

SOCIAL THREAT IDENTIFICATION ON X VIA SENTIMENT ANALYSIS: AN SVM AND LSTM-BASED APPROACH

¹Rajendra G. Pawar, ²Prajwal Chinchmalatpure, ³Suyash Chinchmalatpure, ⁴Ramachandra V. Pujeri, ⁵Dhanashri Wategaonkar, ⁶Nagesh Jadhav, ⁷Suresh Kapare, ⁸Deepali Nilesh Naik

^{1,4,6,7} Department of Computer Science & Engineering, MIT School of Computing, MIT ADT University, Pune, India.

⁵Department of Computer Engineering and Technology, Dr Vishwanath Karad, MIT World Peace University, Pune, India

²Northeastern University, Boston, United States

³University of Toronto, Rotman, Toronto, Canada

⁸Symbiosis Centre for Management and Human Resource Development, Symbiosis International (Deemed University), Pune, India

Email: ¹rgpawar13@gmail.com, ³chinchmalatpure.p@northeastern.edu,
³suyashsvc@gmail.com, ⁵dhanashri.wategaonkar@mitwpu.edu.in,
⁸deepali.bhanage@gmail.com

Abstract

This research explores the use of tweet sentiment analysis for threat detection and identification based on X data. X provides a valuable platform for monitoring public sentiment and identifying potential dangers in real time. By categorizing tweets as positive, negative, or neutral, analysts can gain insights into public opinion and identify potential risks related to specific events or topics. The research reviews prior work on using Tweet sentiment analysis for threat detection, highlighting the advantages, challenges, and limitations. It identifies research gaps in the field and proposes a comprehensive approach that combines sentiment analysis techniques with real-time data gathering to identify and mitigate threats with minimal delay.

Key words: Tweet Sentiment Analysis, Threat Detection, Sentiment Analysis Techniques, Fake News Detection, Machine Learning Algorithms, Natural Language Processing

1. Introduction

The capacity to swiftly and properly detect and identify possible dangers is crucial in today's world of fast

change. X has become a useful tool for identifying and detecting threats on social

media networks. X offers a singular opportunity to monitor public sentiment and spot possible dangers before they materialize because of its billions of users globally and continuous stream of real-time data. With the help of Tweet sentiment analysis, tweets can be automatically categorized as favorable, negative, or neutral.

Analysts can learn a great deal about public opinion and spot potential dangers by examining the sentiment

of tweets about a specific topic or event. This method has been used to track public opinion in the wake

of terrorism attacks, political campaigns, and even natural disasters. The use of Tweet sentiment analysis

for threat detection and identification is explored in this context. The objective is to comprehend how this

technology can be used to track public opinion and spot potential risks immediately. The discussion will dig

into a complete approach to using techniques of sentiment analysis and real time data gathering to identify

threats and contain them with the least possible buffer time.

2. Literature Review

Sentiment analysis is gaining popularity as a viable method in this area since it includes analyzing the emotional

tone of text data. The use of Tweet sentiment analysis for threat detection and identification has been the

subject of some research. However, there are still several research gaps that need to be filled. With a focus

on the advantages, difficulties, and restrictions of employing Tweet sentiment analysis for threat detection

and identification, this literature review attempts to give a general overview of the body of knowledge already

available on the subject. Also, it will point out any gaps in the literature in this area and suggest possible

directions for further study.

2.1. Cybersecurity threat detection

The research in [1], cybersecurity data from Tweet is processed using deep neural networks. Bidirectional

566

long short-term memory network is used to extract named entities from these tweets

to form a security alert

or to fill an indicator of compromise, and prior to that, a convolutional neural network (CNN) [16] is used to

identify tweets containing security-related information relevant to assets in an IT infrastructure. The proposed

deep-learning NER (named entity recognition) methodology is compared in the paper to the Stanford CRF

NER methodology, Support Vector Machine, Multi-Layer Perceptron, and a linear and nonlinear approach.

TensorFlow was used to create two deep architecture models that were discussed in the study, and training

was carried out using the Adam optimizer's TensorFlow implementation. Over three case study infrastructures,

the suggested pipeline obtains an average 94% true positive rate, a 91% true negative rate, and a 92% average

F1-score for the task of recognizing named entities. These findings demonstrate the tool's efficiency in locating

and compiling pertinent cybersecurity data from X for certain IT infrastructure assets.

