

Implementation of Hybridized Meta-Heuristic Model for Feature Selection in Sentiment Analysis

Alok Kumar Jena¹, K Murali Gopal², Abinash Tripathy³, Sachikanta Dash⁴

¹*School of Computer Science and Engineering GIET University, Department of CSE, Siksha O Anusandhan (Deemed to be) University, Odisha, India, alok.jena@giet.edu*

²*School of Computer Science and Engineering GIET University, Odisha, India, kmgopal@giet.edu*

³*School of Computer Application KIIT (Deemed to be) University, Odisha, India, abinash.fca@kiit.ac.in*

⁴*Department of Computer Science & Engineering GIET University, Odisha, India, sachikant@giet.edu*

With the prevalence of social media and its ease of use, scrutinizing social media sentiment can provide effective direction to topics and products. Sentiment analysis is important for understanding the opinions of individuals published on platforms such as social media and product review blogs. Instead, sentiment analysis (SA) is gaining popularity and becoming a buzzword among researchers. Supervised machine learning is popular for its ability to simulate human behavior as an artificial intelligence. Various classification strategies have been developed for this purpose. ANN-based data classification has played an important role in classification. On the other hand, issues such as overfitting, pairwise classification, and parameter regularization affect classifier performance. To overcome this, a group of algorithms known as metaheuristic algorithms iteratively update candidate solutions and identify the best possible solution by maximizing the objective function. Genetic Algorithm (GA) and Firefly Algorithm (FA), which are used to optimize SVM and ANN parameters in this article, outperform SVM, ANN, and ANN. The IMDB film reviews the datasets created for the experiment. In order to properly analyze emotions, this study applied trajectory-based and population-based optimization methods and conventional machine learning methods, and found the relatively best results by comparing all the results.

Keywords: Sentiment Classification, Genetic Algorithm, Firefly optimization, Performance evaluation parameter, Meta-Heuristic algorithm.

1. Introduction

Almost all purchases are motivated by emotion. Analyzing consumer sentiment has become crucial, since emotions influence customer behavior in a significant way. A human being may

have different kinds of emotions like happiness and sadness, interest and disinterest, and positive and negative feelings. These sentiments can be analyzed with various models and can capture the emotions. The sentiment analysis can be grouped as per the categories of i. Analysis of fine-grained emotions ii. Sentiment analysis based on aspects iii. Discovery of emotions and iv. Analysis of intent.

The granular sentiment analysis approach can be used to find the accuracy of polarity. The following polarity scales can be used to achieve sentiment analysis: very good, good, neutral, bad, or very bad. For the analysis of public opinions in the form of reviews and ratings, a fine-grained analysis is helpful. In fine-grained sentiment analysis, “one to five” a range is used, where ‘one’ is considered as very poor and ‘five’ is very positive. The fine-grained sentiment analysis identifies the general polarity whereas the aspect-based analysis delves further. It aids in identifying the specific aspects that are being required. One can identify emotions with the use of emotion detection. This can include temper, depression, joy, delight, anxiety, worry, fear, etc. A list of phrases that express particular emotions is frequently used by emotion detection systems. Additionally, some sophisticated classifiers use powerful machine learning (ML) techniques. Companies can utilize the resources effectively by accurately determining consumer intent. Numerous people put their thoughts and opinions in an unorganized manner. Consumers, celebrities, sportsmen, social media users, weather forecasters, and trade groups all had a say in shaping the final verdict. On a daily basis, millions upon millions of individuals use social media platforms like as Facebook, Twitter, and RenRen. Social media platforms produce a veritable mountain of sentiment data in the form of tweet IDs, status changes, reviews, authors, content, and more. It is quite crucial to analyze and classify the real-world sentiment, as day by day online platforms are used vastly to keep their emotions on the real-time basis. As a result, these reviews serve as a source of knowledge for new consumers, producers, or sales managers. They get the chance to learn in-depth details about the product’s quality, which enables them to make the

best choice regarding whether to purchase, produce, or sell the goods. People also comment on the quality of movies, as is the case with books. In general, the reviews are in text format, so careful and efficient methods are required to extract the proper information. In this regard, the Sentiment analysis helps to analyze and classify the user reviews. The study of Fonseca and Fleming tells that the evolutionary algorithms play a major role in achieving the best solution for the challenging problems [1]. For better categorization, a number of methods have been extensively researched in the literature, including genetic algorithms, simulated annealing, and nature-inspired algorithms. For the multi-object feature selection and entity identification, Ekbal et al., implemented NSGA-II and able to get a better result in comparisons with the traditional systems [2]. Goffe et al., have done a research to avoid the confusion on local optima and global optima with the use of very popular ‘simulated annealing’ algorithm [3]. Sangita et al. studied the behavior of swarm intelligence algorithms in the robotic system and suggested the swarm robotic application in various fields of engineering [4]. Swarm intelligence is a distributed system in which social agents, who are not identified, interact with one another and their immediate environment in a way that leads to cooperative global behavior. In 2010, Xin-She Yang developed what is now known as the firefly algorithm, an algorithm inspired by nature and based on the way fireflies emit light [5]. Comparatively, this new population-based meta-heuristic algorithm performs better than the other swarm

intelligent algorithms. In a search space, each firefly provides one or more potential solution for the issue. The feature selection algorithm used by the suggested system is called firefly optimization. The Fig. 1 indicates the sentiment extraction process adopted in the present paper. This study focused on classification of movie reviews as per their sentiment, using different types of machine learning approaches. Polarity dataset was taken into consideration for classification. The study emphasizes more on the findings of the document-level sentiment analysis. The research connected to "GA and Firefly as feature selection" from the dataset. Section 1 lays out the process flow for this research. The literature review is explained in Section 2, the various methodologies are covered in Section 3, the application of machine learning techniques for classification and feature selection is briefly explained in Section 4, the proposed approach is presented in Section 5, the performance is evaluated in Section 6, the results and discussion are presented in Section 7, and finally, the conclusion and future work are explained in Section 8.

2. Literature Survey

In this section; the survey was prepared with four sections from Section 2.1 to Section 2.5 based on their categories. The section 2.1 explains the "Document-level sentiment classification", the Section 2.2 explains "Aspect-based sentiment analysis", the Section 2.3 explains "sentiment Classification utilizing hybrid machine learning approach" and Section 2.4

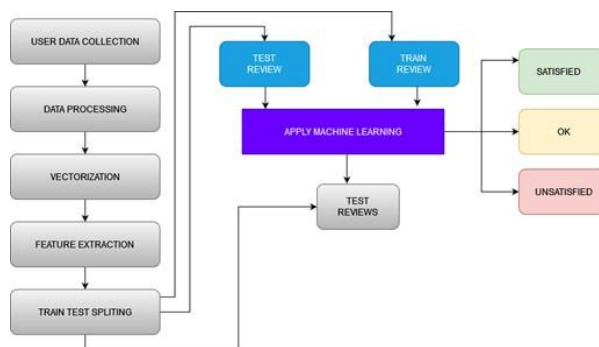


Fig. 1. Sentiment extraction process explains "GA as feature selection", finally, 2.5 discuss the motivation for the proposed approach.

