

Sarcasm in the Digital Age: An Opinion Mining Approach to Social Network Conversations

**Sakshi Sobti¹, Dr. Varsha Agarwal², Dr. Angad Tiwary³, Preeti Naval⁴,
Bhuvana Jayabalan⁵, Jagmeet Sohal⁶**

¹*Centre of Research Impact and Outcome, Chitkara University, Rajpura- 140417, Punjab, India sakshi.sobti.orp@chitkara.edu.in <https://orcid.org/0009-0003-9901-0056>*

²*Associate Professor, Department of ISME, ATLAS SkillTech University, Mumbai, Maharashtra, India, Email Id- varsha.agarwal@atlasuiversity.edu.in, Orcid Id- 0000-0001-8406-9213*

³*Professor, Department of Management, ARKA JAIN University, Jamshedpur, Jharkhand, India, Email Id- dr.angad@arkajainuniversity.ac.in, Orcid Id- 0009-0000-3277-456X*

⁴*Assistant Professor, Maharishi School of Engineering & Technology, Maharishi University of Information Technology, Uttar Pradesh, India, Email Id- er.preetinaval09@gmail.com, Orcid Id- 0000-0003-2988-7082*

⁵*Associate Professor, Department of Computer Science and Information Technology, JAIN (Deemed-to-be University), Bangalore, karnataka, India, Email Id- j.bhuvana@jainuniversity.ac.in, Orcid Id- 0000-0002-8372-6311*

⁶*Chitkara Centre for Research and Development, Chitkara University, Himachal Pradesh- 174103 India jagmeet.sohal.orp@chitkara.edu.in <https://orcid.org/0009-0000-6950-9022>*

Social networking services are becoming increasingly popular as everyone writes blogs, micro posts and other things expressing their unique perspectives. One of the areas of Artificial intelligence (AI) that is developing the quickest is sentiment analysis, which divides viewpoints into positive, negative and neutral attitudes. Sarcasm is one of these sentiment analysis components. On social media platforms, sarcasm is becoming prevalent. It is typical to communicate ambiguous emotions with phrases that betray contempt, making it hard to determine the true meaning of statements. In this paper, we proposed the use of Twitter data to forecast a remark in the categories of sarcastic or no sarcastic using the Improved Gray Wolf Optimized Enhanced Random Forest (IGWO-ERF). We begin with data from social media-based research on sarcasm detection. Tokenization and stop word removal were used in the preprocessing of the data. Term inverse document frequency (TF-IDF) was used to examine various feature extraction settings. To the comparison, the IGWO-ERF technique performed better across a wide variety of performance metrics, such as accuracy (96.10%), precision (94.23%), recall (92.33%) and F1-score (90.12). Our result shows that sarcasm detection is crucial in the ever-changing world of social media communication. Understanding and properly identifying sarcasm is essential for successful communication, sentiment analysis and ethical usage of AI-driven technologies, as online platforms

affect public conversation.

Keywords: Social networking, Sarcasm detection, Improved Gray Wolf Optimized Enhanced Random Forest (IGWO-ERF), Twitter.

1. Introduction

The age of social media has a profound effect on how people communicate across the globe. A person today can exhibit easily and quickly with the help of social media. It's a typical way for social media users to express their feelings about a post saw or an image liked [1]. It rejects surprise that daily sarcasm can be discovered on social media sites like Facebook and Twitter. Using the linguistic strategy known as sarcasm to convey disdain or other unfavourable feelings is common. It is an act of pretense that is meant to be polite, but it has the unintended effect of making people upset [2]. Sarcasm gives the impression of cruel and has a little air of dishonesty about it.

Recent studies have shown that those who tease others have the misconception that their words are not as hurtful as their victim believes them to be. Due to the general belief that political parties and celebrities have great influence, individuals are the targets of disparaging remarks and tags [3]. There is a correlation between the use of sarcasm and mental illnesses such as anxiety and depression. The individuals who were experiencing feelings of depression or anxiety during the epidemic used sarcasm throughout their online chats [4]. When conversing with someone in person, by paying attention to the gestures, tone of voice and feelings of the speaker, one can quickly discern sarcasm. Sarcasm can be difficult to determine in written communication since none of these traits are obvious in sarcastic writing [5]. It is even more difficult to recognize sarcasm in photographs published on social media sites, given that the background information is either included in the images or mentioned elsewhere [6]. The ability to recognize sarcasm is vital for a variety of jobs, including the identification of cyber bullies along with online trolls, the mining as well as analysis of opinions, the identification of fake news and opinion mining. It is essential to recognize sarcasm in online chat boards, social networking applications and other forms of electronic communication [7]. Research on the ability to identify sarcasm is seeing a surge in interest. The ability to spot fake news is impacted by sarcasm as well. Streaming data from social media platforms presents several additional difficulties in addition to those already mentioned. Sarcasm identification is an extremely important part of the company's feedback system, as it enables the business to understand the consumers' genuine intentions about the product [8].

