



International Journal of Engineering Research and Science & Technology

www.ijerst.org

ISSN : 2319-5991

Vol. 22 No. 3 (2026)



ijerst.editor@gmail.com
editor@ijerst.com

*Research Paper***A SYSTEMATIC ANALYSIS ON THE NATURAL CALAMITY
DETECTION BASED ON THE TWITTER DATA****Dr. Angela Deepa^{1*}, Shwetha K², Aswathy M²**¹ Faculty,

Westford Education, India

² Information Technology,

Women's Christian College, Chennai, India

Corresponding Author:

Dr. Angela Deepa

E-mail: rians.bliss@gmail.com

Abstract

The increasing frequency of natural disasters, such as earthquakes, floods, and hurricanes, highlights the importance of timely monitoring and early warning systems for effective disaster management. Social media platforms, particularly Twitter, have become valuable sources of real-time information during such events. This study explores the use of Natural Language Processing (NLP) and machine learning techniques to automatically classify tweets as either disaster-related or non-disaster content, aiming to improve the speed and accuracy of identifying potential crisis situations as they occur. In this work, multiple machine learning approaches, including Random Forest and DistilBERT, are applied to process and categorize tweet data. Random Forest, an ensemble learning method, utilizes multiple decision trees to enhance classification performance and handle high-dimensional feature spaces. DistilBERT, a distilled variant of BERT, is employed to capture contextual information and sequential dependencies within the text, efficiently processing large volumes of tweets while maintaining deep understanding of word relationships. The models are trained on a labeled dataset of tweets marked as "disaster" or "non-disaster." Preprocessing steps, such as text cleaning, tokenization, and padding, are performed to ensure consistent and uniform input data. Experimental results demonstrate that both approaches are effective for real-time disaster tweet detection, with Random Forest performing well in feature-driven classification tasks and DistilBERT excelling at capturing complex contextual nuances in longer tweets. The proposed framework has potential applications in crisis response systems, enabling rapid dissemination of alerts to authorities and supporting proactive disaster preparedness and management.

Keywords: Social Media, Deep Learning Algorithms, Disaster Management System, Machine Learning, Natural Language Processing

Introduction

Natural disasters, including earthquakes, floods, and hurricanes, continue to pose significant threats to societies worldwide, often resulting in substantial human and economic losses. With the increasing frequency and intensity of such events, the need for rapid detection and real-time response has become critical (Hernandez-Suárez et al., 2019). While traditional monitoring systems, such as seismic

networks and weather stations, have advanced over the years, they may still provide delayed alerts or lack coverage in certain regions (Zhang and Vucetic, 2016). In recent times, social media platforms, particularly Twitter, have emerged as valuable sources of immediate, on-the-ground information during crises. Users frequently share firsthand observations and experiences, offering timely insights into ongoing events (Flint and Kellerman, 2020).

Natural Language Processing (NLP) technologies have proven effective in analyzing the vast volumes of unstructured text generated on social media. By leveraging NLP, it is possible to automatically filter and identify tweets containing disaster-related information. Applying these methods to Twitter data enables early detection of natural hazards and can facilitate faster issuance of warnings, thereby helping to minimize response times and enhance community protection (Mahajan et al., 2024). Initial efforts to classify disaster-related tweets primarily relied on simple keyword-based methods. Tweets containing terms such as “earthquake,” “flood,” or “tsunami” were marked as relevant, offering a straightforward approach but limited adaptability. These systems often failed to capture indirect references to disasters and struggled to differentiate genuine alerts from unrelated posts using similar terminology (Zisad and Hasan, 2025). Moreover, the lack of contextual understanding frequently led to misclassifications. Subsequent research introduced machine learning models like Support Vector Machines and Naive Bayes, which learned patterns from labeled datasets and improved identification of relevant tweets. However, these models faced challenges with ambiguous language, sarcasm, and context-dependent expressions common on social media. Additionally, many approaches did not fully account for the sequential or narrative nature of tweets, limiting their ability to accurately interpret disaster situations. This study proposes a hybrid approach that integrates both keyword-based and context-aware techniques to enhance the detection of disaster-related tweets. By combining the complementary strengths of these methods, the aim is to develop a more accurate and robust classification system capable of overcoming the limitations of previous models.

The research is centered around the following aspects:

Focus on Twitter Data: Analysis is confined to Twitter, taking advantage of its real-time nature

for disaster tweet detection. Other social media platforms are not considered.

Natural Calamities Only: The study targets natural disasters, including earthquakes, floods, hail, and wildfires, distinguishing related tweets from unrelated content.

Textual Data Processing: Only textual content from tweets is analyzed using NLP techniques; image or video data are not included.

Machine Learning Models: The research focuses on applying Random Forest and DistilBERT for classification. Other algorithms or hybrid models are outside the scope.

