

Twitter-Based Women's Safety Analysis (ML)

¹ S.Akhila, ² B.Ramya,

¹Assistant Professor, Megha Institute of Engineering & Technology for Women, Ghatkesar.

² MCA Student, Megha Institute of Engineering & Technology for Women, Ghatkesar.

Abstract—

— It is very important to ensure the protection of women in public places in today's society. Twitter has become a leading social media platform due to the increasing popularity of informative database that tracks incidents and public sentiment around women's safety in real-time. This research use machine learning methods to analyze Twitter data in order to get insight into the nature and frequency of debates around women's safety. The proposed method includes mining relevant hashtags and phrases to compile a large dataset of tweets concerning the safety of women. We will use Natural Language Processing (NLP) methods to preprocess the text data and extract important features. Then, we will use sentiment analysis to categorize tweets as positive, negative, or neutral. To further examine the use of Named Entity Recognition in tweet sentiment and context prediction, we will train machine learning models, such classification algorithms, using labeled datasets. The study's overarching goal is to identify problem areas, trends, and patterns related to women's safety. Also, we will use time series analysis to see how the frequency of events varies over the years. The study's goal is to build a prediction system that can identify potential safety risks by analyzing data from Twitter. The results of the research might help policymakers, community groups, and law enforcement agencies take preventative actions to make women's lives safer.

Keywords—

topics such as women's safety, sentiment analysis, social media, machine learning, and Twitter

I. INTRODUCTION

A popular social networking platform, Twitter allows users to openly share their thoughts, join discussions, and express themselves. Still, it's not completely immune to issues with threats, harassment, and abuse, similar to any other online space. Women often face unique challenges and concerns about their safety while utilizing Twitter. A hostile and threatening environment for women may often be

magnified by the platform's ease of communication and anonymity, leading to hazardous behaviors. Many customers use social media sites like Twitter to broadcast their thoughts, feelings, and assumptions to the world. It is possible to take these tweets and run them through an extreme expressions test if one has a good grasp of how to rank women's well-being in certain areas. By way of illustration, a considerable number of people use Twitter to convey their thoughts and feelings to others all over the globe. A good understanding of how to rate female security in their assigned region allows for their rapid removal and subsequent harsh trial of terms. If a tweet was made on Twitter containing the phrases "lady harassments," "lady security," or the hashtags "lady harassment" and "lady well-being," we could get it all via Twitter's API. We collect all the tweets, sort them into datasets, and then isolate the ones with the same polarity using our approach. The safety of women is a major issue in today's society, and there is a wealth of information at our fingertips thanks to the proliferation of social media platforms like Twitter. In particular, the popular microblogging site Twitter offers a unique opportunity to analyze the discourse around women's safety, identify trends, identify problems, and brainstorm solutions. Twitter data may teach researchers and lawmakers a lot about popular sentiment, concerns, and experiences about women's safety.

In recent years, social media platforms such as Twitter have become indispensable for communication, idea sharing, and joining worldwide discussions. However, concerns about online safety, especially for women, have grown in tandem with the popularity and impact of these networks. A growing number of people are worried about the safety of women on Twitter, which highlights the challenges and risks that women face when using the internet. The platform has both empowered women and encouraged harassment, abuse, and sexism as a result of its massive user base and unrestricted access. This article delves into the topic of women's safety on Twitter, discussing the many forms of harassment they face, its impact on their well-being, and the efforts being made to address these issues. already present; The different font styles are shown throughout the article and are indicated by italicized

text inside parentheses after each example. Images, tables, and multi-leveled equations are optional additions, even though there are a lot of table text styles to choose from. The formatter is responsible for creating these components and ensuring that they meet all of the standards specified below.

II. RELATED WORK

TABLE I.

Sr.no	Paper Name	Authors	Year	Methodology Used
1.	Women's Safety Analysis on social media using Machine Learning	Ms V Bhavani Boddu Pavan Ganesh,Bonthu Sai Kiran,Bonthu VenkataNaveen, Kallakuri Baladitya	2020	TF-IDF(term frequency-inverse document frequency),Decision Tree
2.	Analysis Of Women Safety In Indian Cities Using Machine Learning On Tweet	Sheema Nargis,Saadiya,Asma Namera	2022	Linear Regression,Random Forest
3.	Women Safety Analysis based on tweets using Machine Learning	Mrs. Chandini U1, Harika M2, Chandana K P3,G Menaka4, Harshitha s5	2021	Naive Bayes Algorithm,Support Vector Machine(SVM)
4.	Predicting Safeness of Women in Indian Cities using Machine Learning	Nanditha P1, Sindhu B M2, R Geetha3	2022	Naive Bayes,Bag Of Words
5.	Robust Sentiment Detection on Twitter from Biased and Noisy Data	Luciano Junlan Feng AT&T Labs - Research	2020	Support Vector Machine(SVM), WEKA(Machine)

