

Can LLMs Follow the Pulse of a Crisis? Evaluating Crisis Sentiment in Bangladesh’s July Uprising

Md. Samiul Alim^{1,*} Mahir Shahriar Tamim^{1,*} Tanvir Ahmed Khan¹ Sharjil Khan¹
Rafia Ferdous Duti¹ Shahriyar Zaman Ridoy¹ Mohammad Ali Moni²

¹North South University, Dhaka, Bangladesh

²Charles Sturt University, Australia

samiul.alim01@northsouth.edu, mahir.tamim@northsouth.edu, khan.tanvir01@northsouth.edu,

sharjil.khan@northsouth.edu, rafia.duti@northsouth.edu, shahriyar.zaman01@gmail.com, mmoni@csu.edu.au



Figure 1: Graffiti from the July 2024 student-led uprising in Bangladesh, illustrating public expressions of resistance, solidarity, generational identity, and political sentiment captured on urban walls.

Abstract

Crisis sentiment analysis is especially challenging for low-resource languages such as Bangla, where language, context, and public reaction shift rapidly. We introduce UNRESTSENT200K, a Bangla crisis sentiment dataset with $\approx 200K$ Facebook and YouTube comments from the July–August 2024 Bangladesh uprising. The dataset covers five event-aligned phases, from early escalation and internet blackout to regime transition and a later flood crisis. Each comment is linked to its parent post, enabling evaluation with and without discourse context. All comments are annotated through a fully human process involving 14 native Bangla-speaking annotators and senior validation, achieving substantial agreement ($\kappa = 0.73$, $\alpha = 0.71$) and 94.2% blind-audit agreement. We benchmark fine-tuned encoders, prompted LLMs, and LoRA-tuned LLMs. Results show that parent-post context consistently improves performance, while temporal shift across phases causes large performance drops. Strong LLMs perform well, but still struggle with sarcasm, implicit political references, and phase-dependent meaning. UNRESTSENT200K provides a benchmark for studying context-aware and temporally robust sentiment analysis in low-resource crisis discourse.

* These authors contributed equally to this work.

1 Introduction

During a political crisis, public language can change almost overnight. A short social-media comment that appears positive at one moment may sound sarcastic, fearful, or critical after a major event. Its meaning may also depend on the news post to which it responds. Yet most sentiment benchmarks treat language as stable over time, evaluate each utterance in isolation, and focus mainly on English (Socher et al., 2013; Maas et al., 2011; Go et al., 2009). These assumptions are especially fragile in crisis settings, where sentiment systems are increasingly used to support situational awareness, journalism, and policy decisions (Tufekci, 2017; STEINERT-THRELKELD, 2017). The central problem is therefore not only a lack of data. It is also a lack of evaluation settings that show whether models can follow a rapidly changing public conversation.

The July–August 2024 Bangladesh uprising (Rana et al., 2026) offers a rare opportunity to study this problem (Figure 1). In only seven weeks, public discussion on Facebook and YouTube moved through five distinct phases: early mobilisation, a nationwide internet blackout, regime collapse, political transition, and an overlapping flood crisis. The language, platforms, and broad user population remained largely the same, but the

social and political context changed sharply. This creates a natural stress test: can a model trained during one phase still understand sentiment in the next, and can it interpret a comment without seeing the post that prompted it?

For Bangla, these questions have been difficult to answer. Existing resources such as SentNoB, SentiGOLD, and BnSentMix (Islam et al., 2021, 2023; Alam et al., 2025) have made important progress, but they are mostly static, decontextualised, and focused on consumer or general social-media text. They do not divide a single unfolding event into meaningful time periods, link comments to their parent posts, or support controlled comparisons across phases. As a result, they cannot reveal how quickly model performance changes as an event develops.

To fill this gap, we introduce UNRESTSENT200K, the largest Bangla sentiment dataset to date, with approximately 200K comments ($2.8 \times$ SentiGOLD), as summarised in Table 1. Every comment is timestamped, assigned to one of five event-aligned phases, and paired with its parent post. The labels were produced entirely by 14 native Bangla-speaking undergraduate annotators who had first-hand exposure to the events. Our five-stage **Pilot Label Annotation (PiLA)** framework includes calibration, double annotation, senior adjudication, and an independent blind audit. It achieves substantial agreement (Cohen’s $\kappa = 0.73$ and Krippendorff’s $\alpha = 0.71$) and 94.2% agreement in the blind audit.

Our experiments tell a consistent story. We evaluate four fine-tuned encoders (BanglaBERT, mBERT, XLM-R, and BanglaElectra), 15 zero- and few-shot LLMs, and two LoRA-tuned LLMs. Giving an encoder the parent post improves performance by 7–11 percentage points, showing that context is often essential rather than optional. However, context alone does not solve the problem: training on earlier phases and testing on later ones causes drops of 19–28 F1 points. Sentiment changes occur in sharp, event-linked phases rather than along a smooth trend. GPT-4o-mini, the strongest prompted model, reaches approximately 76% few-shot accuracy but shares several failure modes with the encoders. LoRA-tuned Gemma-4B performs best overall at 78.4% accuracy, yet even this result remains below 80%. Together, these findings show that larger models and stronger prompting help, but do not remove the need for contextual grounding and temporal adaptation.

Contributions.

- **A large Bangla crisis benchmark.** We introduce UNRESTSENT200K, containing approximately 200K comments from the July–August 2024 Bangladesh uprising.
- **A context- and time-aware design.** Each comment is linked to its parent post and an event phase, enabling direct tests of contextual understanding and temporal robustness.
- **Careful human annotation.** PiLA provides a five-stage process with native Bangla-speaking annotators, senior adjudication, and an independent blind audit.
- **Evidence across model families.** Experiments with fine-tuned encoders, prompted LLMs, and LoRA-tuned LLMs show clear gains from parent-post context and large losses under temporal shift.

2 Related Work

Sentiment benchmarks and the low-resource gap. Sentiment analysis is driven by large static English benchmarks (Socher et al., 2013; Maas et al., 2011; McAuley et al., 2015; Go et al., 2009; Hu and Liu, 2004). In Bangla, SentNoB (Islam et al., 2021) introduced noisy multi-domain social-media comments (15K), SentiGOLD (Islam et al., 2023) scaled to 70K with fine-grained labels, and BnSentMix (Alam et al., 2025) targeted code-mixed text. Pretrained Bangla encoders (Bhattacharjee et al., 2022) and recent task-specific resources (Kundu et al., 2026; Nayeem et al., 2026; Islam and Hossen, 2025; Hasan et al., 2024; Talukder, 2025; Rashid et al., 2024) have raised the ceiling, but all are static, decontextualised, and consumer-oriented; none is anchored to a single event whose internal phases can be used to probe within-episode shift (Hossen et al., 2025; Olsen et al., 2024).

Temporal and contextual robustness. Prior work has studied concept drift and temporal generalisation in social-media classification (Tufekci, 2017; STEINERT-THRELKELD, 2017), but often across years, topics, or domains, where temporal change is confounded with demographic and topical shifts. In contrast, UNRESTSENT200K isolates temporal shift within a single seven-week crisis, while keeping language, platform, and population largely fixed. It also operationalises contextual grounding by pairing each comment with its

Table 1: Comparison of Bangla social-media datasets across sentiment, hate-speech, and offensive-language benchmarks.

Dataset	Size	Task	Domain
SentNoB (Islam et al., 2021)	15K	Sent.	Multi
SentiGOLD (Islam et al., 2023)	70K	Sent.	Multi
BnSentMix (Alam et al., 2025)	20K	Sent.	Code-mix
BD-SHS (Romim et al., 2022)	50K	HS	Social
Bengali Tweets (Das et al., 2022)	10K	Hate	Code-mix
TB-OLID (Raihan et al., 2023)	5K	Offensive	Transliterated
BanTH (Haider et al., 2025)	37K	HS	Transliterated
BIDWESH (Fayaz et al., 2025)	9K	HS	Dialectal
BanglaMultiHate (Hasan et al., 2026)	50K	HS	YouTube
UNRESTSENT200K (Ours)	≈200K	Sent.	Crisis

parent post, enabling direct comment-only versus post+comment evaluation.

Human annotation for sensitive discourse. Politically sensitive crisis discourse requires annotators with strong linguistic and cultural grounding. Following best practices in corpus annotation (Pustejovsky, 2012; Sabou et al., 2014; Mukta et al., 2021), we use a fully human five-stage PiLA protocol with guideline piloting, double annotation, senior adjudication, and blind audit. The process involved 14 native-Bangla annotators with first-hand familiarity with the events and achieved substantial reliability: $\kappa = 0.73$, $\alpha = 0.71$, and 94.2% blind-audit agreement. Thus, UNRESTSENT200K contributes not only a benchmark, but also a replicable annotation protocol for low-resource, politically sensitive NLP.

3 Dataset

We construct UNRESTSENT200K, a Bangla sentiment dataset based on public Facebook and YouTube discussions about the July–August 2024 Bangladesh uprising (Rana et al., 2026). The dataset contains ≈200K annotated comments. Each comment is linked to its parent post, which allows us to evaluate sentiment classification in both comment-only and post+comment settings. Each comment is assigned one of three sentiment labels: POSITIVE, NEUTRAL, or NEGATIVE. For each instance, we retain the comment text, parent post, timestamp, platform, and event phase.

3.1 Data Collection

We collected public posts and comments from Facebook and YouTube. The posts were selected from verified news pages, public-interest groups, and high-engagement discussion threads related to the July–August 2024 uprising. We used the

Apify platform¹, the Facebook Graph API, and the YouTube Data API v3 for data collection. The collection period covers July 5 to August 31, 2024. We divide this period into five phases: P1 Pre-Escalation (July 5–15), P2 Crisis and Blackout (July 16–August 4), P3 Post-Blackout (August 5–10), P4 Post-Revolution (August 11–20), and P5 Flood Crisis (August 21–31). These phases are used for temporal analysis and split construction. They are not used as sentiment labels. In total, we collected more than 2,000+ source posts and their comment threads. For each comment, we keep the corresponding parent post so that the dataset can support both comment-only and post+comment sentiment classification.