Real-Time Processing: The system is designed for live classification of tweets during ongoing disasters rather than retrospective event analysis.

Limited to English Language: Only tweets in English are analyzed; multilingual or cross-lingual approaches are excluded.

Disaster-related and Non-Disaster-related Classification: The study classifies tweets into two categories-disaster-related and non-disaster-related-without differentiating between specific disaster types or providing detailed event analysis.

Literature Review

The relevant studies in this area include the following:

The study conducted (Gulesan et al., 2021) investigated the use of Twitter for real-time earthquake detection. Tweets were analyzed using Naive Bayes, Support Vector Machines (SVM), and lexicon-based sentiment analysis, achieving an accuracy of 85%. Naive Bayes estimated the likelihood of a tweet being earthquake-related, while SVM distinguished relevant from irrelevant tweets using an optimal separating hyperplane. Lexicon-based sentiment analysis captured emotional cues, such as urgency or panic, associated with seismic events. Despite its promise, the approach was limited by ambiguous language and a relatively small, less diverse dataset, which affected the detection of minor tremors. Deep learning techniques, including CNN, LSTM, and Word2Vec embeddings, have been applied to detect natural disasters in real-time

from social media and satellite data. The system designed by (Kaur et al., 2021), showed a high detection rate for floods and hurricanes, but it encountered difficulties in accurately interpreting regional dialects and slang. Moreover, the system had limited capability in detecting localized disasters that did not garner widespread social media attention. The authors applied CNN (Convolutional Neural Networks) to learn spatial patterns in tweet data, effectively capturing disaster-related content, including any images that might be present. LSTM (Long Short-Term Memory) networks were used to handle the sequential nature of tweets, enabling the model to understand the flow of events in real-time disaster discussions. Word2Vec embeddings helped the model understand the semantic meaning of words, improving its ability to recognize diverse disaster-related terms across different contexts. An automated system has been developed to monitor and detect natural disasters in real-time using social media data. The system developed (Sufi and Khalil, 2024) integrates AI-driven location intelligence with sentiment analysis to identify and track events such as hurricanes, earthquakes, and floods, leveraging both geospatial and temporal information. While the approach demonstrated effectiveness in real-time detection, its accuracy was affected by high volumes of irrelevant or ambiguous posts on social media. Additionally, the precision of location data was occasionally limited, particularly when users provided vague or indirect location references.

A study (Suleman et al., 2023) focused on developing a model to predict earthquake and flood events by analyzing social media data. Leveraging the real-time nature of platforms such as Twitter, the approach aimed to identify early indicators of these disasters, enabling timely alerts and supporting emergency response efforts. The model employed BERT, a pre-trained transformer capable of capturing contextual information from both preceding and following words in a sentence. A notable limitation was the model's difficulty in correctly classifying posts that mentioned both

earthquake and flood events simultaneously, which occasionally led to misclassification.

This study (Imran et al., 2014) focused on developing a robust system for real-time detection of natural disasters using large-scale social media data. By analyzing content from platforms such as Twitter, the system aimed to identify and classify events like earthquakes, floods, and wildfires, extracting relevant information from the vast user-generated content during crises. The predictive model was designed to filter through the noise of social media, providing timely insights for disaster response. Despite its effectiveness, the system faced challenges with scalability during high-volume events, requiring substantial computational resources. Additionally, accurately distinguishing genuine disaster reports from expressions of concern or speculation remained a key limitation.

This study (Zhang et al., 2019) focussed on the use of Twitter data for real-time detection of wildfires. Leveraging the platform's immediate updates and widespread user engagement, the approach aimed to identify wildfire-related discussions and events as they occurred. A key challenge encountered was the high incidence of false positives, where posts containing the term "fire" were unrelated to actual wildfire events, complicating accurate detection.

This study (Sakaki et al., 2010) focus the use of continuous Twitter data streams for real-time earthquake detection, addressing a critical aspect of disaster monitoring and response. Twitter generates large volumes of user-generated content during seismic events, including personal experiences, observations, and speculative posts. By analyzing this data in real-time, earthquakes can be detected faster than through conventional seismic monitoring, potentially enhancing emergency response and saving lives. Recurrent Neural Networks (RNNs) were employed to model the time-series nature of the data. However, the system faced challenges due to noisy content, such as unrelated posts, jokes, or figurative expressions, which could lead to false alarms and reduce detection accuracy.