A. Document-level sentiment classification

According to Chen et al., determining the polarity of a text requires giving different degrees of importance to each sentence [6]. To solve this problem, they present a model of document-level sentence categorization based on deep neural networks, where gate mechanisms determine the relevance levels of phrases in texts automatically. When it came time to classify documents, they determined that sentences should be taken into account. Using the N-gram method, Pang et al. [7] used three distinct ML algorithms: Support Vector Machines (SVM), Naive Bayes, and Maximum Entropy (ME). To categorize the reviews, use a bigram, a unigram, or both. The SVM approach has the highest level of classification accuracy of these techniques. The text has been divided into groups that are subjective and objective. After

taking into account subjective groupings, they used a variety of techniques based on the minimum cut formulation. By analyzing and extracting evaluation groups based on the expression of the adjectives found in the text [8]. Matsumoto et al., have studied syntactic relationships between words gives a great impact of document-level sentiment analysis [9]. Sentences are parsed to find out the word sequence and level of dependency. After which the SVM algorithm uses as features. For classification, they've combined the bigram and the unigram techniques. Read classified the polarity dataset using the NB and SVM algorithm. Then depending upon the emotions of the language, a source of training data prepared to find out the exact sentiment of the word at any instant. Moraes et al., compares the results of document-level sentiment analysis using SVM and ANN [10]. The feature selection method and the bag-of-words (BOW) model applied at a time with various weights included in the reviews. They demonstrate that the ANN-based outcome is better to the SVM result. According to the performance of word2vec and SVM perf, Zhang et al., suggested classifying Chinese comments [11]. Their suggested strategy consists of two elements. The word2vec tool is utilized in the first section. This tool gathers the reviews' semantic attributes in a few chosen fields then in the next section lexicon-based strategy is employed. On the topic of sentiment classification, Liu and Chen, have made several alternative multi-label classification proposals [12]. Eight separate evaluation matrices are utilized to compare the results of the multi-label classification, which was performed using eleven multilevel classification methods. Again, they take into account three separate sentiment dictionaries when classifying. In order to summarize the huge amount of text data into low dimensional space, Luo et al., have presented a method where each word has a clear, consistent meaning for all time the text have been considered for emotional analysis [13]. According to Ekman and Friesen, six different sorts of human emotions are classified [14]. Niu et al. [15] suggested a new dataset for emotion analysis that incorporates Twitter-sourced image-text pairs annotated by humans. Two main schools of thought exist within their methodology: statistical learning and lexicon-based learning. They have considered the words or phrases with the specified sentiment score for lexicon-based analysis. On the other hand, a wide variety of textual components and a plethora of machine learning techniques are considered in statistical learning. Tripathy et al. [16] demonstrated sentiment analysis on the IMDb dataset using four distinct ML algorithms and n-gram methods like unigram, bigram, and trigram. They discovered that SVM with a bigram+unigram strategy produces results which is found out to be best.

B. Aspect-based sentiment analysis

Shan et al., 2022, proposed aspect category sentiment classification (ACSC) to find the polarity of sentiment in sentences under specific aspect [17]. The majority of models fail to completely understand the good interactions between words associated with certain aspects by ignoring syntactic information. To address the issues a BiGAT model for ACSC has proposed, in two graphs he expressed the context information and phrase grammatical structure. Feng et al., 2022, suggested an Aspect-based sentiment analysis (ABSA) which is a specific job that identifies the polarities of the sentiment for some features in a sentence [18]. The majority of modern ABSA models use graph-based techniques due to the development of graph convolution networks (GCNs). These techniques treat each word as a separate node in a dependency tree that is built for each sentence. They classify according to the aspect rather than the representation of the sentence and used it with GCNs. Kit and Mokji, 2022, proposed

Language task improvement for successfully using pre-trained language models like BERT and mBERT [19]. Majorly of the work involved in putting the models into practice is spent adjusting the BERT to get the desired outcomes. However, depending on the dataset, this strategy is resource-intensive and necessitates a lengthy training period. To minimize phrase embedding, he did not consider the fine-tuning of the features and also not focus on the reduction of features in the BERT model. As per Bie et al., suggestion, aspect sentiment prediction and aspect term extraction are the two components of aspect-based sentiment analysis (ABSA) [20]. When handling the sub-tasks in a pipeline fashion, which is how most approaches carry out the ABSA task, performance issues and issues with actual application start to show up. To overcome the drawback that current approaches do not completely use the text information, they presented a back-to-back ABSA model in their study called SSi-LSi. Fei et al., 2021, has inspired by fine-grained sentiment analysis [21]. They considered two aspects i.e., they newly created the document with user's data and the product's information to the existing document in a variation module and then a module for the classification. Rao et al., has proposed a model of ANN with two hidden layers [22]. It is a two-step process. Initially in the first layer extract the sentiment whereas the second layer finds the level of dependency on a sentence. It is a model of a neural network that can identify the semantic relationships between words and sentences. To carry out the document-level sentiment analysis tasks, they have used a mixture of review datasets. Liu et al., have developed the AttDR-2DCNN, a unique neural network model that primarily comprises two components [23]. The first layer extracts the sentence's feature vector and in the second layer it learns the two-dimensional matrix representations. They used two-dimensional convolution and max pooling, to find out the level of dependency of a sentence with word.

C. Sentiment classification with hybrid machine learning approach

In Section 2.1 discussed about document-level sentiment classification using various machine learning algorithms. Few authors have combined two or more machine learning algorithms, known as hybrid approaches, to do the classification. This section discusses studies that have been published in literature and have taken the hybridization approach into account. Salur and Aydin introduced a model with several word embedding techniques and various deep learning techniques to collect characteristics from several word embedding deep learning techniques [24]. The model provides improved sentiment classification. In their study, Hassan et al. [25] looked at the best ways to extract emotions from text using machine learning and natural language processing on various social media sites in order to determine a person's depression level. In order to uncover the hidden semantic information, Du et al. [26] suggested a hybrid neural network model that relies on capsules. To accomplish features that are dependent on one another over long distances, this model employs a bidirectional gated recurrent unit (BGRU). Additionally, it can extract more detailed textual information to enhance expression. The effectiveness is assessed using two small text review datasets. According to the author, the capsule-based model has the maximum accuracy and beats other comparable models using movie review data. Three different classifications were used by Balage Filho et al., to arrive at the categorization outcome [27]. Then they have applied the aforementioned strategies sequentially, one after the other. The results of one strategy are fed into the second strategy, and so on. Classification and expression are the two major sections in the suggested strategy. An innovative hybrid strategy for categorizing reviews has been put out by Wang et al. [28].

They used the SVM technique to classify the reviews after taking into account the words' capacity for category distinction and the information-gain approach. Tripathy et al., performed a hybrid document-sentiment analysis on the movie review dataset [29]. Where they combined SVM and ANN. After removing stop words and completing stemming operations, they employed SVM to determine the emotion values of each word in the lexicon. They then determined the sentiment values for every word then they considered those words are having a sentiment value greater than 0.009. The words that meet the requirement are then fed into an ANN for additional classification. They have altered the hidden nodes that are part of the ANN's hidden layer in order to do further analysis and determine which hidden node their suggested strategy performs best. For Amazon and Flipkart, Dadhich et al. have put up the idea of automatically identifying the sentiment, in English text using Random Forest and the K Nearest Neighbor technique [30]. Based on five important criteria, they gave a comparison of the available sentiment analysis algorithms. They prioritize sentiment polarity and feature reduction when doing their analysis. To improve the efficacy and accuracy of sentiment analysis, AlBadani et al., 2022 proposed a new method that combines SVM with "universal language model fine-tuning" (ULMFiT) [31]. He suggests a novel deep-learning approach to Twitter sentiment analysis in order to infer people's opinions about specific products from their comments. In the train-test model, they primarily concentrated on accurate and complex classification for this large dataset. Sentiment analysis of movie reviews on IMDb allows Kumar et al. to discover the overarching opinion or sentiment expressed by reviewers on a film [32]. Many researchers explore the research work to visibly distinguish between favorable and unfavorable reviews in sentiment analysis. In comparison to traditional classifiers, they demonstrate that the use of hybrid features, which combines both the machine learning features with lexical features and produces better results. Ahuja et al., examined the SS- Tweet dataset using TF-IDF, N-Gram, and some feature extraction techniques [33]. They utilized six classification algorithms and concluded that TF-IDF performing better than N-gram. For text sentiment categorization, Onan et al., worked on a hybrid ensemble pruning approach. They used the clustering and randomized search in their study [34].

