



## DETECTION OF EMOTIONAL ABUSE IN DIGITAL TEXT USING NATURAL LANGUAGE PROCESSING

Adewuyi Joseph Oluwaseyi  
Computer Science Department,  
School of Computing Babcock University  
Ilishan-Remo, Nigeria

Agi Vine Ozioma  
Computer Science Department,  
School of Computing, Babcock University  
Ilishan-Remo, Nigeria

Koiti Oloolade Joseph  
Computer Science Department,  
School of Computing, Babcock University  
Ilishan-Remo, Nigeria

**Abstract:** Emotional abuse in digital communication is difficult to detect because it relies on manipulation, gaslighting, invalidation, and coercive control rather than explicit profanity or hate speech. These patterns are psychologically harmful but linguistically subtle, and they pass through conventional content moderation systems without being flagged. This paper presents the development and evaluation of a Natural Language Processing (NLP) system for detecting emotionally abusive language in short-form digital text. Three classification models were compared: Multinomial Naive Bayes, Linear Support Vector Machine (SVM), and a fine-tuned RoBERTa transformer. The models were trained on a combined dataset of 85,487 labeled samples from the Davidson et al. [2] Hate Speech and Offensive Language dataset and the Jigsaw Toxic Comment Classification dataset. The SVM achieved the highest accuracy at 95.41% with a precision of 0.9558 and an F1-score of 0.8991. RoBERTa achieved 92.61% accuracy and an F1-score of 0.8852, while Naive Bayes recorded 92.30% and an F1-score of 0.8792. A qualitative error analysis identified systematic model failures on sarcasm, gaslighting, and coercive framing. The system was deployed as a real-time web application using Streamlit and FastAPI.

**Keywords:** Emotional Abuse Detection, Natural Language Processing, RoBERTa, Support Vector Machine, TF-IDF, Text Classification, Transformer Models.

### I. INTRODUCTION

Digital communication platforms have grown rapidly in recent years. Statista [1] revealed that over 4.9 billion people use social media and messaging platforms globally. While these platforms have made communication easier, they have also created new channels that can inflict psychological harm on their users. Emotional abuse, which is a pattern of behaviour designed to control or degrade someone through language, has increased in common digital spaces where victims find it difficult to seek help and automated systems fail to detect the problem.

Emotional abuse is different from explicit hate speech. It relies on coded strategies including gaslighting, guilt-tripping, invalidation, and coercive framing. These strategies rarely or barely contain the offensive keywords or profanity that trigger existing moderation filters. Lees et al. [15] showed that emotionally abusive messages frequently receive low toxicity scores from automated systems because they lack surface-level offensive markers. Adebayo and Ogunleye [16] found significant emotional abuse on WhatsApp and Instagram among Nigerian university students, with most students reporting little or no confidence in automated reporting tools.

The research gap is well documented. Zhang and Chen [13] reviewed 147 peer-reviewed NLP papers published between 2018 and 2023 and found that emotional abuse was severely overlooked, with most projects focusing on explicit

toxicity or hate speech. This paper fills this gap by building and testing an NLP framework for identifying emotional abuse in short digital texts. The main contributions of this work are: (1) a comparative evaluation of three classification models on an 85,487-sample combined dataset; (2) A deep-dive qualitative error analysis that identifies specific linguistic blind spots in automated models regarding implicit abuse. and (3) a deployed real-time web interface accessible to non-technical users.

### II. BACKGROUND AND RELATED WORK

#### A. Traditional Approaches to Abusive Language Detection

Early work on abusive language detection relied on rule-based pattern matching. Spertus [5] built Smokey, one of the first systems for recognizing hostile messages in online forums, using syntactic pattern rules. Razavi et al. [6] combined keyword lexicons with Naive Bayes classifiers for forum comment classification. These systems were foundational but failed on context-dependent or implicit abuse because they depended on predefined offensive terms.