A system (Caragea et al., 2016) detect and classify tweets related to natural disasters. By leveraging Twitter's widespread use during emergencies, the system aimed to track user-generated updates, observations, and news, providing authorities with timely insights to support faster emergency responses. The approach combined Naive Bayes, Convolutional Neural Networks (CNNs), and topic modeling to efficiently process and categorize disaster-related content. Naive Bayes served as a probabilistic classifier to distinguish disaster-related tweets from unrelated posts, while CNNs, typically used in computer vision, were applied for text classification tasks, demonstrating effectiveness in natural language processing. A notable limitation of the system was its handling of linguistic diversity and regional variations in disaster reporting. Users from different regions often employ distinct languages, dialects, slang, and locally specific expressions, which may reduce the accuracy of models trained primarily on standard English-language data.

This study (Kireyev et al., 2009) focused on using social media, particularly Twitter, for real-time disaster monitoring and alert generation. The objective was to develop an automated system capable of analyzing tweets to detect events such as floods, hurricanes, earthquakes, and wildfires, and to issue timely alerts for emergency response teams. Given the high volume of real-time updates shared on social media during disasters, the system aimed to leverage this data to enhance disaster management and reduce response times. A key component of the approach was Latent Dirichlet Allocation (LDA), a generative probabilistic topic modeling technique. LDA identifies underlying topics within a collection of documents-in this case, tweets-by assuming that each tweet is a mixture of multiple topics, with each topic represented as a distribution over words.

This study (Hettiarachchi et al., 2020) focus on the use of social media, particularly Twitter, for real-time tsunami detection. Early

identification of tsunamis is critical for reducing damage and saving lives, and social media offers rapid information dissemination that can support timely alerts. The approach leveraged large volumes of real-time Twitter data to detect tsunami events as they unfolded, enabling faster responses from emergency management systems. Support Vector Machines (SVM) were employed to classify tweets as tsunami-related or non-tsunami-related, benefiting from their ability to handle high-dimensional text data. A major limitation of the system was the low density of Twitter users in some tsunami-prone regions, particularly in small island nations or rural coastal areas, which reduced the coverage and effectiveness of real-time detection.

Materials And Methods

Methodology

The methodology refer **Figure 1** involves collecting and preprocessing a comprehensive real-time Twitter dataset for natural disaster prediction. Tweets are classified into real and fake categories, with relevant features being selected or engineered for model training. The dataset is stored and tracked in an SQL database for efficient querying and retrieval. Two separate models are used for prediction: Random Forest, an ensemble learning algorithm, and DistilBERT, a transformer-based model for natural language processing. Both models are trained individually, hyperparameters are fine-tuned, and performance is evaluated using key metrics such as accuracy, precision, recall, and F1-score. Each model's performance is compared to determine the most effective approach. Interpretability techniques are applied to understand the models' decisions. After validating the accuracy of both models in classifying real and fake tweets, they are deployed for real-time natural disaster prediction, ensuring both reliable and interpretable outputs.

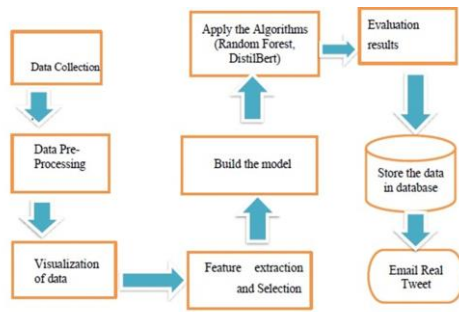


Figure 1: Methodology Diagram

The following are the steps in detail:

Data Collection:

The data collection process for natural disaster tweet prediction involves gathering a comprehensive dataset consisting of tweet details, including unique tweet ID, tweet content, location, text, and target labels. The dataset is obtained from Kaggle. The primary goal is to classify tweets as either related to a real disaster or fake disaster events in real-time. This classification helps in identifying and distinguishing genuine disaster-related tweets from misleading or false ones, enabling timely and accurate disaster predictions and responses based on social media data. The dataset is continuously updated to ensure real-time prediction and monitoring of disaster-related tweets.

Data preprocessing:

The data preprocessing steps for predicting disaster-related tweets involve several key processes to ensure that the dataset is ready for modeling. One of the initial steps is handling missing values, where any null or missing data points in the dataset are identified and addressed. Common strategies for handling missing values include imputation techniques, such as replacing missing values with the mode for categorical data or using more advanced methods like KNN imputation for both categorical and numerical values. Next, textual data needs to be processed effectively for machine learning models. This includes cleaning the tweet text by removing any irrelevant characters such as punctuation, URLs, special symbols, and stopwords. Tokenization is performed to break the text into individual words or tokens, and these tokens are

converted into numerical representations using techniques like TF-IDF (Term Frequency-Inverse Document Frequency) or word embeddings (e.g., Word2Vec, GloVe). These representations help in capturing the semantic meaning of the text for better prediction. For categorical variables, such as the location or target (real or fake disaster tweet), encoding is necessary. Label encoding or one-hot encoding can be used to convert these categorical variables into numerical format, making them compatible with machine learning algorithms. Another important preprocessing step is scaling numerical features, if any, to a consistent range. While textual data is typically processed through embeddings, if there are numerical features (e.g., tweet length or other meta-data), they are scaled using techniques like Min-Max scaling or standardization to ensure that all features contribute equally to the model. Finally, normalization ensures that all features follow a consistent distribution or scale, making them more suitable for model training. By performing these preprocessing steps, the quality of the disaster tweet dataset is improved, making it more suitable for training machine learning models to effectively classify tweets as real or fake disaster-related content. These preprocessing steps are essential to handle the text and categorical data efficiently and ensure the features are in the optimal format for training machine learning models.

