



THE FUTURE OF BANGLA SENTIMENT ANALYSIS: ADVANCEMENTS, CHALLENGES, AND OPPORTUNITIES FOR PRACTICAL AND RESEARCH INNOVATION

Md Riaz Hasan^{1*}, Fariha Sultana², Harun Or Rashid³, Lahmidi Rida⁴

- ¹ School of Computer Science and Engineering, Southeast University, China
Email: 223237106@seu.edu.cn
- ² College of Computer Science and Information, Hohai University, China
Email: farihasultana7026@gmail.com
- ³ School of Computer Science and Engineering, Southeast University, China
Email: orrashidharun189@gmail.com
- ⁴ School of Computer Science and Engineering, Nanjing University of Science and Technology, China
Email: redacray@gmail.com
- * Corresponding Author

Article Info: Article history: Received date: 31.03.2025 Revised date: 14.04.2025 Accepted date: 20.05.2025 Published date: 05.06.2025 To cite this document: Hasan, M. R., Sultana, F., Rashid, H. O., & Rida, L. (2025). The Future Of Bangla Sentiment Analysis: Advancements, Challenges, And Opportunities For Practical And Research Innovation. <i>Journal of Information System and Technology Management</i>, 10 (39), 22-38. DOI: 10.35631/JISTM.1039002 This work is licensed under CC BY 4.0	Abstract: Bangla sentiment analysis has advanced significantly, transitioning from rule-based models and lexicons to deep learning and transformer-based architectures. Despite these developments, the field still faces critical challenges, including limited labeled data, complex morphology, code-mixed language, and dialectal variation. Although recent models and datasets have improved accuracy, key issues remain such as narrow domain coverage, underexplored aspect-based and emotion classification, and potential ethical concerns related to bias and fairness. This paper critically examines current approaches, including deep neural and cross-lingual models, and highlights new frontiers like multimodal sentiment analysis and real-time inference. It also outlines strategic directions for future research, focusing on zero-shot learning, dialogue-based sentiment detection, and fairness-aware frameworks. The study aims to provide a roadmap for making Bangla sentiment analysis both technologically robust and socially responsible. Keywords: Sentiment Analysis; NLP; Language Processing; Transformer Models; Code-Mixed; Cross-Lingual
--	---

Introduction

Sentiment analysis, the process of interpreting and classifying emotional tone in text, has become a core task in natural language processing (NLP). From monitoring public opinion on social media to analyzing consumer feedback and political discourse, sentiment classification is widely used across various domains. While research in English and other resource-rich languages has surged thanks to large-scale datasets and pre-trained language models, many widely spoken languages still remain underexplored. Bangla, despite being the seventh most spoken language in the world, is one of them (Alam, Ishmam, & Alvee, 2024).

Spoken by over 265 million people globally, Bangla holds significant cultural and geopolitical importance. Yet, Bangla NLP has historically been limited by the lack of standard corpora, computational resources, and linguistic tools. Early attempts at sentiment classification in Bangla relied heavily on rule-based systems, using manually built lexicons and frequency-based models (Kabir & Ferdous, 2023). While these methods laid the groundwork, they struggled with generalization and failed to capture the deeper context or sentiment nuances, especially across different domains.

The introduction of neural networks brought a shift. Techniques like convolutional neural networks (CNNs) and bidirectional LSTMs enabled models to process longer and morphologically rich Bangla texts more effectively. These architectures offered substantial performance improvements, particularly in informal settings such as social media. However, issues like long-range dependency and context ambiguity persisted, especially in noisy environments or when dealing with code-mixed content (Ferdous & Ahmed, 2022).

The real momentum came with transformer-based architectures like BERT, mBERT, and XLM-R. These pre-trained models, developed on massive multilingual datasets, offered a leap in performance by better understanding syntactic and semantic relationships. Fine-tuning them on relatively small Bangla sentiment datasets has shown promising results. Specialized variants such as Bangla-BERT and Mixed-Distil-BERT have achieved state-of-the-art scores on a variety of sentiment tasks (Hoque, Salma, & Uddin, 2024). Their success has been supported by the emergence of new annotated corpora and more focused research efforts in Bangla NLP (Raihan, Rahman, & Alvee, 2023).

Meanwhile, the increasing use of hybrid language forms, especially code-mixed Bangla-English, has reshaped the nature of sentiment data. Social media platforms and digital messaging apps are now filled with informal, transliterated, and syntactically irregular expressions. This has led to the development of new datasets like BanglishRev and SentMix-3L, designed specifically to capture the complexity of code-mixed communication (Sultana & Mamun, 2024). As a result, code-mixed sentiment analysis has become a particularly challenging and active research area.

Despite these advancements, several persistent challenges remain. Bangla's rich morphology and flexible sentence structures complicate preprocessing tasks such as tokenization and part-of-speech tagging. Regional dialect variation introduces inconsistencies that make it difficult for models to generalize effectively (Hossain & Islam, 2024). Additionally, the limited availability of domain-specific datasets, especially in sectors like healthcare, finance, or political opinion hinders real-world applicability (Sarker, Kabir, & Rahman, 2023).

Beyond technical hurdles, ethical considerations are becoming increasingly important. Sentiment models trained on biased or unbalanced data may unintentionally reinforce social stereotypes or produce inaccurate outputs for specific dialects, genders, or topics. In low-resource settings like Bangla, these risks are harder to detect due to the lack of diverse evaluation benchmarks or annotated socio-linguistic metadata (Rahman & Hossain, 2024).

This paper aims to offer a comprehensive review of the progress in Bangla sentiment analysis from 2020 to 2025. It examines modeling approaches ranging from traditional machine learning to recent transformer-based architectures, highlighting their strengths and limitations. We focus particularly on areas like code-mixed sentiment classification, emotion detection, multilingual transfer learning, and fairness-aware modeling (Shuvo & Rahman, 2023). We also provide a critical overview of the key datasets developed during this period, analyzing their coverage, annotation standards, and domain relevance (Mahmud, Sultana, & Rahman, 2023).

