

Student Number: 101064415

Multi-Modal Anomaly Detection using Machine Learning and Data Fusion for Intrusion Detection Systems

Mridul Lal

Supervisor: Dr. Konstantinos Markantonakis

Submitted as part of the requirements for the award of the
MSc in Information Security
at Royal Holloway University of London.



Information Security Group
Royal Holloway, University of London
August, 2024

Table of Contents

Executive Summary.....	3
Acknowledgments.....	4
List of Tables and Figures.....	5
List of Abbreviations & Acronyms.....	6
1. Introduction.....	7
1.1 Objectives of the Research:.....	9
1.2 Structure of the Dissertation:.....	10
2. Literature Review on Intrusion Detection Systems.....	11
2.1 Overview of Intrusion Detection Systems.....	11
2.2 Traditional Approaches to Intrusion Detection Systems.....	13
2.3 Classification of Intrusion Detection Systems.....	16
2.4 Machine Learning in Intrusion Detection Systems and the need for Data Fusion.....	18
2.5 Similar work on Machine Learning in Intrusion Detection Systems.....	20
3. Methodology.....	25
4. Dataset Selection.....	27
5. Feature Selection.....	31
6. Data Fusion.....	34
7. Machine Learning Implementation.....	37
8. Evaluation Metrics.....	42
9. Experimental Setup and Implementation Challenges.....	47
10. Results and Discussion.....	49
11. Conclusion and Future Work.....	55
References.....	57

Executive Summary

This dissertation explores the application of multi-modal data fusion techniques in the development of an advanced Intrusion Detection System (IDS). In an era of increasingly sophisticated cyber threats, traditional single-source IDSs often fall short in detecting complex, multi-vector attacks. Our research addresses this challenge by integrating diverse data sources, including system logs, user behavior patterns, and network traffic data, to create a more comprehensive and robust intrusion detection framework. The study begins with a thorough analysis of two prominent datasets in the field of network security: NSL-KDD and UNSW-NB15. We then propose and implement a novel approach to fuse these datasets, creating a multi-modal representation of network activities. This fused dataset forms the foundation for our machine learning-based IDS, which we develop and evaluate using a range of classification algorithms including Support Vector Machines (SVM), Decision Trees, Random Forests, K-Nearest Neighbors (KNN), and Naive Bayes. Our experimental results demonstrate significant performance improvements across all evaluated metrics: accuracy, precision, recall, and F1-score, when using the fused dataset compared to single-source approaches. Notably, models trained on the fused dataset show enhanced generalization capabilities and a better balance between false positives and false negatives.

The Decision Tree algorithm emerged as the top performer, achieving an F1-score of 0.9548 on the fused dataset, a substantial improvement over its performance on individual datasets. We provide a comprehensive analysis of the benefits and challenges associated with implementing multi-modal intrusion detection systems. The benefits include improved detection accuracy, enhanced ability to identify complex attack patterns, and increased robustness against evolving threats.

Challenges such as data integration complexities and increased computational requirements are also discussed. To ensure reproducibility and foster further research in this area, we have made our code and methodology publicly available. This dissertation contributes to the field of cybersecurity by demonstrating the efficacy of multi-modal data fusion in intrusion detection and paving the way for more effective, adaptive, and reliable security solutions.

Our findings have significant implications for the future of network security, suggesting that multi-modal approaches could be key to developing next-generation intrusion detection systems capable of addressing the complex and dynamic nature of modern cyber threats. The research concludes by outlining potential avenues for future work, including real-time implementation, exploration of advanced fusion techniques, and integration with automated response systems.

Acknowledgments

I want to extend gratitude to the people who supported me during this dissertation process, with special thanks to my project supervisor

This project could not have come together without the support and motivation of my mother Preetha, my uncle Pradeep, my aunt Rozaliva, my sister Sneha, and my girlfriend Malavika. All of you have helped me excel in so many different ways and I can't thank you enough for all the love and support.

I also want to thank my friends Cazel and Abigail for taking time out to review my paper and give me many helpful suggestions to improve this dissertation and take it to a higher level.

List of Tables and Figures

Table 1	Comparisons of intrusion detection methodologies	Page 12 - 13
Table 2	Comparison of IDS technology types based on their positioning within the computer system	Page 14
Figure 1	Working Design of Methodology	Page 22
Table 3	NSL-KDD and UNSW-NB15 list of features	Page 25
Table 4	Similarities of the features in NSL-KDD and UNSW-NB15	Page 26
Table 5	Identified features from feature selection on NSL-KDD and UNSW-NB15	Page 29
Table 6	Features of fused dataset as a result of Data Fusion of UNSW-NB15 and NSL-KDD	Page 32
Figure 2	Performance Comparison of NSL-KDD, UNSW-NB15, and Fused Dataset	Page 44
Table 7	Comparative Results of Data Fusion and Machine Learning Implementation on all Datasets	Page 45
Figure 3	ROC Curve for UNSW-NB15 Dataset	Page 48
Figure 4	ROC Curve for NSL-KDD Dataset	Page 48
Figure 5	ROC Curve for Fused Dataset	Page 49

List of Abbreviations & Acronyms

IDS	Intrusion Detection System
IDPS	Intrusion Detection and Prevention System
SIDS	Signature-based Intrusion Detection System
AIDS	Anomaly-based Intrusion Detection System
DF	Data Fusion
HIDS	Host-based Intrusion Detection System
NIDS	Network-based Intrusion Detection System
SVM	Support Vector Machines
KNN	K-Nearest Neighbors
FA	False Alarms
API	Application Program Interface
SNMP	Simple Network Management Protocol
TCP	Transmission Control Protocol
UDP	User Datagram Protocol
ICMP	Internet Control Message Protocol
MIB	Management Information Base
FNR	False Negative Ratio
FPR	False Positive Ratio
APT	Advanced Persistent Threats
DBN	Deep Belief Network
ACC	Accuracy
SRDLM	Semantic Re-encoder and Deep Learning Model
CNN	Convolutional Neural Networks
OSS	One Side Selection
RFE	Recursive Feature Elimination
DoS	Denial of Service
PCA	Principal Component Analysis
ROC	Receiver Operating Characteristic
TP	True Positive
TN	True Negative
FP	False Positive
FN	False Negative

1. Introduction

Over the past few decades, we have seen how computer systems and interconnected networks have made their way into every facet of our lives. Almost every individual and organization has a digital identity which has become an essential part of their everyday lives. These digital identities are further used to communicate and interact with institutions varying from government and financial agencies to small businesses. This has led to a massive increase in network traffic that contains different types of data from many different sources. The increased connectivity and availability provided by interconnected networks have also opened up avenues for malicious actors to steal and defraud.

These cyber attacks have become increasingly prevalent with over 343 million people losing their sensitive data to data breaches in 2023, which is actually a 16 percentage point decrease compared to the attacks in 2022 [1]. Over half of businesses in the UK and around a third of charities have claimed to have fallen victim to some form of cyber attack over the last 12 months. With the average cost of a data breach being around \$4.45 million and companies losing close to \$2.7 billion in 2022 alone to business email compromises, the financial and reputational impact of cyber attacks is large enough to cause entire businesses and organizations to cease operations completely. The frequency of such attacks has also increased with the number of data breaches increasing by 72 percent from 2021 to 2023 [2]. Traffic on the internet is only going to increase exponentially and with it the number of possible cyber attacks. World Economic Forum has projected cybercrime to cost companies worldwide an estimated \$10.5 trillion annually by 2025 [3].

The frequency and impact of cyber attacks have brought in response a range of mechanisms and protocols to help safeguard networks, systems, and data from unauthorized access and manipulation. Technologies like user authentication and authorization, cryptography, and encryption, and firewalls have been implemented to protect against possible cyber attacks. These technologies protect data being transmitted but do not completely provide the ability to detect and react to possible intrusions into interconnected networks. The intrusion detection and prevention system (IDPS) is the software or hardware system that automates the intrusion detection process and aims to stop any possible incidents [4]. An intrusion detection and prevention system continuously monitors and analyzes network traffic and activities on the target networks and individual devices. They make use of algorithms and pattern-matching techniques to identify suspicious behavior or activity that deviates from the “normal” network behavior.