The research in [2] introduces a novelty detection unsupervised machine learning method for identifying

cyber threat events on X. A preprocessing phase to eliminate unnecessary terms in input documents is

present, after which conversion of input documents into lowercase takes place while stripping out punctuation,

numbers, hyperlinks, mentions, and hashtags. Stop words are removed using the Natural Language Toolkit.

The framework was evaluated using a purpose-built data set of tweets from 50 influential Cyber security

related accounts over twelve months (in 2018). The classifier achieved an F1-score of 0.643 for classifying

Cyber threat tweets and outperformed several baselines including binary classification models. The analysis of

the classification results suggest that Cyber threat-relevant tweets on X do not often include the CVE

(Common Vulnerabilities and Exposures) identifier of the related threats.

The study in [3] suggests a framework for automatically gathering Tweet data on cyber threats by

employing a novelty detection model. A new, unseen tweet is categorized as either normal or anomalous by the

model after it learns the characteristics of cyber threat intelligence from the threat descriptions published in

public repositories like CVE. Using a specially created data set of tweets from 50 influential accounts related to

cyber security over a 12-month period, the proposed framework was assessed (in 2018). The findings demonstrate

that the classifier outperforms several baselines, including binary classification models, and receives a good F1-

score for categorizing tweets about cyber threats. According to an analysis of the classification results, tweets

about cyber threats on X don't frequently include the threats' CVE numbers.

2.2. General Sentiment Analysis for Threat Detection

The review paper [4] published in 2016 that conducts a survey on all the prominent techniques of sentiment

analysis for X. In the early days of research around Tweet sentiment analysis, researchers used binary

classification techniques intended to classify the stream of text into positive, negative, or neutral. Different

models we used such as Naive Bayes, Max Entropy, Support Vector Machines. A significant highlight is the

use of a two phase analyzer used by [5] In the first phase, tweets are characterized as subjective and objective,

after which they were characterized similarly as others. It was concluded in [6]. that the Stochastic Gradient

Descent based model outperforms other methods when used with a particular learning rate. The combination of

part-of-speech and prior polarity helps achieve better results with the tree kernel model, as proved in [7]. Many

568 other approaches have been examined, such as K-Next Neighbors for classification and Mutual Information,

Chi square, Lemmas, Polarity Lexicons, Valence Shifters, etc. for feature extraction. Ensemble methods such

as weighted combinations have also been employed. Furthermore, [8] emphasized the difficulties and effective

methods for extracting opinions from tweets on X. X makes it difficult to retrieve opinions because

of spam and the vastly different languages used there.

The study in [9] employs a hybrid methodology to examine political opinions on X using a sentiment

analyzer with machine learning. The effectiveness of Naive Bayes and Support Vector Machines (SVM), two supervised machine learning methods, in sentiment analysis, is also contrasted in this research. The

authors gathered information from X's public accounts and utilized Opinion Lexicon to determine the

proportion of positive, neutral, and unfavorable tweets. Based on gathered hashtags connected to

opinions about political parties, they also used Tweepy API to gather public opinion. The paper presents the

results of experiments conducted to compare the accuracy of different sentiment analyzers such as TextBlob,

SentiWordNet, and W-WSD. The authors also compared the performance of two supervised machine learning

algorithms, Naive Bayes and support vector machines (SVM), in sentiment analysis. The results showed that

the hybrid approach that involved a sentiment analyzer with machine learning performed better than the other

analyzers. The SVM algorithm also outperformed Naive Bayes in sentiment analysis.

The study in [10] used sentiment analysis to determine the polarity of tweets during the Trump vs Clinton

elections. Specifically, the paper uses linear regression to model the relationship between the dependent variables

(polarity) and the independent variables (words, sentences, or entire documents). The paper also uses 10-fold

cross-validation to improve the accuracy of the system. The sentiment analysis system achieved an accuracy of

81.51% when tested on a data set of 14,000 tweets collected from X. The performance of the system is

compared to other approaches, such as support vector machine and naive Bayes, and found to be better.

