

LLM-Based Multi-Task Bangla Hate Speech Detection: Type, Severity, and Target

WARNING: This paper contains examples which may be disturbing to the reader

Md Arif Hasan¹, Firoj Alam², Md Fahad Hossain³, Usman Naseem⁴,
Syed Ishtiaque Ahmed¹

¹University of Toronto, Canada, ²Qatar Computing Research Institute, Qatar
³Daffodil International University, Bangladesh ⁴Macquarie University, Australia
{arif, ishtiaque}@cs.toronto.edu, fialam@hbku.edu.qa

Abstract

Online social media platforms are central to everyday communication and information seeking. While these platforms serve positive purposes, they also provide fertile ground for the spread of hate speech, offensive language, and bullying content targeting individuals, organizations, and communities. Such content undermines safety, participation, and equity online. Reliable detection systems are therefore needed, especially for low-resource languages where moderation tools are limited. In Bangla, prior work has contributed resources and models, but most are single-task (e.g., binary hate/offense) with limited coverage of multi-facet signals (type, severity, target). We address these gaps by introducing the *first multi-task* Bangla hate-speech dataset, *BanglaMultiHate*, one of the largest manually annotated corpus to date. Building on this resource, we conduct a comprehensive, controlled comparison spanning classical baselines, monolingual pretrained models, and LLMs under zero-shot prompting and LoRA fine-tuning. Our experiments assess LLM adaptability in a low-resource setting and reveal a consistent trend: although LoRA-tuned LLMs are competitive with BanglaBERT, culturally and linguistically grounded pretraining remains critical for robust performance. Together, our dataset and findings establish a stronger benchmark for developing culturally aligned moderation tools in low-resource contexts. For reproducibility, we will release the dataset and all related scripts.

1 Introduction

The rise of social media has increased the spread of harmful online content (Walther, 2022), with hate speech emerging as a critical societal issue given its potential to perpetuate discrimination (Gelber, 2021), harassment, and violence. Given the large volume of user-generated content, manual moderation is neither scalable nor consistent, highlighting the urgent need for reliable and scalable auto-

mated hate speech detection systems. Although substantial progress has been achieved in high-resource languages such as English (Albladi et al., 2025), research effort for low-resource languages like Bangla are relatively limited (Sharma et al., 2025; Das et al., 2022a; Haider et al., 2025; Romim et al., 2022).

Identifying hate speech in Bangla imposes unique challenges due to its rich morphology, free word order, and code-switching with English and other regional dialects, making it difficult for models trained on other languages to generalize effectively. Furthermore, the scarcity of annotated datasets and the lack of high-quality pretrained resources exacerbate the difficulty of building accurate classification systems (Al Maruf et al., 2024). Existing studies often rely on classical machine learning models (Kiela et al., 2020; Mridha et al., 2021; Romim et al., 2022), deep learning models (Romim et al., 2022; Keya et al., 2023), and adapt pretrained models designed primarily for English (Mridha et al., 2021). However, these approaches often fail to capture the cultural, social, and linguistic nuances that shape how hate is expressed in Bangla (Al Maruf et al., 2024), such as context-dependent slurs, metaphorical insults, or region-specific idiomatic usage (Jahan et al., 2022). Addressing these challenges requires not only improved datasets and resources but also approaches that are sensitive to the sociolinguistic realities of Bangla discourse, ensuring that models move beyond surface-level understanding and engage with the deeper structures of the language.

In recent years, the rapid advancement of large language models (LLMs) such as GPT-5, Claude, Gemini (Comanici et al., 2025), Llama (Dubey et al., 2024), and Qwen (Yang et al., 2025) has shown remarkable success across a variety of downstream NLP tasks, often demonstrating strong generalization abilities in zero-shot or few-shot scenarios. This has raised important questions about

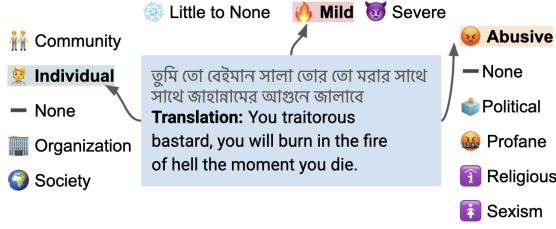


Figure 1: An example of hateful comment with its English translation showing type, severity and target of hate.

their applicability in sensitive domains such as hate speech detection, particularly for underrepresented languages. However, the zero-shot performance of LLMs in low-resource contexts is often limited and their inability to capture context-dependent information (Zahid et al., 2025). Moreover, hate speech is highly context-dependent and culturally nuanced, making it difficult for LLMs pretrained on high-resource languages to detect, demonstrating the need for targeted adaptation strategies in low-resource and sensitive tasks.

To address these challenges, we developed the first multi-tasks hate speech dataset named *Bangla-MultiHate* for Bangla. This dataset is specifically designed to support a variety of classification tasks: (i) identifying different types of hate speech, (ii) the severity of hate, and (iii) determining the target of hate. An example of a hateful comment with type, severity and target is demonstrated in Figure 1. This study also conducts a comprehensive evaluation of hate speech detection in Bangla across three tasks, utilizing SVM, BanglaBERT, zero-shot settings (Llama3 and Qwen3), and LoRA fine-tuned on these LLMs. Our contributions can be summarized as follows:

- We developed the first multi-task hate speech dataset for Bangla and one of the largest manually annotated hate speech datasets, which includes type of hate, severity and target.
- We provide comprehensive comparisons of classical, monolingual pretrained, and zero-shot and LoRA fine-tuned approaches using LLMs for Bangla hate speech detection.
- We assess the effectiveness of zero-shot inference and LoRA fine-tuning for LLMs, offering insights into their adaptability in low-resource tasks.
- We highlight key limitations and trends, demonstrating that while fine-tuned LLMs are comparable to the performance of BanglaBERT, emphasizing the continued

importance of culturally and linguistically grounded pretraining for combating online hate speech in low-resource languages.

Our findings are summarized as follows:

- Fine-tuned monolingual BanglaBERT yields superior performance.
- Zero-shot learning failed to perform better than the majority baseline as well as SVM.
- SVM performs comparatively better than fine-tuned LLMs on severity and target of hate tasks, while Llama3 performs slightly better on the type of hate task.
- Model performance varies significantly with the complexity of the task.

2 Related Work

The identification of offensive language and hate speech has become increasingly important due to the extensive use of social media, which has created an environment in which harmful content can spread rapidly (Jiang and Zubiaga, 2024). Research on hate speech identification has progressed rapidly over the past decade (Fortuna and Nunes, 2018), moving from lexicon-based classifiers to transformer models and, more recently LLMs (Albladi et al., 2025).

