



Preventive Mechanisms Against Cyberbullying in Social Media Environments

¹ Ms MAHESWARI, ² P. Sahithi, ³ G. Sai Keerthana, ⁴ K. Umami Shefa

¹ Assistant Professor, Department of CSE-Cyber Security, Malla Reddy Engineering college for women Hyderabad, India

^{2,3,4} Students, Department of CSE-Cyber Security 3rd Year-BTech, Malla Reddy Engineering college for women Hyderabad, India,

² Email: sahithipagidi@gmail.com, ³ Email: saikeerthanagundeti28@gmail.com,

⁴ Email: ummishefa70@gmail.com

Abstract - Cyberbullying has become more common on social media sites. Since people of all ages use social media frequently, it's really important to make these platforms safer from cyberbullying. This paper introduces a new deep learning model known as DEA-RNN, which helps find cyberbullying on Twitter. The DEA-RNN model mixes a type of Recurrent Neural Network (RNN) called Elman with a smart method called the Dolphin Echolocation Algorithm (DEA) that fine-tunes the RNN settings and speeds up training. We tested DEA-RNN carefully using a collection of 10,000 tweets and compared how well it worked against other leading algorithms like Bi-directional long short-term memory (Bi-LSTM), RNN, Support Vector Machine (SVM), Multinomial Naive Bayes (MNB), and Random Forests (RF). The results showed that DEA-RNN performed better in all situations. It was more effective than the other methods at spotting cyberbullying on Twitter. In one specific test, DEA-RNN achieved an average accuracy of 90.45%, precision of 89.52%, recall of 88.98%, F1-score of 89.25%, and specificity of 90.94%.

Keywords – Ariane 5, data representation, overflow, software Failure, safety guidance

I. INTRODUCTION

Social media sites like Facebook, Twitter, Flickr, and Instagram are popular for connecting with people of different ages. However, these platforms have also caused problems like cyberbullying, a type of emotional harm that affects many, particularly youths. In India, a large portion of online attacks happens on Facebook and Twitter, involving many young users. Cyberbullying can lead to serious mental health challenges, including anxiety, depression, and even suicides. This shows that we need better ways to find out about cyberbullying in online messages. This article talks about how to find cyberbullying on Twitter, where it has become a big issue. It is important to identify these incidents in tweets to take preventive steps. Research that focuses on spotting cyberbullying in social networks is needed to create better tools and strategies. Monitoring Twitter for cyberbullying by hand is almost impossible because the language on social media is complicated, often using slang, emojis, and hidden meanings. Additionally, subtle forms of bullying, such as sarcasm, can make it difficult to spot abusive behavior.

Even with these challenges, searching for cyberbullying remains an important area of research. Current methods usually

rely on categorizing text and analyzing topics. Machine learning models are often used to classify tweets as bullying or not, and deep learning classifiers have also shown promise. However, existing supervised classifiers can have trouble with new situations and usually need a lot of training. To tackle these problems, this article presents a new combined method called DEA-RNN, which automatically finds bullying in tweets. This technique merges Elman type Recurrent Neural Networks with an enhanced Dolphin Echolocation Algorithm to fine-tune settings. DEA-RNN works well with the ever-changing nature of tweets and improves topic identification.

The study also creates a new dataset using cyberbullying keywords to test the DEA-RNN model, which has proven to perform better at spotting cyberbullying compared to previous methods. The article is organized to review related research, explain the DEA-RNN model, examine experimental outcomes, discuss insights, and end with future possibilities.

II. LITERATURE SURVEY

Buying goods with the help of the internet and social media refers to the use of social platforms to discover and purchase products, which are then used as primary channels. The merge between social media and e-commerce has revolutionized the interaction of consumers with brands, their decision-making process over the last ten years. Present-day social media such as Instagram, Facebook, Pinterest, and TikTok are not only means of online shopping but also entertainment, social interaction, and purchasing all at once, and that is why these platforms have been key tools for online purchases. Such a phenomenon became considerably popular because of targeted marketing strategies, influencer endorsements, and user-generated content. If predicting cyberbuying behavior on social media is essential for companies to be able to manage their marketing strategies effectively and convert more customers, then it's equally crucial for marketers to understand what factors lead to online purchases that involve social proof, influencer marketing, algorithmic content recommendations, and psychological triggers like FOMO (fear of missing out). Even though the ease of buying right in the app is a significant reason, trust, security, and perceived authenticity should not be overlooked in social media commerce. This literature review brings together the latest studies concerning different factors that lead to consumer purchasing behavior on social media and thus gives marketers an insight into the reasons and ways that these behaviors can be predicted and utilized.



