



The Impact of Linguistic Features on Emotion Detection in Social Media Texts

Nuk Ghurroh Setyoningrum^{1}, N. Nelis Febriani SM², Alam³, Arif Muhamad Nurdin⁴, Dede Rizal Nursamsi⁵, Mae B. Lodana⁶*

^{1,2,3,4,5}Faculty of Science and Technology, Universitas Cipasung Tasikmalaya, Tasikmalaya 46466, Indonesia

⁶College of Information and Communication Technology, STI West Negros University, Bacolod City, Negros Occidental, Philippines

¹nuke@uncip.ac.id, ²nelis.sm@uncip.ac.id, ³alam@uncip.ac.id, ⁴arif@uncip.ac.id, ⁵dederizalnursamsi@uncip.ac.id, ⁶mae.lodana@wnu.sti.edu

ARTICLE INFORMATION

Article History:

Received: September 15, 2024

Last Revision: March 18, 2025

Published Online: April 1, 2025

KEYWORDS

Emotion Detection,
Naïve Bayes,
Count Vector,
Linguistic Features,
Social Media Texts

CORRESPONDENCE

Phone: +6281321578291

E-mail: nuke@uncip.ac.id

ABSTRACT

Emotions in text play an important role in various fields, including sentiment analysis, mental health monitoring, and marketing strategies. The present study aims to detect emotions in text using the Naïve Bayes method with attribute weighting using count vector to improve classification accuracy. The dataset used comes from the Emotion in Text public dataset, which contains 21,459 entries, with six emotion categories: happy, sad, love, fear, surprise, and anger. The text preprocessing procedure entailed tokenization, normalization, and TF-IDF representation, which was subsequently employed in the model. The experimental outcomes demonstrate that the model attained 86% accuracy, 88% precision, 76% recall, and 80% F1-score. The model demonstrated a high level of accuracy in detecting positive emotions (F1-score 0.88) and negative emotions (F1-score 0.91) but encountered challenges in recognizing emotions such as surprise (F1-score 0.63) and love (F1-score 0.72) due to imbalanced data in these categories. While the Naïve Bayes model demonstrates proficiency in terms of computational speed and interpretability, its efficacy is hindered in scenarios involving data imbalance and emotions with overlapping meanings. Potential enhancements include dataset balancing and the exploration of deep learning-based methods, such as LSTM or BERT, to capture more intricate emotional contexts present in textual data.

1. INTRODUCTION

In the era of digital communication, social media platforms have become a primary medium for expressing opinions, sharing experiences, and communicating emotions. The vast amount of textual data generated on these platforms presents an opportunity to analyze emotions for various applications, such as sentiment analysis, mental health monitoring, and marketing strategies [1], [2].

Recent advances in natural language processing (NLP) have enabled researchers to explore the linguistic features embedded in social media texts to enhance emotion detection accuracy. These features, which include lexical, syntactic, semantic, and pragmatic aspects, play a crucial role in identifying nuanced emotional expressions. However, the informal and dynamic nature of social media language, characterized by slang, emojis, abbreviations, and code-mixing, poses challenges in achieving precise emotion classification [3], [4].

Despite significant progress in emotion detection, research gaps remain in understanding the specific contributions of various linguistic features to emotion classification models. Most studies have focused on traditional features or employed general-purpose NLP techniques, which may not fully capture the complex emotional patterns in social media texts. Furthermore, the lack of standardized datasets and methodologies hinders the development of universally applicable models. This study aims to address these gaps by systematically investigating the impact of diverse linguistic features on emotion detection in social media texts. By doing so, it seeks to provide deeper insights into how linguistic nuances can improve the performance of emotion detection systems, offering potential advancements for applications in both academic research and industry practices [5], [6].

Emotions can be expressed through facial expressions, voice, and written comments on social media. However, writing often fails to accurately convey emotions, so

messages are often misinterpreted [3], [7]. Emotions play an important role in everyday communication. This research produces a dataset for text-based emotion classification, which is divided into the categories of happy, sad, fear, love, shock, and anger. These emotions are used as keywords in this research search [8], [9].

Appropriate emotions at the right time and situation can influence the outcome of human activities. Emotions are usually expressed implicitly and triggered by specific events [10], [11]. Text describes situations that trigger emotions explicitly and is the main medium in computer communication such as email, blogs and social media. Emotions are classified into two: negative emotions (anger, sadness, fear, shock) and positive emotions (joy, love). In business, emotion analysis helps evaluate products and identify consumer complaints. In politics, it is useful for tracking support on political issues. In psychology, it monitors the emotional state of individuals and detects depression and anxiety disorders [12]. This shows the importance of detecting emotions.

Text mining automates the extraction of information from text using machine learning algorithms to detect emotions [13], [14]. The advantages of text mining include efficient analysis of large amounts of text data, uncovering hidden patterns, and providing important insights. It supports better decision-making, improves understanding of customers, optimizes business strategies, and detects trends and sentiment in real-time [15], [16].

