

Journal of Engineering Technology and Applied Physics

Sentiment Analysis Using Facebook Comments

Shakib Hossain *, Md. Abu Yousuf, Md. Jakir Hossen, Jalal Uddin, Sadman Sadik Khan, and Abdus Sattar
Student, Dept. of Computer Science and Engineering, Daffodil International University, Dhaka, Bangladesh.
Student, Faculty of Engineering and Technology, Multimedia University, Melaka, Malaysia.

Associate Professor, Faculty of Engineering and Technology, Multimedia University, Melaka, Malaysia.

Student, Dept. of Computer Science and Engineering, Daffodil International University, Dhaka, Bangladesh.

Student, Dept. of Computer Science and Engineering, Daffodil International University, Dhaka, Bangladesh.

Associate Professor, Dept. of Computer Science and Engineering, Daffodil International University, Dhaka, Bangladesh.

*Corresponding author's email address: jakir.hossen@mmu.edu.my, ORCID: 0000-0002-9978-7987

<https://doi.org/10.33093/jetap.2026.8.1.10>

Manuscript Received: XX Mon 20XX. Revised: XX Mon 20XX. Accepted: XX Mon 20XX. Published: XX Mon 20XX

Abstract—Sentiment analysis of Bangla text on social media platforms like Facebook offers valuable insights into public opinion, sentiment trends, and behavior. In spite of growing interest, challenges such as intricate morphology, script variations, and lack of well-annotated datasets restrict precise analysis. This study is interested in categorizing Bangla Facebook comments into positive, negative, and neutral sentiments using a dataset of 46,655 comments. Various machine learning and deep learning techniques like Logistic Regression, Random Forest, XGBoost, AdaBoost, SVM, and Neural Networks were employed in conjunction with text preprocessing methods like cleaning, tokenization, and feature extraction to enhance model performance. The research highlights the effectiveness of these techniques in managing linguistic complexity and improving prediction accuracy. The findings are of benefit to computational linguistics and social media analysis, offering practical applications to businesses, policymakers, and further research in data-driven decision making. The study highlights the promise of sentiment analysis as a tool for understanding public opinion in Bangla and addressing core challenges in NLP for low-resource language.

Keywords—component, formatting, style, styling, insert (key words)

I. INTRODUCTION

Since individuals have the option to express their thoughts, opinions, and feelings via social media platforms, those represent the majority in the sundry means that could be willing to communicate. Among these social media platforms, Facebook emerges as one of the most popularly used in Bangladesh, with

users always engaged in making conversations and sharing their sentiment through comments. The comments can therefore, to look for, facilitate insightful retrieval of public opinion, general sentiment trends, acceptance levels, etc. into the behavioral patterns. Sentiment analysis, a subsection of natural language processing (NLP), classifies texts into respective broad sentiment categories for ease of annotation, namely: positive, negative, and neutral. This study is set to focus on Bangla sentiment analysis of comments harvested from Facebook. The dataset on which the research is working consists of 46,655 comments and serves to present a substantial ground for answering how user's sentiment in Bangla is treated. Though a rise in interest toward sentiment analysis is a fresh approach to Bangla text, yet very often certain challenges arise due to complex morphology, unavailability of a well-annotated dataset, or so forth with the prominent variations of scripting. Conventional machine-learning models severely hamper performance on those data complexities, inviting an extended range of pre-processing and encoding mechanisms for enhancement of prediction accuracy. The combination of various machine learning and deep learning techniques, They include Logistic Regression, Random Forest, XGBoost, AdaBoost, Support Vector Machines (SVM), and Neural Networks for Bangla Facebook comment classification. The research was based on other techniques like text cleaning, tokenization, and feature extraction to enhance the model performance. The effective development of a sentiment analysis model for Bangla text should therefore add value to the field of computational linguistics and social

media analytics, especially with results useful for businesses and policy issues as well as providing grounds for further research toward public opinion, trends, and better data-driven decision-making.

