

Investigation of Feminism Trends Through Sentiment Analysis Using Machine Learning and Natural Language Processing

Sana Sohail¹, Furqan Amjad²

¹ Institute of Policy Studies, Islamabad, Pakistan

² TalTech University, Tallinn, Estonia

*Correspondence: Sana Sohail waa992@gmail.com, furqan.amjad123@gmail.com

Citation | Sohail. S, Amjad. F, "Investigation of Feminism Trends Through Sentiment Analysis Using Machine Learning and Natural Language Processing", IJIST, Vol. 6 Issue. 1 pp 301-319, Mar 2024

DOI | <https://doi.org/10.33411/ijist/202461301319>

Received | Mar 04, 2024 : **Revised** | Mar 24, 2024 : **Accepted** | Mar 26, 2024 : **Published** | Mar 29, 2024.

Introduction/Importance of Study:

One of the recent changes seen in Pakistan is the growing awareness among people, to end gender discrimination and bring equality, across various spheres. "Aurat March", is a series of rallies that began in 2018 to mark International Women's Day. People across the nation, comment on these rallies through social media.

Novelty Statement:

The response to the "Aurat March", held annually since 2018, was mixed and no analysis of Twitter data had been previously done to investigate the polarity of comments, through Machine learning and Natural Language Large Language Models.

Material and Method:

For this, Sentiment analysis was performed, using Machine learning and NLP techniques, on the pre-processed data. Lexical rule-based VADER and transformer-based pre-trained Large Language Models were used to check the polarity of Twitter comments.

Results and Discussion:

The best results were achieved through RoBERTa-LARGE, which were closest to the Human Labelled Data, hence validating the accuracy of the LLM. On the other hand, VADER results were clearly far from the manually labeled results. The sentiment analysis that was applied the first time on "Aurat March" tweets, gave us satisfactory results, and we were further able to validate our research, by comparing the models' accuracy with human-labeled results. Consequently, by analyzing the sentiments expressed on Twitter, we were able to discern the general mindset of users and gain insights into prevailing trends.

Concluding Remarks:

This analysis provided us with a reasonably accurate gauge to assess the perception of feminism over the past few years, allowing us to evaluate whether it has garnered fame or faced defamation in the public discourse.

Keywords: Machine Learning; Natural Language Processing; Large Language Models; feminism; sentiment analysis; textual data.



Introduction:

Gender discrimination refers to the unjust treatment of individuals or groups based on their gender identity. This discrimination can manifest in various forms, at the suppression of either females or males, within specific contexts or circumstances. There might be diverse motives behind the unfair treatment, but the consequences are unavoidable. According to the UNDP's new Index, approximately half of the global population believe that men make better leaders, particularly political leaders, and over 40% believe that men are inherently suited for corporate CEOs, meaning that men have a superior position when jobs are scarce. Alarming, a significant 28% of individuals surveyed find it acceptable for a man to resort to violence against his partner [1].

Despite possessing qualifications equal to or more than those of male colleagues, women often demonstrate themselves more modestly and are less likely to perceive themselves as qualified for top positions. Studies have shown that women are required to be approximately two-and-a-half times as productive as men to receive the same recognition and advancement opportunities in the workplace. Moreover, women are burdened with the essential responsibility of caring for their children, which often acts as a significant barrier to wholeheartedly pursuing their professions. This challenge is particularly pronounced for women in societies like Pakistan, where cultural norms dictate that women are primarily responsible for managing household affairs and child care, since there is a stereotype concept and a traditionally held belief, making the society accept that men are solely “bread-winners and women duties are restricted to “house makers”.

Motherhood is less negotiable than fatherhood in such a social setting [2]. Although the concept of work-life balance originated in the West and has been extensively researched in this context, it is now spreading to the Eastern Hemisphere as a result of globalization, impacting organizational dynamics worldwide [3]. In emerging and underdeveloped countries, barriers to women's employment are more prevalent. Nations must utilize this potential workforce for the betterment of economic and societal goals, which necessitates a greater awareness of women's challenges, both large and small, in emerging countries. According to Madsen and Scribner [4], there has always been a lack of gender research in specific populations around the world. Pakistan, among other underdeveloped countries, is underrepresented in such studies, hindering efforts to address gender disparities effectively.

Pakistan is a nation with the largest percentage of young people (between 15-29 years) ever recorded in its history, with women comprising half of the young population (aged 15-29) [5]. Having said this, a large number of women in Pakistan, are forced to abandon their careers due to domestic and household activities, as well as the non-cooperative environment, in their workplace. Despite the fact that policies such as harassment legislation exist to improve female employment, most female employees are unwilling to report such incidents [6], for fear of being labeled as immodest. This reluctance stems not only from individual choices but also from broader societal norms, such as gender segregation beliefs ingrained in females from a young age. Addressing these macro-level contextual factors is essential for creating a more inclusive and supportive environment for women's participation in the workforce [1].

Most studies overlook the macro-level factors that drive gender discrimination at the meso-organizational level, and prejudice in work-related activities such as recruitment, selection, training, career development, and equal compensation and benefits persist [7]. Male CEOs and employees are frequently overwhelmed by sociocultural aspects from society's larger context and apply them at work, resulting in a gender imbalance in the workplace [8]. In Pakistan, there is a growing awareness among people to address gender discrimination and promote equality, regardless of gender, ethnicity, or any other bias that impacts individuals in society. This shift reflects a broader movement aimed at acknowledging and supporting marginalized groups within the nation.

This movement can be identified as the fourth wave of feminism that is continuing, and sweeping across the nation and in the previous few years, is even stronger than before. One main difference between this fourth wave and the previous three waves is the role of social media and technology, which had not been this evident in previous ones. Also, the fourth wave is complex. It incorporates many movements that both complement and clatter with each other [7].

Problem Area:

Increased participation of females in various spheres of society has brought about pressures and challenges, particularly in balancing career and personal responsibilities, due to the disproportionate burden of household duties [1]. In Pakistan, the "Aurat March" has emerged as a significant platform for advocating women's rights since its inception in 2018. Women's collectives organized the inaugural Aurat Marches on International Women's Day, coinciding with the "Pakistani #MeToo campaign." On March 8, 2018, the first march took place in Karachi. Marches were organized in 2019 by "Hum Aurtein" (a women's collective) in Lahore and Karachi, as well as elsewhere in the country by Women Democratic Front (WDF), Women's Action Forum (WAF), and other organizations in Islamabad, Hyderabad, Quetta, Faisalabad, and other towns [9].

The march, which included representatives from a number of women's rights organizations, was backed by the Lady Health Workers Association [10]. Aurat Marches were also held in a number of Pakistani cities in 2020 and 2021, with the "Aurat Azadi March" taking place in Islamabad, Sukkur, and Multan. Pakistani ladies select a theme each year that matches the year's most popular issue. The organizers of the Lahore march, for example, chose "Women's Health Crisis" as the march's theme in 2021 to raise attention to the COVID-19 pandemic's consequences, and chose one large poster to portray the women's health problems in the country. Through these rallies, Pakistani women seek to raise awareness about the challenges faced by marginalized communities and advocate for gender equality and social justice [5]. Nevertheless, the resilience and determination of Pakistani women in advocating for their rights and those of other marginalized groups signify a growing momentum towards positive change in the country. Whether these marches will gain more prominence or eventually fade away remains to be seen, but their significance in challenging societal norms and advocating for equality cannot be overstated.

Problem Statement:

There is no comprehensive dataset available to analyze feminism trends in Pakistan. This research aims at developing the first comprehensive dataset on Aurat March tweets, exploring feminism trends, through sentiment analysis, using Machine Learning and NLP techniques, hence comparing the results to observe the accuracies of different AI Models applied for sentiment analysis.

Research Question:

What are the sentiments of Pakistan Twitter users on feminism, from 2018 to 2022?

The objective of the Study:

According to the United Nations Development Program (UNDP), index, biases against women persist in 90 percent of the population. These invisible yet formidable barriers contribute to women feeling pressured in their quest for equality [11]. The purpose of this study is to investigate the increasing popularity that feminism has gained in recent years in Pakistan and to analyze the trend that has penetrated through the nation, during the past few years. Leveraging data from Twitter, the study seeks to fill the gap in sentiment analysis regarding feminism-related tweets by employing Machine Learning and Natural Language Processing (NLP) techniques. Additionally, the accuracy of the models utilized in sentiment analysis is evaluated by comparing them with human-generated results.