The study in [11] proposes a method for detecting changes (especially sudden) in Tweet sentiment. The

approach involves two core parts:

- Collecting and characterizing tweets according to their sentiment using the Lexicon approach.

- Detecting significant changes in the time series of sentiment scores using change detection algorithms that

do not require historical information.

Since the method uses CUSUM, it is extremely lightweight and memory efficient. The paper applied the proposed methodology on a Tweet data set streamed from 15-03-2018 to 24-03-2018 using the hashtag "theresamay". The data set included 15491 posts after discarding non-English language posts. The results

showed that the proposed approach can detect meaningful sentiment changes across a hashtag's lifetime.

The review in [12] goes through all prominent techniques used most recently in the field of Tweet

sentiment analysis. It also includes mathematics involving them.

The study in [13] uses Bag-Of-Words, n-gram and TF-IDF to detect extremist content on X. By using combinations of multiple layers of LSTM and CNN, they are able to classify tweets as extremist or not

extremist with an accuracy of 92%.

2.3. Language based sentiment analysis

To detect threatening language in Urdu tweets, the study in [14] introduces a new data set. The

data set includes 3,564 tweets that have been manually classified as either threatening or non-threatening by

human experts. The paper uses a two-step process to determine whether a particular tweet is threatening

or not, and then it determines whether it is being used to threaten an individual or a group. In the study,

machine learning and deep learning classifiers are used in several experiments to

compare three different types

of text representation. The findings demonstrate that for the task of target identification, SVM using fastText

pretrained word embedding outperformed other classifiers, while the MLP classifier with the combination of

word n-gram features outperformed other classifiers in detecting threatening tweets. The paper also uses two

neural network-based models, 1-Dimensional Convolutional Neural Network (1D-CNN) and Long Short-Term

Memory Networks (LSTM), for threatening language detection.

The study on the challenges involving sentiment analysis for tweets in Indian Languages, in comparison

to tweets in English is done in [15].

3. Research Gap

The literature survey helps us identify the most potent machine learning algorithms for sentiment analysis.

Among traditional machine learning approaches, the Naïve Bayes algorithm and the Support Vector Machine

Classifier have proven to be the most effective. Among Neural Networks, Long-Short Term Memory Networks

and Bidirectional Long-Short Term Memory Networks are most promising. However, an overview of the recent

events involving breakouts directly or indirectly caused by X suggests an implementation that consists of

sentiment analysis paired with effect-based monitoring with metrics to benchmark the monitoring thresholds.

Although there are multitudes of machine learning algorithms which can be honed to produce extremely accurate results, their practical implementation is incomplete if factors such as interpretation and context are accounted for. There are still certain research gaps that need to be filled, notwithstanding the current increase in interest in the use of Tweet sentiment analysis for threat detection and identification. The absence of a consistent approach for conducting sentiment analysis is one of the major gaps in this subject. Although there are many tools for sentiment analysis, their accuracy and usefulness can change based on the context and subject of the analysis. This makes it challenging to establish a consistent methodology for sentiment analysis and to compare findings across investigations.

Furthermore, context is lost in translation and transliteration across languages, which makes it impossible to use a single data set. The difficulty of handling the vast amount of data created on X represents another area of research that needs improvement. Finding relevant and useful information among the millions of tweets that are posted

everyday is difficult to undertake. This calls for the use of cyclic repetition in the approach. Furthermore, to develop an industry-ready product, we must ensure manual supervision. Hence, a dashboard with manual control over flagging and tagging is necessary, along with the power of data, which will enable the model to learn by itself.

3.1. Research questions

Following are the research questions realized from the study:

- How effective are different machine learning algorithms for sentiment analysis in detecting and classifying threatening tweets on X?
- What is the impact of incorporating context-based sentiment analysis in threat detection using Twitter data?
- How does the inclusion of effect-based metrics, such as likes, retweets, and replies, enhance the accuracy and reliability of threat detection on X?
- What are the limitations and challenges in accurately identifying and categorizing different types of threats using Tweet sentiment analysis techniques?
- How can threat detection models adapt and evolve to effectively identify emerging threats and evolving patterns of harmful content on X?
- What are the potential applications and practical implications of Tweet sentiment analysis for threat detection, such as in public safety, cybersecurity, or disaster management?

