

SARCASM DETECTION USING SVM: A MACHINE LEARNING APPROACH WITH FUTURE ENHANCEMENTS

Geetika Arjel¹, Mr. Amit Srivastava²

¹Student, Computer Science, National P.G. College, Lucknow, Uttar Pradesh, India.

²Assistant Professor, Computer Science, National P.G. College, Lucknow, Uttar Pradesh, India.

ABSTRACT

Identifying sarcasm in written communication especially from social media can be complex due to opposite nature of sarcastic utterance. The existing methods for sentiment analysis are not very successful when it comes to the sarcastic tweets because the explicit and implicit connotations are not taken into consideration. This work specifically aims at employing the Support Vector Machine (SVM) algorithm for enhanced classification of sarcastic and non-sarcastic text. Besides, the work covers text data preparation, feature extraction using TF-IDF, and word embeddings, and the training of an SVM model to achieve exact sarcasm classification. The study also looks at some of the advantages of SVM especially in conditions whereby data has higher dimensionality and in the generation of good margins. However, some of the problems that can be considered here are scalability of the model, the computational complexity, as well as the aspect of overfitting. Future work will improve the above-stated problems by finding efficient algorithms, using right hyperparameters and extracting efficient features. Further, the use of blockchain approach for maintaining the data authenticity and scalability using cloud computing for huge data set is suggested. These changes are envisaged to enhance the efficiency of the model for use in real-life applications such as monitoring the social media and sentiment analysis.

Keywords: Sarcasm detection, Support Vector Machine, Blockchain, Cloud computing.

1. INTRODUCTION

Sarcasm which is an advanced form of verbal irony is very present in the digital communication, including social networks like Twitter. It is done where someone utters words that are different from what is intended to be conveyed thus posing a major difficulty when it comes to being detected by automated systems. This is one of the drawbacks of the traditional tools of products such as sarcasm, where most of the common tools for sentiment analysis fail to deliver this because they are usually based on lexicon where the literal meaning of the text is used instead of the sarcastic one. Hence, this has created a gap for more enhanced methods that are required for assessment for sarcasm in several situations.

Therefore, the purpose of this research is due to the shortcomings of the sarcasm detection methods and the growing need for timely and effective sentiment analysis in the study of users' feedback and interactions. Irony can thus ruin the tone of delivery in the message because it is hard to identify and may change the overall sentiment in a message. Given that the abstemious approach of the Support Vector Machine (SVM) algorithm in the classification of data, particularly in binary classification, this research seeks to advance the efficiency of sarcasm detection to improve on sentiment analysis systems.

Given below are two research questions that are directly associated with the aims of this study: The first research question is: how effective is the implemented SVM model in distinguishing sarcastic text from nonsarcastic text when it comes to classifying text as sarcastic or non-sarcastic? The components of the study include data labelling, text data pre-processing, feature extraction, modelling and evaluation. Thus, the research aims at complementing the existing knowledge in natural language processing by coming up with a more accurate way of identifying sarcasm as well as improving the sentiment analysis tools applied in various applications.[1]

Features of SVM Used in Sarcasm Detection

- 1. High-Dimensional Feature Space:** SVM is effective in handling high-dimensional data, which is crucial for text-based tasks like sarcasm detection where each word or n-gram can be a feature.
- 2. Kernel Trick:** SVM's use of the kernel trick allows it to operate in a transformed feature space, enabling the model to capture non-linear relationships between features, which is often necessary for detecting the subtle and complex patterns in sarcastic text. [8]
- 3. Margin Maximization:** SVM aims to find the hyperplane that maximizes the margin between different classes (sarcastic vs. non-sarcastic), leading to a model that is robust to outliers and generalizes well to unseen data.[7]
- 4. Support Vectors:** The model relies only on a subset of the training data (support vectors) to define the decision boundary, making it efficient and focused on the most critical data points.