1.1 Research Gap

The last decade has seen an explosion in the number of studies devoted to the study and analysis of social network data. One of the most difficult yet interesting problems in natural language processing is identifying sarcasm. Not all ironic information is always bad or good, and vice versa. In reality, sarcastic material is typically clear and unclear, making it hard to discern. To gauge consumer sentiment around a product, sarcasm detection is crucial. Business decision-making relies heavily on the ability to recognize sarcasm. The large number of sarcastic texts on social networks further highlights the need for in-depth research and

analysis. However, standard sentiment analysis methods, such as rule-based methods, fail miserably when used in sardonic writing. Therefore, there is a pressing need for a well-designed model for sarcasm detection jobs in particular. With the abundance of sarcastic information on social media, it is essential to reliably recognize sarcasm to understand customer mood and make sound business choices. This is where the IGWO-ERF technique comes in. To overcome the limitations of traditional sentiment analysis algorithms, IGWO-ERF was developed to perform very well in sarcasm detection tasks.

1.2 Significance of the study

The study of sarcasm detection on social networks has a number of potential benefits, including the development of an accurate natural language processing models, the improvement of sentiment analysis and the facilitation of more efficient content moderation on social media platforms as well as a better understanding of users' feelings and perspectives. Sarcasm detection based on social networks is the importance and relevance of the research in understanding and developing algorithms to identify sarcastic language in online interactions. These algorithms have the potential to improve communication analysis, sentiment analysis and the overall quality of content moderation on social media platforms.

1.3 Objective of the Study

- More and more people are turning to online social networking sites to share their thoughts and opinions.
- Sentiment analysis, which classifies opinions as optimistic, pessimistic, or agnostic, is one of the most rapidly expanding fields.
- In this paper, we propose the Improved Gray Wolf Optimized Enhanced Random Forest (IGWO-ERF) for detecting the sarcastic based on social networks.
- By calculating the Accuracy, Precision, Recall and F1-Score, we evaluate our proposed method with conventional approaches.

The remaining portion of the study is divided into section 2, which provides relevant studies; Section 3, which describes the methodology; Results and discussion are covered in sections 4 as well as 5 and the study concludes with some suggestions for future research in section 6.

2. Related works

According to the author of, [9] investigated whether or not sarcasm in political tweets affects the output of computer approaches applied to huge datasets using Twitter mining to Judge Brett Kavanaugh's confirmation to the Supreme Court in 2018. The study finds any alterations in public perception of the Court that could have happened in the aftermath of the confirmation hearing. The burgeoning field of Social Opinion Mining, which seeks to determine the subjectiveness of user-generated material across a variety of social media platforms and in text, image, video and audio formats can be analyzed for polarity of sentiment, emotion, affect, sarcasm and irony. Social opinion mining allows for the analysis of unstructured text according to the numerous characteristics of human opinion [10]. To examined that it is actively working to improve the effectiveness of sentiment analysis by creating new methods for identifying

sarcasm in written materials. To that end, the study provided a survey of research on sarcasm and sentiment analysis [11]. In addition, the article instructs medical personnel on how to base their choice on the patient's values and preferences. To opposed the past lone investigations, which have identified 21 issues with online social networks and provided references to relevant research [12]. To provided information to elucidate the methods that are currently in use and trending for sarcasm detection. Sarcasm was a wonderful tool for humor and display of foolishness [13]. Sarcasm can be communicated orally and by certain body language cues, such as eyebrow lifting or eye-rolling. There were many approaches used to identify sarcasm. To primarily examined several low-level features, such as those retrieved by deep learning and crowd-sourced terms [14]. Sentiment analysis of comments about commercial items on Facebook and Twitter is the application's goal. Author [15] determined the political leanings of users of social media sites like Facebook and Twitter. The study identified the unresolved problems in sentiment analysis and research difficulties linked to election outcome prediction. The study made some recommendations for future developments in the field of election prediction utilizing material from social media. To examine synthetic minority oversampling, this can negatively impact classifier performance [16]. The research used two different-sized datasets and five unique versions of synthetic minority oversampling. The proposed a method based on pattern recognition that uses Twitter data to detect sarcasm [17]. There was a great deal of specific sarcasm in tweets, which are categorized as sarcastic or not, based on four sets of qualities that were provided. The recommended feature sets' supplementary cost categories are analyzed. The study [18] investigated the challenge of identifying sarcasm across social media and other digital platforms in real-time conversations. To do so, build an interpretable deep learning model that used multi-head self-attention and Gated recurrent units (GRUs). Recurrent teams employ the multi-head self-attention module to acquire long-range connections between cue words to recognize important sarcastic cue words in the input text and categorize them.

3. Methods

The dataset obtained from the twitter is used in the suggested method and the goal is to determine whether or not each tweet utilized contains sarcasm. The classification of tweets is accomplished by first extracting from them a variety of attributes and using an appropriate technique from the field of ML. During extracting characteristics, various sarcastic phrases can be constructed from the tweets' constituent parts. Figure 1 shows the methodology for sarcasm detection for social media networks.