D. Motivation for the proposed approach

As per the study of Xu et al., in the field of sentiment analysis, the use of hybridized machine learning techniques up to the year 2022 is not that much significant [35]. As Fig. 2 shows that sentiment analysis with some traditional machine learning algorithm is 55 percent whereas the use of hybridized technique in sentiment analysis is only 10 percent. So, it provides a scope to perform a thorough study of sentiment analysis with optimized machine learning techniques. The literature review that was discussed above provides insight into several potential areas for further study. In order to conduct further research, the following factors have been taken into account. When selecting features, many authors take into account the frequency of a term in reviews and part-of-speech (POS) tags. The encoding process that assigns a piece of text (or corpus) to a specific part of speech is called POS (part of speech), "grammar encoding," or "word class." It does this by analyzing the text's description and the context of the surrounding and associated sentences, words, and paragraphs. For the POS to be as accurate as possible and to produce classification models that make sense, it is a good idea to tag the grammar structures of each word. Simple past verbs rather than past adjectives are frequently used in subjective literature. Negative sets, in contrast to positive sets, frequently use past tense verbs

since different authors have varied ways of expressing negative emotions. But a word's POS tag isn't steady; it might change depending on the context in which it's used. For instance, the verb "honor" can have the POS "noun" when used as "It is a great honor", whereas the POS verb is used as the honor of the request. The use of a term in a text depends on the author; some prefer to use a word as many times as possible, while others may only use it once. Therefore, it's never certain whether these should be taken into account when employing feature selection. A specific machine learning technique has been taken into consideration by some researchers for the selection of features and classification. In that situation, the bias influences the end result to shift in favor of a specific polarity value. In this study, machine learning utilizes two alternative processing methodologies, for the analysis and to reduce the influence of the bias of a particular machine learning technique. Here, the features are selected with the help of the FireFly algorithm whereas the SVM helps to classify them. Since most movie reviews are text-based, they must be converted into a matrix of numbers before being evaluated by any machine learning method. TF-IDF and CV are the two most often used functions for this transformation task. However, because TF-IDF also takes word frequency into account, the majority of authors prefer it. At the initial stage, each word of the review is considered for the analysis. Thus, the feature set is discovered to be very large in size and that increases the processing time when all reviews are taken into account for analysis. In this situation, a feature selection technique is used to choose the best features from the entire dataset. Subsequent processing is then done based on these features. The best features are gathered using FireFly from a larger feature list in this manuscript and sent into an SVM for classification.

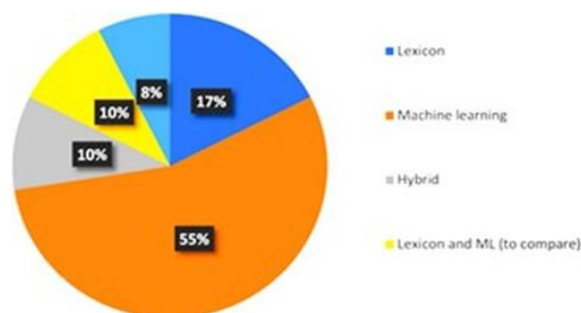


Fig. 2. Sentiment Analysis with Traditional ML

3. Methodology

A. Types of Sentiment Analysis

Depending on their characteristics, the reviews under evaluation can be divided into a variety of classes. The different class of problems is explained as follows: Binary Classification: It divides the reviews into positive and negative categories. The reviewers employ this strategy when they require a result in a yes-or-no manner, i.e. it explains whether the reviewer is favouring or not favouring the product/services. Multi-Class Categorization: This strategy is used for ranking the product, where the reviews may be divided into more than two classes. One way to categorize movies is into "bad, average, fine, good, and great" categories. The

current manuscript prefers the binary classification strategy to categorizing movie reviews based on their quality.

B. Data Set

Any dataset must be used to assess the effectiveness of a machine learning strategy to determine whether the suggested approach yields a result that is generally considered to be of high quality. Therefore, one of the widely used datasets, the “IMDB”.

TABLE I DATA SET DETAILS

Total Number of Reviews	Number of Positive Reviews	Number of Negative Reviews
50000	2500	2500

The polarity” movie review dataset, is analyzed with this current work. The details of the IMDB dataset is shown in Table I. A conventional method of training, the training dataset and then testing, the testing dataset using the information gleaned from it will not be effective in processing the 1000 reviews that have been classified both negatively and positively. To analyze the reviews, the current manuscript takes into account a second, relatively new method called cross-validation. In a situation when there is no separation between the training and test datasets, the cross-validation strategy is recommended over the conventional one. In this method, the K-fold methodology is employed, where (100 - K) percentage of reviews are taken into account for training purposes and the remaining K percentage are taken into account for testing depending on the knowledge learned during training. The procedure is repeated several times to increase precision. In this publication, 10-fold cross-validation is taken into account for analysis, with 90 percentage of the reviews taken into account being utilized for training and the remaining 10 percentage being taken into account for testing. To get a somewhat better result, this technique is repeated ten times.

Fig. 3 shows the preprocessing steps for the sentiment analysis process.

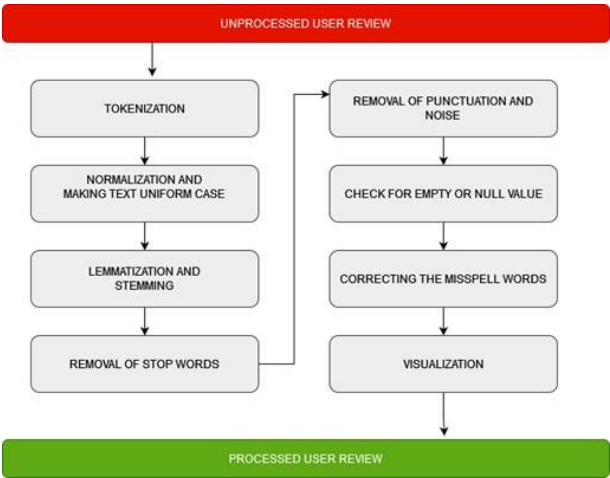


Fig. 3. Pre-processing Steps

- Tokenization: Tokenization implies, breaking down the text into tokens before converting it into vectors. Filtering out useless tokens is also made simpler. Take a document *Nanotechnology Perceptions* Vol. 20 No. S6 (2024)

for instance, and break it up into paragraphs or sentences. We're turning the reviews into words in this instance.