II. LITERATURE REVIEW

Malliga Subramanian et al.[1] This paper surveys recent advancements in hate speech detection and sentiment analysis using machine learning and deep learning. Challenges include language nuances, context dependency, and evolving slang. Further study is needed to understand cultural context, integrate multimodal data, and develop real-time detection methods. Future research should focus on improving model accuracy through context-aware models, cross-lingual approaches, and multimodal data utilization.

Rezaul Haque et al.[2] The paper introduces a CNN-LSTM (CLSTM) model for multi-class sentiment analysis of Bengali social media comments and compares various ML and DL models. It examines if ML algorithms with preprocessing and feature extraction outperform previous research and if DL methods like LSTM, GRU, BiLSTM, BiGRU, and CLSTM achieve better accuracy, macro-precision, macro-recall, and macro-F1 scores. Prior studies focused on ternary classification with limited datasets and tools, showing no DL models excelling in four sentiment types or using datasets over 40,000 instances. No online tools exist for multi-class sentiment analysis in Bengali.

Kholoud Khalil Aldous et al.[3] This study examines user engagement with ~3,000,000 news postings and ~50,000,000 comments across five social media platforms over eight months. It investigates the impact of sentiments and topics on views, likes, comments, and cross-platform posting, using language and metadata features to predict engagement levels with an 83% maximum average F1-score.

Zuzana Sokolová et al.[4] The paper introduces "SentiSK," a Slovak sentiment analysis dataset, and compares it with existing ones. Using Scikit-learn, various ML models were trained, including Random Forest, MLP, Logistic Regression, SVC, KNN, Multinomial Naive Bayes, Multilingual BERT, and SlovakBERT. The study highlights a lack of Slovak datasets and imbalance in the corpus, with different training methods improving results by up to 24%. Future work should explore embeddings, balanced datasets, automatic annotators for hate speech, a web interface, and cross-language model comparisons using translated datasets.

Lisa M. Gandy et al.[5] This paper evaluates VADER, TEXT2DATA, and LIWC-22 tone for classifying YouTube

comments on the opioid crisis, recommending LIWC-22 tone as the most accurate. Analyzing 8,761 comments from top CNN and Fox News videos (2017-2018), VADER and LIWC_tone classified all comments, while T2D missed about 9%. LIWC-22 tone achieved the highest F1 score compared to manual coding.

Shanmugavadivel et al.[6] The study proposed a deep learning system for sentiment analysis and offensive language identification on Tamil-English code-mixed data, with adapter-BERT performing best. It achieved 65% accuracy in sentiment analysis and 79% in offensive language identification, addressing semantic extraction challenges with word embeddings. The research highlights a gap in studies on low-resource code-mixed data compared to monolingual data.

Babu et al.[7] This paper reviews the use of sentiment analysis on social media data to detect depression and anxiety using various artificial intelligence techniques, finding that multi-class classification with deep learning shows higher precision. Use of social media data (text, emoticons, emojis) for multi-class classification of sentiment polarity using machine learning and deep learning techniques.

Nuzhah Gooda Sahib et al.[8] The authors created a dataset of 1300 Kreol Morisien social media comments and proposed a sentiment analysis framework using SVM and MNB algorithms. SVM achieved 66.15% accuracy, outperforming MNB. They emphasized the need for further

NLP tool development for Kreol Morisien and improving sentiment analysis techniques for handling code-switching in social media comments.

Nurul Hidayah Watimin et al.[9] The paper proposes a method to detect pre-crisis situations by monitoring social media engagement patterns, focusing on post types and sentiment polarity in real-time. It analyzes Facebook posts related to the Mariamman Temple incident in Malaysia (Nov 26-30, 2018), using sentiment analysis to categorize comments. The study highlights challenges such as limited sample size due to Facebook search restrictions and the need for better tools to handle Malay language data. It underscores the importance of advancing machine learning for automated pre-crisis detection and improving data quality through effective filtering and collaboration among stakeholders.