A classification approach is used to detect emotions by grouping data based on similar criteria. Each emotion in this classification is assigned a specific class. Common methods used include C4.5, Naïve Bayes, and K-Nearest Neighbor (KNN) [17]. One of the frequently used classification methods is the Naïve Bayes Classifier. Naïve Bayes classifies data based on Bayes' theorem by utilizing conditional probabilities. This method is fast, simple, and has high accuracy, but it cannot measure the accuracy of the prediction. To overcome this, count vector can be used to weight attributes and improve the accuracy of Naïve Bayes Classifier [18], [19].

This research uses a public dataset containing 21,459 data to analyze the performance of Naïve Bayes and attribute weighting using count vector in classifying emotions based on posts on social media. As a contribution, this research highlights how linguistic elements affect emotion detection accuracy, offering actionable insights to refine text mining algorithms and methodologies. In addition, this research also contributes to the combination of the Naïve Bayes algorithm with count vector-based attribute weighting to improve the accuracy of emotion classification and overcome the limitations that exist in previous studies. It is hoped that this method will help future research in choosing the right method for text mining application development, improve maximum performance, and produce useful datasets for text-based emotion classification in the future.

Although much research has been conducted in the field of emotion detection in social media texts, most studies tend to use general feature-based approaches or pre-trained models without considering in depth the specific contributions of linguistic features, such as lexical, syntactic, semantic, and pragmatic. In addition, most previous studies often ignore the complexity of informal

language typical in social media, including the use of slang, emoticons, abbreviations, and code-mixing. This makes existing emotion detection models less able to capture complex and diverse emotional nuances. The lack of research that focuses on systematically analyzing the influence of specific linguistic features on model performance is also an obstacle in developing more accurate and adaptive solutions. Therefore, more in-depth exploration of the role of linguistic features is needed to fill this research gap and improve the capabilities of emotion detection systems in various practical applications.

The purpose of this study is to systematically analyze the influence of various linguistic features, including lexical, syntactic, semantic, and pragmatic, on the accuracy of emotion detection in social media texts. This research seeks to identify linguistic features that have significant contributions in improving the performance of emotion detection models, particularly in the context of informal language frequently used on social media. The results of this study can be used to improve emotion detection models in digital marketing applications and mental health analysis. The article begins with a literature review, followed by the research methodology, results and discussion, and conclusions.

2. RELATED WORK

Several related studies have been used as references by the author as a basis for conducting this research. Among these studies, one of the most recent and relevant is a study published by Fera Fanesia et al [15]. The purpose of this research is to combine various features in emotion detection analysis. This research produces a dataset that can be used for emotion detection purposes, which uses the Naïve Bayes method with testing involving a combination of N-gram features. The results show an accuracy rate of 0.555 with N-gram features that include orthographic features and linguistic features.

Another research conducted by Akhmad Fadjeri et al also an important reference [3], [20]. This research aims to evaluate the use of appropriate emotions using the Naïve Bayes method supported by count vector weighting [21]. The dataset used comes from various sources and is classified into two models, namely positive and negative. The results show that the use of count vector weighting can improve accuracy, with an accuracy value reaching 87.7%.

Meanwhile, the third research conducted by Erfian Junianto combines text mining techniques with Naïve Bayes and Particle Swarm Optimization (PSO) to optimize attributes and improve emotion detection accuracy. This research uses a dataset that has been previously classified in four emotion categories, namely anger, fear, joy, and sadness, with results that show satisfactory performance in emotion classification [22].

Anzum and Gavrilova introduced a novel emotion detection approach for Twitter data using a new input representation called SSEL, which combines stylistic, sentiment, and linguistic features. Their method employed a genetic algorithm to optimize feature selection and trained an ensemble model of XGBoost, Random Forest, and SVM, achieving state-of-the-art accuracy (96.49%) across six emotion categories: sadness, joy, love, anger, fear, and surprise. Unlike prior studies that primarily relied

on linguistic features like TF-IDF or Word2Vec, their approach addressed the integration of stylistic and sentiment features, which are often overlooked. However, challenges such as limited semantic capture, class imbalances in datasets, and underutilized feature combinations persist [23].

Abdullah Al Maruf et al. conducted a comprehensive survey on text-based emotion detection (TBED), exploring state-of-the-art systems, techniques, and datasets. The study highlighted the critical role of machine learning and deep learning algorithms in identifying emotional states from textual data, emphasizing the psychological, social, and commercial significance of emotion detection. While the survey provided an extensive review of models, feature extraction methods, and datasets, it identified gaps in performance analysis, the scalability of current methods, and the integration of advanced architectures like transformers for domain-specific tasks [24].

De León Languré and Zareei proposes a critical analysis of sentiment analysis (SA) and text emotion detection (TED), highlighting gaps in current research that predominantly focus on algorithmic comparisons while neglecting the role of emotion models, training corpora, and validation data. Their findings reveal that this lack of standardization creates challenges in comparing results and advancing real-world applications [1], [25]. The study underscores the need for a unified assessment framework to address these gaps, enabling meaningful performance evaluation across diverse scenarios. This research identifies a significant gap in holistic evaluation approaches and emphasizes the importance of bias minimization and generalizability.