3.2 Preprocessing

We applied several preprocessing steps before annotation. First, we normalized the text using Unicode NFC normalization and removed irregular whitespace. URLs and user mentions were replaced with special tokens. Emojis were kept because they often carry sentiment in social media comments. We removed comments with fewer than three tokens, comments consisting mainly of non-Bangla text, and obvious spam entries. We also removed near-duplicate comments within the same thread using a Jaccard similarity threshold of 0.85. After preprocessing, the final dataset contains ≈200K comments.

3.3 Annotation Guidelines

We prepared annotation guidelines for three-way sentiment classification. Annotators were asked to assign one label to each comment.

Positive. A comment is labeled positive if it expresses support, approval, hope, praise, relief, or solidarity toward the event, actor, or situation discussed in the parent post.

Neutral. A comment is labeled neutral if it is factual, unclear, mixed, question-like, or does not express a clear positive or negative sentiment.

Negative. A comment is labeled negative if it expresses criticism, anger, frustration, sadness, fear, ridicule, sarcasm, blame, or disappointment. Annotators were instructed to use the parent post as context, but the label was assigned only to the comment. The guidelines included examples for

¹<https://apify.com>

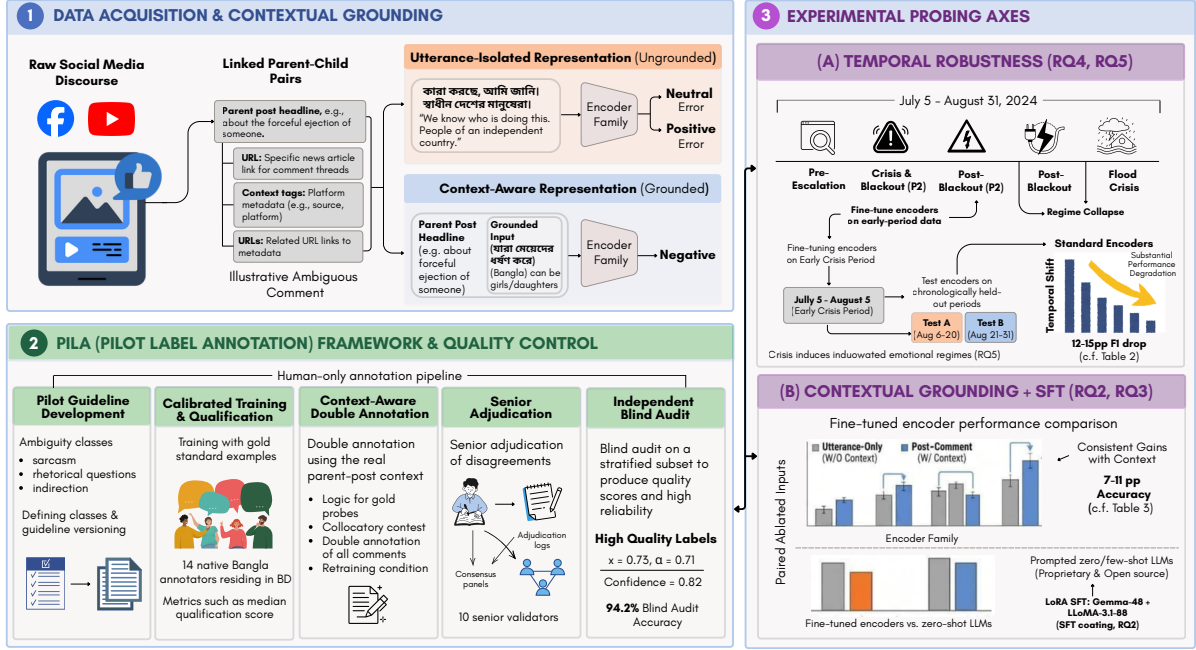


Figure 2: Overview of the UNRESTSENT200K construction and evaluation framework. **(1) Data acquisition and contextual grounding:** public Bangla crisis discourse is collected from Facebook and YouTube, comments are linked to their parent posts and metadata, and utterance-only and context-aware inputs are constructed. **(2) PiLA annotation and quality control:** a five-stage, fully human workflow covers guideline development, annotator calibration, context-aware double annotation, senior adjudication, and an independent blind audit. **(3) Experimental probing axes:** the benchmark evaluates temporal robustness across event phases and the effects of contextual grounding and supervised fine-tuning across encoder and LLM families.

difficult cases, including sarcasm, rhetorical questions, religious expressions, political references, code-switched profanity, and emoji-based sentiment. This follows standard practice in corpus annotation, where guidelines are refined through pilot annotation and disagreement analysis (Pustejovsky, 2012; Sabou et al., 2014).

3.4 Manual Annotation

The annotation was carried out by 14 native Bangla-speaking undergraduate annotators from Bangladesh. All annotators were familiar with the July–August 2024 events and completed pilot training before the main annotation stage. A separate group of 10 senior validators supervised the process. Each comment was annotated independently by two annotators. If both annotators selected the same label, that label was accepted. If they disagreed, the comment was reviewed by a senior validator. Difficult cases were resolved through adjudication based on the annotation guidelines. The annotation process had five stages: pilot guideline development, annotator training and qualification, double annotation, senior adjudication, and blind audit. The full process took approximately eight

months.

3.5 Annotation Agreement

We measured annotation reliability using Cohen’s κ and Krippendorff’s α . The annotation process achieved a pairwise Cohen’s κ of 0.73 and Krippendorff’s α of 0.71, indicating substantial agreement. We also conducted a blind audit on a stratified subset of the dataset. The audit achieved 94.2% exact-label agreement with the final gold labels. These results show that the annotation process was reliable despite the context-dependent nature of the comments.

3.6 Dataset Split

We split the dataset into training, validation, and test sets using a 70/15/15 ratio. The split was stratified by both sentiment label and event phase so that each partition preserves the overall class and temporal distribution. The training set contains 66,382 negative, 29,249 neutral, and 43,950 positive examples. The validation set contains 14,225 negative, 6,268 neutral, and 9,419 positive examples. The test set contains 14,223 negative, 6,267 neutral, and 9,416 positive examples. The full phase-

Table 2: Stratified 70/15/15 train/validation/test split of UNRESTSENT200K by event-aligned phase and sentiment class. Splits are jointly stratified on phase and sentiment to preserve the class and temporal distribution within each partition.

Phase	Period	Train (70%)			Validation (15%)			Test (15%)			Total
		Neg	Neu	Pos	Neg	Neu	Pos	Neg	Neu	Pos	
P1	Pre-Escalation of Revolution (Jul 5–15)	2,281	853	1,576	489	183	338	488	182	337	6,727
P2	Crisis & Blackout (Jul 16–Aug 4)	9,180	4,700	9,698	1,967	1,007	2,078	1,968	1,008	2,078	33,684
P3	Post-Blackout (Aug 5–10)	14,153	7,482	10,704	3,033	1,603	2,294	3,032	1,603	2,293	46,197
P4	Post-Revolution (Aug 11–20)	23,664	10,532	14,479	5,071	2,257	3,103	5,070	2,257	3,103	69,536
P5	Flood Crisis (Aug 21–31)	17,104	5,682	7,493	3,665	1,218	1,606	3,665	1,217	1,605	43,255
Overall		66,382	29,249	43,950	14,225	6,268	9,419	14,223	6,267	9,416	199,399

wise distribution is shown in Table 2. Overall, the dataset contains 94,830 negative (47.6%), 62,785 positive (31.5%), and 41,784 neutral (21.0%) comments.

4 Experiments and Analysis

The July–August 2024 Bangladesh uprising was not a static event (Rana et al., 2026). Public discussion changed as the crisis moved from early protest activity to internet blackout, regime collapse, political transition, and later the flood crisis. This makes UNRESTSENT200K useful for studying more than standard sentiment classification. It allows us to ask whether sentiment models can handle three problems that often appear in real crisis discourse: context dependence, temporal change, and abrupt shifts in public reaction. We organize the experiments around these problems. First, we evaluate recent LLMs to see how far prompting alone can go on Bangla crisis sentiment. Second, we test whether supervised adaptation with mid-scale LLMs improves over prompting. Third, we measure the effect of adding the parent post as context. Finally, we evaluate temporal robustness and examine whether sentiment changes gradually or through phase-level shifts. We evaluate four fine-tuned encoder models: BanglaBERT (Bhattacharjee et al., 2022), BanglaElectra, mBERT (Devlin et al., 2019), and XLM RoBERTa (Conneau et al., 2020). We also evaluate 15 LLMs under zero-shot and few-shot prompting. Across experiments, we report accuracy, macro-F1, and chance-corrected metrics such as Cohen’s κ and MCC where applicable.

4.1 Prompted LLMs

Motivation. We first ask whether recent LLMs can classify sentiment in UNRESTSENT200K without task-specific training. This is an important

baseline because LLMs are often used directly on new events, especially when labeled data is limited.

Models and prompting setup. We evaluate open-source models from the Gemma, Qwen, and LLaMA families, along with proprietary models including GPT-4o-mini, GPT-4.1-mini, and GPT-4.1-nano. Each model is tested with zero-shot and few-shot prompts written in Bangla. In the few-shot setting, the prompt includes seven examples covering positive, negative, neutral, sarcastic, and politically implicit comments. The full prompt templates are shown in Appendix B.1.

Input format. Each test input is given in the same format: [post] [SEP] [comment]. This is the same input format used in the post+comment encoder setting, making the LLM and encoder results comparable.

Main results. Table 3 shows that the strongest prompted model is GPT-4o-mini in the few-shot setting. It reaches 76.0% accuracy, 75.0 macro-F1, and $\kappa = 0.55$. This is slightly above the strongest post+comment encoder, BanglaBERT, which reaches 74.72% accuracy and 0.7167 macro-F1. However, the improvement is small. This suggests that strong LLMs can benefit from the provided context, but prompting alone does not fully solve the task.