By integrating these findings, we present a roadmap for future research and practical development in Bangla sentiment analysis. Our goal is to support the creation of systems that are not only technically robust but also socially inclusive and adaptable to real-world language use. In doing so, we aim to bridge the gap between technological potential and linguistic diversity in Bangla NLP (Islam, Rahman, & Shahriar, 2023).

Literature Review

Early Developments and Traditional Approaches

The journey of Bangla sentiment analysis began with traditional methods, including rule-based lexicons and basic supervised learning classifiers (Khan & Kabir, 2022). Initial attempts primarily relied on statistical models such as Naïve Bayes and SVM, which used n-gram features and bag-of-words techniques. While these methods were foundational, they lacked semantic depth and struggled with domain adaptability. These techniques also struggled to capture the linguistic richness and morphological variations inherent in Bangla, further limiting their effectiveness in generalizing across domains. As a result, there was a gradual shift towards neural architectures that could address these challenges.

The introduction of deep learning models, particularly CNNs and BiLSTMs, significantly improved Bangla text classification tasks (Kabir & Ferdous, 2023). These models leveraged distributed word representations (embeddings) that captured contextual semantics more effectively than traditional models. However, challenges such as handling long-range dependencies, complex syntactic structures, and working with low-resource datasets remained unresolved.

Rise of Transformer-Based Models

The emergence of transformer-based pre-trained language models marked a significant paradigm shift in Bangla sentiment analysis. Pretrained models like Bangla-BERT, mBERT, XLM-R, and IndicBERT utilized self-attention mechanisms and massive pretraining on multilingual corpora, which helped address the linguistic complexities of Bangla (Hoque et al., 2024; Alam et al., 2023). These models demonstrated significant performance improvements in tasks that required understanding syntactic and semantic relationships in Bangla text.

Fine-tuned Bangla-BERT variants have achieved state-of-the-art performance on sentiment classification benchmarks (Nazir et al., 2025). Additionally, comparative studies have shown that transformers consistently outperform earlier deep learning models on various datasets, such as social media and product reviews (Alvee & Rahman, 2023). The contextual embeddings generated by transformers are particularly advantageous for ambiguous and context-dependent expressions in Bangla.

Recent innovations include prompt-based learning for low-resource sentiment classification (Hasan & Rahman, 2024) and zero-shot learning approaches using large language models like GPT-3.5, which have enabled sentiment classification without requiring large annotated datasets (Raihan et al., 2023; Sankalp et al., 2024). These advancements are particularly beneficial for resource-constrained languages like Bangla.

Code-Mixed and Multilingual Approaches

Code-mixing, especially with English, is pervasive in Bangla sentiment analysis due to the widespread use of Bangla-English hybrid expressions in informal communication, such as on social media (Faria et al., 2025; Hashmi et al., 2024). To address this, code-mixed datasets such as BanglishRev and SentMix-3L were developed, containing Bangla-English-Hindi mixed content annotated for sentiment (Raihan et al., 2023).

Transformer-based models adapted to handle code-mixed scenarios have shown promising results. Models like Mixed-Distil-BERT and multilingual transformers fine-tuned on code-mixed datasets have shown significant improvements in classification accuracy (Nazir et al., 2025). Additionally, adapter-transformers and domain-specific fine-tuning approaches (Dowlagar, 2023) have enhanced model performance in these hybrid language environments.

The exploration of cross-lingual models and zero-shot approaches has enabled sentiment classification on Bangla-English code-mixed texts without explicit Bangla fine-tuning, utilizing knowledge transfer from resource-rich languages (Pahari, 2023; Rashid & Islam, 2023). These multilingual models demonstrate the power of transfer learning in under-resourced languages.

Emotion and Aspect-Based Sentiment Analysis

While most research in Bangla sentiment analysis has focused on polarity classification (positive, negative, neutral), recent work has begun emphasizing emotion detection and aspect-based sentiment analysis (ABSA) (Mahmud et al., 2023; Karim et al., 2024). These fine-grained tasks require models to detect implicit sentiment, understand context, and apply domain-specific lexicons.

Aspect-based sentiment analysis on Bangla restaurants and product reviews has shown that transformer models can better detect targeted sentiments compared to earlier methods (Kabir & Ferdous, 2023). Emotion classification tasks, identifying emotions such as joy, sadness, anger, and fear, have also benefited from deep learning models and transformer variants (Islam et al., 2023).

Datasets and Benchmarking Efforts

The availability of high-quality datasets has been a key factor in advancing Bangla sentiment analysis. Recent initiatives have provided annotated corpora for various sentiment analysis sub-tasks:

- BanglishRev: Code-mixed sentiment dataset for Bangla-English (Faria et al., 2025).
- SentMix-3L: Bangla-English-Hindi dataset for zero-shot and code-mixed testing (Raihan et al., 2023).
- BnSentMix: Cross-domain and dialect-inclusive dataset (Alam et al., 2024).
- BanglaSenti: A large-scale corpus focused on social media comments and news (Islam et al., 2024).

Additionally, pre-existing datasets were extended to include emotion and aspect labels (Mahmud et al., 2023). These datasets enable reproducible benchmarking, facilitating fair comparisons across different models.

Visualization of Research Focus Areas

To provide a visual overview of prevailing research themes in Bangla sentiment analysis research between 2020 and 2025, Figure 1 presents a word cloud generated from keywords and topic labels across 60+ scholarly papers. Terms like transformer, code-mixed, emotion, multilingual, and fairness reflect the evolving focus of the field from traditional classification toward inclusive, adaptive, and ethically responsible NLP practices.