The research by Shelke et al. presents an advanced approach to text-based emotion recognition on social media data using the Leaky ReLU Activated Deep Neural Network (LRA-DNN). The methodology involves preprocessing, feature extraction, ranking, and classification, achieving notable improvements in accuracy (94.77%), sensitivity (92.23%), and specificity (95.91%) over prior methods such as CNN, ANN, and DNN. However, the study addresses challenges like misclassification errors but overlooks the potential impact of multilingual or context-specific nuances in social media texts. This highlights the need for further exploration of linguistic and cultural diversity in emotion analysis [26].

In their study, Salloum et al. (2024) explored the clustering of emotional content in tweets using a hybrid approach that combines K-means clustering and dimensionality reduction via 2D Principal Component Analysis (PCA). By analyzing a dataset of 40,000 tweets annotated with 13 emotions, the study demonstrated the effectiveness of reducing high-dimensional emotional data into interpretable clusters, improving computational efficiency and visualization. However, while their methodology enhanced clustering outcomes, it did not fully address the contextual nuances of emotions, such as sarcasm or idiomatic expressions, leaving room for further exploration in incorporating semantic-rich NLP models like transformers. This research highlights the potential for improved emotion classification in sentiment analysis [27].

Chowanda et al. conducted a comprehensive exploration of machine learning techniques for text-based emotion recognition in social media conversations. Using

the AffectiveTweets dataset with over 7,000 labeled utterances, they tested seven algorithms ranging from traditional methods like Naïve Bayes and Decision Tree to advanced models such as Generalised Linear Model (GLM) and Artificial Neural Networks (ANN). Their findings highlighted GLM as the most effective model, achieving the highest accuracy (90.2%) and F1 score (90.1%). However, the study's reliance on pre-defined feature sets and the exclusion of contextual nuances in text presents a gap, suggesting potential improvements through the application of advanced deep learning models like transformers or context-aware approaches [28].

Yang and Zhang explored public emotions and visual preferences associated with East Coast Park in Singapore using a deep learning framework that integrates Transformer BERT and CNN-VGG models to analyze text reviews and images from Google Maps. Their findings revealed a predominance of joy in public emotions and diverse visual preferences for park features, alongside challenges in linking negative emotions to specific environmental factors due to limited data on dissatisfaction. The research highlights the potential of combining social media data with advanced computational models but identifies a gap in understanding fine-grained negative emotional responses and their environmental correlates [29]. On other research, reviewed emotion-based methods for misinformation detection, emphasizing the integration of emotion, sentiment, and stance as key factors to identify fake news and rumors. They analyzed various machine learning and deep learning approaches, including advanced fusion methods, and highlighted their applications in improving the accuracy of misinformation detection. However, their work identifies a research gap in leveraging fine-grained emotional analysis and integrating multilingual datasets to address the challenges of misinformation detection across diverse cultural and linguistic contexts [30].

Abu-Salih et al. conducted a comparative study of eight emotion detection APIs, including IBM Watson, ParallelDots, and NLP Cloud, by analyzing their performance on two datasets: a publicly available NLP emotion dataset and a newly created Twitter Conversations dataset. The study evaluated the APIs based on classification accuracy, precision, recall, and F1 score, revealing significant variability in their performance and emphasizing the need for thorough API selection. A notable research gap identified is the insufficient empirical comparison of APIs using untrained and multilingual datasets, highlighting the need for more robust and context-aware emotion detection tools [31].

Purtiwi et al. examined the influence of social media usage and emotions on risk perception and preventive behaviors related to COVID-19 among university students, utilizing the Preventive Behavior Model. Their findings reveal that factors such as anxiety about social safety, predictions of COVID-19 spread, and increased social media exposure significantly impact preventive behaviors, both directly and indirectly through negative emotions. However, the study identifies a gap in addressing the negative influence of perceived risk and media exposure on preventive behaviors, highlighting the need for more nuanced analyses of these factors in pandemic response frameworks [32].

Li et al. advanced social media sentiment analysis by transforming it from single-label polarity classification to multi-label emotion classification, leveraging Plutchik's eight emotion categories and the AC-BiLSTM model. Their work introduced correlation constraints and emotion dictionaries to enhance classification accuracy and visualized results using emotion graphs for better traceability and reasoning. However, a significant research gap exists in the availability of diverse, robust datasets and comparisons with state-of-the-art models, limiting the generalizability of the proposed framework [33].

3. METHODOLOGY

The research flow on the Impact of Linguistic Features on Emotion Detection in Social Media Texts begins with the collection of data on social media texts to be analyzed. After that, linguistic features in the text, such as word choice, sentence structure, and punctuation, are analyzed to understand how these features affect emotion detection. The next step is the application of emotion detection methods or algorithms by considering the impact of the identified linguistic features. Next, the accuracy and effectiveness of emotion detection based on a particular combination of linguistic features will be evaluated. The evaluation results will be used to draw conclusions about the importance of linguistic features in improving the accuracy of emotion detection in social media texts, as well as provide recommendations for the development of better emotion detection methods in the future.