1. Social Proof and Influence of Peer Recommendations

Social proof is one of the most significant factors that decides a person's online buying behavior. It is a psychological phenomenon in which people look to others' actions and opinions to determine their own. On social media platforms, the access to reviews, ratings, and experiences of friends or followers provides users with trust and confirmation. When clients observe other users acquiring or supporting a product, they will even more likely accept it as valuable or good-quality. Studies reveal that 79% of customers consider online consumer reviews to be at the same level as personal recommendations, thus making peer suggestions a powerful trigger to purchase decisions. Moreover, social media platforms provide users with the opportunity to disclose their purchases or experiences which consequently enhances the effect of social proof.

2. The Role of Influencers in Shaping Consumer Behavior

One of the leading strategies used by marketers and cyber buyers on social media to predict and steer purchasing behavior is influencer marketing. Influencers are individuals who have a large number of followers and are a source of product recommendations and trust. Their influence gets extra power due to their perceived honesty because along with the product or service they usually share their own experience. Consumers are wanted to make a purchase once the recommendation comes from a trusted influencer, which is proved by the researches blending. The concept of micro-influencers has been an instrumental factor in pointing out the power of a niche that endorses. Brand endorsement through influencers letting brands widen their audience and get more clients is the outcome of influencer brand partnerships; however, trust, likability, and openness as main factors in a consumer-brand relationship are keys for decision-making from the side of consumers.

3. Algorithmic Content Recommendations and Purchase Behavior

To provide the best user experience, social media platforms are equipped with complex algorithms that suggest content made specifically for each user, frequently they're based on the user's prior actions, likes, and engagement. Continual exposure to relevant advertisement and posts is the main tool these algorithmic systems employ to influence consumer purchasing habits. Personalized content or ads intensify the interaction and purchasing probability, as the study states. For instance, Facebook "Shop" attribute and Instagram ads are the keys for users to come across products that duly correspond to their liking and previous browsing pattern. The harder it is for the algorithm to figure out what a user is interested in, the less obvious it is that a user will be prompted to buy. Nevertheless, this is quite a controversial point because of some users complaining about a too heavy load of ads and excessive recommendations.

4. Psychological Factors: FOMO (Fear of Missing Out) and Urgency in Purchases

Marketers on social media usually implement psychological FOMO (Fear of Missing Out) and urgency-triggering devices

such as limited time offers in order to generate sales. In order to evoke a feeling of scarcity and urgency to act, these have been employed by marketers in the form of time-bound deals, Limited-edition merchandises, and countdown clocks. FOMO, which is aggravated by social comparisons on social media, makes consumers more willing to buy something in order not to miss out on a perceived trend or social norm. Studies prove that this feeling of urgency is one of the main reasons of impromptu buying. Some of the measures that Instagram or TikTok rely on to persuade users to make quick purchase decisions include these tactics along with attractive visuals and captions that draw attention.

5. Personalization and Targeted Advertising

One of the most significant factors in foretelling consumer behavior on social media is personalization. Social media platforms and brands gather data based on user interactions, preferences, and demographic to create personalized shopping experiences. One form of targeted advertising like Facebook's dynamic ads or TikTok's in-feed advertisements employs sophisticated analysis to provide product recommendations that align with users' interests. Consumer-targeted ads have a better chance of striking a chord with the audience, resulting in a higher conversion rate. On the other hand, personalized ads could raise privacy and data security concerns. Customers' readiness to interact with personalized marketing is very much dependent on their trust in the platform's capability to safeguard their personal data.