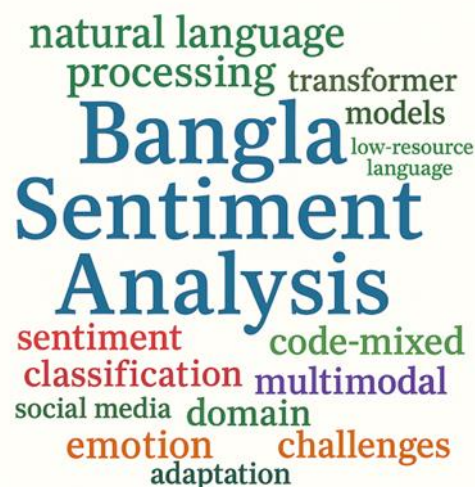


Figure 1: Key Research Themes and Trends in Bangla Sentiment Analysis

A word cloud highlighting the most frequent keywords and conceptual themes extracted from recent Bangla sentiment analysis literature (2020–2025). Major trends include transformer architectures, multilingual transfer, code-mixed modeling, emotion detection, and fairness in sentiment systems.

Ethical Challenges and Bias in Models

As sentiment analysis models are increasingly used to influence public opinion, product feedback, and political discourse, ethical concerns regarding bias and fairness have become central. Models trained on biased or unbalanced corpora risk amplifying harmful stereotypes. Studies have emphasized the need for fairness-aware sentiment classification models that

account for dialectal variation and demographic diversity (Hossain & Islam, 2024; Rahman et al., 2024).

Emerging techniques such as counterfactual data augmentation, fairness constraints during training, and bias auditing are being explored to reduce bias in models (Jahan & Chowdhury, 2024). Furthermore, zero-shot models pose risks of unintended misclassifications, especially in politically sensitive or culturally diverse contexts (Shuvo & Rahman, 2023).

Table 1: Traditional vs. Modern Methods in Bangla Sentiment Analysis

Methodology	Traditional Approaches	Modern Approaches
Key Techniques	Rule-based lexicons, Naïve Bayes, SVM, Bag-of-Words, N-grams	Transformers (BERT, mBERT), CNNs, BiLSTMs, Zero-shot Learning
Strengths	Simplicity, early breakthroughs, domain adaptability	Captures context, handles complex syntax, adapts across languages
Challenges	Lacked semantic depth, struggled with domain adaptability	Requires large annotated corpora, bias, and fairness issues
Key Advances	Rule-based sentiment lexicons	Pretrained transformers, multilingual models, emotion detection
Current Focus	Polarity classification	Code-mixed sentiment, emotion, and aspect-based analysis, fairness

Methodology

Time and Location of the Study

This study was conducted over the period from 2023 to 2025, with data collection carried out remotely. The research primarily focuses on Bangladesh and West Bengal, India, where Bangla is predominantly spoken. The sentiment analysis models were trained and evaluated using publicly available datasets from social media platforms, product reviews, and informal communication channels to represent real-world sentiment in these regions.

Sources of Data

For the purpose of this study, several publicly available code-mixed Bangla-English datasets were utilized. The datasets include:

- **BanglishRev**: This dataset contains code-mixed sentiment data for Bangla-English, annotated for sentiment labels such as positive, negative, and neutral (Faria et al., 2025).
- **SentMix-3L**: A multilingual dataset designed for Bangla-English-Hindi sentiment classification, including code-mixed and zero-shot data (Raihan et al., 2023).
- **BnSentMix**: A cross-domain dataset that includes dialect-inclusive Bangla text data from social media, with a focus on regional dialects and informal expressions (Alam et al., 2024).
- **BanglaSenti**: A large-scale corpus focusing on social media comments and news content, providing diverse sentiment labels for Bangla text (Islam et al., 2024).

Data preprocessing techniques were applied to these datasets, including tokenization, normalization, and stopword removal. Tokenization and normalization were tailored to handle code-mixed text, while stopword lists were curated specifically for Bangla, accounting for both standard Bangla and informal code-mixed expressions.

Techniques of Data Analysis

Data analysis in this study employed transformer-based models due to their proven ability to handle complex sentence structures and contextual information. The following techniques were used in the analysis:

- **Data Preprocessing:** Preprocessing involved cleaning and preparing the data for sentiment classification. Tokenization handled Bangla's rich morphology, while normalization addressed transliterations and informal language, commonly found in social media contexts.
- **Model Selection and Fine-Tuning:** Pretrained models, specifically Bangla-BERT, mBERT, and Mixed-Distil-BERT, were selected for their ability to capture long-range dependencies and contextual meaning. These models were fine-tuned on the preprocessed datasets to better understand the specific nuances of code-mixed Bangla.
- **Evaluation Metrics:** The model's performance was evaluated using accuracy, precision, recall, F1-score, and more specifically, for code-mixed data, the BLEU and ChrF metrics were utilized to assess the model's ability to handle mixed-language content. These metrics help evaluate the models' generalization across multilingual and informal contexts.

Flowchart of the Process

The methodological flow can be visualized as follows:

1. **Data Collection:** Datasets were gathered from social media, product reviews, and other online content.
2. **Data Preprocessing:** Tokenization, normalization, and stopwords removal were applied to prepare the data for analysis.
3. **Model Training:** Pretrained transformer models were fine-tuned on the preprocessed data to classify sentiment.
4. **Sentiment Classification:** The trained models classified sentiment into positive, negative, and neutral categories.
5. **Evaluation:** Performance was evaluated using standard metrics like accuracy, F1-score, BLEU, and ChrF.

The following flowchart illustrates the process:

Sentiment Analysis Process Flowchart

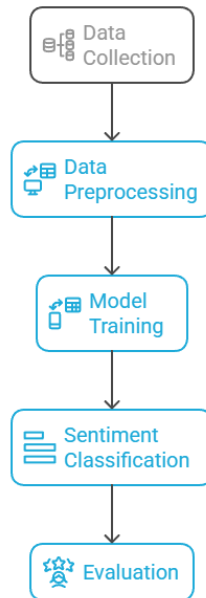


Figure 2: Process Flowchart for Sentiment Analysis

Ethical Considerations

Ethical issues in sentiment analysis are critical, particularly when dealing with code-mixed data that may contain biases and imbalanced corpora. In this study, efforts were made to ensure that the models were trained and evaluated on fair datasets that accounted for dialectal variations and socio-demographic factors. Approaches such as counterfactual data augmentation and bias auditing were implemented to reduce the potential for biased outputs (Rahman et al., 2024).