The following is the working structure applied in this research, which consists of several stages. These are presented in Figure 1 below:

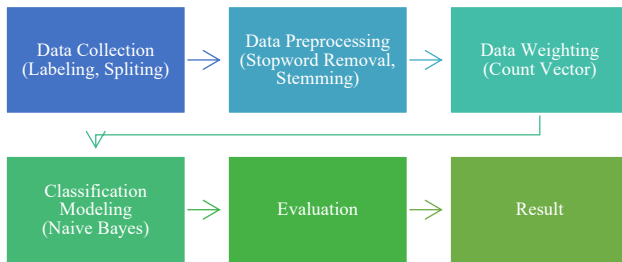


FIGURE 1. RESEARCH METHODOLOGY

3.1 Text Mining

Text mining is an intensive process that aims to gain insights by collecting and analyzing documents using specialized analysis tools. As an emerging field, text mining aims to address the problems that arise from the overload of available information. This is achieved using different techniques from data mining, natural language processing, and information retrieval. Text mining focuses on extracting relevant and useful information from various data sources, both structured and unstructured.

Text mining, in its course, recognizes as well as explores patterns in data from large collections of documents. These patterns are not integrated into conventional database records but remain stored in unstructured text data formats. This allows researchers to discover hidden linkages and gain new insights that may not be detected through traditional data analysis methods [3], [34].

One of the main advantages of text mining is its ability to handle large amounts of text data, automatically identifying important trends and patterns. This has significant benefits in a variety of sectors, including the business industry, education, healthcare, and scientific research. With text mining, organizations can tap into information hidden in text documents such as emails, articles, reports, and social media, allowing them to make more efficient and timely decisions [25], [35], [36].

3.2 Naives Bayes

Naïve Bayes is a simple and efficient classification algorithm that is suitable for large datasets. Using the Bayesian theorem, this algorithm utilizes the probability of documents against categories (prior) and classifies text based on the maximum posterior probability. The first step in using Naïve Bayes is to calculate the probability of documents in each category (prior) [4]. The formula applied for the calculation is as follows [6]:

$$P(w_j) = \frac{Nw_j}{Nc} \quad (1)$$

Nw_j is the number of training data that fall into emotion category w_j , while Nc is the total training data. The equation for the posterior probability is as follows:

$$P(w_j|d) = \frac{P(d|w_j) \cdot P(w_j)}{P(d)} \quad (2)$$

In this context, $P(X_i|W_j)$ represents the likelihood probability, $P(W_j)$ is the prior, and $P(X_i)$ is the probability of the word occurring. Equation 3 simplifies the posterior calculation by ignoring the likelihood of the word occurrence, as this factor does not affect the comparison of classification results between categories.

$$P(W_j|X_i) = P(X_i|W_j) \cdot P(W_j) \quad (3)$$

The likelihood probabilities are calculated from $i=1$ to $i=k$, which allows the posterior value calculation to be performed as shown in equation 4.

$$P(W_j|X_i) = P(X_i|W_j) \times P(X_2|W_j) \times \dots \times P(X_k|W_j) \quad (4)$$

3.3 Data Preprocessing

In the research on “The Impact of Linguistic Features on Emotion Detection in Social Media Texts”, an important first step is data preprocessing. Preprocessing is a critical stage in text analysis that involves a series of steps to clean, tidy, and prepare data to be ready for further analysis. In the context of this research, the preprocessing process aims to prepare social media text data that will be used in emotion detection analysis. This stage includes steps such as punctuation removal, text normalization, removal of irrelevant words, and tokenization of text into smaller units such as words or phrases. By doing the right preprocessing process, it is expected that the resulting data will be cleaner, structured, and ready to be analyzed in the context of emotion detection.

Data that has been labeled will undergo a preprocessing stage. This process involves a series of steps as follows:

3.3.1 Class/Label Creation based on Category

This process involves grouping texts based on the types of emotions contained in them, so that each text will be labeled according to the emotions represented. This step is important because it is the foundation for the development of a classification model that can distinguish and identify the various emotions contained in social media texts. By defining these classes or labels, the research can move to the next stage of analyzing the impact of linguistic features on emotion detection in more detail, enabling more precise sentiment analysis, improving model interpretability, and facilitating deeper understanding of public emotional responses in various online contexts.

Table 1 shows the process of forming classes or labels based on the categories used in this study, including the definition of each label, representative examples of social media texts for each category, and the contextual characteristics that distinguish one label from another to support accurate classification and interpretation.

TABLE 1. CLASS / LABEL CREATION

Text	Keyword	Emotion (Label)
I feel pretty pathetic most of the time.	Pathetic	Sadness
I feel selfish and spoiled.	Selfish	Anger
I feel romantic too.	Romantic	Love
I am now nearly finished the week detox, and I feel amazing.	Amazing	Surprise
I feel very happy and excited since I learned so many things.	Excited	Happy

3.3.2 Stopword Removal

This process involves the elimination of common words or words that have no special meaning in a particular language, such as “and”, “or”, “of”, and so on, which often appear in the text but do not make a significant contribution to emotion analysis. By removing these stop words, it is expected to improve the quality of emotion analysis by focusing attention on words that are more meaningful and relevant in expressing emotions in social media texts. Table 2 presents an example of the application of the stopword removal process, showing the difference between the original text and the text that has been processed by removing stopwords to improve the efficiency of linguistic analysis.