2.1 Models

Various classical models (such as logistic regression (LR), SVM, and random forest), deep learning models (e.g., LSTM), and transformer-based models (e.g., BERT, XLM-R, MuRIL, AraBERT, etc.) have been studied in the literature. Sharif et al. (2021) demonstrated that transformer-based pretrained language models (e.g., Indic-BERT, XLM-R, mBERT) outperform classical models (e.g., LR, SVM). The multi-task learning approach using AraBERT has been studied in Arabic for the identification of offensive language and hate speech (Djandji et al., 2020), while random forest, k-nearest neighbors, and MLP classifiers have been studied for offensive language identification from Dravidian code-mixed texts (B and A, 2021). Pelicon et al. (2021) employed mBERT and LASER models for zero-shot cross-lingual transfer learning, demonstrating promising results in languages such as German, Spanish, Indonesian, and Arabic. Similarly, (Saumya et al., 2021a) explores the impact of cross-cultural transfer learning, showing how biases across cultures affect model performance, examining the impact of cross-cultural

Dataset	Size	Type	Labels / Tasks	Source
BD-SHS (Romim et al., 2022)	50,281	Comments	3-level: HS vs. non-HS; target; HS type	Social media
Bengali Tweets (Das et al., 2022a)	10,000	Code-mixed	Hate/offense detection (binary)	Twitter/X
TB-OLID (Raihan et al., 2023)	5,000	Transliterated, code-mixed	OLID: Offensive vs. Not; target type (Indiv./Group/Untargeted)	Facebook
BanTH (Haider et al., 2025)	37,350	Transliterated	Multi-label target; HS	YouTube
BIDWESH (Fayaz et al., 2025)	9,183	Dialectal	Hate vs. non-hate; ~13 types; 7 targets	Social media
BanglaMultiHate (Ours)	50,746	Comments	Type, severity, target	YouTube

Table 1: Overview of existing datasets and ours.

transfer learning.

Kiela et al. (2020) utilizes SVM, CNN, and LSTM models to evaluate performance on hateful content. SVM, naive bayes, and random forest, along with transformation methods have been studied for multi-label hate speech identification (Ibrohim and Budi, 2019). Mridha et al. (2021) employed L-Boost, a modified AdaBoost algorithm combining BERT embeddings with LSTM models, to identify offensive texts in Bangla and Banglish social media content. SVM, LSTM, and Bi-LSTM models have also been analyzed by Romim et al. on Bangla YouTube and Facebook comments, with results showing that SVM outperforms LSTM and Bi-LSTM. Furthermore, combining BERT and GRU architectures for hate speech detection has been explored in Bengali social media texts (Keya et al., 2023). Explainable hate speech identification has recently attracted attention in the literature (Yang et al., 2023; Piot and Parapar, 2025; Sariyanto et al., 2025).

2.2 Existing Hate Speech Datasets

There has been effort to develop datasets in the past. Gupta et al. (2022) introduced a 150K-comment dataset for abusive speech detection in five Indic languages, while Sharif et al. (2021) studied offensive language detection in multilingual code-mixed text. These work establish important baselines for future research on code-mixed offensive text detection in Dravidian languages (Saumya et al., 2021b; Chakravarthi et al., 2022).

Some of the notable resources for hate and abusive content on Bangla include 10,178 tweets labeled as hate/offensive/normal (Das et al., 2022b), a 30K comments dataset with 10K hate speech examples (Romim et al.), 3K transliterated Bangla-English abusive comments (Sazzed, 2021), 50K offensive comments from online social networking (Romim et al., 2022), and 10K Bangla posts consisting of 5K actual and 5K Romanized Bengali tweets (Das et al., 2022a). Moreover, a multi-label transliterated Bangla hate speech dataset has

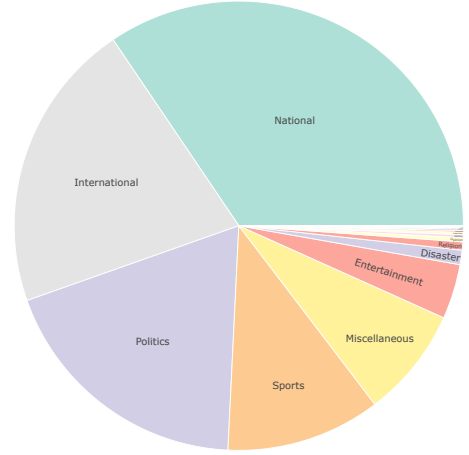


Figure 2: Distribution of the MultiHate dataset across different categories.

been developed by Haider et al. (2024) utilizing a translation-based LLM prompting approach.

Building on these efforts, Table 1 provides an overview of existing resources. Our contribution extends this landscape by introducing a larger dataset that not only supports multiple tasks but also incorporates a richer topical hierarchy, spanning 19 topics and 120 sub-topics.

3 Dataset

3.1 Data Collection

We collected public comments from YouTube videos using the YouTube API¹, primarily from So-moy TV, which is a popular Bangla News channel. The comments belong to 19 different categories, including *Business*, *Celebrities*, *Disaster*, *Entertainment*, *Fashion*, *Geopolitics*, *Health*, *History*, *International*, *Lifestyle*, *Literature*, *Miscellaneous*, *National*, *Opinion*, *Politics*, *Religion*, *Science*, *Sports*, and *Technology*, as well as 120 subcategories. In total, we collected approximately 55,000 comments associated with various Bangla news videos. We then removed all entries containing only emojis and URLs, as well as duplicate entries. Additionally, we excluded all Banglish comments (Bangla text written using the English alphabet) from the initial dataset. After applying these filtering and duplicate-removal steps, the dataset contained 50,746 entries. The category-wise data distribution is presented in Figure 2, with more than 90% of the comments concentrated in five categories.

¹<https://developers.google.com/youtube/v3>

3.2 Data Annotation

3.2.1 Annotation Guidelines

We developed an annotation guideline to facilitate the annotation of data. Our annotation setup was a multitask annotation. Therefore, each instance could be assigned multiple labels to capture the overlapping and nuanced nature of the content, such as identifying the type of hate, the severity of hate, and the target of hate simultaneously. The guidelines provided clear definitions, decision criteria, and illustrative examples to ensure consistency across annotators. Below, we briefly discuss the annotation guidelines for each annotation task.