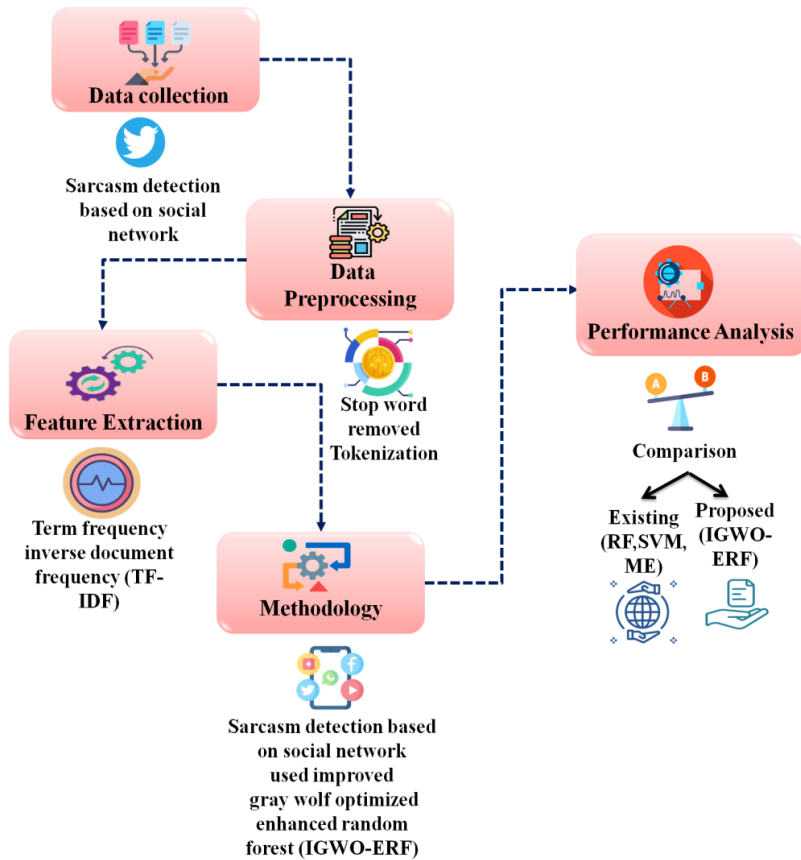


Figure 1: Methodology of Sarcasm Detection

3.1 Datasets

The hashtags #sarcasm and #not are used in 9,104 tweets that might be interpreted as sarcastic. A collection of tweets in English and Hindi from Twitter was used to test the Framework. To solve the issue [19], the most crucial and time-consuming phase is the preparation of the data. Since the data represents the project's input, improving its precision will enhance its accuracy. Tweets with media or link descriptions or URLs that don't go to useful resources are deleted. The tweets include a large number of expressions that are not relevant to the topic at hand, such as "Yeah, right!" in addition to sarcasm. To do a preprocessing step that involves removing the hashtags #sarcasm and #not from the 9104 sarcastic tweets before we grade them. The Twitter Sample API collects tweets that are not intended to be rude and sarcastic.

3.2 Data Preprocessing Using Stop Word Removal and Tokenization

Data preparation is a mining technique to modify raw data into usable information. The input data is collected and hashtags are found during the data preparation. The data input no longer includes these hashtags. Field selection tokenization is included in the next module. The following is how the collected data is processed:

3.2.1 Data Cleaning

The purpose of data cleaning is to remove irrelevant information and make the data far more useful for the research at hand. Since the data includes numerous unusual characters, the primary goal of cleaning is to eliminate these unwanted elements. One popular Python library that helps with data cleaning is called "re," which is related to regular expression. An example of the study's data cleansing is shown below. Unclearness Text: Hey, can you please ask a few of my Twitter followers to subscribe to my YouTube channels? Right now, I don't have any! "#podcast #social media #irony #NFL Clean Text: Hello, will you please ask a few of my Twitter followers to become subscribers to my YouTube channel? The NFL teasing is all in good humor, but I'm excited about my podcast and social media posts!

3.2.2 Stop Word Removal

Words that are used seldom or not at all in the English language are examples of stop words. These terms are eliminated because they are considered useless words and because they take up space in the database; as a result, removing these words is preferable for analytical purposes. Here, some stop word removal techniques are 'i,' 'me,' 'myself,' 'we,' 'our,' 'ours,' 'ourselves,' 'you,' 'yourself,' and others. Stop words are eliminated during preprocessing to increase the flexibility of the processed analysis; the result shown here is what was attained after deleting stop words from the column labeled "tweets" so that better categorization could be achieved.

Unclearness Data: The human mind is so efficient at digesting information that it frequently fails to detect sarcasm.

Clean Data: Sometimes, even with its amazing powers, the human brain has trouble recognizing sarcasm.

3.2.3 Tokenization

Tokenization is a subset of text processing that includes identifying suitable delimiters to cut the text into manageable chunks or tokens. In lexically analyzing the text, it is one of the most important aspects. Tweets go through the tokenization process to be broken down into understandable modules derived from a text.

