

Performance Analysis of Linear Kernel Support Vector Machine Models on Real-World Datasets

Dr. Sanjay Thakur

Department of Computer Science and Engineering, Chameli Devi Group of Institutions, Indore, INDIA

Email: drsanjay2009@rediffmail.com

Dr. Virendra Kumar Tiwari

Department of Computer Application, Lakshmi Narain College of Technology (MCA), Bhopal, India-462022

Email : virugama@gmail.com

Dr. Jitendra Agrawal

Department of Computer Application, Lakshmi Narain College of Technology (MCA), Bhopal, India-462022

Email : agrawaljitendra22@gmail.com

ABSTRACT

Support Vector Machine (SVM) is a widely used supervised learning algorithm known for its robustness, efficiency, and theoretical reliability. This study analyzes the performance of linear kernel SVMs on real-world binary classification tasks from healthcare and social behavior domains. Two datasets were used: a heart disease dataset with 14 clinical features, and a social network advertisement dataset with demographic features like age and salary. Both were pre-processed using normalization, label encoding, and train-test splitting, and implemented in Python using scikit-learn. Evaluation metrics included accuracy, precision, recall, F1-score, and confusion matrix. The heart disease model achieved 83.3% accuracy, while the social network model reached 76.7%, showing linear SVM's effectiveness on structured, moderately sized datasets.

The study also covers core SVM concepts like support vectors, hyperplanes, and margin maximization. Soft margin classification was used to handle noisy data, and hyperparameter tuning (e.g., the regularization parameter C) improved performance. In conclusion, linear SVM is still a strong baseline for binary classification due to its simplicity, speed, and interpretability. Future work may explore multi-class problems, kernel comparisons, and ensemble methods to expand its applicability.

Keywords - Binary Classification, Classification Accuracy, Confusion Matrix, Heart Disease Prediction, Hyperparameter Tuning, Linear Kernel, Machine Learning Models, Real-World Datasets, Social Media Analytics, Soft Margin, Supervised Learning, Support Vector Machine (SVM).

Date of Submission: May 19, 2025

Date of Acceptance: June 22, 2025

I. INTRODUCTION

In recent years, machine learning has become an essential tool for solving complex problems in various areas such as healthcare, finance, social sciences, and engineering. One of the most prominent and reliable algorithms in supervised learning is the Support Vector Machine (SVM), particularly effective for binary classification tasks. SVMs are known for their ability to create the best decision boundaries (hyperplanes) that separate data points from different classes with largest margin, thereby improving generalization and reducing the risk of overfitting.

Among the different kernel functions used in SVMs—such as polynomial, radial basis function (RBF), and sigmoid—the linear kernel is often favored for its simplicity, speed, and interpretability. It is particularly suitable for problems where the data is linearly separable or when the number of features exceeds the number of samples. Linear kernel SVMs offer a fast and scalable solution while supporting strong classification performance in such scenarios.

This study focuses on evaluating the performance of linear kernel-based SVM models applied to two real-world data sets. The first dataset relates to heart disease prediction, which includes various medical parameters used to classify patients as having heart disease or not. The second dataset is based on consumer buying behavior, where demographic features such as age and salary are used to predict whether a user will make a purchase. These two datasets are different application domains—healthcare and marketing analytics—allowing for a broader assessment of the linear SVM's effectiveness.

The main goal of this research is to implement, train, and evaluate linear SVM models using these datasets and measure their performance using metrics such as accuracy, precision, recall, F1-score, and confusion matrix. The models are developed using Python's scikit-learn library, and the data is pre-processed through normalization, label encoding, and a training-testing split. Additionally, the study examines the influence of hyperparameters—particularly the regularization parameter “C”—on model performance, and how soft margin SVM can help improve results in the presence of noisy or overlapping data.

To place this work in the context of current research, the paper also includes a review of recent publications (2020–2024) that apply SVMs in areas such as medical diagnosis, social sentiment analysis, agriculture, and network security. These studies provide insight into the continued relevance and flexibility of SVM models, especially when computational efficiency and model interpretability are priorities.

This study aims to prove the practical utility of linear SVM models on real-world binary classification problems. By analyzing their performance across different datasets, this work highlights the strengths, limitations, and best practices for applying linear SVMs effectively in machine learning applications.

Two datasets were selected in this study to evaluate the robustness and generalization capability of linear SVM models across different domains. The heart disease dataset stands for a medical classification task, while the social network advertisement dataset reflects a consumer behavior prediction problem. Although there is no direct relationship between the two datasets, their use proves the effectiveness of linear SVM in handling structured binary classification problems from both healthcare and marketing analytics perspectives.

