

A Machine Learning-Based Incremental Majority Voting Method for Intrusion Detection Using Stacking Classifiers

Anwary Abdul Rehman
Mohammad Aslam
PG Scholar, Department of
Computer Science and
Engineering, Terna College of
Engineering, Osmanabad
(Dharashiv), Maharashtra,
INDIA - 413501.

Email:

anwaryar22@gmail.com

Sujata Gaikwad
Professor, Department of
Computer Science and
Engineering, Terna College of
Engineering, Osmanabad
(Dharashiv), Maharashtra,
INDIA - 413501.

Sujata Gaikwad
Head Of Department,
Department of Computer
Science and Engineering,
Terna College of Engineering,
Osmanabad (Dharashiv),
Maharashtra, INDIA - 413501.

Abstract: With the rapid growth of digitalization and the increasing volume of data, the cybersecurity threat landscape is expanding at an alarming rate. Intrusion Detection Systems (IDS) have become crucial in conjunction with firewalls to safeguard networks from malicious activities. In this work, four well-known cybersecurity datasets—CIC IDS 2017, NSL KDD, KDD Cup, and CIC IDS 2018—are employed to evaluate the effectiveness of various techniques for intrusion detection. Feature selection is performed using Mutual Information to enhance the relevance of selected features. Data sampling techniques are also explored, including Original Data, Random Under Sampling, Random Over Sampling, and a combination of both under and over-sampling to address data imbalance. To further improve the detection performance, a refined approach utilizing a Stacking Classifier combining Random Forest (RF) and Decision Tree (DT) with a Bagging Classifier is implemented. The results show that this approach achieves high performance across all datasets and sampling techniques, demonstrating its effectiveness in accurately detecting network intrusions in dynamic cybersecurity environments.

“Index Terms –Incremental learning, network intrusion detection, machine learning, majority voting classifier, random sampling, Stacking classifier, Cyber Security”.

1. INTRODUCTION

Intrusion detection plays a pivotal role in the current cybersecurity landscape as the number of evolving attacks, such as Distributed Denial of Service (DDoS), ransomware, and advanced persistent threats (APTs), continues to grow on a daily basis [1], [2], [3]. These attacks have become more sophisticated, causing significant damage to organizations' digital infrastructure, financial systems, and sensitive data. As the cyber threat landscape evolves, traditional security mechanisms often struggle to cope with these new and adaptive threats. Therefore, security systems are in dire need of robust components that can effectively prevent potential attacks and safeguard network integrity [4]. Intrusion Detection Systems (IDS) have emerged as indispensable tools for protecting systems against unauthorized access and mitigating the risks associated with malicious activities, such as unauthorized network access and data exfiltration [5]. IDS are critical in ensuring that systems remain protected, continuously monitoring network traffic and identifying anomalies that may indicate an ongoing attack.

To address these emerging threats, the field of intrusion detection has seen numerous studies employing both machine learning and deep learning approaches [6], [7], [8]. These studies have demonstrated favorable results in detecting a wide range of attack types, from common malware and phishing attacks to more sophisticated threats like zero-day exploits and APTs. Machine learning models, such as decision trees, support vector machines, and neural networks, have shown promise in identifying malicious activity by analyzing network traffic, user behaviors, and other system logs. Additionally, deep learning models, particularly those based on neural networks, have been used to extract high-level features from large-scale datasets and improve the classification of malicious activity. While these approaches offer enhanced detection capabilities, they also face significant challenges. Firstly, most existing IDS systems are static, meaning they lack the ability to adapt and learn in real-time. This hinders their ability to detect new, previously unseen attack vectors, and necessitates expensive and time-consuming retraining processes to keep up with the evolving threat landscape.

Another critical issue with current IDS is their reliance on large, labeled datasets for training, which can be both labor-intensive and costly to obtain. Additionally, these datasets often require significant storage space, creating logistical challenges for organizations with limited resources. The need for such vast datasets can delay the deployment of effective intrusion detection systems, especially in environments with rapidly changing data distributions and attack patterns. These challenges emphasize the need for dynamic, resource-efficient intrusion detection approaches capable of handling emerging threats without requiring extensive labeling efforts or large-scale data storage. In this context, incremental learning approaches have gained significant attention [10], [11] due to their ability to continuously learn and adapt without the need for a complete, ready-to-use dataset [12]. Incremental learning methods enable models to be initially trained with a small amount of data and updated as new data becomes available, offering a more flexible solution to the challenges of traditional IDS systems. This adaptability makes incremental learning particularly suitable for dynamic environments where both the nature of the data and the threats are continuously evolving.

2. RELATED WORK

A significant body of research has focused on advancing intrusion detection systems (IDS) to address the increasing complexity and volume of cybersecurity threats. M. Data and M. Aritsugi [13] proposed AB-HT, an ensemble incremental learning algorithm for IDS, which is designed to dynamically update and improve intrusion detection models without the need for retraining from scratch. The key feature of AB-HT is its ability to learn continuously from incoming data, making it adaptable to new and evolving attack patterns. This method integrates incremental learning with ensemble techniques, enhancing the system's detection capabilities while reducing computational costs and storage requirements.

B. A. Tama and S. Lim [14] conducted a systematic mapping study to explore the role of ensemble learning in IDS. Their study highlights the potential of combining multiple classifiers to improve detection accuracy and robustness. By evaluating various ensemble methods across different benchmark datasets, they emphasize the effectiveness of models such as Random Forest and Adaboost, which improve the generalization capability of IDS by reducing bias and variance. This approach also offers resilience against overfitting, a common issue in cybersecurity applications where attacks often exhibit high variability.

M. Torabi et al. [15] reviewed feature selection and ensemble techniques for IDS, underscoring the importance of identifying relevant features to optimize the performance of machine learning models. They discuss how feature selection methods, such as Information Gain and Chi-square, can be paired with ensemble learning algorithms like Random Forest and Bagging, leading to improved detection rates. The paper suggests that by reducing dimensionality, feature selection not only enhances computational efficiency but also improves the robustness of the detection system by focusing on the most influential attributes in attack data.

A. M. Bamhdi, I. Abrar, and F. Masoodi [16] introduced an ensemble-based approach using majority voting for effective intrusion detection. In this framework, multiple classifiers are combined, and the final decision is made based on the majority vote from the individual models. The study demonstrates that this technique can achieve superior performance by leveraging the strengths of various classifiers such as K-Nearest Neighbors (KNN), Decision Trees, and Support Vector Machines (SVM). Majority voting helps mitigate the limitations of individual models, particularly in terms of false positive rates and detection accuracy.

D. R. Patil and T. M. Pattewar [17] proposed a majority voting and feature selection-based IDS, which integrates a feature selection process with a voting mechanism for classification. Their approach focuses on selecting a minimal but highly informative set of features to train a variety of classifiers, followed by a voting strategy to make the final decision. This hybrid method aims to improve the scalability and adaptability of IDS by reducing the data and model complexity while maintaining high detection performance. It also enhances real-time learning capabilities, a key requirement for IDS to handle new types of attacks without large retraining efforts.

H. Xu and Y. Wang [18] introduced a continual few-shot learning method for IDS using meta-learning techniques. This approach is particularly useful in environments where labeled data is scarce or constantly evolving. The system learns to detect novel attacks by leveraging a small number of labeled examples, significantly reducing the need for large datasets. Meta-learning helps the system generalize better to new, unseen attack types by learning from a small number of training instances, thereby addressing the challenge of data scarcity in real-time intrusion detection.

T. Wang et al. [19] proposed a few-shot class-incremental learning approach for IDS, focusing on the ability of the model to learn and adapt

incrementally from small sets of labeled data. This approach is particularly beneficial for detecting emerging attack types where there is little available training data. The system can continuously update itself as new attack data becomes available, ensuring that it stays current with evolving threat landscapes. This methodology enhances the flexibility and scalability of IDS, making it suitable for real-world applications where the nature of threats constantly changes.

J. Zheng et al. [20] developed an ensemble learning-based two-level network intrusion detection method, which integrates multiple classifiers in two stages to improve detection accuracy. The first level consists of a set of base classifiers that operate independently to make initial predictions, while the second level involves combining these predictions through a meta-classifier to make the final decision. This layered approach boosts performance by incorporating diversity in the classifiers and making the final model more robust to different attack types.

In summary, these studies contribute significantly to the development of more adaptable, accurate, and efficient intrusion detection systems. They highlight the importance of incremental learning, ensemble techniques, feature selection, and few-shot learning for enhancing IDS performance in dynamic and evolving cybersecurity environments. The combination of these approaches provides a promising direction for creating IDS that can effectively handle the increasing complexity of modern cyber threats.