Srabon Bhowmik Shanto et al.[10] This paper used LSTM and GRU deep learning models to detect cyberbullying in Bangla Facebook comments, achieving an accuracy of 83.55% with the GRU model. It included data preprocessing steps such as text cleaning, punctuation removal, tokenization, stop word removal, and stemming before feeding the data into the models for prediction.

Md. Tabil Ahammed et al.[11] The study developed a prototype to analyze attitudes towards women's social difficulties using machine learning on a dataset of Facebook posts. It collected data using a Python-based Facebook scraper, cleaned it with the NLTK toolkit, and applied machine learning

III. Methodology

We have followed a step-by-step process to maintain a structured methodology for this research and every step is described below.

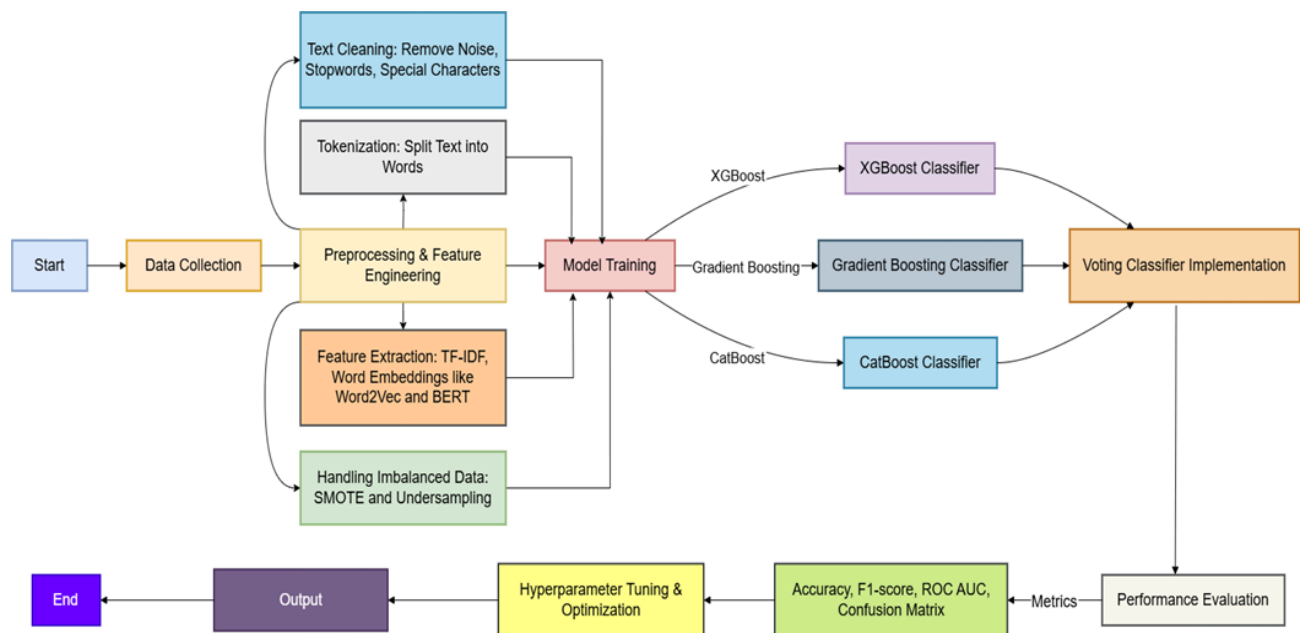


Figure 1: Proposed Methodology

	comment	label
0	ওই হালার পুত এখন কি মদ খাওয়ার সময় রাতের বেলা...	sexual
1	ঘরে বসে স্টুট করতে কেমন লেগেছে? ক্যামেরাতে কে ছি...	not troll
2	অরে বাবা, এইটা কোন পাগল????	troll
3	ক্যাপ্টেন অফ বাংলাদেশ	not troll
4	পটকা মাছ	troll
...	...	...
46649	প্রাপ্তি আপনাদের কাছেই চলে গেছে আপুরা	not troll
46650	জয় আমাদের হবে ই ইনশাআল্লাহ	religious
46651	ধন্যবাদ দিদি আমাদের দেখার সুযোগ করে দেওয়ার জন্য	not troll
46652	নিশা আপু ও ওনার পুরোটিম কে অনেক ধন্যবাদ।	not troll
46653	আসসালামুয়ালাইকুম। সকল আপুরা অনেক ভালো থাকবেন।	religious