Visualization of Data

Visualization plays a crucial role in understanding the disaster tweet dataset and evaluating model performance. Several common techniques are employed to gain insights into the data and assess the effectiveness of the prediction models. Histograms and boxplots are used to explore the distribution of numerical features such as tweet length, the number of words, or any metadata, helping identify outliers and skewness in the data. Word clouds are especially useful for visualizing the most frequent terms in real versus fake disaster tweets, providing insight into the language used in each category. After model evaluation, a confusion matrix is

essential for assessing the model's ability to classify tweets accurately, showing the number of true positives, false positives, true negatives, and false negatives, and helping evaluate key metrics like accuracy, precision, recall, and F1-score. Finally, learning curves help identify overfitting or underfitting by showing how the model's performance improves as more training data is provided or the number of epochs increases. Overall, these visualizations help in identifying trends, assessing model performance, and communicating insights effectively, improving the understanding and accuracy of the disaster tweet classification model.

Feature Extraction and Selection

Feature extraction and selection in the disaster tweet classification task are performed in two stages. First, text-based features from the tweet content, such as term frequency and sentiment, are extracted and transformed using techniques like TF-IDF or word embeddings to make the data suitable for the DistilBERT model. It is a transformer-based model wellsuited for understanding textual data. Second, a new set of features is created by combining these extracted features with predictions from the Random Forest model. This enables the DistilBERT model to leverage the Random Forest's ability to capture non-linear relationships in the data. By combining the strengths of both models-DistilBERT's capacity to understand the nuances of language and Random Forest's ability to handle structured data-this approach aims to enhance the accuracy and robustness of disaster tweet prediction

Building the model

Once the data has been pre-processed and relevant features are extracted, the next step is to build a model capable of predicting disaster-related tweets. In this approach, two separate models are utilized: Random Forest and DistilBERT. Each model plays a critical role in enhancing the overall accuracy and robustness of the prediction system. The Random Forest model, an ensemble learning technique, is employed to leverage multiple decision trees for classification tasks. This model is

particularly effective in handling non-linear relationships between features and is known for its robustness against overfitting. The Random Forest model works by aggregating the results of individual decision trees to produce a final prediction, making it highly effective in classification tasks where multiple variables and interactions are at play. It excels in handling both categorical and numerical features, which is useful for the diverse nature of the tweet dataset. On the other hand, DistilBERT (a smaller, more efficient version of the BERT model) is used for natural language processing tasks, especially for understanding the intricate patterns in tweet text. DistilBERT is designed to process sequences of words and capture contextual meanings in a much more efficient manner compared to traditional models. By utilizing this transformer-based model, we can effectively capture semantic relationships and contextual information in the tweets, enabling the model to distinguish between real and fake disaster-related content. The model architecture includes attention mechanisms that allow it to weigh the importance of different words in a tweet, providing deeper insights into the content of the text. Both models-Random Forest for handling structured features and DistilBERT for processing textual data-are trained independently on the dataset. While Random Forest captures complex relationships between features, DistilBERT focuses on understanding the meaning and context of the text. The combined results of both models are then analyzed to determine which approach offers better accuracy in distinguishing real disaster tweets from fake ones. This dual-model approach, integrating traditional machine learning (Random Forest) and state-of-the-art natural language processing (DistilBERT), enables the system to leverage the strengths of both. The Random Forest model handles non-linear relationships and feature interactions, while DistilBERT captures complex language patterns in tweets. The fusion of these techniques allows for more accurate disaster tweet predictions and enhances the overall robustness of the system. By incorporating both

machine learning and deep learning approaches, this model is better equipped to classify disaster-related tweets effectively in real time, providing reliable predictions that can be used for timely disaster response.

Algorithm Description

In this system, various machine learning algorithms are employed to classify tweets as disaster-related or non-disaster-related. These algorithms work by analyzing the text data of tweets and identifying patterns or features that indicate relevance to natural calamities. The models are trained using labeled data, where each tweet is marked as either disaster-related or not. The algorithms process the tweet's text, extract relevant features, and use them to make predictions.