3. MATERIALS AND METHODS

The proposed system aims to enhance intrusion detection capabilities by evaluating various machine learning algorithms on well-known cybersecurity datasets, including CIC IDS 2017 [17], NSL KDD [21], KDD Cup [17], and CIC IDS 2018 [8]. Feature selection [15] is carried out using Mutual Information to ensure the most relevant features are utilized for training. Data sampling techniques, such as Original Data, Random Under Sampling, Random Over Sampling, and a combination of under and over-sampling, are applied to address class imbalance. The system incorporates several machine learning algorithms, including K-Nearest Neighbors (KNN) [6], Softmax Logistic Regression (LR) [7], Random Forest (RF) [8], HAT/Decision Tree (DT), and a Voting Classifier that combines KNN, LR, RF, and DT to boost classification performance. Additionally, a Stacking Classifier combining RF and DT with a Bagging Classifier is used to improve predictive accuracy. This comprehensive approach aims to optimize network intrusion

detection across various datasets and sampling techniques, providing a robust solution for cybersecurity.

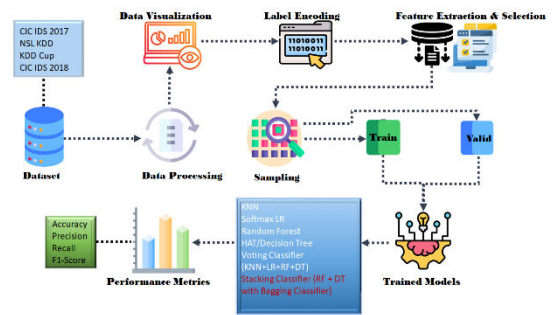


Fig.1 Proposed Architecture

This system (fig. 1) processes four datasets (CIC IDS 2017, NSL-KDD, KDD Cup, and CIC IDS 2018) through data cleaning and visualization. Label encoding prepares data for feature extraction and selection. Four sampling techniques create training and testing sets. Models (KNN, Softmax LR, RF, DT, Voting Classifier and Stacking Classifier) are trained and evaluated using performance metrics (Accuracy, Precision, Recall, F1-Score).

A) Dataset Collection:

a) CIC IDS 2017:

The dataset used for this project is the CIC IDS 2017 [17], which contains 2,044,217 entries with 78 features, capturing network traffic data such as packet lengths, flow statistics, and flag counts. After feature selection, the dataset was reduced to 44,697 entries and 20 relevant features. The "Label" column indicates whether a sample corresponds to a specific type of intrusion or attack, making it suitable for intrusion detection system (IDS) analysis.

Protocol	Flow Duration	Total Fwd Packets	Total Backward Packets	Fwd Packets Length Total	Bwd Packets Length Total	Fwd Packet Length Max	Fwd Packet Length Min	Fwd Packet Length Mean
0	6	4	2	0	12	0	6	6.00000
1	6	1	2	0	12	0	6	6.00000
2	6	3	2	0	12	0	6	6.00000
3	6	1	2	0	12	0	6	6.00000
4	6	609	7	4	484	414	233	69.14286

5 rows x 78 columns

Fig.2 Dataset Collection Table - CICIDS2017

b) NSL-KDD:

The NSL-KDD [21] dataset, containing 125,972 entries with 43 features, is a benchmark dataset for intrusion detection systems. It includes features such as protocol type, service, flag, source and destination bytes, and various traffic statistics. After feature selection, the dataset is reduced to

26,047 entries with 11 significant features, including "attack" as the target label. It is widely used for evaluating anomaly detection and network intrusion classification methods.

	duration	protocol_type	service	flag	src_bytes	dst_bytes	land	wrong_fragment	urgent	hot	...	dst_host_same_srv_rate
0	0	udp	other	SF	146	0	0	0	0	0	...	0.00
1	0	tcp	private	S0	0	0	0	0	0	0	...	0.10
2	0	tcp	http	SF	232	8153	0	0	0	0	...	1.00
3	0	tcp	http	SF	199	420	0	0	0	0	...	1.00
4	0	tcp	private	REJ	0	0	0	0	0	0	...	0.07

5 rows x 43 columns

Fig.3 Dataset Collection Table – NSL-KDD

c) KDD Cup:

The KDD Cup [17] dataset consists of 125,973 entries with 42 features, used for intrusion detection in networks. It includes attributes like protocol type, service, flag, and traffic details. After feature selection, it is reduced to 26,047 entries with 11 critical features, including "labels" as the target. This dataset is widely employed for evaluating machine learning models in anomaly detection and intrusion prevention systems.

	duration	protocol_type	service	flag	src_bytes	dst_bytes	land	wrong_fragment	urgent	hot	...	dst_host_srv_count
0	0	tcp	ftp_data	SF	491	0	0	0	0	0	...	25
1	0	udp	other	SF	146	0	0	0	0	0	...	1
2	0	tcp	private	S0	0	0	0	0	0	0	...	26
3	0	tcp	http	SF	232	8153	0	0	0	0	...	255
4	0	tcp	http	SF	199	420	0	0	0	0	...	255

5 rows x 42 columns

Fig.4 Dataset Collection Table - KDDCUP

d) CIC IDS 2018:

The CIC IDS 2018 dataset [8] contains 3,550,129 entries and 78 features, designed for analyzing network traffic and detecting intrusions. Key features include packet lengths, flow durations, and flag counts. After feature selection, it is reduced to 30,594 entries with 20 essential features, including "Label" as the target. This dataset is widely used in machine learning research for building robust intrusion detection systems, focusing on performance optimization and attack pattern recognition.

	Protocol	Flow Duration	Total Fwd Packets	Total Backward Packets	Fwd Packets Length Total	Bwd Packets Length Total	Fwd Packet Length Max	Fwd Packet Length Min	Fwd Packet Length Mean
0	6	141385	9	7	553	3773.0	202	0	61.444443
1	6	281	2	1	38	0.0	38	0	19.000000
2	6	279624	11	15	1086	10527.0	385	0	96.727272
3	6	132	2	0	0	0.0	0	0	0.000000
4	6	274016	9	13	1285	6141.0	517	0	142.777771

5 rows x 78 columns

Fig.5 Dataset Collection Table - CICIDS2018

B) Pre-Processing:

In the pre-processing step, we focus on preparing the dataset for modeling. This includes cleaning the data, visualizing key relationships, encoding

categorical labels, performing feature extractionandsampling techniques to ensure high-quality input for the prediction model.

a) **Data Processing:** The data preprocessing involves handling missing and duplicate entries. Initially, the dataset is assessed for null values, which are then removed to ensure the data remains clean and complete. Duplicate entries are identified to eliminate redundancy and prevent bias in analysis. By carefully addressing both issues, the dataset is refined for consistency and accuracy, supporting more reliable and efficient machine learning model training. These steps enhance data quality, ensuring robust outcomes in subsequent analysis and predictions.

b) **Data Visualization:** The data visualization focuses on analyzing the distribution of the target variable, providing insights into the balance between the classes within the dataset. By presenting the counts of each class, it helps identify any potential class imbalance, which is critical for ensuring fair and effective machine learning model training. Understanding the representation of each target class aids in evaluating the dataset's characteristics, guiding strategies such as resampling or adjusting evaluation metrics to enhance the model's performance on imbalanced data.

c) **Label Encoding:** Label encoding is applied to transform categorical variables into numerical values, making them suitable for machine learning algorithms. In this process, the target variable, along with other categorical features like protocol type, service, and flag, are converted into integer representations. This technique assigns a unique integer to each category, facilitating the model's ability to handle these variables. By encoding categorical data, the dataset is prepared for more efficient processing, improving model compatibility and performance, especially when dealing with algorithms that require numerical inputs.

d) **Feature Extraction:** Feature extraction involves selecting the relevant input data (X) and the target variable (y) from the dataset. In this case, the features are extracted by removing the 'labels' column from the data, which is the target variable, while the remaining columns are stored as input features (X). The target variable (y) is then set as the 'labels' column, which the model will predict.

e) **Feature Selection:** Feature selection using mutual information helps identify the most relevant features for the prediction task. In this process [15], the mutual information classifier evaluates the relationship between each feature and the target

variable. By using the SelectPercentile method, the top 25% of features with the highest mutual information scores are selected. The selected features are then transformed into a reduced dataset, and the relevant columns are extracted and listed, ensuring that only the most informative features are retained for model training.

f) Sampling the data: Sampling the data involves modifying the dataset to address class imbalances. The original data is examined first to assess the distribution. Random under-sampling reduces the majority class by randomly removing instances, while random over-sampling increases the minority class by duplicating instances. Combining both techniques creates a balanced dataset by adjusting both classes, ensuring that the model receives an equal representation of each class. This approach helps improve the model's ability to learn from both classes without being biased toward the majority class.

C) Training & Testing:

The dataset is divided into training and testing sets to evaluate the model's performance. A portion of the data, typically 20%, is reserved for testing, while the remaining 80% is used for training the model. This split ensures that the model learns patterns from the majority of the data while being tested on unseen examples to assess its generalization capability. The random state is fixed to maintain consistency in the data partitioning across different runs of the model.