Figure 2: Dataset Sample

A. Data Preprocessing

We have collected raw data from Facebook, which has to go through several processing and cleaning steps, such as removing newline characters, Unnecessary symbols, and Bengali stopwords.

Null Value Removal: Comments from the rows where at least one entry is null or missing were captured and removed. This is because this data will not give any meaningful insight for sentiment analysis, but instead ruin the training process. Thus, removing these ensures that the dataset includes only developed, usable data and contributes to more stable and accurate model performance.

Unnecessary Column Removal: In the course of data collection or CSV file generation, it is possible that additional columns would have been added from unnamed index columns, for instance, Unnamed: 2-that were not relevant to the classification task. These columns were dropped out to reduce the clutter of the dataset and avoid any redundancy from interfering with the preprocessing or model training stage.

Text Cleaning: The comments underwent extensive cleaning against noise, permitting only sentiment-relevant content into the analysis. Newline characters that broke up text were discarded along with numerals and all special characters, such as punctuation signs and emojis, having no significance

for the sentiment classification. Non-Bengali characters were deleted as well to retain only linguistically relevant content. Such tasks were executed using regular expressions ('re' module in Python) so that any cleaned text remained in a uniform and understandable format for further processing and analysis.

Stopword Removal: Stopwords are an array of frequently used words in any language that do not contribute weight to a sentence, such as conjunctions, prepositions, or articles. A curated list of Bengali stopwords was used to remove these from the dataset. With stopwords taken out, our model pays attention only to sentiment-bearing words, thus improving the quality of features extracted, which consequently translates into better classification.

Short Text Filtering: Comments that contained no meaningful words or very few tokens were removed from the dataset. These short or empty entries are usually spam, typos, or data entry errors and do not add much to the model training. All such comments were discarded by using a filter condition to retain only those with a token length greater than zero for a more focused and informative dataset

B. Data Analysis

An in-depth analysis of the dataset has been performed to understand the critical aspects for decision-making in designing the model. Exploratory Data Analysis (EDA) is done to explore structural patterns, discover instances of class imbalances, and reveal aspects of the data, which is not limited to but includes the following:

Class Distribution Analysis: A dataset was segmented into five classes of sentiment: Not Troll, Troll, Sexual, Religious, and Threat. The frequency of each class was computed and visualized to discover possible imbalances. An imbalanced dataset in which one class dominates all the other classes yields a biased model training and poor generalization. This analysis helped to see whether class balancing methods like sampling or class weights would be needed.

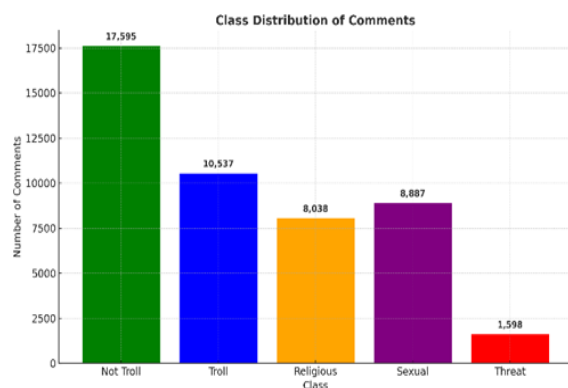


Figure 3: Class Distribution

Word Frequency Analysis: Grasping what words appear most often in the dataset gives a notable understanding of the context-wise aspect of each one of the sentiment classes. The following words are the most frequently used in the dataset based on total frequency:

Top 10 Words by Frequency (Counts):

