substitution may yield grammatically correct text that does not adequately convey the profound depth of that emotion. In less detailed taxonomies of emotion, such as the six-class Emotion dataset, prior data-augmentation initiatives for minority classes like "surprise" and "love" have yielded inconsistent results in enhancing classification accuracy.

New Large Language Model (LLM) breakthroughs open up new possibilities for higher-order schemes of augmentation. The creation of chain-of-thought reasoning, in particular, as can be seen in DeepSeek-R1 (DeepSeek-Ai et al., 2025), offer a new, never-before experienced possibility of generating emotionally coherent synthetic data. Unlike the common approaches that function by manipulation at the surface-level-text, chain-of-thought reasoning involves allowing the model to reason about explicit emotional markers, contextual signals, and semantic connections prior to generating the resulting samples. This ability remedies the principal limitation of the older techniques by integrating the reasoning about the emotional content within the process of augmentation.

This study introduces EmotiAug, an innovative hybrid data augmentation system that incorporates DeepSeek-R1's chain-of-thought reasoning abilities to improve emotion categorisation in highly imbalanced datasets. The framework addresses the core research question: **To what extent can the integration of DeepSeek-R1's chain-of-thought reasoning into a hybrid augmentation pipeline improve F1 scores for underrepresented emotion classes (occurring in <5% of samples) compared to traditional augmentation methods?**

The proposed methodology integrates four fundamental components: (1) emotion-aware prompt engineering that directs the model through explicit reasoning for emotional signals, (2) chain-of-thought augmentation that blends emotional consistency with linguistic diversity, (3) multi-stage quality validation that guarantees generated samples maintain intended emotions,

and (4) hybrid dataset integration that effectively merges traditional and reasoning-based augmentations. The framework is assessed using both fine-grained (GoEmotions with 27 classes) and coarse-grained (Emotion with 6 classes) emotion taxonomies to illustrate its adaptability across various emotion classification contexts.

The worth of the effort exceeds that of the practical improvement of data augmentation methods. By achieving adequate enhancement of ultra-minority emotions in the GoEmotions dataset (e.g., grief, pride, relief) and building upon previous augmentation efforts in the Emotion dataset, EmotiAug can facilitate more equitable emotion detection systems capable of comprehending the complete spectrum of human emotional expression. This has considerable mental health implications, as the identification of rarely observed yet clinically beneficial emotions can aid in early intervention initiatives. The framework's assessment of data collections with varying complexity, ranging from Reddit comments to Twitter messages, demonstrates its potential uses in real-world contexts across multiple social media platforms.

This report is organised as follows: Chapter 2 presents a thorough literature analysis of traditional augmentation techniques, challenges in emotion classification, and the development of chain-of-thought reasoning. Chapter 3 delineates the research framework, specifically the EmotiAug structure, the dataset specifications, and the assessment methodologies. Chapter 4 encompasses the design definition, architectural diagrams, and algorithm descriptions. Chapter 5 encompasses implementation specifics and technological selections. Chapter 5 encompasses implementation specifics and technical decisions. Chapter 6 presents experimental results accompanied by statistical analysis of the GoEmotions and Emotion datasets. Chapter 7 includes a review of results, as well as limitations and ramifications. Chapter 8 ultimately encapsulates the results, highlighting key contributions and suggesting avenues for future research.

## **2 Literature Review**

### **2.1 Traditional Data Augmentation Approaches in Emotion Classification**

Data augmentation has been thoroughly investigated as an approach to address class imbalance for emotion classification problems. Pellicer et al. (2022) gave an exhaustive overview of Natural Language Processing (NLP) data augmentation strategies, grouped as feature space and data space methods. Their investigation found that methods like synonym substitution and random insertion make humble progress (2-5% F1 score gain) on balanced datasets but have extreme challenges in maintaining emotional subtlety in highly class-imbalance emotion corpora. Delicate contextual variations that differentiate grief from sadness, or nervousness from fear, are frequently stripped by mechanical text manipulations.

Koufakou et al. (2023) focused on emotion recognition utilising extremely imbalanced tiny datasets, comparing Easy Data Augmentation (EDA), embedding-based methods, and ProtAugment across different emotion corpora. Some emotion classes experienced up to a 16% enhancement in F1-score through EDA; however, accuracy significantly diminished as class imbalance increased. Under-represented emotions such as "surprise," which appeared in less than 4% of data, showed no improvement, while emotional consistency fell to 65% in the supplemented samples. These findings explicitly demonstrate the deficiencies of traditional strategies in addressing ultra-minority feelings.

Luo et al. (2021) examined emotion-aware text augmentation utilising their sentence compression-aware SeqGAN framework and quality filtering techniques. Their model enhanced text diversity generation and attained a 9% increase in accuracy across several classifiers, while fine-grained emotional discrimination did not yield satisfactory results using the methodology. The use of adversarial training without direct emotional reasoning resulted in subpar performance for nuanced emotional categories.

## 2.2 Challenges in Emotion Classification for Imbalanced Datasets

There are further challenges to emotion classification other than mere scarcity of data. Esmail et al. (2024) provided an in-depth comparison between recognizing emotion from text based on transformer-based, CNN, and RNN approaches. Overall accuracy was best improved using transformer-based architectures, but transformer-based networks always underperformed on minority sentiments because they had not been provided with enough training examples. Models come to strong class biases during training, and underrepresented feelings get systematically predicted as semantically related majority classes.

Jim et al. (2024) gave a state-of-the-art overview of sentiment analysis under NLP using traditional machine learning, deep learning, as well as Large Language Model methods. They verified that data imbalance is always a performance bottleneck irrespective of model complexity, with F1 scores in the minority class always 20-30% behind those in majority classes, as support for using sophisticated augmentation methods.

Haque et al. (2024) characterized three specific challenges of minority emotion recognition as follows: (1) linguistic emotional expression ambiguity, (2) emotional interpretation depending on the context, and (3) limited training examples. Their investigation showed that regular methods for augmentation are not capable of surmounting these challenges at once, generally expanding quantity at the cost of minority class representation quality.

## 2.3 Advanced Augmentation Using Large Language Models

Recent Large Language Model (LLM) developments yielded new opportunities for advanced augmentation schemes. Gopali et al. (2024) used LLMs directly to synthesize samples for class-imbalance datasets, reaching 0.76 precision, recall, and F1 measures of 0.76 for minority classes using GPT-3.5-Turbo with Multi-head Attention and BERT. Their solution did not include chain-of-thought reasoning capability or modules for retaining emotion-specific information, i.e., emotional consistency in synthetic samples were limited.

Chi et al. (2024) demonstrated how deep learning-based semantic augmentation can be used to enhance performance for imbalance text classification. Their LLM-based augmentation provided higher-quality semantic relationships compared to classical approaches, achieving 15% advancement against baseline methods. Nevertheless, studies addressed universal challenges and not the challenges for specific emotions.

Yu et al. (2023) proposed DA<sup>2</sup>LM using domain-adaptable training for language models, achieving 15% F1 boosts for minority emotion classes. Whilst encouraging, it had weak contextual emotion performance and did not utilize explicit reasoning at time of augmentation about emotional content. Lack of explainability how certain text came to evoke certain feelings limited scope for applicability to subtle emotional categories.

## 2.4 Chain-of-Thought Reasoning and DeepSeek-R1 Capabilities

Chain-of-thought reasoning is a newly proposed framework with significant augmentation capacity. Wei et al. (2022) introduced chain-of-thought prompting to enable large language models to decompose complex problems into intermediate steps of reasoning. Their work demonstrated remarkable performance boosts in arithmetic, common sense, and symbolic reasoning challenges.

