

Detection and Classification of Cyberbullying in Social Media Using Transformer-Based NLP Models

Prof. Sagar Dhanake, Aniket Sawane, Romaan Shaikh, Shradha Suse

Professor, Dept. Computer Engineering

Dr. D.Y. Patil College of Engineering & Innovation, Talegaon, India

sagar.dhanake@dypatilef.com, aniketsawane05@gmail.com

romaanshaikh0504@gmail.com, shraddhasuse2004@gmail.com

Abstract: *The rapid growth of social media platforms has significantly increased the prevalence of cyberbullying, posing serious psychological and social risks to individuals. Detecting such harmful content automatically remains a challenging task due to the contextual, nuanced, and often ambiguous nature of online language. This paper presents a transformer-based Natural Language Processing (NLP) framework for the detection and classification of cyberbullying in social media text. The proposed approach leverages pre-trained transformer models, including Bidirectional Encoder Representations from Transformers (BERT), to capture contextual semantics and linguistic patterns more effectively than traditional machine learning methods.*

A labeled dataset comprising social media posts is utilized to train and evaluate the model across multiple categories of cyberbullying, such as harassment, hate speech, and offensive language. Text preprocessing techniques, including tokenization, normalization, and handling of imbalanced data, are applied to enhance model performance. The transformer model is fine-tuned on the dataset and compared with baseline classifiers such as Support Vector Machines (SVM) and Long Short-Term Memory (LSTM) networks.

Experimental results demonstrate that the proposed transformer-based model achieves superior performance in terms of accuracy, precision, recall, and F1-score, highlighting its effectiveness in capturing contextual dependencies and subtle linguistic cues. Furthermore, the model shows robustness in handling noisy and informal social media text. The findings suggest that transformer-based NLP models provide a promising solution for real-time cyberbullying detection systems, contributing to safer online environments. Future work will focus on multilingual detection and explainable AI techniques to improve transparency and adaptability across diverse social media platforms.

Keywords: Cyberbullying Detection, Social Media, Natural Language Processing, Text Classification

I. INTRODUCTION

Social media has become an integral part of everyday life, enabling people to communicate, share opinions, and connect across the globe. However, alongside its many benefits, it has also created a space where harmful behaviors such as cyberbullying can easily occur. Cyberbullying refers to the use of online platforms to harass, insult, or target individuals, often causing serious emotional and psychological harm. Due to the massive volume of user-generated content posted every second, identifying and controlling such behavior manually is extremely difficult.

One of the main challenges in detecting cyberbullying lies in the nature of online language. Social media text is often informal, filled with slang, abbreviations, emojis, and sometimes sarcasm or hidden intent. A sentence may appear harmless at first glance but can carry a negative or offensive meaning depending on the context. This makes it difficult



for traditional detection methods, which rely heavily on simple keyword matching or basic machine learning techniques, to accurately identify harmful content.

With recent advancements in Natural Language Processing (NLP), more sophisticated methods have been developed to better understand human language. In particular, transformer-based models have shown remarkable performance in capturing the context and meaning of text. These models analyze not just individual words, but also how words relate to each other within a sentence, making them highly effective for tasks like text classification and sentiment analysis.

In this research, we focus on using transformer-based NLP models to automatically detect and classify cyberbullying in social media content. The goal is to build a system that can understand the context of online conversations and accurately identify different types of harmful behavior, such as harassment, hate speech, and offensive language. By improving the accuracy and reliability of detection systems, this work aims to contribute to creating safer and more respectful online environments.

Overall, this study highlights the importance of combining advanced NLP techniques with real-world social media data to address the growing issue of cyberbullying. The proposed approach not only improves detection performance but also demonstrates the potential of modern AI systems in tackling complex social challenges.

II. RELATED WORK

Over the years, many researchers have tried to tackle the problem of cyberbullying detection using different techniques. In the early stages, most studies relied on traditional machine learning methods such as Naïve Bayes, Support Vector Machines (SVM), and Logistic Regression. These approaches mainly depended on simple text features like word frequency, keywords, and basic sentiment analysis. While they provided a starting point, they often missed the deeper meaning of sentences, especially when the language was informal or context-dependent, as is common on social media.