না (3107), নাস্তিক (2706), কি (2210), এই (1701), ভাই (1590), থেকে (1558), তুই (1478), করে (1428), হিরো (1321), অনেক (1283)

Top 10 Words by Document Occurrence (i.e., Number of Comments):

না (2891), নাস্তিক (2517), কি (2075), এই (1654), ভাই (1536), থেকে (1536), করে (1363), তুই (1334), আল্লাহ (1243), আলম (1229)

Words --> Documents:

না	2891
নাস্তিক	2517
কি	2075
এই	1654
ভাই	1536
থেকে	1536
করে	1363
তুই	1334
আল্লাহ	1243
আলম	1229

Words --> Index:

ওয়ারলেস	27010
মহাখালী	27009
হয়েছেন	27008
এডজাস্ট	27007
ওয়েদারের	27006
ওখানের	27005
হৈম	27004
প্রেক্টিস	27003
সাপ্রিহিজের	27002
নাজমুন্নাহার	27001

Figure 4: Word Frequency

Words --> Documents:

না	2891
নাস্তিক	2517
কি	2075
এই	1654
ভাই	1536
থেকে	1536
করে	1363
তুই	1334
আল্লাহ	1243
আলম	1229

These common terminologies are likely to indicate the frequent presence of emotionally and ideologically loaded language in the dataset, particularly in the comments assigned as Troll, Religious, or Threat. It also helps to identify the words for a better feature selection, thus improving the interpretability of the model outcomes.

Unique Word Count: The dataset is very rich in vocabulary and diverse, with a total of 27,010 unique tokens identified after preprocessing. This vocabulary was generated through the Keras' Tokenizer, which assigned to each word various distinct integers. Such an immense variety of vocabulary reveals itself as lexical diversity, one aspect helpful in training deep-learning models that can capture semantic nuance. Moreover, less frequent words are usually

domain-specific, slang, or rare names and tend to appear at the end of the index list.

Comment Length Analysis: The length in words of each comment was also analyzed. This analysis tapped into the divergent nature of comment lengths, with many counts being short and some being fairly long. To maintain uniformity and computational efficiency, a maximum sequence length of 30 tokens was fixed for padding purposes in the model's input preparation. Longer comments were truncated, empty comments or ones with very few words were deleted during preprocessing. This was also to ensure a standardized set of input dimensions for the neural network models and to get rid of any noise introduced through irrelevant or overly informative inputs.

Summary of Findings: This study of the data presents certain critical attributes necessary for modeling. The data analysis in the sentiment classes mostly reflected the serious imbalance of the Threat category, with standards far under that of others: advanced evaluation metrics may be needed, or maybe even the ensemble or class-weighted classification models, and they would be needed to evaluate the fairness performance across all classes. The word frequency analysis revealed that there are many sentiment-rich terms, which are specific to the context, that were helpful in feature extraction as indicative terms for sentiment interpretation. The dataset was in a very large and heterogeneous vocabulary; it contained 27,010 unique tokens, which all entail a rich linguistic base that highly suits natural language processing affinity with deep learning techniques. To standardize inputs for accommodating both such modeling training and meaningful retention in time, a fixed sequence length of 30 tokens would then be adopted, striking a balance between computational efficiency and completeness of information captured.