An Overview of Random Forest algorithm

The Random Forest algorithm is an ensemble learning method that combines multiple decision trees to make predictions. It is particularly effective for both classification and regression tasks. Random Forest works by constructing a multitude of decision trees during training and outputting the class that is the majority vote (for classification) or the average prediction (for regression) of the individual trees. Each tree in the forest is built using a random subset of the data and a random subset of features, which helps to ensure diversity among the trees, preventing overfitting and improving generalization. Random Forest is highly valued for its accuracy, flexibility, and ability to handle large datasets with high-dimensional features. It also provides estimates of feature importance, which can be useful for feature selection and understanding the underlying structure of the data.

An Overview of DistilBERT algorithm

DistilBERT is a smaller, faster, and more efficient version of BERT (Bidirectional Encoder Representations from Transformers), which is a deep learning model designed for natural language processing (NLP) tasks. DistilBERT retains 97% of BERT's language understanding capabilities while being 60% faster and requiring 40% fewer parameters.

This is achieved by applying knowledge distillation, a technique where a smaller model (DistilBERT) learns to mimic the behavior of a larger, more complex model (BERT). DistilBERT is widely used for various NLP tasks, including text classification, sentiment analysis, and named entity recognition. It is particularly beneficial for scenarios where computational efficiency and speed are crucial, while still maintaining high accuracy. The DistilBERT model operates using the transformer architecture, which is based on attention mechanisms. This allows it to understand the contextual relationships between words in a sentence, regardless of their position. DistilBERT is particularly useful for tasks where processing speed and memory efficiency are important but high performance is still required. In essence, DistilBERT provides a balance between efficiency and accuracy, making it ideal for real-time applications in resource-constrained environments, such as mobile devices or situations where quick responses are necessary (like disaster detection from social media).

Evaluation Metrics

Evaluation metrics are essential in measuring the performance of the disaster tweet classification models. The following are the metrics used to provide valuable insights into the strengths of the model for reliable predictions in real-time disaster scenarios.

Confusion Matrix:

The confusion matrix offers a detailed breakdown of the model's predictions, showing the number of true positives (correctly identified disaster tweets), true negatives (correctly identified non-disaster tweets), false positives (non-disaster tweets incorrectly labeled as disaster-related), and false negatives (disaster tweets incorrectly labeled as non-disaster). It is crucial for understanding where the model is making errors, particularly in distinguishing between real and fake disaster tweets, and helps identify imbalances in the classification process.

F-Score (F1 Score):

The F1 score is the harmonic mean of Precision and Recall and provides a balanced measure of the model's ability to correctly classify disaster tweets while minimizing false positives and false negatives. The F1 score is particularly useful in situations where there is an imbalance between the classes and both false positives and false negatives need to be taken into account.

Precision:

Precision measures the accuracy of positive predictions, focusing on the proportion of true positive disaster tweets out of all the tweets predicted as disaster-related. A high precision means that when the model predicts a tweet to be related to a disaster, it is more likely to be correct. Precision is critical when minimizing false positive predictions is important, as in disaster response where irrelevant tweets could lead to unnecessary actions.

Recall (Sensitivity):

Recall measures the model's ability to correctly identify real disaster-related tweets from all the actual disaster tweets in the dataset. It calculates the ratio of true positives to the sum of true positives and false negatives. A high recall means that the model is effectively capturing the true disaster tweets. Recall is especially important when it is crucial not to miss any real disaster-related tweets, such as in urgent disaster response situations. Each of these metrics provides valuable insights into different aspects of the model's performance, and their selection depends on the goals and characteristics of the disaster tweet classification task. For example, a high precision may be important in avoiding false alarms, while a high recall may be prioritized when ensuring that no real disaster-related tweets are overlooked. Log loss helps assess the confidence of predictions, and the confusion matrix provides an overall view of how the model performs across different categories. Together, these evaluation metrics enable a comprehensive assessment of the model's effectiveness in classifying disaster-related tweets accurately in real-time scenarios

Experimental Setup**Dataset:**

The dataset used for this research consists of Twitter data that has been labeled to facilitate the classification of tweets as either disaster-related or non-disaster-related. By leveraging machine learning algorithms on this comprehensive dataset, we aim to develop a robust classification model for enhancing crisis management and response efforts.

Dataset Description

The dataset used in this study for disaster-related tweet classification comprises a diverse and comprehensive collection of information relevant to natural calamities. It includes both tweet content and associated metadata to ensure a holistic representation of factors influencing disaster detection. The dataset captures a wide range of natural disasters, such as earthquakes, floods, hurricanes, and wildfires, alongside non-disaster-related content, enabling the creation of an accurate and reliable classification model.

Attribute Information

The dataset tailored for disaster-related tweet classification encompasses a comprehensive set of features essential for developing predictive models aimed at identifying tweets related to natural calamities. Here are further insights into each attribute:

Id: This is a unique identifier assigned to each tweet in the dataset. It serves as a reference for individual records and allows for easy tracking and analysis of specific tweets. The id column is useful for ensuring the uniqueness of tweets in the dataset, as well as linking tweets with additional metadata if needed in further analysis.