Machine learning approaches expanded the field considerably. Warner and Hirschberg [7] applied support vector machines with n-gram features to hate speech detection. Davidson et al. [2] produced one of the most cited datasets in the field by distinguishing hate speech from general offensive language using logistic regression on approximately 25,000 annotated tweets. Their work revealed that context and community norms complicate classification

substantially. Nobata et al. [17] demonstrated that combining lexical, syntactic, and semantic features improves detection across multiple abuse types.

### B. Transformer-Based Approaches

The field underwent many changes because of the introduction of the transformer models. Devlin et al. [3] introduced BERT, which enabled contextual word representations through bidirectional self-attention mechanisms. Liu et al. [4] proposed RoBERTa as an optimised variant with stronger pretraining, and subsequent

research confirmed that fine-tuned transformer models significantly outperform traditional classifiers on toxicity detection tasks. This study uses RoBERTa as its primary model based on these findings.

### C. Implicit and Emotional Abuse Detection

Recent research has moved towards implicit and emotionally manipulative language. Table I summarises six closely related studies published between 2021 and 2024 that are relevant to the present work.

TABLE I. Review of Related Works in Emotional Abuse Detection (2021 to 2024)

| Author / Year           | Problem                                          | Methodology                                         | Results                                          | Limitations                           |
|-------------------------|--------------------------------------------------|-----------------------------------------------------|--------------------------------------------------|---------------------------------------|
| Vidgen et al. [8]       | Detecting implicit hate speech                   | Adversarial Human-in-the-loop dataset generation    | Improved robustness on implicit hate benchmarks  | Does not address emotional abuse      |
| Rajamanickam et al. [9] | Improving recall for emotionally abusive content | Multi-task transformer with emotion classification  | Improved recall for emotionally harmful language | Requires simultaneous emotion labels  |
| Giorgi et al. [10]      | Detecting gaslighting in digital text            | Psychological framework integrated into transformer | Better gaslighting detection over baselines      | Limited to gaslighting only           |
| Pradhan et al. [11]     | Modeling abuse across conversation turns         | Temporal emotion modeling with transformer          | Detects love-bombing and devaluation cycles      | Requires full conversation history    |
| Kumar et al. [12]       | Relational abuse across multiple messages        | Graph neural networks on conversation data          | Captures coercive patterns at conversation level | High compute; not for single messages |
| Zhang and Chen [13]     | Research gaps in emotional abuse NLP             | Systematic review of 147 papers (2018-2023)         | Emotional abuse severely underrepresented        | Review only; no model proposed        |

Source: Compiled by the authors (2025)

The reviewed literature demonstrates and confirms that transformer models outperform traditional classifiers on implicit and context-dependent content. However, none of the studies specifically target emotional abuse in the Nigerian digital communication context, nor address code-switched multilingual communication; gaps the present study partially addresses through dataset design and deployment context.

## III. METHODOLOGY

### A. Dataset

The dataset was compiled from two publicly available benchmark sources. The Davidson et al. [2] Hate Speech and Offensive Language dataset contains approximately 25,000 annotated tweets categorised into hate speech, offensive language, and neither. For this study, hate speech and offensive language were merged into a single abusive label and neither was treated as non-abusive. The Jigsaw Toxic Comment Classification dataset contains approximately 160,000 Wikipedia talk page comments labeled with multiple toxicity categories. Any comment with at least one active toxicity label was assigned an abusive label, and all others were assigned a non-abusive label.

The combined raw dataset contained 102,163 samples. After preprocessing, the final dataset consisted of 85,487 samples comprising 57,970 non-abusive messages (67.8%) and 27,516 abusive messages (32.2%). The class imbalance was addressed during model training through balanced class weight configuration in the SVM and oversampling in the RoBERTa pipeline.