1. **Type of Hate:** The purpose of this task is to identify the type of hate from YouTube comments. The annotators classified whether the comments are *Abusive*, *Sexism*, *Religious Hate*, *Political Hate*, *Profane*, or *None* based on the criteria discussed in the Appendix A. Depending on the nature of the annotation, the annotator proceeds with different subsequent tasks. If a comment is marked as *None*, annotators automatically assign the labels *Little to None* for the severity of hate and *None* for the target of hate; otherwise, they proceed with the regular annotation process.
2. **Severity of Hate:** This task aims to assess the degree of hate expressed in a given comment. Annotators evaluate whether the comment reflects *Little to None*, *Mild*, or *Severe* forms of hate, taking into account factors such as the intensity of derogatory expressions, the use of slurs, and the presence of threats or incitement to violence. The objective is to capture not only the presence of hateful content but also its strength and potential impact. In the annotation guidelines, clear criteria and illustrative examples are provided to ensure consistency and reliability among annotators.
3. **Target of Hate:** This task focuses on identifying the specific *Individuals*, *Organizations*, *Communities*, *Society*, or *None* that is the target of hateful expression. Annotators classify whether the hate is directed toward protected characteristics such as organizations, communities, or society, or if it is aimed at individuals without reference to group identity. In cases where no explicit target is present, annotators assign the label *None*. The goal of this task is to capture the social dimension of hateful language, enabling analysis not only of the pres-

ence of hate but also of who or what is being targeted.

3.2.2 Manual Annotation

The annotation team comprised 35 native Bangla-speaking undergraduate students, including both male and female annotators. The team was trained and supervised by expert annotators to ensure reliability and consistency in the labeling process. Each comment was independently annotated by three annotators, and periodic quality checks of randomly selected samples were conducted, followed by feedback sessions to maintain annotation standards. The final label for each comment was determined by majority agreement among the annotators. In instances of persistent disagreement, consensus meetings were organized to resolve discrepancies and establish the final annotation.

Annotation Agreement We evaluated the inter-annotator agreement (IAA) of the manual annotations using Fleiss’ Kappa coefficient (κ) to assess the reliability of the annotation process across all tasks. The obtained κ scores were 0.71, 0.84, and 0.79 for the type of hate, severity of hate, and target of hate tasks, respectively, indicating substantial to perfect agreement.² We also observed that an increase in the number of annotation classes introduces greater challenges for annotators, which is reflected in lower IAA scores for more fine-grained tasks.

3.3 Analysis and Statistics

We present the detailed class label distribution in Table 2. The table reports the frequency of labels across the three tasks, for the training, development, and test splits. This breakdown highlights the inherent class imbalance across tasks, with *None* being the most frequent label, whereas categories such as *Sexism* or *Religious Hate* are underrepresented. Such skewed distributions demonstrate the challenge of building robust models capable of handling rare but socially significant cases of hate speech. Moreover, Table 5 presents the distribution of class labels across word length bins for the three tasks. The majority of samples in all splits fall within the ≤ 20 , indicating that most instances of hate speech in Bangla are expressed concisely.

In Figures 3 and 4, we present the relationship between hate type and severity, and between hate

²According to Landis and Koch (1977), values of κ between 0.61–0.80 represent substantial agreement, while values between 0.81–1.0 represent almost perfect agreement.

Split	Type of Hate	Severity of Hate			Total	Target of Hate					Total
	Class	LN	Mild	Severe		Comm.	Indiv.	Org.	Society	None	
Train	Abusive	2,254	3,867	2,091	8,212	1,521	3,340	1,510	1,118	723	8,212
	Political Hate	978	2,237	1,012	4,227	404	935	1,896	745	247	4,227
	Profane	127	396	1,808	2,331	399	1,182	385	183	182	2,331
	Religious Hate	150	301	225	676	283	119	51	147	76	676
	Sexism	26	52	44	122	28	70	4	12	8	122
	None	19,954	–	–	19,954	–	–	–	–	19,954	19,954
	Total	23,489	6,853	5,180	35,522	2,635	5,646	3,846	2,205	21,190	35,522
Dev	Abusive	333	507	273	1,113	207	427	251	141	87	1,113
	Political Hate	141	292	141	574	43	130	270	101	30	574
	Profane	25	63	254	342	52	171	58	23	38	342
	Religious Hate	16	36	26	78	28	18	4	17	11	78
	Sexism	4	11	4	19	8	9	1	1		19
	None	2,898	–	–	2,898	–	–	–	–	2,898	2,898
	Total	3,417	402	698	5,024	338	755	584	283	3,064	5,024
Test	Abusive	646	1,075	591	2,312	441	891	465	306	209	2,312
	Political Hate	263	702	255	1,220	124	263	541	231	61	1,220
	Profane	41	116	552	709	127	354	127	51	50	709
	Religious Hate	29	93	57	179	62	47	17	34	19	179
	Sexism	7	15	7	29	5	16	2	3	3	29
	None	5,751	–	–	5,751	–	–	–	–	5,751	5,751
	Total	6,737	2,001	1,462	10,200	759	1,571	1,152	625	6,093	10,200

Table 2: Class label distribution of the dataset. LN: Little to None, Comm.: Community, Indiv.: Individual, Org.: Organization

type and target, respectively. The *None* category was excluded from the hate type for a concise visual representation. Figure 3 shows that *Abusive* content is most prevalent, peaking at mild severity with a notable severe presence, while *Political Hate* follows a smaller but similar trend. *Profane* content stands out, concentrated in the severe category, suggesting profanity as a marker of high severity. *Religious Hate* appears at low frequency across all severities, and *Sexism* is rare overall. Mild severity emerges as the dominant category across types. Figure 4 highlights that individuals and organizations are the main targets, with abusive expressions disproportionately directed at individuals.

3.4 Data Split

The dataset was partitioned into training, development, and test sets, comprising 70%, 10%, and 20% of data, respectively, for our experiments. We applied stratified sampling (Sechidis et al., 2011) to ensure a balanced class label distribution across all splits. We provide the detailed data distribution across the splits in Table 3. As shown, the dataset is highly imbalanced across all three annotation dimensions. For the *type of hate* task, the majority of samples fall under the *none* class, while categories such as *sexism* and *religious hate* are comparatively

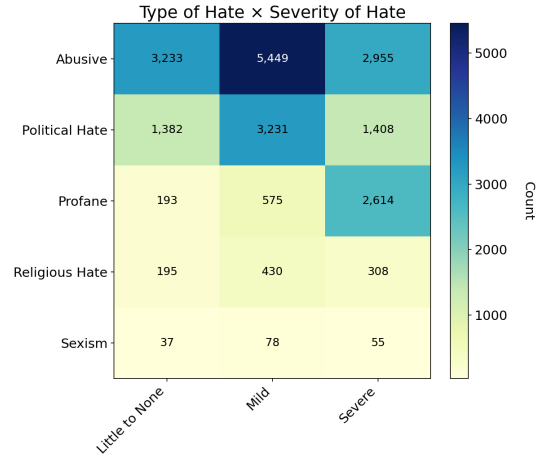


Figure 3: Heatmap demonstrating the relationship between type of hate and severity.

underrepresented. A similar pattern is observed in the *severity of hate* task, where most comments are labeled as *little to none*, followed by *mild* and *severe*. For the *target of hate* task, the *none* class again dominates, whereas labels such as *society* and *community* appear far less frequently. This imbalance highlights the inherent challenges of training reliable models on underrepresented classes and demonstrate the importance of stratification for fair evaluation.