Unclearness Data: Changing our focus from instant gratification to a lifelong pursuit of love, character and bravery, where pain and fear no longer define us but serve to fortify our resolve and grow our affection for one another.

Clean Data: 'lifelong,' 'love,' 'character,' 'bravery,' 'pain,' 'fear,'

3.3 Feature Extraction using Term Frequency-Inverse Document Frequency (TF-IDF)

In a classification algorithm, feature selection is a critical and productive step. Sarcasm has two distinct meanings, making study on the topic difficult. Therefore, researchers strive to choose appropriate characteristics to improve the accuracy of the sarcasm detection algorithm. Team has isolated sentence characteristics such as punctuation, interjection words and emotions expressed via punctuation and emojis. The TF-IDF matrix displays the frequency with which terms appear in a text or phrase. The occurrence count of a term inside a text is measured using the term frequency (TF (s)) measure, whereas document frequency (DF(s))

quantifies whether a given word appears in a given set of documents. One way to measure how significant a word is by its Inverse Document Frequency (IDF (s)). The (IDF (s)) will be small if the term appears in text and large otherwise. So, we have a mathematical definition for it in Equation (1).

$$\text{IDF}(s) = \frac{M}{\text{DF}(s)} \quad (1)$$

Where, M is the total number of documents, $DF(s)$ is the frequency of those documents and $IDF(s)$ is the inverse frequency. Since M can be very big, the value of $IDF(s)$ or $DF(s)$ can be zero at query time in Equation (1). Therefore, the former is solved by adding 1, whereas the latter is solved using the log function. If this is the case, we can rewrite the $IDF(s)$ Equation (2).

$$\text{IDF}(s) = \log \frac{M}{\text{DF}(s)+1} \quad (2)$$

Equation (3) is the mathematical expression of TF-IDF derived from Equation (3).

$$\text{TF-IDF}(s, c) = \text{TF}(s, c) * \log \frac{M}{\text{DF}(s)+1} \quad (3)$$

3.4 Sarcasm Detection Based On Social Network Used Improved Gray Wolf Optimized Enhanced Random Forest (IGWO-ERF)

Artificial intelligence (AI) has a long way to go before it can imitate human intuition as well as behavior; sarcasm and humor are two of the most fundamental human traits yet to be mastered. Although there are several machine learning algorithms designed for this purpose, there are numerous unique challenges associated with classifying text based on sentiment. Figure 2 shows the sarcasm detection based on social media. The sarcasm detection proposed an Improved Gray Wolf Optimized Enhanced Random Forest (IGWO-ERF).



Figure 2: Sarcasm (Source: https://www.researchgate.net/figure/Tag-Cloud-of-Twitter-features-to-detect-Sarcasm_fig1_319620213)

3.4.1 Enhanced Random Forest (ERF)

An assembly of random trees is produced by using a classification method known as the random forest algorithm. The algorithm for creating classification and regression trees, called CART, is the decision tree approach that represents the norm. A CART tree comprises several nodes, including a central node, branching nodes, terminal nodes and connecting edges. The most crucial processes are selecting the most appropriate variables to serve as nodes and determining the most proper points of division between those vertices. This guarantees that the progeny nodes will be more precise than their parents. The CART technique uses the Gini index as a means of determining the degree of impureness associated with each node. Provided that it has a node s as well as the anticipated probability of each class, an expression for the Gini index at node s is $o(d|s)$ ($s = 1, \dots, S$) below in Equation (4):

$$H(s) = \sum_{d_1+d_2} o(d_1|s)o(d_2|s) = 1 - \sum_{d=1}^d o^2(d \cdot s) \quad (4)$$

Let t function as the branching point of the node s , which divides the node into two distinct parts, each of which has a percentage, o_Q , t is the number of samples that corresponds to (s_Q) and a proportion, o_K , is assigned to s_K , i.e., 1. As a result, the decline in the Gini index of impurity is due to $o_Q + o_K$ specified in this way as shown in Equation (5)

$$\Delta H(t, s) = H(s) - o_Q H(s_Q) - o_K H(s_K) \quad (5)$$

The ideal variation t^* and the ideal dividing point i^* that result in the greatest reduction in the Gini impurity are ascertained as shown in Equation (6).

$$t^*, i^* = \arg \max_{t, i} \Delta H(t, s) \quad (6)$$

The CART algorithm makes repeated calls to the above function to produce a tree. Bagging theory and random subspace theory are combined in an ensemble model known as a random decision forest that is derived from the CART method. To be explicit, the CART method educates a collection of decision trees via the use of each non-leaf node containing a bootstrap sample and a set of variables that were randomly selected. Every one of the forest's trees has the potential to develop its maximum height to an endless level as long as the leaf nodes remain unadulterated. Random Forest algorithm is shown in Algorithm 1.