TABLE 2. STOPWORD REMOVAL EXAMPLE

I	or	a	was	with
am	to	it	as	has
and	be	is	on	did
of	an	too	have	but
the	on	are	been	now

After determining several words that are considered meaningless or often referred to as stopwords, here are the results of the stopword removal process. The results of the stopword removal process, which shows a comparison of the text before and after the removal of common words that do not contribute significantly to the analysis, thus improving the data quality for the emotion detection process, are presented in Table 3. This refinement step plays a crucial role in reducing noise, enhancing feature extraction accuracy, and optimizing the overall performance of the classification model by focusing only on meaningful textual content.

TABLE 1. STOPWORD REMOVAL RESULT

Class	Label
fell pretty pathetic most time	anger
feel selfish spoiled	anger
feel romantic	anger
finished week detox feel amazing	anger
feel happy excited learned many things	anger

4. Result and Discussion

This section presents the outcomes of the analysis conducted, covering data processing, classification method implementation, and model performance evaluation. The results are thoroughly analyzed to address the research objectives and assess the effectiveness of the applied approach. The discussion includes comparisons with previous studies and relates the findings to relevant theories. Therefore, this section plays a crucial role in interpreting the meaning behind the processed data and highlighting the scientific and practical contributions of this research.

4.1 Data Collection

In this study, researchers employed the Emotion in text public dataset, which contains 21,459 data points, to analyze the performance of Naïve Bayes and attribute weighting using count vector in classifying emotions based on social media posts. The attribute employed in this study is Emotion, which will be shown in Table 4 below:

TABLE 4. EMOTION IN TEXT DATASET

No.	Text	Emotion
0	I didn't feel humiliated	Sadness
1	I can go from feeling so hopeless to so damned...	Sadness
2	I'm grabbing a minute to post I feel greedy wrong	Anger
3	I am ever feeling nostalgic about the fireplace...	Love
4	I am feeling grouchy	Anger
...	...	...
21454	Melissa stared at her friend in dims	Fear
21455	Successive state elections have seen the govern...	Fear
21456	Vincent was irritated but not dismay	Fear
21457	Kendall-Hume turned back to face the dismayed ...	Fear
21458	I am dismayed, but not surprise	Fear

4.2 Calculation Using Naïve Bayes Method

The present study utilizes the Multinomial Naïve Bayes algorithm for the purpose of text classification based on emotions. Initially, the text data undergoes processing using TF-IDF, resulting in the attainment of a numerical representation. The model is then trained using training data and tested with testing data, where the prediction results are evaluated with metrics such as accuracy, precision, recall, and F1-score. To further enhance understanding, the confusion matrix is presented in graphical form to illuminate the distribution of classification errors. Finally, the efficacy of the model is assessed by evaluating its performance on a single text sample, thereby providing a comprehensive evaluation of the emotion prediction results.

TABLE 5. THE PERFORMANCE RESULTS OF EMOTION DETECTION

Emotion	Precision	Recall	F1-Score	Support
Anger	0.92	0.77	0.84	617
Fear	0.87	0.77	0.82	531
Happy	0.82	0.96	0.88	1381
Love	0.91	0.60	0.72	318
Sadness	0.87	0.94	0.91	1277
Surprise	0.87	0.49	0.63	168
Accuracy			0.86	4292
Macro Avg.		0.88	0.80	4292
Weighted Avg.		0.86	0.85	4292

The classification results from Table 5, obtained through the implementation of the Naïve Bayes method, demonstrated an accuracy of 86%, with the highest performance observed in the categories of happiness and sadness, which attained f1-scores of 0.88 and 0.91, respectively. The model's superior performance in these categories can be attributed to the availability of more data compared to other categories. Conversely, the model demonstrated a lower level of performance in the surprise (f1-score 0.63) and love (f1-score 0.72) classes, suggesting a challenge in detecting emotions with limited data. Furthermore, the recall value for the surprise category was found to be 0.49, suggesting that a significant proportion of surprise data was misclassified into other categories.

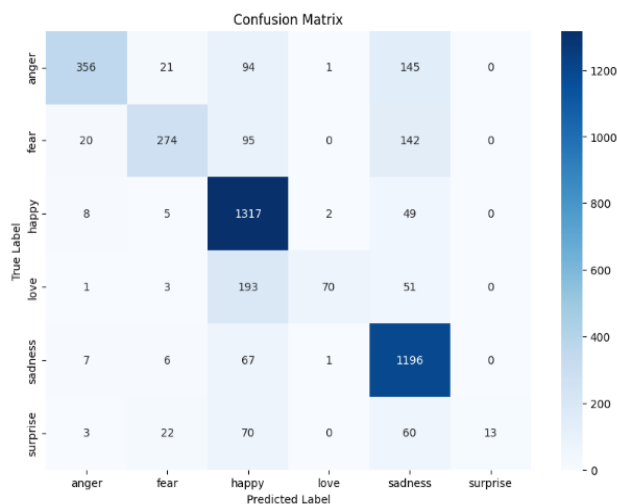


FIGURE 2. RESULTS EMOTION OF THE CONFUSION MATRIX

As illustrated in Figure 2, the results of the emotion classification using Naïve Bayes are displayed. The Y-axis, designated as the "True Label," corresponds to the actual class, while the X-axis, labeled as the "Predicted Label," represents the class predicted by the model. The diagonal values from the top left to the bottom right indicate the number of correct predictions for each class, while the numbers outside the diagonal indicate misclassification.