Literature Review:

The condition of a nation is intricately tied to the condition of its women, as the development and progress of any society are contingent upon ensuring the rights and opportunities of all its citizens without discrimination based on gender. The concept of feminism has been developed just because of the unequal distribution of social, economic, and political rights among males and females. Where feminist theory scrutinizes gender disparities across various aspects of life, highlighting biased working environments that disadvantage women globally. These unfriendly working conditions are playing a vital role as a barrier to women's professional growth. There is a need to identify barriers to providing equal working conditions for both genders [6].

In conducting our literature review, our focus has been twofold. Firstly, we examined the previous and present research related to the issues of females in Pakistan, and secondly, the methodologies, particularly those involving the analysis of datasets, to gain insights into how sophisticated analyses can be conducted to shed light on these important issues. The research encompassed a comprehensive examination of various aspects pertaining to issues faced by women, the concept of feminism internationally and in Pakistan, followed by the research papers and work done previously through Data Science relevant to Sentiment Analysis or Textual Data, etc. Among all, the following studies are worth mentioning and most relevant to our work:

Pakistan has a total population of 208 million people, with women comprising 49 percent [12]. A significant portion of the population, approximately 29%, falls within the age bracket of 15 to 29 years, with half of this demographic being women. With rural regions accounting for 63 percent of Pakistan's populace, the bulk of woman youngsters are likewise rural. The urgent need to focus on empowering young girls in Pakistan is underscored by the country's dismal performance in the World Gender Gap Index 2020, where it ranks 151 out of 153 countries. While it is essential to mobilize and provide opportunities for all youngsters to contribute to Pakistan's development as a strong, dynamic, and innovative nation, addressing the specific challenges faced by young girls is paramount due to the pervasive gender inequality prevalent in the country [13].

Pakistan's Constitution acknowledges the importance of women's rights and provides for affirmative action and special measures to promote gender equality and empower women. Pakistan has committed itself to multiple conventions like the Beijing Platform for Action (PFA), ILO Conventions, Child Rights Convention, and Sustainable Development Goals (SDGs) in order to increase participation of women/girls in economic, social, and political spheres, protection of human, sexual and reproductive rights of women/girls, and eradicating all forms of violence, ensuring women's rightful place in society [14]. Concern for the youth in Pakistan has led to surveys, reports, and policies with recommendations, especially with the view of capitalizing on the youth bulge and reaping a demographic dividend. Vision 2025, a strategic document outlining Pakistan's pathway for the future, recognizes the importance of the youth bulge and aims to capitalize on this demographic dividend [11]. Youth policies have been devised and adopted at the provincial and national levels. The last national policy for youth (15–29 years) was formulated in 2008. Policies for Punjab, Khyber Pakhtunkhwa, and Sindh were approved in 2012, 2016, and 2018, respectively. The draft policy developed in Baluchistan (2015) is yet to be finalized [5].

Research Synthesis:

With more than 100 million users and over 500 million tweets sent on a daily basis, Twitter has established itself as a unique blogging platform. Twitter's vast viewership has always piqued user's interest in expressing and sharing their ideas, opinions, and perspectives on any subject, brand, organization, or other topic selected [3]. Consequently, a variety of businesses, organizations, and enterprises use Twitter as a source of information. Sentiment analysis approaches have been able to use the public data provided by social networking websites for

analyzing sentiments, in sectors such as economy, politics, and finance, thanks to the widespread use of the websites [15]. People using Twitter are allowed to write a tweet of 280 characters. These tweets may have words, informal language, punctuation, emojis, etc. Sentiment analysis is definitely a very effective way to observe and evaluate a number of insights and trends of any topic of interest, as it would help detect the general feelings and mindset of society.

Machine learning has been used for sentiment analysis on textual data to observe the polarity of a sentence. Different ML approaches, either ruled based, or hybrid are applied, to note if a particular sentence is positive, negative, or neutral [16]. Transformers and Large Language Models (LLM) have superseded the majority of the top NLP implementation models. The next step would be to simply decide which transformer-based technique would be best suited for this purpose if you can pinpoint exactly what NLP task you need to conduct. All that's left to do after selecting the best strategy for your voice or text dataset is to fine-tune it or locate a model that has already been tweaked. LLMs are ideal for implementing NLP, due to this and a number of other advantages [17].

Observation and Research Gap:

In the realm of data analysis utilizing machine learning and natural language processing techniques, the systematic identification, extraction, quantification, and study of people's affective and subjective states from various communication platforms have become increasingly prevalent. However, despite the effectiveness of sentiment analysis in discerning societal trends and attitudes, certain gaps and limitations have been identified through previous research endeavors.

Lack of Data Availability:

One notable gap is the absence of a comprehensive dataset specifically focused on feminism in Pakistan. The dearth of such data inhibits thorough analysis and understanding of feminist trends and sentiments within the country.

Untapped Potential of Sentiment Analysis:

While sentiment analysis has proven to be a valuable tool in examining public opinions and reactions across diverse situations and events, there appears to be a notable oversight in its application to "Aurat March" tweets. The absence of sentiment analysis on these tweets hinders the exploration of feminist trends and sentiments within the context of Pakistan's socio-political landscape.

Methodology:

Tweets from the year 2018 to 2022 were scrapped from Twitter, making data sets on a yearly basis. These datasets were then cleaned through pre-processing techniques. Sentiment analysis was performed, using Machine learning and NLP techniques, on the cleaned and pre-processed data. Algorithms like VADER and NLP transformer-based pre-trained large language models like Pysentimiento and RoBERTa-Large, were used to check the polarity of Twitter comments. At a coding level, the various steps were included:

- Reading/Loading the Dataset
- Pre-processing and Cleaning textual tweets
- Performing Sentiment Analysis on five-year data
- Comparative Analysis of Different Models
- Comparative Analysis of Yearly Trends
- Sentiments of Human Labelled Data
- Comparison of Machine Results with Human Results
- Conclusive Inference

General Steps of Sentiment Analysis:

Sentiment analysis (opinion mining) is a text mining technique that automatically analyzes text for the writer's sentiment (positive/negative/neutral) using machine learning and natural language processing. It gets more challenging when the data collected/scraped, is

unstructured and raw. Unstructured data are datasets that have not been structured in a predefined manner. Unstructured data is typically textual, like open-ended survey responses and social media conversations, or in our case, "Tweets".

Figure 1 below, gives a simplistic and generic view of the basic steps taken in sentiment analysis.



Figure 1: Steps of Sentiment Analysis

Raw data that is scraped through API/Scripts is pre-processed to keep quality data and useful information, important for us. Important features are hence extracted for sentiment analysis. Sentiment analysis is performed through Machine learning and NLP models namely, VADER, Psysentimiento, and RoBERTa-LARGE. The data gathered illustrated the users' comments and hence points of view, so as to mention the number of tweets that were positive/negative or just neutral, that were sent in reaction to each year's demonstrations. Finally, comparative analysis is done on the results achieved for each year, to reach a final conclusion.

Flow of Methodology:

The data flow diagram of our study was drawn as given below:

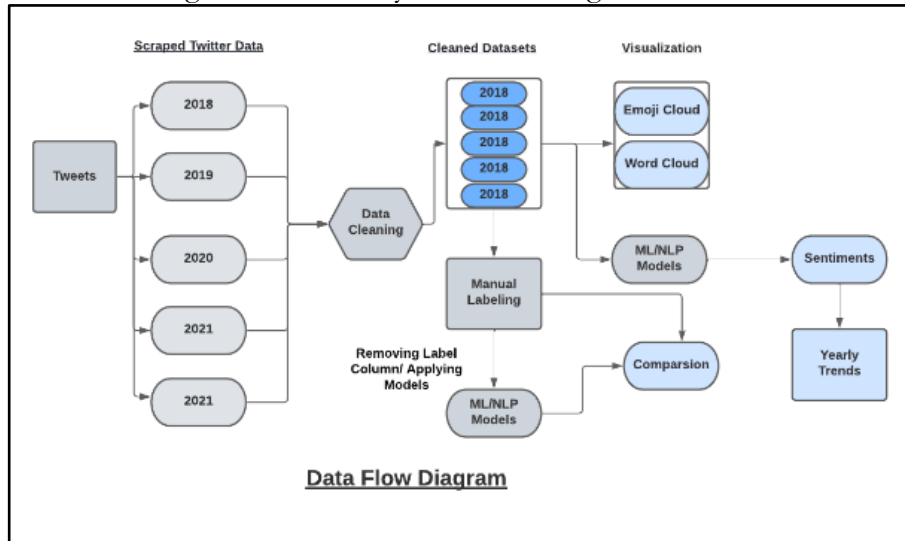


Figure 2: Flow of Methodology

The block diagram is made to make the stages even more understandable. Tweets on the related topic were scraped, for all the years, when the rallies took place. Since these kinds of marches were initiated in 2018, hence the tweets were collected from 2018 to 2022, so as to compile five datasets of textual tweets on "Aurat March". Data pre-processing is an essential step in preparing text data for sentiment analysis, ensuring that it is clean and ready for analysis. It is dependent on the type of analysis that is to be done on the collected/scraped data [18]. In our case, we removed all tags for scraping text out of the source, including HTML entities, punctuation, non-alphabets, and any other characters that aren't part of the language.