Algorithm 1: Random forest algorithm

```

Forj = 1 toA, do
DrawabootstrapsampleofsizeMfromthetrainingdata;
whilenodesize != minimumnodesizedo
Randomlyselectasubsetofnpredictorvariablesfromtotalo;
For ← 1 tondo
ifthepredictoroptimizesthesplittingcriterion, then
Splittheinternalnodeintotwochildnodes;
```

```

break;
end
end
end
end

```

3.4.2 Improved Gray Wolf Optimization (IGWO)

Grey wolf packs served as an influence for IGWO due to its social organization and hunting prowess. The algorithm reaches its optimal state by simulating the behavior of a collection of improved grey wolves as follow prey, surround prey, hunt prey and attack prey. The improved grey wolf's hunting process consists of three stages: stratification based on social hierarchy, encircling the prey and attacking.

Inspiration: Improved Grey wolves have a stratified social order. There is a male and female leader. Decisions on where to hunt, where to sleep, when to get up are delegated to the alpha. The pack must go where the leader decides to go. The grey wolf pack structure includes a secondary group known as the beta pack. The betas are the wolves in the pack that are second in command and help the alpha with decision-making and pack management. In the event of the death or incapacitation of an alpha wolf, the pack will elevate the status of the beta wolf to that of the pack leader. The improved gray wolf ranks dead last among mammals. The omega serves as a hapless spectator in this scenario. When several alphas are in the pack, the alpha wolf must take a back seat. The wolves that are fed right now are the very last of their species. The people in the fourth category are called subordinates (sometimes called deltas), depending on the source. Delta wolves are submissive to alphas and betas, but exercise authority over omegas. Among the Deltas, there are scouts, sentinels, elders, hunters and caretakers, among many other specialized duties. **Mathematical Modeling:** The pack must first surround the target to chase it. The following Equation (7-8) is used to represent encircling behavior mathematically.

$$\vec{W}(s+1) = \vec{W}_o(s) + \vec{B}, \vec{C} \quad (7)$$

$$\vec{C} = |\vec{D}, \vec{W}_o(s) - \vec{W}(s)| \quad (8)$$

The $\vec{B}, \vec{D}$ calculations for vectors follow the Equations. (9 and 10)

$$\vec{B} = 2\vec{B}_1\vec{q}_1 - \vec{b} \quad (9)$$

$$\vec{D} = 2\vec{q}_2 \quad (10)$$

The mechanism of $\vec{b}$ is reduced linearly from iteratively reducing a random vector from (2, 0) to (0, 1), where q_1, q_2 are random numbers. During a hunt, the alpha acts as the leader. The beta and the delta sometimes go on hunting expeditions together. To mathematically model the hunting behaviour of grey wolves, we need to know where their prey is located and this information is believed to be included in the alpha (best candidate solution), beta and delta. Positional adjustments are necessary for other search agents, including the omegas, so that are

in line with the top three solutions found so far. In Equations (11, 12 and 13) should reflect any changes to the wolves' locations.

$$\vec{C}_\alpha = |\vec{D}_1 \vec{W}_\alpha - \vec{W}|, \vec{C}_\beta = |\vec{D}_2 \vec{W}_\beta - \vec{W}|, \vec{C}_\delta = |\vec{D}_3 \vec{W}_\delta - \vec{W}| \quad (11)$$

$$\vec{W}_1 = |\vec{W}_\alpha - \vec{A}_1 \cdot \vec{C}_\alpha|, \vec{W}_2 = |\vec{W}_\beta - \vec{A}_2 \cdot \vec{C}_\beta|, \vec{W}_\delta = |\vec{W}_\delta - \vec{A}_3 \cdot \vec{C}_\delta| \quad (12)$$

$$\vec{W}(s+1) = \frac{\vec{W}_1 + \vec{W}_2 + \vec{W}_3}{3} \quad (13)$$

Finally, the GWO parameter has been updated $\vec{b}$, determining how much of a price to pay for exploration vs exploitation. The parameter $\vec{b}$, using the Equation (14), ranges from 2 to 0 and it is updated linearly with each cycle.

$$\vec{b} = 2 - s \cdot \frac{2}{\text{Max}_{\text{iter}}} \quad (14)$$

Where, s the number of iterations and maxiter is the maximum number of iterations that can be used for optimization.

4. Results

4.1 Experimental Setup

The study used a Windows 8.1 operating system and a setup with a 2.33 GHz CPU and 4 GB of RAM. Python was used in the testing process. The objective of sarcasm detection in social media is to discover as well as categorize text or communications with ironic or mocking purposes in a funny or critical fashion and to tell them apart from straightforward, nonsarcastic statements. Social media sarcasm detection using IGWO-ERF methods was offered. The ability of the model to identify sarcasm has been evaluated using a variety of metrics, including accuracy detection, precision, recall and f1-score, with findings demonstrating that it beats opinion mining methods. When compared to state-of-the-art methods like K Nearest Neighbor (KNN) [20], Maximum Entropy (ME) [20] and Support Vector Machine (SVM) [20], our suggested technique, IGWO-ERF, performs much better.