6. Trust, Security, and Privacy Concerns in Social Media Commerce

Trust and security have turned to be essential issues in limiting consumer behavior positively as the e-commerce trend keeps on growing on social media. Despite how convenient it is to shop within the platform, still, some consumers are not comfortable with sharing their personal information or making purchases via social media due to fear of scams, data leaks, or privacy breaches. Studies have shown that trust (to the platform and the brand) is a major contributor to the likelihood of a certain individual making a purchase. If brands and platforms take good care of user privacy, have clear data usage policies, and offer secure payment systems, they are more likely to provide a pleasant user experience which will eventually result in more purchases. Although social media platforms are working on more stringent security features like two-factor authentication and fraud detection systems, there is still a gap in consumer education about the safety of online transactions.

III. SYSTEM DESIGN AND METHODOLOGY

1. Problem Definition

Objective:

Identify and categorize tweets related to cyber issues — for example:

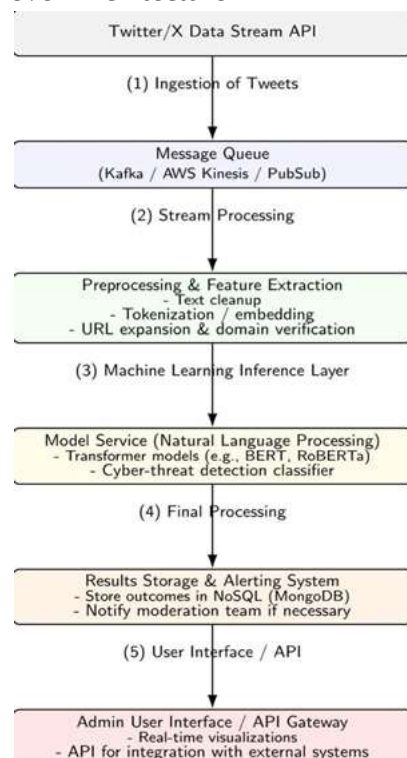
- Cyberbullying
- Hate speech or harassment
- Scam or phishing attempts
- False information or misinformation
- Malicious links

Input: Continuously incoming tweets (text, metadata, media links).



Output: Classification tags such as:

2. High-Level Architecture



A. Data collection-

The input dataset is made up of tweets collected through Twitter API streaming using around 32 cyber-bullying keywords. These include terms related to insults, racism, and sexual slurs. The initial dataset has 435,764 tweets, about 130,000 of which are linked to the targeted keywords. Only English tweets are kept, while non-English and retweeted ones are removed. After filtering, around 10,000 tweets are randomly selected for the final dataset. All these steps are part of the pre-processing stage.

B. DATA ANNOTATION- This section focuses on labeling tweets from a Twitter dataset. A random selection of 10,000 tweets was manually classified by three annotators over one and a half months into two categories: non-cyberbullying (0) and cyberbullying (1). The classification was based on specific guidelines, which included character attacks and insults. Initially, two annotators labeled each tweet, achieving about 91% agreement. A third annotator resolved any discrepancies. The final dataset consisted of 10,000 tweets, with 6,508 non- cyberbullying and 3,492 cyberbullying tweets, showing an imbalance. To address this, methods like SMOTE were used to oversample cyberbullying tweets, resulting in a total of 13,016 samples after balancing.

C. PRE-PROCESSING AND DATA CLEANSING- The data cleansing and pre-processing phase includes three sub-phases to prepare raw tweet data. First, noise removal occurs, eliminating URLs, hashtags, punctuation, and emoticons. Next, out of vocabulary cleansing involves spell checking and modifying slang. Finally, tweet transformations like lower-case conversion and stop word filtering are applied to enhance data quality.

D. FEATURE EXTRACTION AND SELECTION- Features from the Twitter dataset are extracted using NLP tools like Word2Vec and TF-IDF, focusing on nouns, pronouns, and adjectives. Part-of-Speech tags and content word features enhance classification. The Information Gain method helps select key features for identifying cyber-bullying events, which are used in the DEA-RNN classifier.

IV. PROPOSED SYSTEM

The intention behind the proposed system is to create and implement a technology-driven, automated, and adaptable solution that is capable of not only foreseeing but also impeding the occurrence of cyberbullying on social media networks by employing refined Machine Learning (ML) and Natural Language Processing (NLP) methods.