The model demonstrated a high degree of accuracy in its predictions for the "happy" class, with a total of 1,317 correct predictions. However, misclassifications occurred, with 49 data points classified as "sadness" and 95 data points classified as "fear." The sadness class demonstrated adequate classification, with 1,196 accurate predictions,

though minor errors were observed in the happy (67 data) and love (70 data) classes.

Conversely, the surprise class exhibited the poorest performance, with a mere 13 accurate predictions, while 60 surprise records were classified as sadness and 70 as happy, indicating a pronounced challenge for the model in differentiating between these emotions. Anger and fear demonstrated moderate accuracy, exhibiting some misclassification as either happy or sad, suggesting a similarity in their features.

4.3 Sample Data Test Using Naïve Bayes Method

An investigation will be conducted into the use of the Naïve Bayes method for the purpose of detecting emotions in text through the utilization of sample data. The Naïve Bayes algorithm is a probabilistic classification technique that is frequently employed in Natural Language Processing (NLP), encompassing sentiment analysis and emotion detection. Utilizing a probabilistic approach, this method forecasts the emotion category of a text based on word patterns that have been acquired from the training dataset.

The testing process involves the provision of example sentences as input, which are then analyzed and classified by the model to assess its emotional detection capabilities. The prediction results are then compared with the actual labels to evaluate the extent to which the model can accurately recognize emotions. This chapter will also discuss the analysis of prediction results, potential classification errors, and suggestions for improving the model to optimize its performance in detecting emotions in text.

```
[ ] # Sample prediksi 1 data
sample_text = ["I am so happy today!"]
sample_text_cleaned = [clean_text_v2(text) for text in sample_text]
sample_vectorized = tfidf_vectorizer.transform(sample_text_cleaned)
sample_prediction = model.predict(sample_vectorized)
sample_prediction_label = label_encoder.inverse_transform(sample_prediction)

print(f'Sample Prediction: {sample_text[0]} -> {sample_prediction_label[0]}')
```

Sample Prediction: I am so happy today! -> happy

FIGURE 3. TESTING EMOTION TEXT DETECTION

Figure 3 explains that to test the Naïve Bayes model with one example text and see the resulting emotion prediction. First, the text "I am so happy today!" is entered as input, then cleaned using the `clean_text_v2` function, which removes non-alphabetic characters, converts them to lowercase, and performs lemmatization. After that, the cleaned text is converted into numeric form using TF-IDF Vectorizer, so that it can be processed by the model. The model then performs predictions based on the patterns it has learned, resulting in a numerical value that represents the emotion category.

This value is then returned to its original form (e.g., "happy", "sadness") using `LabelEncoder`. Finally, the prediction result is displayed in a format that shows the input text along with the predicted emotion, e.g. "I am so happy today!" happy. This syntax helps in evaluating the extent to which the model can recognize emotions from new text not included in the training data.

5. CONCLUSIONS

In this research, an analysis of emotion detection in text was conducted to evaluate the effectiveness of a model in recognizing different categories of emotions. The process involved text preprocessing techniques such as tokenization, lemmatization, and TF-IDF vectorization, followed by classification using the Naïve Bayes model. The model achieved an accuracy of 86% with good precision, recall, and F1-scores, but struggled with classifying "surprise" and "love" emotions due to limited data. The model's strengths lie in its computational speed and interpretability, but it faces challenges with imbalanced data and distinguishing emotions with overlapping meanings. To improve the model, techniques like dataset balancing and exploring deep learning-based models such as LSTM or BERT could be considered to better capture emotional context in text. In conclusion, Naïve Bayes can be used for emotion detection in text, but further enhancements are possible.