	comment	label	cleaned length
0	ওই হাঙ্গার পুত এখন কি মদ খাওয়ার সময় রাতের বেলা...	sexual	ওই হাঙ্গার পুত এখন কি মদ খাওয়ার সময় রাতের বেলা...
1	ঘরে বসে শুট করতে কেমন সেগেছে? ক্যামেরাতে কে ছি...	not troll	ঘরে বসে শুট করতে কেমন সেগেছে ক্যামেরাতে কে ছিলেন
2	আরে বাবা, এই টি কোন পাপল????	troll	আরে বাবা এই টি কোন পাপল
3	ক্যাপ্টেন অফ বাংলাদেশ	not troll	ক্যাপ্টেন অফ বাংলাদেশ
4	পটিকা মাছ	troll	পটিকা মাছ
...	...	...	...
46626	প্রপণি আপনাদের কাছেই চলে গেছে আপুরা	not troll	প্রপণি আপনাদের কাছেই চলে গেছে আপুরা
46626	জয় আমাদের হবে ই ইনশাআল্লাহ	religious	জয় আমাদের হবে ই ইনশাআল্লাহ
46627	বনাবাদ সিদ্দি আমাদের দেখার সুযোগ করে দেওয়ার জন্য	not troll	বনাবাদ সিদ্দি আমাদের দেখার সুযোগ করে দেওয়ার জন্য
46628	নিশা আপু ও ওনার পুরাতীম কে অনেক বনাবাদ	not troll	নিশা আপু ও ওনার পুরাতীম কে অনেক বনাবাদ
46629	আসসালামুয়ালাইকুম । সকল আপুরা অনেক ভালো থাকবেন।	religious	আসসালামুয়ালাইকুম সকল আপুরা অনেক ভালো থাকবেন

46630 rows x 4 columns

figure 5: Cleaned data

C. Model Selection

In this project, various supervised machine learning algorithms were used and compared to find out the best model for categorizing Bengali Facebook comments into sentiment categories. The qualifying criteria for all classifiers include their theoretical foundations, handling capability for multiclass classification, and proven efficacy in text classification works. Here's a brief on the models employed:

Logistic regression is one of the examples of a probabilistic linear classifier, which is widely used to distinguish between two classes or among multiple classes. It gives a probability that does input belongs to any specific class by logic using the logistic (sigmoid) function. In this project, it is also used as a baseline model to test other advanced classifiers. By a linear combination of the features, it calculates the probability of the sample for class membership and transforms it into values ranging from 0 and 1 by using the logistic function. It gets trained by minimizing the error in terms of predicted and actual labels using Maximum Likelihood Estimation (MLE).

Decision trees are built on the recursive partitioning of the dataset concerning feature values in a hierarchy. The internal node represents a test condition on a feature, whereas each leaf node is associated with a predicted class label.

For this project, the Decision Tree was useful in explaining how certain keywords or linguistic features affect sentiment categories. Because they can be prone to overfitting, they are very transparent and can be valuable for the interpretation of model inference.

Naive Bayes is a probabilistic classifier based on Bayes' Theorem with the "naive" simplification that attributes are conditionally independent given the class. Despite this simplification, Naive Bayes works quite well with text classification. It has been employed in this project to provide a rapid and computationally inexpensive baseline with which to compare. It uses the frequency of words in each class to calculate the probability that the comment is in some particular sentiment class.

Random Forest is one of the ensemble learning methods that learn multiple decision trees from the training examples and outputs the mode of classification predictions. The trees are learned from a bootstrap sample of the data and feature selection in each node is one of the sources of randomness that assists in generalizing. Random Forest in this work worked efficiently and was very robust, especially in handling user-generated text variability and noise.

SVM is a robust classifier that seeks to discover the optimal hyperplane to distinguish between the classes in feature space with maximum margins. Despite its use of kernel functions, SVM can classify in high-dimensional space efficiently. For this research, SVM was evaluated in light of its ability to separate between the sentiment classes with good margins, and though computationally more expensive, it yielded similar accuracy.

The MLP is a multi-layer feedforward neural network that learns to model intricate, non-linear relationships between input characteristics and output categories by employing backpropagation. The MLP in this work has been utilized as a deep model capable of learning hidden semantic relationships in text data with dense layers.

XGBoost is a high-performance, optimized gradient-boosting algorithm that builds trees sequentially with each new tree attempting to correct the errors of previously built trees. XGBoost has parallel computation support with efficiency and uses regularization penalties in order to avoid overfitting. XGBoost's ability to handle sparse data and class

imbalance was also one important aspect that necessitated the use of XGBoost in this project.