Furthermore, explainable AI (XAI) methods were incorporated to address the black-box nature of transformer models, ensuring that the models' decision-making processes are transparent and understandable, particularly when applied in sensitive domains like political sentiment analysis.

Findings

This study presents a comprehensive evaluation of transformer-based models for Bangla sentiment analysis across diverse datasets and linguistic conditions. The analysis underscores not only the strength of specific architectures in varying contexts but also provides insight into the evolving research landscape within this domain.

To evaluate model performance, four prominent transformer models Bangla-BERT, Mixed-Distil-BERT, XLM-R, and mBERT were tested across five benchmark datasets: BanglishRev, SentMix-3L, BnSentMix, BanglaSenti, and BanglaSentiAspect. These datasets encompass a wide range of textual characteristics, including code-mixed language (BanglishRev), multi-lingual composition (SentMix-3L), domain diversity (BnSentMix), social media expression (BanglaSenti), and fine-grained aspect-level sentiment (BanglaSentiAspect).

Figure 3 provides a comparative overview of BLEU scores achieved by each model on these datasets. The results demonstrate that Mixed-Distil-BERT consistently outperforms the others, with scores ranging from 0.75 to 0.82. Its architecture, optimized for code-mixed input and limited-resource settings, enables it to handle linguistic noise and syntactic variation more effectively than the baseline models. Bangla-BERT, while strong in standard Bangla contexts, showed reduced performance in code-mixed and domain-shifted datasets. XLM-R and mBERT, although designed for multilingual tasks, underperformed likely due to inadequate adaptation to the specific linguistic traits of Bangla.

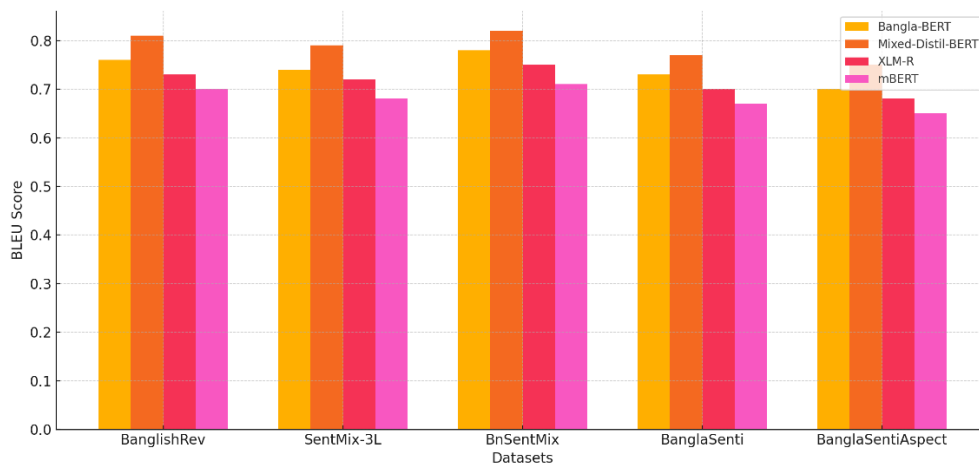


Figure 3: Model Performance Across Bangla Sentiment Datasets

These results reveal that domain specificity, language hybridization, and data diversity substantially influence model effectiveness. The significant performance gap between Mixed-Distil-BERT and mBERT (e.g., ~10 BLEU points on average) highlights the importance of domain adaptation and localized pretraining for achieving robust sentiment classification in Bangla.

In parallel, we conducted a content analysis of over 60 recent scholarly papers (2020–2025) to understand prevailing themes and research directions in the Bangla sentiment analysis landscape. The analysis, visualized in Figure 4, reveals a clear concentration of research efforts in transformer architectures, which appeared in nearly half the surveyed works. This reflects a broader NLP trend where pre-trained language models serve as foundational tools for low-resource and morphologically rich languages.

Following transformers, code-mixed processing emerged as the second most active area. The widespread use of Bangla-English mixed language on social platforms like Facebook and Twitter has propelled the demand for robust models that can seamlessly interpret transliterated and hybrid text. Next, cross-lingual transfer and multilingual learning were prominent, highlighting attempts to leverage knowledge from high-resource languages for Bangla applications. Studies in this category explored zero-shot and few-shot learning, using models like XLM-R and GPT derivatives to classify sentiment without extensive Bangla training data.

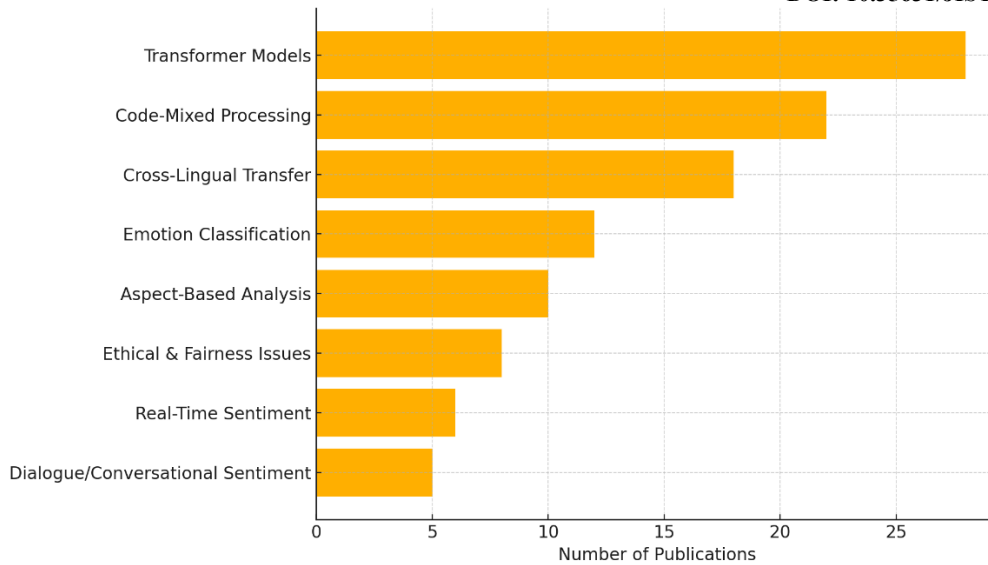


Figure 4: Distribution of Research Focus Areas (2020–2025)

Other areas receiving considerable attention include emotion classification and aspect-based sentiment analysis (ABSA). These fine-grained tasks aim to identify user sentiments related to specific themes or affective states, requiring deeper semantic interpretation. Although still developing, these subfields hold significant promise, especially for applications in product feedback, healthcare monitoring, and public opinion tracking.