We utilized stemming lemmatization to reduce a word's inflectional and occasionally derivationally related forms to a single base form. As a result, stemming/lemmatizing assisted us in reducing the number of overall terms to a few "root" terms. For instance, organizer, organizes, organization, and organized are all reduced to a single root phrase, such as "organize." ML and NLP Models were applied, to observe the response of each model. Through our literature review, we know that the result of the manually labeled data is more accurate than that of a machine [19]. Therefore, a segment of the data was labeled manually, and the sentiments of

this data were evaluated. Next, ML and NLP models were applied for sentiment analysis on the same manually labeled data, in order to have a comparative analysis between manually labeled data results and that of the machines to see which model results were closest to the manually labeled data results.

About the Dataset:

Our research developed special scripts using snscape which is a social network toolkit for scrapping data to overcome limitations imposed by Twitter's API restrictions. Thus, we collected data from 2018, the year "Aurat march" initiated in Pakistan, followed by tweets after the rallies of years 2019, 2020, 2021, and 2022. The keyword included "aurat march", is a union of demonstrations by different organizations, that would have otherwise shown the feminism empowerment in Pakistan individually. Rallies were based on several socio-political contexts each year. The general overall environment was studied and noted to be as follows:

In the year 2018, the women of Pakistan came out to raise their voices for their rights for the first time. The marches were held substantially in Karachi, Lahore, and Islamabad. They were planned and organized entirely by a different group of women belonging to different races, classes, and sections of society. Aurat March 2018 that at one place stressed issues that women are facing in the world, also raised debate on several taglines like "Apna Khana Khud Garam Kero" (Do toast your food yourself) and "Mera Jism Meri Marzi" (My body my will), that remained a center of discussion on social media more than rest. The events in all three metropolises were open to the public and everyone was welcome to join the walk, hence some men and a large representation from different seminaries also showed up in rallies.

Compared to 2018, 2019 rallies grew both in the number of participants, as well as the supporting organizations. It was substantially organized by the "Hum Aurtein Movement", which is a group of feminist women, transgender individuals, and nonages of all classes and religious persuasions who see patriarchal structures as the reasons for women's exploitation. The 2020 Aurat March was held on 8 March in Karachi, Hyderabad, Lahore, Quetta, Islamabad, Sukkur and Multan. The taglines included "Saying 'Mashallah' doesn't make your importunity halal", "Domestic violence kills further than nimbus", "I march so one day my daughters will not have to", "Imagine not loving the women in your life enough to endorse for their rights". Men held signs saying, "I'm girdled by the contrary gender and I feel safe. I want the same for them", "Proud hubby of a feminist, proud father of a feminist, proud feminist" etc. still, some taglines were largely condemned again there were all kinds of response to the marches, both pro and anti-march sentiments were openly shared in mainstream media, and social media followed suit.

The associations organizing the marches provided a specific theme that was used to plan the marches in 2021. For instance, the Lahore march's organizers chose the topic "Women's Health Crisis" to highlight the effects of the COVID-19 scourge on Pakistani women. They put together bills showing ladies in poor health. Additionally, the organizers of the Karachi march provided a sit-in opposing the patriarchal violence as its main focus. More options for female and ambisexual representation on hospital croaker-legal crews were needed by a decree. Around 2000 women reportedly participated in the 2022 Lahore "Aurat March," as per a news item in the Daily Times [10]. The march organizers set up an exhibit at the starting point of the march called "Journalism Must Be Ethical," which featured cardboard cutouts of various media intelligencers who were alleged to have misrepresented or misreported the Aurat March, worn-out marching women, or posted images of the marchers with malicious intent. The Friday Times reports that despite some police-maintained fencing Men from the "Haya March" (Morality March), which passed in close proximity to the Aurat March in Lahore at a mere 200-meter distance, attacked the ladies there, prompting the authorities to order the women in the Aurat March to end their march abruptly.

Data Pre-processing/Cleaning:

Data pre-processing is essential for data mining and principally involves converting raw data into an accessible format for NLP models. Real-world data is frequently partial, inconsistent, and/or lacking in certain actions or trends, and is likely to contain numerous breaches. Data pre-processing is a proven system of resolving similar issues. This helps in achieving reliable results through the classification of algorithms. Important stages that were performed, in data pre-processing are given below.

Tokenization:

This is the process of separating a mass of text into tokens, which can be words, phrases, symbols, or other significant components. For further processing, the list of tokens serves as input. With the help of the NLTK Library's "word tokenize" and "sent tokenize" functions, you can easily turn a long block of text into a set of words or finds.

Word Stemming/ Lemmatization:

Both procedures are intended to reduce each word's inflectional forms to a single base or root. Stemming and lemmatization are closely linked. The distinction is that a stemmer only processes a single word without taking the context into account and cannot distinguish between words with varying meanings based on the part of speech. Stemmers are often quicker and easier to deploy, therefore for some applications, the decreased accuracy may not be a factor. Table 1 gives an idea of stemming and lemmatization.

Table 1: Concept of Stemming & Lemmatization

Form	Stem	Lemma
Studies	Studi	Study
Studying	Study	Study
beautiful	beauti	beautiful
beautifully	beauti	beautifully

Special Character Removal:

A series of characters that creates a pattern known as a hunt is a regex or regular expression. Regex can be used to determine whether a string includes a given hunt pattern. Regular Expressions can be utilized with Python's built-in "re" package, which is available for use. The "re" module provides a number of functions that let us look for matches in a string. The following functions are included in the set:

- **Findall:** Produces a list of matching words as a result.
- **Search:** returns a match object, if there is a match somewhere in the string,
- **Split:** Produces a list in which each match separates the string.
- **Sub:** Uses a string to replace one or more matches.

Thus, after eliminating all special characters, we reached the conclusive point for additional processing. Figure 3 below is a screenshot taken from the pre-processing path of our study, which depicts special character removal.

```
tweet = re.sub(r"https://t.co/[a-zA-Z0-9]*s", " ", tweet)
tweet = re.sub(r"shttps://t.co/[a-zA-Z0-9]*s", " ", tweet)
tweet = re.sub(r"shttps://t.co/[a-zA-Z0-9]*$", " ", tweet)
tweet = tweet.lower()
tweet = re.sub(r"that's", "that is", tweet)
tweet = re.sub(r"there's", "there is", tweet)
tweet = re.sub(r"what's", "what is", tweet)
tweet = re.sub(r"where's", "where is", tweet)
tweet = re.sub(r"it's", "it is", tweet)
tweet = re.sub(r"who's", "who is", tweet)
tweet = re.sub(r"i'm", "i am", tweet)
tweet = re.sub(r"she's", "she is", tweet)
tweet = re.sub(r"he's", "he is", tweet)
tweet = re.sub(r"they're", "they are", tweet)
tweet = re.sub(r"who're", "who are", tweet)
tweet = re.sub(r"ain't", "am not", tweet)
tweet = re.sub(r"wouldn't", "would not", tweet)
tweet = re.sub(r"shouldn't", "should not", tweet)
tweet = re.sub(r"can't", "can not", tweet)
tweet = re.sub(r"couldn't", "could not", tweet)
```

Figure 3: Screenshot of Data Cleaning

Data Labeling:

Data labeling refers to the process of adding some kind of markers, known as tags or labels to textual data. This labeled data is further fed to train the model for analysis of the labeled data. There are varied labeling approaches. Depending upon the problem statement, the time frame of the task, and the number of people who are companies with the work. Data labeling is served through distinctive approaches.