4.2 Detection Accuracy

An accurate machine learning or natural language processing model can recognize and categorize sarcastic detection in social media posts. This statistic usually measures the proportion of properly identified strange remarks to the total number of statements. Accuracy is determined by comparing original data categories to anticipated categories. A reliable prediction occurs when the expected and actual categorizations coincide. To measure precision, divide the number of correct forecasts by the total number of suggestions. Table 1 and Figure 3 show the comparison of accuracy detection.

Table 1: Comparison of Detection Accuracy [Source: Author]

Methods	Detection Accuracy (%)
RF [20]	83.1
KNN [20]	81.5
ME [20]	77.4

GWO-ERF [Proposed]	96.1
--------------------	------

The following table compares the detection accuracy rates of many popular methods. In comparison to the results obtained by RF [20], KNN [20] and ME [20], the suggested GWO-ERF approach achieved an outstanding 96.1%.

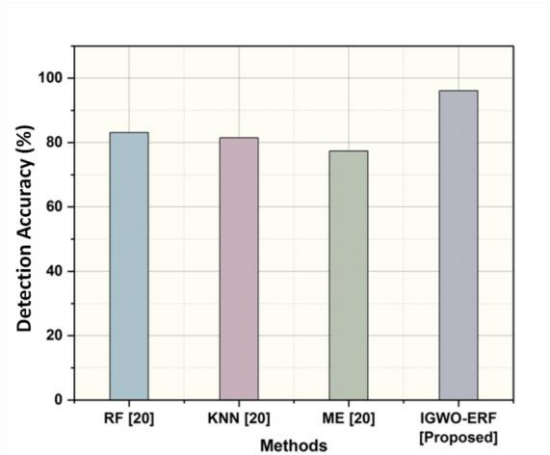


Figure 3: Detection Accuracy [Source: Author]

4.3 Precision

Precision in social media sarcasm detection is the model or system's ability to recognize actual sarcasm among the total number of occurrences it labels sardonic. It assesses how the algorithm avoids misclassifying non-sarcastic utterances as sarcastic, increasing the confidence that the recognized sarcastic comments are sarcastic. Performance is sometimes measured by the way that outcomes fit a category. A model determines the percentage of effective positive predictions from the total number of actual positive forecasts. Table 2 and Figure 4 show the comparison of precision

Table 2: Comparison of Precision [Source: Author]

Methods	Precision (%)
RF [20]	91.1
KNN [20]	88.9
ME [20]	84.6
GWO-ERF [Proposed]	94.23

The suggested GWO-ERF algorithm outperformed RF [20] by 91.1%, KNN [20] by 88.5% and ME [20] by 84.6% to reach the maximum accuracy.

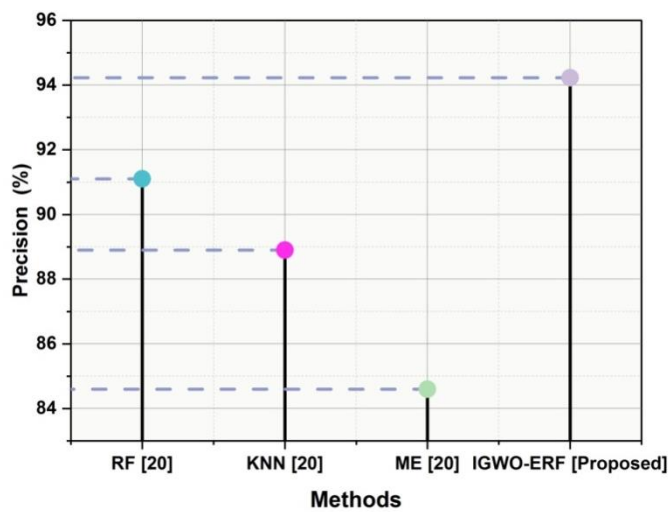


Figure 4: Precision [Source: Author]

4.4 Recall

Multiplying the proportion of accurate positive predictions by the percentage of incorrect negative predictions yields the recall value. The algorithm's accuracy in training data classification is shown. Sarcasm detection via social media requires a model or system to identify and extract the instances of acidic material from a dataset or stream of social media postings. The percentage of accurate recognitions measures the detection system's performance in detecting sarcastic postings. Table 3 and Figure 5 show the comparison of Recall.

Table 3: Comparison of Recall [Source: Author]

Methods	Recall (%)
RF [20]	73.4
KNN [20]	72
ME [20]	67
GWO-ERF [Proposed]	92.33

Recalls are 73.4% for RF [20], 72% for KNN [20] and 67% for ME [20] were overtaken by the 92.33% attained with the proposed GWO-ERF.