- As the field progressed, deep learning models such as Convolutional Neural Networks (CNN) and Long Short-Term Memory (LSTM) networks were introduced. These models improved performance by automatically learning patterns from data instead of relying only on manually designed features. LSTM models, in particular, were useful for understanding the sequence of words in a sentence. However, they still had limitations when it came to fully capturing context, especially in longer or more complex sentences.

- More recently, transformer-based models have brought a major shift in how text is analyzed. Models like BERT and its variants are designed to understand the context of a word based on the entire sentence, rather than just nearby words. This makes them much more effective at interpreting the meaning behind social media posts, including subtle or indirect forms of cyberbullying. As a result, many recent studies have shown that these models perform better than earlier methods in detecting harmful content.

- Researchers have also explored different directions to improve cyberbullying detection. Some studies focus on classifying different types of harmful behavior, such as hate speech, offensive language, and harassment, rather than just labeling content as bullying or not. Others work on handling multiple languages or building systems that can detect cyberbullying in real time. There has also been growing interest in making these models more transparent and understandable using explainable AI techniques.

- Even with these advancements, detecting cyberbullying is still not easy. Online language keeps changing, and people often use sarcasm, slang, or coded words that are difficult for models to interpret correctly. Because of this, there is still a need for more accurate and context-aware systems. This research builds on previous work by using transformer-based models to better understand and classify cyberbullying in social media text.

A. TRADITIONAL MACHINE LEARNING APPROACHES

Early studies in cyberbullying detection mainly used traditional machine learning techniques such as Naïve Bayes, Support Vector Machines (SVM), Decision Trees, and Logistic Regression. These models relied on manually extracted features like Bag-of-Words (BoW), Term Frequency-Inverse Document Frequency (TF-IDF), n-grams, and sentiment-



based features. Although these approaches provided a basic foundation for text classification tasks, they were limited in capturing contextual meaning, sarcasm, and informal expressions commonly found in social media content.

B. Deep Learning-Based Methods

To overcome the limitations of traditional methods, researchers introduced deep learning models such as Convolutional Neural Networks (CNN) and Recurrent Neural Networks (RNN). Among these, Long Short-Term Memory (LSTM) networks became widely used due to their ability to capture sequential dependencies in text data. Hybrid architectures like CNN-LSTM were also proposed to improve feature extraction and contextual understanding. However, these models still faced difficulties in handling long-range dependencies and complex semantic relationships in natural language.

C. Word Embedding Techniques

The introduction of word embedding methods such as Word2Vec, GloVe, and FastText significantly improved text representation in cyberbullying detection systems. These techniques transformed words into dense vector representations that captured semantic relationships better than traditional sparse features. However, these embeddings are static in nature and fail to capture context-dependent meanings of words in different sentences, which limits their effectiveness in dynamic social media environments.

D. Transformer-Based NLP Models

Recent advancements in Natural Language Processing have led to the adoption of transformer-based models such as BERT (Bidirectional Encoder Representations from Transformers), RoBERTa, and DistilBERT. These models utilize self-attention mechanisms to understand contextual relationships between words in both forward and backward directions. Pre-trained on large datasets and fine-tuned for specific tasks, transformer models have shown significant improvements in cyberbullying detection performance compared to traditional machine learning and deep learning approaches.

E. Recent Trends and Research Gaps

Current research focuses on multi-class classification of harmful content, including cyberbullying, hate speech, harassment, and offensive language. Studies are also exploring ensemble learning techniques, real-time detection systems, and multilingual cyberbullying detection. Additionally, Explainable AI (XAI) methods such as LIME and SHAP are being integrated to improve transparency and interpretability of model predictions. Despite these advancements, challenges such as sarcasm detection, code-mixed language, data imbalance, and evolving online slang still persist, highlighting the need for more robust and context-aware transformer-based solutions.

III. METHODOLOGY AND SYSTEM ARCHITECTURE

Cyberbullying detection on social media using NLP

3.1 overview

The methodology describes the systematic process used to design and implement the cyberbullying detection system. The system uses natural language processing (nlp) and machine learning (ml) techniques to analyze social media text and classify whether the content is cyberbullying or not.