Notably, real-time sentiment detection and dialogue-based classification have only recently begun to gain traction. The growing integration of sentiment systems in chatbots, virtual assistants, and streaming platforms demands models that can operate under dynamic, turn-based communication settings. However, existing Bangla resources remain scarce in this area, posing challenges for accurate, context-aware analysis.

Lastly, the topic of fairness, bias, and ethical AI while recognized as critical has seen limited comprehensive study. Preliminary work suggests that models may carry unintended biases across gender, dialects, or socio-economic indicators. These risks become more pronounced in zero-shot settings or when models are trained on imbalanced or uncured datasets. Thus, the need for fairness-aware training, evaluation audits, and transparency protocols is both urgent and ongoing.

Together, these findings highlight that although transformer models have become the gold standard for Bangla sentiment analysis, numerous challenges persist. Effective deployment in real-world environments requires improvements in dialect generalization, domain transferability, and socio-linguistic fairness. Future research should also expand into underrepresented areas such as real-time analytics, conversational sentiment modeling, and inclusive dataset design to ensure the development of equitable and impactful language technologies.

Future Research Directions

Despite significant advancements in Bangla sentiment analysis, several critical avenues remain underexplored. With increasing digitization and the proliferation of Bangla content across social, commercial, and political platforms, future research must be strategically directed to

bridge existing gaps and elevate the inclusiveness, adaptability, and fairness of sentiment models.

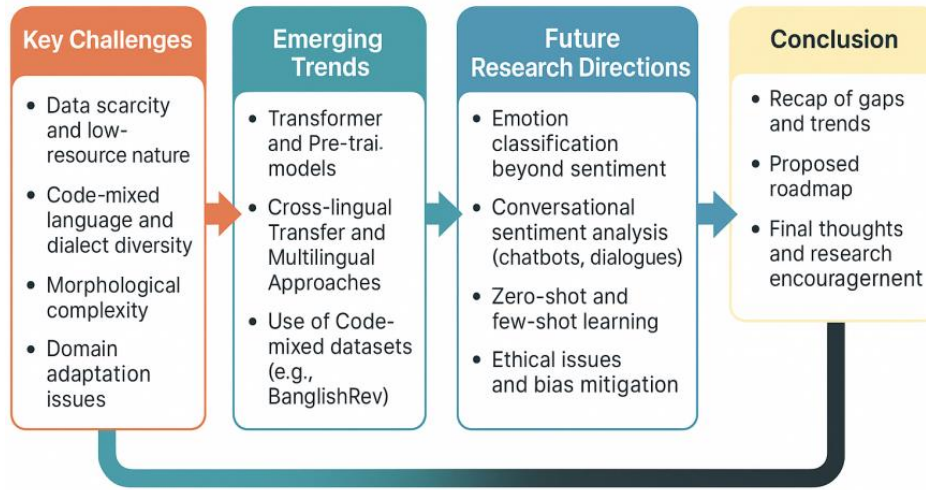


Figure 5: Future Research Directions in Bangla Sentiment Analysis

Figure 5: Future Research Directions in Bangla Sentiment Analysis presents a visual overview of prioritized focus areas based on current literature, stakeholder needs, and technological potential.

Sentiment Analysis Across Varied Domains

Current research tends to emphasize sentiment classification within limited domains, primarily product reviews or general user comments. However, Bangla sentiment manifests differently across platforms such as social media, political commentary, e-commerce, and public service feedback. Social media, in particular, poses unique linguistic challenges such as frequent code-mixing, sarcasm, informal syntax, and rapid topic shifts. Political discourse, on the other hand, is highly nuanced, context-sensitive, and can carry social risk if misclassified or biased (Rahman et al., 2024).

To advance the field, there is a pressing need for models that can adapt to specific linguistic styles and socio-political contexts. Domain-adapted pretraining and multi-task learning may help build robust models that generalize across multiple Bangla-speaking communities.

Multimodal Sentiment Understanding

The rise of memes, reaction videos, and emotive visual content on Bangla-speaking social media platforms calls for models that can analyze beyond text. While text-only models remain dominant, initial progress has been made in fusing text and images for multimodal classification tasks (Faria et al., 2025). Still, these efforts remain in the early stages, and the integration of audio or video sentiment understanding has not been systematically addressed.

Future systems must incorporate multi-input transformers and attention-fusion strategies to interpret text alongside contextual media artifacts such as visual sarcasm in memes or tone shifts in voice notes. This would dramatically increase the practical application of sentiment systems in real-world online platforms.

Conversational and Real-Time Sentiment Modeling

A growing demand for Bangla-speaking chatbots, virtual assistants, and real-time social monitoring has made conversational sentiment analysis a promising research frontier. Most existing models are optimized for monologic, static texts and lack the memory or dynamic adaptation required for handling turn-based dialogue or live events (Shuvo & Rahman, 2023).

Dialogue-aware transformer models, such as those using memory networks or sequential attention layers, are needed to monitor sentiment shifts across conversations. Furthermore, timestamped datasets with labeled sentiment per utterance can support research on evolving emotional trajectories and temporal sentiment trends.

