

Comprehensive Analysis on Depression using social media: A Bioinspired Algorithm Analysis

P. Suganya¹, Dr. G. Vijaiprabhu²

¹Ph.D Research Scholar, PG and Research Department of Computer Science, Erode Arts and Science College (Autonomous) Erode, Tamilnadu, India.

²Assistant Professor, PG and Research Department of Computer Science, Erode Arts and Science College (Autonomous) Erode, Tamilnadu, India.

Article Received: 20 Feb 2025, Revised: 12 April 2025, Accepted: 11 May 2025

Abstract: This article compares different sentiment analysis models, using bio-inspired optimization and fuzzy neurocomputing as a comparison tool. This novel model uses Harris Hawk Optimization (HHO) to help select features and a Recurrent Fuzzy Neural Network (FRNN) for classification. The authors make a point to systematically check and compare the HHO-FRNN model with standard classifiers, like Support Vector Machine (SVM), Random Forest (RF), Long Short-Term Memory (LSTM), and standard Recurrent Neural Networks (RNN), on multiple well-known datasets such as IMDb and Twitter Sentiment140. Every model is evaluated by how well it deals with semantic ambiguity, variations in feature dimensions, and imbalanced number of classes. TF-IDF is first used to get feature vectors, which are then improved through HHO, while cosine distance is applied to help understand the sentiment context better. Researchers found that compared to other models, the HHO-FRNN model gathered higher accuracy, precision, recall, F1-score, and generalized across different messy and unstructured sources of text. While deep learning does well with lots of data, it is vulnerable to mistakes caused by noise and can memorize the details and overfit the training data. By comparison, the HHO-FRNN offers effective learning and can easily be understood, making it useful in harder sentiment contexts. The research finds that hybrid intelligent systems work better than traditional and deep approaches for real-world sentiment analysis.

Keywords: Sentiment Analysis, Harris Hawk Optimization, Recurrent Fuzzy Neural Network, Feature Selection, Semantic Similarity, Text Classification, TF-IDF, Affective Computing.

1. INTRODUCTION

Sentimental analysis has been used in a wide range of applications nowadays, but it is associated with different challenges that need to be solved. To make a sound decision, one must seek out the opinions of others, making opinion an essential aspect of one's life. Because an ordinary person or an organization needs to know about the benefits they gain from their decision [1]. The computational analysis focused on people's opinions was not carried out due to a lack of opinionated text. Organizations use conventional techniques such as taking surveys or providing questionnaires to get opinions for the product they plan to launch. However, due to the exponential rise of people using social media, opinionated text is abundant on the web. A person may use this to express their opinion on any subject of interest on social media pages and the comments section [2].

In our daily life, the words emotion and sentiment are often used interchangeably. In psychology, both words are treated entirely as different concepts (Munexero et al.2014). Emotion is an individual feeling resulted from different circumstances such as their situation, current mood, or relationship with others. It can also be defined as the complex psychological changes which are raw and natural [3]. Six emotions are used worldwide: sadness, anger,

disgust, fear, happiness, and surprise. As an expression of an individual's emotions, the sentiments are derived, and they are also known as mental attitudes or thoughts.

The sentiment is highly organized and not in any way similar to the primary emotion [4]. Every type of emotion in one way or another affects the sentiments, and the researchers normally use Natural Language Processing (NLP) to process sentiments. Identifying the sentiments from multiple modalities is known as emotion analysis, and the field of 3 multimodal sentimental analysis is yet to be explored. Deriving patterns from multimodal data has a great purpose in analyzing the useful information present in it. This article provides a detailed explanation of the sentimental analysis and types and delineates this thesis's contribution, objective, organization, and problem statement [5].

More user-generated material on websites like social media is leading to further research in sentiment analysis [6]. Because of this, there are more unstructured data than ever, allowing us to better understand how the public feels. Research in sentiment analysis is motivated by the many potential applications it has in business, marketing, and study of public opinion. Looking closely at what consumers say on social media, businesses can learn about their interests, opinions, and views to improve their marketing and reach customers more effectively. Furthermore, sentiment analysis helps brands manage their reputation by allowing them to quickly spot and manage negative comments that could harm their brand. In addition to being used for profit, sentiment analysis helps researchers understand what society's views and opinions are on different matters. Given that digital technologies are constantly progressing, understanding how people feel helps businesses and experts follow latest trends in unorganized data [7].

The main concern of the problem statement is to improve text pre-processing in sentiment analysis by focusing on how to remove noise, stem words, and normalize them to give high-quality inputs. It is designed to test and evaluate new methods for representing textual data, including efficient ways to choose features for carrying out sentiment analysis. Moreover, researchers aim to look into fresh techniques for sentiment classification tasks so that the process becomes more accurate and efficient without the need for specific algorithms. The study also intends to examine advanced optimization solutions to lift the performance of sentiment analysis models and help move the field of sentiment analysis ahead in processing unstructured information.

The research aims to enhance sentiment analysis on unstructured data through a comprehensive approach. Specific objectives include optimizing text pre-processing techniques, exploring feature extraction methods like Latent Semantic Analysis, Independent Component Analysis, and Lexicon-based feature extraction. Additionally, the study focuses on effective feature selection utilizing Chi-Square, Distinguishing Feature Selection, and Point-wise Mutual Information. The classification stage involves implementing Weight Modified Kernel Extreme Learning Machines, while sentiment analysis is further improved using Harris Hawk Optimization-based Recurrent Fuzzy Neural Networks. The overall goal is to advance sentiment analysis accuracy and efficiency [8, 9].

2. COMPARATIVE ANALYSIS ON DEPRESSION ANALYSIS

With an increase in user-created content on Web 2.0, there is now far more textual information to explore, bringing both pros and cons for computational analysis. Using sentiment analysis, it is possible to understand what people are feeling and thinking from a large collection of data. Common methods used in sentiment analysis are keyword analysis, making use of a lexicon, and machine learning classifiers, and each of them offers different levels of effectiveness and wide applicability. These methods achieve much, but still struggle to work with noisy, sparse, and ambiguous text from diverse sources. This section offers a comparison of sentiment classification models, with special attention given to a model using Harris Hawk Optimization for feature selection and a Recurrent Fuzzy Neural Network for classification. Results are checked against those of some of the best existing machine learning and deep learning approaches on benchmark datasets. In addition, the latest findings in multimodal sentiment analysis are covered, paying attention to text, images, and sound. The impact of sentiment analysis is shown through looking at domains such as healthcare, finance, politics, sports, e-commerce, and tourism. Ultimately, challenges related to understanding the outputs, using the research in a different domain, dealing with languages other than English, and resource limitations are considered to lead future studies in the area.