- In-house data labeling
- Crowdsourcing
- Outsourcing

Our data was labeled through the outsourcing route, where the task of data reflection is outsourced to an association or an existent. One of the advantages of outsourcing to individualities is that they can be assessed on the particular content before the work has been handed over. This approach of assembling annotation datasets is integral for systems that don't have big backing, yet demand a meaningful quality of data annotation. We took assistance from IPS (Institute of Policy Studies), Islamabad, which is an independent, not-for-profit, civil society association, devoted to promoting policy-acquainted exploration. IPS provides a forum for informed discussion and dialogue on public and transnational issues. The benefactions gauging over about forty times and the overall impact signifies the importance of realistic exploration of policy issues. The Institute highlights the part of think- tanks in ultramodern popular policy. It especially aims to address policy-acquainted issues concerning Pakistan, and the world community through applied exploration, dialogue, and publications [20]. In order to perform a comparative analysis between human and model results, we handed the data to IPS for manual labeling. Meetings with the IPS team were carried out, to make sure that the labeling had to be done through individuals, who had prior knowledge of the scenario that is the Women's March and all about the rallies, etc. Each tweet was labeled by three individuals, the three results were compared, and the final label assigned was the one that the majority selected. The reason for this procedure was to label the tweets as unbiased as possible

(Removed “Significant Libraries”)

Models Applied:

We used three models for sentiment analysis of the five-year data, collected on “Aurat March”, from the Twitter platform. The models were VADER, Pysentimiento and RoBERTa-Large. Details of the working of these three models are mentioned as follows:

VADER:

Valence Aware Dictionary and sEntiment Reasoner or VADER for short, is a lexicon and plain rule-predicated model for sentiment analysis. It can efficiently handle vocabulary, abbreviations, capitalizations, repeated punctuations, feelings (😞, 😊, 🤔), etc. generally embraced on social media platforms to express one's sentiment, which makes it a skilled fit for social media sentiment text analysis. VADER has the advantage of assessing the sentiment of any given textbook without the need for former training as we might have to go for Machine literacy models. To work out whether these words are positive or pessimistic, the team of VADER developers employed Amazon's Mechanical Turk to pick up the vast majority of their estimations.

When VADER examines a corpus, it verifies whether any of the words in the content are available in the lexicon. For the case, the ruling “The food is startling, and the climate is stunning”, has two words in the wordbook (startling and stunning), with the grading of 1.9 and 1.8, individually (according to the algorithm computation). VADER produces four sentiment magnitudes from word grading, which you can see below. The original three, positive, neutral, and negative, address the extent of the content that falls into those groups. It can be observed that the model judgment was graded as 45% positive, 55% neutral, and 0% negative. The last

dimension, the emulsion score, is the total quantum of the wordbook grades (1.9 and 1.8 for this situation), which have been regularized to run between -1 and 1 . For this situation, our model judgment has a grade of 0.69 , which is firmly on the positive side. Their sum should be equal to 1 or close to it with pier operation.

Emulsion corresponds to the sum of the valence score of each word in the wordbook and determines the degree of the sentiment rather than the real value as opposed to the prior ones. Its value is between -1 (extreme negative sentiment) and 1 (extreme positive sentiment). Applying the compound score can be enough to judge the beginning sentiment of a text. Figure 4 [21] shows the procedure for calculating the VADER score.

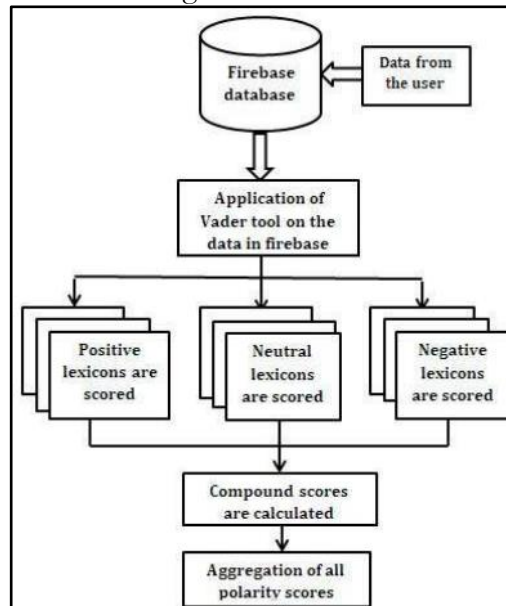


Figure 4: Procedure for Calculating VADER Score

In social media texts like as Twitter and Facebook, the massive usage of emoticons and terminologies with sentiment values also makes text assessment hard. For example, a "😊" indicates a smiley and generally relates to happy or positive sentiment while on the other hand, "😞" denotes sadness.

Also, acronyms similar to "LOL," "OMG" and constantly used cants are also effective pointers of sentiment in a judgment. In dealing with social media manuals, reviews, film reviews, and product reviews, VADER has been discovered to be relatively effective. This is because VADER not only informs us about the score of Positivity and Negativity but also informs us how positive or negative a feeling is. VADER developers have used Amazon's Mechanical Turk to get the utmost of their scores. In other words, it doesn't need any training data but is fabricated from a generalizable, valence-based, mortal-cured gold standard sentiment lexicon [22].

Transformer Based Large Language Models:

There are two primary approaches to sentiment analysis, machine learning and deep learning. From the perspective of machine learning, a corpus must be used to train a classification model for the text sentiment analysis methodology. Nevertheless, pre-trained, self-supervised models are currently taking the world of NLP by storm. A transformer is a new type of neural network architecture. What makes this transformer architecture stand out from the rest (CNN & RNN) is how it caters to machine learning. RNN used to be a good option for assessing ongoing streams of things that encompass strings of textbooks. However, the effects were not that smooth when the subject of interest was to model the relationship between words in case of either a lengthy sentence or a para. This is why the need for a natural language

processing model that would manage the said limitations rose. It was likewise that the researchers at Google discovered that they could achieve way better results by bringing in transformer configuration. Transformer has the potential to find patterns within the existing data. These models possess hundreds of millions of parameters [23]. These models employ a novel architecture that creates representations of its input and output almost entirely through self-attention. Sequence-aligned recurrent neural networks are used to solve sequence-to-sequence challenges (RNN).

RNN is frequently utilized in NLP because of its capacity for structuring effective and precise language models. RNN in conjunction with transformers offers some of the quickest solutions for enforcing NLP outcomes for text and speech processing. Transformer architecture Large Language Models outperform other RNN-exercised models in terms of further state-of-the-art results and efficiency. Figure 5 [24] demonstrates the internal architecture of the transformer

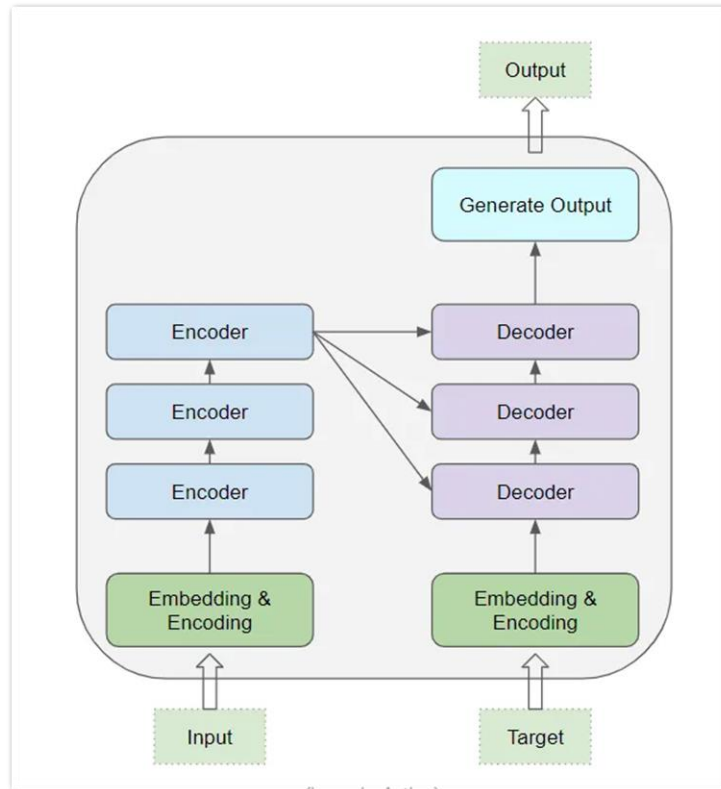


Figure 5: Internal Architecture of Transformers

Step-step working of internal architecture can be learned in simple ways as follows:

- The word embeddings of the input sequence are passed to the first encoder.
- These are then transformed and propagated to the next encoder.
- The output from the last encoder in the encoder stack is passed to all the decoders in the decoder stack.