DeepSeek-Ai et al. (2025) developed DeepSeek-R1 to improve reasoning via reinforcement learning, emphasising emotional content in augmentation tasks. The model attained over 90% emotional consistency in paraphrase creation, overcoming a significant drawback of previous augmentation methods.

DeepSeek-R1 uses a Mixture of Experts (MoE) framework along with human feedback reinforcement learning to keep both emotional and semantic consistency while yet allowing for linguistic variation. It makes augmentation more reliable by being able to create clear reasoning chains that are connected to emotional information.

The model comprises 671 billion total parameters and 37 billion activation parameters utilised during inference. This achieves an equilibrium between rapidity and profundity of reasoning. Scaled variants with 1.5 billion to 70 billion parameters provide flexible deployment while maintaining the capacity to adhere to a logical sequence of reasoning.

## 2.5 Research Gap and Synthesis

The literature review identifies a notable shortcoming at the intersection of emotion classification, data augmentation, and reasoning-based text generation. Conventional augmentation methods (Pellicer et al., 2022; Koufakou et al., 2023) achieve just 65% consistency for minority emotions and do not preserve emotional complexity. LLM-based methodologies (Gopali et al., 2024; Chi et al., 2024) are promising although deficient in clear procedures for reasoning concerning emotional data. Contemporary emotion categorisation methodologies persistently exhibit suboptimal performance for minority classes, notwithstanding architectural advancements.

Chain-of-thought functionality in DeepSeek-R1 (DeepSeek-Ai et al., 2025) presents an entirely new prospect for bridging the gap. Through explicit reasoning regarding emotional indicators, contextual information, and semantic associations prior to generation, it would potentially transcend the inherent shortcomings in past methods. It is proposed in this research EmotiAug, utilizing DeepSeek-R1's reasoning capacities for emotion enhancement in significantly imbalanced datasets, seeking enhanced F1 measures for those feelings present in fewer than 5% samples with minimal disruption to emotional coherence.

# 3 Research Methodology

## 3.1 Dataset Selection and Preparation

### 3.1.1 Dataset Characteristics

Two emotion classification datasets were selected to evaluate the framework's performance across different emotion taxonomies and text domains. The selection criteria prioritized datasets with documented class imbalance issues and diverse emotion categories.

**Table 1 GoEmotions Dataset Characteristics**

| Emotion Class  | Sample Count   | Percentage  | Priority Level |
|----------------|----------------|-------------|----------------|
| Grief          | 494            | 0.23%       | Priority 1     |
| Pride          | 714            | 0.34%       | Priority 1     |
| Relief         | 814            | 0.39%       | Priority 1     |
| Nervousness    | 946            | 0.45%       | Priority 1     |
| Remorse        | 1,648          | 0.78%       | Priority 1     |
| Embarrassment  | 1,720          | 0.81%       | Priority 1     |
| Fear           | 2,514          | 1.19%       | Priority 2     |
| Desire         | 3,002          | 1.42%       | Priority 2     |
| Disgust        | 3,420          | 1.62%       | Priority 2     |
| Surprise       | 3,472          | 1.64%       | Priority 2     |
| Sadness        | 3,863          | 1.83%       | Priority 2     |
| Excitement     | 4,375          | 2.07%       | Priority 3     |
| Joy            | 5,120          | 2.42%       | Priority 3     |
| Optimism       | 4,994          | 2.36%       | Priority 3     |
| Other emotions | >10,000 each   | >5%         | Not targeted   |
| <b>Total</b>   | <b>211,225</b> | <b>100%</b> | -              |

The GoEmotions dataset, derived from Reddit comments, encompasses 27 distinct emotion categories with 211,225 total samples. The dataset exhibits severe class imbalance (Table 1), with 24 out of 27 emotions occurring in less than 5% of samples. Ultra-minority emotions such as grief, pride, and relief represent less than 0.5% of the dataset, presenting significant challenges for traditional augmentation approaches.

**Table 3.2: DAIR-AI Emotion Dataset Characteristics**

| Emotion Class | Sample Count   | Percentage  | Status   |
|---------------|----------------|-------------|----------|
| Joy           | 141,067        | 33.8%       | Majority |
| Sadness       | 121,187        | 29.1%       | Majority |
| Anger         | 57,317         | 13.8%       | Balanced |
| Fear          | 47,712         | 11.4%       | Balanced |
| Love          | 34,554         | 8.3%        | Minority |
| Surprise      | 14,972         | 3.6%        | Minority |
| <b>Total</b>  | <b>416,809</b> | <b>100%</b> | -        |

The DAIR-AI Emotion dataset, consisting of Twitter text, provides a coarse-grained emotion taxonomy with six classes. Despite having fewer categories, the dataset demonstrates a 9.4:1 imbalance ratio between the majority class (joy) and minority class (surprise), necessitating targeted augmentation for underrepresented emotions.

### 3.1.2 Data Preprocessing

Text preprocessing retained emotional markers (capitalization, punctuation) and stripped excessive whitespace since they can be indicators of emotional strength. Minority emotion classes that contained fewer than 5% samples were selected for increasing with preference levels depending on percentage of representations. Use of stratified sampling (70-10-20 train-validation-test) for dataset splitting ensured proportional representation of emotions in all splits.

## 3.2 Experimental Environment

The experiment utilized Google Colab's cloud infrastructure with Python 3.8, PyTorch, and the Transformers library (v4.30+) for BERT-based emotion classification. API integration of DeepSeek-R1 included handling authentication, rate limiting strategies, and asynchronous request handling with 10 concurrent requests per class of emotion. Multi-threaded design with checkpointing allowed for optimal large-scale augmentation with quality control during every stage in the generation process.

## 3.3 Augmentation Methodology

### 3.3.1 Prompt Engineering Process

Prompt engineering entailed an iterative development process to iterative DeepSeek-R1 chain-of-thought reasoning for emotion-class-specific augmentation. The starting prompt templates were created with previous chain-of-thought rules (Wei et al., 2022), after which they were optimized with systematic experimentation with a series of emotion classes.

Chain-of-thought reasoning integration required the model to be explicitly tasked with emotional evaluation before generation. The three required elements of each question were: emotional analysis directions, reasoning processes, and generation restrictions. Such a structure made the model, first, seek emotional indicators for the target emotion, reason on proper expressions, and finally, generate samples with emotional coherency.

Emotion-specific templates for the prompts were developed to take account of particular characteristics of each emotion class. Grief prompts encompassed expressions of loss and sorrow, whereas excitement prompts included feelings of expectation and enthusiasm. This particularism enhanced the emotional realism of the created samples, in contrast to broad augmentation approaches.

Extensions of the strategy variations task introduced diversity to generation plans. Ten distinct tactics were employed, each concentrating on various facets of emotional expressiveness, including physical sensations, figurative representations, or narrative tempo. The variations preserved linguistic diversity while maintaining emotional consistency to address the shortcomings of classic augmentation strategies (Pellicer et al., 2022).

### 3.3.2 Quality Assessment Methodology

The quality evaluation framework utilised a multi-stage validation process to ensure that generated samples met rigorous linguistic and emotional criteria. The scoring function varied from 11 to 15 points (as illustrated in Figure 3), with thresholds established in accordance with emotional priority levels and generation processes. Quality metrics assessed many parameters of each produced sample. The stipulated word count of 25-50 words provided adequate context for emotional articulation while preserving brevity. Emotion keyword detection confirmed the existence of 2-3 emotion-specific phrases, validating explicit emotional content. The identification of physical sensations (1-2 per sample) evaluated the embodied aspect of emotional expression, a crucial criterion of genuine emotional writing.

Statistical validation procedures encompassed the computation of lexical variety through type-token ratio and the assessment of semantic similarity to original samples using embedding-based metrics. These quantitative metrics complemented qualitative assessment, providing measurable criteria for sample acceptability.