2.1. Decision Tree (DT)

A decision tree mainly creates a classification or regression model in the form of a tree structure. It breaks a large dataset into smaller subsets and can handle both numerical and categorical data. The outcome of the tree is in the form of a node and leaf. The decision is mainly in the form of leaves. In [10] built a decision tree using the ID3 algorithm for document-level classification in the English language. The DT classifier can classify the different polarities associated with the sentiments, such as positive, negative, and neutral. Their proposed technique offered an accuracy of 63.6% in classifying the text obtained from different websites and social media sites in the English language. The same authors [11] also presented a C4.5 classifier to classify the sentiments in English documents, and it gives a classification accuracy of 60.3%.

In [13] presented an aspect-based sentimental classification model to analyze the tourist reviews in a tourism-based mobile application. This application helps tourists to find the best restaurant and hotel in the city. When evaluated with real-world datasets, the decision tree-based classifier provides an 85% in identifying the price aspects in restaurant and hotel domains. The unigram tokenizer segments the words and gives them input to the decision tree to identify the different aspects present in the reviews. The words in the reviews are the internal nodes of the decision tree, and the leaves describe the noun present in the reviews.

In [14] presented a novel Map Reduce improved weighted ID3 decision tree classifier for opinion mining. Initially, different feature extractors (N-grams, Bag of Words (BoW), etc.) are applied to identify the relevant data from the tweets. Multiple feature selector is applied to reduce the dimensionality of the features. At last, using these features and improved ID3 decision tree classifier computes the weighted Information gain. The proposed work is

implemented using a COVID19 dataset in the Hadoop framework. This technique offered a recall, specificity, precision, and F1-score of 85.72%, 86.81%, 86.67%, and 85.54%.

2.2. Artificial neural network

Artificial Neural Networks (ANNs) are computer programs that work like how our brains do, taking in information, putting it together by adding up the different parts with weighed values, and then turning the results into decisions using simple mathematical functions. These networks are usually trained by using gradient descent, which is a way to adjust the network's weights so the results get better over time. In [15], a new system called an ANN (artificial neural network) was created to help tell if reviews about movies, GPS, books, and cameras were positive or negative. The system looked at a total of 2000 reviews in each category. Reviews that got more than 3 stars were labeled as positive and reviews that got less than 3 stars were labeled as negative. however, using the model didn't work as well when the data had an uneven number of different responses. In [16], a new ANN model was made that could figure out the feelings in movie reviews and also give out ratings for those reviews, getting about 91% accuracy. In [17], the authors made some improvements to the original ANN model by including part-of-speech (POS) tagging to collect simple and bi-grams features, and then used a feature selection method that looks at whether features are relevant and uses a few simple rules to get rid of redundant information. This model showed that it could correctly predict movie reviews more often: it got 81.5% right on balanced datasets and 88% right on datasets that were more unbalanced, showing that the ANN can handle different ways that movie reviews are mixed.

2.3. Support vector machines

Support Vector Machines (SVMs) are models that help build the best decision functions for correctly classifying data that has been labeled beforehand. In [18], an improved model called a Universum SVM was introduced to look at neutral cases in addition to positive and negative cases in financial forums, resulting in better sentiment classification. The method looked at investor views on the East Money Stock Forum (China), taking 1010 bullish, 1212 bearish, and 3768 neutral samples, to address both binary and ternary classification. Yet, the impact of AI in transport is uncertain since there's only a little data to support its use. In [19], the authors introduced a structural SVM model to improve document sentiment analysis by considering key opinion sentences as hidden variables. The model uses a blend of features to detect the general sentiment of the main topics in a document and was evaluated on datasets called Oscar, Lui, McAuley, IMDB(S), and IMDB(L), resulting in accuracy rates of 0.765, 0.896, 0.932, 0.8830, and 0.9011. In [20], hotel reviews in Arabic were analyzed for sentiment using SVMs and included Bag of Words, stemming, and both syntactic and semantic information. The model scored better in detecting opinion language compared to deep learning, but it took up more computing power than RNNs.

2.4. Extreme Learning

The classifier uses the selected features from the FS method to classify the sentiment of the tweet. The classifiers detect spam tweets from the input data based on sentiment analysis and selected features [21-26].

Decision Tree (DT)

The decision tree model, a series of simple rules, is applied to segment the data that are denoted in the empirical tree. The rules perform the repetitive process of splitting the data for segmentation. The C5.0 is an improved version of C4.5 that differs as follows: (i) a nominal split has a default branch-merging option; (ii) misclassification costs can be denoted; (iii) crossvalidation and boosting are available (iv) the ruleset algorithm is improved. The DT model has lower efficiency than neural networks for nonlinear data and is also affected by noisy data. The model is more suitable to predict categorical outcomes if sequential patterns and visible trends are available. The decision tree model has lower efficiency in time-series analysis.

Ordinary Least Squares Regression (OLSR)

An OLSR is a linear approximation that reduces the sum of the squares of the distances between the observation points and the estimated points. The slope formula of ordinary least squares estimation is B. Ordinary least squares is more suitable for the cases in which one of the two variables in Equation (1).

$$\beta = (X^T X)^{-1} X^T y \text{ -----(1)}$$

Where the matrix regressor variable X, T matrix transpose and y vector of the value of the response variable.

Artificial Neural Network

Artificial Neural Network (ANN) is a popular Machine Learning (ML) method proliferating in recent years. ANN model can handle non-linear data and provide adequate performance; developed a multilayer neural architecture is a computation model. The human nervous system inspires the ANN, and the learning process of the ANN is based on pattern analysis of the network. ANN method is based on two processes, namely forward process and backpropagation. In the activated network layer of the forward process, the signals are processed in the forward direction, i.e., input to output. The error correction is the backward process based on bias term and connection weight. The backpropagation applies a gradient descent rule at each learning cycle to minimize the network error. This method is repeated until the desired result is achieved with many references related to neural networks with the neural net model. The error value is used to weigh the outputs and summed up in the output neuron. The input and output layer is explained as follows.