Pysentimiento:

A transformer-based library, Pysentimiento is a Python toolkit for sentiment analysis and text classification. While carrying out sentiment analysis, we are required to fabricate a model, take care of model type, seek the smart hyper-parameter tuning, fit the data into the model, and train and test the model. Luckily, Pysentimiento comes to save us from all these hard-working operations. This toolkit works by installing the library, and also importing its “sentiment analyzer” since Pysentimiento gives results for two languages, that are English and Spanish, therefore it is needed to choose the language before putting the text, for analysis. The results are shown as probability; positive, neutral, or negative. This toolkit is relatively sensitive

towards the text selection, for illustration if say a text having the word “love”, results in a 99.4% positivity, also, on replacing the word with “like”, would respond to a lower probability of positivity, i.e. 98.7%. The installation is a one-line code that is:

```
!pip install pysentimiento
```

After installation, a “sentiment analyzer” is imported, through which, each sentence is analyzed to evaluate sentiments. For instance, if say a sentence, “I love burgers, but I hate the pickle in it” is passed, the simple way Pysentimiento would be:

Input:

```
sentence = 'i love burger but i hate pickle in it'
sentiment_analyzer_en.predict(sentence)
```

Output:

```
SentimentOutput(output=Pos, probas={Pos: 0.483, Neu: 0.367, Neg: 0.150})
```

In short, a model for sentiment analysis involves, the selection of the best hyper-parameter, fitting the data into the model, training and hence testing the model. Fortunately, Pysentimiento comes to save us from all these hard-working processes.

Roberta-LARGE:

The self-supervised method introduced by Google in 2018 is known as a robustly optimized method for pretraining Natural Language Processing (NLP) systems that enhances Bidirectional Encoder Representations from Transformers, or BERT. BERT is a ground-breaking method that achieves cutting-edge outcomes on a variety of NLP tasks by relying on unannotated text extracted from the internet rather than a language corpus that has been specifically labeled for a given job. Since then, the style has gained popularity as a framework for both the ultimate goal and an NLP exploratory baseline. Google's open release, enables us to manage a replication study of BERT, exposing opportunities to enhance its performance. BERT also demonstrates the collaborative nature of AI research. On the widely used NLP benchmark, General Language Understanding Evaluation, our optimized technique, Roberta, yielded a cutting-edge result (GLUE).

Working of RoBERTa-Large:

Roberta expands on BERT's language masking technique, which teaches the computer to predict purposefully hidden passages of text inside unannotated language instances. By delaying BERT's next-sentence pretraining idea and training with much larger mini-batches and literacy rates, Roberta, which was implemented in PyTorch, modifies key hyperparameters in BERT. As a result, Roberta performs better on subsequent tasks than BERT in terms of the masked language modeling aim.

Removing the Next Sentence Prediction (NSP) Objective:

Through an auxiliary Next Sentence Prediction (NSP) loss, the model is trained to predict whether the attached document portions come from the same or different documents in the next sentence prediction task. Postponing the NSP loss corresponds to or barely improves downstream task interpretation.

Training with Bigger Batch Sizes & Longer Sequences:

BERT is initially trained with a batch size of 256 sequences over 1M steps. The model was trained using RoBERTa-Large with 125 steps of 2K sequences and 31K steps of batch size 8K sequences. Dealing with huge batches improves end-task accuracy as well as confusion on masked language modeling aim. Additionally, distributed parallel training makes it simpler to parallelize large batches.

Dynamically Changing the Masking Pattern:

The masking in the BERT architecture is done just once, during the data preparation, producing a single static mask. Training data is replicated and masked 10 times, each time using a different mask strategy over 40 epochs, resulting in 4 epochs with the same mask. This avoids using the same static mask. This system is contrasted with dynamic masking, which generates a

new mask each time data is fed into the model. In addition to a notable performance enhancement on the GLUE benchmark, the model now provides state-of-the-art performance on the MNLI, QNLI, RTE, STS-B, and RACE tasks.

Roberta equals the performance of the previous leader, XLNet-Large, by scoring 88.5 to take the top spot on the GLUE benchmark. These findings shed light on previously unrecognized design considerations in BERT training and assist in sorting out the relative merits of data size, training duration, and pretraining goals [25].

(Details of data on which Roberta is trained were removed)

Results:

The models discussed above were applied to our dataset, and results were obtained, that showed quite interesting insights about the trend taking place in Pakistan, regarding “Aurat March”, on the Twitter platform.

VADER Results:

VADER sentimental analysis relies on a dictionary that maps lexical features to emotion intensities known as sentiment scores. The sentiment score of a text was therefore obtained by summing up the intensity of each word in the text. This trained NLTK lexicon, when applied to each year’s dataset, gave the following results shown as bar graphs in Figure 6, it demonstrates positive, negative, and neutral tweets, regarding “Aurat March”:

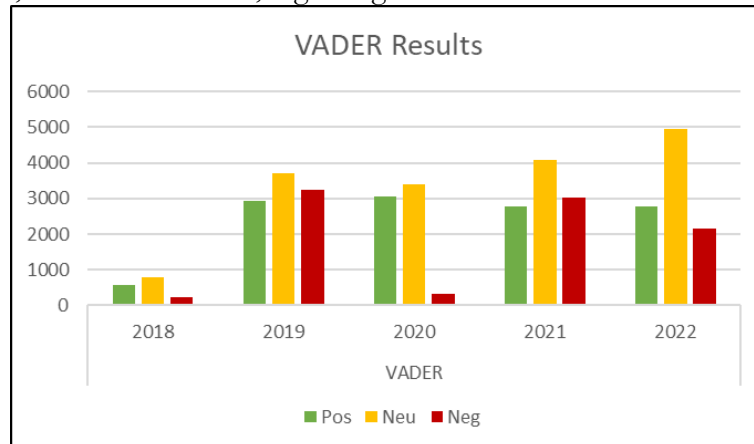


Figure 6: VADER Results 2018-2022

We may observe that VADER results show neutral sentiments to be the highest. However, the positive sentiments are seen to exceed the negative ones in 2022, which is not seen in the rest of the previous years. Overall, VADER results show an increase in positive sentiment, lately regarding “Aurat March”, on the Twitter platform.

PySentimiento Results:

The results obtained from applying PySentimiento are shown in Figure 7:

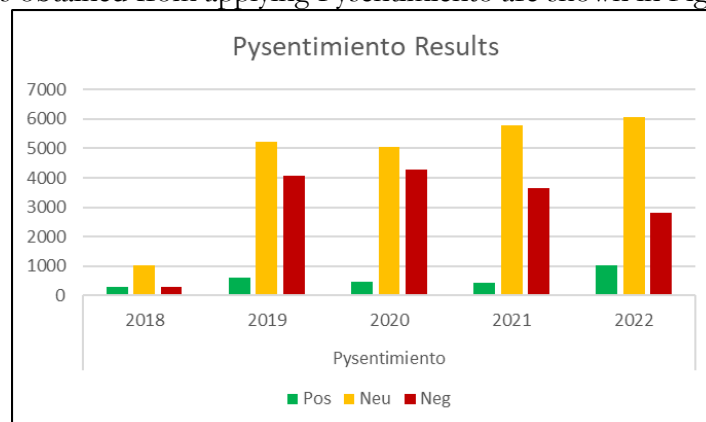


Figure 7: PySentimiento Results 2018-2022

Pysentimiento suggests that throughout, the neutral tweets have been more than both, positive as well as negative tweets. On the other hand, this model shows, positive tweets to be more in 2018, as compared to 2019, 2020, and 2021, where the positive tweets keep decreasing, and finally in 2022, according to Pysentimiento, the positive tweets seem to increase.

Roberta-LARGE Results:

Applying RoBERTa-Large, gave the following sentiment analysis, for the five-year Twitter data regarding “Aurat March”.