REFERENCES

- [1] A. D. L. Langure and M. Zareei, "Breaking Barriers in Sentiment Analysis and Text Emotion Detection: Toward a Unified Assessment Framework," *IEEE Access*, vol. 11, pp. 125698–125715, 2023, doi: 10.1109/ACCESS.2023.3331323.
- [2] Y. Hu and H. Wang, "A Review of Mining User Needs Based on Text Sentiment Analysis Technology and KANO Model," in *Proceedings of the 2024 9th International Conference on Intelligent Information Processing*, New York, NY, USA: ACM, Nov. 2024, pp. 257–264. doi: 10.1145/3696952.3696987.
- [3] A. Nurlaila and R. Saptono, "CLASSIFICATION OF CUSTOMERS EMOTION USING NAÏVE BAYES CLASSIFIER (Case Study: Natasha Skin Care)."
- [4] B. Li, H. Fei, F. Li, T. Chua, and D. Ji, "Multimodal Emotion-Cause Pair Extraction with Holistic Interaction and Label Constraint," *ACM Transactions on Multimedia Computing, Communications, and Applications*, Aug. 2024, doi: 10.1145/3689646.
- [5] W.-C. Huang and Y.-L. Hsueh, "An Emotional Dialogue System Using Conditional Generative Adversarial Networks with a Sequence-to-Sequence Transformer Encoder," *ACM Transactions on Asian and Low-Resource Language Information Processing*, vol. 23, no. 12, pp. 1–23, Dec. 2024, doi: 10.1145/3698394.
- [6] Y. Jing and X. Zhao, "DQ-Former: Querying Transformer with Dynamic Modality Priority for Cognitive-aligned Multimodal Emotion Recognition in Conversation," in *MM 2024 - Proceedings of the 32nd ACM International Conference on Multimedia*, Association for Computing Machinery, Inc, Oct. 2024, pp. 4795–4804. doi: 10.1145/3664647.3681599.
- [7] Y. Cai, R. Ye, J. Xie, Y. Zhou, Y. Xu, and Z. Wu, "Robust Representation Learning for Multimodal Emotion Recognition with Contrastive Learning and Mixup," in *Proceedings of the 2nd International Workshop on Multimodal and Responsible Affective Computing*, New York, NY, USA: ACM, Oct. 2024, pp. 93–97. doi: 10.1145/3689092.3689418.
- [8] A. Buker and A. Vinciarelli, "Emotion Recognition for Multimodal Recognition of Attachment in School-Age Children," in *International Conference on Multimodal Interaction*, New York, NY, USA: ACM, Nov. 2024, pp. 312–320. doi: 10.1145/3678957.3685747.
- [9] L. L. V and D. K. Anguraj, "A DENSE SPATIAL NETWORK MODEL FOR EMOTION RECOGNITION USING LEARNING APPROACHES," *ACM Transactions on Asian and Low-Resource Language Information Processing*, Aug. 2024, doi: 10.1145/3688000.
- [10] S. J. Cantrell, R. M. Winters, P. Kaini, and B. N. Walker, "Sonification of Emotion in Social Media: Affect and Accessibility in Facebook Reactions," *Proc ACM Hum Comput Interact*, vol. 6, no. CSCW1, Apr. 2022, doi: 10.1145/3512966.
- [11] F. Hasan, Y. Li, J. R. Foulds, S. Pan, and B. Bhattacharjee, "DoubleDistillation: Enhancing LLMs for Informal Text Analysis using Multistage Knowledge Distillation from Speech and Text," in *International Conference on Multimodal Interaction*, New York, NY, USA: ACM, Nov. 2024, pp. 526–535. doi: 10.1145/3678957.3685705.
- [12] A. Fadjeri, K. Hidayat, and D. R. Handayani, "DETEKSI EMOSI PADA TEKS MENGGUNAKAN ALGORITMA NAÏVE BAYES," vol. 1, no. 2, pp. 1–4, 2021, doi: 10.53863/juristik.v1i02.365.
- [13] Y. Xu, Y. Zhou, Y. Cai, J. Xie, R. Ye, and Z. Wu, "Multimodal Emotion Captioning Using Large Language Model with Prompt Engineering," in *Proceedings of the 2nd International Workshop on Multimodal and Responsible Affective Computing*, New York, NY, USA: ACM, Oct. 2024, pp. 104–109. doi: 10.1145/3689092.3689403.
- [14] T. Thebaud *et al.*, "Multimodal Emotion Recognition Harnessing the Complementarity of Speech, Language, and Vision," in *International Conference on Multimodal Interaction*, New York, NY, USA: ACM, Nov. 2024, pp. 684–689. doi: 10.1145/3678957.3689332.
- [15] F. Fanesya and R. Cahya Wihandika, "Deteksi Emosi Pada Twitter Menggunakan Metode Naïve Bayes Dan Kombinasi Fitur," 2019. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [16] C. Strapparava and R. Mihalcea, "SemEval-2007 Task 14: Affective Text," 2007.
- [17] L. Bulla and M. Mongiovi, "Adequate Prompting Improves Performance of Regression Models of Emotional Content," Association for Computing Machinery (ACM), Sep. 2024, pp. 135–142. doi: 10.1145/3677525.3678653.
- [18] "308-Article Text-377-1-10-20181116".
- [19] C. Cherry, S. M. Mohammad, and B. De Bruijn, "Binary Classifiers and Latent Sequence Models for Emotion Detection in Suicide Notes," *Biomed Inform Insights*, vol. 5s1, p. BII.S8933, Jan. 2012, doi: 10.4137/bii.s8933.
- [20] Md. A. Akber, T. Ferdousi, R. Ahmed, R. Asfara, R. Rab, and U. Zakia, "Personality and emotion—A comprehensive analysis using contextual text embeddings," *Natural Language Processing Journal*,