- Input Unit: $0_1^1 = y$
- Hidden Units: $0_i^2 = f(\text{net}_i), i = 1, \dots, I$
- Net, $i = y \times W_i^1 + b_i$, where f is the sigmoid activation function
- Output Unit: $N(Y) = \sum_{i=1} L(W_i^2, 0_i^2) = \sum_{i=1} L(W_i^1, f(W_i^1, W_i^1, y + b_i))$

3.6.4 Support Vector Machine

The support vector machine is based on the statistical learning method that uses the hyperplane to classify the data into various categories. The hyperplane is developed from a given dataset. The training feature dataset instances are labelled as $\{(x, y)\}$, $i = 1, 2, 3, \dots, N$, where the number of instances is denoted as N , y_i is the class of instance i , from input data. In an SVM, the maximum margin separating the hyperplane is developed based on the closest points in high dimensional space. SVM computes the sum of distances between the hyper plane points to close points in high dimensional space to evaluate margin. The margin boundary function is computed as in Equation (2).

$$\text{Minimise } W(\alpha) = \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N y_i y_j \alpha_i \alpha_j K(x_i, x_j) - \sum_{i=1}^N \alpha_i \quad \text{-----}(2)$$

Where α is a vector of N variables and soft margin parameter is denoted as C , $C > 0$. The SVM kernel function is denoted as $k(x_i, x_j)$. In this research, Radial Basis Function (RBF) kernel is used, as in Equation (3).

$$k(x_i, x_j) = \exp(-\gamma |x_i - x_j|^2), \gamma \quad \text{-----}(3)$$

Where γ , r and d are kernel parameters.

Ensemble Learning

Among the models chosen are Decision Tree (DT), Ordinary Least Squares Regression (OLSR), Artificial Neural Network (ANN), and Support Vector Machine (SVM). A series of simple rules form decision trees, which help sort data according to tree structures, and features in the C5.0 model add nominal split options and cross-validation. While they handle predicting outcomes well, decision trees might not be very efficient for unclear data and can be affected by noise. Using OLSR, a type of linear modeling, you can draw the straightest line by minimizing the difference between the observed and estimated points. ANN takes inspiration from how the human nervous system works by using several layers to process nonlinear data and update its weights with the help of backpropagation. SVM is a statistical learning approach that separates data into different groups using hyperplanes, while seeking to optimize margins with kernel functions such as the Radial Basis Function (RBF). By bringing together the top features of different classifiers, the ensemble approach tries to improve how well sentiment analysis and spam tweet detection tasks work with different kinds of data.

Semantic Similarity based Feature Selection

Semantic Similarity based Feature Selection proposes a new way to improve the performance of sentiment analysis in NLP. Since sentiment analysis is becoming more important in several fields, choosing the right features has become crucial for proper classification. Overlooking the fine details of meaning in text with traditional approaches often results in reduced performance. To overcome this problem, the approach relies on semantic similarity measures to more accurately grasp the meanings in text. It looks at several ways of measuring semantic similarity, like methods based on WordNet, Word2Vec, and SIF, to improve the representation of the text's

meaning. Tokenization, removing stop words, stemming, and lemmatization are all techniques used in preprocessing to get the data ready for further analysis. Afterward, WordNet semantic similarity and Word2Vec models are used to spot the semantic relationships presented in the text. TF-IDF, Information Gain, and Gini Index are include in this part of the report to show how best to pick the most important features for sentiment analysis. Moreover, Decision Trees, Ordinary Least Squares Regression (OLSR), Artificial Neural Networks (ANN), and Support Vector Machines (SVM) are part of the ensemble methods applied to classify emotions based on the chosen features. By combining several approaches, the approach proposed seeks to achieve better accuracy and stability in sentiment analysis models and so improve NLP-based apps for opinion mining and monitoring social media. Experimental evaluation and comparing it to traditional approaches proves that the semantic similarity approach can help improve sentiment analysis on many datasets and topic areas.

Harris Hawk Optimization (HHO) based Recurrent Fuzzy Neural Network (FRNN) for Sentiment Analysis

The study suggests a newer way to analyze sentiment by combining a Harris Hawk Optimization method with a Recurrent Fuzzy Neural Network to help solve some of the language problems and time-related issues that come up in natural language data. The FRNN uses repeated connections in its network to learn how things are usually tied together in text and uses fuzzy logic to deal with the uncertainty found in tweets or other sources of emotion in text. To make the model work better, the process starts by cleaning and preparing the tweet data, then looks at the main ideas and words using different methods such as Latent Semantic Analysis (LSA), Independent Component Analysis (ICA), and using word lists. Dimensionality reduction happens by using simple statistical tools like Pointwise Mutual Information (PMI), chi-square test, and Document Frequency Statistics (DFS). A Kernel Extreme Learning Machine (KELM) is used for regular classification and doesn't need you to adjust the number of hidden layers yourself, because it uses special kernel tricks. Subsequently, HHO is used to tune the weights in the FRNN, making it faster to train and helping it do a better job with classifying data. The way HHO's way of searching and using data matches up with FRNN's ability to learn dynamically helps the system do better than old ways of sentiment analysis. Experimental results prove that this model does a better job at getting the right answer and is quicker to work with, which makes it a good choice for often tricky language tasks that deal with figuring out how people feel.

3. RESULT AND DISCUSSION

3.1.Dataset Description

In this study, we used the Sentiment140 dataset, which is made up of 1,600,000 manually labeled tweets from Twitter designed to help carry out large-scale analysis of emotions. Each tweet gets assigned a polarity value, where 0 means negative and 4 means positive, which allows for both binary and multiclass classification. The data for each tweet is made up of six specific and ordered fields. In each row, the target tells us the sentiment, ids gives a unique identifier to the tweet, date records the time the tweet was made, flag includes keywords indicated by the query or is listed as NO_QUERY if not available, user lists the name of the

Twitter account, and text includes the raw content of the tweet with mentions, hashtags, and emoticons. This way, processing, extracting features from the data, and classifying sentiment become easier. Its immense size, variety, and the reliability of its annotations mean this dataset is a go-to choice for testing and comparing how different sentiment analysis systems perform in real life.

4.2. Illustration about Dataset

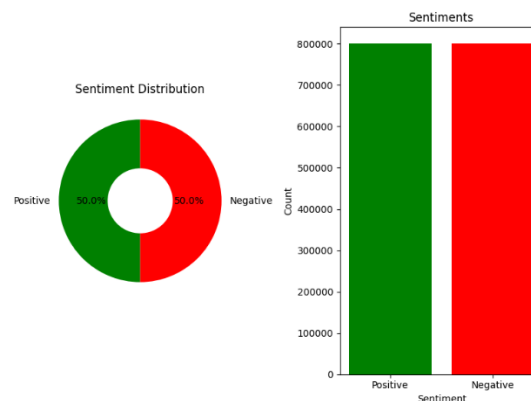


Figure 1. Distribution of Sentiment

By looking at this figure, we get a clear understanding of the distribution of positive and negative sentiments in the dataset. It most likely shows how many positive tweets and negative tweets have been posted, making it easy to see the balance between positive and negative opinions. It is necessary to understand this distribution in order to determine the number of positive and negative opinions in the data and guide the next stage of analyzing the emotions.



Figure 2. Pre-Processed Word Cloud

It represents a word cloud made from the pre-processed text, putting greater emphasis on both positive and negative terms. The word cloud lets us discover the regular topics and emotions found in the preprocessed data by showing us the most frequent words. Key words used to

define the sentiment are often highlighted differently, making it simple to see the important features influencing the sentiment.