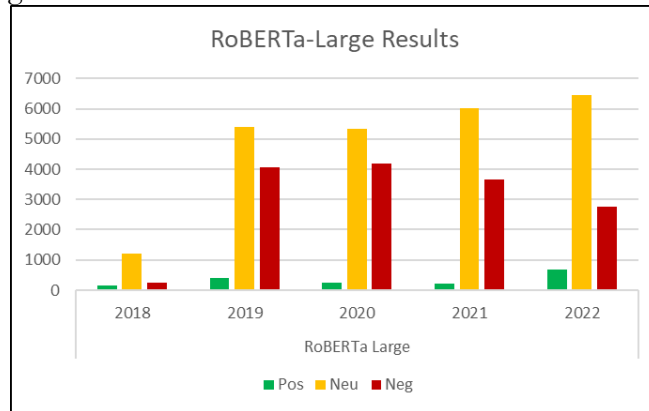


Figure 8: Roberta Results 2018-2022

Manually Labelled Data Results:

A total of 10,933 tweets were manually labeled with the assistance of individuals at IPS (Institute of Policy Studies). Tweets from the five years' data, were shuffled and thus results obtained were not for a specific year, but for the analysis of all possible sentences manually, that could be tweeted in different passages of time. Each tweet was labeled by three persons, and the final label was to be the one selected by the majority, thus avoiding any bias. The sentiments obtained by labeling the dataset manually can be observed in Figure 9.

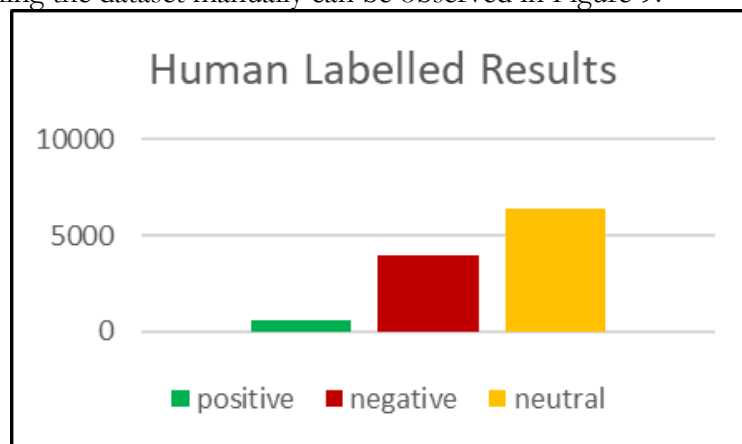


Figure 9: Sentiment Analysis on Manually Labelled Data

Table 2: Sentiments of Manually Labelled Data

Sentiments	Positive	Neutral	Negative
No. of Tweets	578	6363	3992
Percentage	5.29%	58.20%	36.51%

Analysis of the manually labeled data clearly showed that the neutral tweets were prevalent (58.20%), followed by the negative tweets (36.51%), while positive tweets were only 5.29%. This analysis gave us a pretty clear idea that Twitter users tweeting on Aurat March, were mostly neutral about the rallies, also negative comments were tweeted more than positive ones.

Comparative Analysis: The next step involved the application of the three models, namely VADER, Pysentimiento, and RoBERTa-Large on the manually labeled data, by dropping the

column that had been labeled. The reason was to observe and compare the results with the human-labeled results. For a better comparative analysis, results achieved by different models, along with the human-labeled dataset, are shown in Figure 10.

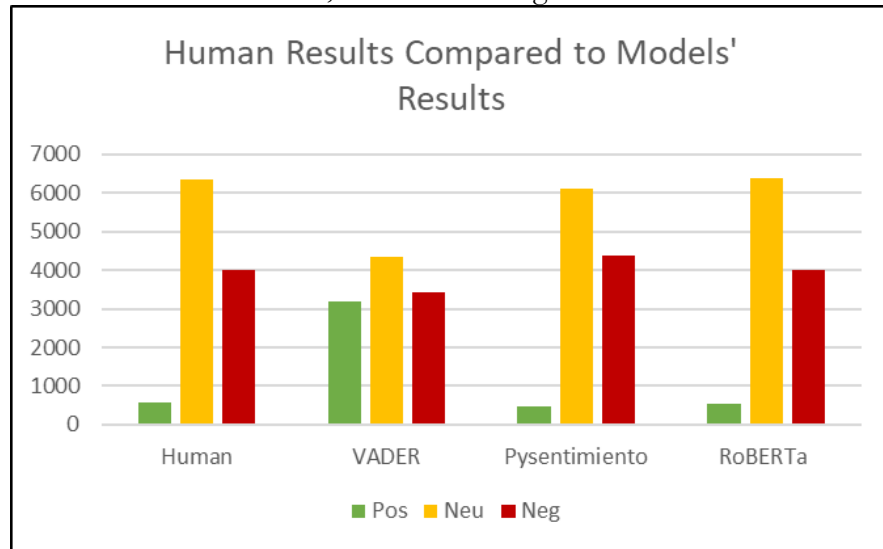


Figure 10: Comparison of Human & Model Analysis

Observing the bar graphs, one can clearly see that VADER analysis, in the case of our data, is quite different than what Pysentimiento and RoBERTA-LARGE have made. Table 3 gives the numeric comparison between the four models, through which we analyzed the Twitter sentiments, so as to see even which model gave the closest resemblance with the human-labeled results.

Table 3: Models versus Humans

Mode of Analysis	Positive	Neutral	Negative
Manual	578	6363	3992
	5.29%	58.20%	36.51%
VADER	3173	4344	3416
	29.02%	40%	31%
Pysentimiento	360	6095	4378
	4.21%	56%	40%
Roberta	551	6384	3998
	5.04%	58.39%	36.56%

- The best analysis results were achieved through RoBERTa-LARGE, which is closest to the Human Labelled Data.
- On the other hand, VADER results were clearly far from the manually labeled results, meaning that they have the least similarity to the results, which we consider to be correct.
- Pysentimiento, has also shown results, pretty close to Human Labelled ones, however, RoBERTa-LARGE stands top of the line among the three Models.

Yearly Trends Through Graphical Visualization:

Results achieved from sentiment analysis using different models were compiled as shown in Table 3, to observe the trend of tweets, from the 2018 to 2022 dataset, through line graphs. The visualization for each model results is shown below:

Vader Yearly Trends:

VADER is a model used for text sentiment analysis that is sensitive to both polarity (positive/negative) and intensity (strength) of sentiment. It can be applied directly to unlabelled textual data. VADER sentimental analysis relies on a dictionary that maps lexical features to

emotion intensities known as sentiment scores. The graphical visualization in Figure 11, gives an idea of how sentiments fluctuated throughout the five years, regarding “Aurat March”.

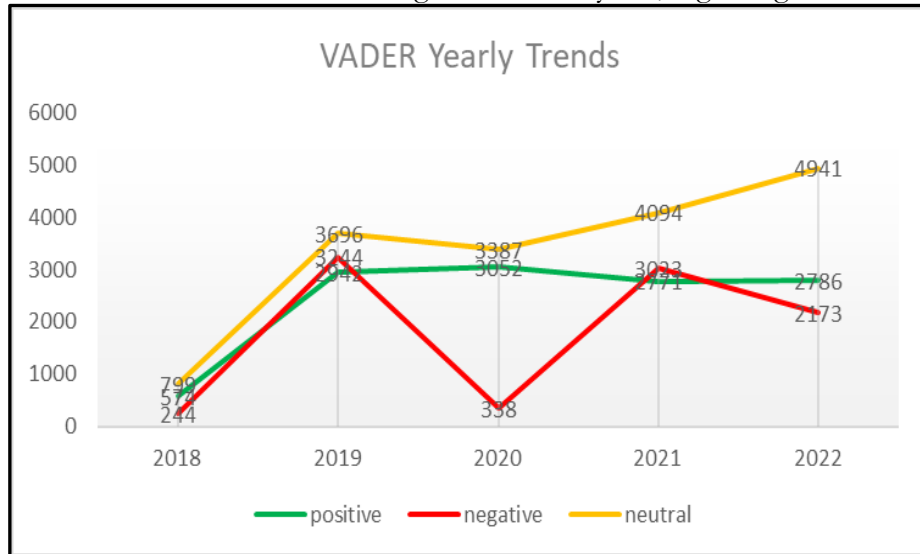


Figure 11: VADER Yearly Trends

Positive sentiments surpassed negative ones in 2018, but over the subsequent four years, the graph depicted a stable trend with slight fluctuations. Analysis of the negative sentiment curve reveals its lowest point in 2018, followed by a notable increase in 2019, before sharply declining in 2020. Interestingly, neutral sentiments remained predominant throughout the five-year period, consistently outnumbering both positive and negative ones. This observation underscores the resilience of neutral opinions in the face of fluctuating trends and highlights the complexity of sentiment dynamics over time.