- **Normalization and Make Text Uniform case:** Words that seem different because of the case or may have been spelled differently but have the same meaning must undergo correct transformation. These terms are treated consistently because of normalization processes. Like making everything uppercase or converting numbers to their word forms. All reviews will be converted to lowercase during this pre-processing stage.
- **Lemmatization and Stemming:** Lemmatization is a technique for counting the occurrences of each word's natural form without considering the grammar tense of each word. A word's base or dictionary form can be restored by the process of lemmatization, which involves precisely carrying out responsibilities using a vocabulary and morphological examination of words to delete all inflectional endings. Another way to restore words to their original form is by stemming. Although stemming extracts a word's linguistic root, lemmatization returns a word to its original lemma. For instance, lemmatization yields "appl" when we conduct stemming on the word "apples," whereas stemming produces "apple". Lemmatization is used instead of stemming since it is much simpler to understand.
- **Removal of Stop Words:** The most frequent words that don't make sense in the context of the facts and don't give the statement a deeper meaning should be eliminated from the text are stop words, or words lacking emotion in this situation.
- **Removal of punctuation and noise:** The only remaining characters are alphabetic when stand-alone punctuation marks, special characters, and number tokens are dropped because they don't add to the sentiment. Tokenized words are required to go with this stage since they have been suitably processed for removal. The text may have mark-up's or wrappers in HTML or XML, for instance, if it was scraped from the internet. With the help of regular expressions, these can be eliminated. Since we were able to retrieve the specific review from the XML file, fortunately our reviews do not require this step.
- **Check for empty/ null values:** Since empty values are not filled with nulls during the cleaning process, there were no null values discovered. Both were examined, and empty values were eliminated. The words "10," "LOL," and "Why?" were found in the text review and all of these had been turned into null values throughout the cleaning procedure. They didn't add anything to the sentiment analysis, so eliminated them.
- **Correcting Misspell words:** Many spelling errors are not amusing or harmless, yet some are. Not only can a simple misspelling make us seem less clever than we actually are, but it may also be embarrassing. Additionally, incorrect spelling might result in a lack of understanding and clarity.
- **Explore with word cloud:** Word clouds have become a popular and eye-catching way for text visualization. By reducing text to the words that occur most frequently, they are utilized in a variety of contexts as a way to give an overview. However, the majority of systems employ them very statically and only offer a few limited inter- action options. The Word Cloud Explorer looks into the use of word clouds as the core of text analysis. It just uses word clouds as a visualization technique. It provides them with sophisticated natural language processing and con- text understanding. To facilitate the visual text analysis, flexible filter mechanisms

consider the polarity for feature selection. Generally, sentiment analysis methods consist of two steps: i. Determining the part of the document that will provide either positive or negative emotions and ii. Combine these sections in a way that will excel the likelihood that the document will fall into either of this polarity. By analyzing terms inclusion and exclusion in each polarity class, feature selection assigns terms a ranking. Terms that come up frequently in class receives a high score.

A. Support Vector Machine

The SVM is an effective classification method that also shows as an example of supervised learning based on the idea of the lowest structural risk. The algorithm will develop a new hyperplane during training to categorize the data as positive or negative. The point on the hyperplane where each sample must be placed is then specified, and the new samples are then categorized. The SVM is a supervised machine learning technique that analyses the data and can further recognize the patterns worked in the classification. Due to the high computing costs associated with developing SVM classifiers, which require solving a quadratic optimization problem, their application is challenging when working with large datasets. Even though the SVM may create a hyperplane to separate two distinct classes prominently but the major issue is that it uses a quadratic programming technique even for a medium- sized dataset. Although there are issues that make the application of the SVM algorithm challenging, it is extremely generalizable. These problems stanch from the selection of kernel function's type and parameters. The accuracy of the data classification is largely dependent on the selection of these parameters' values. The SVM was developed with the goal of enhancing generalization. It is based on the principle of statistical learning. Using the two class labels -1 and 1, it predicts newer cases using training examples. In the simplest example, by scanning the grid, the SVM classifier's optimal parameter values can be found. This process, however, is time- consuming and does not guarantee a high-quality classification. Using evolutionary optimization algorithms, it is possible to find the best values for the regularization parameter and the kernel function parameters for an SVM classifier with a fixed kind of kernel function i.e. genetic algorithm.

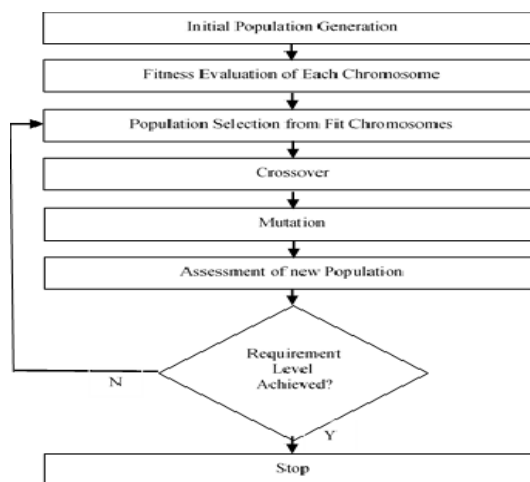


Fig. 6. Genetic Algorithm flowchart

Step 4. Different distance metrics, including Euclidean, Manhattan, cosine, etc., can be employed to determine the closeness in between objects and the kNN.

B. k-Nearest Neighbors

Usually, the k-Nearest Neighbors (KNN) classification and

$$\text{Score}(x, p_i) = \sum_{x_j \in \text{KNN}(x)} \text{Cal}(x, x_j) \delta(x_j, c_i) \quad (1)$$

regression algorithms work on either a Memory-based or instance-based learning approach although it also depends on the training document's class levels, which is similar to the test document. It implies that it is unable to produce a proper classification. A neighbor's similarity-based vote categorizes an instance, with the instance being assigned to the class with high popularity among its k closest neighbors (where k = +ve number). When the value of k is equals to 1 then it simply returns the class of its closest neighbor. The KNN algorithm locates the k training documents' nearest neighbors given a test document d. The weight of the classes in each closest neighbor document in comparison to the test document depends on how similar each closest neighbor document and the test document are to one another. The kNN algorithm typically includes the following steps in order to determine the association class of any object in a certain number of k-nearest-neighbors.

Step 1. Determine the distance $d(z, z_i)$ between each item z_i (the known class) and object z , and arrange the calculated distances according to their increasing values.

Step 2. Choose the k closest objects z_i (k-nearest-neighbors) to object z .

Step 3. Identify the class to which each of object z 's k closest neighbors belongs. The class that is more prevalent among object z 's k nearest neighbors is designated as its association class.

where Eq. (1) is used to calculate the closeness between objects and the kNN.

Step 5. The document set of k documents that are document d's closest neighbors is referred to as KNN (x). If x_j and c_i are related, then (x_j, c_i) equals 1, else 0. The class with the largest weighted sum should include a test document (x).

Step 6. The kNN classifier, which is the most basic metric classifier, uses a voting system to determine how similar a given object is to its nearest neighbors, a total of k objects. In some situations, the kNN classifier can boost the overall quality of data categorization while taking a little bit longer to construct a hybrid classifier overall.

C. Genetic Algorithm

Booker et al., has proposed GA in 1980 [36]. GA is based on the principle of Charles Drawin's natural evaluation. This approach finds out fit chromosomes, and allow those for the reproduction of the next generation whereas the unfit chromosomes goes through operations of GA like crossover and mutation to transform their phase from unfit to fit. The activities performed in GA can be summarized using the following flowchart in Fig. 6. In this work, the

kNN is used to find out the fitness of the chromosomes. If they found out to be fit, they move forward for further analysis. But the unfit chromosomes again go through the crossover and mutation process, and then again the fitness is checked. This process is continued until the required level is achieved.