D) Algorithms:

KNN: K-Nearest Neighbors is applied to classify data based on the proximity of feature values to labeled examples. It [6] helps detect patterns and classify instances by majority voting from the nearest neighbors.

$$y = \arg \max_c \sum_{i \in N} 1(y_i = c)$$

Softmax LR: Softmax Logistic Regression is used to handle multiclass classification [7], transforming logits into probability distributions, enabling the model to predict the likelihood of each class.

$$P(y = k | x) = \frac{e^{w_k^T x + b_k}}{\sum_{j=1}^K e^{w_j^T x + b_j}}$$

Random Forest: A decision tree [8] ensemble method that aggregates multiple trees to improve accuracy, robustness, and handle overfitting. It is

used for reliable classification and feature importance evaluation.

$$y^{\wedge} = \arg \max_k \sum_{t=1}^T 1(h_t(x) = k)$$

HAT/DecisionTree: Decision Tree or Hierarchical Agglomerative Tree (HAT) is utilized for classification tasks, providing an interpretable, tree-based model that splits data based on feature thresholds.

$$H(X) = -\sum p(x) \log_2 p(x)$$

Voting Classifier: A combination of KNN, LR, RF, and DT, where each model votes and the class with the most votes is selected. This ensemble approach improves classification performance by combining strengths of different models.

$$y^{\wedge} = \arg \max_c \sum_{i=1}^N w_i 1(h_i(x) = c)$$

Stacking Classifier: A meta-model built from RF and DT with Bagging Classifier, which leverages predictions from base learners to improve accuracy. It allows the model to correct biases and errors made by individual classifiers.

$$y^{\wedge} = H\left(\sum_{i=1}^n w_i h_i(X)\right)$$

4. RESULTS & DISCUSSION

Accuracy: The accuracy of a test is its ability to differentiate the patient and healthy cases correctly. To estimate the accuracy of a test, we should calculate the proportion of true positive and true negative in all evaluated cases. Mathematically, this can be stated as:

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (1)$$

Precision: Precision evaluates the fraction of correctly classified instances or samples among the ones classified as positives. Thus, the formula to calculate the precision is given by:

$$Precision = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}} \quad (2)$$

Recall: Recall is a metric in machine learning that measures the ability of a model to identify all

relevant instances of a particular class. It is the ratio of correctly predicted positive observations to the total actual positives, providing insights into a model's completeness in capturing instances of a given class.

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

F1-Score: F1 score is a machine learning evaluation metric that measures a model's accuracy. It combines the precision and recall scores of a model. The accuracy metric computes how many times a model made a correct prediction across the entire dataset.

$$F1\ Score = 2 * \frac{Recall \times Precision}{Recall + Precision} * 100 \quad (4)$$

The below **tables (1.1 to 1.4)** present the performance evaluation of various machine learning models on the CICIDS2017 dataset and multiple sampling techniques. The results show that the Stacking Classifier achieved the highest accuracy across all sampling techniques.

Table.1.1 Performance Evaluation for Original Data – CICIDS2017

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.920	0.921	0.920	0.920
Softmax LR	0.304	0.661	0.304	0.404
RandomForest	0.977	0.984	0.977	0.979
HAT/DecisionTree	0.970	0.975	0.970	0.972
Majority Voting	0.972	0.976	0.972	0.974
Stacking Classifier	0.979	0.992	0.979	0.984

Table.1.2 Performance Evaluation for Random UnderSampling – CICIDS2017

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.795	0.797	0.795	0.796
Softmax LR	0.331	0.653	0.331	0.424
RandomForest	0.866	0.868	0.866	0.867
HAT/DecisionTree	0.840	0.930	0.840	0.864
Majority Voting	0.863	0.874	0.863	0.866
Stacking Classifier	0.913	0.929	0.913	0.915

Table.1.3 Performance Evaluation for Random Over Sampling – CICIDS2017

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.958	0.958	0.958	0.958
Softmax LR	0.198	0.620	0.198	0.275
RandomForest	0.938	0.951	0.938	0.939
HAT/DecisionTree	0.941	0.947	0.941	0.941
Majority Voting	0.942	0.954	0.942	0.943
Stacking Classifier	0.952	0.957	0.952	0.952

Table.1.4 Performance Evaluation for Random Under/Over Sampling (Combine) – CICIDS2017

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.993	0.993	0.993	0.993
Softmax LR	0.331	0.676	0.331	0.404
RandomForest	0.989	0.989	0.989	0.989
HAT/DecisionTree	0.960	0.966	0.960	0.961

Majority Voting	0.995	0.995	0.995	0.995
Stacking Classifier	0.994	0.994	0.994	0.994

In tables (2.1 to 2.4), the performance analysis of multiple machine learning models on the CICIDS2018 dataset. The findings indicate that the Stacking Classifier consistently achieved the highest accuracy across all data sampling methods.

Table.2.1 Performance Evaluation for Original Data – CICIDS2018

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.903	0.903	0.903	0.903
Softmax LR	0.470	0.650	0.470	0.529
RandomForest	0.940	0.944	0.940	0.941
HAT/DecisionTree	0.930	0.931	0.930	0.930
Majority Voting	0.929	0.935	0.929	0.929
Stacking Classifier	0.948	0.950	0.948	0.948

Table.2.2 Performance Evaluation for Random UnderSampling – CICIDS2018

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.784	0.804	0.784	0.783
Softmax LR	0.392	0.653	0.392	0.434
RandomForest	0.922	0.938	0.922	0.921
HAT/DecisionTree	0.647	0.850	0.647	0.664
Majority Voting	0.863	0.901	0.863	0.862
Stacking Classifier	0.961	0.964	0.961	0.960

Table.2.3 Performance Evaluation for Random Over Sampling – CICIDS2018

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.939	0.939	0.939	0.939
Softmax LR	0.423	0.672	0.423	0.471
RandomForest	0.959	0.962	0.959	0.959
HAT/DecisionTree	0.867	0.896	0.867	0.867
Majority Voting	0.954	0.957	0.954	0.954
Stacking Classifier	0.953	0.954	0.953	0.953

Table.2.4 Performance Evaluation for Random Under/Over Sampling (Combine) – CICIDS2018

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.995	0.995	0.995	0.995
Softmax LR	0.429	0.568	0.429	0.453
RandomForest	0.997	0.997	0.997	0.997
HAT/DecisionTree	0.996	0.996	0.996	0.996
Majority Voting	0.996	0.996	0.996	0.996
Stacking Classifier	0.996	0.996	0.996	0.996

The tables (3.1 to 3.4) presents the performance evaluation of several machine learning models on the KDD Cup dataset. The results highlight that the Stacking Classifier achieved the highest accuracy across all data sampling strategies.

Table.3.1 Performance Evaluation for Original Data – KDDCUP

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.981	0.981	0.981	0.981
Softmax LR	0.730	0.937	0.730	0.820
RandomForest	0.995	0.995	0.995	0.995
HAT/DecisionTree	0.986	0.986	0.986	0.986
Majority Voting	0.991	0.991	0.991	0.991
Stacking Classifier	0.997	0.998	0.997	0.997

Table.3.2 Performance Evaluation for Random UnderSampling – KDDCUP

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.846	0.851	0.846	0.848
Softmax LR	0.404	0.942	0.404	0.530
RandomForest	0.942	0.946	0.942	0.942
HAT/DecisionTree	0.865	0.887	0.865	0.867
Majority Voting	0.923	0.929	0.923	0.924
Stacking Classifier	0.981	0.983	0.981	0.981

Table.3.3 Performance Evaluation for Random Over Sampling – KDDCUP

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.989	0.989	0.989	0.989
Softmax LR	0.372	0.604	0.372	0.447
RandomForest	0.998	0.998	0.998	0.998
HAT/DecisionTree	0.994	0.994	0.994	0.994
Majority Voting	0.996	0.996	0.996	0.996
Stacking Classifier	0.999	0.999	0.999	0.999

Table.3.4 Performance Evaluation for Random Under/Over Sampling (Combine) – KDDCUP

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.997	0.997	0.997	0.997
Softmax LR	0.386	0.584	0.386	0.463
RandomForest	0.999	0.999	0.999	0.999
HAT/DecisionTree	0.992	0.992	0.992	0.992
Majority Voting	0.998	0.998	0.998	0.998
Stacking Classifier	1.000	1.000	1.000	1.000

The *tables (4.1 to 4.4)* display the performance assessment of various machine learning models on the NSL-KDD dataset. The findings indicate that the Stacking Classifier delivered the highest accuracy across all sampling methods.