Keyword: This column contains disaster-related keywords that are used to identify tweets related to a particular type of disaster (e.g., "earthquake," "flood," "wildfire"). These keywords serve as primary indicators for the relevance of the tweet to a natural calamity.

Location: This column provides geographical information about where the tweet was posted. It may include information about the city, region, or country, though it can also be empty or vague for many tweets, as users may not always provide precise location data.

Text: This is the core content of each tweet, containing the actual message shared by the user.

DATA COLLECTION:

The data collection process for natural disaster tweet prediction involves gathering a comprehensive dataset consisting of tweet details, including unique tweet ID, tweet content, location, text, and target labels. The dataset is obtained from Kaggle. The primary goal is to classify tweets as either related to a real disaster or fake disaster events in real-time. This classification helps in identifying and distinguishing genuine disaster-related tweets from misleading or false ones, enabling timely and accurate disaster predictions and responses based on social media data. The dataset is continuously updated to ensure real-time prediction and monitoring of disaster-related tweets. It includes complex expressions of sentiment or commentary on the disaster. Text preprocessing techniques such as tokenization, stemming, and removal of stop words will be applied to this column to prepare it for analysis.

Target: This is the target or label column, indicating whether the tweet is related to a disaster or not. A value of 1 indicates that the tweet is disaster-related, while 0 represents a non-disaster-related tweet. The target column is critical for training and evaluating the classification models. The models will learn to map tweet content (from the text column) and additional features (such as keyword) to this binary target.

Sample data: The careful selection and curation of the dataset refer **Figure 2** aims to provide a realistic representation of the complexities associated with disaster-related tweet classification. By incorporating diverse and relevant information, the dataset serves as a foundation for training, validating, and testing the proposed machine learning models, such as Random Forest and DistilBERT, for real-time disaster detection in dynamic social media environments.

id	keyword	location	text	target
2	ablate	New York City	Arsonist sets cars ablaze at dealership	1
3	ablate	Morgantown, WV	Arsonist sets cars ablaze at dealership	1
5	ablate	OC	If this child was Chinese, this tweet would have gone viral. Social media would be ablaze. SNL	0
6	ablate	London, England	Several houses have been set ablaze in Ngembasa village, Oku sub division in the North West F	1
7	ablate	Bharat	Asansol: A BJP office in Salanpur village was set ablaze last night. BJP has alleged that TMC is b	1
8	ablate	Accra, Ghana	National Security Minister, Kan Dapaah's side chic has set the internet ablaze with her latest pc	0
9	ablate	Searching	This creature who's soul is no longer clarent but blue ablaze This thing Carrying memories &	0
13	ablate	HYDERABAD	Under #MamataBanerjee political violence & vandalism continues to unabated in West B	1
14	ablate	Reno, NV	AMEN! Set the whole system ablaze, man. https://t.co/JOAHDCGd0	0
18	ablate	Worldwide	paulzickaphoto: iKceHunde Ablaze!K! Wishing you all a good evening... https://t.co/DONMEIH	0
22	ablate	Italy	#ThankfulTuesday! Allah 43:2 When you pass through the waters, I will be with you; and when	0
23	ablate	'É--ÉÉ: 7mi Kwaht 4 43 CW	When you walk through the fire, 4 43 Éyou will not be scorched, and the 4 43 Éflames wil	0
24	ablate	Oklahoma	Originally they were intended to be fired at borders or the opposing ship's crew from gunners	1
27	ablate	Havana 3am	Another public market in #Haiti mysteriously set ablaze. The independent merchants, who are	1
31	ablate	Washington, DC	Marivan, Kurdistan Province Monday, Jan 13th, 2020 Protesters set the propaganda banner of i	1
32	ablate	India	How can you turn a blind eye to the incident of setting ablaze more than 50 houses of H&K? htt	1
33	ablate	Salta, Argentina	This love is so completely crazy. You've been fucking with my dreams, ripped me like your torn	0
34	ablate	Wherever socks go	Terms in A Demon Burning Dark: The Ruined: People who cannot use magic or interact with i	0
35	ablate	Kuala Lumpur	8P* Heartfelt appreciation to Prime Minister YAB Tun Dr. wife, YABhg. Tun Dr. Siti Haasmah Mo	0
37	ablate	Sydney, New South	86! 86! 86! he gave us everything... He had a horrible foot infection once so wore one thong!e	0
38	ablate	Nigeria	8P* yeah! His new swag is on point 100%, since the accident! Like this is a totally transformed 8	0
39	ablate	Davanagere, India	This is cool and all these days I have been doing "git push origin CURRENT_BRANCH_NAME". Y	0

Figure 2: Sample Dataset

It illustrates the sample dataset, showcasing features like tweet content, disaster keywords, location data, and the target label, all contributing to the predictive power of the deep learning models for classifying disaster-related tweets.