Traditional intrusion detection systems often rely on a single data source, such as network traffic or system logs to detect any anomalous behavior. However, with the increasing complexity of computers, networks, and architectures, relying on a single data source can prove inadequate for effective detection of intrusions. Intrusion detection systems used in the industry are also known to work on predefined rules, configured signatures, or observe singular modalities like system logs or network traffic. This makes intrusion detection systems weak when it comes to detecting multi-faceted attacks or sophisticated attacks like zero-day attacks.

The evolving landscape of cyber threats presents a significant challenge to conventional intrusion detection systems. Advanced Persistent Threats (APTs), polymorphic malware, and sophisticated social engineering attacks often bypass traditional security measures. These complex attack vectors exploit multiple vulnerabilities across different system components, making them particularly difficult to detect using single-source data analysis. Furthermore, the rapid growth of Internet of Things (IoT) devices and cloud computing has expanded the attack surface, introducing new vulnerabilities and increasing the volume and variety of data that must be monitored for potential threats.

To address these challenges, researchers and cybersecurity professionals are turning to more advanced, intelligent approaches to intrusion detection. Machine learning and artificial intelligence techniques have shown promising results in improving the accuracy and efficiency of intrusion detection systems. These approaches can adapt to new threats, learn from past incidents, and identify subtle patterns that might be missed by rule-based systems. However, even these advanced techniques can be limited when relying on a single data source.

This dissertation aims to address this limitation of existing IDS solutions by investigating the effectiveness of a Multi-modal intrusion detection system that makes use of a machine learning concept called Data Fusion to improve detection [5]. Data Fusion is a machine learning concept that can be used in intrusion detection to deal with the problem of detecting and analyzing multiple sources of information simultaneously to identify potential intrusions [6].

The concept of multi-modal intrusion detection leverages the strengths of data fusion to create a more comprehensive and robust security solution. By integrating data from various sources such as network traffic, system logs, user behavior patterns, and even external threat intelligence feeds, multi-modal IDS can provide a more holistic view of the system's security state. This approach not only improves the accuracy of threat detection but also reduces false positives, a common issue in traditional IDS that can lead to alert fatigue among security personnel. Moreover, multi-modal IDS has the potential to enhance the interpretability of detected threats. By correlating information from multiple sources, security analysts can gain deeper insights into the nature and context of potential intrusions, facilitating more effective and timely responses. This capability is particularly crucial in today's complex IT environments, where the speed and accuracy of threat response can significantly impact the extent of damage caused by a cyber attack.

This dissertation aims to contribute to the growing body of research on multi-modal intrusion detection systems. By exploring the application of data fusion techniques and advanced machine learning algorithms, we seek to develop a more effective and adaptive approach to network security. The findings of this study have the potential to inform the design of next-generation intrusion detection systems, capable of meeting the evolving challenges of cybersecurity in an increasingly interconnected world. Through rigorous experimentation and analysis, we hope to demonstrate the tangible benefits of multi-modal IDS and pave the way for its wider adoption in real-world security infrastructures.

1.1 Objectives of the Research:

The primary objectives of this research are:

1. To explore and evaluate the effectiveness of multi-modal data fusion in the context of intrusion detection.
2. To develop a machine learning-based IDS that integrates system logs, user behaviour, and network traffic data.
3. To compare and display the performance improvement of the proposed multi-modal Intrusion Detection System with solutions that make use of singular modalities.
4. To write code for the proposed methodology and make it publicly available so that the results stated and discussed in this dissertation can be reproduced and used for further research.
5. To provide a comprehensive analysis of the benefits and challenges associated with implementing multi-modal intrusion detection systems.

1.2 Structure of the Dissertation:

The dissertation is organized as follows:

- **Chapter 1** presents the introduction to the project, the objectives of the research and the structure of the dissertation.
- **Chapter 2** presents a detailed literature review on IDS, machine learning applications in IDS, and multi-modal data fusion techniques.
- **Chapter 3** outlines the methodology, including the datasets used, data pre-processing techniques, feature extraction, machine learning models employed, and the representation and evaluation of findings.
- **Chapter 4** describes the steps and criteria involved in the process of dataset selection for datasets that cover attack data of different modalities
- **Chapter 5** discusses the steps and methodology involved in the feature selection process along with its results
- **Chapter 6** details the steps and methodology involved in the data fusion process along with its results
- **Chapter 7** presents the machine learning implementation with a focus on explaining the algorithms along with their relevance and importance to the project.
- **Chapter 8** describes the evaluation metrics used in this project along with their relevance to intrusion detection in general.
- **Chapter 9** details the experimental setup used for running the project code and achieving the results mentioned
- **Chapter 10** discusses the results and findings of the project and discusses them along with graphical representations of the findings.
- **Chapter 11** is the conclusion to the dissertation with an emphasis on future work in the field of intrusion detection

2. Literature Review on Intrusion Detection Systems

2.1 Overview of Intrusion Detection Systems

Ever since the inception of the internet in the 1980s, the size and number of interconnected devices, networks, and users have exponentially increased. The increased efficiency and decreased processing times that this interconnectedness provides have resulted in the increased digitalization of everyday processes. The concept of having a digital identity has become essential to individuals and organizations alike. The devices around us now have an increasing number of sensors that collect a wide range of data. This has resulted in a day and age where extensive amounts of personable and sensitive information are exchanged between networks and devices daily. This in turn has opened up many different avenues and opportunities for threat actors to gain illegal access to and exploit networks and devices for obtaining and using this data for malicious purposes.

One of the earliest known instances of cyber attack is probably the Morris Worm. A self-replicating worm that could spread itself to all connected devices in a network causing systems to become sluggish or completely unusable due to the excessive load from multiple infections. Around 6000 systems, which was a significant portion of the internet at the time, were affected which had an extensive financial impact on businesses and organizations. The worm exploited the vulnerabilities in applications on Unix systems to infect individual systems and spread itself to other systems in the network. The Morris Worm attack and the impact it had led to widespread awareness of security mechanisms and consideration of security when drafting organization-wide policies. The biggest security-related change that the Morris Worm attack brought forth is the firewall, making it an essential part of any organization's network presence [7]. Subsequent cyber attacks have led to the creation and implementation of various security tools and mechanisms that protect the confidentiality, integrity and availability of data, systems, and networks like authentication and authorization protocols and tools, cryptographic encryption, and security policies that incorporate security into every step of organizational processes. These tools are considered the first line of defense and are extremely essential to ensuring security and must be implemented to build resilient networks and processes as they help mitigate the associated risks but these tools often lack the reactive ability to detect and prevent any cyber attacks as they occur. For example, if a password is weak and has been leaked, user authentication cannot prevent unauthorized access and use [8]. Firewalls work mainly on the principle of filtering out unwanted traffic and not letting that traffic travel into the network. The desirable or undesirable packets have to be configured into firewalls for them to function but in the case that some undesirable packets have bypassed the firewall to the network, the firewall has no capability to stop the packet or to contain the impact the packet might have on the network. Firewalls are also prone to a range of vulnerabilities caused by errors in configurations.

The idea and need for intrusion detection and prevention systems come from the requirement to respond to security incidents or cyber attacks as they occur. This can only be done by monitoring the data flow and systems in the networks, detecting any anomalous behaviour or packet flow and preventing its propagation. The goal of intrusion detection and prevention systems is to monitor network assets to

detect and prevent any kind of anomalous behaviour or misuse within the target networks. The concept of intrusion detection is one that has been around for around thirty years or more but has now seen a dramatic rise in popularity and incorporation because of the capability it provides to detect ongoing attacks and prevent further infection of systems and networks [9].

The origins of intrusion detection as a concept originated in the 1980s from James Anderson's paper, Computer Security Threat Monitoring and Surveillance. This paper written for a government organization stated that audit trails contained valuable information that could be vital when it comes to identifying, tracking, or detecting any kind of misuse and for learning trends of user behaviour. This paper led to tremendous improvements in auditing processes for every operating system and the hypothesis provided in this paper became the foundation for intrusion detection systems design and development in the future [10]. This work is considered to be the basis of the development of Host-based intrusion detection and IDS in general.