3.2. Research objectives

- To develop a practical approach to detect cybersecurity, terrorist, public violence etc. threats from X data.
- To use metrics and effect based dynamic visualization to track threats mentioned above.
- To ensure that the developed approach accounts for the context in terms of language.

4. Proposed methods and materials

Considering the complexity of the problem, the following scenarios are possible:

- Even if a tweet is offensive, it might not gain enough traction to be simply ignored by the public. To solve

Thus, we will create benchmarks for metrics such as likes, retweets, replies, and a

priority queue. This

will ensure that even if the sentiment analysis algorithm detects the tweet as offensive, if it does not gain enough traction, it will not be prioritized.

- Secondly, there is the case where the tweet is detected as offensive but is not offensive (a false

positive). Finally, there is also the case where a tweet is offensive but has not been detected (a false negative). A combined approach of metrics and recursive checking can ensure that the false positives are not prioritized, and that the false negatives end up on the true negatives' list eventually. The false positives will get eliminated based on metrics, and even if the metrics fail, running the algorithm on the replies will reduce its priority. Similarly, a false negative will be retested since the metrics threshold will be crossed, and once the algorithm is run on the replies, it is highly unlikely that any model will fail multiple times in a row. Hence, it will be prioritized. Based on the above concept, the following architecture is proposed refer fig.1:

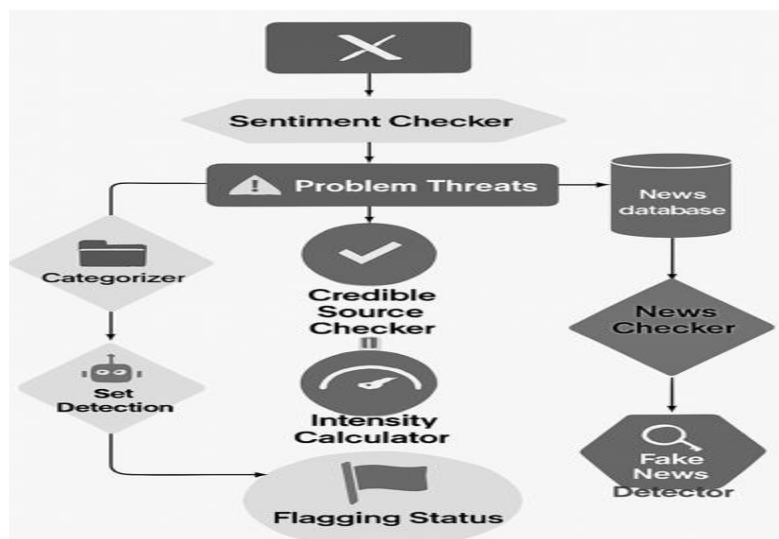


Figure 1. Architecture of machine learning model

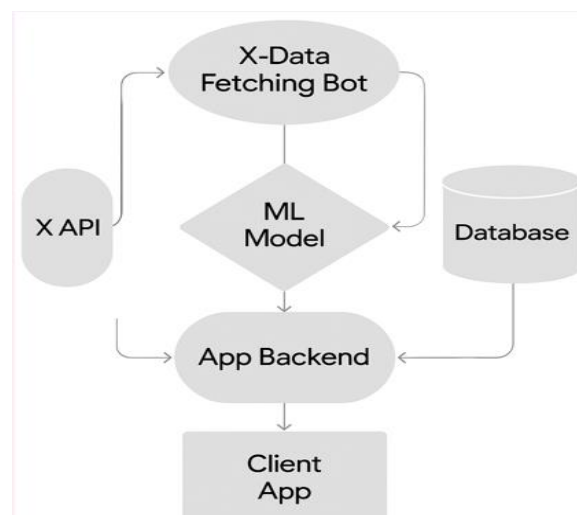
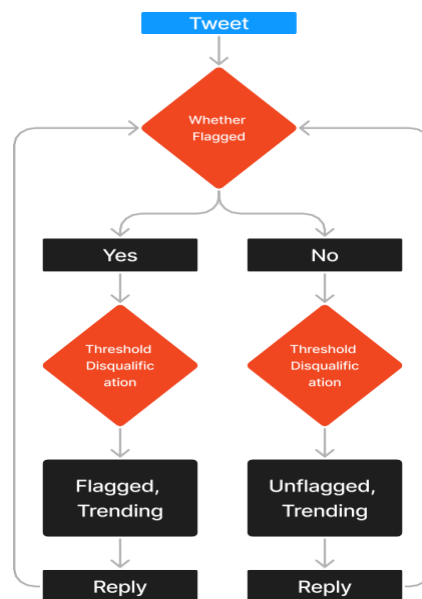


Figure 2. Back-end Overview.