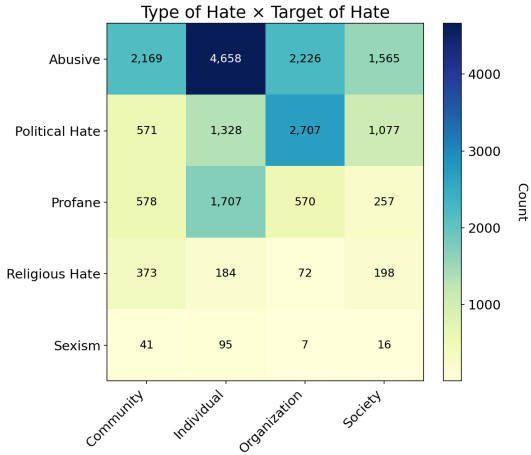


Figure 4: Heatmap demonstrating the relationship between *type* of hate and *target*.

Class	Train	Dev	Test	Total
Type of Hate				
Abusive	8,212	1,113	2,312	11,637
Political Hate	4,227	574	1,220	6,021
Profane	2,331	342	709	3,382
Religious Hate	676	78	179	933
Sexism	122	19	29	170
None	19,954	2,898	5,751	28,603
Total	35,522	5,024	10,200	50,746
Severity of Hate				
Severe	5,180	698	1,462	7,340
Mild	6,853	909	2,001	9,763
Little to None	23,489	3,417	6,737	33,643
Total	35,522	5,024	10,200	50,746
Target of Hate				
Community	2,635	338	759	3,732
Individual	5,646	755	1,571	7,972
Organization	3,846	584	1,152	5,582
Society	2,205	283	625	3,113
None	21,190	3,064	6,093	30,347
Total	35,522	5,024	10,200	50,746

Table 3: Class label distribution across three tasks of the *BanglaMultiHate* dataset.

4 Methodology

4.1 Models

We experiment with classical model such as SVM, monolingual pretrained language model such as BanglaBERT (Bhattacharjee et al., 2022), and large language models such as Llama-3.2-3B-Instruct³ and Qwen3-4B-Instruct-2507⁴. We choose models

³Llama-3.2-3B-Instruct

⁴Qwen3-4B-Instruct

from different model families to provide extensive evaluation with this dataset.

Baseline. We used a majority-class baseline that always predicts the class with the highest frequency in the training data and a random approach. These methods have been widely used as a baseline technique in numerous prior studies (e.g., (Rosenthal et al., 2017)).

Classical models. We employed SVM (Platt, 1998) with TF-IDF representation which has been extensively utilized in prior research and remains prevalent in low-resource production settings. Our setup employed 1–5 n-grams with TF-IDF weighting and a regularization parameter of $C = 1$.

Pretrained Language Model (PLM). Given that PLMs have demonstrated significant success in the past years and are also computationally reasonable choices for many downstream NLP tasks, we fine-tuned the monolingual BanglaBERT model (Wolf et al., 2020). Following the procedure of Devlin et al. (2019), we trained each model with default hyperparameters for 3 epochs. To mitigate training instability, we performed ten runs with different random seeds and selected the best model based on development set performance. All experiments were conducted independently for each task.

LLMs. Recent advances in LLMs have received significant attention from researchers to evaluate the performance of these models, especially for low-resource languages in various downstream NLP tasks. We experiment with Llama-3.2-3B-Instruct (Dubey et al., 2024) and Qwen3-4B-Instruct (Yang et al., 2025). We adopt a zero-shot learning setup for all models. To ensure reproducibility, we apply a consistent prompt, response format, output token limit, and decoding configuration (such as temperature set to 0) across models. The prompts were crafted using concise instructions, as detailed in Appendix B. We also demonstrate the efficacy of *BanglaMultiHate* dataset by fine-tuning both Llama and Qwen models. We choose PEFT using LoRA (Hu et al., 2022) to reduce the computational cost. We trained the model in full precision (FP16) using the Adam optimizer. The learning rate was set to 2×10^{-4} , with LoRA parameters $\alpha = 16$ and $r = 64$. The maximum sequence length was fixed at 512, and training was performed with a batch size of 8. Fine-tuning was conducted for three epochs without additional hyperparameter tuning. Moreover, both zero-shot and fine-tuned approaches were studied in a multi-task setup due to computational resource constraints.

4.2 Instructions Dataset

We employed a template-based approach to generate diverse English instructions, obtaining 10 hate speech classification task templates per language from GPT-4.1 and Claude-3.5 Sonnet.⁵ During fine-tuning and inference, one template was randomly selected and combined with the comment. We report examples of prompts in Appendix B.

4.3 Evaluation Measures

Across all experimental settings, we evaluate performance using accuracy, micro-F1 score, as well as weighted precision and recall, with the weighted metrics chosen to account for class imbalance.

5 Results and Discussion

In Table 4, we present the performance of different models across all three tasks. Results are reported in terms of accuracy, precision, recall, and micro-F1.

5.1 Comparison with Baselines

Across all three tasks, both the SVM and pretrained language models substantially outperform the majority and random baselines. Although the majority baseline achieves relatively high accuracy in the hate severity task, this is largely due to label imbalance. Zero-shot results consistently surpass the random baseline across all tasks, yet they still fall behind the majority baseline, demonstrating their limitations without task-specific adaptation. In contrast, PEFT with LoRA yields notable gains. In particular, fine-tuned Llama3 outperforms both baselines across all three tasks, demonstrating the effectiveness of fine-tuning the model. Qwen3 with LoRA shows modest improvements, performing slightly above the majority baseline.