### B. Data Preprocessing

Four key preprocessing steps were carried out. First, duplicate messages were removed to prevent the model from

learning redundant patterns. Second, messages shorter than 10 characters were filtered out because they typically lack sufficient content for classification. Third, messages exceeding 500 characters were removed to maintain consistency and reduce computational overhead. Fourth, samples with unclear or ambiguous labels were dropped, leaving only clean binary labels. The cleaned dataset was split into three stratified subsets: 70% for training, 15% for validation, and 15% for testing. All three models were evaluated using the same test to make a fair comparison.

### C. Feature Extraction

Two feature extraction strategies were used. For Naive Bayes and SVM, TF-IDF vectorisation was applied to cleaned and lowercased text. The vectoriser was configured with a maximum vocabulary of 10,000 features and an n-gram range of (1, 2), capturing both unigrams and bigrams. For RoBERTa, the tokeniser from the Hugging Face Transformers library was applied directly to the original raw text without prior cleaning. This was a deliberate design decision because RoBERTa was pretrained on raw English text and its byte-pair encoding tokeniser treats punctuation, casing, and special characters as meaningful signals.

### D. Model Implementation

Three classification models were implemented. The Multinomial Naive Bayes classifier was trained on TF-IDF feature vectors using Scikit-learn. The Linear SVM was implemented using LinearSVC with balanced class weights. The RoBERTa model [4] was fine-tuned using the Hugging Face Trainer API with a PyTorch backend, loaded from the roberta-base checkpoint containing approximately 125 million parameters. Table II presents the hyperparameters used during fine-tuning.

**TABLE II. RoBERTa Fine-Tuning Hyperparameters**

| Hyperparameter   | Value | Rationale                                                      |
|------------------|-------|----------------------------------------------------------------|
| Learning Rate    | 2e-5  | Standard for RoBERTa fine-tuning; preserves pretrained weights |
| Train Batch Size | 16    | Balanced for Google Colab GPU memory constraints               |
| Eval Batch Size  | 32    | Safe for inference; reduces evaluation time                    |
| Epochs           | 3     | Sufficient convergence; no overfitting observed                |
| Weight Decay     | 0.01  | L2 regularisation to reduce overfitting                        |
| Optimiser        | AdamW | Default Hugging Face Trainer optimizer                         |
| Precision        | FP16  | Reduces GPU memory; accelerates training                       |

Source: Compiled by the authors (2025)

Training was executed using the AdamW optimiser with FP16 mixed precision. The best-performing checkpoint was restored automatically using the `load_best_model_at_end` flag. The final model was uploaded to Hugging Face Hub and served through a FastAPI backend connected to a Streamlit web interface for real-time classification.

### E. Evaluation Metrics

Performance was measured using accuracy, precision, recall, and F1-score on the test dataset. Accuracy measures the proportion of correctly classified samples. Precision measures the proportion of predicted abusive messages that

were genuinely abusive. Recall reflects the model's ability to identify true abusive messages. The F1-score is the harmonic mean of precision and recall, it provides balanced measure by comping both precision and recall. Confusion matrices were also generated to examine true positives, true negatives, false positives, and false negatives for across all models.

## IV. RESULTS

### A. Quantitative Performance

Table III presents the classification performance of all three models evaluated on the held-out test set.

**TABLE III. Model Performance Comparison on Held-Out Test Set**

| Model       | Accuracy        | Precision | Recall | F1-Score |
|-------------|-----------------|-----------|--------|----------|
| Naive Bayes | 0.9230 (92.30%) | 0.8880    | 0.8706 | 0.8792   |
| SVM         | 0.9541 (95.41%) | 0.9558    | 0.8991 | 0.8991   |
| RoBERTa     | 0.9261 (92.61%) | 0.8852    | 0.8852 | 0.8852   |

Source: Compiled by the authors (2025)

The SVM achieved the highest accuracy at 95.41%, correctly classifying 16,334 out of 17,098 test samples. RoBERTa followed at 92.61% and Naive Bayes at 92.30%. The SVM also recorded the highest precision at 0.9558, meaning that when it flagged a message as abusive, it was correct in 95.58% of cases. This is particularly valuable in real-world moderation settings where falsely flagging benign content erodes user trust and results in unjustified content removal.