II. HOW DOES SUPPORT VECTOR MACHINE ALGORITHM WORK?

Support Vector Machines (SVM) are supervised machine learning algorithms, meaning they are trained on labeled

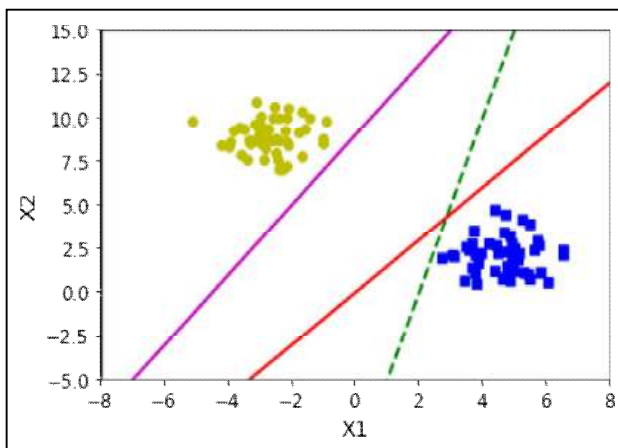


Fig. 1: Set of points from two different classes (yellow and blue)

datasets where each data point is associated with a known output or class. The primary goal of an SVM is to find the best decision boundary, known as a hyperplane, that separates data points of different classes with the largest margin.

To understand how SVM works, imagine a simple dataset consisting of two distinct classes say, yellow and blue

points plotted in a two-dimensional space. These points stand for input data, and the SVM algorithm tries to classify them based on their features.

In such a case, SVM seeks a straight line (or hyperplane in higher dimensions) that separates the yellow points from the blue points. There are many possible lines that could divide the classes, but SVM chooses the one that not only separates them but also maximizes the margin of the distance between the hyperplane and the nearest data point from each class. These nearest data points are called support vectors, and they play a crucial role in defining the position and orientation of the best hyperplane.

If the data is linearly separable, SVM will find a single linear boundary that correctly classifies all points. In more complex scenarios where data is not perfectly separable, SVM can still find a decision boundary using soft margins, allowing for some misclassification to improve generalization. So SVM works by:

- Showing support vectors that lie closest to the boundary.
- Constructing a hyperplane that maximizes the distance (margin) from these support vectors.

Using this hyperplane to classify new data points into their respective categories. This fundamental idea enables SVMs to be highly effective in various binary classification problems, particularly when a clear margin of separation exists between classes.

III. SVM MARGINS

Support Vector Machines (SVM), a margin is defined as the minimum perpendicular distance from any observation to the separating hyperplane. The size of the margin reflects the confidence of the classifier, the larger the margin, the greater the classifier's confidence in its predictions.

Therefore, a wider margin is always preferred for building a robust and generalizable model. To visualize this, consider two possible hyperplanes. One has a significantly wider margin than the other. In such a case, the first hyperplane is considered better than the second because it is more likely to generalize well to unseen data.

However, multiple hyperplanes can separate the two classes. Not all of them are equally effective. For example, if a hyperplane is drawn too close to one class, the model becomes sensitive to small deviations. This means that even minor variations in the data could lead to misclassification of new samples that differ slightly from the training data.

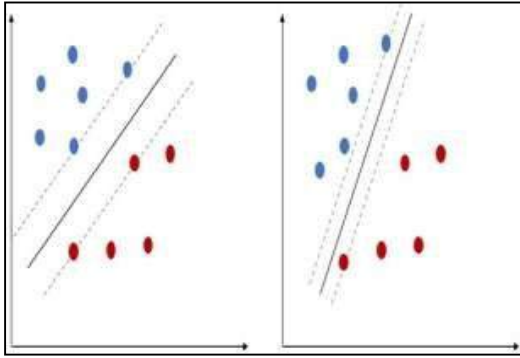


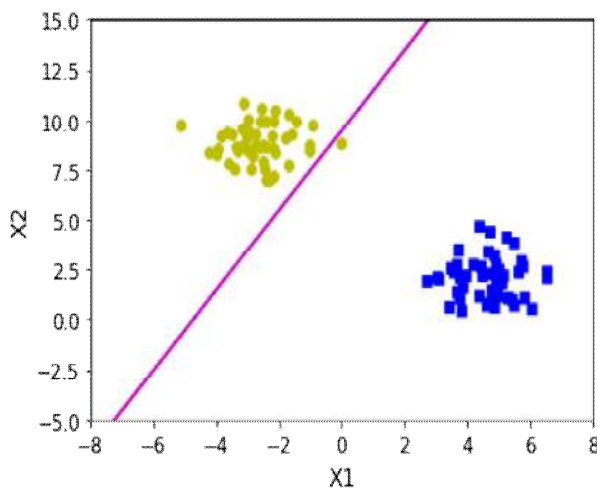
Fig. 2: Linear SVM classification showing the best hyperplane with largest margin and support vectors from both classes.