Results And Discussion

Experimental analysis:

Timely disaster tweet detection plays a crucial role in effective disaster response and mitigating the spread of misinformation. The combination of DistilBERT and Random Forest models in this study enhances the accuracy and reliability of tweet classification, ensuring that real disaster-related tweets are identified correctly. This methodology, when integrated into real-time disaster monitoring systems, facilitates swift action, improves coordination, and enhances the accuracy of emergency communications. The model using DistilBERT achieves an impressive test accuracy of 95%, showcasing its superior ability to understand and classify disaster-related content refer Table 1.

The classification report reveals that for real disaster-related tweets (class 1), the model demonstrates a precision of 92% and a recall of 87%, reflecting its strong ability to identify true disaster-related content while minimizing false positives. The F1-score for real disaster tweets stands at 89%, indicating a solid balance between precision and recall, making it highly reliable for disaster tweet classification.

In comparison, the Random Forest model, which is used for handling the structured features in the dataset, achieves an accuracy of 85%. For real disaster-related tweets, the precision is 82%, and the recall is 90%, indicating that while the model does well in identifying true disaster tweets, it still has a slightly higher false positive rate than DistilBERT refer **Figure 4**. The F1-score for the Random Forest model is 86%, which demonstrates good performance but slightly lower than that of DistilBERT.

Overall, the study highlights the advantages of using DistilBERT for natural language processing tasks in disaster tweet detection, achieving higher accuracy and a better F1-score. The Random Forest model, while slightly

less accurate, still performs robustly, showing strong recall and utility in handling the feature set. This combination of models enables a comprehensive approach to disaster tweet classification, with DistilBERT excelling in text understanding and Random Forest providing solid performance in processing structured data, together offering a powerful solution for real-time disaster monitoring.

The efficiency of the proposed system Natural Calamity detection in real time is encountered by evaluating the performance parameters. Parameters like Precision, recall, accuracy, Confusion Matrix are calculated in order to evaluate the performances of the Natural calamity detection system in real time.

Table 1: Accuracy Range of different Classifications

Model	Accuracy Range	Best Use Case
SVM	70-85%	Binary classification (disaster vs nondisaster)
Naive Bayes	65-80%	Simple classification tasks, fast processing
CNN	80-90%	Text classification, feature extraction
DistilBERT	85-95%	Contextual understanding
Random Forest	80-90%	Structured/tabular data

Confusion Matrix by labels Given: The Predicted results for Natural Calamity Detection in real time using AI by labels given for two different classification algorithms refer **Figure 3**.

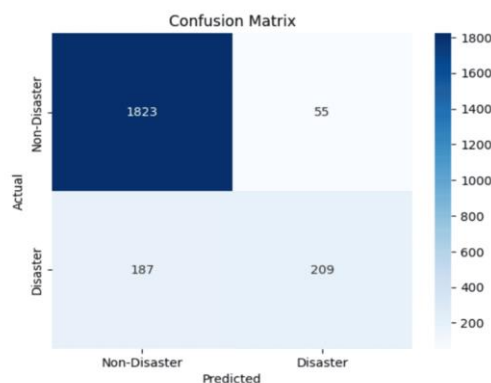


Figure 3: Confusion Matrix of the Random Forest Model

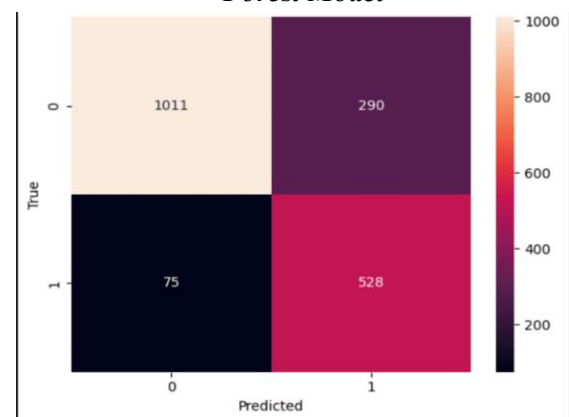


Figure 4: Confusion Matrix of the DistilBert Model

Conclusion

The integration of machine learning models for disaster tweet classification has significantly improved real-time disaster response efforts. By effectively identifying and categorizing real disaster-related tweets, the system enhances the ability of emergency responders to act quickly and make informed decisions. This methodology ensures the timely dissemination of critical information and minimizes the impact of misinformation, which is crucial in disaster scenarios. This system leverages advanced machine learning techniques, including DistilBERT and Random Forest, to create a powerful system for disaster tweet classification. The DistilBERT model excels at understanding the context and nuances of tweet text, achieving high accuracy and an impressive F1-score, while the Random Forest model provides strong support for handling structured features. Together, these models offer a robust solution for real-time disaster monitoring and effective decision-making. Through thorough data preprocessing, feature extraction, and model training, the system delivers highly accurate and reliable predictions. The ability to classify tweets as either real or fake disaster-related content enhances situational awareness, supporting efficient and effective responses to disasters. Additionally, the system's integration with a SQL database for tweet tracking ensures long-term monitoring and the potential for future analysis.