At L0, data will be fetched using the X API at the backend and will go through the profanityChecker (refer fig.2). Any flagged tweets will go through to L1, where the categorizer will flag them into categories of hatespeech, cyberbullying, crime, or threat, and further flag them into the kind of abuse using the canonical form of text. At the same time, we will confirm, using a database of credible news sources, if there is a form of news asserted in the text and whether it is a credible source. If the source is not credible, we will look for similar publications from credible news sources. After going through all this, if the tweet gets flagged as fake news, it will be tagged. Bot detection will be triggered at both instances of L1. Furthermore, we will detect the effect that the tweet has had on the community. For this, replies to the flagged and classified tweets will go through L0 and L1. If they too get flagged, they will be tagged and assigned priority according to the amount of traffic and number of flagged tweets in their thread.

**Figure 3.** Flagged tweet status updating

Considering the various implementation gaps of ML alone, an effect-based approach has been implemented. It ensures using metrics and recursive implementation, that toxic threads will be detected before they cause any large-scale violence issues. If a tweet is flagged and classified, apart from simultaneously implementing fake news, its top replies are immediately tested for toxicity as well. If a tweet is getting popular, that is, it has many likes, replies, and retweets, it will be assigned higher priority. Furthermore, if its replies are also flagged and have a high intensity score, its priority increases further. This will help the administrator to clearly gauge threats. Further, once action is taken, even if the administrator does not unflag the tweet, the recursive checking of metrics will ensure that if it is not relevant anymore, its priority decreases, and other threads of more priority take its

place.

5. Results

5.1. Fake news detection

Multiple studies were conducted by us using various algorithms and data sets. At first, we used SVM and Naive

Bayes algorithms to check for fake news (refer fig.4 and fig.5). We made a custom data set consisting of about 400,000 real and fake

news headlines. The results are as follows:

5.1.1. SVM

Figure 4. (a) SVM precision recall curve. (b) SVM confusion matrix.

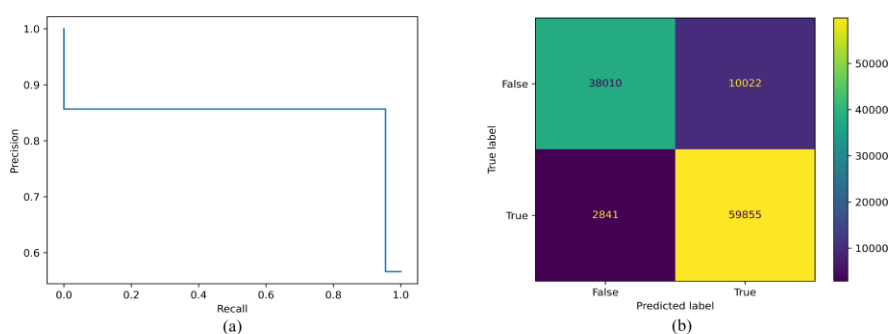


Table 1. SVM accuracy.

	Precision	Recall	F1-Score	Support
0	0.93	0.79	0.86	48032
1	0.86	0.95	0.90	62696
Accuracy			0.88	110728
Macro Average	0.89	0.87	0.88	110728
Weighted Average	0.89	0.88	0.88	110728

5.1.2. Naive Bayes

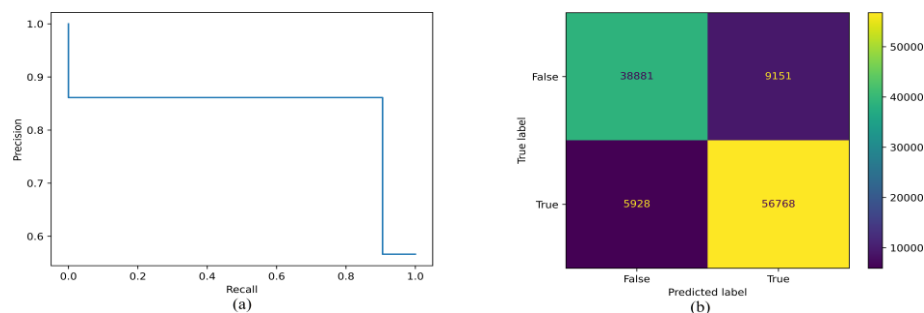


Figure 5. (a) Naive Bayes recall curve. (b) Naive Bayes confusion matrix.