Effect of few-shot examples. Few-shot examples help the stronger models more than the smaller ones. For GPT-4o-mini, κ improves from 0.47 to 0.55. For GPT-4.1-mini, it improves from 0.34 to 0.40. In contrast, several smaller open-source models show little improvement from few-shot examples. This suggests that examples are useful only when the model already has enough Bangla language ability and event-level reasoning capacity to use them.

Table 3: Zero-shot and few-shot performance (%) of LLMs with expert prompting (Bangla-native) on UnRest-Sent200K. Best result per section is **bolded and underlined**.

Model	Zero-Shot				Few-Shot			
	Acc	F1	MCC	κ	Acc	F1	MCC	κ
Proprietary Models								
GPT-4o-mini	71.0	69.0	<u>50.0</u>	<u>47.0</u>	<u>76.0</u>	<u>75.0</u>	<u>58.0</u>	<u>55.0</u>
gpt-4.1-mini	<u>73.0</u>	<u>71.0</u>	36.0	34.0	73.0	71.0	45.0	40.0
gpt-4.1-nano	58.0	56.0	22.0	13.0	69.0	61.0	31.0	29.0
Open-Source Models ($\leq 3B$)								
gemma4-2b	<u>61.0</u>	<u>59.0</u>	<u>40.0</u>	<u>38.0</u>	<u>64.0</u>	<u>63.0</u>	<u>45.0</u>	<u>43.0</u>
gemma3-1b	43.0	37.0	18.0	15.0	52.0	47.0	26.0	22.0
qwen2.5-3b	49.0	44.0	25.0	20.0	55.0	52.0	31.0	28.0
qwen2.5-1.5b	46.0	36.0	20.0	11.0	49.0	40.0	23.0	15.0
llama-3.2-3b	43.0	29.0	15.0	5.0	45.0	33.0	17.0	8.0
qwen2.5-0.5b	38.0	22.0	0.2	0.0	39.0	23.0	2.0	0.4
Open-Source Models ($> 3B$)								
gemma4-4b	52.0	49.0	16.0	15.0	<u>66.2</u>	<u>65.0</u>	<u>48.7</u>	<u>46.3</u>
gemma3-4b	49.0	40.0	27.0	16.0	<u>59.0</u>	<u>54.0</u>	<u>39.0</u>	<u>33.0</u>
qwen3-4b	<u>56.0</u>	<u>52.0</u>	<u>35.0</u>	<u>30.0</u>	58.0	56.0	<u>38.0</u>	<u>34.0</u>
qwen3-4b-instruct	56.0	52.0	35.0	30.0	59.0	57.0	23.0	14.0
qwen2.5-7b	47.0	45.0	23.0	21.0	61.0	60.0	42.0	39.0
llama-3.1-8b	53.0	52.0	<u>37.0</u>	11.0	58.0	55.0	37.0	33.0
LoRA SFT								
Fine-Tuned Models								
Gemma-4B	<u>78.4</u>	<u>77.6</u>	<u>65.3</u>	<u>63.1</u>	–	–	–	–
Llama-3.1-8B	<u>76.8</u>	<u>75.9</u>	<u>62.4</u>	<u>60.2</u>	–	–	–	–

Error patterns. A manual inspection of errors shows three recurring patterns. First, models often miss sarcasm, especially when a comment appears positive on the surface but is negative in context. Second, neutral comments are often pushed toward negative labels. Third, politically charged comments become harder when their interpretation depends on the phase of the uprising. These errors are also seen in the encoder models, suggesting that they are properties of the task rather than failures of one model family.

4.2 Supervised Adaptation with Mid-Scale LLMs

Prompting provides a useful starting point, but UNRESTSENT200K also allows us to test whether in-domain supervised training improves performance. We therefore fine-tune two mid-scale LLMs, Gemma-4B and LLaMA-3.1-8B, using

LoRA. Both models are trained with $r=64$, $\alpha=32$, and learning rate 3×10^{-4} . The fine-tuned models outperform all prompted models in Table 3. Gemma-4B achieves 78.4% accuracy, 77.6 macro-F1, MCC 65.3, and κ 63.1. LLaMA-3.1-8B achieves 76.8% accuracy and 75.9 macro-F1. Gemma-4B also improves over GPT-4o-mini by 2.4 percentage points in accuracy and over grounded BanglaBERT by 3.7 points. These results show that supervised adaptation is still important for low-resource crisis sentiment analysis. The smaller Gemma-4B model also performs better than the larger LLaMA-3.1-8B model, suggesting that the choice of backbone and fine-tuning recipe matters more than parameter count alone. At the same time, the best model remains below 80% accuracy, leaving room for future work on crisis-aware and context-aware sentiment models.

Table 4: **Context ablation on UNRESTSENT200K**. Fine-tuned encoder performance when the model receives only the *comment* versus the *parent post concatenated with the comment* via [SEP]. Corpus, splits, optimiser, schedule, and hyperparameters are held fixed; the only variable is the input. Adding the parent post yields a consistent +7–11pp accuracy gain across every architecture. Best per column **bolded and underlined**.

Model	Comment-only input				Post + Comment input			
	[comment]				[post] [SEP] [comment]			
	Acc. (%)	Prec.	Rec.	F1	Acc. (%)	Prec.	Rec.	F1
BanglaBERT	<u>64.19</u>	0.6195	0.6158	0.6170	<u>74.72</u>	0.7245	0.7116	0.7167
mBERT	61.06	0.5869	0.5881	0.5874	69.90	0.6705	0.6790	0.6731
BanglaElectra	62.02	0.6193	0.6202	0.6159	68.98	0.6637	0.6705	0.6649
XLM-RoBERTa-Base	63.41	0.6287	0.6312	0.6299	72.96	0.7289	0.7296	0.7012

Table 5: Temporal generalization performance of supervised fine-tuned (SFT) transformer models under distribution shift on UnRestSent200K. Models are trained on comments collected between July 5 and August 5, 2024, and evaluated on two chronologically non-overlapping future test partitions: Test-A (August 6–20) and Test-B (August 21–31). Results highlight the degradation of model performance under evolving crisis discourse and temporal non-stationarity.

Model	Test	Acc. (%)	Prec.	Rec.	F1 Score
BanglaBERT	A	53.70	0.5204	0.5236	0.5217
	B	<u>54.83</u>	0.5027	0.5222	0.5064
BanglaElectra	A	40.11	0.4195	0.4275	0.4006
	B	41.18	0.4164	0.4343	0.3942
mBERT	A	40.96	0.4234	0.4380	0.4111
	B	41.24	0.4191	0.4354	0.3948
XLM-RoBERTa	A	47.35	0.4708	0.4784	0.4695
	B	47.28	0.4480	0.4664	0.4436

4.3 Qualitative Effect of Fine-Tuning

Table 6 illustrates the three failure patterns flagged in §4.1 sarcasm, neutral/positive collapse to negative, and politically charged surface cues and shows the same (post, comment) pairs being recovered after LoRA-SFT on Gemma-4B. The pattern is qualitative: SFT exposure to the gold PiLA labels lets the model pick up domain-specific pragmatic cues that the zero-shot model misses despite seeing identical input.

4.4 Effect of Parent-Post Context

Many comments in UNRESTSENT200K are difficult to interpret without their parent post, as shown in Table 6. This is common in social media discussions during political events. A short comment

may look factual or neutral by itself, but become clearly positive or negative once the surrounding news post is known. For example, the comment “*We know who is doing this.*” is hard to label in isolation. When paired with a post about a family being forced from their home, the same comment becomes a negative reaction. This shows why the parent post is not extra information; in many cases, it is part of the meaning of the comment. We test this directly using a paired ablation. In the comment-only setting, the encoder receives only the comment. In the post+comment setting, the encoder receives the parent post and comment separated by [SEP]. All other training settings are kept fixed. The results in Table 4 show consistent gains from context. BanglaBERT improves from 64.19% to 74.72% accuracy. XLM-RoBERTa improves by 9.55 points, mBERT by 8.84 points, and BanglaElectra by 6.96 points. The gain appears across all encoder families. This result has a direct implication for future dataset design. In crisis discourse, comment-only sentiment classification can remove information needed for the correct label. For this reason, UNRESTSENT200K supports both comment-only and post+comment evaluation, allowing future work to measure how well models use discourse context.

4.5 Temporal Robustness

We next evaluate whether models trained on earlier phases of the uprising generalize to later phases. This setting reflects a realistic use case: during an unfolding crisis, models may be trained on available data and then applied as the event continues to change. We train the encoder models on comments from July 5 to August 5, 2024. We then

Table 6: Same (post, comment) pair, two models. **Zero-shot** is Gemma-4B with no fine-tuning; **LoRA SFT** is the same backbone after fine-tuning on UNRESTSENT200K ($r=64$, $\alpha=32$, $lr 3 \times 10^{-4}$). ✓ marks a correct prediction against the PiLA gold label; ✗ marks an error. Selected examples illustrate the three failure modes discussed in §4.1.