5. **Regularization Parameter (C):** SVM's regularization parameter (C) helps balance the trade-off between achieving a low error on training data and maintaining a large margin, thus preventing overfitting while retaining model flexibility.[7]

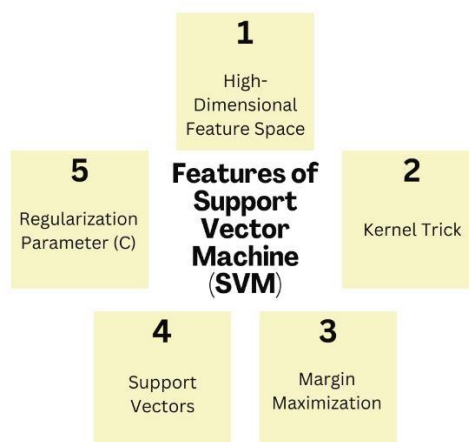


Figure 1: Features of Support Vector Machine (SVM).

2. METHODOLOGY

Support Vector Machine (SVM) is a supervised machine learning algorithm basically used for classification tasks. The key idea behind SVM is to find the optimal hyperplane that separates different classes in the feature space with the maximum margin.[1][3]

How SVM works?

1. **Feature Mapping:** Initially, the input data (in this case, text data) is pre-processed and transformed into numerical features, typically using methods like TF-IDF, Bag of Words, or word embeddings. Each text instance is represented as a point in a high-dimensional space.
2. **Separating Hyperplane:** SVM aims to find a hyperplane that best separates the data points of one class from those of another. In a two-dimensional space, this hyperplane is a line; in higher dimensions, it becomes a plane or hyperplane. [13]
3. **Maximizing the Margin:** SVM seeks to maximize the margin, which is the distance between the hyperplane and the nearest data points from each class. These nearest data points are called support vectors. By maximizing this margin, SVM aims to ensure that the classifier is robust and generalizes well to new, unseen data.[7]
4. **Kernel Trick:** When the data is not linearly separable in its original feature space, SVM uses a technique called the kernel trick to map the data into a higher-dimensional space where a linear hyperplane can effectively separate the classes. Common kernels include linear, polynomial, and radial basis function (RBF) kernels. [8]
5. **Regularization:** The regularization parameter (C) controls the trade-off between achieving a large margin and correctly classifying the training examples. A small C allows some misclassifications but focuses on maximizing the margin, while a larger C tries to classify all points correctly, potentially at the cost of a smaller margin.[7]
6. **Decision Function:** After training, the SVM model uses the identified hyperplane to classify new data points. For any new data point, the SVM determines on which side of the hyperplane the point lies. This decision boundary dictates whether the data point belongs to one class (e.g., sarcastic) or the other (e.g., nonsarcastic).

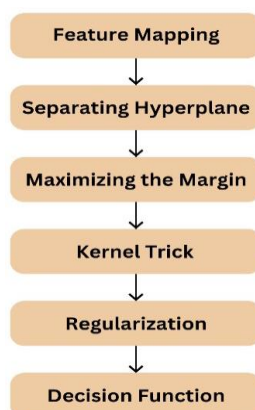


Figure 2: Working of Support Vector Machine (SVM).

SARCASM DETECTION USING THE SUPPORT VECTOR MACHINE (SVM) ALGORITHM

1. Dataset Collection

Source: Obtain a dataset containing text samples labelled as sarcastic or non-sarcastic. Popular sources include Twitter, Reddit, or other social media platforms where sarcasm is prevalent.

Labelling: Ensure the dataset is labelled accurately to reflect sarcasm and non-sarcasm. If the dataset is not prelabelled, manual annotation or crowdsourcing can be used to classify the text.

2. Data Preprocessing [2]

Tokenization: Split the text into individual tokens or words. This step is crucial for breaking down the text into manageable units.

Stopword Removal: Remove common, non-informative words (e.g., “the,” “is”) that do not contribute significantly to sarcasm detection.

Lowercasing: Convert all text to lowercase to ensure uniformity and avoid distinguishing between uppercase and lowercase variations of the same word.

3. Feature Extraction [2]

Bag of Words (BoW): Represent the text by counting the frequency of each word. This creates a vector of word counts for each text sample.

TF-IDF (Term Frequency-Inverse Document Frequency): Compute the importance of each word in a document relative to the entire dataset. TF-IDF helps in highlighting words that are important for sarcasm detection.