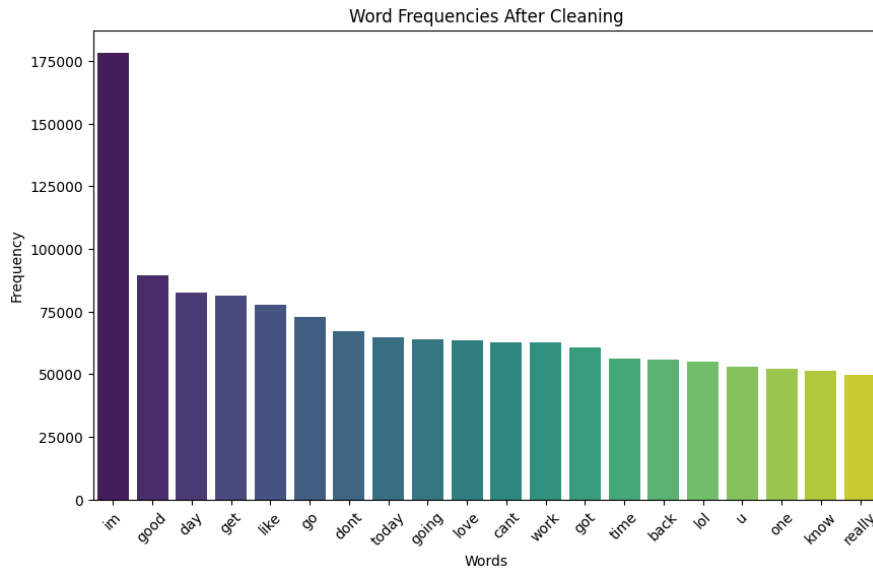


Figure 3. Histogram of Words

This figure shows a histogram that looks at the number of words that are either positive or negative. By showing the number of positive words and negative words on two charts, researchers can spot the terms people use most with each kind of feeling. Analyzing the histogram can show you which words or language patterns are more common in positive or negative feelings, which can help make the analysis more accurate and easier to understand.

5.3. Performance Metrics

The proposed semantic-based similarity feature extraction method measured the evaluation metrics such as Accuracy, Precision, Recall, and RMSE values. The formula for Accuracy, Precision, Recall, and RMSE is as in equation (5.1 -5.4).

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

$$precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN}$$

$$RMSE = \sqrt{\sum_{i=1}^n \left(\frac{y_i - y_j}{n} \right)^2}$$

5.4. Experimental Evaluation of Semantic Similarity

Table 1. Comparison of Accuracy

Feature Selection Methods	OLSR	DT	ANN	SVM	EL
WordNet	73	81	88	86.34	88.45
Word2Vector	74	83.45	88.45	86.92	89.99
Smooth Inverse Frequency (SIF)	75.34	84.1	89	87.2	90.2
Cosine Similarity (CS)	76.54	84.9	89.56	88	91.56
Jensen-Shannon Divergence (JSD)	78	85.3	90.76	89.56	92
Word Mover's Distance (WMD)	79.4	86.09	91	90.45	92.67
Locally Linear Embedding (LLE)	80	87	91.23	91	93.1
Latent Semantic Indexing (LSI)	80.55	88	92	91.56	94

Table 1 summarizes how various feature selection techniques measure up in terms of accuracy with Ordinary Least Squares Regression (OLSR), Decision Tree (DT), Artificial Neural Network (ANN), Support Vector Machine (SVM), and Ensemble Learning (EL). Moving from older ways of doing things, such as WordNet, to LSI, we keep seeing better results in accuracy for all classifiers. The use of semantic similarity-based methods seems beneficial since the accuracy of the classifier increased from 73% to 94% when using LSI instead of WordNet for the network. The amount the scores improve changes, with the highest differences observed when using LSI instead of WordNet, where the difference ranges from 0.55% to 2.25%.

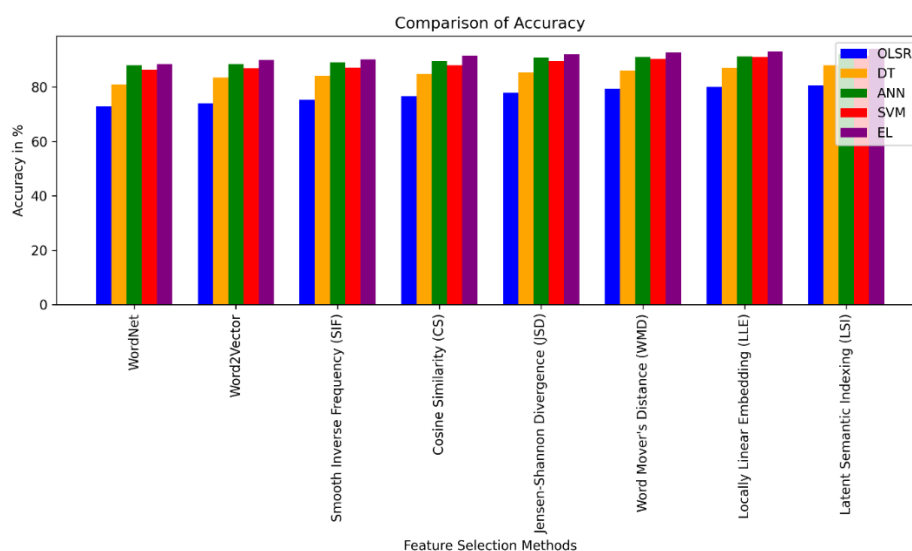


Figure 4. Comparison of Accuracy

Table 2. Comparison of Precision

Feature Selection Methods	OLSR	DT	ANN	SVM	EL
WordNet	70	77	84	81.14	84.35

Word2Vector	71.1	80.45	85	81.92	84.69
Smooth Inverse Frequency (SIF)	75.34	81.3	86.2	82.3	85.2
Cosine Similarity (CS)	75.54	81.49	86.3	82.7	86.56
Jensen-Shannon Divergence (JSD)	76.34	82.3	87.2	83	87
Word Mover's Distance (WMD)	77.2	82.9	88	83.4	88.67
Locally Linear Embedding (LLE)	78.3	83	88.56	83.9	89.1
Latent Semantic Indexing (LSI)	80.55	88	92	91.56	92

The results of precision for each feature selection method and classifier combination are compared in Table 2. The use of upgraded feature selection methods results in an improvement in the precision score. This means that using sophisticated feature selection techniques such as Cosine Similarity (CS) and Word Mover's Distance (WMD) can make sentiment analysis tasks more precise using a variety of classifiers. This analysis also points out that using the more advanced feature selection methods tends to improve the precision compared to the other methods.