#	Bangla (Post / Comment) + Translation	Gold	Zero-shot	LoRA SFT
1	Post: আমার বাসায় কাজ করা পিয়ন এখন ৪০০ কোটি টাকার মালিক: প্রধানমন্ত্রী “The peon who worked at my house now owns 4 billion taka: Prime Minister.” Comment: আহ কি গর্বের কথা! “Ah, what a matter of pride!” <i>Sarcasm. Surface “proud” inverts to negative under the PM-corruption headline; zero-shot follows the literal meaning while LoRA SFT captures the sarcastic intent.</i>	Negative	Positive ✗	Negative ✓
2	Post: দেশে ইন্টারনেট সেবায় বীরগতি, ফেসবুক-মেসেঞ্জারে প্রবেশেও বিঘ্ন “Internet service is slow nationwide; access to Facebook and Messenger is also disrupted.” Comment: এটা হাসিনার কাজ “This is Hasina’s doing.” <i>Neutral collapse. The zero-shot model interprets the comment as a factual attribution, whereas LoRA SFT learns its accusatory and negative pragmatic meaning from the task-specific data.</i>	Negative	Neutral ✗	Negative ✓
3	Post: বন্যায় এখন পর্যন্ত ১৩ জনের মৃত্যু, ক্ষতিগ্রস্ত ৪৫ লাখ মানুষ “13 dead in floods so far; 4.5 million affected.” Comment: আল্লাহ আপনি সবাইকে হেফাজত করেন। “May Allah protect everyone.” <i>Surface-cue confusion. Negative disaster-related terms dominate the post and mislead the zero-shot model, while LoRA SFT correctly recognizes the prayer as a positive expression of solidarity.</i>	Positive	Negative ✗	Positive ✓

test them on two later windows: Test-A, covering August 6–20, and Test-B, covering August 21–31. These windows do not overlap with the training period. The results in Table 5 show a large performance drop. BanglaBERT remains the best model, reaching 54.83% accuracy on Test-B, but this is far below its in-distribution post+comment accuracy of 74.72%. Other encoders also degrade under the time-based split. Overall, the models lose between 19 and 28 macro-F1 points compared with their in-distribution results. This drop is important because the language, platforms, and broad event remain the same. What changes is the stage of the crisis. After the blackout and regime collapse, public discussion shifts toward new actors, new concerns, and new emotional frames. Appendix A.8 further shows measurable lexical change across phases using Jensen–Shannon divergence. These results show why random splits are not enough for crisis sentiment analysis. A model that performs well on a random test set may still fail when the event enters a new phase. UNRESTSENT200K therefore provides a useful setting for studying temporal robustness in low-resource NLP.

4.6 Sentiment Change Across Phases

The temporal experiment shows that performance drops over time. We then ask whether this change is gradual or concentrated around specific moments in the uprising. To measure this, we compute Cohen’s d for all pairwise comparisons between the five event phases. The full results are

shown in Table 7. Seven of the ten comparisons show negligible effect sizes. The strongest changes are centered around the Post-Blackout phase. The comparison between Crisis and Blackout and Post-Blackout has a medium effect size ($d = -0.556$), and the comparison between Post-Blackout and Flood Crisis also has a medium effect size ($d = 0.567$). This suggests that sentiment did not change smoothly across the full period. Instead, the largest shift occurred around the moment when internet access returned, the regime collapsed, and people began reacting to a new political situation. Sentiment then shifted again as the flood crisis became the dominant concern. For future research, this makes UNRESTSENT200K useful for studying change-point-aware sentiment models. Models for crisis monitoring may need to detect when the discourse has entered a new phase, rather than assuming that sentiment changes slowly over time.

4.7 Findings

The experiments lead to five findings:

Prompted LLMs are strong, but not sufficient. GPT-4o-mini gives the best prompted result, but it only slightly improves over the strongest post+comment encoder. This shows that prompting alone does not solve the benchmark.

Supervised adaptation gives the best performance. LoRA fine-tuning on UNRESTSENT200K gives the strongest overall result, with Gemma-4B reaching 78.4% accuracy. This suggests that in-domain training remains important for low-

Table 7: Pairwise sentiment-shift comparisons between the five event phases of the 2024 Bangladesh uprising using Cohen’s d effect size. Larger absolute values indicate stronger changes in sentiment distribution between phases, with the strongest shifts observed around the immediate post-blackout period.

Period 1	Period 2	d	Effect
Pre-Escalation	Crisis/Blackout	0.138	Negligible
Pre-Escalation	Post-Blackout	−0.411	Small
Pre-Escalation	Post-Revolution	−0.034	Negligible
Pre-Escalation	Flood Crisis	0.144	Negligible
Crisis/Blackout	Post-Blackout	−0.556	Medium
Crisis/Blackout	Post-Revolution	−0.174	Negligible
Crisis/Blackout	Flood Crisis	0.003	Negligible
Post-Blackout	Post-Revolution	0.381	Small
Post-Blackout	Flood Crisis	0.567	Medium
Post-Revolution	Flood Crisis	0.178	Negligible

resource crisis sentiment analysis.

Parent-post context improves prediction. All encoder models improve by 7–11 percentage points when the parent post is included. This shows that many comments cannot be reliably interpreted in isolation.

Temporal shift causes large degradation. Models trained on earlier phases perform much worse on later phases, even though the language, platforms, and broad event remain the same. This shows that random splits can overestimate performance in crisis settings.

Sentiment changes through phase-level shifts. The largest changes occur around the post-blackout period, rather than gradually across the full timeline. This makes UNRESTSENT200K useful for studying change-point-aware sentiment models.

5 Discussion

Our results show that crisis sentiment analysis depends strongly on both context and time. Models perform much better when the parent post is included, with consistent gains of 7–11 percentage points across encoder models. At the same time, performance drops substantially when models are trained on earlier phases of the uprising and tested on later phases. This shows that standard random splits can overestimate model performance in an unfolding crisis. The results also show that strong LLMs do not remove these challenges. GPT-4o-mini performs well under few-shot prompting, but it only slightly improves over the best post+comment encoder and remains below the best fine-tuned LLM. Its errors also overlap with those of smaller models, especially on sarcasm, implicit political references, and comments whose meaning depends on the event phase. These

findings suggest that future sentiment benchmarks, especially for crisis settings, should include temporal splits and discourse context. For low-resource languages such as Bangla, this is particularly important because models may appear reliable under standard evaluation while failing when the public conversation changes. UNRESTSENT200K provides a setting for studying these issues in a controlled way.

6 Conclusion

We introduce UNRESTSENT200K, a Bangla crisis sentiment dataset containing $\approx 200K$ comments from Facebook and YouTube discussions around the July–August 2024 Bangladesh uprising. Each comment is linked to its parent post and assigned one of three sentiment labels: POSITIVE, NEUTRAL, or NEGATIVE. The dataset is annotated through a fully human process under the PiLA framework. Experiments with fine-tuned encoders, prompted LLMs, and LoRA-tuned LLMs show three main results. First, parent-post context consistently improves sentiment classification. Second, temporal shift across phases of the uprising leads to large performance drops. Third, supervised adaptation gives the strongest results, but the benchmark remains challenging. Overall, UNRESTSENT200K supports research on context-aware and temporally robust sentiment analysis in low-resource crisis discourse.

7 Limitations

UNRESTSENT200K has several limitations. First, it covers only Facebook and YouTube. These platforms were important sources of public discussion during the uprising, but they do not represent all Bangla speakers or all forms of political communication. Other sources, such as X/Twitter, Telegram, private messaging, and offline discourse, are not included. Second, the dataset uses three sentiment labels. This makes large-scale annotation reliable, but it does not capture finer distinctions such as anger, fear, hope, grief, or stance. Future work can extend the label schema to include emotion and stance labels. Third, the dataset is based on one event in Bangladesh. The results should therefore not be treated as universal claims about all crises or all low-resource languages. Cross-event and cross-language studies are needed to test how well the findings generalize. Finally, the models evaluated in this paper were not specifically designed for tem-

poral adaptation or change-point detection. Developing models that can use context and adapt to new crisis phases is an important direction for future work.

References

- Sadia Alam, Md Farhan Ishmam, Navid Hasin Alvee, Md Shahnewaz Siddique, Md Azam Hossain, and Abu Raihan Mostofa Kamal. 2025. BnSentMix: A diverse Bengali-English code-mixed dataset for sentiment analysis. In *Proceedings of the First Workshop on Language Models for Low-Resource Languages*. Association for Computational Linguistics.
- Abhik Bhattacharjee, Tahmid Hasan, Wasi Ahmad, Kazi Samin Mubasshir, Md Saiful Islam, Anindya Iqbal, M. Sohel Rahman, and Rifat Shahriyar. 2022. BanglaBERT: Language model pretraining and benchmarks for low-resource language understanding evaluation in Bangla. In *Findings of the Association for Computational Linguistics: NAACL 2022*, Seattle, United States. Association for Computational Linguistics.
- Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2020. Unsupervised cross-lingual representation learning at scale. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics.
- Mithun Das, Somnath Banerjee, Punyajoy Saha, and Animesh Mukherjee. 2022. Hate speech and offensive language detection in Bengali. In *Proceedings of the 2nd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics and the 12th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, Minneapolis, Minnesota. Association for Computational Linguistics.
- Azizul Hakim Fayaz, MD Uddin, Rayhan Uddin Bhuiyan, Zakia Sultana, Md Samiul Islam, Bidyarthi Paul, Tashreef Muhammad, and Shahriar Manzoor. 2025. Bidwesh: A bangla regional based hate speech detection dataset. *arXiv preprint arXiv:2507.16183*.
- Alec Go, Richa Bhayani, and Lei Huang. 2009. Twitter sentiment classification using distant supervision. *CS224N project report, Stanford*, 1(12):2009.
- Maarten Grootendorst. 2022. Bertopic: Neural topic modeling with a class-based tf-idf procedure. *arXiv preprint arXiv:2203.05794*.
- Fabiha Haider, Fariha Tanjim Shifat, Md Farhan Ishmam, Md Sakib Ul Rahman Sourve, Deeparghya Dutta Barua, Md Fahim, and Md Farhad Alam Bhuiyan. 2025. BanTH: A multi-label hate speech detection dataset for transliterated Bangla. In *Findings of the Association for Computational Linguistics: NAACL 2025*, Albuquerque, New Mexico. Association for Computational Linguistics.
- Mahmudul Hasan, Md Rashedul Ghani, and KM Azharul Hasan. 2024. Aspect based sentiment analysis datasets for bangla text. *Data in Brief*, 57:111107.
- Md Arif Hasan, Firoj Alam, Md Fahad Hossain, Usman Naseem, and Syed Ishtiaque Ahmed. 2026. LLM-based multi-task Bangla hate speech detection: Type, severity, and target. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, San Diego, California, United States. Association for Computational Linguistics.
- Md. Sabbir Hossen, Md. Saiduzzaman, and Pabon Shaha. 2025. Social media sentiments analysis on the july revolution in bangladesh: A hybrid transformer based machine learning approach. In *2025 17th International Conference on Electronics, Computers and Artificial Intelligence (ECAI)*, pages 1–8.
- Minqing Hu and Bing Liu. 2004. Mining and summarizing customer reviews. In *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '04, page 168–177, New York, NY, USA. Association for Computing Machinery.
- Ariful Islam and Md Rifat Hossen. 2025. Banglasentnet: A hybrid deep learning framework for multi-aspect sentiment analysis in bangla e-commerce reviews. In *Data Science, AI and Applications*. Springer Nature Switzerland.
- Khondoker Ittehadul Islam, Sudipta Kar, Md Saiful Islam, and Mohammad Ruhul Amin. 2021. SentNoB: A dataset for analysing sentiment on noisy Bangla texts. In *Findings of the Association for Computational Linguistics: EMNLP 2021*. Association for Computational Linguistics.
- Md. Ekramul Islam, Labib Chowdhury, Faisal Ahamed Khan, Shazzad Hossain, Md Sourave Hossain, Mohammad Mamun Or Rashid, Nabeel Mohammed, and Mohammad Ruhul Amin. 2023. Sentigold: A large bangla gold standard multi-domain sentiment analysis dataset and its evaluation. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, KDD '23, page 4207–4218. Association for Computing Machinery.