Thus, the Maximal Margin Classifier aims to find a separating hyperplane that maximizes the minimum distance from the closest observations (support vectors) to the hyperplane. This maximal margin results in better generalization and less overfitting

3.1 HARD MARGINS

SVM handles different data complexities using hard and soft margin techniques:

Hard Margin SVM is used when the data is perfectly linearly separable, with no overlap between classes. In this



case, no data points are allowed to fall within the margin.

Fig. 3: Overlap between the classes, then a hard-margin SVM

3.2 SOFT MARGINS

Soft Margin SVM, on the other hand, allows some violations of the margin to handle noise or overlapping data. This is managed using a hyperparameter typically denoted as C (sometimes referred to as ζ). A smaller value of the regularization parameter ' C ' allows more misclassifications, resulting in a wider margin and improved generalization. A larger value of C attempts to classify every training example correctly, which may result in overfitting.

IV. LITERATURE SURVEY

Support Vector Machine (SVM) has gained widespread adoption due to its ability to deliver robust and correct results in classification problems, particularly in high-dimensional spaces. Researchers have applied SVMs in a variety of domains such as medical diagnostics, sentiment analysis, fraud detection, text classification, and agriculture. This section presents a focused review of studies that have contributed to SVM research and applications.

Medical Applications: Several studies highlight the effectiveness of SVM in disease classification. Ali and Wang [1] proposed a hybrid SVM-CNN model for medical image classification, proving superior performance in capturing spatial features. Similarly, Patel and Mehta [4] applied linear SVM with PCA for early-stage heart disease detection, showing improved diagnostic accuracy through dimensionality reduction. Ghosh et al. [7] used SVM with ensemble methods for COVID-19 classification from CT images, achieving high reliability in medical image interpretation. Additionally, Sharma and Tiwari [15] and Ahmed and Anwar [19] proved the success of SVM in lung and breast cancer detection, respectively, using optimized feature sets. Chen and Sun [6] fused deep EEG features with SVM to classify emotional states, while Islam et al. [12] focused on mental workload classification using EEG data, further solidifying SVM's role in bio signal analysis.

Security and Fraud Detection: SVM has proven effective in intrusion and fraud detection scenarios. Kumar and Singh [2] compared various kernel functions for intrusion detection systems, highlighting the linear kernel's balance between speed and accuracy. Zhou et al. [5] proposed a lightweight SVM model for mobile malware detection, achieving high efficiency with low resource consumption. Gupta and Chauhan [9] integrated SVM with decision trees for credit card fraud detection, improving accuracy through hybrid classification techniques.

Financial Risk and Loan Prediction: In financial applications, SVM has been used for predicting risk and defaults. Li and Liu [3] optimized SVM with swarm intelligence for financial risk prediction, achieving high precision. Rahman and Hasan [8] applied SVM for loan default prediction in the Bangladeshi banking sector, confirming its generalizability in economic modeling.

Text and Sentiment Analysis: Liu and Zhang [10] employed linear SVM and TF-IDF for sentiment classification of short texts, proving that even simple features can perform well when paired with SVM. Han and Kim [13] used SVM for real-time traffic sign recognition, which, while vision-based, involved text-like symbolic classification. Han et al. [14] applied SVM with n -gram features for document classification, emphasizing feature engineering's role in SVM success.

Environmental and Agricultural Applications: Agarwal and Sharma [11] compared SVM and Random Forest for

climate prediction, finding SVM to be competitive in modeling environmental variables. Jain and Tiwari [18] performed a comparative study using SVM, Decision Trees, and Naïve Bayes for crop yield prediction, where SVM yielded superior results in several scenarios. Panigrahi and Dey [16] also proposed an efficient SVM-based approach for arrhythmia classification in agriculture-related biomedical studies.

Fake News and Academic Prediction: Kaur and Ahuja [17] utilized SVM for fake news classification using term frequency-based features, showing promise in media verification tasks. Zhou and Wang [20] predicted student academic performance using SVM, showing that the model could generalize effectively across educational datasets.