Table.4.1 Performance Evaluation for Original Data – NSL-KDD

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.981	0.981	0.981	0.981
Softmax LR	0.712	0.916	0.712	0.800
RandomForest	0.994	0.994	0.994	0.994
HAT/DecisionTree	0.988	0.988	0.988	0.988
Majority Voting	0.988	0.989	0.988	0.989
Stacking Classifier	0.993	0.993	0.993	0.993

Table.4.2 Performance Evaluation for Random UnderSampling – NSL-KDD

ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.846	0.843	0.846	0.841
Softmax LR	0.558	0.692	0.558	0.587
RandomForest	0.923	0.928	0.923	0.923
HAT/DecisionTree	0.981	0.982	0.981	0.981
Majority Voting	0.885	0.883	0.885	0.883
Stacking Classifier	0.904	0.909	0.904	0.903

Table.4.3 Performance Evaluation for Random Over Sampling – NSL-KDD

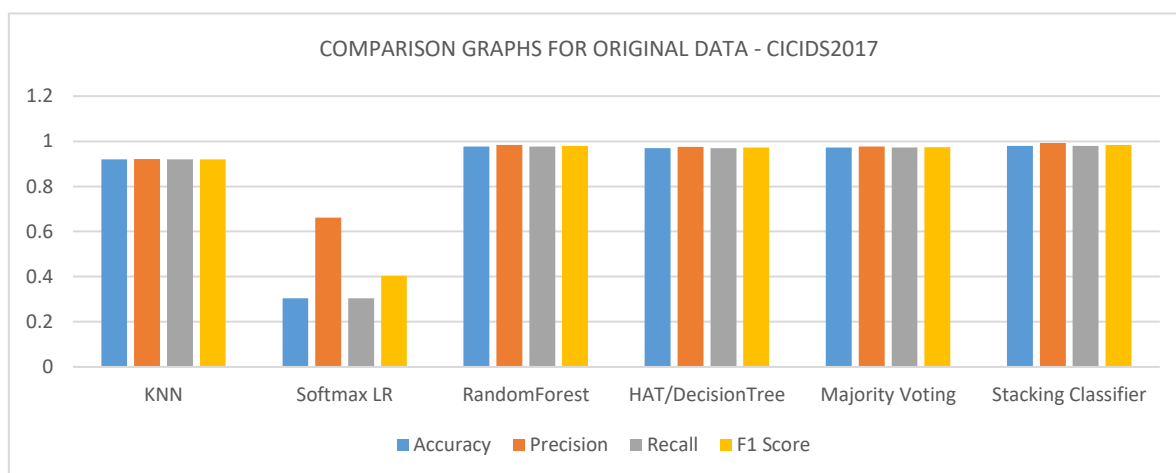
ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.989	0.989	0.989	0.989
Softmax LR	0.396	0.652	0.396	0.486
RandomForest	0.998	0.998	0.998	0.998
HAT/DecisionTree	0.991	0.991	0.991	0.991
Majority Voting	0.995	0.995	0.995	0.995
Stacking Classifier	0.999	0.999	0.999	0.999

Table.4.4 Performance Evaluation for Random Under/Over Sampling (Combine) – NSL-KDD

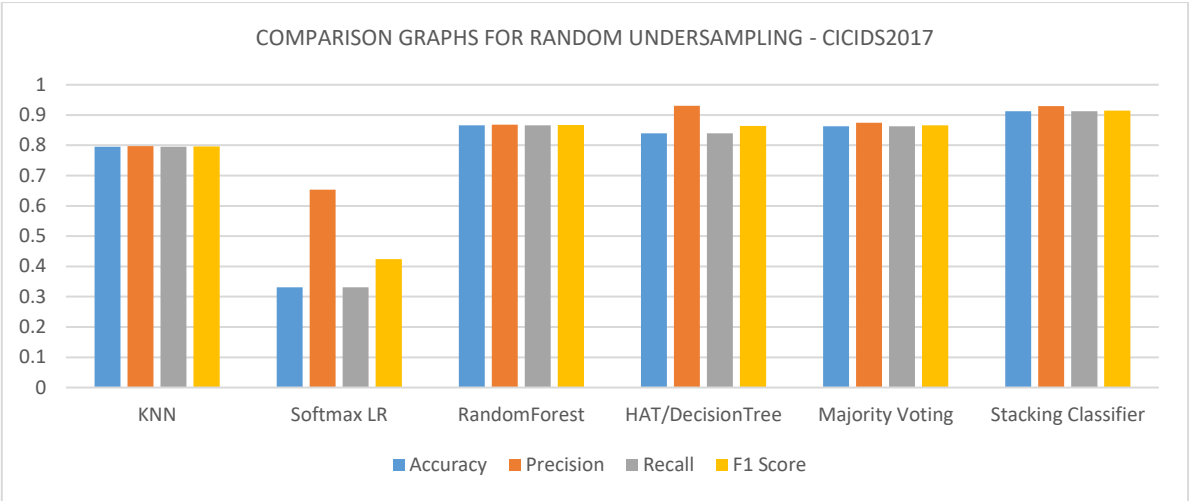
ML Model	Accuracy	Precision	Recall	F1 Score
KNN	0.998	0.998	0.998	0.998
Softmax LR	0.403	0.617	0.403	0.485
RandomForest	0.998	0.998	0.998	0.998
HAT/DecisionTree	0.995	0.995	0.995	0.995
Majority Voting	0.997	0.997	0.997	0.997
Stacking Classifier	1.000	1.000	1.000	1.000

The accuracy is represented in blue, precision in orange, recall in gray, and F1-score in yellow in *Graphs (1.1 to 1.4)*. Compared to other models, the Stacking Classifier demonstrates superior performance across all sampling techniques, achieving the highest values. These results are visually represented in the graphs above, which are based on the CICIDS 2017 dataset..

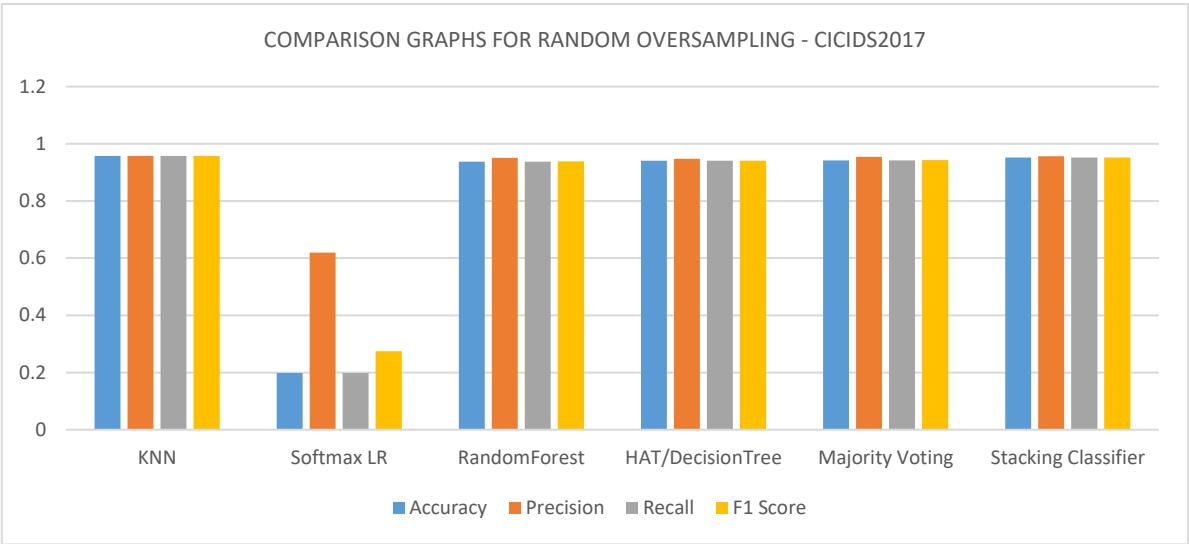
Graph.1.1 Comparison Graphs for Original Data – CICIDS2017