Additional Information

Disclosures

Human subjects: All authors have confirmed that this study did not involve human participants or tissue. **Animal subjects:** All authors have confirmed that this study did not involve animal subjects or tissue. **Conflicts of interest:** In compliance with the ICMJE uniform disclosure form, all authors declare the following: **Payment/services info:** All authors have declared that no financial support was received from any organization for the submitted work. **Financial relationships:** All authors have declared that they have no financial relationships at present or within the

previous three years with any organizations that might have an interest in the submitted work.

Other relationships: All authors have declared that there are no other relationships or activities that could appear to have influenced the submitted work.

Data Availability Statements

The datasets (and/or code) supporting this study are available from the corresponding author upon reasonable request.

The data (and/or code) supporting this study are openly available in Kaggle at <https://www.kaggle.com/datasets/shailjakanttiwari/nlp-with-disaster-tweets?select=train.csv>, reference shailjakanttiwari/nlp-with-disaster-tweets.

References

- [1] Caragea, C., Silvescu, A., & Tapia, A. H.: *Identifying informative messages in disaster events using convolutional neural networks*. 2016.
- [2] Flint, K., & Kellerman, Z.: *NLP disaster tweets classification*. Proceedings of Student Research and Creative Inquiry Day, 2020, 4.
- [3] Gulesan, O. B., Anil, E., & Boluk, P. S.: *Social media-based emergency management to detect earthquakes (and other events) using SVM, KNN & Naive Bayes*. 2021.
- [4] Hernandez-Suárez, A., Sánchez-Pérez, G., Toscano-Medina, K., et al.: *Using Twitter data to monitor natural disaster social dynamics: A recurrent neural network approach with word embeddings and kernel density estimation*. Sensors, 2019, 19:1746. DOI: 10.3390/s19071746.
- [5] Hettiarachchi, H., Adedoyin Olowe, M., Bhogal, J., et al.: *Embed2Detect: Temporally clustered embedded words for event detection in social media*. arXiv, 2020.
- [6] Imran, M., Castillo, C., Diaz, F., et al.: *AIDR: Artificial intelligence for disaster response*. 2014, 159–162. DOI: 10.1145/2567948.2577034.
- [7] Kaur, S., Gupta, S., Singh, S., et al.: *A review on natural disaster detection in social media and satellite imagery using machine learning and deep learning*. International Journal of Image and Graphics, 2021,

22:2250040. DOI: 10.1142/S0219467822500401.

[8] **Kireyev, K., Palen, L., & Anderson, K.:** *Applications of topic models to analysis of disaster-related Twitter data.* 2009.

[9] **Mahajan, P., Raghuwanshi, P., Setia, H., et al.:** *A multi-model approach for disaster-related tweets: A comparative study of machine learning and neural network models.* Journal of Computers, Mechanical and Management, 2024, 3:19–24. DOI: 10.5281/zenodo.1234567.

[10] **Sakaki, T., Okazaki, M., & Matsuo, Y.:** *Earthquake shakes Twitter users: Real-time event detection by social sensors.* 2010, 851–860. DOI: 10.1145/1772690.1772777.

[11] **Sufi, F. K., & Khalil, I.:** *Automated disaster monitoring from social media posts using AI-based location intelligence and*

sentiment analysis. IEEE Transactions on Computational Social Systems, 2024, 11:4614–4624. DOI: 10.1109/TCSS.2022.3157142.

[12] **Suleman, M. K., Riegler, M., Pogorelov, K., et al.:** *Floods relevancy and location from Twitter posts using BERT.* 2023.

[13] **Zhang, H., Li, Z., & Wang, C.:** *Fusing social media, remote sensing, and fire dynamics to track wildland-urban interface fire.* ISPRS International Journal of Geo-Information, 2019, 8:264. DOI: 10.3390/ijgi8060264.

[14] **Zhang, S., & Vucetic, S.:** *Semi-supervised discovery of informative tweets during the emerging disasters.* arXiv, 2016. DOI: 10.48550/arXiv.1610.03750.

[15] **Zisad, S. N., & Hasan, R.:** *Comparative analysis of transformer models in disaster tweet classification for public safety.* arXiv, 2025.