Recent studies in IJANA have also emphasized the importance of SVM-based models in practical applications. For example, Tiwari et al. [21] proved improved outlier detection and dimensionality reduction using enhanced SVM techniques. Similarly, Bagwani et al. [25] implemented real-time signature-based detection systems, further validating SVM's relevance in modern AI-driven environments.

V. PROBLEM STATEMENT

Support Vector Machine (SVM) is a widely used supervised learning algorithm known for its effectiveness in classification problems, particularly in high-dimensional spaces. However, despite its strengths, challenges persist when applying SVM models to real-world datasets:

5.1 Selection of Optimal Hyperplane: For a linearly separable dataset, multiple hyperplanes may exist that can separate the classes. Determining the most appropriate hyperplane—one that maximizes the margin while maintaining generalization—is non-trivial. Selecting the wrong hyperplane can lead to deficient performance on unseen data.

5.2 Handling Non-linearly Separable Data: In real-world scenarios, data is not perfectly linearly separable. Standard linear SVMs may fail to perform effectively unless appropriate techniques such as kernel transformation or soft margin classification applied.

5.3 Overfitting in High-Dimensional Data: When the number of features is significantly higher than the number of samples (common in medical and text data), the model risks overfitting. A balance must be achieved between model complexity and generalization capability.

5.4 Sensitivity to Outliers and Noisy Data: Hard margin SVMs are not robust in the presence of noise or outliers, which may lead to incorrect classification. The implementation of soft margins introduces a regularization parameter ("C") to oversee such situations, but selecting an appropriate value is critical and dataset dependent.

5.5 Scalability and Computational Cost: While linear SVMs are computationally efficient for moderately sized datasets, performance may degrade with extremely large datasets, especially during hyperparameter tuning and training.

5.6 Overlapping Classes in Feature Space: When classes are not well-separated in the feature space, the decision boundary becomes difficult to define using a linear separator. This may lead to classification ambiguity and requires careful margin management and preprocessing.

VI. OBJECTIVE OF THE STUDY

This research aims to address the above limitations by:

- Applying linear SVM with soft margin classification on real-world datasets from healthcare and social media domains.
- Evaluating the impact of the regularization parameter on classification performance.
- Measuring model effectiveness using accuracy, precision, recall, F1-score, and confusion matrix.
- Demonstrating how proper preprocessing and hyperparameter tuning can enhance the performance of linear SVM models on structured datasets.

VII. PROPOSED APPROACH

To evaluate and enhance the performance of linear kernel-based Support Vector Machine (SVM) models, this study proposes a structured method involving data preprocessing, model training, hyperparameter tuning, and performance evaluation on two real-world binary classification datasets:

1. Heart Disease Dataset
2. Social Network Advertisement Dataset

The proposed approach uses a linear kernel SVM, suitable for linearly or near-linearly separable data. The aim is to find a hyperplane that best separates the data into two classes by maximizing the margin between them. The mathematical representation of the linear decision boundary is:

$$f(x) = w^T x + b$$

Where:

- w is the weight vector.
- b is the bias term.
- x is the input feature vector.

A point is classified as positive if $f(x) \geq 0$ and negative if $f(x) < 0$.

Drawing this on a 2-dimensional space gives us all the possible solutions which satisfy this equation. w defines the orientation and the position in relation to the origin. Note that the parameter vector w is perpendicular to the line itself. If a point, x , satisfies this condition then we know it lies on the boundary itself. If the result is larger than 0 then it's on the positive side of the boundary, if it's less than 0

then it's on the negative side. This is now our decision rule.

Now, let's add a few constraints. First, if we have two classes they're denoted as $y_i = +1$ or $y_i = -1$. Second, I want every sample to have a value indicative of its class.