- vol. 9, p. 100105, Dec. 2024, doi: 10.1016/j.nlp.2024.100105.
- [21] Y. Shang and T. Fu, "Multimodal fusion: A study on speech-text emotion recognition with the integration of deep learning," *Intelligent Systems with Applications*, vol. 24, Dec. 2024, doi: 10.1016/j.iswa.2024.200436.
- [22] C. O. Alm, D. Roth, and R. Sproat, "Emotions from text: machine learning for text-based emotion prediction," 2005. [Online]. Available: <http://l2r.cs.uiuc.edu/>
- [23] F. Anzum and M. L. Gavrilova, "Emotion Detection From Micro-Blogs Using Novel Input Representation," *IEEE Access*, vol. 11, pp. 19512–19522, 2023, doi: 10.1109/ACCESS.2023.3248506.
- [24] A. Al Maruf, F. Khanam, M. 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, 2024, doi: 10.1109/ACCESS.2024.3356357.
- [25] A. De Leon Langure and M. Zareei, "Evaluating the Effect of Emotion Models on the Generalizability of Text Emotion Detection Systems," *IEEE Access*, vol. 12, pp. 70489–70500, 2024, doi: 10.1109/ACCESS.2024.3401203.
- [26] N. Shelke, S. Chaudhury, S. Chakrabarti, S. L. Bangare, G. Yogapriya, and P. Pandey, "An efficient way of text-based emotion analysis from social media using LRA-DNN," *Neuroscience Informatics*, vol. 2, no. 3, p. 100048, Sep. 2022, doi: 10.1016/j.neuri.2022.100048.
- [27] S. Salloum, K. Alhumaid, A. Salloum, and K. Shaalan, "K-means Clustering of Tweet Emotions: A 2D PCA Visualization Approach," *Procedia Comput Sci*, vol. 244, pp. 30–36, 2024, doi: 10.1016/j.procs.2024.10.175.
- [28] A. Chowanda, R. Sutoyo, Meiliana, and S. Tanachutiwat, "Exploring Text-based Emotions Recognition Machine Learning Techniques on Social Media Conversation," in *Procedia Computer Science*, Elsevier B.V., 2021, pp. 821–828. doi: 10.1016/j.procs.2021.01.099.
- [29] C. Yang and Y. Zhang, "Public emotions and visual perception of the East Coast Park in Singapore: A deep learning method using social media data," *Urban For Urban Green*, vol. 94, Apr. 2024, doi: 10.1016/j.ufug.2024.128285.
- [30] Z. Liu, T. Zhang, K. Yang, P. Thompson, Z. Yu, and S. Ananiadou, "Emotion detection for misinformation: A review," Jul. 01, 2024, *Elsevier B.V.* doi: 10.1016/j.inffus.2024.102300.
- [31] B. Abu-Salih *et al.*, "Emotion detection of social data: APIs comparative study," *Heliyon*, vol. 9, no. 5, May 2023, doi: 10.1016/j.heliyon.2023.e15926.
- [32] R. G. C. Pertiwi, F. Artwodini, and R. Nadlifatin, "Analysis of the Role of the Use of Social Media and Emotions in Risk Perception and Prevention Behavior of Covid-19 Using the Preventive Behavior Model," in *Procedia Computer Science*, Elsevier B.V., 2024, pp. 1145–1153. doi: 10.1016/j.procs.2024.03.110.
- [33] Y. Li, J. Chan, G. Peko, and D. Sundaram, "Mixed emotion extraction analysis and visualisation of social media text," *Data Knowl Eng*, vol. 148, Nov. 2023, doi: 10.1016/j.datak.2023.102220.
- [34] C. Strapparava and R. Mihalcea, "SemEval-2007 Task 14: Affective Text," 2007.
- [35] R. Pan, J. A. García-Díaz, M. Á. Rodríguez-García, and R. Valencia-García, "Spanish MEACorpus 2023: A multimodal speech-text corpus for emotion analysis in Spanish from natural environments," *Comput Stand Interfaces*, vol. 90, Aug. 2024, doi: 10.1016/j.csi.2024.103856.
- [36] Q. Zhao, Y. Xia, Y. Long, G. Xu, and J. Wang, "Leveraging sensory knowledge into Text-to-Text Transfer Transformer for enhanced emotion analysis," *Inf Process Manag*, vol. 62, no. 1, Jan. 2025, doi: 10.1016/j.ipm.2024.103876.

AUTHORS



Nuk Ghurroh Setyoningrum

Faculty of Science and Technology,
Department of Information System,
Universitas Cipasung Tasikmalaya.



N. Nelis Febriani SM

Faculty of Science and Technology,
Department of Information System,
Universitas Cipasung Tasikmalaya.



Alam

Faculty of Science and Technology,
Department of Software Engineering,
Universitas Cipasung Tasikmalaya.



Arif Muhamad Nurdin

Faculty of Science and Technology,
Department of Information System,
Universitas Cipasung Tasikmalaya.



Dede Rizal Nursamsi

Faculty of Science and Technology,
Department of Information System,
Universitas Cipasung Tasikmalaya.



Mae B. Lodana

College of Information and Communication
Technology, STI West Negros University,
Philippines.