Table 2. Naive Bayes accuracy.

	Precision	Recall	F1-Score	Support
0	0.87	0.81	0.84	48032
1	0.86	0.91	0.88	62696
Accuracy			0.86	110728
Macro Average	0.86	0.86	0.86	110728
Weighted Average	0.86	0.86	0.86	110728

Concluding that the accuracy is good enough, we can consider the machine learning approach for detecting fake news after making a comprehensive data set consisting of multiple features. However, the number of false positives are extremely high. Hence, we propose to use an ensemble approach consisting of a news handles database as well.

5.2. Sentiment checker

Next, we went through multiple approaches to look for the initial sentiment of the tweet. Out of them the

results for SVM model are as follows:

Table 3. SVM precision, recall, f1-score, and support for initial sentiment.

	Precision	Recall	F1-Score	Support
Positive	0.9185	0.9124	0.9154	3435
Neutral	0.9010	0.9833	0.9404	2694
Negative	0.8926	0.7862	0.8360	1871

576 5.3. Bot detection

Then, we used feature extraction and logistic regression to detect bots. Cresci

served as the data set. The many

types of accounts that are present are divided into sub classes, such as traditional spambots, social spambots,

real accounts, accounts with fake followers, etc. These are the outcomes:

Table 4. Logistic regression confusion matrix.

	Predicted Positive	Predicted Negative
Positive	635	60
Negative	60	1893

Table 5. Logistic regression accuracy.

Accuracy	0.9547
----------	--------

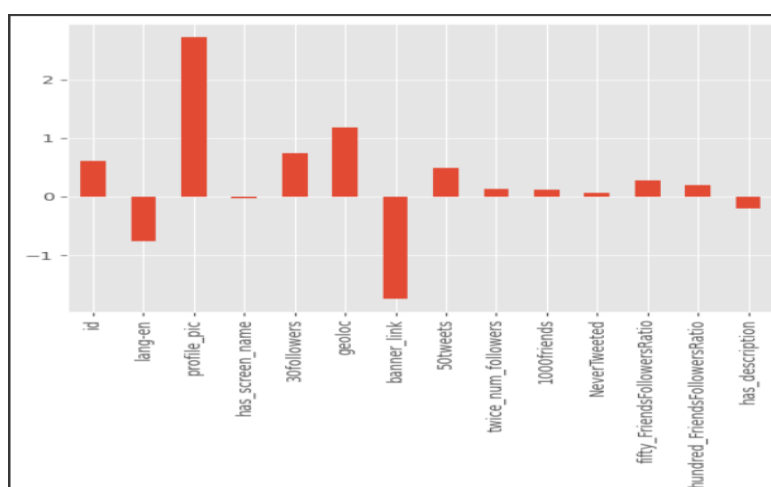


Figure 6. GG plot of newly created data frame based on feature selection.

Table 6. Pipeline approach confusion matrix.

	Predicted Positive	Predicted Negative
Positive	627	49
Negative	57	1915

Table 7. Pipeline approach accuracy.