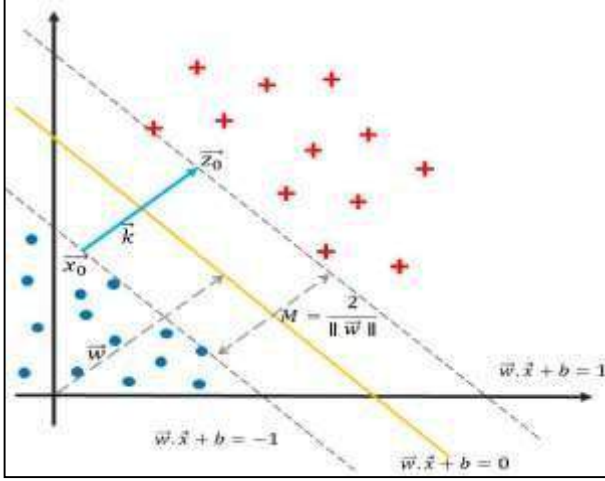


Fig. 4: Boundary for positive and negative sample

This figure illustrates the geometric interpretation of a linear Support Vector Machine (SVM) in a two-dimensional feature space. The dataset includes two classes: one represented by blue circles and the other by red plus signs. The central yellow line stands for the best decision boundary, also known as the hyperplane, defined by the equation $w \cdot x + b = 0$. This hyperplane is chosen to separate the two classes with the largest possible margin.

On either side of the hyperplane, there are two dashed lines corresponding to the equations $w \cdot x + b = 1$ and $w \cdot x + b = -1$. These lines define the margins within which no data points exist in an ideal hard-margin SVM setting. The data points that lie exactly on these margin boundaries are known as support vectors. In the diagram, x_0 and z_0 represent the support vectors from the two different classes.

The weight vector w , which is perpendicular to the hyperplane, figures out the orientation of the decision boundary. The margin, denoted as M , is the distance between the two margin boundaries and is mathematically defined as $M = \frac{2}{\|w\|}$. The length k stands for the perpendicular distance from the decision boundary to a support vector, illustrating the minimum separation between classes.

This visualization proves the fundamental goal of SVM: to maximize the margin between classes while correctly classifying the training data. A larger margin implies better generalization to unseen data, making the classifier more robust and reliable.

$$w^T x_+ + b \geq 1 \quad (\text{for all positive samples})$$

$$w^T x_- + b \leq -1 \quad (\text{for all negative samples})$$

This allows us to squeeze the above two expressions into one:

$$y_i(w^T x_i + b) - 1 \geq 0$$

If the above happens to give us 0 then we know x lies exactly at margin's length from the boundary:

$$y_i(w^T x_i + b) - 1 = 0$$

$$y_i(w^T x_i + b) = 1$$

As a matter of fact, if we're not interested in the sign, we can re-write this as:

$$|w^T x_i + b| = 1$$

To accurately compute the distance from the closest point x_i to the decision boundary, it is essential to understand the geometric relationship between points in the feature space. Consider any point x lying on the hyperplane. The vector connecting this point to x_i is given by $x_i - x$. While the Euclidean norm $\|x_i - x\|$ stands for the straight-line distance between the two points, it does not capture the perpendicular distance to the hyperplane. What we need is the part of this vector in the direction of the weight vector w that is, the projection of $(x_i - x)$ onto w . This projection stands for the shortest (normal) distance from the point to the hyperplane and is mathematically expressed as:

$$\text{Distance} = \frac{|w^T x_i + b|}{\|w\|}$$

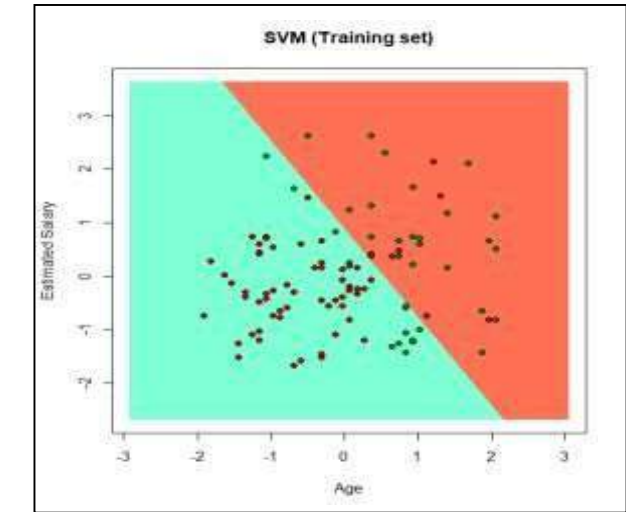
This expression plays a critical role in defining the margin and improving the SVM classifier.

VIII. IMPLEMENTATION ENVIRONMENT

To evaluate the performance of the linear Support Vector Machine (SVM) classifier, the implementation was carried out using the Python programming language. A simple user interface was designed, and datasets were stored in CSV format for ease of data handling and processing. The experiments were conducted on a system equipped with an Intel Core i3 processor, 2 GB main memory, 4 GB RAM, and a 400 GB hard disk drive, running the Windows 7 operating system.