5.2 Classical Model and PLMs

SV trained with lexical features such as n-grams and TF-IDF, provides consistent but modest improvements over both baselines. For instance, in *type of hate* task, the SVM achieves 0.609 micro-F1 compared to 0.564 micro-F1 from the majority classifier. Similar improvements are seen in the *target of hate* task. These results suggest that traditional supervised methods can exploit word-level patterns that naive baselines miss, making them moderately effective for detecting stereotypical hate expressions. However, the performance changes when

⁵claude-3-5-sonnet

Model	Acc.	P.	R.	F1
Type of Hate				
Majority baseline	0.564	0.318	0.564	0.564
Random baseline	0.164	0.385	0.164	0.164
SVM	0.609	0.574	0.609	0.609
BanglaBERT	0.712	0.716	0.712	<u>0.712</u>
Zero-Shot				
Llama3	0.275	0.619	0.275	0.275
Qwen3	0.520	0.542	0.520	0.520
Fine-tuned				
Llama3	0.620	0.725	0.620	0.620
Qwen3	0.595	0.453	0.595	0.595
Severity of Hate				
Majority baseline	0.660	0.436	0.660	0.660
Random baseline	0.327	0.486	0.327	0.327
SVM	0.672	0.607	0.672	0.672
BanglaBERT	0.722	0.727	0.722	<u>0.722</u>
Zero-Shot				
Llama3	0.508	0.729	0.508	0.508
Qwen3	0.589	0.639	0.589	0.589
Fine-tuned				
Llama3	0.685	0.682	0.685	0.685
Qwen3	0.661	0.436	0.661	0.661
Target of Hate				
Majority baseline	0.597	0.357	0.597	0.597
Random baseline	0.204	0.404	0.204	0.204
SVM	0.629	0.568	0.629	0.629
BanglaBERT	0.715	0.716	0.715	<u>0.715</u>
Zero-Shot				
Llama3	0.340	0.465	0.340	0.340
Qwen3	0.434	0.508	0.434	0.434
Fine-tuned				
Llama3	0.610	0.716	0.610	0.610
Qwen3	0.598	0.470	0.598	0.598

Table 4: Performance of different models on Bangla hate speech detection across three tasks. Acc.: Accuracy, P.: Precision, R.: Recall, F1: micro-F1 score. **Bold** indicates results that surpass both baseline methods for the respective task, while Underline denotes the best overall performance across all three tasks.

used with narrow or context-dependent forms of hate speech, such as implicit derogatory references. This limitation stems from their reliance on surface-level features without deeper semantic understanding.

Leveraging pretraining on large-scale Bangla

corpora, BanglaBERT consistently outperforms all models across all tasks. These gains highlight the importance of contextual embeddings from language-specific pretrained models, which enable the system to capture semantic nuances, idiomatic expressions, and subtle markers of abusive tone.

5.3 Zero-shot Learning

Across all three tasks Llama3 and Qwen3 show mixed performance. For *type of hate*, Llama3 achieves micro-F1 score of 0.275, which is only marginally better than the random baseline (0.164), highlighting the difficulty of zero-shot classification in datasets with imbalanced labels. Qwen3 performs substantially better in the same task with an accuracy and F1 of 0.520, performing better than Llama3 and approaching the majority baseline (0.564), demonstrating that model size significantly influence zero-shot performance.

In the *severity of hate* task, Llama3 and Qwen3 achieve micro-F1 score of 0.508 and 0.589, respectively. Both models perform better than the random baseline (0.327) but could not perform better than the majority baseline (0.660), suggesting that the inherent structure of the severity task, requires nuanced understanding of language intensity, limits the effectiveness of zero-shot approaches. For the *target of hate*, the models achieve 0.340 (Llama3) and 0.434 (Qwen3) micro-F1 score, similarly perform better the random but below the majority baseline, indicating that identifying the *target of hate* often requires explicit task-specific knowledge that zero-shot models may not fully possess. Overall, zero-shot learning provides a reasonable starting point, particularly for Qwen3, but generally falls short of majority baseline, emphasizing the need for task adaptation to achieve high performance.

5.4 Fine-tuning LLMs

Fine-tuning with LoRA significantly improves performance across all three tasks. For *type of hate* task, Llama3 achieves micro-F1 score of 0.620, surpassing both zero-shot Llama3 (0.275) and Qwen3 (0.520), and also performs better than majority baseline. Fine-tuned Qwen3 achieves 0.595 micro-F1, slightly below Llama3, however, still demonstrating improvement over its zero-shot experiment. These results show that fine-tuning enables the models to better capture nuanced patterns.

In the *severity of hate* task, Llama3 achieves a micro-F1 score of 0.685, while Qwen3 obtains 0.661. Both models outperform the majority base-

line (0.660), demonstrating the effectiveness of task-specific adaptation in assessing the intensity of hateful content. Similarly, in the *target of hate* task, Llama3 reaches a micro-F1 score of 0.610, with Qwen3 achieving 0.598, marking a clear improvement over zero-shot performance. These results demonstrate that LoRA fine-tuning enables the models to capture subtle contextual cues that indicate the intended target of hate. Overall, LoRA fine-tuning effectively transforms pre-trained LLMs into task-aware classifiers, narrowing the performance gap with classical models like SVM and pretrained models such as BanglaBERT, while providing a scalable approach for hate speech detection in low-resource languages.

5.5 Findings

Does language-specific pretraining improve Bangla hate-speech classification across tasks? *BanglaBERT* achieves the highest scores on all three tasks, indicating that pretraining on linguistically and culturally relevant Bangla data is most effective, especially for fine-grained distinctions such as severity and target.

Are zero-shot LLMs sufficient, or is task-specific fine-tuning required? Zero-shot approaches are insufficient for Bangla hate speech. Task-specific fine-tuning substantially improves LLMs performance; fine-tuned *Llama3* is a promising alternative, whereas *Qwen3* exhibits weaker gains, suggesting differences in pretraining data and alignment.

How do fine-tuned LLMs compare to monolingual PLMs? Fine-tuned LLMs performs reasonably, however, do not surpass *BanglaBERT*. This shows that the continued importance of language-specific pretraining for reliable detection in low-resource settings.

What dataset properties most affect evaluation? Class imbalance inflates baseline performance (e.g., majority-class predictions, particularly for the severity task). Robust evaluation should report macro-F1 and per-class metrics and consider stratified splits.

Does task-specific training help uniformly across tasks? Task-specific training improves performance on all three tasks, with the largest practical benefits for nuanced categories (e.g., severity levels and target groups), where generic zero-shot models struggle.