Gradient Boosting learns by iteratively optimizing a loss function by adding decision trees stage-wise. Unlike Random Forests which build trees in parallel, Gradient Boosting builds trees sequentially to reduce errors. We applied it in this study because it works very well with structured datasets and is flexible in handling complex decision boundaries.

CatBoost is specially designed to cater to the configuration of invoked features of a categorical nature, which makes it a gradient-boosting framework with very good intrinsic performance. In addition to that, it handles multi-class classification as well as handles imbalanced datasets. During this project, CatBoost placed better than other classical classifiers, thus being one of the top models in terms of precision, F1 score, and interpretability.

An ensemble technique known as a Voting Classifier was employed to augment the classification performance. This meta-classifier took the soft votes on the predictions from three of the top classifiers: XGBoost, Gradient Boosting, and CatBoost. The individual models' probabilities are aggregated and the class assigned to the highest mean probability is predicted, thus making the method more robust and generalized across sentiment classes.

In the end, training and evaluating these models were accomplished with consistent data partitioning and preprocessing of each model, and the performance scores were obtained through the commonly used measures of accuracy, precision, recall, F1-score, and ROC-AUC to realize an all-inclusive comparison. The final criterion of model selection depended on the balance among predicted accuracy, computational efficiency, and the capacity to address class imbalance effectively.

iv. RESULT AND DISCUSSION

This section presents a comprehensive evaluation of various machine learning and deep learning models using standard metrics such as training accuracy, test accuracy, precision, recall, F1 score, and ROC AUC score. Among the evaluated models, gradient boosting achieved the highest test accuracy (62.01%) and ROC AUC score (70.04%), making it the most reliable classifier in generalization and class

separation. XGBoost recorded the highest training accuracy (78.76%), which shows a strong learning ability but also indicates overfitting due to relatively low testing accuracy (61.56%). The voting classifier combining XGBoost, gradient boosting, and CatBoost provides a balanced performance with 61.71% test accuracy, 63.73% precision, and a competitive ROC AUC score (69.11%), indicating that ensemble methods can stabilize results across classes.

CatBoost was a close second, although slightly lower in test metrics (58.38% test accuracy, 65.93% ROC AUC), yet the algorithm proved effective in boosting. In contrast, AdaBoost, Random Forest, and Logistic Regression performed significantly less well, with Logistic Regression showing only 43.62% test accuracy and ROC AUC score of 50.52%, revealing its limitations in handling complex, multi-class datasets. MLP (Neural Network) had one of the lowest scores (40.30% test accuracy), with particularly poor accuracy (31.25%) and F1 scores (31.16%), indicating the need for further tuning in architecture and parameters. Naïve Bayes provided the worst results, with both training and testing accuracies exceeding only 34%, confirming its inefficiency on this dataset.

Key observations include the consistent outperformance of boosting algorithms over traditional classifiers and the effectiveness of ensemble methods in enhancing interclass performance. However, challenges such as class imbalance affected the minority class prediction and highlighted the potential use of techniques like SMOTE. Deep learning models, while promising, incurred high computational costs and required significant tuning for optimal performance.

Recommendations for improving model accuracy include implementing advanced feature engineering, adopting K-Fold cross-validation for robustness, and applying hyperparameter optimization to neural networks—such as adjusting learning rates, dropout ratios, and activation functions—to achieve better generalization and precision.

TABLE 1. MODEL ACCURACY

Model	Test Accuracy	Train Accuracy
XGBoost	78.76%	61.56%
Gradient Boosting	72.08%	62.01%
Voting Classifier (XGB + GB + Cat)	74.47%	61.71%
CatBoost	67.17%	58.38%
AdaBoost	50.84%	48.86%
Random Forest	54.66%	47.78%
Logistic Regression	42.75%	43.62%
Naïve Bayes	34.01%	34.16%
MLP	41.66%	40.30%

V. CONCLUSION

In this project, we have developed a system for sentiment analysis of Bangla comments using different machine learning techniques. Among the tested models, the support vector machine (SVM) achieved the highest accuracy. It not only classifies comments as positive or negative but also effectively identifies harmful content such as hate speech and inappropriate language.