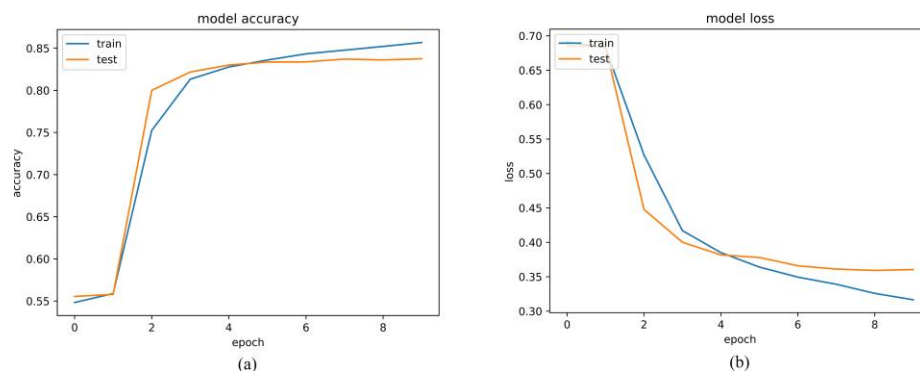
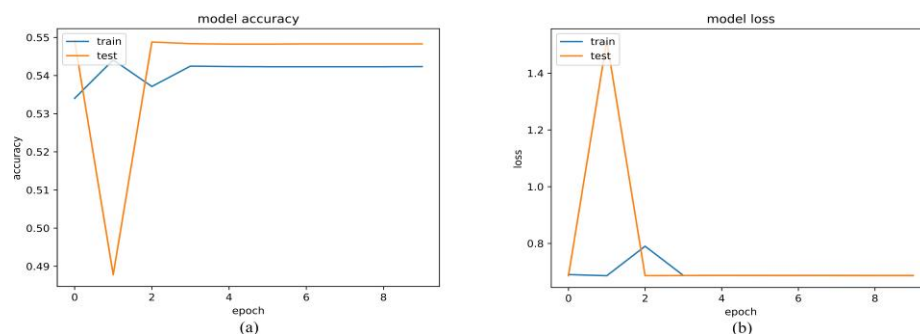
Accuracy	0.9599
----------	--------

5.4. Categorizer

Further, we made another custom data set by curating text from the multiple data sets. This data set was used

to train a classifier that identified the classes of hatred as hate speech and cyber bullying. The results are as

follows:


Figure 7. (a) Accuracy on using single output LSTM. (b) Loss on using single output LSTM.

Figure 8. (a) Accuracy on using multiple output LSTM. (b) Loss on using multiple output LSTM.

5.5. Translation and Transliteration

With reference to the literature review subsection 2.3, a study was conducted to check for context loss.

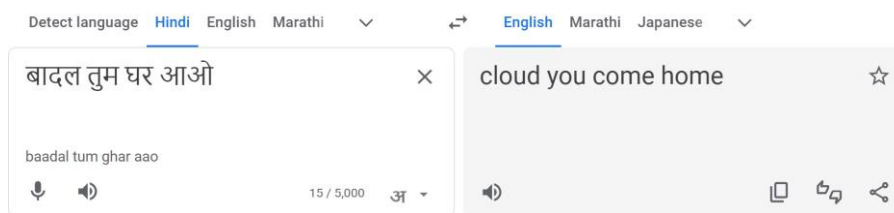


Figure 9. Change of meaning in translation.

In Figure 9, the original sentence clearly asks for a person named '*Badal*' to come home, whereas the translation misinterprets the name of the person and considers it to be addressing a cloud instead.



Figure 10. Loss of context in transliteration.

In Figure 10, the transliteration is not accurate as the stress patterns and syllable significance across various languages are different. Hence it cannot detect sentiment accurately.

Hence, we propose to opt for a script identification layer, after which separate models will be trained for

separate data sets. A single data set will be used for English and Hindi texts that are written in English. The

dictionary for the flagger will consist of these texts as well. This will help maintain accuracy, and with the help

of ensemble approaches, we will be able to detect threats from X across all prominent languages.

6. Conclusion and discussion

The above approach has been tested and integrated within a web app consisting of a back end and a front end.

With reference to the research objectives and questions, the following results have

been obtained:

- Machine Learning and Deep Learning techniques are theoretically very powerful in identifying the underlying sentiment of tweets, but need to be coupled with metrics to ensure that threats are effectively tracked and visualized.
- The use of a different data set for a different script will ensure that the context is not lost while analyzing the sentiment.
- When this approach is implemented recursively to threads and the results are properly categorized for visualization purposes, an extremely effective method for monitoring X threads for threats is constructed.
- This approach has a variety of use cases such as detection of riots, prevention of terrorist attacks, reducing suicide attempts among others.