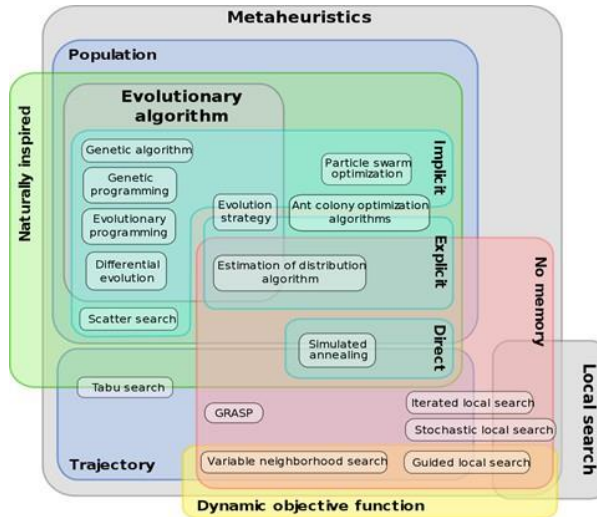


Fig. 7. Type of metaheuristic algorithms

D. Meta-heuristic optimization

The simplest definition of an optimization is a minimization or maximization problem. For nonlinear, multi-modal, multi-variate functions, this is not an easy operation. Moreover, some functions could have discontinuities in nature, which have a major impact on obtaining derivative information. Several conventional methods, such as climbing hills, may encounter numerous challenges. Using meta-heuristic algorithms, meta-heuristic optimization tackles optimization issues. From economics to engineering design, from vacation planning to Internet routing, optimization is practically everywhere. The most effective use of the available resources is crucial because time, money, and resources are all finite. Most real-world optimizations are multi-modal and highly nonlinear while being subject to a variety of intricate constraints. Various goals frequently clash with one another. Sometimes there may not even be an optimal solution for a single objective. Finding the best solution, or even sub-optimal solutions, is not always simple. Some well-known meta-heuristic algorithms are displayed in the Fig. 7, out of which, this study aims to implement the Firefly swarm intelligence meta-heuristic optimization.

E. Firefly Algorithms (FF)

The social interactions of fireflies or lightning bugs in the tropical summer sky are the focus of one meta-heuristic optimization technique called firefly algorithms. At Cambridge University, Slowik in 2020, conducted research on the behavior of fireflies and created a brand-new algorithm known as firefly algorithms in 2007 [37]. It takes its cues from the way in which birds, insects, and fish all tend to cluster together. Firefly is very similar to other SI algorithms, such as PSO, Bacterial Foraging, and Artificial Bee Colony optimization. The

Firefly method makes use of genuine random numbers. The principle behind it is the interplanetary communication of the swarming particles, which are represented by firefly. It appears to be more effective when optimized. The following three rules of the firefly algorithm are designed on the basis of the flashing light of actual fireflies.

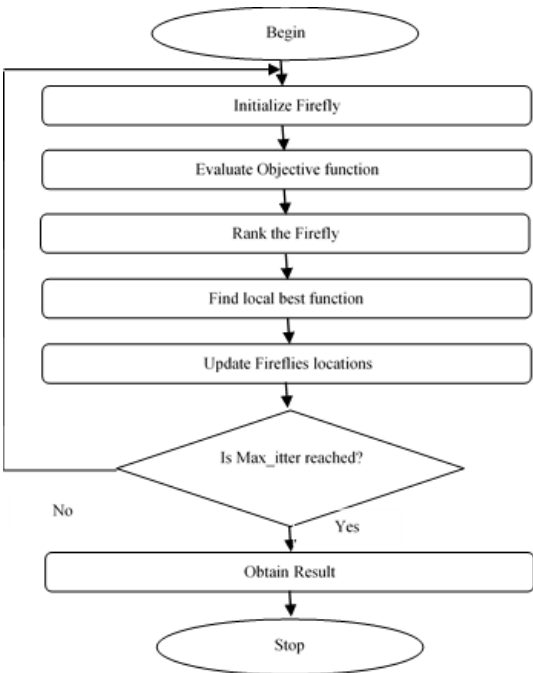


Fig. 8. FireFly Algorithm Flowchart

- Regardless of their sex, fireflies are universally attracted towards brighter or more attractive ones.
- A firefly's level of attraction is related to its brightness, which decreases as it gets farther away from another firefly because the air absorbs the light. It will then move randomly if there are no fireflies that are brighter or attractive.
- A firefly's illumination is dependent on the value of the objective function for a certain problem.

The application of firefly algorithms can be found in a variety of contexts. The following categories include networking, image processing, bench-marking, and optimization. We are aware that queuing theory is employed to examine intricate service systems in manufacturing, networks, and transportation. In this paper, the effectiveness of the firefly method with mutation in solving optimization issues is evaluated. The following Fig. 8 discusses the steps involved in the firefly algorithm.

F. Artificial Neural Network

In addition to GA and Firefly, ANN is also employed in this study. The following is a description of this algorithm's mapping function:

$$\text{Cal} : N^{\text{in}} \rightarrow N^{\text{out}} \quad (2)$$

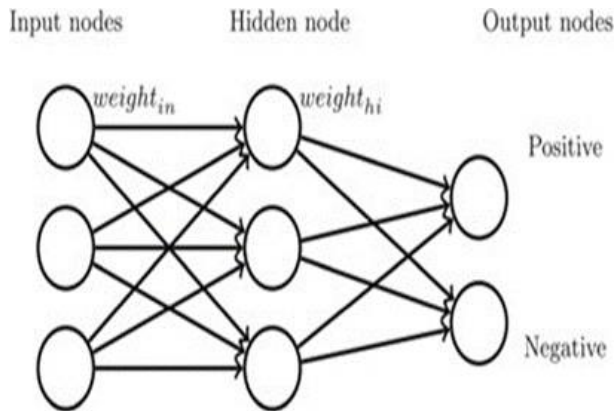


Fig. 9. ANN Architecture

where “out” indicates the output classes of classification and “in” indicates the input dimension or input nodes to the networks as shown in Eq. (2).

The fundamental arrangement of an ANN is represented in Fig. 9. The three major layers in the ANN architecture are the input layer, hidden layer, and output layer. The fit chromosomes acquired by GA are considered to be the ‘in’ number of independent nodes that the ANN considers to make up the input layer in the current publication. Users may choose to have one, two, three, or more levels of hidden nodes; however, in this study, only one layer of hidden nodes is taken into consideration for analysis. In this work, the independent variables are getting processed in the hidden layers and then it comes to the output layer, where it is processed with the dependent variables.

5. Proposed Approach

The proposed method is carried out using the following three phases. In the first phase, the subset of effective features is selected from all the preprocessed features using the conventional machine learning i.e. SVM and KNN approach. Whereas in the second phase, the optimal features set have been selected using the Firefly meta-heuristic algorithm from the selected subset of features in the previous phase. Further in the third phase, the final selected features are classified using the ANN classifier. The two-phase classifiers are used to avoid the misclassification as well as to improve the classification accuracy. On the other side, practically it is quite difficult to analyze such a huge amount of preprocessed data, in the second phase, it tries to select only a few optimal features, which are classified in the next phase of processing. The following Fig. 10 discusses about the operational flowchart of the proposed approach.