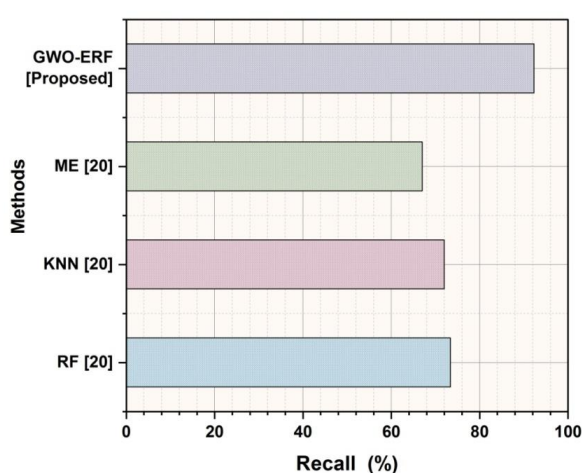


Figure 5: Recall [Source: Author]

4.5 F1-score

The F1-score measures a machine learning model's ability to distinguish sardonic and non-sarcastic material in social media messages. It balances accuracy (the percentage of identified sarcasm that is truly sarcastic) and recall (the proportion of actual sarcasm that is accurately recognized) to assess the model's ability to differentiate it from ordinary material. The F1 score is a standard performance statistic for task categorization. A high-quality F1-score implies a balanced model for accuracy and recall. Table 4 and Figure 6 show the comparison of the F1-score.

Table 4: Comparison of F1-score [Source: Author]

Methods	F1-Score (%)
RF [20]	81.3
KNN [20]	79.6
ME [20]	74.8
GWO-ERF [Proposed]	90.12

With an F1-Score of 90.12%, the suggested GWO-ERF method outperformed RF [20], KNN [20] and ME [20], which had scores of 81.3%, 79.6% and 74.8%, respectively.

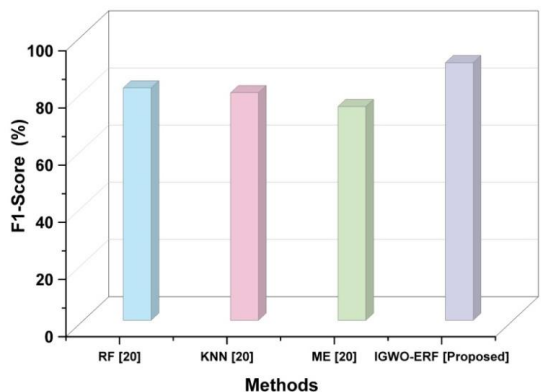


Figure 6: F1-score [Source: Author]

5. Discussion

The paper discusses the study on sarcasm detection, using data collected from social media and online retail platforms like Twitter and Facebook. Context-based, machine learning-based, pattern-based and rule-based methods are discussed in detail, in addition to the many aspects that are used by each of these four separate approaches. The incongruity in the text was the primary factor that led to the formation of sarcasm; nevertheless, to identify sarcasm, the detection method can involve seeing beyond this incongruity. The paper discusses the wide range of difficulties associated with sarcasm detection and it explains recent developments in the techniques employed for sarcasm identification. In the last step of this research project, a comparison study is conducted on four algorithms: K-nearest Neighbors (KNN) [20], Support Vector Machine (SVM) [20] and Maximum Entropy (ME) [20]. These algorithms were used to identify whether or not a given dataset included sarcastic data. Accuracy, recall, precision and F1-measure are the four performance measures used to compare a study.

6. Conclusion

The unboundedness of sarcasm means it is difficult for a computer to understand human emotions in general, particularly sarcasm and we are well aware of this. Algorithms must be processed repeatedly for a computer to get used to these recurring challenges. Reading people's tweets to understand their feelings about a topic is an interesting exercise. As a result, the primary focus of this work is on the problem of sarcasm detection using Twitter data, which is changing. After-sale support and customer satisfaction can be improved by reviewing client comments and complaints with an understanding of their motivations. In this paper, we proposed IGWO-ERF techniques for detecting Twitter-based sarcasm. This study obtained the EDNN-LSTM with 96.10% accuracy, 94.23% precision, 92.33% recall and 90.12% F1-score utilizing our proposed approach. Research into this area can shed light on how a computer could recognize sarcasm and might be used in the future to identify sarcastic tweets as either positive or negative sarcasm. The same method can be used for the elaborate forms of sarcasm

such as satire, pun, banter, comedy, etc., allowing a computer to grasp them better and make the appropriate conclusion.