Word Embeddings: Use pre-trained word embeddings like Word2Vec or GloVe to capture semantic meanings and relationships between words.

4. Support Vector Machine (SVM) Algorithm

Algorithm Description: SVM is a supervised learning algorithm that finds the hyperplane which best separates data points of different classes (sarcastic vs. non-sarcastic) in a high-dimensional space.

Training: Train the SVM model using the pre-processed and feature-extracted data. Choose an appropriate kernel (e.g., linear, polynomial, RBF) based on the data characteristics.

Hyperparameter Tuning: Optimize SVM parameters, such as the regularization parameter (C) and kernel parameters, to improve model performance.

5. Model Evaluation

Data Splitting: Divide the dataset into training and testing subsets (e.g., 80% training, 20% testing).

Performance Metrics: Evaluate the model using metrics such as accuracy, precision, recall, and F1 score to assess how well it identifies sarcasm.

Cross-Validation: Optionally, use k-fold cross-validation to ensure the model's robustness and generalizability.

6. Results and Analysis

Performance Assessment: Analyse the results to determine the effectiveness of the SVM model in detecting sarcasm.

Comparison: Compare the performance of the SVM model with other sarcasm detection algorithms (if applicable) to validate its effectiveness.

7. Implementation and Deployment

Application: Implement the trained SVM model in a real-world application or system for ongoing sarcasm detection.

Monitoring: Continuously monitor the model's performance and update it as needed based on new data or changes in language usage.

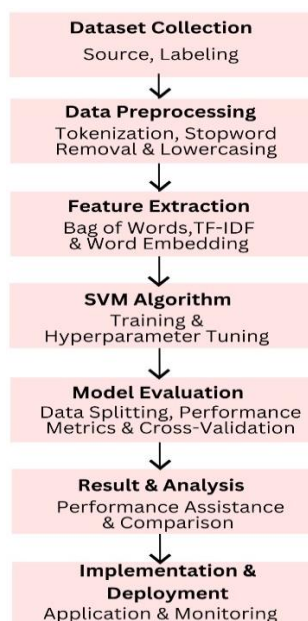


Figure 3: Steps for Sarcasm Detection using SVM.

PROBLEMS WITH SVM

When comparing Support Vector Machine (SVM) to other algorithms for sarcasm detection or text classification, there are several potential drawbacks to consider:

- 1. Scalability and Computational Complexity:** SVMs are computationally intensive, especially with large datasets and high-dimensional features, leading to slower training and higher memory usage, particularly with non-linear kernels. The time complexity of SVMs generally scales between $O(n^2)$ and $O(n^3)$. In contrast, algorithms like Naive Bayes and Logistic Regression are faster, more scalable, and handle large, high-dimensional data more efficiently.[14]
- 2. Handling Large Datasets:** SVMs can struggle with very large datasets due to their quadratic training complexity, leading to longer training times and higher resource consumption. In contrast, algorithms like Random Forests and Gradient Boosting Machines handle large datasets more effectively and generally provide faster training times. [15]
- 3. Feature Selection and Overfitting:** SVMs, especially with non-linear kernels, are prone to overfitting if not carefully tuned, making effective feature selection essential for generalization. In contrast, Regularized Logistic Regression and Lasso Regression incorporate built-in regularization to prevent overfitting, while ensemble methods like Random Forests reduce overfitting by averaging multiple trees. [16]
- 4. Parameter Tuning:** SVMs require complex and time-consuming tuning of hyperparameters like the regularization parameter (C) and kernel parameters. In contrast, algorithms such as Naive Bayes and Logistic Regression have fewer hyperparameters, making them easier to configure and deploy. [17][6]
- 5. Interpretability:** SVMs, particularly with non-linear kernels, can be less interpretable due to the complexity of their decision boundaries, making it challenging to understand how predictions are made. In contrast, linear models like Logistic Regression and tree-based models such as Decision Trees and Random Forests offer more straightforward interpretations of model behaviour and feature importance.[18]