6 Conclusions and Future Work

In this study, we present *BanglaMultiHate*, a Bangla hate-speech dataset that is among the largest manually annotated corpora in a multi-task setting. The dataset comprises approximately 51K instances spanning 19 topics and 120 sub-topics. To demonstrate its utility, we conduct comprehensive experiments comparing classical approaches, pretrained language models (PLMs), and large language models (LLMs). Our findings underscore the importance of language-specific pretraining as well as task-specific fine-tuning for robust performance. As future work, we plan to extend the dataset with reasoning annotations to support task-level interpretability and explanation.

7 Limitations

This study has several limitations. First, our dataset is collected from YouTube comments, which contain examples that may be disturbing or offensive to readers. During the annotation process, annotators were explicitly cautioned about this content and provided with appropriate warnings. Second, from a modeling perspective, the dataset is highly imbalanced across classes, which may affect both training stability and performance evaluation. Addressing these issues, for example, through data augmentation, re-sampling strategies, or collecting additional underrepresented examples, remains an important direction for future work.

Ethics and Broader Impact

Our dataset consists solely of comments and does not include any personally identifiable user information, thereby posing no direct privacy risks. Nonetheless, it is important to acknowledge that annotation is inherently subjective, which can introduce biases into the dataset. To mitigate this, we designed a clear annotation schema and provided detailed guidelines to annotators, aiming to ensure greater consistency and reliability. However, we encourage researchers and practitioners to remain careful of these limitations when using the dataset for model development or further studies.

Despite these issues, the dataset holds significant potential for positive societal impact. Models trained on *BanglaMultiHate* can support social media platforms in identifying and moderating harmful content, thereby contributing to healthier online discourse.

References

- Abdullah Al Maruf, Ahmad Jainul Abidin, Md Mahmudul Haque, Zakaria Masud Jiyad, Aditi Golder, Raaid Alubady, and Zeyar Aung. 2024. Hate speech detection in the bengali language: a comprehensive survey. *Journal of Big Data*, 11(1):97.
- Aish Albladi, Minarul Islam, Amit Das, Maryam Bigonah, Zheng Zhang, Fatemeh Jamshidi, Mostafa Rahgouy, Nilanjana Raychawdhary, Daniela Marghitu, and Cheryl Seals. 2025. Hate speech detection using large language models: A comprehensive review. *IEEE Access*.
- Bharath B and S. Ajith A. 2021. [SSNCSE_NLP@DravidianLangTech-EACL2021: Offensive language identification on multilingual code mixing text](#). In *Proceedings of the First Workshop on Speech and Language Technologies for Dravidian Languages, DravidianLangTech*, pages 313–318. 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*, pages 1318–1327, Seattle, United States. Association for Computational Linguistics.
- Bharathi Raja Chakravarthi, Ruba Priyadarshini, Vigneshwaran Muralidaran, Navya Jose, Shardul Suryawanshi, Elizabeth Sherly, and John P McCrae. 2022. Dravidiancodemix: Sentiment analysis and offensive language identification dataset for dravidian languages in code-mixed text. *Language Resources and Evaluation*, 56(3):765–806.
- Gheorghe Comanici, Eric Bieber, Mike Schaeckermann, Ice Pasupat, Naveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. 2025. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*.
- Mithun Das, Somnath Banerjee, Punyajoy Saha, and Animesh Mukherjee. 2022a. [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)*, pages 286–296, Online only. Association for Computational Linguistics.
- Mithun Das, Somnath Banerjee, Punyajoy Saha, and Animesh Mukherjee. 2022b. [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)*,

- pages 286–296, Online only. 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*, NAACL-HLT '19, Minneapolis, Minnesota, USA.
- Mouadh Djandji, Freddy Baly, Wissam Antoun, and Hady Hajj. 2020. [Multi-task learning using arabert for offensive language detection](#). In *Proceedings of the 4th Workshop on Open-Source Arabic Corpora and Processing Tools, with a Shared Task on Offensive Language Detection (OSACT4)*, pages 97–101. European Language Resource Association.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv e-prints*, pages arXiv–2407.
- Azizul Hakim Fayaz, MD. Shorif 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](#).
- Paula Fortuna and Sérgio Nunes. 2018. A survey on automatic detection of hate speech in text. *Acm Computing Surveys (Csur)*, 51(4):1–30.
- Katharine Gelber. 2021. Differentiating hate speech: a systemic discrimination approach. *Critical Review of International Social and Political Philosophy*.
- Vikram Gupta, Sumegh Roychowdhury, Mithun Das, Somnath Banerjee, Punyajoy Saha, Binny Mathew, Hastagiri Prakash Vanchinathan, and Animesh Mukherjee. 2022. Macd: Multilingual abusive comment detection at scale for indic languages. In *36th Conference on Neural Information Processing Systems (NeurIPS 2022) Track on Datasets and Benchmarks*.
- Fabiha Haider, Fariha Tanjim Shifat, Md Farhan Ishmam, Deeparghya Dutta Barua, Md Sakib Ul Rahman Sourove, Md Fahim, and Md Farhad Alam. 2024. Banth: A multi-label hate speech detection dataset for transliterated bangla. *arXiv preprint arXiv:2410.13281*.
- Fabiha Haider, Fariha Tanjim Shifat, Md Farhan Ishmam, Md Sakib Ul Rahman Sourove, 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*, pages 7217–7236, Albuquerque, New Mexico. Association for Computational Linguistics.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. Lora: Low-rank adaptation of large language models. In *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net.
- Muhammad Okky Ibrohim and Indra Budi. 2019. [Multi-label hate speech and abusive language detection in Indonesian Twitter](#). In *Proceedings of the Third Workshop on Abusive Language Online*, pages 46–57, Florence, Italy. Association for Computational Linguistics.
- Md Saroar Jahan, Mainul Haque, Nabil Arhab, and Mourad Oussalah. 2022. [BanglaHateBERT: BERT for abusive language detection in Bengali](#). In *Proceedings of the Second International Workshop on Resources and Techniques for User Information in Abusive Language Analysis*, pages 8–15, Marseille, France. European Language Resources Association.
- Aiqi Jiang and Arkaitz Zubiaga. 2024. Cross-lingual offensive language detection: A systematic review of datasets, transfer approaches and challenges. *arXiv preprint arXiv:2401.09244*.
- A. J. Keya, M. M. Kabir, N. J. Shammey, M. F. Mridha, M. R. Islam, and Y. Watanobe. 2023. [G-bert: An efficient method for identifying hate speech in bengali texts on social media](#). *IEEE Access*, 11:79697–79709.
- Douwe Kiela, Hamed Firooz, Aravind Mohan, Vedanuj Goswami, Amanpreet Singh, Pratik Ringshia, and Davide Testuggine. 2020. [The hateful memes challenge: Detecting hate speech in multimodal memes](#). In *Advances in Neural Information Processing Systems (NeurIPS) 33*.
- J. Richard Landis and Gary G. Koch. 1977. The measurement of observer agreement for categorical data. *Biometrics*, 33(1).
- Md. Firoj Mridha, Md. Abul Hasnat Wadud, Md. Abdul Hamid, Md. Mostafa Monowar, Md. Abdullah-Al-Wadud, and Atif Alamri. 2021. [L-boost: Identifying offensive texts from social media post in bengali](#). *IEEE Access*, 9:164681–164699.
- Anze Pelicon, Raghav Shekhar, Matjaz Martinc, Blaž Škrlj, Matthew Purver, and Simon Pollak. 2021. [Zero-shot cross-lingual content filtering: Offensive language and hate speech detection](#). In *Proceedings of the EACL Hackashop on News Media Content Analysis and Automated Report Generation*, pages 30–34. Association for Computational Linguistics.
- Paloma Piot and Javier Parapar. 2025. Towards efficient and explainable hate speech detection via model distillation. In *European Conference on Information Retrieval*, pages 376–392. Springer.
- John Platt. 1998. [Fast training of support vector machines using sequential minimal optimization](#). In *Advances in Kernel Methods - Support Vector Learning*. MIT Press.