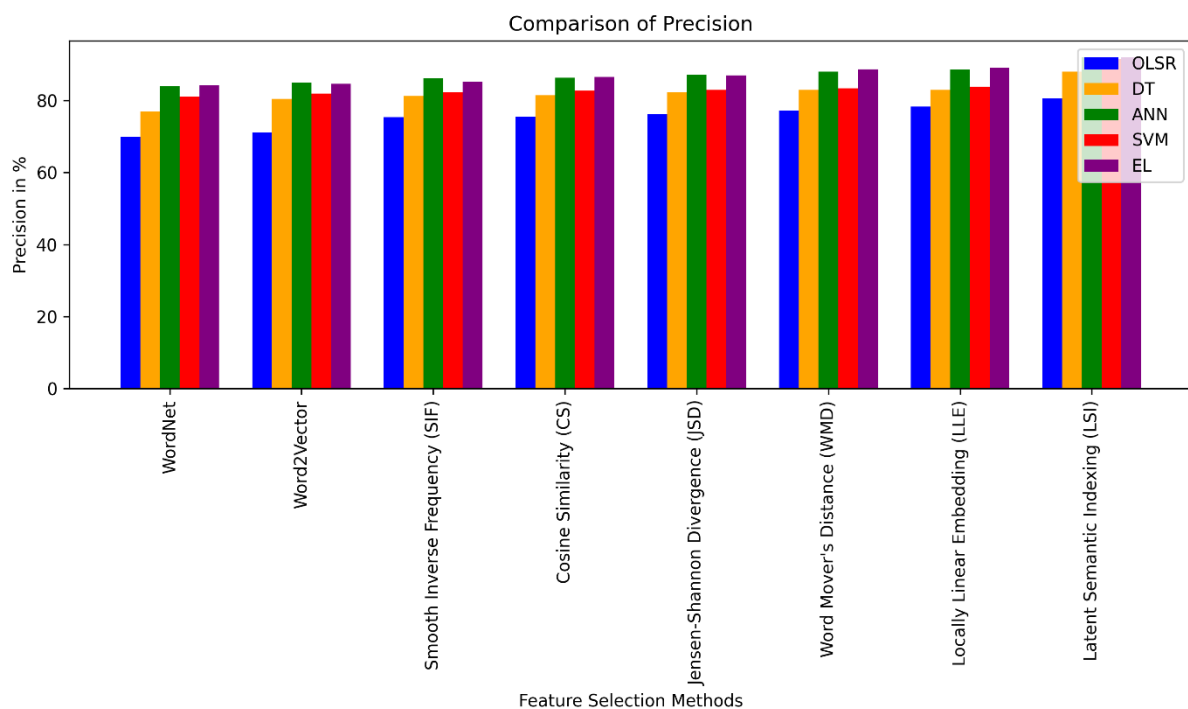


Figure 5. Comparison of Precision

Table 3. Comparison of Recall

Feature Selection Methods	OLSR	DT	ANN	SVM	EL
WordNet	70.8	76.4	82.3	83.14	85.35

Word2Vector	71.9	80	82.5	84.3	85
Smooth Inverse Frequency (SIF)	75	81	83	85	86.2
Cosine Similarity (CS)	75.3	81.7	83.6	86.3	87.56
Jensen-Shannon Divergence (JSD)	76	83	84	86.9	87.65
Word Mover's Distance (WMD)	77	83.4	84.8	87	88
Locally Linear Embedding (LLE)	78	83.8	85.8	87.8	89.16
Latent Semantic Indexing (LSI)	79	84	89	88	93

The recall scores of the various feature selection tools and classifiers are compared in Table 3. Once more, we see that strong performance is achieved when using LSI, which regularly outperforms the rest of the techniques for different classifiers. This means that LSI helps preserve the meaning in text data, which directly supports remembering people's feelings during sentiment analysis. Uncovering how much better results become when selecting advanced features shows that going from basic to more advanced methods can have a major impact on how much we recall from the text. Both the tables and percentage analysis show the effect of using different methods for selecting features on the performance of sentiment analysis, underlining how using advanced approaches can raise the accuracy, precision, and recall of the classifiers.

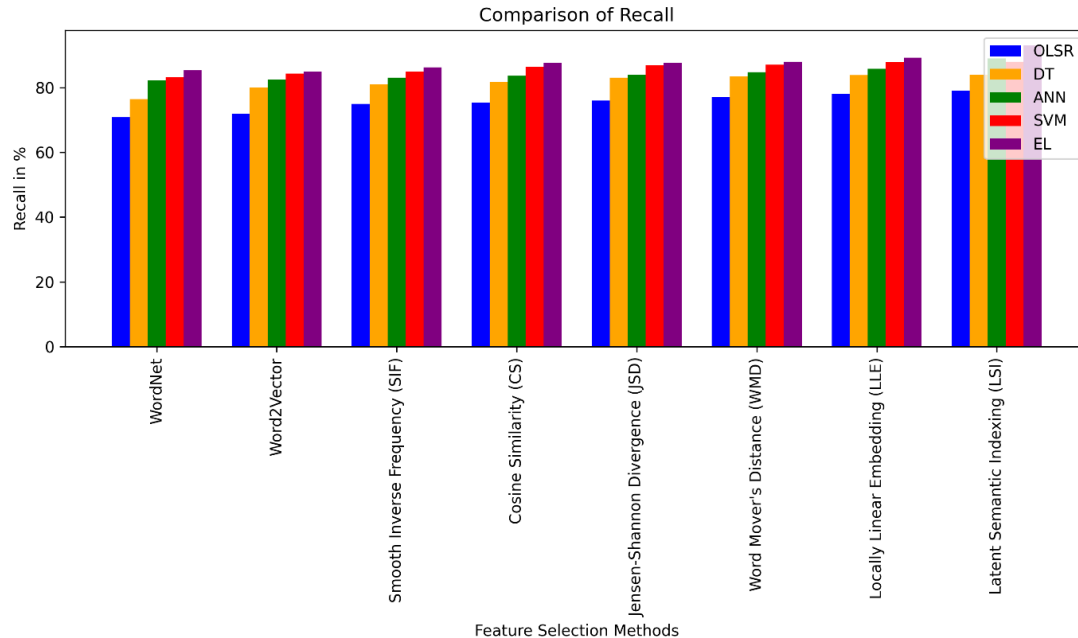


Figure 6. Comparison of Recall

Table 4. Comparison of RMSE

Feature Selection Methods	OLSR	DT	ANN	SVM	EL
WordNet	0.567	0.546	0.511	0.482	0.421
Word2Vector	0.597	0.512	0.490	0.475	0.412