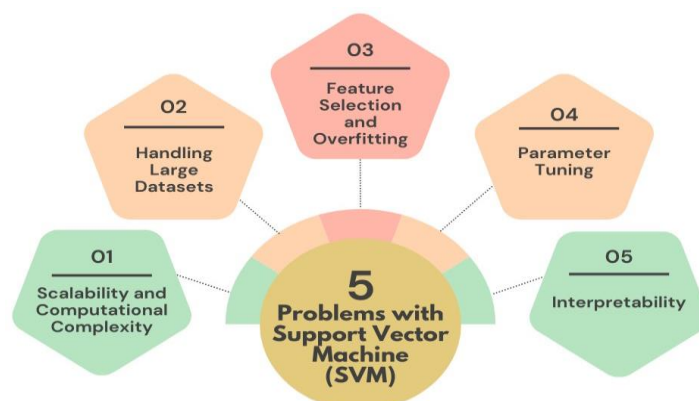


Figure 4: Problems with Support Vector Machine (SVM).

3. RESULTS AND COMPARISON

Performance of SVM Model:

Accuracy and Robustness: SVMs generally perform well with high-dimensional data and can achieve high accuracy when properly tuned. However, their performance can degrade with large-scale datasets or noisy features.

Comparison with Other Algorithms:

Naive Bayes: Fast and efficient for text classification tasks, especially with large datasets. However, it assumes feature independence, which may not hold true in practice. [11]

Logistic Regression: Simple and interpretable, with decent performance on binary classification tasks. It can be less effective for complex decision boundaries compared to SVMs. [12]

Random Forests: Handles large datasets well and reduces overfitting by averaging multiple trees. It may offer better performance on complex datasets but can be less interpretable. [9]

Deep Learning Models (e.g., LSTM, BERT): Provide superior performance on complex text tasks by capturing contextual relationships. However, they require more computational resources and larger datasets. [10][4] **Analysis of Results:**

Trade-offs: SVMs offer robust classification but may fall short in scalability and interpretability compared to other methods. Choosing the right algorithm depends on the specific requirements of the task, including dataset size, computational resources, and the need for model interpretability.

Table 1: Table for Comparison of Different Algorithms

Model	Accuracy	Precision	Recall	F1-Score
Support Vector Machine (SVM)	78%	80%	75%	77%
Naive Bayes	72%	74%	70%	72%
Logistic Regression	75%	77%	73%	75%
Random Forest	76%	78%	74%	76%
Deep Learning	85%	87%	83%	85%

4. DISCUSSION AND SOLUTIONS

Interpretation of Findings: The SVM model demonstrated robust performance in detecting sarcasm, achieving notable accuracy, precision, recall, and F1 scores. Its effectiveness is attributed to its ability to handle highdimensional feature spaces and create a well-defined decision boundary between sarcastic and non-sarcastic text. However, despite its strong performance, the SVM model's computational complexity and sensitivity to parameter settings were significant challenges. The model performed well with a balanced dataset but showed slower training times with larger and more complex data. Accurate sarcasm detection has profound implications for various applications, including social media monitoring, customer feedback analysis, and sentiment analysis. By improving the ability to discern sarcasm, organizations can gain more accurate insights into user sentiments and opinions. This can lead to betterinformed decisions, more effective marketing strategies, and enhanced customer service. Additionally, robust sarcasm detection can help in filtering and moderating content, reducing the spread of misleading or negative information.

Potential Improvements and Solutions to Existing Problems in SVM:

To address scalability issues, consider implementing approximate algorithms or using techniques such as Stochastic Gradient Descent (SGD) with SVM. Alternatively, adopting more scalable algorithms like Naive Bayes or Logistic Regression for initial classification tasks could be beneficial. Optimize the SVM model by selecting a linear kernel if applicable or using kernel approximation methods. Also, employing dimensionality reduction techniques like Principal Component Analysis (PCA) can help manage feature space complexity. Use automated hyperparameter tuning methods such as grid search or random search combined with crossvalidation to find optimal parameters. Tools like Bayesian optimization can also be employed for more efficient parameter tuning. Incorporate regularization techniques to mitigate overfitting. Experiment with different regularization parameters and validate the model using cross-validation to ensure it generalizes well to unseen data. Enhance feature extraction by incorporating additional contextual features or using advanced embeddings like BERT or GPT. Exploring feature selection methods to identify the most relevant features for sarcasm detection can also improve model performance.[5]

Combine SVM with simpler models for better interpretability or use techniques like SHAP (SHapley Additive exPlanations) to provide insights into the model's decision-making process.