Still, there were some challenges. The dataset was small and manually labeled. Bengali is a complex language with many variations, and the system is not yet ready for real-time use. It also struggled to detect certain types of sentiment, such as threats or sexual content.

Looking ahead, we can improve this system by using powerful models like BanglaBERT and collecting more data by making it work in multiple languages to make it more useful in real-world situations.

REFERENCES

- [1] Subramanian, M., Easwaramoorthy Sathiskumar, V., Deepalakshmi, G., Cho, J., & Manikandan, G. (2023). A survey on hate speech detection and sentiment analysis using machine learning and Deep Learning Models. *Alexandria Engineering Journal*, 80, 110–121. <https://doi.org/10.1016/j.aej.2023.08.038>
- [2] Haque, R., Islam, N., Tasneem, M., & Das, A. K. (2023). Multi-class sentiment classification on Bengali social media comments using machine learning. *International Journal of Cognitive Computing in Engineering*, 4, 21–35. <https://doi.org/10.1016/j.ijcce.2023.01.001>
- [3] Aldous, K. K., An, J., & Jansen, B. J. (2022). What really matters?: Characterising and predicting user engagement of news postings using multiple platforms, sentiments and topics. *Behaviour & Information Technology*, 42(5), 545–568. <https://doi.org/10.1080/0144929x.2022.2030798>
- [4] Sokolová, Z., Harahus, M., Juhár, J., Pleva, M., Staš, J., & Hládek, D. (2024). Comparison of machine learning approaches for sentiment analysis in Slovak. *Electronics*, 13(4), 703. <https://doi.org/10.3390/electronics13040703>
- [5] Gandy, L., Ivanitskaya, L. V., Bacon, L., & Bizri-Baryak, R. (2024). An Evaluation of Automated Sentiment Analysis Methods: YouTube Comments on the Opioid Crisis (Preprint). <https://doi.org/10.2196/preprints.57395>
- [6] Shanmugavadivel, K., Sathishkumar, V. E., Raja, S., Lingaiah, T. B., Neelakandan, S., & Subramanian, M. (2022). Deep learning based sentiment analysis and offensive language identification on multilingual code-mixed data. *Scientific Reports*, 12(1). <https://doi.org/10.1038/s41598-022-26092-3>
- [7] Babu, N. V., & Kanaga, E. G. (2021). Sentiment analysis in social media data for Depression Detection Using Artificial Intelligence: A Review. *SN Computer Science*, 3(1). <https://doi.org/10.1007/s42979-021-00958-1>
- [8] Sahib, N. G., Marianne, M. A., & Gobin-Rahimbux, B. (2023). Sentiment analysis of social media comments in Mauritius. 2023 IEEE 13th Annual Computing and Communication Workshop and Conference (CCWC). <https://doi.org/10.1109/ccwc57344.2023.10099291>
- [9] Watimin, N. H., Zauddin, H., & Rahamad, M. S. (2023). Religious and racial tension breakout: An online pre-crisis detection strategy via sentiment analysis for Riot Crime Prevention. *Social Network Analysis and Mining*, 13(1). <https://doi.org/10.1007/s13278-023-01086-9>
- [10] Shanto, S. B., Islam, M. J., & Samad, Md. A. (2023). Cyberbullying detection using Deep Learning techniques on Bangla facebook comments. 2023 International Conference on Intelligent Systems, Advanced Computing and Communication (ISACC). <https://doi.org/10.1109/isacc56298.2023.10083690>
- [11] Ahammed, Md. T., Gloria, A., J, S. D., Oion, Md. S., Ghosh, S., Balaii, P., & Nisat, T. (2022). Sentiment analysis using a machine learning approach in Python. 2022 International Conference on Communication, Computing and Internet of Things (IC3IoT). <https://doi.org/10.1109/ic3iot53935.2022.9768004>