Another path-breaking paper in the field of intrusion detection is "An Intrusion Detection Model" by Dorothy E. Denning in 1983, which sought to create a model of a real-time intrusion detection system that could detect intrusions, break-ins, or any form of computer abuse. The model was based on the hypothesis that any security violations can be detected by monitoring the audit logs of a system for any abnormal patterns of usage. It was a pioneer in the concept of creating profiles that represented the behaviour of subjects with respect to the devices being used. It also aimed at making use of metrics and statistical models to identify acceptable and abnormal behaviour. An additional achievement of this model is how it created a framework for identifying any kind of intrusion or security violation irrespective of the system or the application environment [11].

Intrusion Detection Systems (IDS) have now become critical components in modern network security infrastructure, which have been designed to identify and alert on any suspicious or abnormal activity that may indicate unauthorized access or malicious behaviour in a computer network. As defined by [4], "An intrusion detection system is a device or software application that monitors a network or systems for malicious activity or policy violations". This paper also defines intrusions as "any kind of unauthorized activities that cause damage to an information system. This means any attack that could pose a possible threat to the information confidentiality, integrity or availability will be considered an intrusion". The importance of intrusion detection systems in network security cannot be overstated with the capability it provides to prevent and respond to incidents as they occur serving as a crucial line of defense in the face of evolving cyber security threats.

2.2 Traditional Approaches to Intrusion Detection Systems

Traditional approaches to Intrusion Detection Systems can be classified into 2 main categories or types:

- Misuse-based detection and,
- Anomaly-based detection.

Misuse-based detection, also called Knowledge-based Detection or Signature-based detection, makes use of pattern-matching techniques to identify and differentiate between normal and abnormal activity. To achieve this, a database of intrusion signatures needs to be maintained for comparison with the current set of activities to decide whether an intrusion is occurring or not. In systems with a Signature-based Intrusion Detection System (SIDS), the logs of current activities are matched with the signature entries in the database, and on a successful match, an alarm will be raised which could result in quarantining or stopping processes to prevent further damage to the system and network [12]. Signature-based Intrusion Detection Systems are known to usually give high detection accuracy for previously known attacks or intrusions, but they have difficulty in identifying and protecting against attacks like zero-day attacks which would be unknown and not present in the database of stored intrusion signatures [13]. A signature matching the zero-day attack would have to be stored in the database for the accurate detection and prevention of such attacks. Numerous common tools used by organizations for providing network security and protecting against potential attacks make use of the Signature-based Intrusion Detection System, Snort and NetSTAT being the most popular and commonly used options. The concept of using known signatures, storing them and comparing them to identify and prevent against attacks is one that has been implemented and used for quite some time which is why it is counted among one of the traditional approaches. Such intrusion detection systems will have an amount of memory allocated to it to store intrusion signatures but these systems will not be storing information of activity within networks, intrusion detection systems would scan a chunk of real time activity logs, compare them for possible intrusions and move on to the next chunk if it finds no alarming activity [14]. This makes intrusion detection systems weak to attacks that span several packets over a time period. As attack patterns and methods keep evolving and becoming more sophisticated it may become necessary for extract signature related information over multiple packets over time.

Attacks like zero-day attacks are on the rise and this has make Signature-based Intrusion Detection Systems less effective when it comes to detecting or protecting against these attacks without having a similar signature stored. A potential solution was considered to be an Anomaly-based Intrusion Detection System which works by profiling what is acceptable behaviour or activity and raising an alarm when it finds activity other than the expected.

Anomaly-based intrusion detection systems have garnered a lot of interest from academics and organizations for their ability to overcome the shortcomings of Signature-based intrusion detection systems. In the Anomaly-based intrusion detection systems, models are created using machine learning, statistics based methods or knowledge based methods to create and capture the expected normal behaviour in a computer system. Large deviations from said expected behaviour will be considered an intrusion or an anomaly and such activities would result in raising of alarms and quarantining of

associated processes. The assumption being that malicious activities would differ from normal activity to a great extent and such activity can be classified or categorized as intrusions potentially from threat actors. Anomaly-based Intrusion Detection Systems have two main phases, the training and the testing phase. The training phase makes use of the logs and data of normal activity to train the model and in the testing phase, new data set is used to understand and evaluate the model or the system's capability to detect and prevent intrusions that it previously was unaware of [15].

Anomaly-based Intrusion Detection Systems in this way can overcome the weaknesses in Signature-based Intrusion Detection systems in identifying and preventing intrusions from attacks like zero-day attacks. This is mainly because the intrusions would vary largely from the normal or expected activity in the network and systems and the intrusion detection systems could raise an alarm to alert the concerned parties and quarantine the systems and applications to prevent the further spread of malicious packets. Anomaly-based Intrusion Detection Systems have other benefits as well. It can identify any internal malicious activities which can be used to identify any insider threats where an authorized user commits unexpected activity. Another advantage of Anomaly based intrusion detection systems is that different intrusion detection systems may make use of different datasets for training their model of expected behaviour and this makes extremely difficult for threat actors with malicious intent to understand the model and test the limitations without triggering alerts and losing the access they have into the system. This adds an additional layer of randomness and security from any potential attackers trying to study the system before running exploits in the system [16]. The major disadvantage of Anomaly-based Intrusion Detection Systems is the high false positive rate that these systems have when it comes to detecting intrusions. False positives here means authorized users carrying out expected and normal activity but the intrusion detection system identifying this activity as a possible intrusion and raising false alarms for it. Another drawback is the requirement for extensive training data for such intrusion detection systems, training data is often very large in size and is done in this way so that the model can create correlations between the different features involved and having more data would help make more accurate correlations and analysis as a result.

The table below encapsulates the main differences, advantages and disadvantages of both traditional intrusion detection methodologies discussed.

Table 1: Comparisons of intrusion detection methodologies [17]

		Advantages	Disadvantages
Detection Methodologies	SIDS	• Very effective in identifying intrusions with minimum false alarms (FA). • Promptly identifies the intrusions. • Superior for detecting the known attacks. • Simple design	• Needs to be updated frequently with a new signature. • SIDS is designed to detect attacks for known signatures. When a previous intrusion has been altered slightly to a new variant, then the system would be unable to identify this new deviation of the similar attack. • Unable to detect the zero-day attack. • Not suitable for detecting multi-step attacks. • Little understanding of the insight of the attacks
	AIDS	• Could be used to detect new attacks. • Could be used to create intrusion signature	• AIDS cannot handle encrypted packets, so the attack can stay undetected and can present a threat. • High false positive alarms. • Hard to build a normal profile for a very dynamic computer system. • Unclassified alerts. • Needs initial training.

There also exists Hybrid Intrusion Detection Systems, which aims to combine the advantages of both Signature-based Intrusion Detection Systems and Anomaly-based Intrusion Detection Systems. These systems are known to built by adding signature-based detection components to anomaly-based intrusion detection systems or vice versa, which has shown to have higher intrusion detection and lower false positive rates [18].

2.3 Classification of Intrusion Detection Systems

The previous sections of the literature review aimed at giving an overview of intrusion detection systems and methods used by different intrusion detection systems to identify intrusions as we need to understand how intrusion detection systems have evolved over time. This section aims to classify intrusion detection systems based on different input data sources. It is important to know this classification as based on the different input sources that intrusion detection systems track they can be placed at different levels of the organizations infrastructure and will be able to identify different types of attacks. Intrusion detection systems are mainly classified into two main classes based on their input sources:

- Host-based Intrusion Detection Systems (HIDS)
- Network-based Intrusion Detection Systems (NIDS)

A Host-based Intrusion Detection System inspects data on a single host which will include audit sources like the operating system, network connections, packet flow, firewalls logs, database logs and user activity among other sources. The intrusion detection system would observe these logs and identify intrusions based on signature based or anomaly based mechanisms as discussed earlier. Host-based Intrusion Detection Systems are placed on servers and systems that are crucial to the functioning and business processes of the organization. These systems can be automated to block network connections, terminate suspicious processes and isolate effected applications on detecting an intrusion of any kind [19]. The key advantages of this system is the ability to detect local attacks that might not be visible at the network level and the detailed visibility it provides to activities and user behaviour on particular hosts. These systems also have the added capability of monitoring encrypted network packets or traffic as it is decrypted on the host. The major drawbacks are that Host-based Intrusion Detection Systems can be resource intensive and requires installation and maintenance on individual systems

Comparatively, Network-based Intrusion Detection Systems monitors network traffic across an entire network and all systems and network devices on the network. Network-based intrusion detection systems is able to monitor the external malicious activities that could come from threat actors as they try to gain access to the network or try to run exploits and can stop the spread of said malicious packets from spreading to other systems in the network. Networks are often very large in size with many systems and high bandwidth connections and this can limit the ability of the intrusion detection system to inspect all data and detect potential intrusions. The major benefit of a Network-based intrusion detection system is the network level visibility that these systems have that can be used to detect more sophisticated attacks. Another advantage is that such intrusion detection systems can be placed at different levels of the network topology according to the importance of specific networks or sub-networks according to the security policy [20].