- 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)*, pages 1–6.
- 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*, pages 5153–5162, Marseille, France. European Language Resources Association.
- Nauros Romim, Mosahed Ahmed, Hriteshwar Talukder, and Md Saiful Islam. Hate speech detection in the bengali language: A dataset and its baseline evaluation. *IJCACI 2020*, page 457.
- Sara Rosenthal, Noura Farra, and Preslav Nakov. 2017. Semeval-2017 task 4: Sentiment analysis in twitter. In *Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017)*.
- Happy Khairunnisa Sariyanto, Diclehan Ulucan, Oguzhan Ulucan, and Marc Ebner. 2025. Towards explainable hate speech detection. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 12883–12893.
- S. Saumya, A. Kumar, and J. P. Singh. 2021a. [Offensive language identification in dravidian code mixed social media text](#). In *Proceedings of the First Workshop on Speech and Language Technologies for Dravidian Languages, DravidianLangTech*, pages 36–45. Association for Computational Linguistics.
- Sunil Saumya, Abhinav Kumar, and Jyoti Prakash Singh. 2021b. [Offensive language identification in Dravidian code mixed social media text](#). In *Proceedings of the First Workshop on Speech and Language Technologies for Dravidian Languages*, pages 36–45, Kyiv. Association for Computational Linguistics.
- Salim Sazzed. 2021. [Abusive content detection in transliterated bengali-english social media corpus](#). In *Proceedings of the Fifth Workshop on Computational Approaches to Linguistic Code-Switching*, pages 125–130, Online. Association for Computational Linguistics.
- Konstantinos Sechidis, Grigorios Tsoumakas, and Ioannis Vlahavas. 2011. On the stratification of multi-label data. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 145–158. Springer.
- Omar Sharif, Eftekhar Hossain, and Mohammed Moshuiul Hoque. 2021. [Nlp-cuet@dravidianlangtech-eacl2021: Offensive language detection from multilingual code-mixed text using transformers](#). In *Proceedings of the First Workshop on Speech and Language Technologies for Dravidian Languages, DravidianLangTech*, pages 255–261, Kyiv. Association for Computational Linguistics.
- Deepawali Sharma, Tanusree Nath, Vedika Gupta, and Vivek Kumar Singh. 2025. Hate speech detection research in south asian languages: a survey of tasks, datasets and methods. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 24(3):1–44.
- Joseph B Walther. 2022. Social media and online hate. *Current Opinion in Psychology*, 45:101298.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, et al. 2020. Transformers: State-of-the-art natural language processing. In *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*, pages 38–45.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, et al. 2025. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*.
- Yongjin Yang, Joonkee Kim, Yujin Kim, Namgyu Ho, James Thorne, and Se-Young Yun. 2023. [HARE: Explainable hate speech detection with step-by-step reasoning](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 5490–5505, Singapore. Association for Computational Linguistics.
- Anwar Hossain Zahid, Monoshi Kumar Roy, and Swarna Das. 2025. Evaluation of hate speech detection using large language models and geographical contextualization. *arXiv preprint arXiv:2502.19612*.

Appendix

A Detailed Annotation Guideline

A.1 Definitions

The primary goal of this annotation task is to categorize YouTube comments in the Bangla language into specific categories based on the nature and severity of hate speech they contain, as well as identifying the target of such speech.

Type of Hate Categorization This annotation task involves having annotators annotate Bangla text samples according to the type of hate expressed. Each text is categorized into one of six classes: *Abusive*, *Sexism*, *Religious Hate*, *Political Hate*, *Profane*, or *None*. The goal is to capture the specific nature of hateful content, enabling models to distinguish between different forms of hate speech and non-hateful content.

- **Abusive:** Comments that are directly insulting, intending to belittle or harm someone’s dignity. For example, (*English: You are completely useless*). This comment degrades someone by calling them utterly useless.
- **Political Hate:** Comments that display hostility towards political beliefs, parties, or figures. For example, (*English: All political leaders are thieves*).
- **Profane:** Comments that use swear words or vulgar language, intended to shock or offend without targeting anyone specifically. For example, (*English: Son of a pig, you got guts*). This comment uses profanity to express frustration.
- **Religious Hate:** Comments targeting individuals or groups based on their religion or religious beliefs. For example, (*English: All people in this religion are bad*). This comment generalizes a whole religion as bad.
- **Sexism:** Comments that discriminate or belittle someone based on their gender, often reflecting stereotypes. For example, (*Women should cook only*). This comment reinforces the stereotype expression.
- **None:** Comments that do not exhibit hate or negativity, including neutral or positive comments. For example, (*English: I saw the movie today*). This comment do not reflect any hate content.

Severity of Hate This annotation task involves having annotators label Bangla text samples according to the intensity of hateful content expressed. The severity is categorized into three levels: *Severe*, *Mild*, and *Little to None*. This classification scheme is designed to capture the varying degrees of harmfulness in hate speech.

- **Severe:** Comments that contain threats, extreme prejudice, or are highly offensive. For example, (*English: Get out of the house, I’ll take care of a beast like you*).
- **Mild:** Comments that are derogatory or mildly offensive but do not contain threats. For example, (*English: You are a traitor, you will burn in hell as soon as you die*).

- **Little to None:** Comments that are slightly negative, ambiguous, or completely neutral or positive. For example, (*English: Hanif transport should be stopped*).