Pysentimiento Yearly Trends:

Pysentimiento, which had been created by developers, especially for text classification and sentiment analysis, gave the trends shown in Figure 12 for the five years of tweets.

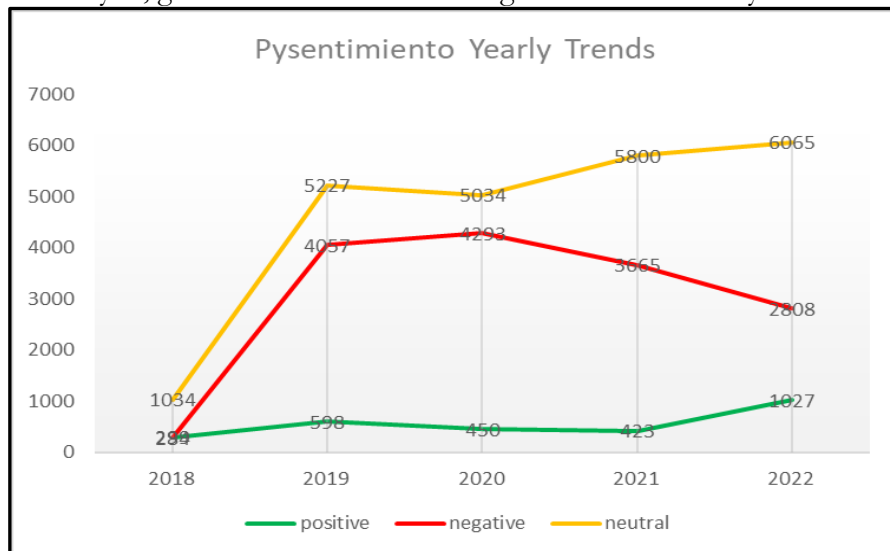


Figure 12: PYSENTIMIENTO Yearly Trends

In 2018, positive sentiments were closely aligned with negative ones, but as time progressed, there was a significant increase observed in both negative and neutral sentiments. When comparing negative and neutral sentiments, it becomes apparent that neutral sentiments consistently outweighed negative ones. Notably, during 2021 and 2022, there was a decrease in negative sentiments, while positive and neutral sentiments experienced an upward trend from

their previous levels. This fluctuation in sentiment distribution highlights the dynamic nature of public opinion over the years, indicating shifts in societal attitudes and perceptions

Roberta-Large Yearly Trends:

Roberta, A Robustly Optimized BERT Pretraining Approach is based on Google's BERT model released in 2018. It modifies key hyperparameters, training with much larger mini-batches and learning rates. The graphical trends for the five years are given in Figure 13.

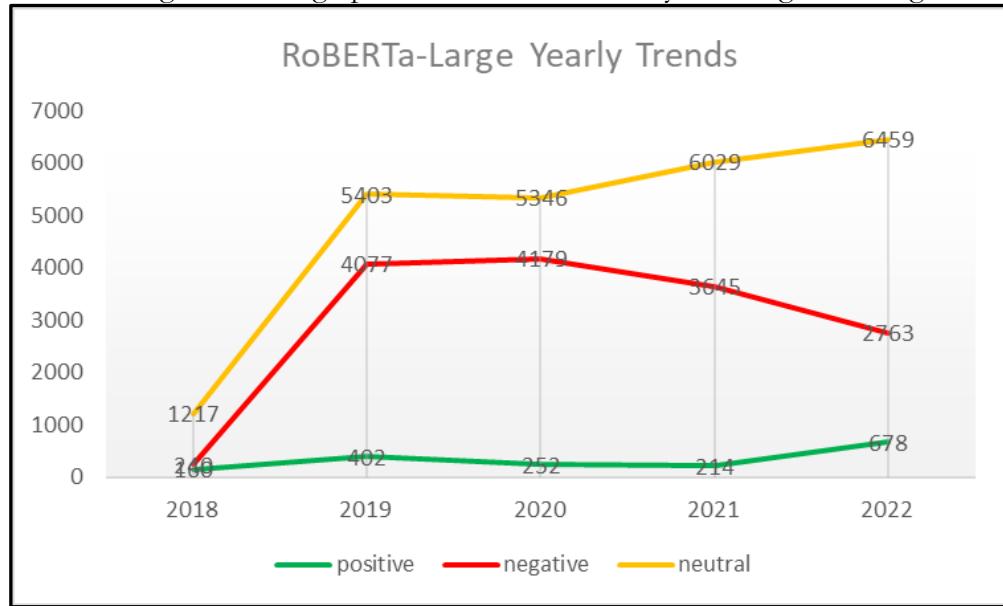


Figure 13: Roberta-Large Yearly Trends

In the initial year, 2018, positive sentiments dominated, but they witnessed a significant decline in the subsequent two years. However, there was a notable resurgence in positive sentiment observed in 2021, which continued into 2022. Conversely, negative sentiments experienced a drastic increase between 2018 and 2019, followed by relative stability before gradually decreasing over time. Interestingly, according to an analysis by RoBERTa-LARGE, neutral sentiments consistently held the highest count on Twitter throughout the five-year period. Nonetheless, akin to negative sentiments, they also experienced an uptick between 2018 and 2019, with this trend persisting in the following years. These trends shed light on the nuanced fluctuations in sentiment dynamics over time, reflecting shifts in public opinion and discourse on the platform.

Conclusion:

Gender equality transcends mere financial objectives, encompassing access to health, education, earning potential, and political representation for women, who constitute half of the global population. Women would gain further negotiating influence in the marketable and governmental sectors as a result of similar reforms, as well as the capability to work as alternate-income earners, performing in increased family support for a woman's career. The emergence of social media platforms like Twitter has generated vast amounts of user-generated data daily, with tweets becoming valuable sources of information for various decision-making processes. Twitter's short texts, known as tweets, have gained a lot of attention as a precious source of information for numerous decision-making processes. The elementary thing of Twitter SA is to assess whether the tweet has a favorable or negative sentiment.

The following are the primary challenges of Twitter sentiment analysis:

- Tweets are generally written in colloquial English
- Brief messages give little suggestions as to the sentiment
- Acronyms and abbreviations are generally employed on Twitter.

Machine learning and NLP are playing a vital role in achieving insights through Twitter data. The sentiment analysis of “Aurat March” tweets, gave us satisfactory results, and we were able to validate our research, by comparing the models that were applied, with human-labelled results. Thus, we were able to observe the general mindset of Twitter users, plus the trends, giving us a pretty reasonable idea, to see either the fame or defame of feminism during the past few years. Sentiment analysis on textual data is a field under widespread exploration and examination worldwide. However, when tweets are processed by machines, they often fail to capture the nuances present in human communication. During our research work, we encountered several constraints that hindered the accuracy of sentiment analysis.

In some cases, tweets had something that was said sarcastically, and the machine understands the literal meaning, and not the spirit behind the statement. Some people had tweeted their point of view, breaking the talk in two tweets. In such cases, the machine responded to the individual tweet separately, not being able to realize the former to be a continuation of the previous tweet.

Future Work:

We are fortunate to witness state-of-the-art advancements taking place in Artificial Intelligence, at such a rapid pace. Architectures like Transformers and BERT are paving the way for even more advanced breakthroughs to happen in the coming years, especially in the field of sentiment analysis. In fostering a collaborative environment, we urge researchers to not only implement diverse models but also to openly share their work, thereby facilitating the dissemination of knowledge. Furthermore, another future direction is to implement these and/or different techniques on the first ever collected dataset on “Aurat March” for more insights and trends. Furthermore, it is vital to study the internal architectures of LLMs and/or other advancements in AI, to have an in-depth knowledge of the technology’s strengths as well as limitations.

Author’s Contribution:

After going through a significant number of research previously done, Ms. Sana Sohail brought up the idea of investigating the mentioned topic through Data Science techniques. Mr. Furqan Amjad assisted with the valuable insights and trends that were achieved from the study.

Conflict of Interest:

There exists no conflict of interest for publishing this manuscript in IJIST among authors.