Table 2: Comparison of IDS technology types based on their positioning within the computer system [17]

	Advantages	Disadvantages	Data Source
<u>HIDS</u>	• <u>HIDS</u> can check end-to-end encrypted communications behaviour. • No extra hardware required. • Detects intrusions by checking hosts file system, system calls or network events. • Every packet is reassembled • Looks at the entire item, not streams only	• Delays in reporting attacks • Consumes host resources • Needs to be installed on each host. • It can monitor attacks only on the machine where it is installed.	• Audits records, log files, Application Program Interface (API), rule patterns, system calls.
<u>NIDS</u>	• Detects attacks by checking network packets. • Not required to install on each host. • Can check various hosts at the same period. • Capable of detecting the broadest ranges of network protocols	• Challenge is to identify attacks from encrypted traffic. • Dedicated hardware is required. • It supports only identification of network attacks. • Difficult to analysis high-speed network. • The most serious threat is the insider attack.	• Simple Network Management Protocol (<u>SNMP</u>) • Network packets (<u>TCP/UDP/ICMP</u>), • Management Information Base (<u>MIB</u>) • Router <u>NetFlow</u> records

2.4 Machine Learning in Intrusion Detection Systems and the need for Data Fusion

In previous sections we have seen the evolution of intrusion detection systems to identify intrusions better but with high false positive rates, less accuracy, increasing size of networks and increasing complexity of attacks, there is still a gap that needs to be bridged to provide higher levels of information and network security to organizations and individuals' data. The solution that bridges this gap would have to provide improved accuracy and detection rates, should be able to handle and adapt to complex and evolving attack patterns and process large amounts of data. Machine learning is one such solution that has been applied on publicly available datasets by academics and other professionals which has shown to bridge the gap to providing higher levels of security to networks. Machine learning has been defined as "a category of artificial intelligence that enables computers to think and learn on their own" [21]. Machine learning gives intrusion detection systems the ability to adapt to complex systems & network activities and modify their actions to attain more accuracy.

There are two main classifications of machine learning methods,

- supervised and
- unsupervised machine learning

Supervised machine learning makes use of labeled data to create a function that maps a particular input to an output. In the context of intrusion detection datasets, labeled data here would mean datasets where network traffic or log entries have been labeled as normal activity or a specific kind of attack. This labeling process is either done manually or can be done automatically using tools. For supervised machine learning algorithms in Intrusion Detection Systems, labeled data is essential to allow the algorithm to learn patterns and characteristics of different classes of data. Supervised machine learning can be further classified into:

- Classification models and,
- Regression models

Classification models as the name says puts the data into specific categories. This is done on a dataset based on the different features of the data entries in the dataset. The model is trained on labeled input and output data to understand what are the deterministic features that make the particular data input entry correlate to the particular output, in context of intrusion detection systems the model will learn the deterministic features that make a particular user behaviour, network or system logs be categorized as an intrusion. Classification models are very useful to detect attacks as a well trained model can categorize incoming traffic as normal activity or as intrusions. If the model is trained well and performs well, the model can categorize the incoming traffic into attack types like Denial of Service, phishing, worms, port scanning etc. Commonly used and well known classification algorithms are decision tree, support vector machine, random forest, k-nearest neighbour and neural networks [22].

Regression models on the other hand is used to understand continuous behaviour or outcomes. The model is trained to make the correlation between an independent variable and the dependent variables. Regression models find patterns and relationships in one set of data and can apply these patterns to

recognize behaviour in other datasets. This is why regression models are using in financial markets to predict trends and even in forecasting weather patterns [23]. Commonly used regression algorithms include linear regression, decision tree, random forest,

One of the main advantages of supervised machine learning is that it uses previously recognized patterns and behaviours to produce outputs with a higher degree of precision thereby improving the models performance. Comparatively one pitfall of supervised learning is how dependent it's performance is on precise and good inputs and outputs to make accurate predictions and analysis. This training is also resource intensive and requires a lot of computation time. It also becomes difficult to classify big data if the dataset is too big, which is where unsupervised machine learning algorithms tends to get preference.

Unsupervised machine learning comparatively is used with unlabeled data. As the name suggests, unsupervised machine learning does not require any kind of intervention manually for training stages. The algorithm does the grouping and classification of data based on the similarities or differences that the data entries in the data set have. Unsupervised machine learning algorithms are particularly effective when dealing with large datasets or big data. Unsupervised machine learning can be classified into:

- Clustering,
- Association and,
- Dimensionality reduction

Clustering is a technique which groups the unlabeled data into groups based on features which are common or uncommon. Association is a method which finds relationships between the different input data in the dataset. Dimensionality reduction on the other hand is a method to reduce the number of features in a dataset and keep only the relevant ones to make processing and training of datasets a much easier process. One of the main advantages of unsupervised learning is that it doesn't require any kind of manual intervention and that it can discover links or differences in larger datasets much quicker than any kind of manual analysis. Unsupervised learning also shares the same pitfall of requiring computation time and power for processing large datasets for its training [24].

Even by applying machine learning algorithms to intrusion detection systems and datasets, intrusion detection systems cannot detect or protect against attacks that use multiple layers or different entry points to gain access and exploit vulnerabilities. This raises the need to consider different modalities and attack vectors simultaneously to determine more accurately if the chain of activities is an intrusion or normal activity. Data Fusion is one such machine learning concept and process of integrating multiple data sources together to produce more consistent, accurate and useful information than a single modality or source of information would provide. Data Fusion can help intrusion detection systems get a holistic view of networks and all the different indicators that help identify a possible intrusion would be used together to train the model. Data Fusion with respect to Network-based intrusion detection systems has been defined as "single source or multi-source data collected from the network is preprocessed to obtain a uniform data format. More refining data of greater quality is obtained through feature fusion and association, which greatly improves the

identification of malicious network behaviors. The initial decisions generated from multisource data are integrated in a decision fusion center to achieve more accurate and comprehensive inferences or decisions.” [6]. This definition is written in the intrusion detection systems where the goal of Data Fusion is to improve efficiency, accuracy rate and robustness while reducing False Negative Rate (FNR) and False Positive Rate (FPR), saving computational resources of the system.

Multi-Modal Data Fusion also has the added advantage of giving Intrusion Detection Systems the capability of detection of Advanced Persistent Threats (APTs) as it gives intrusion detection systems an insight into various system components and the data collected on them at the same time. Data fusion provides better context to intrusion detection system on security events and their associated causes giving these systems and models the ability to understand the potential impact, scope and affected systems based on the analysis. Multi-Modal Intrusion Detection System with Data Fusion can identify new attack patterns by observing correlations across different data sources, even if individual components of the attack are not recognized on their own. This given intrusion detection systems the required adaptability to evolving cyber security threats and attack methods and vectors. Having multiple data sources encapsulated together for intrusion detection purposes also has the added advantage of aiding investigation and forensic analysis processes, enabling analysts to trace the attack path used by attackers, understand the complete scope of the attacks and to gather comprehensive evidence required for their procedures.

2.5 Similar work on Machine Learning in Intrusion Detection Systems

This section aims to shed light on similar work that makes use of machine learning to improve the detection capabilities of intrusion detection systems and to study and identify the different methodologies and results achieved in these papers.

A lightweight algorithm is introduced in [25] for developing a real-time intrusion detection system that combines a deep belief network (DBN) with a support vector machine (SVM). The DBN is utilized to reduce the dimensionality of the data. To identify abnormal data, the sliding window technique leverages pattern behavior, which is stored in a pattern behavior library and regularly updated. Once the DBN reduces the features of the high-dimensional data, the SVM-based model is trained for intrusion detection. The approach is tested using the CICIDS 2017 dataset. However, this method does not incorporate data fusion techniques and relies on a single dataset for evaluation, which may lead to overfitting. Additionally, the study does not address metrics such as accuracy (ACC) and false-positive rate (FPR).