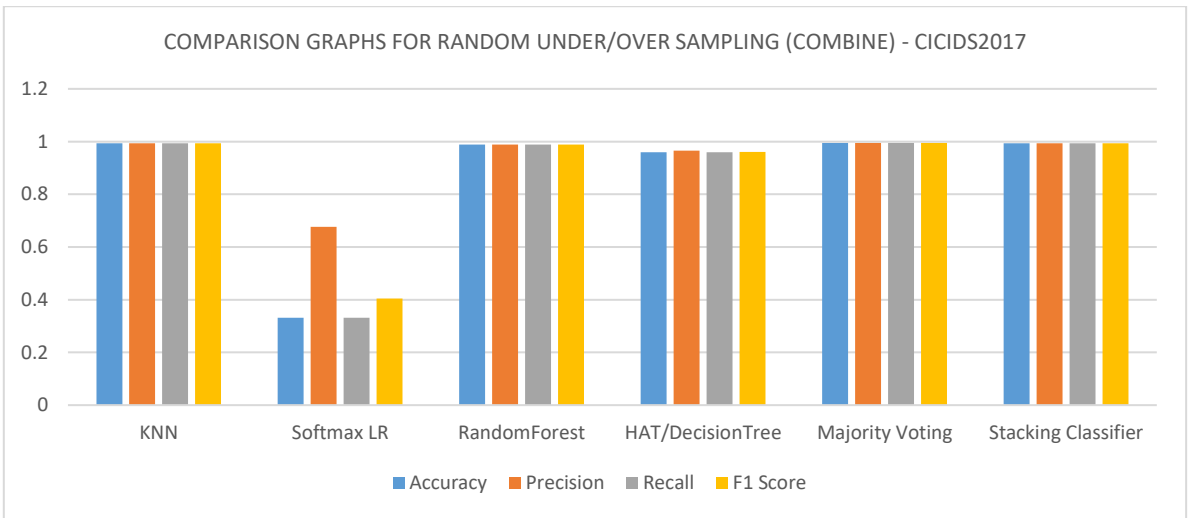
Graph.1.2 Comparison Graphs for Random UnderSampling – CICIDS2017



Graph.1.3 Comparison Graphs for Random Over Sampling – CICIDS2017



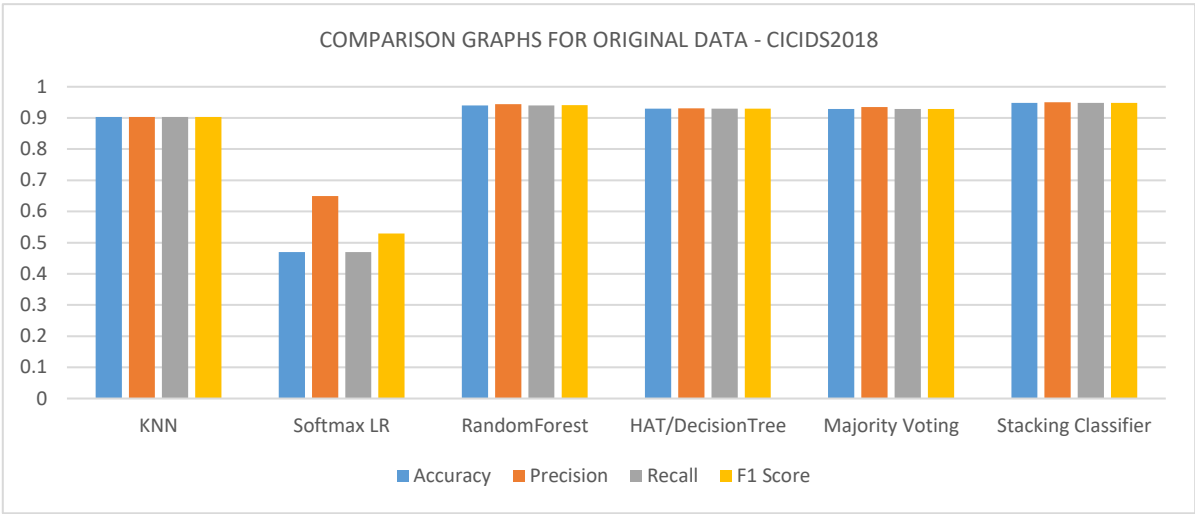
Graph.1.4 Comparison Graphs for Random Under/Over Sampling (Combine) – CICIDS2017



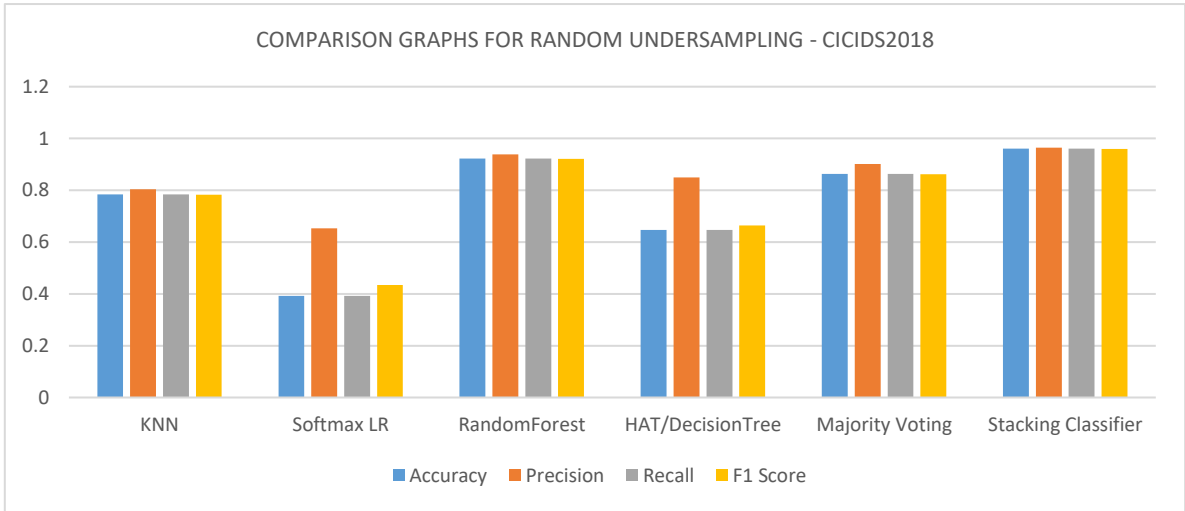
In Graphs (2.1 to 2.4), accuracy is depicted in blue, precision in orange, recall in gray, and F1-score in yellow. The Stacking Classifier outperforms all other models across every sampling technique, consistently achieving

the highest metrics. These outcomes are illustrated in the graphs above, which correspond to the CICIDS 2018 dataset.

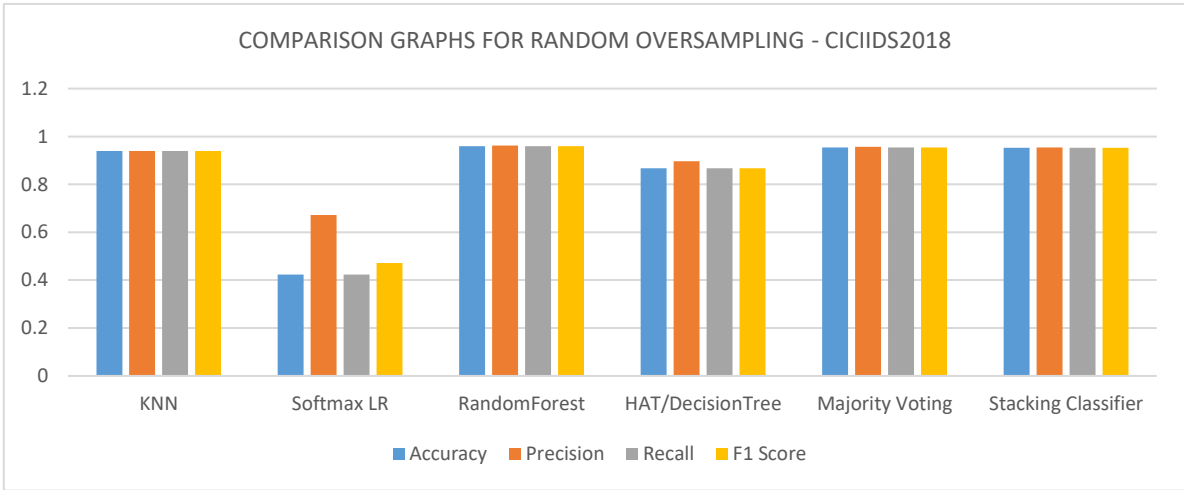
Graph.2.1 Comparison Graphs for Original Data – CICIDS2018



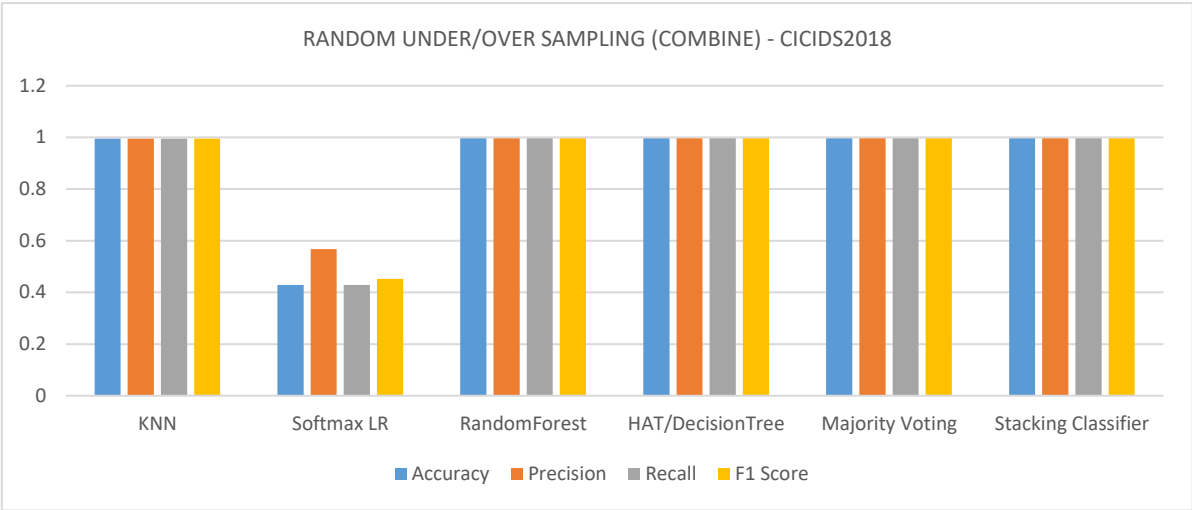
Graph.2.2 Comparison Graphs for Random UnderSampling – CICIDS2018