- Swastika Kundu, Autoshi Ibrahim, Mithila Rahman, and Tanvir Ahmed. 2026. Anubhuti: a comprehensive corpus for sentiment analysis in bangla regional languages. In *2026 IEEE 2nd International Conference on Quantum Photonics, Artificial Intelligence & Networking (QPAIN)*, pages 1–6. IEEE.
- Andrew L. Maas, Raymond E. Daly, Peter T. Pham, Dan Huang, Andrew Y. Ng, and Christopher Potts. 2011. Learning word vectors for sentiment analysis. In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, Portland, Oregon, USA. Association for Computational Linguistics.
- Julian McAuley, Christopher Targett, Qinfeng Shi, and Anton van den Hengel. 2015. Image-based recommendations on styles and substitutes. In *Proceedings of the 38th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '15. Association for Computing Machinery.
- Md. Saddam Hossain Mukta, Md. Adnanul Islam, Faisal Ahamed Khan, Afjal Hossain, Shuvanon Razik, Shazzad Hossain, and Jalal Mahmud. 2021. A comprehensive guideline for bengali sentiment annotation. *ACM Trans. Asian Low-Resour. Lang. Inf. Process.*, 21(2).
- Md. Darun Nayeem, Zarin Rafa, Tasnuva Tasnim Nova, Yasin Rahman, Abdul Mumeet Pathan, and Md. Masudul Islam. 2026. Banglamuse: A multi-modal bangla sentiment dataset of text–audio pairs for speech and sentiment analysis. *Data in Brief*, 65:112458.
- Helene Olsen, Étienne Simon, Erik Velldal, and Lilja Øvrelid. 2024. Socio-political events of conflict and unrest: A survey of available datasets. In *Proceedings of the 7th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2024)*, St. Julians, Malta. Association for Computational Linguistics.
- James Pustejovsky. 2012. [The role of linguistic models and language annotation in feature selection for machine learning](#). In *Proceedings of the Sixth Linguistic Annotation Workshop*, page 1, Jeju, Republic of Korea. Association for Computational Linguistics.
- Md Nishat Raihan, Umma Tanmoy, Anika Binte Islam, Kai North, Tharindu Ranasinghe, Antonios Anastasopoulos, and Marcos Zampieri. 2023. Offensive language identification in transliterated and code-mixed Bangla. In *Proceedings of the First Workshop on Bangla Language Processing (BLP-2023)*, Singapore. Association for Computational Linguistics.
- Md. Sohel Rana, Shadia Sharmin, Mohammad Bin Amin, and Judit Oláh. 2026. From revolt to revolution: exploring motivating factors of bangladesh’s july 2024 uprising. *Research in Globalization*, 12:100341.
- Mohammad Rifat Ahmmad Rashid, Kazi Ferdous Hasan, Rakibul Hasan, Aritra Das, Mithila Sultana, and Mahamudul Hasan. 2024. A comprehensive dataset for sentiment and emotion classification from bangladesh e-commerce reviews. *Data in Brief*, 53:110052.
- Nauros Romim, Mosahed Ahmed, Md Saiful Islam, Arnab Sen Sharma, Hriteshwar Talukder, and Mohammad Ruhul Amin. 2022. BD-SHS: A benchmark dataset for learning to detect online Bangla hate speech in different social contexts. In *Proceedings of the Thirteenth Language Resources and Evaluation Conference*, Marseille, France. European Language Resources Association.
- Marta Sabou, Kalina Bontcheva, Leon Derczynski, and Arno Scharl. 2014. Corpus annotation through crowdsourcing: Towards best practice guidelines. In *Proceedings of the Ninth International Conference on Language Resources and Evaluation (LREC’14)*, Reykjavik, Iceland. European Language Resources Association (ELRA).
- Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D. Manning, Andrew Ng, and Christopher Potts. 2013. Recursive deep models for semantic compositionality over a sentiment treebank. In *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*, Seattle, Washington, USA. Association for Computational Linguistics.
- ZACHARY C. STEINERT-THRELKELD. 2017. [Spontaneous collective action: Peripheral mobilization during the arab spring](#). *American Political Science Review*, 111(2):379–403.
- Md. Ashraful Islam Talukder. 2025. Bangla and banglish e-commerce reviews dataset for aspect-based sentiment analysis. *Mendeley Data*.
- Zeynep Tufekci. 2017. *Twitter and tear gas: The power and fragility of networked protest*. Yale University Press.

A Appendix

A.1 Reproducibility Statement

To support reproducibility, the release package will include the processed and de-identified UNRESTSENT200K corpus; fixed training, validation, and test splits; parent-post context; timestamps; platform and event-phase metadata; annotation guidelines; prompt templates; preprocessing and evaluation scripts; and the configurations used for the encoder and LoRA experiments. Dataset statistics and phase boundaries are reported in Section 3, while detailed preprocessing, annotation, temporal-analysis, topic-modeling, and prompting procedures are provided in the remainder of this appendix.

Availability. The dataset and accompanying resources will be available in the [UNRESTSENT200K repository](#).

A.2 AI Usage Disclosure

Large language models were used as experimental systems to generate benchmark predictions for UNRESTSENT200K under zero-shot and few-shot prompting and supervised fine-tuning. AI-assisted tools were also used for limited non-experimental support, including language editing, figure ideation, and visualization assistance. All AI-assisted material was manually reviewed and verified by the authors, who retain full responsibility for the paper’s content, results, and interpretations.

A.3 Preprocessing Details

This section describes the preprocessing steps used before annotation, as summarized in Section 3.2.

A.3.1 Normalization

- **Unicode normalization:** Text is normalized to NFC form to reduce inconsistent Bangla character encodings.
- **Whitespace cleanup:** Extra spaces, tabs, and irregular line breaks are removed.
- **Entity masking:** URLs and user mentions are replaced with [URL] and [USER].
- **Emoji retention:** Emojis are kept because they often express sentiment in social media comments.

A.3.2 Filtering

We remove comments with fewer than three tokens after normalization. We also remove comments that consist mainly of non-Bangla text or obvious spam. To keep the released corpus primarily Bangla, comments with substantial English or Bangla–English transliteration are excluded. Figure 3 reports the resulting sentiment distribution across word-length bins of the retained corpus: short comments (≤ 10 words) skew positive, while neutral expression increases with comment length and negative sentiment stays relatively flat across lengths.

A.3.3 Deduplication

Near-duplicate comments within the same thread are removed using a Jaccard similarity threshold

Table 8: Table showing topic modeling statistics for each period. C_v denotes the topic coherence score, computed using normalized pointwise mutual information (NPMI) and ranging from 0 to 1, where higher values indicate greater semantic coherence. $Eff.$ represents the effective number of topics, calculated as $\exp(H)$, where $H = -\sum_i p_i \log p_i$ is the Shannon entropy of the topic distribution. This quantity reflects the number of equally prominent topics that would yield the same level of topic diversity.

Period	Docs	Topics	C_v	Eff.
P1: Pre-Escalation	6,727	12	0.405	1.7
P2: Crisis & Blackout	33,684	29	0.781	1.8
P3: Post-Blackout	46,197	18	0.384	1.9
P4: Post-Revolution	69,536	25	0.420	2.0
P5: Flood Crisis	43,255	25	0.422	1.9
Global	199,399	29	0.375	—

of 0.85. Duplicates across different posts are retained when they appear as part of separate public discussions.

A.3.4 Post-Level Framing Labels

In addition to comment-level sentiment labels, each source post is assigned a coarse framing label:

$$\mathcal{Y}^{(p)} = \{\text{HOPE}, \text{DESPAIR}, \text{OUTRAGE}\} \quad (1)$$

These labels describe the dominant framing of the parent post. HOPE captures optimistic or change-oriented framing, DESPAIR captures grief or uncertainty, and OUTRAGE captures anger, blame, or calls for accountability. These post-level labels are used as annotation context and metadata; the main prediction task remains comment-level sentiment classification.

A.4 PiLA: Annotation Protocol Details

This appendix provides additional details about the **Pilot Label Annotation (PiLA)** framework shown in Figure 2. PiLA is a fully human annotation process. No labels are assigned by automated systems. The full annotation effort took approximately eight months and involved 14 primary annotators and 10 senior validators.

A.4.1 Annotator and Validator Pool

The primary annotation team consists of 14 undergraduate annotators from Bangladesh. All annotators are native Bangla speakers and were familiar with the July–August 2024 events through Bangla

news and social media. The validator team consists of 10 senior annotators who did not overlap with the primary annotation pool. All participants provided informed consent, and the study protocol was reviewed by the host institution.

A.4.2 Guideline Document

The annotation guideline defines the three sentiment labels: POSITIVE, NEUTRAL, and NEGATIVE. It also provides rules and examples for common ambiguous cases, including sarcasm, rhetorical questions, religious expressions, political references, code-switched profanity, and emoji-based sentiment. Annotators were instructed to use the parent post as context, but not to copy the post’s sentiment into the comment label. The label was assigned to the comment only. The guideline was refined during the pilot stage and then fixed before the main annotation stage.