The Semantic Re-encoder and Deep Learning Model (SRDLM) for intrusion detection systems (IDS) is explored in [26]. The authors present a model for detecting network anomalies using a large semantic coding space, while disregarding word order. Although semantic re-encoding has performance constraints with network traffic, these are mitigated by integrating it with a deep learning approach. The method employs the ResNet deep learning framework, which is based on convolutional neural networks (CNN), due to its ability to generalize unknown intrusive network traffic. Experiments were conducted

using the NSL-KDD and Hduxss_data1.0 datasets. The study demonstrates strong performance with the Hduxss_data1.0 dataset and improved outcomes with the NSL-KDD dataset. However, it does not incorporate data fusion or feature selection techniques, as each dataset is evaluated separately, and the false-positive rate (FPR) is not assessed for either dataset.

In this study [27], a novel approach for intrusion detection is introduced by integrating hybrid sampling with deep hierarchical networks. Initially, the method employs one-side selection (OSS) and the synthetic minority over-sampling technique (SMOTE) to create a composite dataset for training, which helps reduce the model's training time and facilitates data preprocessing for complex networks. Subsequently, a model classifier is developed using a hierarchical network that combines convolutional neural networks (CNN) with bi-directional long short-term memory (BiLSTM). The UNSW-NB15 and NSL-KDD datasets are utilized for evaluating the model. However, the study does not incorporate a data fusion technique, as each dataset is assessed independently, and the false-positive rate (FPR) is not considered in the analysis.

A new semi-self-taught (SST) network intrusion detection system is introduced in [28], utilizing a semi-supervised discriminant autoencoder (SSDA) that requires minimal human intervention. The SSDA is implemented through a denoising autoencoder combined with a c-mean algorithm for class identification, and the experiment is conducted using the CSE-CIC-IDS2018 dataset. This approach may lead to overfitting due to reliance on a single dataset, and it lacks data fusion techniques.

Additionally, a novel XGBoost-DNN model is proposed for network intrusion detection [29]. The process begins with data preprocessing for normalization, followed by the use of the XGBoost method for feature selection in high-dimensional data. Binary classification is then performed using a deep neural network (DNN), with the Adam optimizer (AO) enhancing the learning frequency during training. Experiments are carried out using the NSL-KDD dataset, but this method also faces overfitting issues due to the use of a single dataset and lacks model evaluation and data fusion techniques.

A novel intrusion detection framework [30], DT-EnSVM, is presented, which combines ensemble learning with a data transformation method. Initially, ratio transformation is applied to create a new, balanced dataset. Support Vector Machine (SVM) classifiers are then selected for training the model. Finally, an ensemble learning-based model is constructed using a non-mixture method to combine these classifiers. Experiments are conducted using the NSL-KDD and Kyoto 2006+ datasets. However, the study does not incorporate data fusion, as each dataset is evaluated separately, and it lacks consideration of true positive rate (TPR) and F-Measure parameters.

An adaptive ensemble learning-based model is also discussed [31], emphasizing the benefits of using different base classifiers in an ensemble approach. The model employs five machine learning classifiers: Decision Tree, Support Vector Machine, Logistic Regression, K-Nearest Neighbor, and Random Forest, using the NSL-KDD dataset for experimentation. This study faces overfitting issues and lacks feature selection and data fusion techniques, as well as consideration of the false-positive rate.

Another study addresses training challenges of feed-forward neural networks (FNN) using locust swarm optimization (LSO) [32], a meta-heuristic optimization algorithm. The combination of FNN and LSO

(FNN-LSO) is claimed to enhance the overall performance of intrusion detection systems (IDS). The NSL-KDD and UNSW-NB15 datasets are used for experimentation, but the study lacks feature selection and data fusion techniques.

Lastly [33], a new model for network intrusion detection systems (NIDS) is proposed, based on a hierarchical neural network that integrates a convolutional neural network (CNN) using the LeNet-5 architecture with a long short-term memory (LSTM) neural network. Experiments are conducted using the CICIDS 2017 and UTC datasets. This study also lacks feature selection and data fusion techniques, and does not consider the false-positive rate.

The MIND (Multi-source Information Network Defense) framework, [78] presents a novel approach to intrusion detection using multi-source data fusion. The methodology involves combining two prominent network security datasets, NSL-KDD and UNSW-NB15, through a feature selection and data fusion process. The framework employs the deep learning ResNet architecture based on convolutional neural networks (CNN) for classification due to its generalization ability. MIND succeeds in improving detection accuracy and reducing false positives by leveraging complementary information from multiple data sources. The authors demonstrate that their fused dataset outperforms individual datasets when applied to machine learning algorithms like K-Nearest Neighbor (KNN) with bagging.

The paper by [79] presents a deep learning ensemble approach for network anomaly and cyber-attack detection. Their methodology involves a two-stage process, Feature engineering using a Deep Sparse AutoEncoder (DSAE) for dimensionality reduction and a stacking ensemble learning approach that combines multiple deep learning models. The paper makes use of multi source data from three distinct datasets: IoT-23, LITNET-2020, and NetML-2020. While the paper doesn't explicitly focus on overcoming overfitting, the ensemble approach inherently helps reduce overfitting by combining predictions from multiple models. The paper doesn't employ traditional data fusion techniques. However, the stacking ensemble method can be considered a form of decision-level fusion, as it combines outputs from multiple deep learning models to make a final prediction.

This paper by [80]. presents a deep learning approach for intelligent intrusion detection systems. The methodology involves using various deep learning models, including Deep Neural Networks (DNN), Recurrent Neural Networks (RNN), and Convolutional Neural Networks (CNN), to detect network intrusions. The authors evaluate their approach on multiple datasets, including KDD Cup 1999, NSL-KDD, UNSW-NB15, and WSN-DS. The paper doesn't explicitly focus on traditional data fusion techniques and neither does it suggest the fusion of datasets, it does combine information from multiple features and datasets to improve detection accuracy. The use of deep learning models inherently performs a form of feature-level fusion by learning hierarchical representations from the input data. The authors demonstrate that their deep learning approach outperforms traditional machine learning methods in terms of accuracy and false positive rates across the different datasets.

The literature on network intrusion detection contains numerous studies that makes use of different machine learning algorithms, tools and methodologies as discussed earlier. Our project addresses several key gaps in the existing literature on intrusion detection systems and pushes the boundaries of current

research in several important ways.

While some existing studies like [78] and [79] incorporate multi-source data, our project takes a more holistic approach by fusing system logs, user behavior patterns, and network traffic data. This comprehensive fusion provides a more complete view of the system state, potentially capturing complex attack patterns that may be missed by approaches focusing on fewer data sources.

Many existing studies, such as [80], evaluate their approaches on a limited number of datasets. Our project's use of both NSL-KDD and UNSW-NB15 datasets, along with their fusion, demonstrates improved generalization across different network environments and attack types. This addresses a significant gap in the literature where models often perform well on specific datasets but struggle with others.

While papers like [79] mention techniques to reduce overfitting, our project explicitly focuses on this issue through the use of data fusion. By combining diverse datasets, our approach inherently reduces the risk of overfitting to specific attack patterns or network characteristics, a common problem in many existing intrusion detection systems.

Unlike some studies that focus primarily on accuracy or detection rate, our project emphasizes balanced performance across accuracy, precision, recall, and F1-score. This comprehensive evaluation provides a more realistic assessment of the system's practical utility, addressing a gap in many existing studies that may overlook important performance aspects. By demonstrating significant performance improvements across various models, including simpler ones like Naive Bayes, our project advocates for the practical applicability of multi-modal data fusion in real-world intrusion detection scenarios. This addresses a gap in the literature where complex models often show good performance but may be challenging to implement in practice.

The literature review reveals a significant evolution in intrusion detection systems, from traditional rule-based approaches to more sophisticated machine learning and deep learning methods. Recent studies have demonstrated the potential of various algorithms, including deep belief networks, convolutional neural networks, and ensemble methods, in improving the accuracy and efficiency of intrusion detection. There is a growing recognition of the limitations of single-source data and the benefits of incorporating multiple data sources for more comprehensive threat detection.