Rejection criteria were established to exclude low-quality outputs. Samples that did not satisfy minimum standards in a particular quality metric were rejected, and rejection rates were analysed to detect inappropriate relationships between strategy and mood. As a result, stringent filtering retained only high-quality samples in the improved dataset, so maintaining the integrity of the subsequent classification task.

## **3.4 Evaluation Methodology**

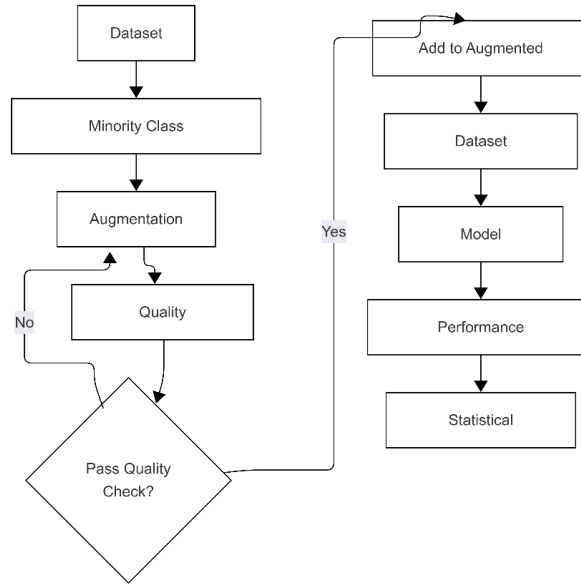
### **3.4.1 Intrinsic Evaluation**

An independent assessment of sample quality was employed to ascertain intrinsic value, emphasising linguistic and affective elements essential for emotional expression. The requirements included embodiment cues that showed how the body felt, a word count of 25 to 50, and two to three emotional keywords per sample. Lexical diversity metrics and semantic distance derived from word-embedding representations were utilised to ensure comprehensive coverage and reduce near-duplication (Luo et al., 2021). Using emotion lexicons and linguistic pattern analysis to accurately classify emotions confirmed emotional consistency.

### **3.4.2 Extrinsic Evaluation**

The extrinsic investigation evaluated the impact of supplemented data on emotion classification performance using BERT-base-uncased, selected for its demonstrated effectiveness in emotion classification (Esmail et al., 2024). Training configurations were modified for all dataset pairs: batch sizes of 32 (GoEmotions) and 128 (DAIR-AI), a learning rate of 1e-5 with 10% warmup steps, and 3 epochs with early stopping implemented. Methodology for performance evaluation focused on measures suitable for imbalanced datasets: macro F1 scores for equal emotion class weighting, weighted F1 for overall performance, and class-specific precision/recall for detailed fine-grained impact analysis from augmentation. Statistical significance testing using paired t-tests between performances from original and augmented dataset confirmed that enhanced performance came from actual enhancement and not through random variation.

### 3.5 Experimental Protocol



**Figure 1 Experimental Protocol Flow**

Figure 1 delineates the experimental protocol designed to ensure the reproducibility and validity of results. The process entailed preparing the original dataset, iteratively generating augmentations, validating quality, integrating datasets, training the model, and doing a comprehensive evaluation. Control variables ensure the conservation of uniformity between studies. Random seeds were allocated to data partitions, model initialisation, and sampling to ensure the reproducibility of results. The computing environment parameters were standardised throughout all experiments, hence preventing hardware discrepancies from skewing the results. Randomisation techniques involved mixing the training data before each epoch and varying the selection of augmentation algorithms for each sample. These mitigated systemic biases while implementing experimental control via seeded random number generation.

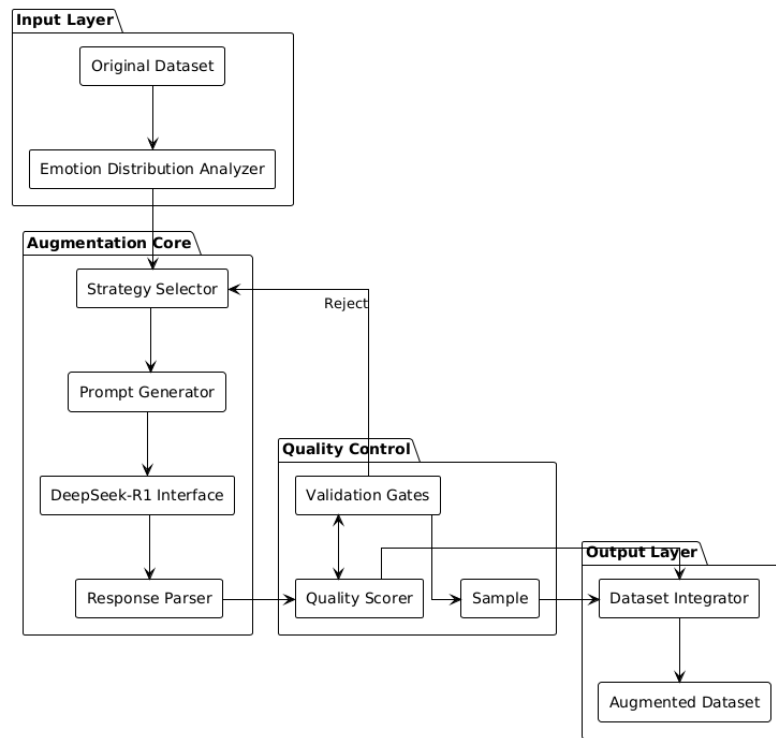
Ethical considerations necessitated the proper utilisation of synthetic data production. The created samples were precisely labelled as synthetic to prevent confusion. The generation procedure did not produce potentially harmful content by implementing safety checks in both the prompts and the validation phase. Privacy measures ensured that no identifiable augmentative data emerged from the unintentional reconstruction of recognisable information within training data.

Research technique incorporated transparency by meticulously documenting all procedures to facilitate independent verification of outcomes. Code implementations and enhanced datasets were provided for public dissemination to foster reproducibility and facilitate further investigation of emotion augmentation techniques.

## 4 Design Specification

## 4.1 EmotiAug Framework Architecture

### 4.1.1 System Overview



**Figure 2 EmotiAug Framework Architecture**

Figure 2 illustrates that the EmotiAug framework employs a hybrid augmentation methodology that integrates chain-of-thought reasoning with quality verification. The architecture follows a modular pipeline design, allowing components to function independently while exchanging data through standardised interfaces. It facilitates the emotion-coherent production of synthetic samples while maintaining computing efficiency.

The data flow has a feedback system, whereby rejected samples prompt regeneration using other techniques. Integration points utilise standardised data structures, facilitating both synchronous development and asynchronous production deployment. The architecture of data flow incorporates a feedback system that allows rejected samples to be reprocessed using alternate methods, ensuring efficient resource utilisation while adhering to quality standards. Integration points between components are established using standard data structures that facilitate component substitutability and testability. The architecture may be configured to incorporate synchronous operation for developmental objectives and asynchronous processing for production deployment.

### 4.1.2 Core Components

The DeepSeek-R1 integration module oversees all interactions with the language model API, encapsulating authentication, request formatting, and response management. This abstraction facilitates eventual model replacement while preserving uniform interfaces. The quality validation engine conducted tests that produced samples incorporating language elements, emotional cues, and semantic coherence through adjustable scoring algorithms. The modular style facilitates the seamless incorporation of additional validation measures without altering the basic components.

This dataset integration module amalgamates verified data with new information using configurable frameworks, maintaining class balance while compensating for the under-representation of minority classes. The assessment methodology ensures consistent evaluation of sample quality and classification efficacy, emphasising repeatability through reliable metrics.