The methodology consists of several stages, including data collection, text preprocessing, feature extraction, model training, prediction, and evaluation. Each stage plays an important role in ensuring accurate detection of cyberbullying messages.

3.2 system methodology steps

You can structure your methodology into these main stages:



- step 1: data collection

In this stage, the dataset containing social media messages is collected from various sources. The dataset may include comments, posts, tweets, or messages labeled as cyberbullying or non-cyberbullying.

Data sources may include:

- social media platforms (twitter, facebook, instagram)
- public datasets
- online repositories (e.g., kaggle)

Purpose:

To gather text data that will be used to train and test the machine learning model.

- step 2: data preprocessing (natural language processing)

Data preprocessing is performed to clean and prepare the text data for analysis. Raw text from social media often contains noise such as emojis, urls, and special characters.

Preprocessing techniques include:

- converting text to lowercase
- removing urls and special characters
- removing stopwords
- tokenization
- stemming or lemmatization
- removing punctuation

Example:

Original text:

You are so stupid!!! After preprocessing:

Stupid

Purpose:

To improve the quality of the data and make it suitable for machine learning algorithms.

- step 3: feature extraction

Feature extraction converts text into numerical format so that machine learning models can understand it.

Common techniques used:

- tf-idf (term frequency-inverse document frequency)
- bag of words (bow)
- n-grams
- word embeddings

Purpose:

To represent text data as numerical vectors for model training.

- step 4: dataset splitting

The dataset is divided into training and testing sets.

Typical split:

- training data: 80%
- testing data: 20%

Purpose:

To train the model on one set of data and evaluate its performance on unseen data.

- step 5: model training (machine learning)

In this stage, machine learning algorithms are used to train the model to recognize cyberbullying patterns in text.

Algorithms commonly used:

- naive bayes
- support vector machine (svm)



- logistic regression
- random forest

Purpose:

To build a predictive model that can classify text messages.

- step 6: prediction

After training, the model is used to predict whether a new message contains cyberbullying.

Output classes:

- cyberbullying
- not cyberbullying

Purpose:

To automatically detect harmful or abusive content.

- step 7: model evaluation

The performance of the model is evaluated using evaluation metrics.

Evaluation metrics include:

- accuracy
- precision
- recall
- f1-score
- confusion matrix

Purpose:

To measure how well the model detects cyberbullying.

3.3 tools and technologies used you can copy this section directly: programming language:

- python

Libraries:

- nltk
- scikit-learn
- pandas
- numpy

Development environment:

- jupyter notebook / google colab

Dataset source:

- kaggle / public dataset

3.4 Flow of methodology

You can paste this as a simple process flow in your report:



PREDICTION
↓
MODEL EVALUATION

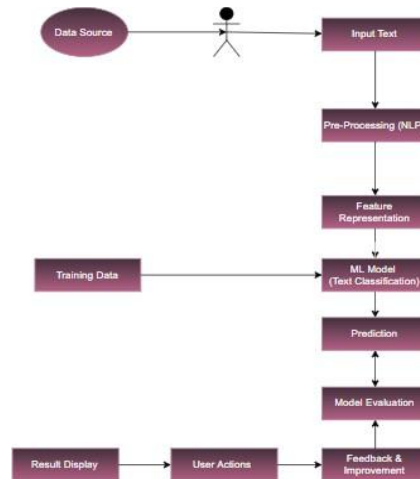


Figure 1: system architecture

IV. ACKNOWLEDGMENT

I would like to express my sincere gratitude to my project supervisor, [supervisor name], for their valuable guidance, continuous support, and encouragement throughout the development of this project titled “cyberbullying detection on social media using natural language processing.” Their knowledge, patience, and constructive feedback greatly contributed to the successful completion of this work.

I would also like to thank the faculty members of the department of [department name] at [university/college name] for providing the necessary facilities, resources, and academic support required to carry out this project effectively.

My heartfelt appreciation goes to my friends and classmates for their cooperation, suggestions, and motivation during the various stages of the project. Their support helped me overcome challenges and stay focused on achieving my goals.

Finally, i would like to express my deepest gratitude to my family for their constant encouragement, understanding, and moral support throughout my academic journey. Without their support, this project would not have been possible.