References

- [1] N. Dionisio, F. D. Alves, P. L. Ferreira, and A. Bessani, Cyberthreat Detection from Twitter using Deep Neural Networks. 2019. Doi: 10.1109/ijcnn.2019.8852475.
- [2] B. D. Le, G. Wang, M. Nasim, and M. A. Babar, Gathering Cyber Threat Intelligence from Twitter Using Novelty Classification. 2019. Doi: 10.1109/cw.2019.00058.
- [3] A. J. Bose, V. Behzadan, C. Aguirre, and W. Hsu, A novel approach for detection and ranking of trendy and emerging cyber threat events in Twitter streams. 2019. Doi: 10.1145/3341161.3344379.
- [4] V. Kharde and S. S. Sonawane, "Sentiment Analysis of Twitter Data: A Survey of Techniques," International Journal of Computer Applications, vol. 139, no. 11, pp. 5–15, Apr. 2016, Doi: 10.5120/ijca2016908625.
- [5] Luciano Barbosa AT&T Labs - Research, "Robust sentiment detection on Twitter from biased and noisy data — Proceedings of the 23rd International Conference on Computational Linguistics: Posters," DL Hosted Proceedings. <https://dl.acm.org/doi/10.5555/1944566.1944571>
- [6] Albert Bifet University of Waikato, Hamilton, New Zealand, "Sentiment knowledge discovery in twitter streaming data, Proceedings of the 13th international conference on Discovery science," Guide Proceedings. <https://dl.acm.org/doi/10.5555/1927300.1927301>

- [7] A. Agarwal, "Sentiment Analysis of Twitter Data," ACL Anthology, Jun. 01, 2011. <https://aclanthology.org/W11-0705/>
- [8] Z. Luo, M. Osborne, and T. Wang, "An effective approach to tweets opinion retrieval," World Wide Web, vol. 18, no. 3, pp. 545–566, May 2015, Doi: 10.1007/s11280-013-0268-7.
- [9] A. N. Hasan, S. Moin, A. Karim, and S. Shamshirband, "Machine Learning-Based Sentiment Analysis for Twitter Accounts," Mathematical and Computational Applications, vol. 23, no. 1, p. 11, Feb. 2018, Doi: 10.3390/mca23010011.
- [10] A. R. Sulthana, A. K. Jaithunbi, and L. S. Ramesh, "Sentiment analysis in twitter data using data analytic techniques for predictive modelling," Journal of Physics, vol. 1000, p. 012130, Apr. 2018, Doi: 10.1088/1742-6596/1000/1/012130.
- [11] S. K. Tasoulis, A. G. Vrahatis, S. V. Georgakopoulos, and V. P. Plagianakos, Real Time Sentiment Change Detection of Twitter Data Streams. 2018. doi: 10.1109/inista.2018.8466326.
- [12] A. Alsaeedi and M. M. Khan, "A Study on Sentiment Analysis Techniques of Twitter Data," International Journal of Advanced Computer Science and Applications, vol. 10, no. 2, Jan. 2019, Doi: 10.14569/ijacsa.2019.0100248.
- [13] S. Ahmad, M. Asgher, F. Alotaibi, and I.-U. Awan, "Detection and classification of social media-based extremist affiliations using sentiment analysis techniques," Human-centric Computing and Information Sciences, vol. 9, no. 1, Jul. 2019, Doi: 10.1186/s13673-019-0185-6.
- [14] M. Amjad, N. Ashraf, A. Zhila, G. Sidorov, A. Zubiaga, and A. Gelbukh, "Threatening Language Detection and Target Identification in Urdu Tweets," IEEE Access, vol. 9, pp. 128302–128313, Sep. 2021, Doi: 10.1109/access.2021.3112500.
- [15] S. J. Soman, P. Swaminathan, R. Anandan, and K. Kalaivani, "A comparative review of the challenges encountered in sentiment analysis of Indian regional language tweets vs English language tweets," International Journal of Engineering and Technology (UAE), vol. 7, no. 2.21, p. 319, Apr. 2018, Doi: 10.14419/ijet.v7i2.21.12394.
- [16] R. Pawar, S. Ghumbre, and R. Deshmukh, "Visual Similarity Using Convolution Neural Network over Textual Similarity in Content- Based Recommender System," IJAST, vol. 27, pp. 137–147, Sep. 2019.