Recall scores were more evenly distributed. The SVM recorded 0.8991, RoBERTa achieved 0.8852, and Naive Bayes recorded 0.8706. The SVM produced the lowest false positive rate at 1.97%, compared to 5.45% for RoBERTa.

False negative rates were 10.14% for SVM and 11.48% for RoBERTa. The combination of high precision and competitive recall confirms the SVM as the strongest performer on this dataset.

### B. Confusion Matrix Analysis

Table IV presents the confusion matrix for the SVM model. The SVM correctly identified 11,366 non-abusive messages (true negatives) and 4,945 abusive messages (true positives). It produced 229 false positives and 558 false negatives. Table V presents results for RoBERTa, which produced 8,221 true negatives, 3,654 true positives, 474 false positives, and 474 false negatives.

**TABLE IV. Confusion Matrix - SVM Model**

|                    | Predicted Non-Abusive  | Predicted Abusive     |
|--------------------|------------------------|-----------------------|
| Actual Non-Abusive | True Negatives: 11,366 | False Positives: 229  |
| Actual Abusive     | False Negatives: 558   | True Positives: 4,945 |

**TABLE V. Confusion Matrix - RoBERTa Model**

|                    | Predicted Non-Abusive | Predicted Abusive     |
|--------------------|-----------------------|-----------------------|
| Actual Non-Abusive | True Negatives: 8,221 | False Positives: 474  |
| Actual Abusive     | False Negatives: 474  | True Positives: 3,654 |

## V. DISCUSSION

### A. Interpretation of Results

The SVM's superior performance is explained by the nature of the training data. Both the Davidson et al. [2] and Jigsaw datasets predominantly contain explicit forms of abusive languages with clear lexical markers. TF-IDF

vectorisation, which assigns higher weights to words that are distinctive within abusive messages, effectively captures these surface patterns. The SVM, operating on high-quality TF-IDF feature vectors, identifies separating hyperplanes in a high-dimensional feature space with strong precision.

RoBERTa's contextual modelling capability did not provide the same benefit in this setting because the abusive language in the dataset is largely explicit rather than implicit or context-dependent. As established in the literature review, transformer models demonstrate their greatest advantage when language is subtle, sarcastic, or manipulative, precisely the characteristics of emotional abuse including gaslighting, invalidation, and coercive control [9], [10]. When abusive content is lexically explicit, contextual modelling provides a comparatively smaller advantage because the classification signals are already present at the surface lexical level.

Dataset size relative to model complexity is a further contributing factor. RoBERTa contains approximately 125 million parameters. Fine-tuning such a model effectively requires substantial domain-specific labeled data. While the combined dataset of 85,487 samples is reasonably large, it may not have been sufficient to fully leverage RoBERTa's representational capacity for a task as nuanced as emotional abuse detection. This is consistent with prior findings that existing datasets are largely derived from public social media posts containing explicit content, whereas emotional abuse frequently manifests through subtler linguistic strategies

[12].

It is important to note that RoBERTa's performance was not poor in absolute terms. At 92.61% accuracy with an F1-score of 0.8852, the model demonstrated strong classification capability and outperformed Naive Bayes. Furthermore, transformer models are specifically designed to detect the kinds of implicit and context-dependent manipulation patterns that constitute emotional abuse, patterns that TF-IDF-based models are fundamentally ill-equipped to capture. The results therefore suggest that RoBERTa's comparative advantage would likely become more pronounced with a dataset specifically annotated for emotional abuse rather than general toxicity and hate speech.

## B. Qualitative Error Analysis

To complement the quantitative evaluation, a qualitative error analysis was conducted to examine how the models perform on subtle and context-dependent instances of emotional abuse. Table VI presents examples of successful detections where the model correctly identified abuse despite the absence of explicit profanity.