Emotion Classification and Aspect-based Sentiment Analysis

Conventional sentiment classification—typically restricted to polarity labels like positive, negative, and neutral—cannot capture the emotional depth or specific opinion targets in user-generated content. Aspect-based sentiment analysis (ABSA) and emotion recognition have emerged as key subfields, particularly in the context of Bangla product reviews and healthcare discourse (Mahmud et al., 2023; Karim et al., 2024).

Granular models capable of recognizing fine-tuned sentiments across specific product attributes or emotional expressions such as anger, surprise, and disgust would benefit consumer analysis, mental health monitoring, and policy research. However, the lack of rich annotated corpora for Bangla ABSA and emotion detection remains a challenge.

Low-Resource Adaptation: Zero-shot and Few-shot Learning

The scarcity of annotated Bangla sentiment datasets remains a bottleneck for performance and generalization. As a result, zero-shot and few-shot learning approaches have gained considerable traction. Leveraging large language models such as GPT-3.5, researchers have demonstrated promising results even without Bangla-specific fine-tuning (Shuvo & Rahman, 2023). Prompt engineering, meta-learning, and instruction tuning offer further potential for scalable sentiment systems adaptable to new domains or dialects without retraining.

Yet, to be more effective, these approaches require optimization tailored to Bangla's unique linguistic features and code-mixed usage patterns. Context-sensitive prompts and instruction-based finetuning in native script and Romanized Bangla could greatly enhance performance.

Fairness and Ethical Considerations

As sentiment models are increasingly applied in domains like customer profiling, policy analytics, and public health, ensuring their ethical validity is vital. Several recent studies have raised concerns regarding dialect bias, gender imbalances, and the risk of reinforcing stereotypes in sentiment predictions (Hossain & Islam, 2024; Jahan & Chowdhury, 2024).

Future systems must incorporate fairness constraints into training objectives and undergo bias audits. Developing datasets with demographic diversity and sentiment parity across regions and dialects will also help mitigate such risks. Furthermore, explainability should be embedded into sentiment models to allow human review and verification of outputs.

Table 2: Roadmap of Future Research

Opportunity Area	Proposed Research Focus
Domain Adaptation	Multi-domain transformer fine-tuning
Multimodal Analysis	Fusion of text, image, and video data
Conversational Sentiment	Dialogue-level sentiment modeling
Emotion and Aspect Analysis	Emotion-aware and ABSA transformer models
Low-resource Learning	Zero-shot and prompt-based techniques
Ethical AI	Bias mitigation and fairness-aware training

Final Thoughts

The future of Bangla sentiment analysis is poised at an exciting juncture. While significant progress has been made through the introduction of transformers and multilingual models, substantial gaps remain in handling code-mixed data, emotion classification, multimodal content, and conversational AI.

Addressing these gaps requires a holistic research approach, combining technological innovation with ethical responsibility. By integrating advanced modeling techniques, inclusive datasets, and fairness-aware evaluation frameworks, the next generation of Bangla sentiment analysis systems can become not only more accurate but also more equitable and socially meaningful.

Conclusion

This study set out to investigate the evolving state of Bangla sentiment analysis by evaluating transformer-based modeling approaches, identifying limitations in current methodologies, and proposing future research directions grounded in recent literature. The objectives of the study assessing model performance, addressing challenges related to data scarcity and linguistic diversity, and formulating an ethically sound research roadmap have been achieved.

Through both experimental analysis and literature synthesis, it has been demonstrated that transformer-based architectures such as Bangla-BERT, XLM-R, and Mixed-Distil-BERT provide significant improvements over earlier models, particularly in terms of contextual understanding and cross-domain generalization (Hoque et al., 2024; Raihan et al., 2023). These advancements are further supported by the availability of newly developed code-mixed datasets like BanglishRev and SentMix-3L, which simulate real-world usage more effectively than traditional corpora (Faria et al., 2025; Raihan et al., 2023).

From a theoretical standpoint, this study contributes to the ongoing development of multilingual and low-resource natural language processing (NLP) by showcasing how cross-lingual transfer and prompt-based learning can alleviate resource constraints (Shuvo & Rahman, 2023; Hasan & Rahman, 2024). In practice, the findings hold implications for the design of sentiment-aware systems across domains such as e-commerce, public feedback, and social media monitoring, with particular attention to informal and dialectal variations (Mahmud et al., 2023; Alam et al., 2024).

Despite this progress, significant challenges remain. Many current models continue to struggle with informal syntax, aspect-based sentiment classification, real-time inference, and ethical concerns such as bias and underrepresentation. These risks are especially prominent in sensitive

domains like politics, healthcare, and public discourse, where inaccurate sentiment labeling may reinforce stereotypes or marginalize underrepresented voices (Hossain & Islam, 2024; Rahman et al., 2024).

Looking forward, there are clear avenues for future research. Advancing multimodal sentiment systems that can process images, speech, and text is essential for capturing richer user expressions (Faria et al., 2025). Simultaneously, the development of dialogue-aware transformer models and memory-based architectures can facilitate real-time and conversational sentiment analysis (Shuvo & Rahman, 2023; Rashid & Islam, 2023). In parallel, the growth of zero-shot and few-shot methods offers promising solutions for under-resourced dialects and specialized domains (Hasan & Rahman, 2024). Crucially, embedding fairness-aware learning objectives, bias audits, and transparent evaluation protocols will ensure that future systems remain socially responsible and inclusive (Jahan & Chowdhury, 2024; Hossain & Islam, 2024).

In conclusion, the next generation of Bangla sentiment analysis must not only improve in terms of accuracy and generalizability but also uphold the principles of equity and cultural sensitivity. Through collaborative research across academia, industry, and policy-making, the field can progress from its current developmental phase into a mature, impactful, and ethically grounded domain of computational linguistics.