Despite these advancements, several gaps remain in the current body of research. Many studies still rely on single datasets, potentially limiting the generalizability of their models to diverse network environments. While some studies use multiple datasets, few have explored advanced data fusion methods to fully leverage the complementary information from different sources. Overfitting concerns persist, as many approaches demonstrate high performance on specific datasets but may not generalize well to real-world scenarios with evolving threat landscapes. Some studies focus primarily on accuracy or detection rate, overlooking other crucial metrics like false positive rates and F1-scores. Additionally, as models become more complex, there's a growing need for interpretable results that can guide security professionals in their decision-making process.

Our proposed methodology aims to address these limitations in several ways. By combining the NSL-KDD and UNSW-NB15 datasets, our approach leverages diverse data sources to create a more comprehensive view of network activities and potential threats. We incorporate sophisticated feature selection techniques to identify the most relevant attributes across datasets, enhancing the model's ability to detect a wide range of attack patterns. Our ensemble learning approach, which employs multiple machine learning algorithms and combines their outputs, aims to reduce overfitting and improve the generalizability of our intrusion detection system.

Our study considers a balanced set of performance metrics, including accuracy, precision, recall, and F1-score, to provide a more holistic assessment of the model's effectiveness. We also prioritize the use of interpretable models alongside more complex ones, aiming to provide actionable insights for security professionals.

This approach is necessary to advance the field of intrusion detection systems for several reasons. It addresses the growing complexity of cyber threats by incorporating diverse data sources and advanced machine learning techniques. It aims to improve the practical applicability of intrusion detection systems by focusing on generalizability and reducing false positives. Our methodology contributes to the development of more adaptive and robust security solutions capable of detecting both known and emerging threats. By making our code and methodology publicly available, we foster reproducibility and further advancement in the field of multi-modal intrusion detection.

In conclusion, our methodology represents a significant step forward in addressing the limitations of current intrusion detection systems and paves the way for more effective, adaptive, and reliable cybersecurity solutions.

3. Methodology

In this section we will discuss in detail the proposed methodology for implementing Multimodal Intrusion Detection using Machine Learning and Data Fusion. The five major steps in this project are dataset selection, feature selection, data fusion, machine learning implementation and representation of results as shown in the image below.

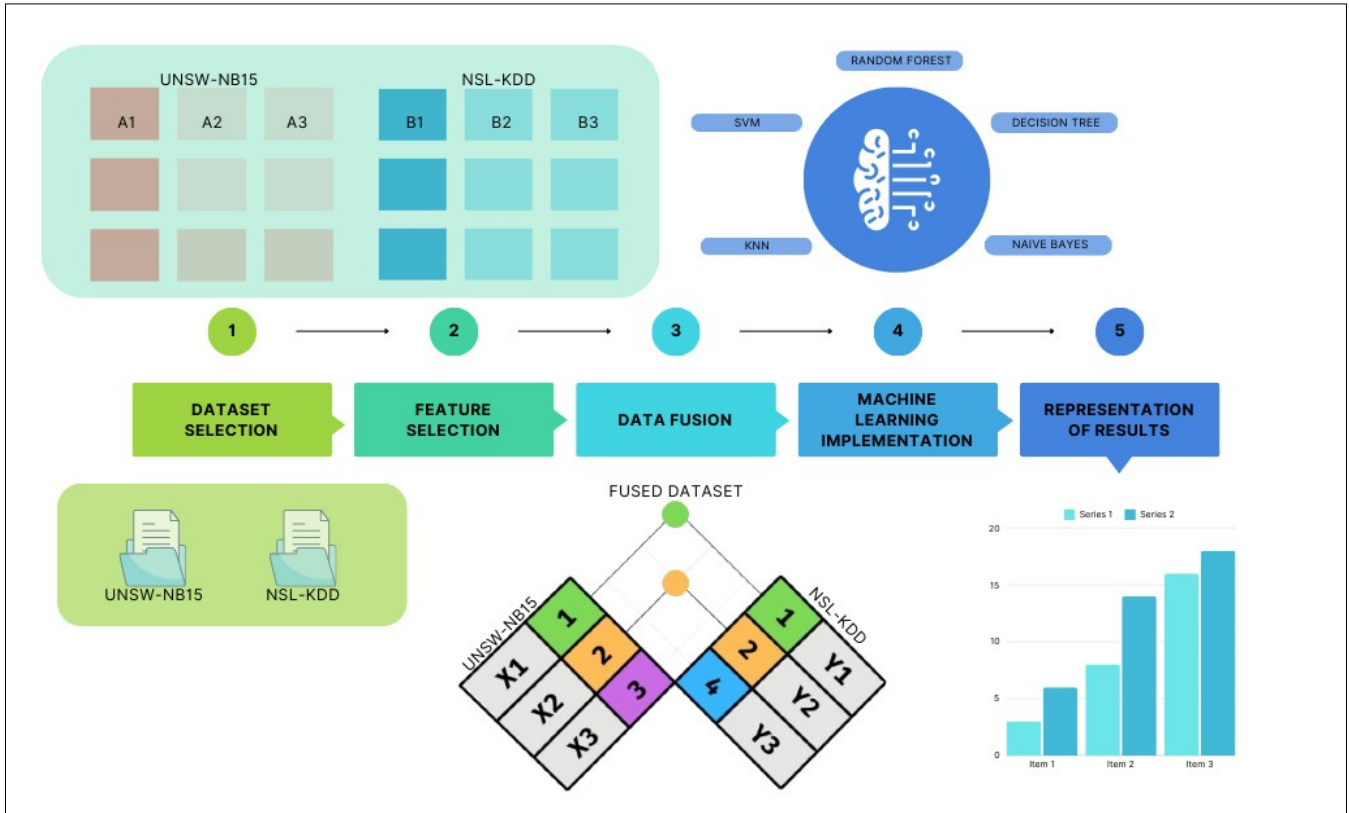


Figure 1: Working Design of Methodology

In the dataset selection step, two publicly available datasets NSL-KDD [35] and UNSW-NB15 [36] are selected for building the model. The reasons for doing so are mentioned in the Dataset Selection section below. The first steps involved in using a dataset for machine learning is always data preprocessing and feature selection. This will be the second step of the process and next we move on to the Data Fusion process which is carried out at the feature level which merges the two datasets into one fused dataset which gives a lot more context and detail to the intrusion detection system to make its decisions. This step is followed by implementation of machine learning algorithms on the fused dataset and the original datasets UNSW-NB15 and NSL-KDD to compare and contrast the effectiveness of data fusion on a range of evaluation metrics like Accuracy, Precision, Recall, F1-Score, Validation Accuracy, Cross-Validation Score and Test Accuracy. The final step of this methodology is to graphically represent the findings and

comparisons using a suitable graphs or charts. The following subsections delves into the different steps of the methodology in detail.

In studies considering the applications of machine learning algorithms in intrusion detection systems, the main concerns and challenges encountered by the models include overfitting of model, generalizability of the intrusion detection system to different environments and threats and handling of false positives and false negatives. This project and the suggested methodology takes concrete steps to address and mitigate these challenges in every step.

The data fusion approach employed in the project, combining the UNSW-NB15 and NSL-KDD datasets, plays a crucial role in addressing overfitting, enhancing generalizability, and managing false positives and negatives in the Intrusion Detection System (IDS). This approach significantly increases the diversity and volume of network traffic patterns and attack scenarios in the training data, which is fundamental in combating overfitting. By exposing the model to a broader range of both normal and malicious traffic patterns from different time periods and network environments, the likelihood of the model memorizing dataset-specific nuances is reduced. Instead, it is encouraged to learn more generalizable patterns of network behavior and intrusion attempts.

The fusion of these datasets enhances the IDS's ability to generalize to different types of network environments and threats. As highlighted in the research in [44], combining diverse datasets helps create a more robust model that can potentially adapt better to various network scenarios. This approach is particularly valuable given the dynamic nature of cyber threats and the diversity of network architectures. By training on a fused dataset that encompasses a wide range of attack types and normal traffic patterns from different sources, the model is better equipped to identify both known and potentially new types of intrusions across various network environments.