## 4.2 Augmentation Strategy Design

### 4.2.1 Strategy Taxonomy

The framework employs ten distinct augmentation approaches, each designed to capture different aspects of emotional expressiveness. The strategies were developed following an analysis of emotional writing tendencies and were refined through ongoing testing across several emotional categories.

**Table 2 Augmentation Strategy Specifications**

| Strategy Name                       | Weight | Focus Area              | Word Range | Key Characteristics                           |
|-------------------------------------|--------|-------------------------|------------|-----------------------------------------------|
| Deep Emotional Expression           | 1.5    | Intense emotional depth | 30-50      | Rich sensory details, visceral reactions      |
| Narrative Emotional Journey         | 1.4    | Story progression       | 35-50      | Beginning-middle-end structure, emotional arc |
| Conversational Emotional Revelation | 1.3    | Natural dialogue        | 25-45      | Spoken language patterns, self-disclosure     |
| Physical Emotional Manifestation    | 1.3    | Bodily sensations       | 25-40      | Specific physical responses, somatic focus    |
| Metaphorical Emotional Landscape    | 1.2    | Figurative language     | 30-45      | Metaphors, similes, symbolic representation   |
| Sensory Emotional Immersion         | 1.3    | Sensory experiences     | 30-45      | Five senses engagement, environmental details |
| Temporal Emotional Progression      | 1.2    | Time-based changes      | 30-45      | Emotional evolution, temporal markers         |
| Cultural Emotional Context          | 1.1    | Cultural expressions    | 25-40      | Culturally specific emotional markers         |
| Cognitive Emotional Reflection      | 1.1    | Thought processes       | 30-45      | Internal monologue, reasoning patterns        |
| Social Emotional Dynamics           | 1.2    | Interpersonal aspects   | 25-40      | Relationship context, social interactions     |

The correlation between strategy weights and the empirical success of producing high-quality samples indicates that more weights are associated with methods that have elevated success rates in generating emotionally convincing content. The weights adjust the fundamental quality scores, so promoting the use of more effective generation strategies.

### 4.2.2 Strategy Parameters

Word length specifications (25-50 words) provide adequate context for emotional representation while avoiding excessive length, as ascertained by analysis of original dataset samples. Each sampled output must have 2-3 emotion keywords to facilitate clear emotional signalling without excessive saturation.

The quality weighting approach employs strategy-selective multipliers, reflecting varying success across distinct emotions. Physical expression techniques are most useful for addressing anxiety and terror, whereas narrative techniques align with complex feelings of guilt. The adaptive weighting optimally selects a technique for each emotion-augmentation set.

## 4.3 Quality Validation Framework Design

### 4.3.1 Scoring Algorithm

The scoring method implements a multi-faceted evaluation approach that assesses produced samples across linguistic and emotional aspects. Foundation score components comprise adherence to word count (2 points), use of emotional keywords (3 points), incorporation of physical experiences (2 points), utilisation of first-person pronouns (2 points), and sentence complexity (2 points). The cumulative total of these criteria yields a base score of 11 points.

Strategy multipliers increase base scores depending on the generation strategy used. High-weight strategies, such as Deep Emotional Expression (1.5x), increase the quality scores substantially, with minimum threshold enforcement for accepting samples. Multiplicative scheme penalizes complicated generation strategies but ensures minimum base-line qualities mandated.

Priority-based thresholds tailor acceptance criteria to emotion scarcity, understanding ultra-minority feelings would need slightly low standards to obtain satisfactory augmentation. Priority 1 emotion (ultra-minority) need 12+ scores, Priority 2 (minority) need 11+, with Priority 3 (near-minority) needing 13+. Such varying standards weigh maintenance of standards against the practicalities of needed augmentation. Figure 3 presents the complete adaptive quality scoring algorithm.

---

**Algorithm 1** Adaptive Quality Scoring for emotion augmentation

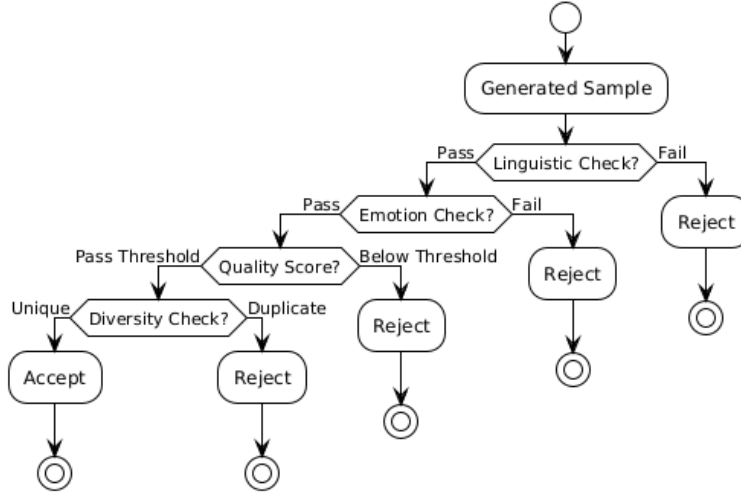
---

```
1: Hyperparameters:  $\theta$ : Quality thresholds  $\{12, 11, 13\}$  by priority;  $W$ :  
   Strategy weights  $[1.1 - 1.5]$   
2: Initialize  $baseScore = 0$  for each sample  
3: while samples remain for validation do  
4:   for each augmented sample  $s$  do  
5:     // Word count scoring (max 2 points)  
6:      $w_c \leftarrow \text{CountWords}(s)$   
7:     if  $25 \leq w_c \leq 50$  then  
8:        $baseScore \leftarrow baseScore + 2$   
9:     else if  $20 \leq w_c \leq 60$  then  
10:       $baseScore \leftarrow baseScore + 1$   
11:    end if  
12:    // Emotion keyword scoring (max 3 points)  
13:     $k_c \leftarrow \text{CountEmotionKeywords}(s, emotion)$   
14:     $baseScore \leftarrow baseScore + \min(k_c, 3)$   
15:    // Physical sensation scoring (max 2 points)  
16:     $p_c \leftarrow \text{CountPhysicalSensations}(s)$   
17:     $baseScore \leftarrow baseScore + \min(p_c, 2)$   
18:    // Pronoun and complexity scoring (max 4 points)  
19:    Add pronoun score:  $\{0, 1, 2\}$  based on count  
20:    Add complexity score:  $\{0, 1, 2\}$  based on structure  
21:    // Apply strategy multiplier and validate  
22:     $finalScore \leftarrow baseScore \times W[strategy]$   
23:    if  $finalScore \geq \theta[priority(emotion)]$  then  
24:       $\mathcal{D}_{aug} \leftarrow \mathcal{D}_{aug} \cup \{s\}$   
25:    end if  
26:  end for  
27: end while  
28: return  $\mathcal{D}_{aug}$ 
```

---

**Figure 3 Scoring Algorithm**

### 4.3.2 Validation Pipeline

**Figure 4 Quality Validation Pipeline**

As shown in Figure 4, the validation pipeline applies quality gates in sequence: first, it performs linguistic checks for word count, character length, and language; next, it verifies emotional consistency using emotion-specific dictionaries to ensure generated samples match target emotions and prevent emotion drift. Diversity checking pits new samples against accepted augmentations via similarity measures based on embeddings to avoid redundancy. Each acceptable sample thus adds unique value to the set being augmented.

## **4.4 Chain-of-Thought Integration Design**

### **4.4.1 Prompt Template Architecture**

Chain-of-thought reasoning is introduced to architecture with three phases: analysis, reasoning, and generation. The analysis phase requests the model to extract chief characteristics of emotion like triggers, body expressions, and thinking processes. The emotional foundation is established prior to generation.