TABLE VI. Successful Detections of Abuse Without Profanity

| Message                                                           | Predicted | Actual  | Why It Is Abuse                                                                          |
|-------------------------------------------------------------------|-----------|---------|------------------------------------------------------------------------------------------|
| "You should be grateful anyone pays attention to you at all."     | Abusive   | Abusive | Conditional worth framing implying the target's value is contingent on others' tolerance |
| "Everyone knows your kind are all the same - lazy and dishonest." | Abusive   | Abusive | Group-targeted inferiority language assigning negative attributes without profanity      |
| "No one else would ever put up with you."                         | Abusive   | Abusive | Isolation and emotional manipulation through conditional worth                           |

Table VII presents failure cases where the model produced false negatives. These examples reveal specific linguistic challenges the model was unable to resolve,

consistent with the technical gap identified in the introduction regarding context-dependent semantics, sarcasm, and single-message analysis constraints.

TABLE VII. Failure Cases and Linguistic Explanation

| Message                                                                                     | Predicted   | Actual  | Why It Failed                                                                                                                                               |
|---------------------------------------------------------------------------------------------|-------------|---------|-------------------------------------------------------------------------------------------------------------------------------------------------------------|
| "Oh wow, you actually did something right for once. Mark the calendar."                     | Non-Abusive | Abusive | Sarcastic backhanded praise; surface sentiment appears positive but carries contempt. The model lacks pragmatic reasoning to detect irony [8].              |
| "I am just asking questions. Do you not think people deserve to know the truth about them?" | Non-Abusive | Abusive | Gaslighting through false neutrality presenting targeted manipulation as innocent inquiry. Without conversational context, the message appears benign [11]. |
| "Nobody is forcing you to stay. But do not come crying to me when you regret it."           | Non-Abusive | Abusive | Coercive framing disguised as choice, embedding a threat of emotional withdrawal. Absence of explicit threat markers caused misclassification [15].         |

The patterns across these cases reveal a consistent theme: the model performs well when abusive intent aligns with identifiable lexical or structural patterns but struggles when abuse is conveyed through pragmatic implication, sarcasm, or context-dependent framing. The failure cases confirm that single-message analysis is insufficient for detecting emotionally abusive language in its most harmful forms. As Pradhan et al. [11] noted, emotional abuse frequently unfolds across sequences of interactions, and models that analyse messages in isolation will systematically miss the cumulative patterns that define coerciveness.

## VI. CONCLUSION

This study demonstrates that while modern deep learning models like RoBERTa provide sophisticated contextual understanding, traditional machine learning approaches remain highly effective for abuse detection when paired with precise feature engineering. The Support Vector Machine emerged as the most accurate model in this experimental setup, achieving an accuracy of 95.41% and a false positive rate of just 1.97%. The comparative analysis revealed that with datasets under 100,000 samples composed primarily of explicit abuse, well-engineered traditional models can outperform large transformer architectures,

offering important guidance for practitioners with constrained computational resources.

The research highlights a critical distinction in digital safety: emotional abuse often lacks the overt profanity that conventional filters rely on. By leveraging NLP, this project successfully bridged the technical gap, demonstrating that computational tools can identify manipulative language patterns such as gaslighting, invalidation, and coercive control. The deployed Streamlit interface demonstrates that such technology can be made accessible to non-technical users, including mental health professionals, educators, and platform moderators.

Based on the findings, several directions for future research are recommended. First, datasets specifically annotated for emotional abuse including gaslighting, guilt-tripping, and coercive control rather than general toxicity would expose models to the subtle linguistic patterns where contextual models like RoBERTa hold their greatest advantage. Second, dialogue-level transformer architectures or graph neural networks that model conversational structure would enable the system to identify cumulative coercive patterns that single-message classification misses. Third, multi-task learning frameworks that jointly optimise for abuse detection and emotion classification have been shown to improve recall for emotionally harmful content [9] and represent a promising direction for future work. Fourth, multilingual transformer models such as XLM-RoBERTa would extend detection capability to the Nigerian communication context where code-switching between English, Yoruba, Igbo, and Nigerian Pidgin is common [16].