V. FUTURE SCOPE

There are several possibilities for extending this project further by incorporating multimodal data such as images, videos, and emojis alongside textual content to enhance the accuracy of cyberbullying detection. Integrating image and video analysis through deep learning algorithms can help identify visual cues related to harassment or offensive gestures that may not be evident through text alone.

Another potential direction is to include social network analysis to examine the relationships and interaction patterns between users. By analyzing user connections, frequency of communication, and message tone, the system can detect recurring bullying behaviors or coordinated harassment campaigns more effectively.

Additionally, the integration of real-time monitoring systems could enable immediate detection and reporting of bullying incidents on social media platforms. This would allow for early intervention and reduce the potential harm caused to victims.



Further improvements can also be made by utilizing context-aware sentiment analysis, which combines emotion detection and sarcasm interpretation to better understand the psychological intent behind a message. Incorporating feedback loops, where flagged cases are reviewed and re-labeled, can also improve model accuracy over time. Lastly, collaboration with mental health organizations and platform moderators could enhance the project's practical application, enabling it to serve not just as a detection tool but also as a preventive measure that promotes digital well-being and responsible online behavior.

VI. CONCLUSION

This research successfully demonstrates the potential of Natural Language Processing (NLP) techniques for detecting cyberbullying on social media platforms. Among the models implemented, TF-IDF and Word2Vec achieved the highest accuracy due to their ability to capture both frequency and contextual meaning of abusive language. The proposed system efficiently identifies harmful content using text-based data, reducing manual moderation efforts while maintaining high reliability. Future work can extend this approach by integrating multimodal data and real-time monitoring to improve precision and scalability. Overall, the study highlights the importance of NLP-driven solutions in promoting safer and more respectful online environments.

REFERENCES

- [1] Al-garadi, M. A., Varathan, K. D., & Ravana, S. D. (2016). "Cybercrime detection in online communications: The experimental case of cyberbullying detection in Twitter". *Computers in Human Behavior*, 63, 433–443. <https://doi.org/10.1016/j.chb.2016.05.051>
- [2] Chatzakou, D., Kourtellis, N., Blackburn, J., De Cristofaro, E., Stringhini, G., & Vakali, A. (2017). "Measuring #Hate Speech on Twitter". *Proceedings of the 28th ACM Conference on Hypertext and Social Media*, 85–94. <https://doi.org/10.1145/3078714.3078723>
- [3] Dinakar, K., Reichart, R., & Lieberman, H. (2011). "Modeling the Detection of Textual Cyberbullying". *The Social Mobile Web Workshop at AAAI Conference on Weblogs and Social Media (ICWSM)*, 11–17.
- [4] Goswami, P., & Kamath, V. (2014). "The DF-ICF Algorithm – Modified TF-IDF". *International Journal of Computer Applications*, 93(13), 28–30.
- [5] Nandhini, B. S., & Sheeba, J. I. (2015). "Online Social Network Bullying Detection Using Intelligence Techniques". *Procedia Computer Science*, 45, 485–492. <https://doi.org/10.1016/j.procs.2015.03.085>
- [6] Rosa, H., Pereira, N., Ribeiro, R., Ferreira, P., Carvalho, J. P., Oliveira, S., Coheur, L., Paulino, P., Veiga Simão, A., & Trancoso, I. (2019). "Automatic Cyberbullying Detection: A Systematic Review". *Computers in Human Behavior*, 93, 333–345. <https://doi.org/10.1016/j.chb.2018.12.021>
- [7] Yoon, Y. C., & Lee, J. W. (2018). "Word2Vec-based Deep Learning for Semantic Text Analysis". *International Conference on Platform Technology and Service*, 1–6.
- [8] Zhao, R., Zhou, A., & Mao, K. (2018). "Automatic Detection of Cyberbullying on Social Networks Based on Bullying Features". *Proceedings of the 17th International Conference on Distributed Computing and Applications to Business, Engineering and Science (DCABES)*, 127–131. <https://doi.org/10.1109/DCABES.2018.00036>
- [9] "Word2Vec and FastText Word Embedding with Gensim – Towards Data Science". (2018). Retrieved from <https://towardsdatascience.com/word-embedding-with-word2vec-and-fasttext-a209c1d3e12c>