A.4.3 Qualification and Quality Checks

Before the main annotation stage, annotators completed a qualification round. Only annotators who reached Cohen’s $\kappa \geq 0.65$ were admitted to the main task. All 14 annotators met this threshold, with scores ranging from 0.66 to 0.81.

During annotation, gold-check examples were inserted at a 2% rate. Annotators whose agreement dropped below the required threshold were paused, retrained on the relevant ambiguity cases, and re-qualified before continuing. Labels from affected batches were reviewed and re-annotated where necessary.

A.4.4 Adjudication

Each comment was independently labeled by two annotators. If both annotators agreed, their shared label was accepted. If they disagreed, the comment was reviewed by a senior validator. Difficult cases were escalated to a small validator panel.

For adjudicated comments, we record the two initial labels, annotator confidence scores, validator decision, and final gold label. These logs allow later analysis of disagreement patterns and label reliability.

A.5 Quality Metrics Details

A.5.1 Pairwise Inter-Annotator Agreement

We report Cohen’s κ for agreement between the two primary annotators assigned to each comment. The overall pairwise score is $\kappa = 0.73$, indicating

substantial agreement. We compute:

$$\kappa = \frac{p_o - p_e}{1 - p_e}, \quad (2)$$

where p_o is observed agreement and p_e is expected agreement under the annotators’ marginal label distributions. Per-phase κ ranges from 0.69 to 0.78.

A.5.2 Multi-Rater Reliability

We also report Krippendorff’s α to measure reliability across the full annotator pool. The overall score is $\alpha = 0.71$, showing that the labels are consistent beyond individual annotator pairs.

A.5.3 Annotator Confidence

Annotators provided a confidence score in $[0, 1]$ for each label. The mean confidence score is $\bar{c} = 0.82$ with standard deviation 0.14. In total, 87.3% of labels have confidence scores of at least 0.7. These scores are included as metadata so that users can filter examples by confidence if needed.

A.5.4 Pre-Adjudication Agreement

The two primary annotators agreed directly on 78.4% of comments. The remaining 21.6% were sent to senior validators for adjudication. Among these, 3.1% required review by a three-validator panel.

A.5.5 Independent Blind Audit

As a final quality check, two senior validators who were not involved in adjudication re-labeled a stratified random sample of 2,000 comments, with 400 comments from each phase. They were blind to the final gold labels. The audit achieved 94.2% exact-label agreement and 96.8% agreement within one step on the sentiment scale. Most disagreements involved rhetorical questions and emoji-mediated sentiment, which were also identified as difficult cases during guideline development.

A.6 Effect of Comment Length

We analyze sentiment distribution across word-level comment length intervals to examine whether affective polarity varies with comment verbosity. As shown in Figure 3, shorter comments are more frequently positive, with the positive share decreasing from 50% in 1–10 word comments to 20% in comments longer than 50 words. In contrast, neutral sentiment increases steadily with length, suggesting that longer comments tend to provide explanation, context, or factual discussion rather than direct affective reactions. Negative sentiment

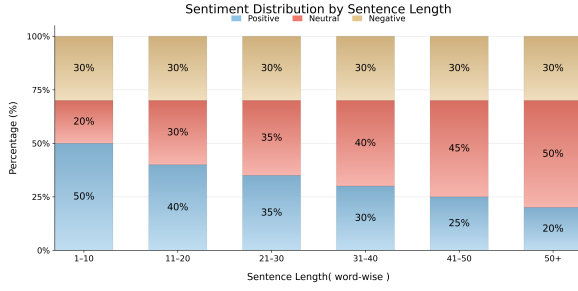


Figure 3: Sentiment distribution across sentence length intervals (word-wise). Short sentences (1–10 words) are dominated by positive sentiment, while the proportion of neutral sentiment increases with sentence length. Negative sentiment remains relatively stable across all length ranges.

remains comparatively stable across all length ranges. This indicates that comment length is associated with sentiment composition, and that models may need to account for verbosity-driven shifts in crisis discourse.

A.7 Quantifying Temporal Non-Stationarity

We measure lexical change across event phases using three signals: Jensen–Shannon (JS) divergence over unigram distributions, Jaccard similarity of top words, and Shannon entropy. The results are shown in Figure 4.

JS divergence. JS divergence values range from 0.3713 to 0.3832 across phase pairs, showing that word distributions change over time.

Jaccard similarity. Top-word Jaccard similarity ranges from 0.4235 to 0.4652. This means that less than half of the most frequent terms are shared across phase pairs, indicating vocabulary turnover as the event develops.

Entropy. Shannon entropy increases from 7.5989 to 7.7872, suggesting that the discourse becomes more lexically diverse in later phases. Together, these results show that the corpus changes measurably across event phases. This supports the temporal-shift results in Table 5: performance drops are associated with changes in the data distribution, not only with model limitations. The corresponding sentiment-level shift, quantified using Cohen’s d on the polarity distributions of all $\binom{5}{2} = 10$ pairwise phase pairs, is reported in Table 7; seven of ten contrasts are negligible, with the strongest effects clustered around the post-blackout transition mirroring the lexical-level drift pattern above.

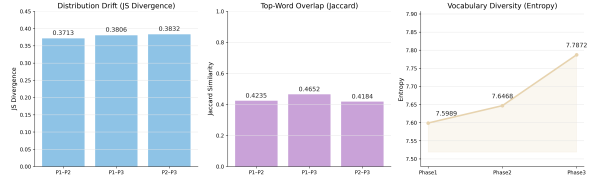


Figure 4: Temporal drift across three chronological phases of Bangla text. Divergence, vocabulary turnover, and rising entropy together reveal substantial non-stationarity and increasing lexical diversity over time.

A.8 Topic Modeling: Extended Methodology and Results

Detailed Methodology

Our topic modeling methodology employs BERTopic with the following specifications: We use a multilingual sentence transformer for embeddings, UMAP for dimensionality reduction, and HDBSCAN for density-based clustering, augmented with a custom Bangla stopwords list (619 terms) to filter discourse markers (158 terms), platform-specific artifacts (151 terms) and special tokens (6 terms). (see Table 8 for period-specific topic statistics).

Hyperparameter selection was guided by a trade-off between granularity and interpretability. In the global model, we use UMAP ($n_components=5$, $n_neighbors=20$, $min_dist=0.1$) to preserve semantic structure while reducing dimensionality, HDBSCAN ($min_cluster_size=20$, $min_samples=5$, $cluster_selection_epsilon=0.1$) with lenient parameters to avoid excessive fragmentation of crisis discourse, and CountVectorizer with $n_grams=(1,3)$, $min_df=10$, $max_df=0.9$ to capture Bangla phrasal compounds while filtering ubiquitous terms. In contrast, period-specific models employ adaptive clustering. The $min_cluster_size$ is scaled with each period’s corpus size, from 25 for the smallest period to 100 for the largest, to maintain comparable topic granularity across corpus sizes. Documents identified as outliers (HDBSCAN $topic=-1$) are reassigned to the nearest topic centroids via cosine similarity, ensuring complete corpus coverage. Finally, topic coherence stability is assessed under reasonable clustering perturbations ($min_cluster_size \pm 20\%$: $\Delta C_v < 0.05$), indicating that the results are robust to sensible hyperparameter variation.

Extended Results

Topic dynamics exhibit bursty, non-stationary patterns. Dynamic temporal tracking reveals that

75% of topics display a coefficient of variation exceeding 1.0, indicating high volatility relative to mean activity. Peak topic intensity occurs on August 9 three days post-collapse with Topic T1 reaching 1,063 daily posts. The average topic lifespan is 44.8 days, with the dominant topic T0 experiencing 15 momentum shifts and volatility of 4,269.78. These results reflect the punctuated attention dynamics during the crisis.

Thematic evolution shows high semantic continuity. Cross-period analysis of 100 topic transformations reveals high semantic similarity (mean = 0.93), indicating that themes tend to *evolve* rather than rupture across event phases. Student protest discourse (ছাত্রলীগের, আন্দোলন) morphs into state violence narratives (পুলিশ, তোলা) with similarity 0.902, while confrontation rhetoric transforms into celebratory discourse (জড়িত, হাসিমাখা) with similarity 0.954 and +462% volume growth post-collapse.

Sentiment is asymmetrically negative and tightly coupled to topic networks. Topic-level sentiment analysis identifies pronounced negativity bias: Only Topic T8 (সিদ্ধান্ত, সংস্কার — decisions, reforms) shows a positive polarity (+0.246) whereas the most negative topics (T15: -0.486; T6: -0.294) dominate volume. Topic co-occurrence analysis identifies 38 significant associations ($r > 0.5$), with decision-making topics tightly coupled to public reaction topics (T3 $\leftrightarrow$ T1, $r = 0.988$), indicating reactive action-response cycles characteristic of crisis communication (Tufekci, 2017).

Collectively, these results show that revolutionary online discourse is characterized by: (i) extreme thematic centralization around core grievances, (ii) crisis-induced semantic crystallization under acute stress, (iii) bursty temporal dynamics with delayed collective sensemaking, (iv) high-continuity thematic evolution across event phases, and (v) an asymmetric sentiment structure dominated by negative expression. These patterns challenge common assumptions about pluralistic fragmentation in online political discourse modeling and suggest that crisis communication follows distinct structural principles that require specialized analytical frameworks.

A.9 Failure Mode Discovery: Topic-Level Sentiment Instability

We go beyond aggregate performance metrics by examining sentiment volatility at the topic level

to uncover systematic failure modes under temporal shift, complementing the lexical-level drift signals reported in Figure 4. Using our sentiment-flip detection framework on the 25 topics extracted via BERTopic, we measure polarity volatility using standard deviation (σ) across the five event phases. The result is pronounced, topic-dependent instability.