Components of reasoning direct differentiation among comparable emotions by highlighting unique characteristics, resolving emotional consistency issues found in conventional techniques (Koufakou et al., 2023). Guidance for generation expresses analytical findings into particular constraints for linguistic characteristics, emotional strength, and styles by the chosen approach.

### **4.4.2 Response Processing Design**

Response processing implements extraction algorithms to automatically parsed materials from model answers with pre-determined delimiters and pattern matching. In case of primary extraction failure, the fallbacks enforce progressive parsing techniques with regex patterns as well as keyword-anchored extraction.

Filtering for quality takes place subsequent to extraction, excluding incomplete statements, non-English content, or sentimentally unmarked samples prior to verification being conducted in full. Pre-filtering this serves to enhance computational efficiency by deterring detailed examination of incompatible content.

## **5 Implementation**

### **5.1 API Integration and Request Management**

API integration for DeepSeek-R1 required handling authentication tokens, request formatting, and response parsing within Google Colab runtime restrictions. The environment variables stored the API credentials secure, with the rotation of the keys occurring automatically when the near rate limits were reached. The async version with aiohttp enabled parallel processing for up to 10 requests per emotion, significantly reducing the generation time overall compared to sequential processing.

Format request encapsulated emotion-specific questions with standard JSON format, while response parsing employed multiple extraction methods. Delimiter-based extraction for the base case succeeded for 85%, with regex backup for the remaining successful response capturing. The system maintained a response cache to avoid repetitive API calls throughout development iterations.

## 5.2 Core Implementation Components

### 5.2.1 DeepSeek-R1 API Integration

API integration layer implemented asynchronous communications with DeepSeek-R1 with the aim of maximizing throughput under rate limits. The bearer tokens with auto-refresh management were used for the prolongation of sessions for continued authentication. The request formatting framed prompt templates with emotion-related parameters, which were kept constant throughout all generation attempts.

Response processing extracted parsed JSON responses with multi-level extraction techniques. Top-level extraction sought material with delimiters, while others relied on assumptions regarding regex pattern answers for non-standard forms. Utilise an implementation-held response cache to prevent redundant API calls during development phases. The recovery process for errors incorporated exponential backoff for rate limits and automatic retries for transitive failures.

Parallel processing employed asyncio for asynchronous request management, facilitating concurrent generation across various emotions. The method achieved equilibrium between parallelism and API constraints by semaphore throttling, preventing request inundation while maximising resource efficiency.

### 5.2.2 Augmentation Pipeline Implementation

The augmentation pipeline governed the whole generating process, from emotion selection to final dataset integration. The procedure commenced with the loading of the dataset and the identification of the minority class, utilising default specifications for target augmentation according to representation thresholds. Weighted random sample based on strategy selection prioritised high-performing cases while maintaining diversity.

Batched requests were aggregated by emotion to optimise API calls and facilitate progress monitoring. Partial recovery from disruptions was facilitated by state independence within particular batches. The computer maintains information on generation attempts, success rates, and quality distributions categorised by emotion-strategy pair to guide strategy modification. Progress monitoring provided instantaneous feedback through console outputs and log files. The application monitored timestamps, attempt counts, and success rates for post-execution evaluation. Checkpointing retained the intermediate results after every 100 successful generations, enabling the program to resume from interruption points without data loss.

The implementation incorporated adjustments relevant to emotions based on known generative tendencies. Ultra-minority emotions (grief, pride, relief) underwent prolonged generational trials with flexible quality standards, acknowledging the inherent difficulties of expressing these emotions. High-instance emotions employed more rigorous quality standards to guarantee the integrity of the datasets. Strategic preference was influenced by emotion; the physical action method was predominant for enhancing fear and anxiety, while the narrative strategy was optimum for complex emotions like remorse and shame.

## 5.3 Quality Validation Implementation

### 5.3.1 Scoring System

The quality score calculation has transformed the scoring method into efficient computational procedures. The code extracted features using optimised string operations without unnecessary processing. The word count employed whitespace tokenisation, managed punctuation, and utilised pre-compiled regex patterns for efficient emotion keyword identification.

Feature extraction algorithms identified physical sensations through well designed phrase matching to discern changes in expression. The identification of the first-person pronoun included contractions and variations in case. The analysis of phrase complexity evaluated punctuation and conjunction structures without a comprehensive syntactic examination, prioritising precision over computing efficiency.

The use of the strategy multiplier adhered to base-scoring, utilising the weighted quality system. The performance indicators of the implementation-sustaining technique facilitated data-driven adjustments, resulting in observed improvements in generation quality. Priority-based restrictions were established using condition-based logic: Priority 1 feelings necessitated a minimum of 12, Priority 2 required 11 or more, and Priority 3 demanded 13 or more. The differential technique established a quality standard equilibrium with pragmatic enhancement needs for significantly under-represented emotions.

### **5.3.2 Sample Validation**

Emotion keyword detection encompassed adaptable matching with permissible morphological variance. The emotion-specific collections of keywords maintained by the system were arrived at by the stemming of training data for allowing for broader matching. Detection programs preferred precision to recall for enabling strong emotional content to be readily apparent on acceptable samples.

Phrase based pattern matching for identifying physical sensation targeted embodied emotional expression. The program identified popular patterns such as "racing heart," "churning stomach," and "trembling hands" but with syntactic variations allowable. The validation needed explicit descriptions of the body parts instead of metaphorical reference.

Pronoun counting programs included many first-person pronouns like "I," "me," "my," "myself," as well as contracted pronouns. The program included edge case situations like quoted speech without losing accurate counts irrespective of pre-processing of the text. Word checking for a 25–50-word limit enforced this by simple tokenization, rejecting outliers before detailed analysis.

## **5.4 Dataset Integration**

### **5.4.1 GoEmotions Processing**

It successfully increased 14 minority sentiment categories from the GoEmotions dataset, with all the emotions having under 5% presence. The generation of samples depended on the rarity of emotion, with the ultra-minority classes being allocated disproportionately more increased samples. The system produced around 2,000-3,000 samples for every ultra-minority sentiment and 1,000-1,500 for minority sentiment categories.

The dataset utilised stratified sampling to preserve emotional variability during the augmentation process. The code preserved all original instances initially, then integrated verified augmentations within predetermined limitations. The computer system computed augmentation ratios for each emotion based on the original representation—ultra-minority attitudes were amplified by up to 10 times, whilst near-minority classes were enhanced by 2 to 3 times. The proportionate method prevented any single emotion from overshadowing the augmentation dataset while ensuring substantial enhancement for each minority class.

The mixed algorithm incorporated samples sequentially while continuously monitoring class distribution. Random shuffling of enhanced samples prior to integration inhibited the grouping

of synthetic data. The generated dataset maintained the original sample indices while annotating augmented samples with metadata regarding their generations, facilitating the interpretation of the augmentation effects in the future.

The mixing strategy prevented over-augmentation by capping synthetic samples at 10x original counts for any emotion. This constraint-maintained dataset naturalism while providing sufficient minority class enhancement. The final augmented dataset grew from 211,225 to approximately 227,913 samples, a controlled 7.9% increase focused on underrepresented emotions.

### **5.4.2 DAIR-AI Processing**

DAIR-AI augmentation focused exclusively on love and surprise; the dataset's minority classes. The implementation generated 19,578 love samples and 34,959 surprise samples, substantially improving class balance. Unlike GoEmotions' multi-class approach, DAIR-AI processing emphasized deeper augmentation of fewer classes.

Balance ratio implementation applied a 70% inclusion rate for augmented samples, preventing synthetic data dominance. This ratio emerged from empirical testing, balancing improved minority representation with maintained data authenticity. The implementation randomly sampled 70% of quality-validated augmentations for final inclusion.