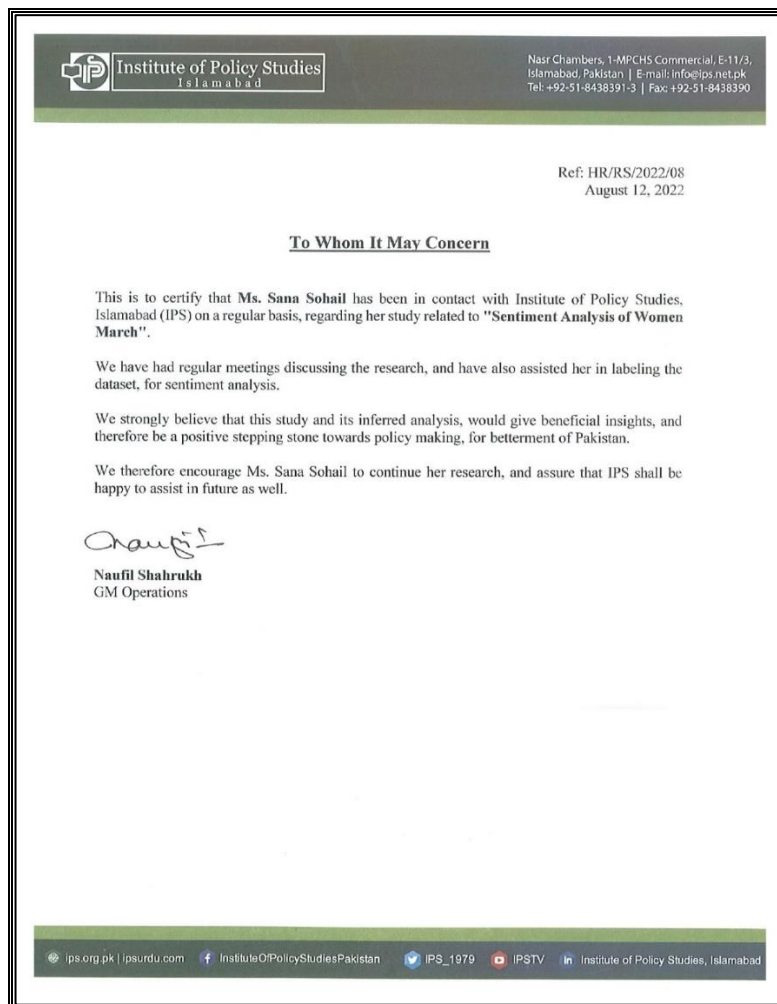
Reference:

- [1] A. Sarwar and M. K. Imran, “Exploring women’s multi-level career prospects in Pakistan: Barriers, interventions, and outcomes,” *Front. Psychol.*, vol. 10, no. JUN, p. 444970, Jun. 2019, doi: 10.3389/FPSYG.2019.01376/BIBTEX.
- [2] A. Grönlund, “More Control, Less Conflict? Job Demand–Control, Gender and Work–Family Conflict,” *Gender, Work Organ.*, vol. 14, no. 5, pp. 476–497, Sep. 2007, doi: 10.1111/J.1468-0432.2007.00361.X.
- [3] “Research Methods For Business Students: Adrian Thornhill / Philip Lewis / Mark N. K. Saunders: 9781292208787: Amazon.com: Books.” Accessed: Mar. 28, 2024. [Online]. Available: <https://www.amazon.com/Research-Methods-for-Business-Students/dp/1292208783>
- [4] S. R. Madsen and R. T. Scribner, “A perspective on gender in management The need for strategic cross-cultural scholarship on women in management and leadership,” *Cross Cult. Strateg. Manag.*, vol. 24, no. 2, pp. 231–250, 2017, doi: 10.1108/CCSM-05-2016-0101/FULL/XML.
- [5] “Young Women in Pakistan – Status Report 2020 | UN Women – Asia-Pacific.” Accessed: Mar. 28, 2024. [Online]. Available: <https://asiapacific.unwomen.org/en/digital-library/publications/2020/11/young-women-in-pakistan-status-report-2020>

- [6] J. Acker, "Inequality Regimes," <http://dx.doi.org/10.1177/0891243206289499>, vol. 20, no. 4, pp. 441–464, Aug. 2006, doi: 10.1177/0891243206289499.
- [7] "Types of Feminism: The Four Waves | Human Rights Careers." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.humanrightscareers.com/issues/types-of-feminism-the-four-waves/>
- [8] R. F. Baumeister, J. D. Campbell, J. I. Krueger, and K. D. Vohs, "Does High Self-Esteem Cause Better Performance, Interpersonal Success, Happiness, or Healthier Lifestyles?," *Psychol. Sci. Public Interes.*, vol. 4, no. 1, pp. 1–44, 2003, doi: 10.1111/1529-1006.01431.
- [9] A. Vera and D. Hucke, "Managerial orientation and career success of physicians in hospitals," *J. Health Organ. Manag.*, vol. 23, no. 1, pp. 70–84, 2009, doi: 10.1108/14777260910942560.
- [10] "Aurat Azadi March - Wikipedia." Accessed: Mar. 28, 2024. [Online]. Available: https://en.wikipedia.org/wiki/Aurat_Azadi_March
- [11] "Almost 90% of Men/Women Globally Are Biased Against Women | United Nations Development Programme." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.undp.org/press-releases/almost-90-men/women-globally-are-biased-against-women>
- [12] A. Khurshid and A. Saba, "Contested womanhood: women's education and (re)production of gendered norms in rural Pakistani Muslim communities," *Discourse Stud. Cult. Polit. Educ.*, vol. 39, no. 4, pp. 550–563, Jul. 2018, doi: 10.1080/01596306.2017.1282425.
- [13] "The future of female doctors", [Online]. Available: <https://www.bmj.com/content/338/bmj.b2223>
- [14] "The Women's March and Its Impact, One Year Later | The Brink | Boston University." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.bu.edu/articles/2018/the-womens-march-and-its-impact/>
- [15] K. Taylor, T. Lambert, and M. Goldacre, "Future career plans of a cohort of senior doctors working in the National Health Service," *J. R. Soc. Med.*, vol. 101, no. 4, p. 182, Apr. 2008, doi: 10.1258/JRSM.2007.070276.
- [16] "What Is Sentiment Analysis (Opinion Mining)? | Definition from TechTarget." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.techtarget.com/searchbusinessanalytics/definition/opinion-mining-sentiment-mining>
- [17] "Using Transformer-Based Language Models for Sentiment Analysis | by Bogdan Kostić | Towards Data Science." Accessed: Mar. 28, 2024. [Online]. Available: <https://towardsdatascience.com/using-transformer-based-language-models-for-sentiment-analysis-dc3c10261eec>
- [18] "DATA PREPROCESSING IN SENTIMENT ANALYSIS USING TWITTER DATA." Accessed: Mar. 28, 2024. [Online]. Available: https://www.researchgate.net/publication/334670363_DATA_PREPROCESSING_IN_SENTIMENT_ANALYSIS_USING_TWITTER_DATA
- [19] "Manual or automated data labeling, how to decide?" Accessed: Mar. 28, 2024. [Online]. Available: <https://www.labellerr.com/blog/manual-or-automated-data-labeling-how-to-decide/>
- [20] "Institute of Policy Studies - Institute of Policy Studies." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.ips.org.pk/>
- [21] S. M. Rao, N. Monica, P. Nikhila, T. Tejasri, and B. Maram, "Positivity Calculation using Vader Sentiment Analyser," *Int. J. Acad. Eng. Res.*, vol. 4, pp. 13–17, 2020, Accessed: Mar. 28, 2024. [Online]. Available: www.ijcais.org/ijaer

- [22] C. J. Hutto and E. Gilbert, "VADER: A Parsimonious Rule-Based Model for Sentiment Analysis of Social Media Text," Proc. Int. AAAI Conf. Web Soc. Media, vol. 8, no. 1, pp. 216–225, May 2014, doi: 10.1609/ICWSM.V8I1.14550.
- [23] X. Gong, W. Ying, S. Zhong, and S. Gong, "Text Sentiment Analysis Based on Transformer and Augmentation," Front. Psychol., vol. 13, p. 906061, May 2022, doi: 10.3389/FPSYG.2022.906061/BIBTEX.
- [24] "Transformers Explained Visually (Part 1): Overview of Functionality | by Ketan Doshi | Towards Data Science." Accessed: Mar. 28, 2024. [Online]. Available: <https://towardsdatascience.com/transformers-explained-visually-part-1-overview-of-functionality-95a6dd460452>
- [25] "A Gentle Introduction to RoBERTa - Analytics Vidhya." Accessed: Mar. 28, 2024. [Online]. Available: <https://www.analyticsvidhya.com/blog/2022/10/a-gentle-introduction-to-roberta/>

APPENDIX:



Copyright © by authors and 50Sea. This work is licensed under Creative Commons Attribution 4.0 International License.