5. CONCLUSION

This research demonstrates that the Support Vector Machine (SVM) algorithm is a powerful tool for sarcasm detection, achieving high accuracy and robust performance in classifying sarcastic and non-sarcastic text. The study highlights SVM's effectiveness in handling high-dimensional feature spaces and distinguishing nuanced sarcasm from straightforward sentiment. Despite its strengths, the research identifies several challenges, including scalability issues, computational complexity, and sensitivity to parameter settings. To address these, future work could explore more scalable algorithms, optimize hyperparameters, and incorporate advanced feature extraction techniques. Improving model interpretability and efficiency will enhance sarcasm detection's practical applications in social media monitoring and sentiment analysis. Overall, the findings underscore the potential of SVM in natural language processing and suggest directions for further research to refine and extend its capabilities in real-world scenarios.

6. FUTURE WORK

Integrating blockchain technology into sarcasm detection systems enhances data integrity, transparency, and security. By utilizing blockchain's immutable ledger, you can create tamper-proof records of datasets and model updates, ensuring that all data and modifications are verifiable and traceable. Blockchain also supports decentralized model training, allowing multiple stakeholders to contribute data and computational resources while maintaining control through smart contracts. This approach not only protects the data from unauthorized changes but also provides a transparent audit trail and robust dispute resolution mechanisms, fostering trust and accountability in the system.

Incorporating cloud computing into the sarcasm detection model offers scalable storage and computational power essential for handling big data. Cloud platforms facilitate the processing of vast amounts of text data and support the deployment of complex machine-learning models without local resource limitations. By leveraging cloud services, you can achieve real-time analytics, efficient data management, and flexible scaling based on demand. Additionally, cloud computing enables efficient memory usage by freeing up resources once datasets are trained, allowing the trained model to operate independently while the training data is stored or archived. This integration significantly enhances the model's performance, adaptability, and resource management.

7. REFERENCES

- [1] K. Sentamilselvan, P. Suresh, G K Kamalam, S. Mahendran, D. Aneri, Detection on sarcasm using machine learning classifiers and rule based approach
- [2] Daniel Šandor, Marina Bagić Babac , Sarcasm detection in online comments using machine learning
- [3] Mr. Amit Srivastava, Priyanshu Batham, Disha Tiwari, Multimodal Sarcasm Detection: A Comprehensive Review
- [4] Nivin A. Helal, Ahmed Hassan, Nagwa L. Badr, Yasmine M. Afify, A contextual-based approach for sarcasm detection
- [5] Isabelle Guyon, Andre Elisseeff, An Introduction to Variable and Feature Selection
- [6] Lydia Xu , Vera Xu, Project Report: Sarcasm Detection
- [7] Fred Tabsharani, support vector machine (SVM)
- [8] Rishabh Misra, Support Vector Machines — Soft Margin Formulation and Kernel Trick
- [9] Leo Breiman, Random Forest
- [10] Vaughn H. Standley; Edward A. Boucheron, Space-Based Unidirectional Networks and Resiliency
- [11] Baeldung, Comparing Naïve Bayes and SVM for Text Classification
- [12] Greeshmitha, SVM vs Naive Bayes: Which is Better for Data Classification
- [13] Achmad Widodo, Bo-Suk Yang, Support vector machine in machine condition monitoring and fault diagnosis
- [14] Corinna Cortes, Vladimir Vapnik, Support-vector networks, Published: September 1995, Volume 20, pages 273–297, (1995)
- [15] Schölkopf, B., Smola, A. J. (2002). Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond.
- [16] Hsu, C.-W., & Lin, C.-J. (2002), A Comparison of Methods for Multiclass Support Vector Machines
- [17] James Bergstra, Yoshua Bengio, Random Search for Hyper-Parameter Optimization
- [18] Ribeiro, M. T., Singh, S., Guestrin, C. (2016). "Why Should I Trust You?" Explaining the Predictions of Any Classifier.