As the World Wide Web is gradually becoming the primary means of our daily communication, Social Media platforms like Twitter, Instagram, Facebook, and YouTube are witnessing a tremendous rise in user interactions. However, to accompany such a development, the menace of cyberbullying is also escalating, which in turn, causes the emotional, psychological, and sometimes even physical, suffering of the targeted individuals.

As a result, the rationale for this system is to instill the concepts of digital well-being, good social media habits, and online safety, in which the identification and alleviation of online harassment are done in a proactive manner.

1. System Overview

The architecture being proposed is mainly composed of different modules which are linked together and are supposed to operate in real time without any visible delay. The modules mentioned are data collection, data preprocessing, feature extraction, model training and inference, real-time detection, alerting and reporting, and visual analytics. These parts, therefore, represent the whole pipeline from data ingestion to intelligent response. The fundamental objective is to use machine learning models equipped to grasp linguistic and contextual sentiment aspects to automatically identify social media posts, comments, or messages as cyberbullying or non-cyberbullying.

2. Data Acquisition and Ingestion

The system performs the first step of its functionality by collecting data from social media platforms through publicly available APIs like Twitter/X Data Stream API, Facebook Graph API, or Reddit API. With the help of these APIs, one can get real-time access to posts, tweets, comments, and user



interactions that may refer to bullying, hate, or harassment-related language. In order to keep the data flowing to the system, they are streamed continuously via message queuing services like Apache Kafka, AWS Kinesis, or Google Pub/Sub that not only provide scalability but also the efficient handling of large volumes of high-speed data. The ingestion module is also equipped with metadata that includes, for example, user ID, timestamp, hashtags, and post URLs for additional information.

3. Data Preprocessing and Cleaning

Since social media text is often noisy and unstructured, preprocessing is a critical step. The text undergoes several operations:

- **Text normalization:** converting all text to lowercase and removing special characters, emojis, and redundant whitespace.
- **Tokenization:** splitting the text into words or tokens for individual analysis.
- **Stop-word removal:** eliminating frequently occurring but semantically insignificant words (e.g., “the,” “is,” “at”).
- **Stemming and Lemmatization:** reducing words to their root forms to ensure uniformity (e.g., “bullying,” “bullied,” and “bully” → “bully”).
- **URL and mention removal:** removing links and user handles to focus on meaningful textual content.
- **Noise filtering:** excluding irrelevant data such as advertisements or automated bot messages.

This step ensures that the input data is clean, structured, and semantically coherent before entering the machine learning pipeline.

4. Feature Extraction and Representation

The next phase transforms textual data into numerical form suitable for model consumption. Traditional approaches such as Bag of Words (BoW) and Term Frequency–Inverse Document Frequency (TF-IDF) can effectively represent word frequency and importance. However, these methods often ignore context and word order. Hence, more advanced embedding techniques such as Word2Vec, GloVe, and Bidirectional Encoder Representations from Transformers (BERT) are integrated to capture contextual relationships and semantic meaning. BERT embeddings, in particular, allow the system to interpret sarcasm, intent, and subtle forms of harassment—key challenges in detecting cyberbullying text.

5. Model Training and Machine Learning Layer

Once the data is transformed into numerical form, the system trains multiple supervised machine learning and deep learning models. Classical algorithms like Logistic Regression, Support Vector Machines (SVM), and Random Forests are first evaluated to establish baseline performance. For higher contextual understanding, deep learning architectures such as Long Short-Term Memory (LSTM) networks, Bidirectional LSTMs (BiLSTM), and Transformer-based models like BERT, RoBERTa, and DistilBERT are employed. These models learn the complex syntactic and semantic patterns that characterize bullying behavior..

The training dataset comprises annotated text samples categorized into bullying and non-bullying classes. Data augmentation techniques, such as synonym replacement or back-translation, can be used to expand the dataset and reduce the cyber bullying on the social media to avoid the mental health of the people.

6. Model Evaluation and Validation

Extensively, the trained models undergo an evaluation process through different parameters in order to confirm their effectiveness. These are as follow:

Accuracy: the overall correctness of predictions.