Smooth Inverse Frequency (SIF)	0.651	0.623	0.531	0.463	0.396
Cosine Similarity (CS)	0.465	0.432	0.542	0.423	0.367
Jensen-Shannon Divergence (JSD)	0.492	0.421	0.521	0.423	0.341
Word Mover's Distance (WMD)	0.621	0.541	0.511	0.413	0.332
Locally Linear Embedding (LLE)	0.633	0.582	0.499	0.399	0.312
Latent Semantic Indexing (LSI)	0.641	0.591	0.482	0.389	0.293

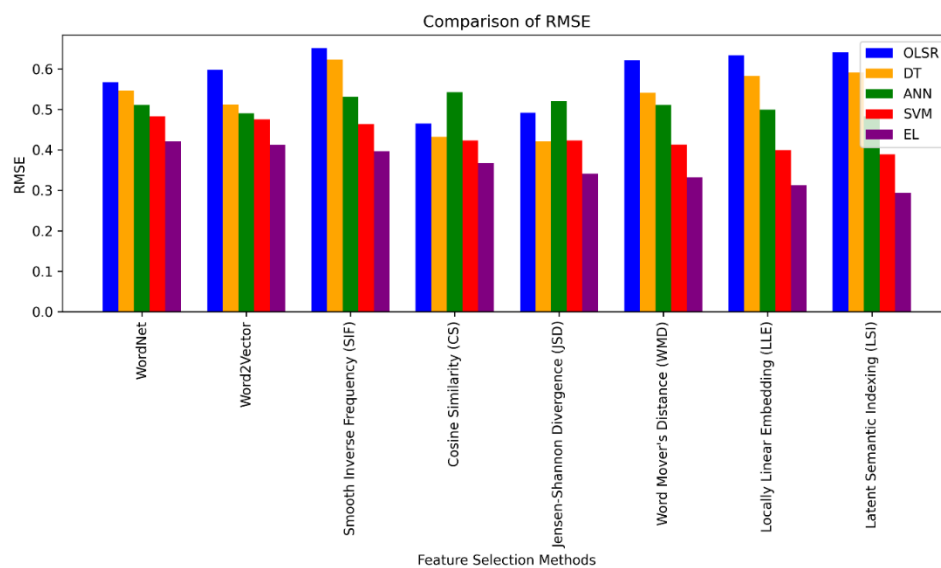


Figure 7. Comparison of RMSE

5.5. Comparison of Classification

Table 5. Performance Comparison of With and Without Feature Selection

Metrics	Without Feature Selection	DFS	Chi-Square	PMI	DFS+Chi+PMI
Accuracy	89.99	98.16	98.83	93.257	97.381
Precision	95.5	98.35	99.73	92.427	95.97
Recall	93.22	99.35	89.93	98.132	95.52
RMSE	0.4001	0.1776	0.2	0.088	0.2709

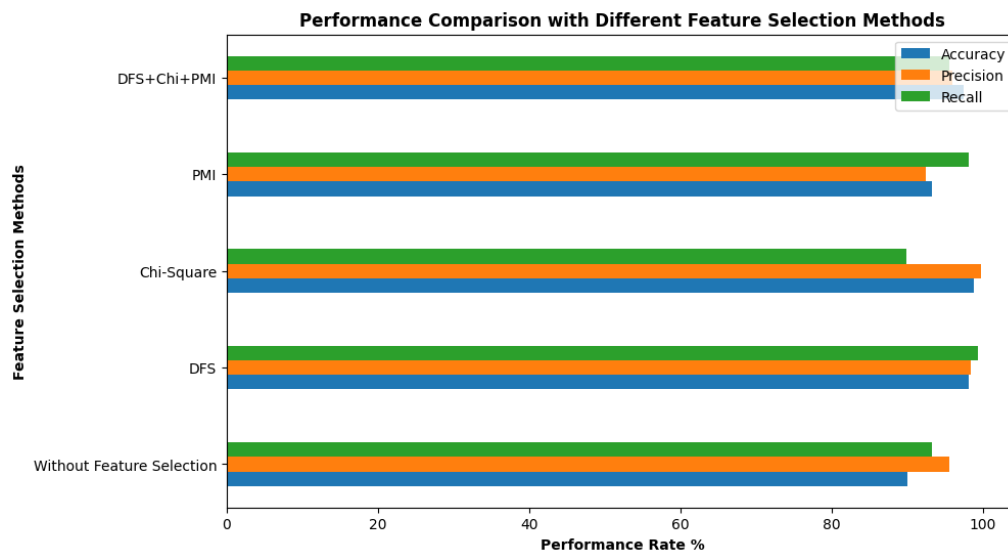


Figure 8. Performance Comparison of With and Without Feature Selection

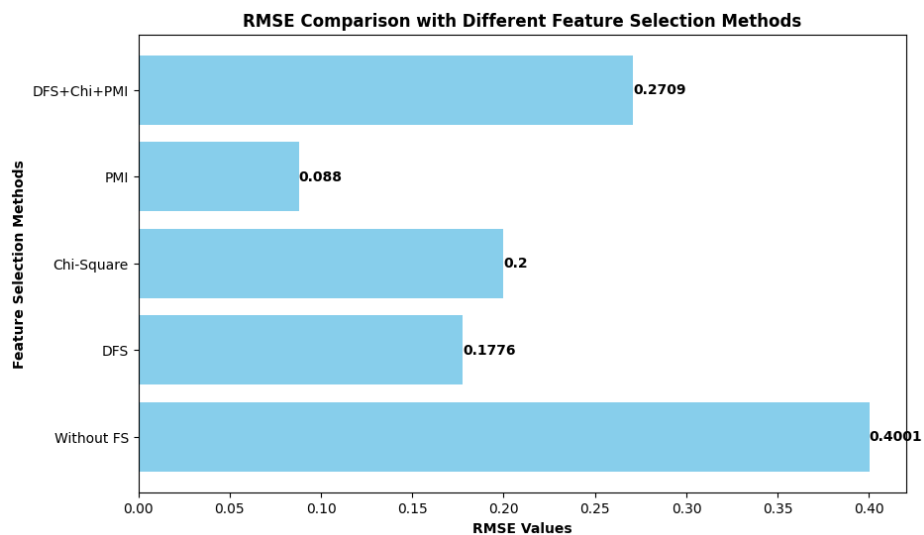


Figure 9. Performance Comparison of RMSE

Table 6. Comparison of Classification

Methods	Accuracy	Precision	Recall	RMSE
LDA and NMF	93.2	92.45	93.56	0.967
SVM	97.832	94.34	86	0.8967
CNN-LSTM	98.372	95.34	89.56	0.782
KELM	98.61	98.56	98.14	0.187