References

1. Salo, Markus, Henri Pirkkalainen, and Tiina Koskelainen. "Technostress and social networking services: Explaining users' concentration, sleep, identity, and social relation problems." *Information Systems Journal* 29, no. 2 (2019): 408-435. <https://doi.org/10.1111/isj.12213>
2. Lüders, Adrian, Alejandro Dinkelberg, and Michael Quayle. "Becoming "us" in digital spaces: How online users creatively and strategically exploit social media affordances to build up social identity." *Acta Psychologica* 228 (2022): 103643. <https://doi.org/10.1016/j.actpsy.2022.103643>
3. Mauchand, Maël, Nikos Vergis, and Marc D. Pell. "Irony, prosody, and social impressions of affective stance." *Discourse Processes* 57, no. 2 (2020): 141-157. <https://doi.org/10.1080/0163853X.2019.1581588>
4. Leung, Janni, Jack Yiu Chak Chung, Calvert Tisdale, Vivian Chiu, Carmen CW Lim, and Gary Chan. "Anxiety and panic buying behaviour during COVID-19 pandemic—a qualitative analysis of toilet paper hoarding contents on Twitter." *International journal of environmental research and public health* 18, no. 3 (2021): 1127. <https://doi.org/10.3390/ijerph18031127>
5. Eke, Christopher Ifeanyi, Azah Anir Norman, Liyana Shuib, and Henry Friday Nweke. "Sarcasm identification in textual data: systematic review, research challenges and open directions." *Artificial Intelligence Review* 53 (2020): 4215-4258. <https://doi.org/10.1007/s10462-019-09791-8>
6. Hasnat, Fahim, Md Mazidul Hasan, Abdullah Umar Nasib, Ashik Adnan, Nazifa Khanom, SM Mahsanul Islam, Md Humaion Kabir Mehedi, Shadab Iqbal, and Annajiat Alim Rasel. "Understanding sarcasm from reddit texts using supervised algorithms." In *2022 IEEE 10th Region 10 Humanitarian Technology Conference (R10-HTC)*, pp. 1-6. IEEE, 2022. <https://doi.org/10.1109/R10-HTC54060.2022.9929882>
7. Akula, Ramya, and Ivan Garibay. "Explainable detection of sarcasm in social media." In *Proceedings of the Eleventh Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis*, pp. 34-39. 2021.
8. Mehndiratta, Pulkit, and Devpriya Soni. "Identification of sarcasm in textual data: A comparative study." *Journal of Data and Information Science* 4, no. 4 (2019): 56-83. <https://doi.org/10.2478/jdis-2019-0021>
9. Sandhu, Mannila, C. Danielle Vinson, Vijay K. Mago, and Philippe J. Giabbanelli. "From associations to sarcasm: mining the shift of opinions regarding the supreme court on twitter." *Online Social Networks and Media* 14 (2019): 100054. <https://doi.org/10.1016/j.osnem.2019.100054>
10. Cortis, Keith, and Brian Davis. "Over a decade of social opinion mining: a systematic review." *Artificial intelligence review* 54, no. 7 (2021): 4873-4965. <https://doi.org/10.1007/s10462-021-10030-2>
11. Godara, Jyoti, Rajni Aron, and Mohammad Shabaz. "Sentiment analysis and sarcasm detection from social network to train health-care professionals." *World Journal of Engineering* 19, no. 1 (2022): 124-133. <https://doi.org/10.1108/WJE-02-2021-0108>
12. Can, Umit, and Bilal Alatas. "A new direction in social network analysis: Online social network analysis problems and applications." *Physica A: Statistical Mechanics and its*

- Applications 535 (2019): 122372. <https://doi.org/10.1016/j.physa.2019.122372>
13. Shah, Bhumi, and Margil Shah. "A Survey on Machine Learning and Deep Learning Based Approaches for Sarcasm Identification in Social Media." In *Data Science and Intelligent Applications: Proceedings of ICDSIA 2020*, pp. 247-259. Springer Singapore, 2021. https://doi.org/10.1007/978-981-15-4474-3_29
 14. Tsapatsoulis, Nicolas, and Constantinos Djouvas. "Opinion mining from social media short texts: Does collective intelligence beat deep learning?." *Frontiers in Robotics and AI* 5 (2019): 138. <https://doi.org/10.3389/frobt.2018.00138>
 15. Chauhan, Priyavrat, Nonita Sharma, and Geeta Sikka. "The emergence of social media data and sentiment analysis in election prediction." *Journal of Ambient Intelligence and Humanized Computing* 12 (2021): 2601-2627. <https://doi.org/10.1007/s12652-020-02423-y>
 16. Banerjee, Arghasree, Mayukh Bhattacharjee, Kushankur Ghosh, and Sankhadeep Chatterjee. "Synthetic minority oversampling in addressing imbalanced sarcasm detection in social media." *Multimedia Tools and Applications* 79, no. 47-48 (2020): 35995-36031. <https://doi.org/10.1007/s11042-020-09138-4>
 17. Dhore, Pratiksha, Priyanka Dudhe, and Mona Mulchandani. "An empirical perspective on sarcasm detection in tweets using machine learning techniques." In *AIP Conference Proceedings*, vol. 2782, no. 1. AIP Publishing, 2023. <https://doi.org/10.1063/5.0155036>
 18. Akula, Ramya, and Ivan Garibay. "Interpretable multi-head self-attention architecture for sarcasm detection in social media." *Entropy* 23, no. 4 (2021): 394. <https://doi.org/10.3390/e23040394>
 19. Pawar, Neha, and Sukhada Bhingarkar. "Machine learning based sarcasm detection on Twitter data." In *2020 5th international conference on communication and electronics systems (ICCES)*, pp. 957-961. IEEE, 2020. <https://doi.org/10.1109/ICCES48766.2020.9137924>
 20. Ahire, L. K., Sachin D. Babar, and Gitanjali R. Shinde. "Sarcasm detection in online social network: Myths, realities, and issues." *Security issues and privacy threats in smart ubiquitous computing* (2021): 227-238. https://doi.org/10.1007/978-981-33-4996-4_15