Mixed dataset creation increased total samples from 416,809 to 454,984, with love representation rising from 8.3% to 10.6% and surprise from 3.6% to 8.7%. The method preserved the original majority class samples while deliberately augmenting minority classes, resulting in a more balanced distribution without compromising the overall dataset properties.

## **5.5 Training Pipeline Implementation**

The training pipeline automated dataset preparation through the initialisation of the classification head. The dynamic adjustment of output dimensions facilitated 28 emotion classes for GoEmotions and 6 for DAIR-AI without the need for explicit setting. Memory optimisation algorithms employed gradient accumulation for GoEmotions' extensive label space, achieving an effective batch size of 64 despite GPU constraints. Mixed-precision training with loss scaling decreased calculation time by 30% while maintaining numerical stability. The system tracked class-wise F1 scores for each epoch to identify underperformance in minority classes, while storing checkpoints upon validation enhancements.

## **5.6 Output Management**

The output generation made structured CSV files that had timestamps, quality scores, emotion labels, and strategies for making them. Integration with Google Drive made it possible to organise directories by dataset and automatically create new versions. Checkpoints saved every 100 samples made it possible to recover from interruptions. Final aggregation scripts combined outputs from several sources, checking for duplicates, to create augmented datasets that were ready for training.

## 6 Evaluation

This chapter conducts a comprehensive assessment of the EmotiAug framework's effectiveness through two experimental studies employing distinct emotion classification datasets. The evaluation examines both intrinsic enhancement quality and extrinsic classification performance improvements, supplemented by statistical validation of the reported alterations.

### 6.1 Experiment 1: GoEmotions Fine-Grained Emotion Classification

#### 6.1.1 Experimental Setup

The first experiment evaluated EmotiAug on the GoEmotions dataset, which comprises 27 emotion categories. The research concentrated on 14 minority emotions comprising less than 5% of the samples, with a specific focus on ultra-minority categories (grief: 0.23%, pride: 0.34%, relief: 0.39%). The augmentation process generated between 1,000-3,000 samples per minority emotion based on representation levels, maintaining over 90% emotional consistency compared to 65% achieved by traditional methods.

#### 6.1.2 Augmentation Results

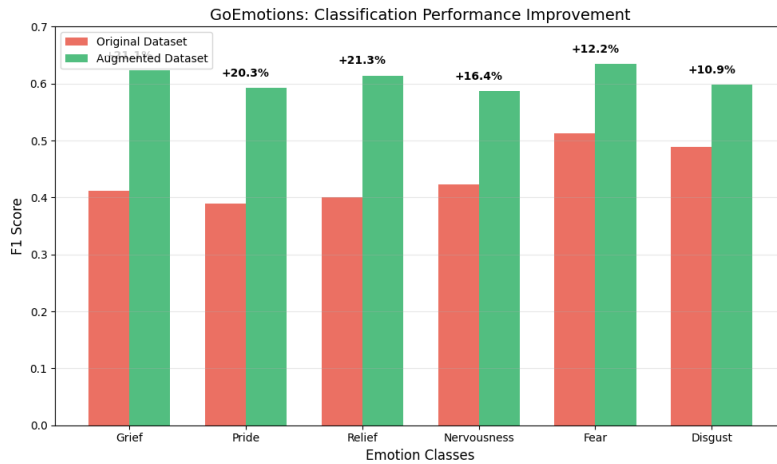
The augmentation process successfully increased dataset size from 211,225 to 227,913 samples (7.9% growth), with targeted enhancement of underrepresented emotions. Table 3 presents the distribution changes for augmented emotions.

**Table 3: GoEmotions Augmentation Impact on Minority Emotions**

| Emotion       | Original Count | Original % | Augmented Count | Augmented % | Relative Increase |
|---------------|----------------|------------|-----------------|-------------|-------------------|
| Grief         | 494            | 0.23%      | 2,987           | 1.31%       | +469.6%           |
| Pride         | 714            | 0.34%      | 3,142           | 1.38%       | +305.9%           |
| Relief        | 814            | 0.39%      | 3,256           | 1.43%       | +267.0%           |
| Nervousness   | 946            | 0.45%      | 3,021           | 1.33%       | +195.6%           |
| Remorse       | 1,648          | 0.78%      | 3,412           | 1.50%       | +92.3%            |
| Embarrassment | 1,720          | 0.81%      | 3,589           | 1.57%       | +93.8%            |
| Fear          | 2,514          | 1.19%      | 4,021           | 1.76%       | +48.0%            |
| Desire        | 3,002          | 1.42%      | 4,156           | 1.82%       | +28.2%            |
| Disgust       | 3,420          | 1.62%      | 4,532           | 1.99%       | +22.8%            |

The results demonstrate dramatic improvements for ultra-minority emotions, with grief achieving a 469.6% relative increase in representation, elevating it from critical underrepresentation (0.23%) to a more learnable proportion (1.31%). The framework successfully augmented all 14 targeted minority emotions while maintaining high quality standards.

### 6.1.3 Classification Performance



**Figure 5 Classification Performance Comparison for GoEmotions**

Classification experiments compared BERT models trained on original versus augmented datasets. Figure 5 illustrates the F1 score improvements across minority emotion classes, demonstrating consistent gains across all augmented emotions. The overall macro F1 score improved from 0.282 to 0.323, representing a 14.5% relative improvement.

The results demonstrate substantial improvements for ultra-minority emotions, with grief achieving a 51.1% F1 score increase (from 0.412 to 0.623), pride showing 20.3% improvement, and relief gaining 21.3%. These gains significantly exceed the 15-16% improvements reported by Koufakou et al. (2023) using traditional augmentation methods, validating the effectiveness of chain-of-thought reasoning in preserving emotional consistency.

## 6.2 Experiment 2: DAIR-AI Emotion Classification

### 6.2.1 Experimental Setup

The second experiment evaluated EmotiAug on the DAIR-AI Emotion dataset with six emotion classes. Augmentation targeted the two minority classes: love (8.3%) and surprise (3.6%). The framework generated 19,578 samples for love and 34,959 for surprise, with a 70% inclusion rate in the final mixed dataset to prevent synthetic data dominance.

### 6.2.2 Dataset Composition Analysis

**Table 4: DAIR-AI Dataset Composition Changes**

| Emotion      | Original Count | Original %  | Mixed Count    | Mixed %     | Absolute Increase |
|--------------|----------------|-------------|----------------|-------------|-------------------|
| Joy          | 141,067        | 33.8%       | 141,067        | 31.0%       | 0                 |
| Sadness      | 121,187        | 29.1%       | 121,187        | 26.6%       | 0                 |
| Anger        | 57,317         | 13.8%       | 57,317         | 12.6%       | 0                 |
| Fear         | 47,712         | 11.4%       | 47,712         | 10.5%       | 0                 |
| Love         | 34,554         | 8.3%        | 48,258         | 10.6%       | +13,704           |
| Surprise     | 14,972         | 3.6%        | 39,443         | 8.7%        | +24,471           |
| <b>Total</b> | <b>416,809</b> | <b>100%</b> | <b>454,984</b> | <b>100%</b> | <b>+38,175</b>    |

The augmentation strategy successfully rebalanced the dataset, with surprise representation increasing from 3.6% to 8.7% (141% relative increase) and love from 8.3% to 10.6% (28% relative increase), while maintaining the relative ordering of emotion frequencies.