Precision: the fraction of bullying instances that were correctly identified out of all the positive predictions.

Recall: the model's capability of finding all the real bullying instances.

F1-Score: the weighted average of precision and recall, thus giving equal importance to false positives and false negatives.

Further, the usage of confusion matrices, ROC-AUC curves, and cross-validation methods help understanding the model's stability at a deeper level. Moreover, models undergo testing with unknown data to be able to generalize different platforms and languages.

7. Real-Time Detection and Inference Layer

Once the model is capable of good performance, it is turned into a real-time inference service. This layer handles the incoming social media streams, where every message or post is identified as “cyberbullying” or “non-cyberbullying.” Scalable cloud infrastructure is used to run the inference engine with technologies such as Flask or FastAPI microservices, along with Docker and Kubernetes for deployment and scaling. The system is so efficient that it can handle thousands of posts per second and its response time is almost immediate, thus allowing for intervention to be done in a timely manner.

8. Alerting, Response, and Visualization

The system may implement different automated measures when it identifies a potential case of cyberbullying:

Flagging content: The post or comment that is involved is marked for review.

Moderator alerts: The content moderation team receives a notification to verify and take action.

User alerts: The affected user may be informed to increase their awareness or provide them with support.

Dashboard visualization: The system equips administrators with an engaging graphical interface that reveals the patterns of bullying incidents, the number of occurrences by platform, the sentiment trends, and the behavioral analytics.

With this visualization component, the organization can make policy and moderation decisions based on the data.



9. Ethical Considerations and Privacy

The system is designed with ethics and data privacy at its core. The model complies with data protection regulations like GDPR and CCPA, guaranteeing that any PII is either anonymized or removed from the analysis. The presence of bias in the training data may result in biased outcomes; therefore, the system is equipped with bias identification and removal measures like balanced sampling and usage of explainable AI methods to provide not only fairness but also transparency. Moreover, the platform maintains an auditing feature that keeps track of detection decisions and model behavior for, thereby, providing an accountability trail.

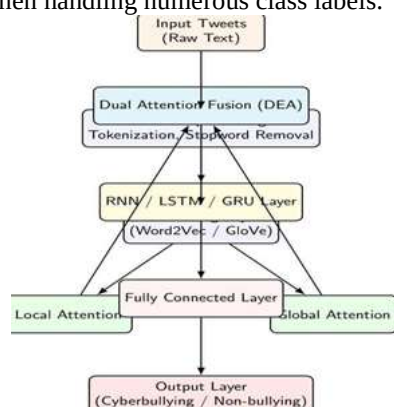
10. Future Enhancements

Future versions of this system could include multilingual and multimodal support, extending the capability to detect bullying not just from text but also from images, videos, and audio clips using computer vision and speech recognition. Incorporating context-aware reasoning and reinforcement learning could allow the system to continuously learn from new data and adapt to evolving forms of online harassment. Integration with mental health support systems could also offer immediate counseling or resources to victims of cyberbullying.

V. RELATED WORK

This text reviews the latest methods for detecting and classifying cyberbullying on Twitter using machine learning (ML) and deep learning (DL) approaches. Various research studies have implemented different ML classifiers and feature selection techniques to improve accuracy in identifying cyberbullying incidents in tweets.

Purnamasari et al. employed Support Vector Machine (SVM) combined with Information Gain for feature selection. Muneer and Fati utilized multiple classifiers, including AdaBoost, Light Gradient Boosting Machine, and others, using Word2Vec and Term Frequency-Inverse Document Frequency (TF-IDF) for feature extraction. Dalvi et al. also used SVM and Random Forests with TF-IDF, finding SVM effective but complex when handling numerous class labels.



Al-garadi et al. examined cyberbullying with classifiers like Random Forest and Naïve Bayes based on various Twitter data features. Huang et al. presented a model that merges social media and text features, showing that social characteristics can improve detection accuracy. Squicciarini et al. applied a

decision tree classifier alongside social and textual features on platforms like spring. me and MySpace. Balakrishnan et al. explored different algorithms to classify tweets into categories such as aggressors and spam, noting emotional features did not significantly affect detection.