Regarding false positives and false negatives, the data fusion approach provides several advantages. Firstly, it allows the model to learn from a more comprehensive set of examples, potentially reducing false positives by exposing it to a wider range of normal network behaviors that might otherwise be flagged as suspicious. Similarly, it can help reduce false negatives by including a more diverse set of attack patterns in the training data. The use of multiple evaluation metrics in our project, including precision (related to false positive rate) and recall (related to false negative rate), indicates a focus on balancing these two types of errors.

Moreover, the data fusion approach, combined with the Recursive Feature Elimination (RFE) technique using Random Forest, helps in selecting the most relevant features across both datasets. This feature selection process is crucial in reducing noise that could lead to false alerts while focusing on the most informative indicators of intrusion. By basing feature selection on a fused dataset, the chosen features are likely to be more robust and generalizable across different network environments, further contributing to the reduction of both false positives and false negatives.

4. Dataset Selection

Datasets are crucial for evaluating intrusion detection systems, but real-time network traffic data is often inaccessible for research due to privacy concerns. Consequently, researchers rely heavily on public datasets such as KDD CUP99 [37], NSL-KDD [35], Kyoto 2006 [39], UNSW-NB15 [36], NGIDS-DS, ISCX [41], and CICIDS 2017[40]. For this project, NSL-KDD and UNSW-NB15 were chosen for experiments. Data fusion requires at least two datasets, and a key requirement for fusion is having at least one common column between them. These two datasets were selected because they meet this criterion.

The NSL-KDD dataset is a public dataset developed to address issues found in the earlier KDD Cup '99 dataset. A statistical analysis of the KDD Cup '99 dataset [37] revealed significant concerns that could affect the accuracy of intrusion detection systems (IDS) and lead to misleading evaluations. One major issue with the KDD dataset was the large number of duplicate records, with approximately 78% of the training data and 75% of the test data being duplicates. This redundancy could bias machine learning models towards normal instances, hindering their ability to learn from irregular, potentially more harmful instances. To address these issues, the NSL-KDD dataset was created by removing duplicate records. It contains 125,973 records in the training set and 22,544 records in the test set, making it feasible to use the entire dataset without random sampling. This consistency allows for comparable results across different studies. The NSL-KDD dataset includes 22 types of intrusion attacks and 41 attributes, with 21 attributes related to the connection itself and 19 describing the nature of connections within the same host. The intrusion data is categorized into remote to local (R2L), denial of service (DoS), user 2 root (U2R), and Probe.

The UNSW-NB15 dataset [36] is a comprehensive network intrusion dataset designed to evaluate intrusion detection systems (IDS). Developed by the Cyber Range Lab of UNSW Canberra, it was created using the IXIA PerfectStorm tool to simulate a mix of real-world normal activities and synthetic attack behaviors. The dataset captures 100 GB of raw network traffic, which includes both benign and malicious activities, over a two-day period. It incorporates nine types of attacks: Fuzzers, Analysis, Backdoors, Denial of Service (DoS), Exploits, Generic, Reconnaissance, Shellcode, and Worms. These attacks are representative of contemporary threats, making the dataset highly relevant for testing IDS. The dataset's structure includes 49 features extracted using tools like Argus and Bro-IDS, which are categorized into Basic, Content, Time, and additional generated features. This detailed feature set allows for a robust analysis of network traffic, providing a rich basis for developing and testing machine learning models for intrusion detection. UNSW-NB15 is considered superior to older datasets like KDD Cup '99 due to its realistic representation of modern network traffic and diverse attack types. Unlike KDD Cup '99, which suffers from issues such as duplicate records leading to potential overfitting, UNSW-NB15 offers a cleaner and more balanced dataset. This makes it more suitable for developing models that generalize well to real-world scenarios. The dataset's comprehensive nature and the inclusion of both normal and various attack traffic make it an excellent resource for researchers aiming to improve IDS technologies. Its ability to provide a realistic testing ground for IDS models ensures that the models can be effectively evaluated for their ability to detect and respond to a

wide range of network intrusions.

While both NSL-KDD and UNSW-NB15 are widely used in intrusion detection research, they offer different strengths and represent different eras of network traffic. NSL-KDD, being an improvement over the older KDD Cup '99 dataset, provides a good baseline for comparison with newer techniques. It's particularly useful for evaluating the detection of traditional attack types. On the other hand, UNSW-NB15 represents more modern network traffic patterns and attack types, making it more relevant for current cybersecurity challenges. By using both datasets, our research benefits from a historical perspective (NSL-KDD) as well as a contemporary view (UNSW-NB15) of network intrusions.

The characteristics of these datasets significantly influence the development and evaluation of intrusion detection systems. NSL-KDD's balanced nature and removal of redundant records make it ideal for training models that can distinguish between normal and anomalous behavior without bias towards overrepresented attack types. UNSW-NB15's diverse feature set and modern attack representations allow for the development of more sophisticated detection algorithms capable of identifying subtle and complex intrusion patterns. By combining these datasets, we aim to create a more robust and versatile intrusion detection model that can handle both well-established and emerging threat patterns.

To go a little deeper into the features, columns and parameters, we need to first identify and list down the parameters in both datasets. This step would help us further identify the common features and columns which would be crucial for the data fusion process which requires the datasets in question to have common columns for the successful execution of the process.

Table 3: NSL-KDD and UNSW-NB15 list of features [42]

NSL-KDD		UNSW-NB15	
Feature	Name	Feature	Name
f_{1-1}	duration	f_{2-1}	Dur
f_{1-2}	protocol_type	f_{2-2}	Proto
f_{1-3}	service	f_{2-3}	Service
f_{1-4}	flag	f_{2-4}	State
f_{1-5}	src_bytes	f_{2-5}	Spkts
f_{1-6}	dst_bytes	f_{2-6}	Dpkts
f_{1-7}	land	f_{2-7}	Sbytes
f_{1-8}	wrong_fragment	f_{2-8}	Dbytes
f_{1-9}	urgent	f_{2-9}	Rate
f_{1-10}	hot	f_{2-10}	Sttl
f_{1-11}	num_failed_logins	f_{2-11}	Dttl
f_{1-12}	logged_in	f_{2-12}	Sload
f_{1-13}	lnum_compromised	f_{2-13}	Dload
f_{1-14}	lroot_shell	f_{2-14}	Sloss
f_{1-15}	lsu_attempted	f_{2-15}	Dloss

f_{1-16}	lnum_root	f_{2-16}	Sinpkt
f_{1-17}	lnum_file_creations	f_{2-17}	Dinpkt
f_{1-18}	lnum_shells	f_{2-18}	Sjit
f_{1-19}	lnum_access_files	f_{2-19}	Djit
f_{1-20}	lnum_outbound_cmds	f_{2-20}	Swin
f_{1-21}	is_host_login	f_{2-21}	Stcpb
f_{1-22}	is_guest_login	f_{2-22}	Dtcpb
f_{1-23}	count	f_{2-23}	Dwin
f_{1-24}	srv_count	f_{2-24}	Tcprt
f_{1-25}	serror_rate	f_{2-25}	Synack
f_{1-26}	srv_serror_rate	f_{2-26}	Ackdat
f_{1-27}	rerror_rate	f_{2-27}	Smean
f_{1-28}	srv_rerror_rate	f_{2-28}	Dmean
f_{1-29}	same_srv_rate	f_{2-29}	trans_depth
f_{1-30}	diff_srv_rate	f_{2-30}	response_body_len
f_{1-31}	srv_diff_host_rate	f_{2-31}	ct_srv_src
f_{1-32}	dst_host_count	f_{2-32}	ct_state_ttl
f_{1-33}	dst_host_srv_count	f_{2-33}	ct_dst_ltm
f_{1-34}	dst_host_same_srv_rate	f_{2-34}	ct_src_dport_ltm
f_{1-35}	dst_host_diff_srv_rate	f_{2-35}	ct_dst_sport_ltm
f_{1-36}	dst_host_same_src_port_rate	f_{2-36}	ct_dst_src_ltm
f_{1-37}	dst_host_srv_diff_host_rate	f_{2-37}	is_ftp_login
f_{1-38}	dst_host_serror_rate	f_{2-38}	ct_ftp_cmd
f_{1-39}	dst_host_srv_serror_rate	f_{2-39}	ct_flw_http_mthd
f_{1-40}	dst_host_rerror_rate	f_{2-40}	ct_src_ltm
f_{1-41}	dst_host_srv_rerror_rate	f_{2-41}	ct_srv_dst
		f_{2-42}	is_sm_ips_ports