1) To verify the legitimacy of a certain tweet, this study employs two algorithms. In order for people to make informed choices about the protection of women, by verifying that tweets are genuine. Another innovative feature included in this research is the ability to extract the average frequency of each keyword from tweets and display it as a numerical

vector. The Decision Tree method will utilize this TFIDF vector to determine the authenticity of a tweet. To determine whether a tweet is genuine, one may use a Decision Tree algorithm that has been taught to differentiate between real and fake terms [1]. There are millions of tweets and text messages sent every single day on Twitter, and our research study used a number of machine learning approaches to help us sift through and evaluate all of that data. Machine learning methods like the SPC algorithm and linear algebraic factor model techniques are very useful and efficient when it comes to analyzing large volumes of data. Further meaningful categorization of the data is achieved with their help. Support vector machines are another popular machine learning technique for discovering how safe Indian cities are for women based on data retrieved from Twitter [2]. 4) In order to assess data and generate intelligent classifications, this research looked at the various sentiment analysis methodologies. For data categorization, the article makes use of the Sentiment approach, a Lexicon-based technique. We can do further classification using a hybrid strategy or an approach based on machine learning to get the classification's subgroups. 5) A few well-known classifiers that may be used for thorough attribute-based categorization while considering factors impacting women's safety [3]. 6) Machine learning methods have been brought up at various points in the project. With the help of machine learning algorithms, Twitter is able to better organize and analyze the daily data set that contains millions of tweets and messages. The SPC and linear algebraic algorithms are two examples of effective methods for sorting and extracting useful information from large datasets. Consequently, we may improve women's safety via sentiment analysis and the use of machine learning algorithms if we enhance awareness. More complete integration of the proposed current philosophy into the Twitter application interface will increase safety by enabling sentiment analysis to be applied to millions of tweets [4]. 7) A dependable and effective strategy for identifying the tone of Twitter postings was presented in a study article. The model is built using biased and noisy labels. The following facts are true and contribute to this performance: (1) a method generates a less concrete picture of these communications than earlier methods that relied on raw word representations; and (2) data sources provide reasonable-quality labels while being noisy and skewed. Their distinct biases meant that combining them also yielded some useful results [5].

III. EXISTING WORK

1. Collect Data: Compile a collection of text records that include a variety of feelings (positive, negative, and neutral). A varied and inclusive dataset is ideal, standing in for the domain of the application.

2. Preprocessing Data: Remove any typos from the text data by excising non-essential letters, symbols, and numerical values. Break down the phrases into their component words and change their case to lowercase. Stemming or lemmatization should be performed after removing stop words. Third, identify the data with a positive, negative, or neutral sentiment depending on the overall sentiment of the statements.

4. Dataset Splitting: To evaluate the performance of the model, divide the dataset into two parts: one for training and another for testing.

5. The Extraction of Features: For sentiment analysis, use the TextBlob library, which is based on the NLTK library. When it comes to sentiment analysis and other common NLP tasks, TextBlob provides an easy-to-use API.

Another option for sentiment analysis is to utilize the NLTK library's Naive Bayes algorithm. Apply the labels to the training dataset in order to train the Naive Bayes classifier. TextBlob comes pre-trained for sentiment analysis, so there's no need to explicitly train it for model 6. Nonetheless, if necessary, you may adjust its behavior with the use of bespoke training data. Use the labeled training dataset to train the classifier for Naive Bayes. Use the trained TextBlob model or the Naive Bayes classifier to predict the sentiment of phrases in the testing dataset for sentiment classification, which is the seventh application.

8. Model Evaluation: Evaluate the efficacy of sentiment analysis models by measuring their performance using measures such as accuracy, precision, recall, and F1 score. Step 9: Deploy the model for real-world applications if you're satisfied with its performance. This will allow automatic sentiment analysis and categorization in textual data. Keeping an eye on the model's progress and adding fresh data as needed will ensure its correctness as time goes on.

IV. PROPOSED METHODOLOGY

1. Collecting Data: Make sure your dataset is diverse by include both statements that are related to women and those that are not. Achieving a harmonious contain comments that span positive, negative, and neutral sentiments in the collection. Make sure to identify the dataset with the gender of the people to whom each statement is addressed. To clean up the text data, remove any unnecessary letters, special

symbols, and numbers as part of the data preparation step two. To make sure everything is consistent, convert the sentences to words using tokens and lowercase them. Remove stop words and use stemming or lemmatization to reduce the dataset's dimensionality. Third, identify the data with a positive, negative, or neutral emotion based on how the terms are generally felt. To make it easier to find the parts of phrases that pertain to women, you may add a label to them.