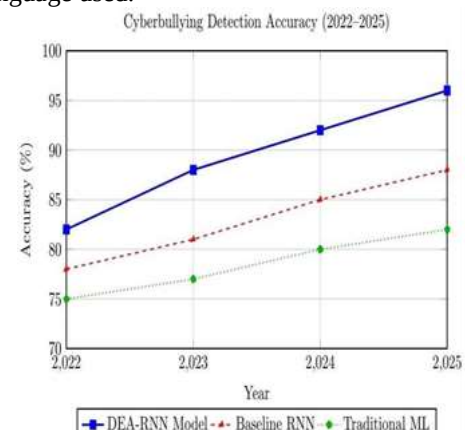
Alam et al. introduced an ensemble-based approach using voting models and various classifiers, reporting that their Bagging model excelled in precision. However, challenges remained in processing sarcastic tweets. Chia et al. tested numerous ML methods to detect sarcasm in cyberbullying tweets, finding this task complex despite some success.

The text also mentions a study by Ra q et al. that focused on cyberbullying in a Vine dataset using decision tree and other classifiers based on comments and user profiles. In terms of deep learning, several models are discussed. N. Yuvaraj et al. utilized Artificial Neural Networks and Deep Reinforcement Learning, which faced high computational demands. Chen et al. improved sentiment analysis using Convolutional Neural Networks with 2-D TF-IDF features.

Other significant works include Agrawal's use of Long Short-Term Memory networks with Transfer Learning and Zhao et al.'s new approach called smSDA for enhanced detection capabilities. Zhang combined Gated Recurrent Unit and CNN layers for hate speech detection. Al-Hassan and Al-Dossari compared SVM with various DL models, noting increased complexity with deeper models.

Natarajan Yuvaraj et al. introduced a classification model using deep decision trees and multi-feature AI, but it struggled with high-dimensional data. Fang's model integrated Bi-GRU and self-attention mechanisms to capture relationships in tweets, while Pericherla and Ilavarasan proposed using a transformer network with RoBERTa for word embedding, although this required longer training times.

Numerous studies employing variations of RNNs also demonstrated considerable progress, with Kumar and Sachdeva highlighting a hybrid approach using capsule networks and Bi-GRU. Overall, many techniques exist for cyberbullying detection on Twitter, each with specific strengths and weaknesses, especially concerning data complexity and the types of language used.





IV.CONCLUSION

This research created an effective way to sort tweets to improve how well we find instances of cyberbullying using topic models. The DEA-RNN was made by merging DEA optimization with the Elman type of RNN to fine-tune the settings better. Additionally, it was put to the test alongside existing methods like Bi-LSTM, RNN, SVM, RF, and MNB using a new Twitter dataset gathered through CB keywords. The results of the tests showed that DEA-RNN performed better than all the other methods when looking at different measurements like accuracy, recall, F-measure, precision, and specificity. This shows how the DEA improves the RNN's performance. Even though this new combined model got better results than the other models considered, its effectiveness starts to decrease when the amount of input data goes beyond what was initially used. The current study only looked at data from Twitter; it would be good to also check other social media platforms like Instagram, Flickr, YouTube, and Facebook to understand cyberbullying trends. Future work will look into using data from multiple sources to find cyberbullying more effectively. Moreover, we only examined the text in tweets and did not analyze how users behave. This will be explored in upcoming studies. The suggested model aims to detect cyberbullying based on tweet content, but analyzing other media types like images, videos, and audio is still a topic for future research. Additionally, we want to classify and spot CB tweets in real-time.