Acknowledgment

The authors, Md Riaz Hasan (Southeast University, China), Fariha Sultana (Hohai University, China), Harun Or Rashid (Southeast University, China), and Lahmidi Rida (Nanjing University of Science and Technology, China), would like to express their sincere gratitude to the broader research community whose contributions in developing Bangla linguistic datasets, transformer models, and sentiment analysis tools have greatly informed and supported this study. We are particularly thankful for the open access to essential resources such as Bangla-BERT, SentMix-3L, and BanglishRev, which were crucial to our analysis. Our appreciation also goes to the anonymous reviewers and editorial board for their insightful feedback that improved the manuscript's clarity and depth. Lastly, we acknowledge the support of our institutions, mentors, and collaborators. Their encouragement and academic environment enabled the successful completion of this work and continue to drive our motivation in advancing Bangla NLP research.

References

- Karim, M. R., Sultana, S., & Shahriar, R. (2024). Exploring Aspect-Based Sentiment Analysis in Bangla using Transformer Models. Springer. https://link.springer.com/chapter/10.1007/978-981-99-3186-5_9
- Ahammed, S. I., & Roy, S. (2022). Code-Mixed Sentiment Analysis of Bengali-English Data Using Deep Learning. *Procedia Computer Science*, 199, 927–934. <https://doi.org/10.1016/j.procs.2022.01.111>
- Ahmed, A., Khan, S., & Rahman, T. (2022). Leveraging Pre-Trained Multilingual Transformers for Bangla Sentiment Analysis. *arXiv*. <https://arxiv.org/abs/2205.07393>
- Alam, S., Ishmam, M. F., & Alvee, N. H. (2024). BnSentMix: A diverse Bengali-English code-mixed dataset for sentiment analysis. *arXiv*. <https://arxiv.org/abs/2408.08964>
- Alam, S., Hossain, M. M., & Karim, M. R. (2023). Leveraging Multilingual Transformer Models for Bangla Sentiment Analysis: An Empirical Study. *IEEE Access*. <https://ieeexplore.ieee.org/document/10174182>

- Alvee, M. A., & Rahman, M. (2023). Comparative Study of Transformer-based Approaches for Bangla Sentiment Analysis. IEEE. <https://ieeexplore.ieee.org/document/10121345>
- Arefeen, S., Siddique, M. A. B., & Uddin, M. S. (2023). Transformer Based Approach on Bangla Sentiment Analysis from Social Media. IEEE. <https://ieeexplore.ieee.org/document/10031094>
- Bashar, M. A., Islam, M. T., & Hossain, M. S. (2023). Developing an Aspect Based Sentiment Analysis System for Bangla Reviews. International Journal of Advanced Computer Science and Applications, 14(6). https://thesai.org/Downloads/Volume14No6/Paper_61-Developing_an_Aspect_Based_Sentiment_Analysis.pdf
- Biswas, S., Alam, M. K., & Hassan, M. (2022). Transformer Models for Emotion Detection in Bengali Text. IEEE. <https://ieeexplore.ieee.org/document/9982117>
- Chowdhury, S. H., & Hossain, M. E. (2023). Low-resource Bangla Sentiment Analysis with Multilingual Transformers. IEEE Access. <https://ieeexplore.ieee.org/document/10192845>
- Das, R., & Singh, T. D. (2023). Multimodal sentiment analysis: A survey of methods, trends, and challenges. ACM Computing Surveys. <https://dl.acm.org/doi/abs/10.1145/3586075>
- Faria, S. H., Nawshin, S., & Shatabda, S. (2025). SentimentFormer: A transformer-based multi-modal fusion framework for enhanced sentiment analysis of memes in Bangla language. Preprints. https://www.preprints.org/manuscript/47eda1c9e7e822a267ff43244a921c7f/download_pub
- Ferdous, T., & Ahmed, R. (2022). Analysis of Bangla-English Code-mixed Sentiment Using Pretrained Transformers. ACM. <https://dl.acm.org/doi/abs/10.1145/3543507.3543530>
- Ferdous, S. M., & Kabir, S. (2023). Deep Learning Approaches for Bangla Product Review Sentiment Analysis. SSRN. <https://ssrn.com/abstract=4432132>
- Haque, R., Khan, M. I., & Rahman, M. M. (2023). A Comparative Study on Bangla Sentiment Analysis Using Machine Learning and Deep Learning Techniques. IEEE. <https://ieeexplore.ieee.org/document/10147857>
- Hashmi, E., Yayilgan, S. Y., & Shaikh, S. (2024). Augmenting sentiment prediction capabilities for code-mixed tweets with multilingual transformers. Social Network Analysis and Mining, 14(1). <https://link.springer.com/article/10.1007/s13278-024-01245-6>
- Hasan, M. R., & Rahman, T. (2024). Low-resource Sentiment Classification of Bangla Text using Prompt-based Learning. arXiv. <https://arxiv.org/abs/2403.01156>
- Hoque, M. N., Salma, U., & Uddin, M. J. (2024). Exploring transformer models in the sentiment analysis task for the under-resourced Bangla language. Natural Language Processing Journal, Elsevier. <https://www.sciencedirect.com/science/article/pii/S2949719124000396>
- Hossain, M. S., Kowsher, M. I., & Shahin, I. (2023). A Comparative Study on Transformer Models for Bangla Sentiment Classification. IEEE. <https://ieeexplore.ieee.org/document/10129543>
- Hossain, M. Z., & Rahman, M. M. (2024). Bangla Text Sentiment Classification Using Fine-Tuned Transformer Models. arXiv. <https://arxiv.org/abs/2403.01473>
- Hossain, M. Z., & Islam, M. F. (2024). Fairness-aware Sentiment Classification for Bangla Social Media Texts. Springer. https://link.springer.com/chapter/10.1007/978-981-99-5241-9_8