Pronounced heterogeneity in temporal robustness. Volatility clusters are asymmetric: 6 topics are highly volatile ($\sigma > 30$), 14 are medium ($15 \leq \sigma \leq 30$), and 5 are stable ($\sigma < 15$). Topic 22 (কোটা — quota reform) is especially unstable ($\sigma = 45.76$, range = 90.48), flipping from +57.1% to -33.3% polarity across the P3→P4 transition (Figure 5). By contrast, humanitarian topics remain relatively steady ($\sigma \approx 12$ –15). This shows volatility is topic-specific, not uniform.

Critical transitions concentrate misclassification risk. Of 15 topics with dramatic flips ($\Delta > 40\%$), the largest shifts cluster around two transitions: (i) P2→P3 (post-blackout), where four topics flip by >60% as information flow and sense-making resume and (ii) P3→P4 (regime collapse), where five topics reverse by >45% as celebratory and uncertain framings collide. Topic 11 (সমস্যা — problems) shifts from -15.0% to -45.6% ($\Delta = 64.2\%$) during P2→P3. A classifier trained on P2 would therefore underestimate negativity in P3 by more than 30 percentage points for this topic.

Implications for temporal model robustness. Temporal distribution shift is not uniform. Models fail at different rates depending on the topic, with politically charged themes (quota reform, league politics) showing 3× volatility than humanitarian concerns. This motivates topic-aware temporal recalibration where model confidence is dynamically adjusted based on topic volatility profiles. Doing so could meaningfully improve robustness over time. Future work should explore selective fine-tuning strategies that prioritize highly volatile topics during temporal adaptation.

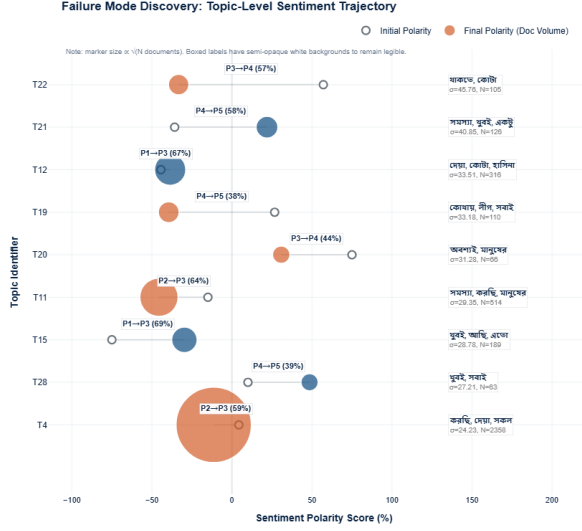


Figure 5: Topic-level sentiment trajectories across the five phases (P1–P5) of the July Revolution. Hollow circles mark initial polarity, while filled circles mark final polarity and scale with document volume. Arrows and percentage labels show the direction and magnitude of each shift; Bangla topic names and document counts appear along the right margin. Volatility varies substantially across topics: T22 (কোটা) shifts from +57.1% to −33.3% ($\sigma = 45.76$), whereas flood-relief topics remain relatively stable ($\sigma \approx 12$ –15). The largest reversals cluster around P2→P3, following the internet blackout, and P3→P4, during regime collapse, with several shifts exceeding 60%. These patterns reveal topic-specific temporal distribution shifts that challenge robust sentiment modeling.

A.10 Failure Mode Discovery through Topic Modeling

To investigate systematic weaknesses beyond overall performance metrics, we applied BERTopic to the 2,789 misclassified comments in our dataset. The model identifies 33 semantically coherent topics ($C_v=0.9499$), using a smaller cluster configuration tailored to the error corpus ($\text{min_cluster_size}=15$, $\text{n_grams}=(1,2)$).

Catastrophic failure of the neutral class. The error analysis reveals a critical failure of the neutral class. Although the negative and positive classes show 83.71% and 63.10% accuracy respectively, the neutral class achieves only 42.54% accuracy, corresponding to a 57.46% error rate. Neutral→negative errors alone account for 895 instances, representing 32.1% of all misclassifications and revealing a strong systematic bias toward negative sentiment.

Error concentration in semantic clusters. Misclassifications are not evenly distributed but instead concentrate in distinct topics. Topic 0 alone

Table 9: Top misclassification topics revealing systematic failure modes. Topic 0 dominates with 676 errors (24.2% of all misclassifications), while the neutral class shows catastrophic failure across all topics with 57.46% error rate. High prediction confidence (mean=0.977) on errors indicates overconfident misclassification rather than uncertainty.

Topic	Errors	Conf.	Neu→Neg	Neu→Pos	Neg↔Pos
T0: আন্দোলনের স্পিরিটের	676	0.968	261	61	104
T1: ঘটেছে বৈষম্যবিরোধী	369	0.971	186	20	58
T2: বিদ্যুৎকেন্দ্র	134	0.966	29	5	23
T3: আগের শিক্ষাক্রমে	110	0.968	40	18	29
T4: শিক্ষার্থীদের স্বপ্নকে	92	0.970	14	13	38
T5: আটক মির্জাপুরে	89	1.000	0	85	4
T6: ছাত্রলীগের হামলায়	75	0.980	40	7	10
T7: চিরকালই বিদ্রোহী	75	0.982	19	9	17
T8: শিক্ষার্থীদের অসন্তি	72	0.992	10	16	27
T9: দূর আগের	70	0.977	10	11	27
All Topics (33)	2,789	0.977	895	356	531

contributes 676 errors. Of these, 322 involve neutral samples predominantly mislabeled as negative (81.1%). These clusters center on political discourse (বৈষম্যবিরোধী আন্দোলনের) and educational policy (আগের শিক্ষাক্রমে), suggesting that BanglaBERT relies heavily on lexical negativity cues while failing to capture pragmatic neutrality, which is a subtle but pervasive feature of Bangla social media discourse. LoRA-SFT on a mid-scale LLM recovers several of these pragmatically dependent cases (qualitative examples in Table 6).

Class-level bias and skew. Table 9 highlights the scale of this failure: 1,251 of 2,177 neutral samples (57.46%) are misclassified, with a 2.5:1 bias towards negative labels (895 neutral→negative vs. 356 neutral→positive). This skew is consistent across topics. T1, T3, and T6 each show 4–5× higher neutral→negative rates than neutral→positive, reinforcing the model’s tendency to interpret ambiguous Bangla discourse as negative.

Overconfident misclassification. The model is highly certain even when it is wrong. The mean confidence across all errors is 0.977. In Topic 5, confidence reaches a perfect 1.000 while 85 neutral samples are mislabeled as positive. Such extreme overconfidence indicates a failure of uncertainty calibration. This makes standard probability thresholds ineffective for error detection and highlights the need for topic-aware confidence recalibration.

Extreme sentiment flips. Apart from the collapse of the neutral class, 531 samples exhibit neg↔pos flips, accounting for 19.0% of all errors. The topics 0 (104 flips) and 1 (58 flips) dominate these failures, both centered on বৈষম্যবিরোধী আন্দোলনের dis-

course where the political framing strongly influences sentiment valence. Topic 4 (শিক্ষার্থীদের স্বপ্নকে students’ dream) shows a particularly high flip rate of 41.3% (38/92 errors), revealing the weakness of BanglaBERT when it comes to differentiating between aspirational and critical framings of educational reform.

A.11 Structural Diagnosis: Where Do Models Fail?

Aggregate metrics tell us *how much* models fail, but not *where*. To pinpoint failure structure we apply BERTopic (Grootendorst, 2022) to the corpus with dynamic temporal tracking, and to the model error set separately, yielding two complementary views: (i) the latent thematic structure of crisis discourse itself, and (ii) the concentration of model errors within that structure. The full pipeline (multilingual sentence-transformer embeddings, UMAP + HDBSCAN, period-adaptive hyperparameters, 619-term Bangla stopword list) is detailed in Appendix A.8; we report the headline findings here and defer extended results to Appendices A.8–A.8.

Table 8 summarises the period-stratified output. Two phenomena are immediately load-bearing for our discussion of model failures.

Crisis discourse is hegemonic, not pluralistic. The global model recovers 29 topics ($C_v = 0.375$, silhouette = 0.402), but one topic (T0: লীগ, সরকার, কোটা league, government, quota) absorbs 56.3% of all documents ($n = 105,226$), exhibits remarkable temporal stability ($CV = 0.014$), and carries net-negative polarity (-0.224). Shannon-entropy effective topic counts of just 1.7–2.0 across phases confirm that apparent diversity masks deep concentration. We term this pattern *discourse hegemony*: under crisis, collective attention converges on a single political grievance rather than fragmenting across separable concerns. The implication for sentiment modelling is direct the dominant training signal is one centred, persistent topic, and any model that overfits its lexical surface will inherit T0’s negativity prior on out-of-topic comments. This matches the systematic neutral-to-negative collapse observed in our error analysis (Appendix A.8).

Crisis induces semantic crystallization, not noise. A counter-intuitive pattern emerges from period-specific modelling: topic coherence is not driven by corpus size. P2 (Crisis & Blackout, $n = 33,684$, 16.9% of corpus) achieves by far the highest coherence ($C_v = 0.781$, against 0.384–0.422

in every other period) while producing as many topics (29) as the global model. The largest period, P4 (Post-Revolution, $n = 69,536$, 34.9%), is markedly less coherent ($C_v = 0.420$), and the smallest, P1 (Pre-Escalation, $n = 6,727$, 3.4%), lower still ($C_v = 0.405$). We call this the crisis coherence paradox: extreme conditions do not fragment discourse, they crystallise it. Under information scarcity (the blackout) and acute existential stakes, the public converges sharply on a small set of immediately salient concerns; after the crisis resolves, discourse diversifies and coherence drops.

These two phenomena hegemony and crystallization explain a structural feature of model failure that aggregate metrics hide: errors do not distribute uniformly over discourse but concentrate around the dominant topic and its adjacent themes. The topic-conditional error analysis in Appendix A.8 formalises this, showing that politically charged topics (quota reform, party politics) exhibit $3\times$ the sentiment volatility of humanitarian ones and that 33% of all errors fall within T0’s orbit. The actionable consequence is that temporal robustness must be measured *topic-conditionally*, and that recalibration strategies for crisis sentiment models should be steered by topic volatility profiles rather than by aggregate confidence scores.