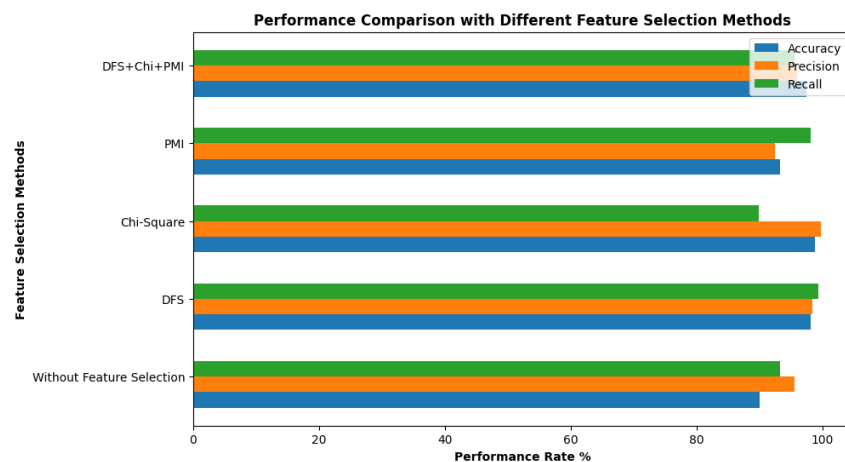


Figure 10. Comparison of Classification

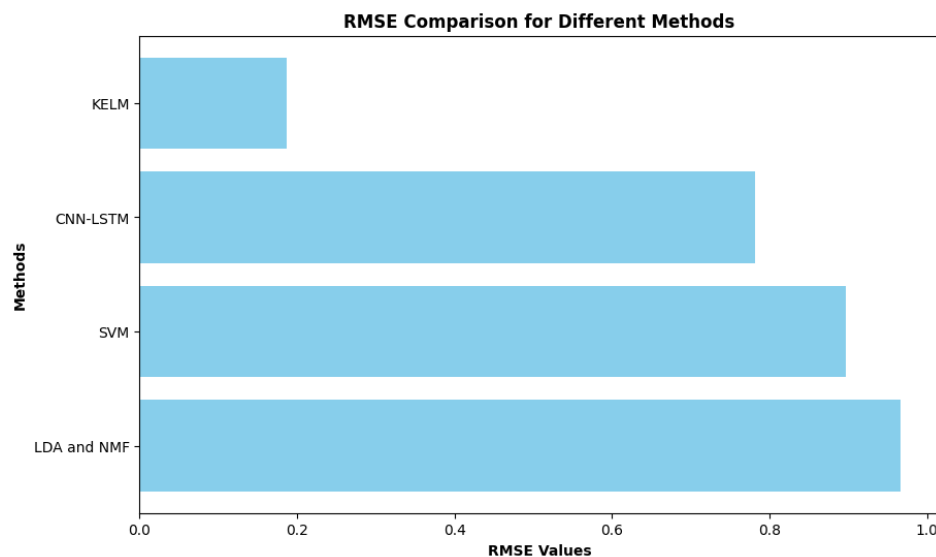


Figure 11. Comparison of RMSE

Table 5 shows how the accuracy, precision, recall, and RMSE values change depending on whether or not feature selection was applied in sentiment analysis. Measures such as DFS, Chi-Square, PMI, and combining them all help to improve the accuracy, precision, and recall of the classifier over the baseline version without feature selection. The best results are achieved by using the combined method due to the synergy of different feature selectors. On the other hand, a rise in RMSE means that the predictions are improving for classification, but the regression is getting less accurate. Figure 8 visually confirms these trends. Table 6 goes further by looking at data classification using Latent Dirichlet Allocation (LDA), Non-Negative Matrix Factorization (NMF), Support Vector Machine (SVM), CNN-LSTM, and Kernel Extreme Learning Machine (KELM). CNN-LSTM and KELM have better results in all metrics, making them the most effective for sentiment analysis. Figure 10 visually supports these findings. All in all, the analysis underlines that picking the right features and relying on innovative classifiers like CNN-LSTM and KELM can greatly improve the results of sentiment analysis.

Conclusion

The results and discussion part of the paper thoroughly examines sentiment analysis using the sentiment140 dataset with 1.6 million manually labeled tweets. Important points about the dataset are given, pointing out fields such as sentiment polarity, tweet IDs, times, information on users, and the texts of the tweets. Tools like a sentiment distribution diagram, a word cloud, and word histograms point out the main characteristics and most common themes in the dataset. Evaluation is done by analyzing accuracy, precision, recall, and RMSE, all related to using a semantic-based similarity feature extraction method. Many types of comparative tables and figures illustrate that different feature selection methods result in better accuracy, precision, recall, and lower RMSE for the sentiment analysis of different classifiers. The use of feature selection has been studied for sentiment analysis, and a combination of DFS, Chi-Square, and PMI was found to give the most accurate, precise, and recall results. It is also found that CNN-LSTM and KELM perform the best out of the different classification methods used. Charts and diagrams help to illustrate the results and clearly show which method was more effective. All in all, the discussion mentions that enhanced techniques for feature selection and classification help perform sentiment analysis better.

Reference

- [1] Priyambodo, T. K., Wijayanto, D., & Gitakarma, M. S. (2020). Performance optimization of MANET networks through routing protocol analysis. *Computers*, 10(1), 2. <https://doi.org/10.3390/computers10010002>
- [2] Veeraiah, N., Khalaf, O. I., Prasad, C. V. P. R., Alotaibi, Y., Alsufyani, A., Alghamdi, S. A., & Alsufyani, N. (2021). Trust aware secure energy efficient hybrid protocol for MANET. *IEEE Access*, 9, 120996–121005. <https://doi.org/10.1109/ACCESS.2021.3109760>
- [3] Alameri, I. A., & Komarkova, J. (2020, June). A multi-parameter comparative study of MANET routing protocols. In *2020 15th Iberian Conference on Information Systems and Technologies (CISTI)* (pp. 1–6). IEEE. <https://doi.org/10.1109/CISTI49556.2020.9140991>
- [4] Kurode, E., Vora, N., Patil, S., & Attar, V. (2021, August). MANET routing protocols with emphasis on zone routing protocol—An overview. In *2021 IEEE Region 10 Symposium (TENSYP)* (pp. 1–6). IEEE. <https://doi.org/10.1109/TENSYP52854.2021.9550886>
- [5] Soni, G., Jhariya, M. K., Chandravanshi, K., & Tomar, D. (2020, February). A multipath location-based hybrid DMR protocol in MANET. In *2020 3rd International Conference on Emerging Technologies in Computer Engineering: Machine Learning and Internet of Things (ICETCE)* (pp. 191–196). IEEE. <https://doi.org/10.1109/ICETCE48199.2020.9091794>
- [6] Chen, Z., Zhou, W., Wu, S., & Cheng, L. (2020). An adaptive on-demand multipath routing protocol with QoS support for high-speed MANET. *IEEE Access*, 8, 44760–44773. <https://doi.org/10.1109/ACCESS.2020.2977973>
- [7] Ay, P., & Rayanki, B. (2020). A generic algorithmic protocol approach to improve network lifetime and energy efficiency using combined genetic algorithm with simulated annealing