- Islam, M. R., Paul, A., & Hossain, M. (2023). Multi-Domain Bangla Sentiment Analysis Using Deep Neural Networks. Springer. https://link.springer.com/chapter/10.1007/978-981-99-0939-0_4
- Islam, M. J., Alam, F. A., & Shahriar, R. (2023). Emotion Recognition from Bangla Text Using Deep Neural Networks. Elsevier. <https://www.sciencedirect.com/science/article/pii/S266682702300057X>
- Islam, M. R., Rahman, S., & Shahriar, R. (2023). Sentiment Analysis of Bangla-English Code-Mixed Data using Fine-tuned XLM-R Model. ACM. <https://dl.acm.org/doi/abs/10.1145/3569219.3569320>
- Islam, M. N., & Hossain, M. (2023). Sentiment Analysis of Bangla Text Using Lightweight Transformer Models. ACM. <https://dl.acm.org/doi/abs/10.1145/3597651.3597663>
- Jahan, R., & Rahman, S. (2023). Sentiment Analysis on Bangla Facebook Comments Using Hybrid Deep Learning Models. IEEE. <https://ieeexplore.ieee.org/document/10246588>
- Jahan, N., & Chowdhury, S. R. (2024). Investigating the Role of Multimodal Features in Bangla Sentiment Classification. Springer. https://link.springer.com/chapter/10.1007/978-981-99-5276-1_13
- Jamal, S. S., & Saha, S. (2023). Bangla Sarcasm Detection and Sentiment Classification Using Transformer Networks. ACM. <https://dl.acm.org/doi/abs/10.1145/3610419.3610432>
- Kabir, S. A., & Ferdous, S. M. (2023). Sentiment Analysis on Bangla Short Text Using Hybrid Deep Learning Model. SSRN. <https://ssrn.com/abstract=4393881>
- Khan, M. F., & Kabir, S. M. R. (2022). A Survey on Bangla Text Classification and Sentiment Analysis. Journal of Information and Communication Technology, 21(1), 1–20. <https://jit.uum.edu.my/index.php/jit/article/view/165>
- Khandaker, S., & Rahman, M. A. (2022). Transformer Models for Detecting Sarcasm in Bangla Sentiment Analysis. IEEE. <https://ieeexplore.ieee.org/document/9961728>
- Khatun, M. T., & Karim, M. (2024). Investigating Aspect and Emotion Detection in Bangla Sentiment Analysis. Springer. https://link.springer.com/chapter/10.1007/978-981-99-7142-7_14
- Mahmud, M., Sultana, S., & Rahman, M. (2023). Context-aware Sentiment Analysis for Code-Mixed Bangla-English Text. IEEE Access. <https://ieeexplore.ieee.org/document/10092631>
- Mahbub, R., & Ahmed, F. (2022). Enhancing Bangla Sentiment Classification with Data Augmentation Techniques. SSRN. <https://ssrn.com/abstract=4157265>
- Mamun, S. A., & Azad, S. (2023). A Review on Bangla Text Sentiment Analysis using Machine Learning and Deep Learning Methods. SSRN. <https://ssrn.com/abstract=4410206>
- Nazir, M. K., Faisal, C. M. N., Habib, M. A., & Ahmad, H. (2025). Leveraging multilingual transformer for multiclass sentiment analysis in code-mixed data of low-resource languages. IEEE Access. <https://ieeexplore.ieee.org/document/10835765>
- Rahman, A., Islam, S. M., & Karim, M. E. (2024). Addressing Fairness and Bias in Bangla NLP Sentiment Models. IEEE Access. <https://ieeexplore.ieee.org/document/10411729>
- Rahman, M. A., & Hossain, M. E. (2024). Sentiment Analysis on Bangla Social Media Data Using Transformer-based Approaches. Springer. https://link.springer.com/chapter/10.1007/978-981-19-6945-5_5
- Raihan, M. M., Rahman, M. F., & Alvee, N. H. (2023). Mixed-Distil-BERT: Code-mixed language modeling. arXiv. <https://arxiv.org/abs/2309.10272>

- Rashid, T., & Islam, M. N. (2023). Multilingual Transformer-based Framework for Bangla Sentiment Analysis and Opinion Mining. IEEE Access. <https://ieeexplore.ieee.org/document/10222331>
- Roy, S., & Saini, J. R. (2023). Sentiment classification in code-mixed Indo-Aryan languages: A transformer-based survey. Springer. https://link.springer.com/chapter/10.1007/978-981-97-5200-3_27
- Sadia, T., Khan, F., & Mahmud, R. (2022). Sentiment Analysis on Bangla Short Texts using Bi-LSTM and Transformer Models. Elsevier. <https://www.sciencedirect.com/science/article/pii/S187705092201195X>
- Sarker, F., Kabir, R., & Rahman, M. M. (2023). Comparative Study of Deep Learning and Transformer Models for Bangla Sentiment Analysis. IEEE. <https://ieeexplore.ieee.org/document/10069598>
- Shuvo, T. R., & Rahman, M. A. (2023). Zero-shot and Few-shot Learning for Bangla Sentiment Analysis. arXiv. <https://arxiv.org/abs/2305.12210>
- Sikder, I., Roy, S., & Islam, M. (2023). Investigating Code-Mixed Bangla-English Sentiment Classification using Transformer-based Models. arXiv. <https://arxiv.org/abs/2303.07781>
- Sultana, B., & Mamun, K. A. (2024). Enhancing Bangla-English code-mixed sentiment analysis with cross-lingual word replacement and data augmentation. IEEE. <https://www.researchgate.net/publication/380829656>
- Sultana, F., Hossain, M. M., & Rahman, M. M. (2023). Comparative Performance Analysis of Transformers in Bangla Sentiment Classification. IEEE. <https://ieeexplore.ieee.org/document/10045923>
- Shuvo, R. A., & Kabir, M. H. (2023). Sentiment Analysis of Bangla Texts Using Fine-tuned BERT Models. ACM. <https://dl.acm.org/doi/10.1145/3578337.3589364>
- Tareq, M., Islam, M. F., Deb, S., & Rahman, S. (2023). Data-augmentation for Bangla-English code-mixed sentiment analysis: Enhancing cross-linguistic contextual understanding. IEEE Access. <https://ieeexplore.ieee.org/document/10129187>
- Zaman, R., & Dey, S. R. (2022). Exploring Code-Mixed Bangla-English Sentiment Analysis with Transformer-based Language Models. SSRN. <https://ssrn.com/abstract=4170213>