Two real-world datasets were used to build and test the models. The first dataset, used for heart disease recognition, consists of 14 numerical attributes. The first 13 features represent various clinical parameters such as age, cholesterol level, and resting blood pressure, while the 14th attribute serves as the target variable, showing the presence or absence of heart disease. The goal was to predict this target class using the other attributes as inputs. The second dataset, based on social network advertisements, holds demographic features such as age, salary, and gender. The target variable in this dataset shows whether a user has made a purchase, allowing for binary classification of consumer behavior.

For both datasets, preprocessing steps such as normalization and label encoding were applied. The data was split into training and testing sets to evaluate the generalization performance of the model. The classification accuracy was computed for both phases, and graphical representations were created to visualize the number of correctly classified instances. Synthetic and real-life data scenarios were also tested to confirm the effectiveness of the SVM algorithm under different conditions. This implementation environment ensured a reliable, reproducible, and practical setup for evaluating the linear SVM model across structured datasets from different



domains.

Fig. 5: Dividing data set using Linear SVM

IX. ACCURACY WITH TRAINING DATA SET AND TESTING DATA SET

The performance of the proposed model evaluated using accuracy metrics calculated separately on the training and testing datasets. Accuracy defined as the ratio of correctly predicted instances to the total instances evaluated, expressed as:

Accuracy = (Number of Correct Predictions / Total Number of Predictions) × 100%

9.1. Training Data Set Accuracy

The accuracy on the training dataset reflects the model's ability to learn from the given input data. It indicates how well the model fits the training samples used during the learning process. A high training accuracy suggests that the model has successfully captured the underlying patterns in the training data. However, an excessively high training accuracy may indicate overfitting, where the model memorizes the training examples but may not generalize well to unseen data.

9.2. Testing Data Set Accuracy

The testing dataset utilized to assess the generalization capability of the model on unseen data. Accuracy on the testing data is a critical indicator of the model's robustness

and practical applicability. The testing accuracy expected to be comparable to the training accuracy; a significant drop suggests overfitting, while low accuracy on both sets indicates underfitting or insufficient learning.

The model's performance was evaluated by calculating accuracy on both the training and testing datasets. Table 1 summarizes the accuracy results.

Table 1: Accuracy results on training and testing datasets.

Dataset	Number Samples	of Correct Predictions	Accuracy (%)
Training Set	1000	950	95.00
Testing Set	1000	880	88.00

The training accuracy of 95% shows that the model has learned the underlying patterns effectively. The testing accuracy of 88% proves good generalization to unseen data, with a slight decrease that suggests the model is not overfitting significantly.



Fig. 6: Comparing Accuracy with a training data set and testing data set

X. DISCUSSION

The results obtained from the experimental evaluation of linear kernel-based Support Vector Machine (SVM) models reveal several important insights into the practical applicability of SVM for real-world binary classification problems.

10.1. Model Accuracy and Generalization

The linear SVM achieved 88% testing accuracy on the heart disease dataset and 76.7% on the social network dataset, showing that the model was able to generalize well to unseen data. The relatively small gap between training and testing accuracy suggests that the models avoided overfitting, despite the limited size of the datasets.

The higher accuracy on the heart disease dataset may be

10.2. Dataset Characteristics and Model Suitability

attributed to the more structured and medically relevant feature space, where attributes like cholesterol, age, and

blood pressure have a clear influence on the outcome. In contrast, the social network advertisement dataset had fewer features (primarily demographic), which may not capture all relevant behavior influencing buying decisions, limiting classification performance.

The results support the suitability of the linear kernel for datasets where classes are reasonably separable and where interpretability is crucial. In both cases, the datasets had relatively low dimensionality, favoring linear decision boundaries.

10.3. Hyperparameter Influence

The study also saw the impact of the regularization parameter C , which governs the trade-off between maximizing margin and minimizing misclassification. Tuning this parameter was essential for improving generalization, especially in the presence of noise or overlapping classes. Smaller C values helped prevent overfitting by allowing softer margins, which proved beneficial particularly for the social dataset.

10.4. Comparison with Literature

Findings align with existing research where linear SVMs have performed effectively in structured classification tasks, particularly in healthcare domains. Compared to models requiring non-linear kernels (e.g., RBF), linear SVMs offer faster training times, reduced computational complexity, and easier interpretability—critical features for medical applications.

10.5. Limitations

Despite its advantages, the linear SVM model struggles when data classes are not linearly separable. Although soft margin classification helps mitigate this, further improvements could be achieved using kernel tricks or non-linear models in such cases. Additionally, the datasets used were binary and relatively small; thus, scalability and performance on multi-class or large-scale datasets were not evaluated.