4. Dataset Splitting: To evaluate the performance of the model, divide the dataset into two parts: one for training and another for testing.

5. Extraction of Features: Apply TF-IDF (Term Frequency-Inverse Document Frequency) or bag-of-words format to numerically characterize the textual input. If the remark makes reference to women, then sort the traits by that fact.

6. construct the Model: Apply the features extracted from women-related terms to construct a multinomial naive bayes classification model. Words unrelated to women may be included for sentiment analysis using a model based on neural networks, such Transformer or LSTM.

7. Analyze Sentiment: Make use of the algorithms that have been taught to predict the emotional tone of a text. Sort statements about women into positive, negative, or neutral categories using the Multinomial Naive Bayes model. When doing sentiment analysis, use the model based on neural networks if the comment is unrelated to women.

8. Model Evaluation: Evaluate the performance of each model using the testing dataset. To assess the performance of the models, use metrics such as F1 score, accuracy, precision, and recall. Evaluate the two models using the test data. To assess the models' performance, use metrics like as F1 score, accuracy, precision, and recall.

9. Put it into action: When the model has shown to be satisfactory, put it to work in a real-world setting to automatically detect and classify texts that discuss women.

10. Revisions and Observation: To maintain accuracy over time, monitor the model's performance often and add new data as appropriate.

V. ALGORITHMS USED

An aspect of NLP known as sentiment analysis is multinomial naive bayes (MNB), which is concerned with the detection and categorization of It is inherent to the text's emotional tone. Among its many uses, it is crucial for understanding public sentiment, analyzing consumer reviews, and monitoring online conversations. It is usual practice to classify sentiment as either good, negative, or

neutral.

One of the most popular probabilistic algorithms in machine learning for classification problems is Multinomial Naive Bayes (MNB), which is based on Bayes' theorem. When the input data is described in terms of word frequencies, MNB shows efficacy, especially in the field of sentiment analysis. It is assumed by the algorithm that the presence of a word inside a document does not rely on the existence of any other words, offering a simple yet often useful assumption for modeling. Classifying texts is a typical use case for MNB in the process of sentiment analysis entails classifying the input text according to predetermined criteria. These characteristics are often word frequencies in sentiment analysis. Computing the likelihood that a text belongs to a given emotion category is at the heart of the Multinomial Naive Bayes (MNB) method. The anticipated emotion for the document is then determined by identifying the category with the greatest likelihood. When the input is a representation of a bag of words, MNB excels because it can deal with the high-dimensional and sparse nature of text data. The computational efficiency and unexpectedly good performance of MNB in reality are due in large part to its simplicity and the assumption of feature independence. In general, MNB has gained popularity in sentiment analysis for its simplicity, ease of implementation, and good performance in cases when the input data is mostly textual. An effective algorithm in NLP, it uses transcend beyond sentiment analysis to include a wide range of text categorization tasks. Chapter B. TF-IDF The numerical statistic known as VECTORIZER TF-IDF (Term Frequency-Inverse Document Frequency) indicates the importance of a term within a document that pertains to a larger set of papers called a corpus. As a common method for converting a set of text documents into numerical vectors, TF-IDF vectorization is finding widespread use in sentiment analysis and natural language processing (NLP).

1) The frequency of terms (TF): The measure of how frequently a term or phrase appears in a text is called Term Frequency, or TF for short. One way to find this out is to look at the document's word count and compare it to the frequency of a certain word. According to mathematical definitions, TF is:

$$TF(t,d) = \frac{\text{(Total number of terms in document } d\text{)}}{\text{(Number of times term } t \text{ appears in document } d\text{)}}$$

The second metric is the inverse document frequency (IDF), which is a way to find out how often a key term appears in a collection of documents. By providing expressions that are not prevalent across the whole corpus, they carry more weight, which

helps to find terms that are unique to a certain text. Using mathematics, we can define IDF as: is equal to the logarithm of IDF(t,D) (Total word count in the corpus divided by word count including the phrase t+1)

If a phrase appears in all of the documents, adding 1 to the denominator will prevent division by zero.

3) Vectorization using TF-IDF: Combining Term Frequency (TF) with Inverse Document Frequency (IDF), TF-IDF vectorization creates a numerical representation of a document. How it works requires taking each phrase in the document and multiplying its TF and IDF values. A high-dimensional representation, with each dimension representing a unique phrase in the whole corpus, is produced by the TF-IDF vector. The importance of each phrase in the paper is shown by its value along each dimension.