Target of Hate This annotation task involves having annotators label Bangla text samples based on the intended target of the hateful expression. The labels are divided into five categories: *Community*, *Individual*, *Organization*, *Society*, and *None*. This categorization aims to capture whether the hate speech is directed at a specific person, a collective group, broader societal structures, or institutions, while also accounting for instances where no explicit target is present.

- **Community:** Comments against a specific racial, ethnic, gender, or religious group. For example, (*English: People from this community are not trustworthy*).
- **Individual:** Comments targeting a specific person, either by name or implication. For example, (*English: Sheikh Hasina you are thief, the thief’s mother has a loud voice*).
- **Organization:** Comments aimed at specific companies, governmental bodies, or any formal group. For example, (*English: Somoy television seems to be government’s*).
- **Society:** Comments that critique societal norms, values, or general community practices. For example, (*English: Italians do not get food, they will increase birth rate again*).
- **None:** Comments that are ambiguous, or completely neutral or positive. For example, (*English: It should be like this: it’s hard to bring politics into any game*).

B Prompts

We provide the instructions used to generate prompts for all three tasks in Listings 1, 2, and 3. Additionally, the prompts employed for zero-shot learning, fine-tuning, and inference are presented in Listing 4.

We are creating an English instruction-following dataset for Type of Hate hate speech detection. Read the given text carefully and choose the most appropriate label for the task from the label lists.

Split	Task	Label	Word Length Bins					
			<=10	11-20	21-30	31-40	41-50	51+
Train	Type of Hate	Abusive	4374	2438	801	311	115	173
		Political Hate	1614	1415	625	259	139	175
		Profane	1329	624	214	72	45	47
		Religious Hate	281	233	81	42	19	20
		Sexism	57	36	14	9	0	6
		None	12624	4814	1353	508	265	390
	Severity of Hate	Little to None	14433	5840	1699	683	336	498
		Mild	3207	2176	825	307	153	185
		Severe	2639	1544	564	211	94	128
	Target of Hate	Community	1131	888	336	134	62	84
		Individual	3146	1552	539	202	89	118
		Organization	1755	1259	470	175	86	101
		Society	951	694	293	124	58	85
		None	13296	5167	1450	566	288	423
Dev	Type of Hate	Abusive	599	315	121	37	16	25
		Political Hate	212	197	92	33	16	24
		Profane	190	104	26	6	7	9
		Religious Hate	27	35	9	0	4	3
		Sexism	11	7	0	1	0	0
		None	1809	717	185	87	42	58
	Severity of Hate	Little to None	2079	864	239	106	54	75
		Mild	413	300	118	30	23	25
		Severe	356	211	76	28	8	19
	Target of Hate	Community	141	111	49	18	6	13
		Individual	429	203	76	19	12	16
		Organization	261	196	78	23	11	15
		Society	124	92	33	13	11	10
		None	1893	773	197	91	45	65
Test	Type of Hate	Abusive	1224	711	228	68	40	41
		Political Hate	434	431	183	91	23	58
		Profane	371	207	79	32	9	11
		Religious Hate	73	59	28	12	4	3
		Sexism	17	10	1	1	0	0
		None	3645	1390	373	159	59	125
	Severity of Hate	Little to None	4153	1678	484	204	73	145
		Mild	885	680	243	96	38	59
		Severe	726	450	165	63	24	34
	Target of Hate	Community	329	242	106	46	16	20
		Individual	861	468	136	54	25	27
		Organization	494	399	142	59	20	38
		Society	260	206	91	36	11	21
		None	3820	1493	417	168	63	132

Table 5: Detailed class label distribution with word length bin count.

For the 'Type of Hate' task, the labels are 'Abusive', 'Sexism', 'Religious Hate', 'Political Hate', 'Profane', and 'None'. Select

only one correct label for each task based on the information provided in the text and return your response in the following json format.

```
{"type_of_hate": "Abusive"}
```

Write 10 very diverse and concise English instructions. Only return the instructions without additional text. Do not generate additional text.

Return the instructions in a list format as follows.
['sent1', 'sent2']

Listing 1: Prompt for generating instructions for Type of Hate task.

We are creating an English instruction-following dataset for Hate Severity hate speech detection. Here is an example instruction:
Read the given text carefully and choose the most appropriate label for the task from the label lists.
For the 'Hate Severity' task, the labels are 'Little to None', 'Mild', and 'Severe'. Select only one correct label for each task based on the information provided in the text and return your response in the following json format.
{"severity_of_hate": "Mild"}

Write 10 very diverse and concise English instructions. Only return the instructions without additional text. Do not generate additional text.

Return the instructions in a list format as follows.
['sent1', 'sent2']

Listing 2: Prompt for generating instructions for Severity of Hate task.

We are creating an English instruction-following dataset for the Target of Hate task of hate speech detection. Here is an example instruction:
Read the given text carefully and choose the most appropriate label for the task from the label lists."
For the 'Target of Hate' task, the labels are 'Individuals', 'Organizations', 'Communities', 'Society', and 'None'. Select only one correct label for the task based on the information provided in the text and return your response in the following json format.
{"type_of_hate": "Society"}

Write 10 very diverse and concise English instructions. Only return the instructions without additional text. Do not generate additional text.

Return the instructions in a list format as follows.
['sent1', 'sent2']

Listing 3: Prompt for generating instructions for Target of Hate task.

You are a Bangla AI assistant specialized in the hate speech detection task. Your task is to identify the correct label for the task.

Read the text and assign the correct labels for the type of hate, severity of hate, and target of hate. Return only the answer without any explanation, justification, or additional text. For the 'Type of Hate' task, the labels are 'Abusive', 'Sexism', 'Religious Hate', 'Political Hate', 'Profane', and 'None'. For the 'Severity of Hate' task, the labels are 'Little to None', 'Mild', and 'Severe'. And for the 'Target of Hate' task, the labels are 'Individuals', 'Organizations', 'Communities', 'Society', and 'None'. Select only one correct label for each task based on the information provided in the text and return your response in the following JSON format.

```
{  
  "type_of_hate": "Abusive",  
  "severity_of_hate": "Mild",  
  "target_of_hate": "Society"  
}
```

If you select the 'None' label for the 'Type of Hate' task, the labels for the 'Severity of Hate' and 'Target of Hate' tasks would be 'Little to None' and 'None'.

Listing 4: Sample Prompt for zero-shot learning, model fine-tuning, and inference.

C Data Release

The *BanglaMultiHate* dataset will be released under the CC BY-NC-SA 4.0 – Creative Commons Attribution 4.0 International License: <https://creativecommons.org/licenses/by-nc-sa/4.0/>.