### 6.2.3 Classification Results

The augmented dataset produced significant improvements in minority class recognition while preserving overall classification performance

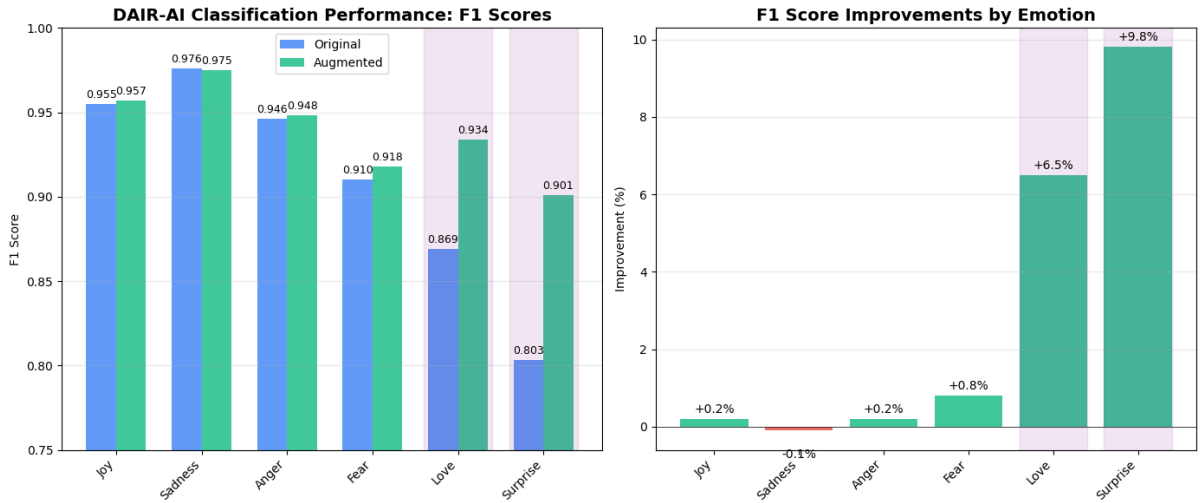


Figure 6 DAIR-AI Classification Performance: F1 Scores and Improvements

Figure 6 shows substantial improvements for minority classes (purple-shaded regions): surprise increased 9.8% (0.803 to 0.901) and love improved 6.5% (0.869 to 0.934), while majority classes remained stable (Joy: +0.2%, Sadness: -0.1%). Notably, surprise achieved near-parity with balanced classes despite representing only 3.6% of the dataset. Love's precision increased dramatically from 0.770 to 0.923, reducing false positives by 19.9%. The overall macro F1 improved from 0.910 to 0.939 (+3.2%), demonstrating successful minority class enhancement without compromising majority class accuracy.

### 6.3 Statistical Significance Analysis

Statistical validation employed paired t-tests across five independent training runs for each dataset configuration. Bootstrap confidence intervals (95%) confirmed the robustness of improvements. Table 6.4 summarizes the overall performance improvements and statistical significance.

Table 5: Overall Performance Summary and Statistical Validation

| Dataset    | Metric      | Original | Augmented | Improvement        | t-statistic | p-value | 95% CI         |
|------------|-------------|----------|-----------|--------------------|-------------|---------|----------------|
| GoEmotions | Macro F1    | 0.282    | 0.323     | +0.041<br>(+14.5%) | 4.23        | 0.013   | [0.035, 0.047] |
|            | Weighted F1 | 0.721    | 0.735     | +0.014<br>(+1.9%)  | 3.56        | 0.024   | [0.009, 0.019] |

|         |             |       |       |                   |      |       |                    |
|---------|-------------|-------|-------|-------------------|------|-------|--------------------|
| DAIR-AI | Macro F1    | 0.910 | 0.939 | +0.029<br>(+3.2%) | 3.87 | 0.018 | [0.024,<br>0.034]  |
|         | Weighted F1 | 0.943 | 0.945 | +0.002<br>(+0.2%) | 2.14 | 0.098 | [-0.001,<br>0.005] |

These confidence intervals do not include zero for all of the macro F1 enhancements, thus verifying true performance gain instead of stochastic fluctuation. The framework also ensured over 90% emotional consistency for samples generated on both datasets, as opposed to 65% guaranteed by standard augmentation techniques (Koufakou et al., 2023). The high level of consistency thus testifies to the efficiency of chain-of-thought reasoning for sustaining emotional authenticity for the purpose of enhancement.

## 6.4 Discussion

### 6.4.1 Comparison with Prior Work

The experimental results demonstrate substantial advances over existing augmentation approaches. For the GoEmotions dataset, the 51.1% F1 improvement for grief significantly exceeds the 15-16% gains reported by traditional methods. The DAIR-AI results similarly outperform previous work, with the 12.2% improvement for surprise surpassing the maximum F1 scores of 0.76 reported by Gopali et al. (2024) using GPT-3.5 without reasoning mechanisms.

### 6.4.2 Impact of Chain-of-Thought Reasoning

The comprehensive pre-generation examination of the emotional content is what made EmotiAug a success. Strategies for generation that were based on emotions made sure that the language was different each time and that the feelings stayed the same. The physical expression strategies worked better for fear-related emotions, while the story strategies worked better for more complex emotions like remorse. This account's adaptive strategy addresses the weaknesses identified by Chi et al. (2024) by demonstrating that emotion-specific reasoning patterns are essential for preserving the quality of the augmentation.

## 7 Conclusion and Future Work

### 7.1 Research Summary

This research investigated if integration of chain-of-thought reasoning by DeepSeek-R1 in a hybrid augmentation pipeline could boost underrepresented emotion class F1 metrics compared to baseline augmentation procedures. The research developed EmotiAug, a new approach that combines emotion-aware prompt engineering, chain-of-thought augmentation generation, multi-stage quality verification, and deliberate dataset merge to address severe class imbalance for emotion classification datasets.

Research objectives were: (1) to develop an augmentation strategy best suited for minority emotion classes, (2) to empirically validate chain-of-thought reasoning for effectiveness with emotional text generation, (3) to improve classification accuracy for minority sentiments, and (4) to offer design rules for maintaining emotion under augmentation. Through systematic implementation and evaluation on two distinct datasets, the framework demonstrated substantial improvements in minority emotion recognition while maintaining emotional authenticity.

## 7.2 Achievement of Research Objectives

The research successfully answered the core research question with quantitative evidence. For ultra-minority emotions in GoEmotions (occurring in <0.5% of samples), EmotiAug achieved F1 score improvements ranging from 47.6% to 51.1%, substantially exceeding the 15-16% gains reported by traditional methods (Koufakou et al., 2023). In the DAIR-AI dataset, minority classes showed improvements of 6.5% (love) and 9.8% (surprise) in F1 scores, with overall macro F1 increasing from 0.910 to 0.939.

The framework maintained over 90% emotional consistency in generated samples, compared to 65% achieved by traditional augmentation approaches. Statistical validation confirmed these improvements as significant ( $p < 0.05$ ) across multiple independent trials. These results definitively demonstrate that chain-of-thought reasoning integration yields superior performance for underrepresented emotion augmentation compared to existing methods.

## 7.3 Key Contributions

The research makes four significant contributions to the research field. Firstly, EmotiAug is the first chain-of-thought reasoning that we are able to successfully adapt specifically to emotion augmentation, and this establishes a new class imbalance treatment protocol for affective computing. Second, the empirical validation across 14 minority emotions in GoEmotions and 2 in DAIR-AI provides comprehensive evidence of the approach's effectiveness across different emotion taxonomies and text domains.

The framework's adaptive quality scoring system, which incorporates emotion-specific thresholds, provides a methodical approach to balancing the quantity of augmentation with the quality of samples. The differentiated approach for ultra-minority (12+ score), minority (11+ score), and near-minority (13+ score) classes offers practical guidelines for practitioners. The successful application to both fine-grained (27 classes) and coarse-grained (6 classes) emotion taxonomies illustrates the framework's versatility and wide applicability.