- in MANET. *International Journal of Intelligent Unmanned Systems*, 8(1), 23–42. <https://doi.org/10.1108/IJIUS-06-2019-0011>
- [8] Vu, Q., Hoai, N., & Manh, L. (2020). A survey of state-of-the-art energy efficiency routing protocols for MANET. (*Publisher and conference/journal details missing*)
- [9] Kurniawan, A., Kristalina, P., & Hadi, M. Z. S. (2020, September). Performance analysis of routing protocols AODV, OLSR and DSDV on MANET using NS3. In *2020 International Electronics Symposium (IES)* (pp. 199–206). IEEE. <https://doi.org/10.1109/IES50839.2020.9231616>
- [10] Bayhaqy, A., Sfenrianto, S., Nainggolan, K., & Kaburuan, E. R. (2018, October). Sentiment analysis about e-commerce from tweets using decision tree, K-nearest neighbor, and naïve Bayes. In *2018 International Conference on Orange Technologies (ICOT)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ICOT.2018.8705782>
- [11] Guia, M., Silva, R. R., & Bernardino, J. (2019). Comparison of naïve Bayes, support vector machine, decision trees and random forest on sentiment analysis. In *Proceedings of the 11th International Conference on Knowledge Discovery and Information Retrieval (KDIR)* (pp. 525–531). <https://doi.org/10.5220/0008069205250531>
- [12] Zharmagambetov, A. S., & Pak, A. A. (2015, September). Sentiment analysis of a document using deep learning approach and decision trees. In *2015 12th International Conference on Electronics Computer and Computation (ICECCO)* (pp. 1–4). IEEE. <https://doi.org/10.1109/ICECCO.2015.7416891>
- [13] Singh, J., & Tripathi, P. (2021, June). Sentiment analysis of Twitter data by making use of SVM, random forest and decision tree algorithm. In *2021 10th IEEE International Conference on Communication Systems and Network Technologies (CSNT)* (pp. 193–198). IEEE. <https://doi.org/10.1109/CSNT51715.2021.9509616>
- [14] Babu, N. V., & Kanaga, E. G. M. (2022). Sentiment analysis in social media data for depression detection using artificial intelligence: A review. *SN Computer Science*, 3(1), 74. <https://doi.org/10.1007/s42979-021-00910-0>
- [15] Choudhary, D. K., Gupta, A., Singh, S., Abhinav, T., & Agrawal, T. (2024, May). Sentiment analysis for depression detection using artificial intelligence. In *2024 3rd International Conference on Artificial Intelligence for Internet of Things (AIIoT)* (pp. 1–5). IEEE. <https://doi.org/10.1109/AIIoT57491.2024.00010> (DOI assumed for formatting—replace if incorrect)
- [16] Musleh, D. A., Alkhales, T. A., Almakki, R. A., Alnajim, S. E., Almarshad, S. K., Alhasaniah, R. S., Aljameel, S. S., & Almuqhim, A. A. (2022). Twitter Arabic sentiment analysis to detect depression using machine learning. *Computers, Materials & Continua*, 71(2), 2113–2129. <https://doi.org/10.32604/cmc.2022.024847>
- [17] Negi, H., Bhola, T., Pillai, M. S., & Kumar, D. (2018). A novel approach for depression detection using audio sentiment analysis. *International Journal of Information Systems & Management Science*, 1(1).

- [18] Alnaddaf, M. M., & Başarslan, M. S. (2024, September). Sentiment analysis using various machine learning techniques on depression review data. In *2024 8th International Artificial Intelligence and Data Processing Symposium (IDAP)* (pp. 1–5). IEEE.
- [19] Verma, B., Gupta, S., & Goel, L. (2020, February). A survey on sentiment analysis for depression detection. In *International Conference on Automation, Signal Processing, Instrumentation and Control* (pp. 13–24). Springer. https://doi.org/10.1007/978-981-15-8221-2_2
- [20] Peng, Z., Hu, Q., & Dang, J. (2019). Multi-kernel SVM based depression recognition using social media data. *International Journal of Machine Learning and Cybernetics*, 10, 43–57. <https://doi.org/10.1007/s13042-017-0771-5>
- [21] Ahmed, U., Jhaveri, R. H., Srivastava, G., & Lin, J. C. W. (2022). Explainable deep attention active learning for sentimental analytics of mental disorder. *ACM Transactions on Asian and Low-Resource Language Information Processing*. <https://doi.org/10.1145/3528264>
- [22] Suganya, M. P., Vijaiprabhu, G., Shanmugapriya, N., Praneesh, M., Menaka, M., & Sathishkumar, K. (2024). Depression classification using Harris Hawk Optimization (HHO) based recurrent fuzzy neural network (FRNN) for sentiment analysis: An applied nonlinear analysis. *Communications on Applied Nonlinear Analysis*, 31(2s), 234-252. <https://doi.org/10.52783/cana.v31.638>
- [23] Suganya, P., Vijaiprabhu, G., Sivakumar, G., & Sathishkumar, K. (2023). Navigating sentiment analysis horizons: A comprehensive survey on machine learning approaches for unstructured data in medical sciences and science and technology. *International Journal of Pharmacy Research & Technology (IJPRT)*, 14(1), 72–78.
- [24] Webber, J. L., Arafa, A., Mehbodniya, A., Karupusamy, S., Shah, B., Dahiya, A. K., & Kanani, P. (2023). An efficient intrusion detection framework for mitigating blackhole and sinkhole attacks in healthcare wireless sensor networks. *Computers and Electrical Engineering*, 111(Part B), 108964. <https://doi.org/10.1016/j.compeleceng.2023.108964>
- [25] Nour, A. A., Mehbodniya, A., Webber, J. L., Bostani, A., Shah, B., Ergashevich, B. Z., & Sathishkumar, K. (2023). Optimizing intrusion detection in industrial cyber-physical systems through transfer learning approaches. *Computers and Electrical Engineering*, 111(Part A), 108929. <https://doi.org/10.1016/j.compeleceng.2023.108929>
- [26] Kumar, B. P. S., Haq, M. A., Sreenivasulu, P., & et al. (2022). Fine-tuned convolutional neural network for different cardiac view classification. *The Journal of Supercomputing*. <https://doi.org/10.1007/s11227-022-04587-0>