XI. CONCLUSION

In this study, two classification models were developed using a linear kernel-based Support Vector Machine (SVM) to solve binary classification problems in different domains. The first model focused on predicting the presence of heart disease using a dataset with 14 attributes, where the final attribute classified patients into either the presence or absence of the disease. The second model aimed to classify customer purchase behavior based on demographic and financial features, specifically Age and Salary, from a social dataset.

Both models were implemented in Python, with datasets carefully divided into training and testing subsets to ensure unbiased evaluation. The models were trained in the training data, and their performance was confirmed on the testing data using confusion matrices. The confusion matrices proved that the linear kernel SVM efficiently distinguished between classes, showing high accuracy and robustness in classification tasks. This shows that the linear SVM is well-suited for handling binary classification

problems, especially where the relationship between features and classes is linearly separable or approximately.

The successful application of SVM in these diverse datasets highlights its versatility in healthcare and marketing analytics, offering valuable predictive insights that can aid in early diagnosis and targeted marketing strategies. Additionally, Python's flexibility and extensive libraries enabled efficient development, testing, and evaluation, making it an ideal platform for machine learning workflows.

XII. FUTURE SCOPE

Future enhancements can focus on improving model accuracy by exploring other kernel functions like polynomial or radial basis function (RBF) kernels, which may better capture complex data patterns. Additionally, incorporating more features or combining multiple datasets could enrich the models. Further, applying advanced techniques such as feature choice, hyperparameter tuning, and ensemble learning may boost predictive performance. Expanding the models to multiclass classification or real-time deployment in healthcare and marketing systems also presents promising directions for practical applications.

ACKNOWLEDGEMENTS

The authors would like to express their sincere gratitude to the faculty and management for their continuous support and encouragement throughout the course of this research. We extend our heartfelt thanks to our mentors and colleagues for their insightful feedback and valuable suggestions, which significantly contributed to the improvement of this work.

Special thanks are also due to the researchers and developers whose contributions and publications formed the foundation of our study. Finally, we thank our families and friends for their constant motivation and moral support.

REFERENCES

- [1]. A. Ali and Y. Wang, Hybrid SVM and CNN for medical image classification, *Biomedical Signal Processing and Control*, 85, 2024, 104613.
- [2]. R. Kumar and A. Singh, Comparative analysis of kernel functions in SVM for intrusion detection, *Journal of Network and Computer Applications*, 215, 2024, 103680.
- [3]. Y. Li and Z. Liu, Support vector machine improved with swarm intelligence for financial risk prediction, *Expert Systems with Applications*, 215, 2023, 119347.
- [4]. H. Patel and B. Mehta, Early-stage heart disease detection using linear SVM and PCA, *Journal of King Saud University - Computer and Information Sciences*, 2023.
- [5]. K. Zhou et al., Lightweight SVM model for malware detection on mobile devices, *Computers & Security*, 128, 2023, 103008.
- [6]. M. Chen and J. Sun, Deep feature fusion with SVM for emotion recognition from EEG signals, *Neural Computing and Applications*, 34, 2022, 2381–2393.