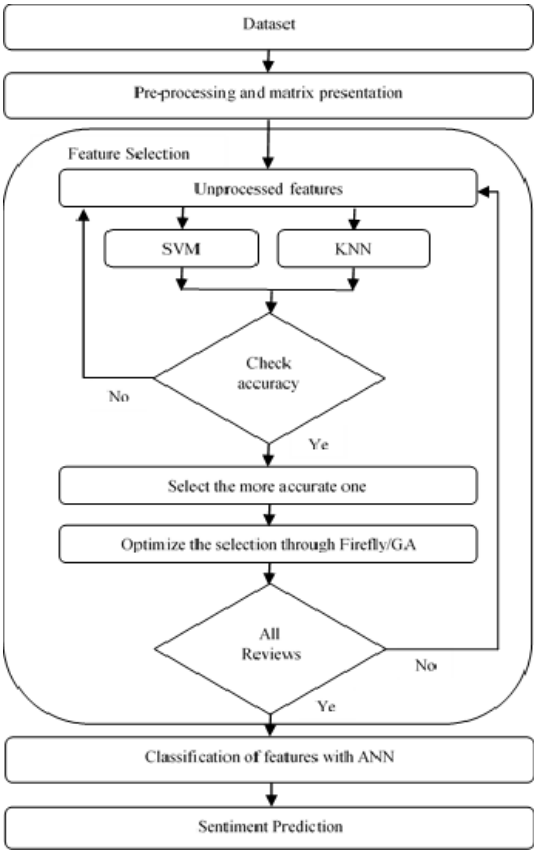


Fig. 10. Operational Flowchart of proposed approach

released varies with firefly proximity. Two important elements that affect the behavior of the firefly algorithm during processing are the intensity of the light and the attractiveness of the fireflies (Wang et al., 2016) [38]. A fitness function-like expression and measurement of the firefly's brightness controls the amount of light that reaches each location. Each firefly's attraction can be calculated using Eq. (3).

$$\beta(r) = \beta e^{-\gamma r^2} \quad (3)$$

the amount of light absorbed by the air is represented by c , and the attraction at a distance (r) = 0, widely believed to be 1, is indicated by β_0 . The value of r indicates the distance between fireflies i and j . The space between two fireflies is a representation of the attraction between them. Throughout the day, fireflies are always on the go. Since the attractiveness of fireflies is proportional to their distance from one another, we can use Eq. (4) to determine the distance between the two insects.

A. Firefly based ANN optimized classification

The behavior of the actual firefly insect served as the model for the firefly algorithm. Through the rhyme-flashing lights that emanate from their bodies, they communicate and convey

information to one another. The emitted light is utilized to draw in more people. The amount of light or intensity that is

$$r_{ij} = \frac{1}{\sum_{k=1}^d (x_{i,k} - x_{j,k})^2} \quad (4)$$

Here it indicates that, a Firefly located in “i” location having kth component $x_{i,k}$ and the dimension of the problem is d. The attraction between two Firefly increases when they will be closer to each other. Let’s say the firefly i is less bright than the firefly j. Type of association is controlled as per the Eq. (5).

$$x_i^{t+1} = x_i^t + \beta \frac{e^{\gamma r_{ij}}}{0} * (x_j^t - x_i^t) + \alpha * (rand - 1/2) \quad (5)$$

Where “t” denotes number of iterations, “a” denotes random value dictating the size of the random walk and “rand” denotes a random number generator. After weighing three factors, the low-brightness firefly flies to the higher one. The first term refers to the low-brightness firefly’s current location. By using the attraction coefficient b, the second term describes the movement toward the firefly with increasing brightness. The last component is an alpha-multiplied version of a random walk that a random number generator has calculated. Making use of firefly behavior, the proposed approach for firefly classification aims to build a classifier and feature selection tool [39] (Mashhour et al., 2018). The suggested classification algorithm has three distinct steps, which are as follows:

i. Selecting features ii. Building the model iii. Model application and prediction In this study, a hybrid strategy developed that combining the Firefly Algorithm (FFA) and Artificial Neural Network (ANN) to predict the opinion. We use the Firefly Algorithm (FFA) to choose the best ANN parameters. The primary goal of the study is to predict the sentiment from the user reviews using the combined method (ANN-FFA). The ANN-FFA follows procedure as per Table II.

TABLE II. PSEUDOCODE FOR ANN-FFA

Step 1. Initialize the algorithm’s parameters, including alpha, beta, gamma, and delta as well as the number of fireflies and the maximum number of generations, and provide the range of values for the SVM parameters.
 Step 2. Evaluate the ANN classifier’s performance using the firefly-derived parameters. The firefly’s optimal position’s light intensity will be used to estimate performance.
 Step 3. Arrange the fireflies in ascending order of light intensity. Step
 4. Repeat step 2 and 3 until it will evaluate for all the values. Step 5.
 Move every firefly to a better location.
 Step 6. Multiply the randomness reduction parameter delta by the randomness parameter alpha to reduce its value.
 Step 7. The optimum parameters generated by the Firefly algorithm is used in the SVM classifier for training and then find out the training accuracy.

6. Performance Evaluation

A. Confusion Matrix

The confusion matrix shows the accuracy of the method and compared with the classification successes from the definitions of the acronym’s TP (True Positive), FP (False Positive), FN (False Negative), and TN (True Negative) in confusion matrix and as follows: i. TP: The percentage of samples for which both the correct class label and the anticipated class label are positive. ii. FP: The percentage of samples predicted class label is positive but the actual class label is negative. iii. FN: The percentage of samples predicted class label is negative but the actual class label is positive. iv. TN: The percentage of samples for which both the correct class label and the anticipated class labels are negative. sentiment classification in this study. It determines measurement values for accuracy, precision, and F1. The following Fig. 11 contains the basic confusion matrix for a two-class classifier.

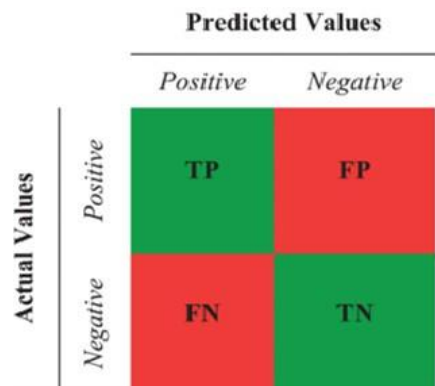


Fig. 11. Confusion matrix

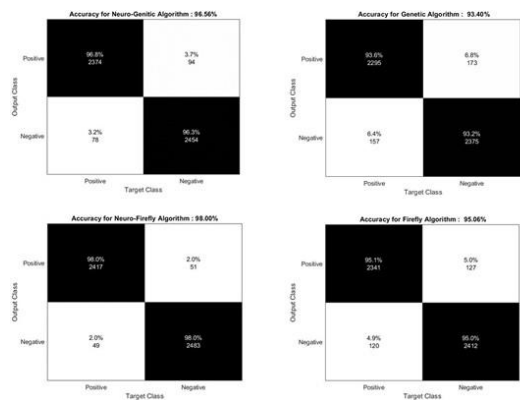


Fig. 12. Confusion Matrix for Neuro-GA, GA, Neuro-FA, FA

The confusion matrix and accuracy for Neuro-Genetic, Genetic (SVM+KNN), Firefly (SVM+KNN) and Neuro-Firefly has been shown for the comparison of results, where it depicts the proposed Neuro-Firefly model is performing comparatively better for the prediction of reviews using Fig. 12.

The error histogram is the histogram of the discrepancies between the target values and the predicted values after a neural network has been trained. These error numbers, which show how the anticipated values and the goal values diverge, may be negative. For the proposed model Fig. 13 shows the training, validation, test and zero error for 20 bins which shows the model obtained low error rate than the other conventional approaches for all iterations.

Using the values obtained from the confusion matrix different performance evaluation parameters are found out, which depict the quality and overall performance of the proposed approach. The parameters are as follows: Precision: The overall estimate of correctly predicted class labels for each class is known as precision. The Eq. (6) is used to compute the precision measure.