A.12 Prompting Techniques

We evaluate the LLMs in our benchmark under two standard in-context learning paradigms, **zero-shot** and **few-shot** prompting. In both settings the model parameters are kept frozen and the only signal the model receives about the task comes through the natural-language prompt, which lets us isolate the contribution of in-context reasoning from any task-specific fine-tuning.

Zero-shot prompting. In the zero-shot setting, the model is given a task description and the input instance but no labelled examples. Each prompt starts with a short *expert-persona* instruction (e.g., “You are an expert in Bengali political sentiment analysis”), followed by the parent post, the target comment, the allowed label set {positive, neutral, negative}, and a constraint that the response must be a single label. Because no demonstrations are provided, zero-shot performance reflects the model’s intrinsic ability to interpret crisis-era Bangla discourse using only its pre-training knowledge, and is therefore especially sensitive to lexical coverage, cultural-political grounding, and instruction-following quality in a low-

resource language.

Few-shot prompting. In the few-shot setting, the same template is preceded by $k = 7$ labelled (post, comment, label) demonstrations drawn from the training split of UNREST-SENT200K. The demonstrations are chosen to jointly cover (i) all three sentiment classes in {positive, neutral, negative}, (ii) both supportive and critical framings of crisis events, and (iii) linguistically challenging phenomena such as sarcasm, rhetorical questions, emoji descriptions, code-mixed Bangla–English expressions, and implicit political references. The model is then asked to label a held-out (post, comment) pair under the same output constraint, which lets us measure how much a small number of in-domain exemplars can compensate for the absence of fine-tuning.

Decoding and output format. For all settings we use greedy decoding (temperature = 0) and constrain the output to a single token from {positive, neutral, negative}. Responses that do not exactly match one of the allowed labels are counted as errors rather than re-prompted, so reported numbers reflect end-to-end prompt reliability rather than post-hoc parsing. The full prompt templates used in our experiments are reproduced in §B.1.

A.13 Expert Prompting with LLMs

Table 3 reports the zero-shot and few-shot performance of open-source and proprietary LLMs using expert prompting with Bangla-native instructions on the UnRestSent200K benchmark.

B Ethics Statement

All data used in this work were collected from publicly accessible Facebook and YouTube content in accordance with platform terms of service. Personally identifiable information, including usernames and profile links, was removed during preprocessing.

B.1 Sample Prompts and Result

ফিউ-শট প্রম্পট: সেন্টিমেন্ট অ্যানোটেশন (Few-Shot: Sentiment Annotation — Bengali)

System: তুমি বাংলা ভাষার রাজনৈতিক sentiment analysis বিশেষজ্ঞ। তোমাকে একটি সংবাদ পোস্ট/শিরোনাম এবং সেই পোস্টের উপর করা একটি মন্তব্য দেওয়া হবে। তোমার কাজ হলো নির্ধারণ করা যে মন্তব্যটি উল্লিখিত পোস্ট, ঘটনা, বা প্রসঙ্গের প্রতি কী ধরনের sentiment প্রকাশ করছে।

User: নিচে কিছু উদাহরণ দেওয়া আছে। উদাহরণগুলো অনুসরণ করে শেষে দেওয়া প্রশ্নটির উত্তর দাও।

নির্দেশনা:

- post শুধুমাত্র contextual information হিসেবে ব্যবহার করবে।
- sentiment নির্ধারণের সময় sarcasm, rhetorical tone, slang, emoji description (যেমন: 'কান্না জড়িত হাসিমুখ প্রতিক্রিয়া'), এবং mixed Bangla-English expression বিবেচনা করবে।
- রাজনৈতিক মতাদর্শ বা ব্যক্তিগত অবস্থান নয়, শুধুমাত্র comment-এর sentiment বিচার করবে।
- যদি comment সমর্থন, প্রশংসা, আশাবাদ, বা ইতিবাচক প্রতিক্রিয়া প্রকাশ করে → positive
- যদি comment সমালোচনা, হতাশা, রাগ, ব্যঙ্গ, আক্রমণ, বা নেতিবাচক প্রতিক্রিয়া প্রকাশ করে → negative
- যদি comment তথ্যভিত্তিক, প্রশংসামূলী, দ্ব্যর্থক, বা স্পষ্ট ইতিবাচক/নেতিবাচক sentiment না প্রকাশ করে → neutral

উদাহরণসমূহ:

উদাহরণ ১

পোস্ট: কোটাবিরোধী আন্দোলন : ঢাকা বিশ্ববিদ্যালয়ে মিছিল শুরু, মধুর ক্যানটিনে জড়ো হয়েছে ছাত্রলীগ

মন্তব্য: সাবধান, আবার রক্ত ঝরাবে।

উত্তর: negative

উদাহরণ ২

পোস্ট: অপরাধে সম্পৃক্ততায় যুবদল--ছাত্রদলের ১৫ জনকে বহিষ্কার

মন্তব্য: নাটক কম করো পিও, জনগণ সব বুঝে।

উত্তর: negative

উদাহরণ ৩

পোস্ট: আজ বৃহস্পতিবার দেশের প্রায় সব মোবাইল অপারেটরের ইন্টারনেট ধীরগতিতে চলছে।

মন্তব্য: শ্লা টাউট! কান্না জড়িত হাসিমুখ প্রতিক্রিয়া

উত্তর: negative

উদাহরণ ৪

পোস্ট: কোটাবিরোধী আন্দোলন : ঢাকা বিশ্ববিদ্যালয়ে মিছিল শুরু, মধুর ক্যানটিনে জড়ো হয়েছে ছাত্রলীগ

মন্তব্য: ছাত্রলীগও কোটা আন্দোলনে অংশ নিয়েছে? ভাবনা মুখ প্রতিক্রিয়া

উত্তর: neutral

উদাহরণ ৫

পোস্ট: অপরাধে সম্পৃক্ততায় যুবদল--ছাত্রদলের ১৫ জনকে বহিষ্কার

মন্তব্য: যাদের পদ নেই তাদের কি শাস্তি হবে?

উত্তর: neutral

উদাহরণ ৬

পোস্ট: কোটাবিরোধী আন্দোলন : ঢাকা বিশ্ববিদ্যালয়ে মিছিল শুরু, মধুর ক্যানটিনে জড়ো হয়েছে ছাত্রলীগ

মন্তব্য: কোটা সিস্টেম নিপাত যাক। মেধাবীরা চাকরি পাক।

উত্তর: positive

উদাহরণ ৭

পোস্ট: অপরাধে সম্পৃক্ততায় যুবদল--ছাত্রদলের ১৫ জনকে বহিষ্কার

মন্তব্য: সঠিক সময়ের সঠিক সিদ্ধান্ত ধন্যবাদ

উত্তর: positive

এখন নিচের উদাহরণটির sentiment নির্ধারণ করো:

পোস্ট: {post}

মন্তব্য: {comment}

শুধুমাত্র নিচের একটি label প্রদান করো: positive, neutral, negative

Few-Shot: Sentiment Annotation — English

System: You are an expert in political sentiment analysis for the Bangla language. You will be given a news post or headline and a comment made on that post. Your task is to determine what kind of sentiment the comment expresses toward the mentioned post, event, or context.

User: Some examples are given below. Follow the examples and answer the final question.

Instructions:

- Use the post only as contextual information.
- While determining sentiment, consider sarcasm, rhetorical tone, slang, emoji descriptions such as “crying laughing face reaction”, and mixed Bangla-English expressions.
- Judge only the sentiment of the comment, not political ideology or personal position.
- If the comment expresses support, praise, optimism, or a positive reaction, label it **positive**.
- If the comment expresses criticism, frustration, anger, sarcasm, attack, or a negative reaction, label it **negative**.
- If the comment is factual, question-like, ambiguous, or does not express a clear positive or negative sentiment, label it **neutral**.

Examples:

Example 1

Post: Anti-quota movement: A procession has started at Dhaka University, and Chhatra League has gathered at Madhur Canteen.

Comment: Be careful, they will shed blood again.

Answer: negative

Example 2

Post: Fifteen members of Jubo Dal and Chhatra Dal expelled for involvement in crimes.

Comment: Stop the drama, people understand everything.

Answer: negative

Example 3

Post: Today, Thursday, internet service is slow across almost all mobile operators in the country.

Comment: Damn fraud! Crying laughing face reaction.

Answer: negative

Example 4

Post: Anti-quota movement: A procession has started at Dhaka University, and Chhatra League has gathered at Madhur Canteen.

Comment: Has Chhatra League also joined the quota movement? Thinking face reaction.

Answer: neutral

Example 5

Post: Fifteen members of Jubo Dal and Chhatra Dal expelled for involvement in crimes.

Comment: Will those who do not hold any position also be punished?

Answer: neutral

Example 6

Post: Anti-quota movement: A procession has started at Dhaka University, and Chhatra League has gathered at Madhur Canteen.

Comment: Down with the quota system. Let meritorious students get jobs.

Answer: positive

Example 7

Post: Fifteen members of Jubo Dal and Chhatra Dal expelled for involvement in crimes.

Comment: The right decision at the right time. Thank you.

Answer: positive

Now determine the sentiment of the following example:

Post: {post}

Comment: {comment}

Provide only one of the following labels: **positive**, **neutral**, or **negative**.

Zero-Shot Prompt: Political Expert (Bengali)

তুমি একজন বাংলা রাজনৈতিক sentiment analysis বিশেষজ্ঞ। নিচের বাংলা পোস্ট এবং মন্তব্যটির sentiment নির্ধারণ করো।

পোস্ট: {post}

মন্তব্য: {comment}

শুধুমাত্র একটি লেবেল দাও: positive, neutral, negative

Zero-Shot: Political Expert — English

You are an expert in Bengali political sentiment analysis. Determine the sentiment of the following Bengali post and comment.

Post: {post}

Comment: {comment}

Respond with only one label: positive, neutral, or negative.