Graph.2.3 Comparison Graphs for Random Over Sampling – CICIDS2018

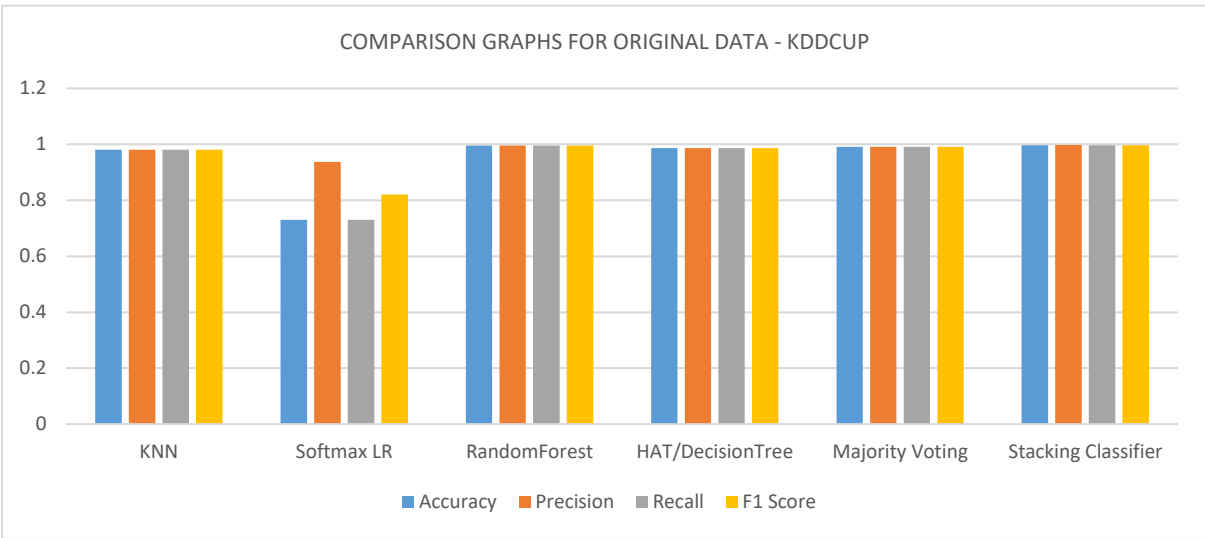


Graph.2.4 Comparison Graphs for Random Under/Over Sampling (Combine) – CICIDS2018

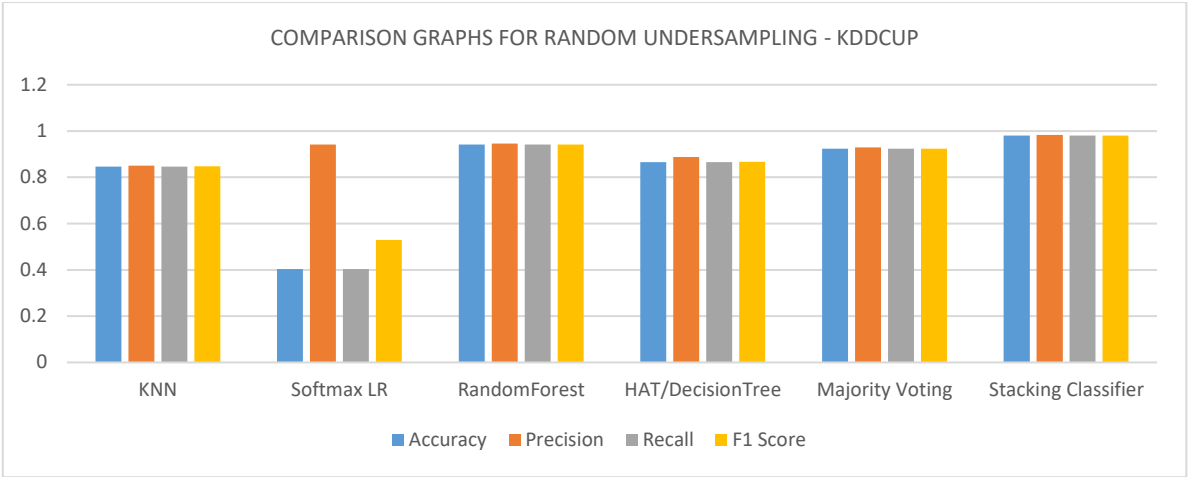


In *(Graphs 3.1 to 3.4)*, the metrics are represented with accuracy in blue, precision in orange, recall in gray, and F1-score in yellow. Across all sampling techniques, the Stacking Classifier consistently achieves the best performance, surpassing the other models. These results are clearly shown in the graphs, which are based on the KDD Cup dataset.

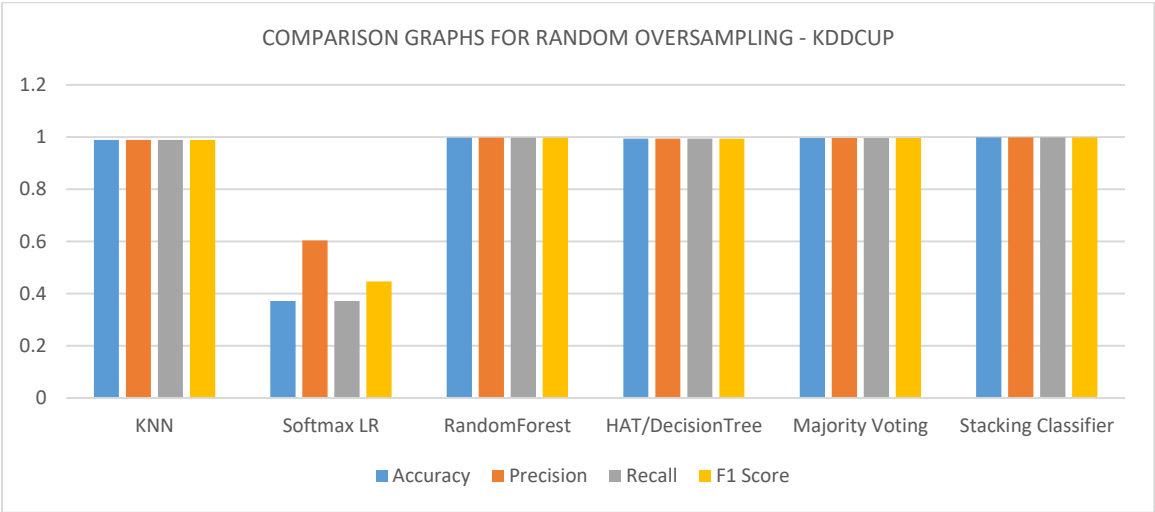
Graph.3.1 Comparison Graphs for Original Data – KDDCUP