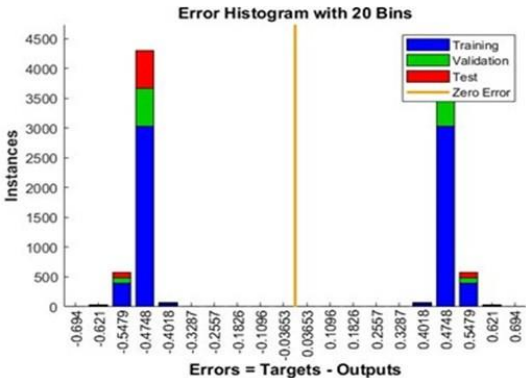


Fig. 13. Error Histogram Bar Graph

$$Precision = \frac{TP}{(TP + FP)} \quad (6)$$

TABLE III. CONFUSION MATRIX ALONG WITH ACCURACY OF PROPOSED APPROACH

Confusion Matrix			Methodology	Accuracy
	Original Labels			
	Positive	Negative		
Positive	2295	137	GA	93.40
Negative	157	2375		
	Original Labels			
	Positive	Negative		
Positive	2375	94	ANN-GA	96.56
Negative	78	2454		
	Original Labels			
	Positive	Negative		
Positive	1341	127	FA	95.06
Negative	120	2412		
	Original Labels			
	Positive	Negative		
Positive	2417	51	ANN-FA	98.0
Negative	49	2483		

Recall: It is calculated as the weighted average of the correctly assigned labels for each class. According to Eq. (7) recall is determined.

$$Recall = \frac{TP}{(TP + FN)} \quad (7)$$

F-Score: To combine precision and recall values into a single measurement, other measures, such as F1, are utilized. If the classifier properly identifies every sample, it assigns the value 1 to this measurement, which ranges from 0 to 1. Eq.(8) provides the F1 measure, and a successful classification is indicated by an F1 value that is close to 1.

$$F1Score = 2 * \frac{Precision * Recall}{(Precision + Recall)} \quad (8)$$

Accuracy: The overall result of the classification is generally indicated by Accuracy. The following Eq. (9) shows the way to calculate the accuracy value.

$$Accuracy = \frac{TP + TN}{(TP + TN + FP + FN)} \quad (9)$$

Table III provides the details of the classification result obtained using the proposed approaches. Testing samples were successfully assigned to the proper classes using a classification methodology built on the Firefly algorithm. The classification procedure was implemented based on firefly distance, firefly intensity, and average intensity. Distance-based classification proved to be the most effective technique for raising classification accuracy. The firefly technique is an effective optimization procedure that considerably raises the accuracy of the model. The evolutionary genetic algorithm model is outperformed by the firefly model of feature extraction, which is less efficient than the baseline model.

7. Conclusion and Future Work

The application of social media has expanded as a result of a shift toward them serving as personal communication and blocking tools in the entertainment sector. These applications have received a lot of user feedback. These days, NLP and hybridized machine learning are essential for uncovering the user sentiments of these huge datasets. In this study, it provides a new hybridized machine learning model which is suggesting for establishing a strategic interaction with the representation of text format data. The suggested method is predicated on the notion that diverse text representation strategies pair well with various machine-learning models. For a higher classification accuracy rate in light of the presented strategy, we recommend for the combining of several text representation techniques. Machine learning is used to resolve a variety of problems relating to the classification of the phrase. The FireFly algorithm is used to determine the parameters that will make for the best predictions, and it can also be used to solve the local optimum issue. Future sentiment analysis challenges involving aspect-based opinion classification and cross-domain sentiment analysis will make use of our approach. The effectiveness of our approach will be enhanced through the investigation of neural networks.

References

1. C. M. Fonseca and P. J. Fleming, "An Overview of Evolutionary Algorithms in Multiobjective Optimization," *Evolutionary Computation*, vol. 3, no. 1, pp. 1–16, Mar. 1995, doi: 10.1162/evco.1995.3.1.1.
2. A. Ekbal, S. Saha, and C. S. Garbe, "Feature Selection Using Multiobjective Optimization for Named Entity Recognition," in *2010 20th International Conference on Pattern Recognition*, Istanbul, Turkey: IEEE, Aug. 2010, pp. 1937–1940. doi: 10.1109/ICPR.2010.477.
3. W. L. Goffe, G. D. Ferrier, and J. Rogers, "Global optimization of statistical functions with simulated annealing," *Journal of Econometrics*, vol. 60, no. 1–2, pp. 65–99, Jan. 1994, doi: 10.1016/0304-4076(94)90038-8.
4. ECE Department, Narula Institute of Technology, WBUT, India, S. Roy, S. Biswas, S. Sinha Chaudhuri, and ETCE Department, Jadavpur University, India, "Nature-Inspired Swarm Intelligence and Its Applications," *IJMECS*, vol. 6, no. 12, pp. 55–65, Dec. 2014, doi: 10.5815/ijmeecs.2014.12.08.
5. X. S. Yang, "Firefly algorithm, stochastic test functions and design optimisation," *IJBIC*, vol. 2, no. 2, p. 78, 2010, doi: 10.1504/IJBIC.2010.032124.
6. X. Chen et al., "Aspect Sentiment Classification with Document-level Sentiment Preference Modeling," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Online: Association for Computational Linguistics, 2020, pp. 3667–3677. doi: 10.18653/v1/2020.acl-main.338.
7. B. Pang, L. Lee, and S. Vaithyanathan, "Thumbs up? Sentiment Classification using Machine Learning Techniques," 2002, doi: 10.48550/ARXIV.CS/0205070.
8. B. Pang and L. Lee, "A Sentimental Education: Sentiment Analysis Using Subjectivity Summarization Based on Minimum Cuts," 2004, doi: 10.48550/ARXIV.CS/0409058.
9. S. Matsumoto, H. Takamura, and M. Okumura, "Sentiment Classification Using Word Subsequences and Dependency Sub-trees," in *Advances in Knowledge Discovery and Data Mining*, vol. 3518, T. B. Ho, D. Cheung, and H. Liu, Eds., in *Lecture Notes in Computer Science*, vol. 3518, Berlin, Heidelberg: Springer Berlin Heidelberg, 2005, pp. 301–311. doi: 10.1007/11430919_37.
10. R. Moraes, J. F. Valiati, and W. P. Gavião Neto, "Document-level sentiment classification: An empirical comparison between SVM and ANN," *Expert Systems with Applications*, vol. 40, no. 2, pp. 621–633, Feb. 2013, doi: 10.1016/j.eswa.2012.07.059.
11. D. Zhang, H. Xu, Z. Su, and Y. Xu, "Chinese comments sentiment classification based on word2vec and SVMperf," *Expert Systems with Applications*, vol. 42, no. 4, pp. 1857–1863, Mar. 2015, doi: 10.1016/j.eswa.2014.09.011.
12. S. M. Liu and J.-H. Chen, "A multi-label classification based approach for sentiment classification," *Expert Systems with Applications*, vol. 42, no. 3, pp. 1083–1093, Feb. 2015, doi: 10.1016/j.eswa.2014.08.036.
13. B. Luo, J. Zeng, and J. Duan, "Emotion space model for classifying opinions in stock message board," *Expert Systems with Applications*, vol. 44, pp. 138–146, Feb. 2016, doi: 10.1016/j.eswa.2015.08.023.
14. P. Ekman and W. V. Friesen, "Constants across cultures in the face and emotion.," *Journal of Personality and Social Psychology*, vol. 17, no. 2, pp. 124–129, 1971, doi: 10.1037/h0030377.
15. T. Niu, S. Zhu, L. Pang, and A. El Saddik, "Sentiment Analysis on Multi-View Social Data," in *MultiMedia Modeling*, Q. Tian, N. Sebe, G.-J. Qi, B. Huët, R. Hong, and X. Liu, Eds., in *Lecture Notes in Computer Science*, vol. 9517, Cham: Springer International Publishing, 2016, pp. 15–27. doi: 10.1007/978-3-319-27674-82.
16. A. Tripathy, A. Agrawal, and S. K. Rath, "Classification of sentiment reviews using n-gram machine learning approach," *Expert Systems with Applications*, vol. 57, pp. 117–126, Sep.