## REFERENCES

- [1] Statista, "Number of social media users worldwide from 2017 to 2025," Statista Research Department, 2025.
- [2] T. Davidson, D. Warmley, M. Macy, and I. Weber, "Automated hate speech detection and the problem of offensive language," *Proc. Int. AAAI Conf. Web and Social Media*, vol. 11, no. 1, pp. 512-515, 2017.
- [3] J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," *Proc. 2019 Conf. North American Chapter of the ACL*, pp. 4171-4186, 2019.
- [4] Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov, "RoBERTa: A robustly optimised BERT pretraining approach," *arXiv preprint arXiv:1907.11692*, 2019.
- [5] E. Spertus, "Smokey: Automatic recognition of hostile messages in online discussions," *Proc. Conf. Human Factors in Computing Systems*, pp. 1058-1065, 1997.
- [6] A. H. Razavi, D. Inkpen, S. Uritsky, and S. Matwin, "Offensive language detection using multi-level classification," *Canadian Conf. Artificial Intelligence*, pp. 16-27, 2010.
- [7] W. Warner and J. Hirschberg, "Detecting hate speech on the world wide web," *Proc. Second Workshop on Language in Social Media*, pp. 19-26, 2012.
- [8] B. Vidgen, T. Thrush, Z. Waseem, and D. Kiela, "Learning from the worst: Dynamically generated datasets to improve online hate detection," *Proc. 59th Annual Meeting of the ACL*, pp. 1667-1682, 2021.
- [9] S. Rajamanickam, P. Mishra, H. Yannakoudakis, and E. Shutova, "Multi-task learning for emotion and abuse detection in online text," *Proc. 60th Annual Meeting of the ACL*, pp. 3456-3472, 2022.
- [10] S. Giorgi, S. C. Guntuku, and L. H. Ungar, "Integrating psychological frameworks into transformer architectures for detecting gaslighting in digital communication," *Proc. ACL*, vol. 61, pp. 2345-2362, 2023.
- [11] R. Pradhan, T. Chakraborty, and N. Goyal, "Modelling emotional trajectories for detecting manipulation and coercive control in digital conversations," *ACM Trans. Intelligent Systems and Technology*, vol. 15, no. 2, pp. 1-24, 2023.
- [12] V. Kumar, R. Singh, and M. Gupta, "Conversation-level models for detecting relational abuse patterns in digital communication," *Proc. ACM Conf. Fairness, Accountability, and Transparency*, vol. 7, pp. 234-251, 2024.
- [13] L. Zhang and Y. Chen, "Emotional abuse detection in digital text: A systematic review of approaches and research gaps," *Computers in Human Behavior*, vol. 150, pp. 107-128, 2024.
- [14] C. Nobata, J. Tetreault, A. Thomas, Y. Mehdad, and Y. Chang, "Abusive language detection in online user content," *Proc. 25th Int. Conf. World Wide Web*, pp. 145-153, 2016.
- [15] A. Lees, A. Borkowski, and M. Johnson, "Coercive control in digital spaces: Linguistic patterns of emotional abuse," *J. Interpersonal Violence*, vol. 37, no. 15-16, pp. 14567-14589, 2022.
- [16] A. Adebayo and T. Ogunleye, "Digital emotional abuse among Nigerian university students: Prevalence and reporting patterns," *J. African Digital Humanities*, vol. 8, no. 2, pp. 45-62, 2023.
- [17] C. Nobata, J. Tetreault, A. Thomas, Y. Mehdad, and Y. Chang, "Abusive language detection in online user content," *Proc. 25th Int. Conf. World Wide Web*, pp. 145-153, 2016.