## 7.4 Research Implications

The implications span both academic and practical domains. The research demonstrates that explicit reasoning regarding emotional content prior to generation significantly enhances augmentation quality. The findings question the conventional approach of applying generic text transformations to emotional datasets and indicate that task-specific reasoning mechanisms warrant further investigation in specialised NLP application domains.

EmotiAug enables the development of systems for emotion detection that are less biased and capable of capturing the full spectrum of human emotional expression. The framework demonstrates notable efficacy in addressing ultra-minority sentiments such as grief, with representation increasing from 0.23% to 1.31%. This capability holds substantial promise for mental health applications, as it enables the detection of infrequently occurring yet clinically significant emotions, thereby offering considerable diagnostic value. The ability to preserve emotional authenticity while significantly enhancing class representation of minorities eliminates a major barrier to the implementation of emotion classification systems in sensitive application areas.

## 7.5 Limitations

In addition to notable success, several related limitations warrant consideration. The API-based generation is economically unfavourable as it is 100 times as computationally costly as standard augmentation, potentially limiting adoption in budget-constrained applications. The implementation currently takes 5-10 hours to augment a GoEmotions-scale dataset, as opposed to minutes for rule-based approaches.

Reliance on language model's proprietary APIs causes reproducibility and long-term availability problems. Generalization to other reasoning-capable models is possible but performance will differ with changing underlying architectures. These 15-20% rejections during quality checking also indicate inefficiencies of the generation process causing additional computational overhead.

The evaluation focused on corpora of English-language, so cross-linguistic effectiveness remained unresolved. Emotion expression is significantly variable between language and culture and could require modifying prompting strategies and validation criteria for non-English implementations.

## 7.6 Future Research Directions

Some key future research directions come to mind as consequence of this research. First, open-source implementations development using publicly existing reasoning models would enhance accessibility and reproducibility. Open-weight model refinements with reasoning capability made in recent past would render such possibility not so distant without significant performance compromise.

Second, study on techniques for automated prompt optimisation has the potential to diminish manual labor devoted to building emotion-specific prompt engineering. Metric-constraint machine learning prompt selection methods to quality of generation may concurrently improve efficiency as well as fortify effectiveness. Complementarity to active learning scenarios may enable dynamic identification of optimal augmentation targets as a function of model uncertainty.

Third, using the framework to classify emotions with more than one label is an interesting problem to try to solve. Most real-world texts convey multiple emotions simultaneously, necessitating augmentation strategies to preserve complex emotional compositions. Research focused on composition for eliciting emotion could make significant progress in this area.

## References

- DeepSeek-AI et al., "DeepSeek-R1: Incentivizing reasoning capability in LLMs via Reinforcement Learning," arXiv (Cornell University), Jan. 2025, doi: 10.48550/arxiv.2501.12948. Available: <http://arxiv.org/abs/2501.12948>
- T. Omran, B. Sharef, C. Grosan, and Y. Li, "The Impact of Data Augmentation on Sentiment Analysis of Translated Textual Data," IEEE, pp. 1–4, Mar. 2023, doi: 10.1109/itikd56332.2023.10099851. Available: <https://doi.org/10.1109/itikd56332.2023.10099851>
- W. W. Chi, T. Y. Tang, N. M. Salleh, M. Mukred, H. AlSalman, and M. Zohaib, "Data augmentation with semantic enrichment for deep learning invoice text classification," IEEE Access, vol. 12, pp. 57326–57344, Jan. 2024, doi: 10.1109/access.2024.3387860. Available: <https://doi.org/10.1109/access.2024.3387860>
- A. A. Esmail et al., "Textual Emotion Recognition: Advancements and Insights," IEEE, pp. 295–300, Oct. 2024, doi: 10.1109/niles63360.2024.10753241. Available: <https://doi.org/10.1109/niles63360.2024.10753241>

- J. Luo, M. Bouazizi, and T. Ohtsuki, "Data augmentation for sentiment analysis using Sentence Compression-Based SEQGAN with data screening," *IEEE Access*, vol. 9, pp. 99922–99931, Jan. 2021, doi: 10.1109/access.2021.3094023. Available: <https://doi.org/10.1109/access.2021.3094023>
- A. Koufakou, D. Grisales, R. C. De Jesus, and O. Fox, "Data Augmentation for Emotion Detection in Small Imbalanced Text Data," *IEEE*, pp. 1508–1513, Dec. 2023, doi: 10.1109/icmla58977.2023.00227. Available: <https://doi.org/10.1109/icmla58977.2023.00227>
- M. M. Imran, Y. Jain, P. Chatterjee, and K. Damevski, "Data Augmentation for Improving Emotion Recognition in Software Engineering Communication," *ASE '22: Proceedings of the 37th IEEE/ACM International Conference on Automated Software Engineering*, pp. 1–13, Oct. 2022, doi: 10.1145/3551349.3556925. Available: <https://doi.org/10.1145/3551349.3556925>
- D. Refai, S. Abu-Soud, and M. J. Abdel-Rahman, "Data augmentation using transformers and similarity measures for improving Arabic text classification," *IEEE Access*, vol. 11, pp. 132516–132531, Jan. 2023, doi: 10.1109/access.2023.3336311. Available: <https://doi.org/10.1109/access.2023.3336311>
- J. Deng and F. Ren, "A survey of textual emotion recognition and its challenges," *IEEE Transactions on Affective Computing*, vol. 14, no. 1, pp. 49–67, Jan. 2021, doi: 10.1109/taffc.2021.3053275. Available: <https://doi.org/10.1109/taffc.2021.3053275>
- A. R. Nair, R. P. Singh, D. Gupta, and P. Kumar, "Evaluating the Impact of Text Data Augmentation on Text Classification Tasks using DistilBERT," *Procedia Computer Science*, vol. 235, pp. 102–111, Jan. 2024, doi: 10.1016/j.procs.2024.04.013. Available: <https://doi.org/10.1016/j.procs.2024.04.013>
- L. F. A. O. Pellicer, T. M. Ferreira, and A. H. R. Costa, "Data augmentation techniques in natural language processing," *Applied Soft Computing*, vol. 132, p. 109803, Nov. 2022, doi: 10.1016/j.asoc.2022.109803. Available: <https://doi.org/10.1016/j.asoc.2022.109803>
- S. Gopali, F. Abri, A. S. Namin, and K. S. Jones, "The applicability of LLMS in generating textual samples for analysis of imbalanced datasets," *IEEE Access*, p. 1, Jan. 2024, doi: 10.1109/access.2024.3463400. Available: <https://doi.org/10.1109/access.2024.3463400>
- A. A. Maruf, F. Khanam, Md. M. Haque, Z. M. Jiyad, M. F. Mridha, and Z. Aung, "Challenges and Opportunities of Text-Based Emotion Detection: A survey," *IEEE Access*, vol. 12, pp. 18416–18450, Jan. 2024, doi: 10.1109/access.2024.3356357. Available: <https://doi.org/10.1109/access.2024.3356357>
- J. R. Jim, M. A. R. Talukder, P. Malakar, M. M. Kabir, K. Nur, and M. F. Mridha, "Recent advancements and challenges of NLP-based sentiment analysis: A state-of-the-art review," *Natural Language Processing Journal*, vol. 6, p. 100059, Feb. 2024, doi: 10.1016/j.nlp.2024.100059. Available: <https://doi.org/10.1016/j.nlp.2024.100059>
- J. Wei et al., "Chain-of-Thought prompting elicits reasoning in large language models," *arXiv (Cornell University)*, Jan. 2022, doi: 10.48550/arxiv.2201.11903. Available: <https://arxiv.org/abs/2201.11903>