- 2016, doi: 10.1016/j.eswa.2016.03.028.
18. Y. Shan, C. Che, X. Wei, X. Wang, Y. Zhu, and B. Jin, “Bi- graph attention network for aspect category sentiment classification,” *Knowledge-Based Systems*, vol. 258, p. 109972, Dec. 2022, doi: 10.1016/j.knosys.2022.109972.
19. S. Feng, B. Wang, Z. Yang, and J. Ouyang, “Aspect-based sentiment analysis with attention-assisted graph and variational sentence represen- tation,” *Knowledge-Based Systems*, vol. 258, p. 109975, Dec. 2022, doi: 10.1016/j.knosys.2022.109975.
20. Y. Kit and M. M. Mokji, “Sentiment Analysis Using Pre-Trained Language Model With No Fine-Tuning and Less Resource,” *IEEE Access*, vol. 10, pp. 107056–107065, 2022, doi: 10.1109/AC- CESS.2022.3212367.
21. Y. Bie, Y. Yang, and Y. Zhang, “Fusing Syntactic Structure Information and Lexical Semantic Information for End-to-End Aspect-Based Senti- ment Analysis,” *Tsinghua Sci. Technol.*, vol. 28, no. 2, pp. 230–243, Apr. 2023, doi: 10.26599/TST.2021.9010095.
22. H. Fei, Y. Ren, S. Wu, B. Li, and D. Ji, “Latent Target-Opinion as Prior for Document-Level Sentiment Classification: A Variational Approach from Fine-Grained Perspective,” in *Proceedings of the Web Conference 2021*, Ljubljana Slovenia: ACM, Apr. 2021, pp. 553–564. doi: 10.1145/3442381.3449789.
23. G. Rao, W. Huang, Z. Feng, and Q. Cong, “LSTM with sentence rep- resentations for document-level sentiment classification,” *Neurocomput- ing*, vol. 308, pp. 49–57, Sep. 2018, doi: 10.1016/j.neucom.2018.04.045.
24. F. Liu, J. Zheng, L. Zheng, and C. Chen, “Combining attention- based bidirectional gated recurrent neural network and two-dimensional convolutional neural network for document- level sentiment classi- fication,” *Neurocomputing*, vol. 371, pp. 39–50, Jan. 2020, doi: 10.1016/j.neucom.2019.09.012.
25. M. U. Salur and I. Aydin, “A Novel Hybrid Deep Learning Model for Sentiment Classification,” *IEEE Access*, vol. 8, pp. 58080–58093, 2020, doi: 10.1109/ACCESS.2020.2982538.
26. A. U. Hassan, J. Hussain, M. Hussain, M. Sadiq, and S. Lee, “Sen- timent analysis of social networking sites (SNS) data using machine learning approach for the measurement of depression,” in *2017 In- ternational Conference on Information and Communication Technol- ogy Convergence (ICTC)*, Jeju: IEEE, Oct. 2017, pp. 138–140. doi: 10.1109/ICTC.2017.8190959.
27. Y. Du, X. Zhao, M. He, and W. Guo, “A Novel Capsule Based Hybrid Neural Network for Sentiment Classification,” *IEEE Access*, vol. 7, pp. 39321–39328, 2019, doi: 10.1109/ACCESS.2019.2906398.
28. P. Balage Filho, L. Avanc,o, T. Pardo, and M. D. G. Volpe Nunes, “NILC USP: An Improved Hybrid System for Sentiment Analysis in Twitter Messages” in *Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval 2014)*, Dublin, Ireland: Association for Computational Linguistics, 2014, pp. 428–432. doi: 10.3115/v1/S14- 2074.
29. S. Wang, Y. Wei, D. Li, W. Zhang, and W. Li, “A Hybrid Method of Feature Selection for Chinese Text Sentiment Classification,” in *Fourth International Conference on Fuzzy Systems and Knowledge Discovery (FSKD 2007)*, Haikou, China: IEEE, 2007, pp. 435–439. doi: 10.1109/FSKD.2007.49.
30. A. Tripathy, A. Anand, and S. K. Rath, “Document-level sentiment classification using hybrid machine learning approach,” *Knowl Inf Syst*, vol. 53, no. 3, pp. 805–831, Dec. 2017, doi: 10.1007/s10115-017-1055-z.
31. Jaipur National University, Jaipur, Rajasthan, India, A. Dadhich, and B. Thankachan, “Sentiment Analysis of Amazon Product Reviews Using Hybrid Rule-based Approach,” *IJEM*, vol. 11, no. 2, pp. 40–52, Apr. 2021, doi: 10.5815/ijem.2021.02.04.
32. B. AlBadani, R. Shi, and J. Dong, “A Novel Machine Learning Approach for Sentiment

- Analysis on Twitter Incorporating the Universal Language Model Fine-Tuning and SVM,” *ASI*, vol. 5, no. 1, p. 13, Jan. 2022, doi: 10.3390/asi5010013.
33. Dash, S., Padhy, S., Parija, B., Rojashree, T., & Patro, K. A. K. (2022). A Simple and Fast Medical Image Encryption System Using Chaos-Based Shifting Techniques. *International Journal of Information Security and Privacy (IJISP)*, 16(1), 1-24.
34. K. Kumar, B. S. Harish, and H. K. Darshan, “Sentiment Analysis on IMDb Movie Reviews Using Hybrid Feature Extraction Method,” *IJIMAI*, vol. 5, no. 5, p. 109, 2019, doi: 10.9781/ijimai.2018.12.005.
35. R. Ahuja, A. Chug, S. Kohli, S. Gupta, and P. Ahuja, “The Impact of Features Extraction on the Sentiment Analysis,” *Procedia Computer Science*, vol. 152, pp. 341–348, 2019, doi: 10.1016/j.procs.2019.05.008.
36. A. Onan, S. Korukoglu, and H. Bulut, “A hybrid ensemble pruning approach based on consensus clustering and multi-objective evolutionary algorithm for sentiment classification,” *Information Processing and Management*, vol. 53, no. 4, pp. 814–833, Jul. 2017, doi: 10.1016/j.ipm.2017.02.008.
37. Q. A. Xu, V. Chang, and C. Jayne, “A systematic review of social media-based sentiment analysis: Emerging trends and challenges,” *Decision Analytics Journal*, vol. 3, p. 100073, Jun. 2022, doi: 10.1016/j.dajour.2022.100073.
38. Dash, S., Padhy, S., Devi, S. A., & Patro, K. A. K. (2023). An Efficient Intra-Inter Pixel Encryption Scheme to Secure Healthcare Images for an IoT Environment. *Expert Systems with Applications*, 120622. <https://doi.org/10.1016/j.eswa.2023.120622>
39. L. B. Booker, D. E. Goldberg, and J. H. Holland, “Classifier systems and genetic algorithms,” *Artificial Intelligence*, vol. 40, no. 1–3, pp. 235–282, Sep. 1989, doi: 10.1016/0004-3702(89)90050-7.
40. A. Slowik, Ed., *Swarm intelligence algorithms. A tutorial*, First edition. Boca Raton: Taylor and Francis, 2020.
41. B. Wang, D.-X. Li, J.-P. Jiang, and Y.-H. Liao, “A modified firefly algorithm based on light intensity difference,” *J Comb Optim*, vol. 31, no. 3, pp. 1045–1060, Apr. 2016, doi: 10.1007/s10878-014-9809-y.
42. E. M. Mashhour, E. M. F. El Houby, K. T. Wassif, and A. I. Salah, “Feature Selection Approach based on Firefly Algorithm and Chi-square,” *IJECE*, vol. 8, no. 4, p. 2338, Aug. 2018, doi: 10.11591/ijece.v8i4.pp2338-2350.