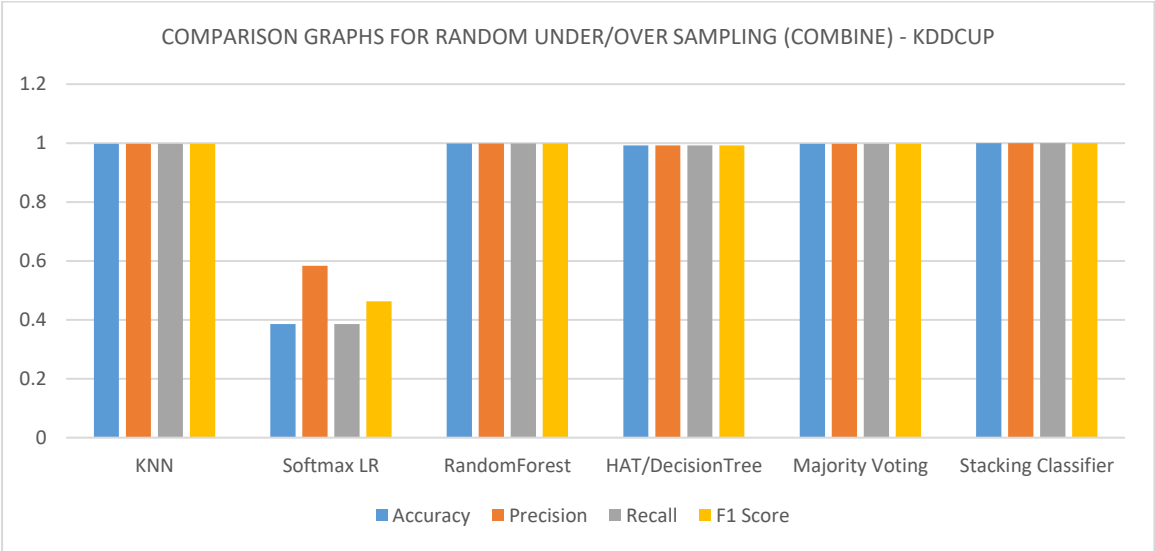
Graph.3.2 Comparison Graphs for Random UnderSampling – KDDCUP



Graph.3.3 Comparison Graphs for Random Over Sampling – KDDCUP

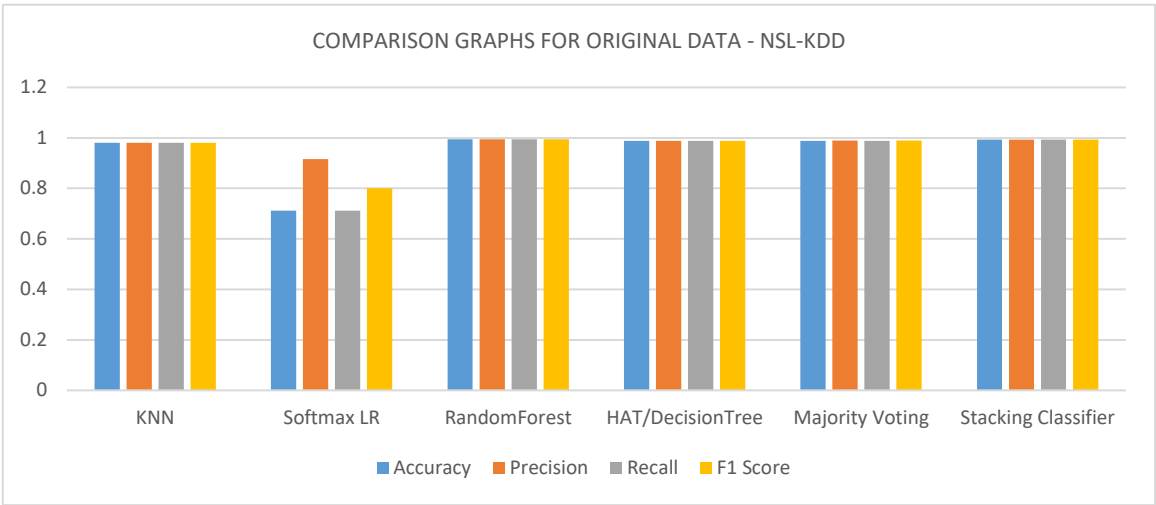


Graph.3.4 Comparison Graphs for Random Under/Over Sampling (Combine) – KDDCUP

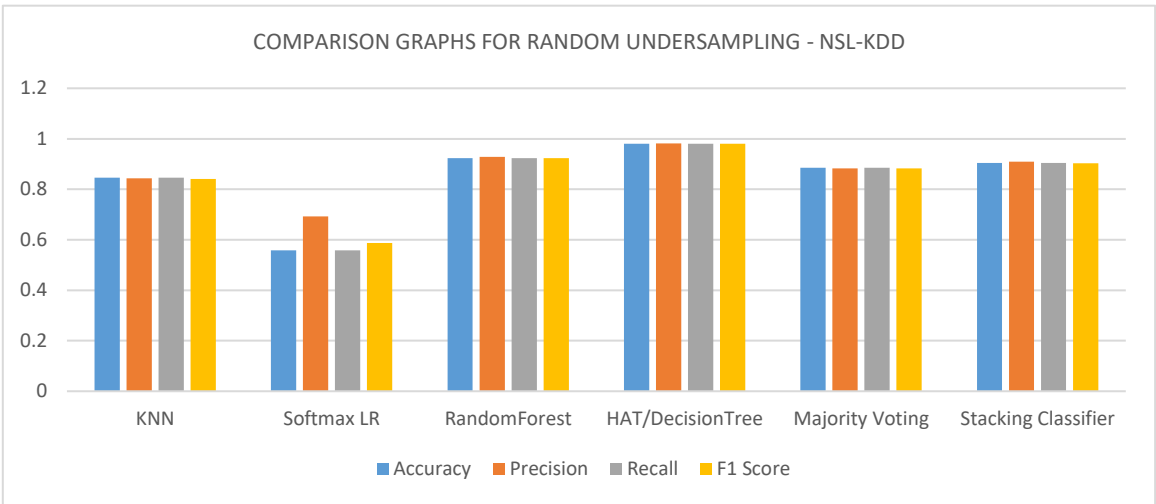


The NSL-KDD dataset *Graphs (4.1 to 4.4)*, where accuracy is shown in blue, precision in orange, recall in gray, and F1-score in yellow. Among all the models tested, the Stacking Classifier stands out, delivering the highest performance across every sampling method. This trend is visually represented in the below graphs.

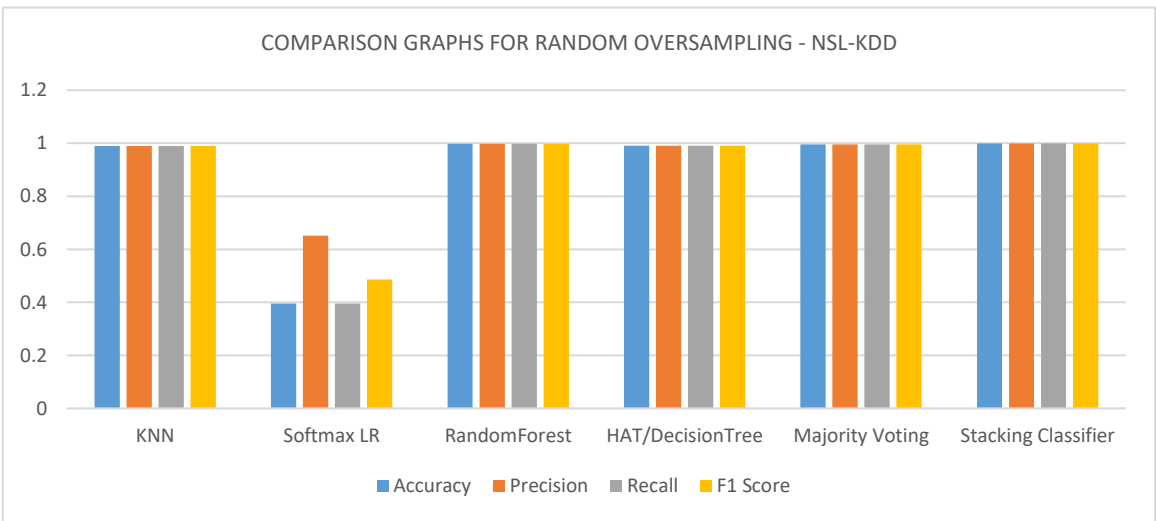
Graph.4.1 Comparison Graphs for Original Data – NSL-KDD



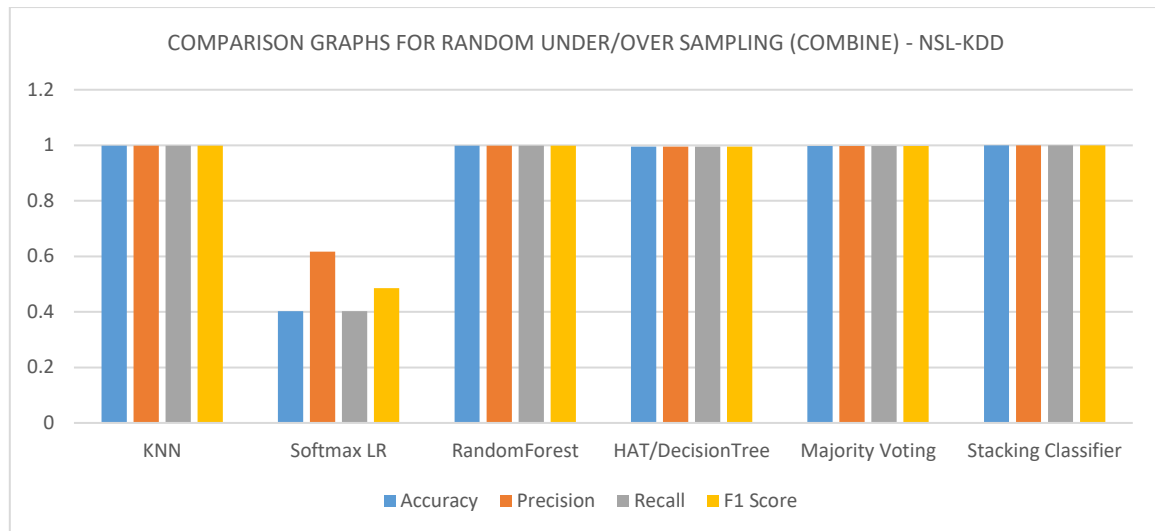
Graph.4.2 Comparison Graphs for Random UnderSampling – NSL-KDD



Graph.4.3 Comparison Graphs for Random Over Sampling – NSL-KDD



Graph.4.4 Comparison Graphs for Random Under/Over Sampling (Combine) – NSL-KDD



Accuracy is represented in blue, precision in orange, recall in gray, and F1-Score in yellow, **Graph(1.1 to 4.4)**. In comparison to the other models, the Stacking Classifier shows superior performance across all metrics, achieving the highest values. The graphs above visually illustrate these findings.