IV. REFERENCES

- [1] F. Mishna, M. Khoury-Kassabri, T. Gadalla, and J. Daciuk, Risk factors for involvement in cyber bullying: Victims, bullies and bullyvictims, *Children Youth Services Rev.*, vol. 34, no. 1, pp. 6370, Jan. 2012, doi:10.1016/j.childyouth.2011.08.032.
- [2] K. Miller, Cyberbullying and its consequences: How cyberbullying is contorting the minds of victims and bullies alike, and the laws limited available redress, *Southern California Interdiscipl. Law J.*, vol. 26, no. 2, p. 379, 2016.
- [3] A. M. Vivolo-Kantor, B. N. Martell, K. M. Holland, and R. Westby, Asystematic review and content analysis of bullying and cyber-bullying measurement strategies, *Aggression Violent Behav.*, vol. 19, no. 4, pp. 423434, Jul. 2014, doi: 10.1016/j.avb.2014.06.008.
- [4] Sruthi. M. V, "Advanced Lung Cancer Diagnosis Using Optimized Deep Learning Models," 2025 2nd International Conference on New Frontiers in Communication, Automation, Management and Security (ICCAMS), pp. 1–6, Jul. 2025, doi: 10.1109/iccams65118.2025.11234121.
- [5] M. Dadvar, D. Trieschnigg, R. Ordelman, and F. de Jong, Improving cyberbullying detection with user context, in *Proc. Eur. Conf. Inf. Retr.,in Lecture Notes in Computer Science: Including Subseries Lecture Notes in Arti cial Intelligence and Lecture Notes in Bioinformatics*, vol. 7814, 2013, pp. 693696.
- [6] A. S. Srinath, H. Johnson, G. G. Dagher, and M. Long, Bul lyNet: Unmasking cyberbullies on social networks, *IEEE Trans. Computat. Social Syst.*, vol. 8, no. 2, pp. 332344, Apr. 2021, doi:10.1109/TCSS.2021.3049232.
- [7] A. Agarwal, A. S. Chivukula, M. H. Bhuyan, T. Jan, B. Narayan, and M. Prasad, Identi cation and classi cation of cyberbullying posts: A recurrent neural network approach using under-sampling and class weighting, in *Neural Information Processing (Communications in Computer and Information Science)*, vol. 1333, H. Yang, K. Pasupa, A. C.-S. Leung, J. T. Kwok, J. H. Chan, and I. King, Eds. Cham, Switzerland: Springer, 2020, pp. 113120.
- [8] Z. L. Chia, M. Ptaszynski, F. Masui, G. Leliwa, and M. Wroczynski, Machine learning and feature engineering-based study into sarcasm and irony classi cation with application to cyberbullying detection, *Inf. Process. Manage.*, vol. 58, no. 4, Jul. 2021, Art. no. 102600, doi: 10.1016/j.ipm.2021.102600.
- [9] N. Yuvaraj, K. Srihari, G. Dhiman, K. Somasundaram, A. Sharma, S. Rajeskannan, M. Soni, G. S. Gaba, M. A. AlZain, and M. Masud, Nature-inspired-based approach for automated cyberbullying classi cation on multimedia social networking, *Math. Problems Eng.*, vol. 2021, pp. 112, Feb. 2021, doi: 10.1155/2021/6644652.
- [10] B. A. Talpur and D. OSullivan, Multi-class imbalance in text classi cation: A feature engineering approach to detect cyberbullying in Twitter, *Informatics*, vol. 7, no. 4, p. 52, Nov. 2020, doi: 10.3390/informatics7040052.
- [11] A. Muneer and S. M. Fati, A comparative analysis of machine learning techniques for cyberbullying detection on Twitter, *Futur. Internet*, vol. 12, no. 11, pp. 121, 2020, doi: 10.3390/12110187.
- [12] R. R. Dalvi, S. B. Chavan, and A. Halbe, Detecting a Twitter cyberbullying using machine learning, *Ann. Romanian Soc. Cell Biol.*, vol. 25, no. 4, pp. 1630716315, 2021.
- [13] R. Zhao, A. Zhou, and K. Mao, Automatic detection of cyberbullying on social networks based on bullying features, in *Proc. 17th Int. Conf. Distrib. Comput. Netw.*, Jan. 2016, pp. 16, doi: 10.1145/2833312.2849567.
- [14] L. Cheng, J. Li, Y. N. Silva, D. L. Hall, and H. Liu, XBully: Cyberbullying detection within a multi-modal context, in *Proc. 12th ACM Int. Conf. Web Search Data Mining*, Jan. 2019, pp. 339347, doi:10.1145/3289600.3291037.