4) Sentiment Analysis Use Case: One common method for representing textual input in a way that machine learning algorithms can understand is TF-IDF, which is used extensively in sentiment analysis. By considering both the frequency of words inside the text and their scarcity throughout the whole corpus, this technique effectively captures the importance of words within a document. By determining the relative significance of words in texts, TF-IDF vectorization becomes an important tool in sentiment analysis and natural language processing (NLP). This method is useful for creating numerical representations that can be fed into machine learning models. These models may then perform tasks such as clustering, regression, or classification.

C. Flow Diagram-

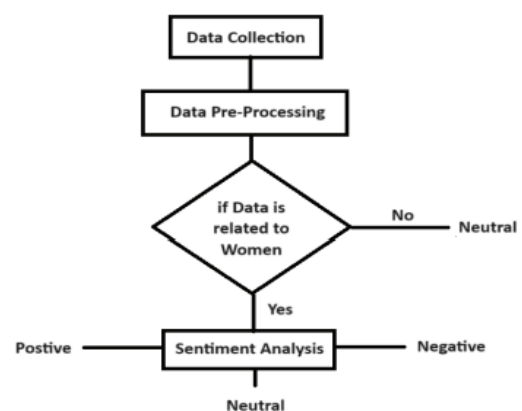


Fig. 1.

VI. RESULT



Fig. 2

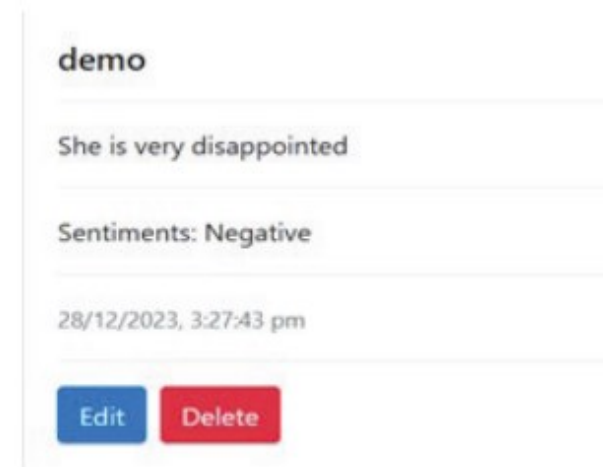


Fig. 3.

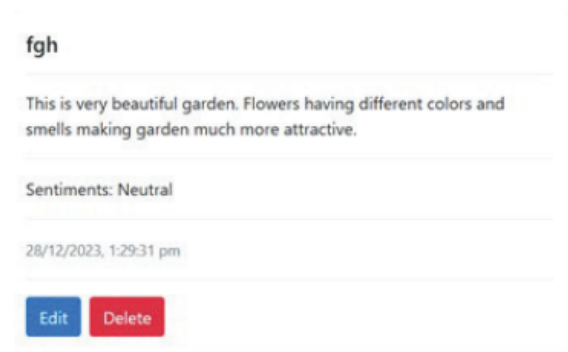


Fig. 4.

VII. CONCLUSION

The public's sentiment toward women's safety may be better understood with the use of machine learning techniques such as TFIDF and Naive Bayes. When it comes to reading people's emotions in their writing, these techniques really shine. When used in tandem, these techniques really shine. When used in tandem, When it comes to the protection of women, we can learn more about what people really care about. A simple and intelligent algorithm, Naive Bayes can tell us if something is good or bad in a flash. TF-IDF aids in the comprehension of spoken language by drawing attention to key words that occur often. Using these techniques, we can learn how most people feel and what they believe about the safety of women. But we must not forget that these instruments are not without their flaws, such as the fact that they are susceptible to biased data and linguistic shifts. To ensure these tools are trustworthy and equitable, they must undergo continuous development and thorough evaluation. All things considered, these machine learning methods provide a practical means by which we may garner insight from public opinion and strive to create spaces that are less hazardous to women.

REFERENCES

- [1] "Women's Safety Analysis on social media using Machine Learning" by Ms V BhavaniBodduPavanGanesh,BonthuSaiKiran,Bonthu VenkataNaveen,KallakuriBaladitya(2020)[1].
- [2] "Analysis Of Women Safety In Indian Cities Using Machine Learning On Tweet" by SheemaNargis,Saadiya,AsmaNamera(2022)[2].
- [3] "Women Safety Analysis based on tweets using Machine Learning" by Mrs. Chandini U1, Harika M2, Chandana K P3,G Menaka4, Harshitha s5(2021)[3].
- [4] "Predicting Safeness of Women in Indian Cities using Machine Learning" by Nanditha P1, Sindhu B M2, R Geetha(2022)[4].
- [5] "Robust Sentiment Detection on Twitter from Biased and Noisy Data" by Luciano ,Junlan Feng AT&T Labs - Research. (2020)[5].