5. CONCLUSION

In conclusion, this study demonstrates the significant potential of advanced machine learning techniques in improving Intrusion Detection Systems (IDS) for network security. By utilizing four widely recognized cybersecurity datasets—CIC IDS 2017, NSL KDD, KDD Cup, and CIC IDS 2018—we have effectively evaluated various approaches for intrusion detection. The application of feature selection via Mutual Information proved to be beneficial in reducing irrelevant features, thereby enhancing model performance. Furthermore, the exploration of data sampling techniques, including Random Under Sampling, Random Over Sampling, and a combination of both, effectively addressed data imbalance issues, ensuring more reliable predictions. The implementation of a Stacking Classifier, which combines the strengths of Random Forest and Decision Tree models with a Bagging Classifier, demonstrated superior performance in accurately detecting intrusions. This approach consistently outperformed other methods across all datasets, showcasing its robustness and adaptability in dynamic cybersecurity environments. The results underscore the importance of integrating multiple strategies, such as feature selection, data sampling, and ensemble learning, to develop effective IDS solutions capable of detecting evolving threats in real-time network traffic.

Future work could focus on enhancing the performance of IDS by incorporating more advanced feature selection methods, exploring deep learning models for anomaly detection, and evaluating additional datasets for broader generalization. Additionally, investigating the use

of real-time data streaming and online learning algorithms could improve IDS responsiveness to emerging threats. Furthermore, integrating threat intelligence feeds and developing hybrid systems combining multiple ensemble techniques may lead to even more robust and accurate intrusion detection capabilities, adapting to evolving cybersecurity challenges.

REFERENCES

- [1] M. Mijwil, O. J. Unogwu, Y. Filali, I. Bala, and H. Al-Shahwani, "Exploring the top five evolving threats in cybersecurity: An in-depth overview," *Mesopotamian J. Cyber Secur.*, vol. 2023, pp. 57–63, Mar. 2023.
- [2] T. Fadziso, U. Thaduri, S. Dekkati, V. Ballamudi, and H. Desamsetti, "Evolution of the cyber security threat: An overview of the scale of cyber threat," *Digitalization Sustainability Rev.*, vol. 3, no. 1, pp. 1–12, 2023.
- [3] R. Dillon, P. Lothian, S. Grewal, D. Pereira, and A. Kuah, "Cyber security: Evolving threats in an ever changing world," in *Digital Transformation in a Post-Covid World: Sustainable Innovation, Disruption and Change*. Boca Raton, FL, USA: CRC Press, 2021, pp. 129–154.
- [4] P. Vanin, T. Newe, L. L. Dhirani, E. O'Connell, D. O'Shea, B. Lee, and M. Rao, "A study of network intrusion detection systems using artificial intelligence/machine learning," *Appl. Sci.*, vol. 12, no. 22, p. 11752, Nov. 2022, doi: 10.3390/app122211752.

- [5] H. Albasheer, M. Md Siraj, A. Mubarakali, O. Elsier Tayfour, S. Salih, M. Hamdan, S. Khan, A. Zainal, and S. Kamarudeen, "Cyber attack prediction based on network intrusion detection systems for alert correlation techniques: A survey," *Sensors*, vol. 22, no. 4, p. 1494, Feb. 2022.
- [6] Z. Ahmad, A. Shahid Khan, C. Wai Shiang, J. Abdullah, and F. Ahmad, "Network intrusion detection system: A systematic study of machine learning and deep learning approaches," *Trans. Emerg. Telecommun. Technol.*, vol. 32, no. 1, p. e4150, Jan. 2021.
- [7] H. Liu and B. Lang, "Machine learning and deep learning methods for intrusion detection systems: A survey," *Appl. Sci.*, vol. 9, no. 20, p. 4396, Oct. 2019, doi: 10.3390/app9204396..
- [8] L. Shahbandayeva, U. Mammadzada, I. Manafova, S. Jafarli, and A. Z. Adamov, "Network intrusion detection using supervised and unsupervised machine learning," in *Proc. IEEE 16th Int. Conf. Appl. Inf. Commun. Technol. (AICT)*, Oct. 2022, pp. 1–7.
- [9] S. Gamage and J. Samarabandu, "Deep learning methods in network intrusion detection: A survey and an objective comparison," *J. Netw. Comput. Appl.*, vol. 169, Nov. 2020, Art. no. 102767. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1084804520302411>
- [10] K. S. Adewole, T. T. Salau-Ibrahim, A. L. Imoize, I. D. Oladipo, M. AbdulRaheem, J. B. Awotunde, A. O. Balogun, R. M. Isiaka, and T. O. Aro, "Empirical analysis of data streaming and batch learning models for network intrusion detection," *Electronics*, vol. 11, no. 19, p. 3109, Sep. 2022.
- [11] D. Bhosale and R. Ade, "Intrusion detection using incremental learning from streaming imbalanced data," *Int. J. Manag. Public Sector Inf. Commun. Technol.*, vol. 6, no. 1, pp. 9–20, Mar. 2015.
- [12] M. R. Mohamed, A. A. Nasr, I. F. Tarrad, and M. Z. Abdulmageed, "Exploiting incremental classifiers for the training of an adaptive intrusion detection model," *Int. J. Netw. Secur.*, vol. 21, no. 2, pp. 275–289, 2019.
- [13] M. Data and M. Aritsugi, "AB-HT: An ensemble incremental learning algorithm for network intrusion detection systems," in *Proc. Int. Conf. Data Sci. Appl. (ICoDSA)*, Jul. 2022, pp. 47–52.
- [14] B. A. Tama and S. Lim, "Ensemble learning for intrusion detection systems: A systematic mapping study and cross-benchmark evaluation," *Comput.Sci. Rev.*, vol. 39, Feb. 2021, Art. no. 100357. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1574013720304573>
- [15] M. Torabi, N. I. Udzir, M. T. Abdullah, and R. Yaakob, "A review on feature selection and ensemble techniques for intrusion detection system," *Int. J. Adv. Comput. Sci. Appl.*, vol. 12, no. 5, pp. 6–7, 2021. [Online]. Available: <https://api.semanticscholar.org/CorpusID:236317529>
- [16] A.M.Bamhdi,I.Abrar,andF.Masoodi,"An ensemble based approach for effective intrusion detection using majority voting," *Telkomnika*, vol. 19, no. 2, pp. 664–671, Apr. 2021.
- [17] D.R.Patil and T.M.Pattewar, "Majority voting and feature selection based network intrusion detection system," *EAI Endorsed Trans. Scalable Inf. Syst.*, vol. 9, no. 6, p. e6, 2022.
- [18] H. Xu and Y. Wang, "A continual few-shot learning method via meta learning for intrusion detection," in *Proc. IEEE4thInt.Conf.CivilAviation Saf. Inf. Technol. (ICCASIT)*, Oct. 2022, pp. 1188–1194.
- [19] T.Wang, Q.Lv, B.Hu, and D.Sun, "A few-shot class-incremental learning approach for intrusion detection," in *Proc. Int. Conf. Comput. Commun. Netw. (ICCCN)*, Jul. 2021, pp. 1–8.
- [20] J. Zheng, X. Ni, L. Li, K. Yu, and J. Zhang, "An ensemble learning-based two-level network intrusion detection method," in *Proc. Int. Conf. Comput. Eng. Artif. Intell. (ICCEAI)*, 2022, pp. 571–575.
- [21] T.Wang,Q.Lv,B.Hu,andD.Sun,"A few-shotclass-incremental learning approach for intrusion detection," in *Proc. Int. Conf. Comput. Commun. Netw. (ICCCN)*, Jul. 2021, pp. 1–8.