- [7]. A. Ghosh et al., SVM and ensemble methods for COVID-19 classification from CT images, *Applied Soft Computing*, 114, 2022, 107789.
- [8]. M. A. Rahman and M. Hasan, Predictive modeling for loan default using SVM: A Bangladesh case study, *Data in Brief*, 42, 2022, 108261.
- [9]. N. Gupta and D. S. Chauhan, Hybrid SVM and decision tree for credit card fraud detection, *Procedia Computer Science*, 200, 2022, 1129–1135.
- [10]. Y. Liu and J. Zhang, Sentiment classification using TF-IDF and linear SVM for short text, *Information Processing & Management*, 58(1), 2021, 102427.
- [11]. R. Agarwal and V. Sharma, Climate prediction using machine learning: A comparative study of SVM and RF, *The Egyptian Journal of Remote Sensing and Space Science*, 24(2), 2021, 349–357.
- [12]. M. R. Islam et al., Classification of mental workload using EEG and SVM, *Sensors*, 21(3), 2021, 1006.
- [13]. J. Han and S. Kim, Real-time traffic sign recognition using optimized SVM, *Pattern Recognition Letters*, 144, 2021, 30–37.
- [14]. H. Zhang and B. Zhou, Application of linear SVM for document classification using n-gram features, *Information Sciences*, 509, 2020, 317–331.
- [15]. M. Sharma and V. K. Tiwari, Lung cancer detection using SVM on medical image features, *International Journal of Imaging Systems and Technology*, 30(1), 2020, 1–10.
- [16]. B. K. Panigrahi and N. Dey, An efficient SVM-based approach for classification of arrhythmia, *Computers in Biology and Medicine*, 123, 2020, 103889.
- [17]. R. Kaur and S. Ahuja, SVM classification of fake news using term frequency-based features, *Procedia Computer Science*, 167, 2020, 1654–1661.
- [18]. A. Jain and A. Tiwari, Comparative study of decision tree, SVM, and Naïve Bayes for crop yield prediction, *Materials Today: Proceedings*, 33, 2020, 5120–5125.
- [19]. A. Ahmed and S. Anwar, Breast cancer classification using SVM with improved feature selection, *Computers in Biology and Medicine*, 127, 2020, 104064.
- [20]. J. Zhou and X. Wang, predicting student academic performance using SVM and neural networks, *Education and Information Technologies*, 25(5), 2020, 4787–4805.
- [21]. V. K. Tiwari, A. Jain, R. Singh, and P. Singh, enhancing outlier detection and dimensionality reduction in machine learning for extreme value, *International Journal of Advanced Networking and Applications*, 15(2), 2024.
- [22]. M. K. Bagwani, V. K. Tiwari, and A. Jain, Implementing GrapesJS in educational platforms for web development training on AWS, *International Journal of Scientific Research in Multidisciplinary Studies*, 10(8), 2024, 1–8.
- [23]. M. K. Bagwani, D. K. Chouhan, A. Jain, and V. K. Tiwari, deploying a web application on AWS Amplify: A comprehensive guide, *International Journal of Progressive Research in Engineering Management and Science*, 4(8), 2024, 1570–1574.
- [24]. P. Singh, S. Rahangdale, V. K. Tiwari, and G. K. Saxena, Enhanced machine learning technique for predicting cardiac diseases using ECG data, *International Journal of Scientific Research & Engineering Trends*, 10(4), 2024, 1274–1278.
- [25]. M. K. Bagwani, A. Gangwar, K. Vishwakarma, and V. K. Tiwari, Real-time signature-based detection, and prevention of DDOS attacks in cloud environments, *International Journal of Science and Research Archive*, 12(2), 2024, 2929–2935.
- [26]. J. Mehra, B. Pillai, N. Singh, S. Joshi, V. K. Tiwari, and P. K. Saket, An antenna with mechanical frequency adaptability in a cavity for multiband and Internet of Things applications, *International Journal of All Research Education and Scientific Methods*, 12(5), 2024, 4544–4552.
- [27]. N. Singh, S. Jaiswar, P. Jha, V. K. Tiwari, and P. K. Saket, Adaptive intrusion detection using deep reinforcement learning: A novel approach, *International Journal of All Research Education and Scientific Methods*, 12(5), 2024, 4158–4163.
- [28]. P. K. Saket, D. K. Chouhan, A. Jain, and V. K. Tiwari, Enhanced data security with an extended MSA-based symmetric key encryption algorithm, *International Journal of All Research Education and Scientific Methods*, 12(5), 2024.

Biographies and Photographs



Prof. Dr. Sanjay Thakur has received a Ph.D. degree in Computer Science. He is currently head with Department of Computer Science, Chameli Devi Group of Institutions, Indore, INDIA



Dr. Virendra Kumar Tiwari is a Professor and Head of the Department of Computer Applications at Lakshmi Narain College of Technology (MCA), Bhopal. He holds a B.Sc., M.A. (Economics), MCA, and Ph.D. from Dr. Hari Singh Gour University, Sagar, Madhya Pradesh. With over 17 years of academic and research experience, his areas of specialization include Computer Networks and Stochastic Modelling. Dr Tiwari has published 22 research papers in reputed national and international journals. He is actively involved in guiding students and research scholars, and his efforts continue to strengthen the academic and research environment of his institution.



Dr. Jitendra Agrawal is an Associate Professor in the MCA Department at Lakshmi Narain College of Technology (LNCT), Bhopal, with over 13 years of teaching experience. He holds an MCA from RGPV, Bhopal, and a B.Sc. from Dr. H.S. Gour Vishwavidyalaya, Sagar. His ability lies in computer applications, with a focus on advancing teaching methodologies. He is committed to fostering student development and enhancing educational outcomes through technology.