The common features for both datasets are listed in the table below:

Table 4: Similarities of the features in NSL-KDD and UNSW-NB15 [42]

Category	NSL-KDD	UNSW-NB15
Common features	$f_{1-1}, f_{1-2}, f_{1-3}, f_{1-5}, f_{1-6}$	$f_{2-1}, f_{2-2}, f_{2-3}, f_{2-7}, f_{2-8}$
Features that use connection flags	$f_{1-4}, f_{1-9}, f_{1-24}, f_{1-25}, f_{1-29}, f_{1-30}, f_{1-38},$ $f_{1-39}, f_{1-40}, f_{1-41}$	$f_{2-4}, f_{2-24}, f_{2-25}, f_{2-26}$
Features that count connections	$f_{1-5}, f_{1-6}, f_{1-23}, f_{1-24}, f_{1-25}, f_{1-26}, f_{1-27},$ $f_{1-28}, f_{1-29}, f_{1-30}, f_{1-31}, f_{1-32}, f_{1-33}, f_{1-34},$ $f_{1-35}, f_{1-36}, f_{1-37}, f_{1-38}, f_{1-39}, f_{1-40}, f_{1-41}$	$f_{2-31}, f_{2-33}, f_{2-34}, f_{2-35}, f_{2-36}, f_{2-40},$ f_{2-41}
Size-based features (transmitted bits, bytes, or packets)	f_{1-5}, f_{1-6}	$f_{2-5}, f_{2-6}, f_{2-7}, f_{2-8}, f_{2-12}, f_{2-13}, f_{2-14}, f_{2-15},$ $f_{2-27}, f_{2-28}, f_{2-30}$
Features that calculates time (e.g., connection duration)	$f_{1-1}, f_{1-23}, f_{1-28}$	$f_{2-1}, f_{2-10}, f_{2-11}, f_{2-18}, f_{2-19}, f_{2-24}, f_{2-25}, f_{2-}$

The process of selecting appropriate datasets for this project presented several significant challenges, each requiring careful consideration and analysis.

Relevance and Currency posed a primary challenge. While the NSL-KDD dataset is widely used in intrusion detection research, it represents older network traffic patterns that may not fully reflect current cyber threats. Balancing this historical perspective with the more current UNSW-NB15 dataset required careful consideration. We needed to weigh the importance of historical attack patterns, which provide a foundation for understanding long-standing threats, against contemporary patterns that represent the evolving nature of cyber attacks. This balance was crucial to ensure our intrusion detection system could identify both well-established and emerging threat patterns.

Dataset Compatibility was another critical factor in our selection process. For effective data fusion, it was essential that the selected datasets had sufficient common features. This necessitated a detailed analysis of both datasets' structures and attributes. We had to identify overlapping features, understand their representations in each dataset, and determine how they could be harmonized for fusion. This process was time-consuming but crucial for ensuring that our multi-modal approach could leverage complementary information from both datasets effectively.

Assessing the Representativeness of these datasets in relation to real-world network traffic and attack scenarios presented a significant challenge. Given the rapid evolution of cyber threats, we had to critically evaluate whether these datasets truly captured the complexity and diversity of current network environments. This involved analyzing the types of attacks represented, the network architectures simulated, and the traffic patterns included. We also had to consider how well these datasets reflected modern network protocols and emerging technologies that might influence attack vectors.

Understanding and accounting for Bias and Limitations in each dataset was necessary to ensure the validity of our results. Every dataset has inherent biases, whether in the types of attacks represented, the network environments simulated, or the data collection methods used. We had to thoroughly investigate these biases and limitations, considering factors such as class imbalance, potential overfitting risks, and any known issues with the datasets. This understanding was crucial for interpreting our results accurately and avoiding misleading conclusions.

Finally, Ethical Considerations played a crucial role in our dataset selection. We had to ensure that the use of these datasets complied with ethical standards and data protection regulations. This involved verifying the datasets' origins, understanding any privacy implications of the data they contained, and ensuring that our use of the data aligned with the terms under which they were released. We also considered the broader ethical implications of our research, such as the potential for bias in intrusion detection systems and the responsible use of machine learning in cybersecurity contexts. Addressing these challenges required a thorough review of literature, careful analysis of dataset characteristics, and alignment with our project's objectives and available resources. The selection of NSL-KDD and UNSW-NB15 represents a balanced approach that addresses these challenges while providing a comprehensive foundation for our multi-modal intrusion detection system.

5. Feature Selection

Feature selection is a crucial process in machine learning and data analysis that involves identifying and selecting a subset of relevant features (or variables) from a larger set of available features. The primary goal of feature selection is to improve the performance of a machine learning model by reducing the dimensionality of the data, which can lead to faster training times, reduced overfitting, and improved model interpretability. By focusing on the most informative features, feature selection helps in removing redundant or irrelevant data that may not contribute to the predictive power of the model. Techniques for feature selection can be broadly categorized into three types: filter methods, which evaluate the relevance of features based on statistical measures; wrapper methods, which use a predictive model to assess the importance of features; and embedded methods, which perform feature selection as part of the model training process.

The feature selection process for the UNSW-NB15 and NSL-KDD datasets involves identifying and selecting a subset of relevant features from each dataset to improve the performance of intrusion detection systems. This process is crucial due to the high dimensionality of the datasets, which can lead to increased training times and resource consumption. In this project's code, feature selection is accomplished using machine learning techniques, specifically Recursive Feature Elimination (RFE) with a Random Forest classifier. Recursive Feature Elimination (RFE) with a Random Forest classifier is a popular technique for feature selection in machine learning, combining the strengths of both RFE and Random Forests to identify the most relevant features for a given task. RFE is an iterative process that works by recursively removing the least important features based on the model's performance, ultimately narrowing down the feature set to the most informative ones. When paired with a Random Forest classifier, this method leverages the inherent feature importance scoring of Random Forests, which are ensemble models that consist of multiple decision trees. Random Forests evaluate the importance of each feature by measuring the decrease in the model's accuracy when the feature is excluded. During the RFE process, the Random Forest model is trained on the dataset, and feature's importance is computed. The least important features are pruned, and the process is repeated until the desired number of features is reached. This combination is particularly effective because Random Forests are robust to overfitting and can handle large datasets with many features [38]. By using RFE with a Random Forest classifier, practitioners can efficiently reduce dimensionality, improve model interpretability, and enhance predictive performance by focusing on the most significant features.

For both datasets, the process begins by separating the features and labels, followed by encoding categorical variables into numerical format using 'LabelEncoder'. Any non-numeric columns are converted to numeric, and missing values are imputed with the mean of the respective columns. RFE is then applied to select the top features based on their importance scores, as determined by the Random Forest model. This method iteratively removes the least important features, refining the feature set to include only those that contribute most significantly to the model's predictive power. The selected features are then used for further analysis and model training, ensuring that the intrusion detection system can efficiently and effectively identify malicious activities with reduced computational overhead. This approach not only enhances model accuracy but also improves generalization by focusing on the

most informative attributes from each dataset.

The Random Forest classifier's ensemble nature provides a more stable and reliable feature importance ranking compared to single-model approaches, reducing the risk of overfitting to dataset-specific nuances. This is crucial when working with datasets from different time periods and network environments, as it helps identify truly relevant features across both datasets. Furthermore, RFE with Random Forest inherently captures non-linear relationships and feature interactions, which are prevalent in intrusion detection scenarios where attack patterns can be complex and multifaceted. This method's ability to handle both numerical and categorical features without extensive preprocessing aligns well with the diverse feature types present in network traffic data.

In comparison to other methods like correlation-based feature selection or principal component analysis (PCA), RFE with Random Forest preserves the original features [43], maintaining interpretability which is vital in security applications. It also outperforms filter methods by considering the collective impact of features on model performance rather than evaluating them in isolation. The iterative nature of RFE allows for fine-tuning the feature set size, enabling you to balance between model complexity and performance, which is particularly beneficial when fusing datasets with different original feature counts. Given the project's goal of creating a robust, generalized intrusion detection model from fused datasets, RFE with Random Forest offers a powerful, adaptable, and interpretable feature selection approach that aligns closely with the objectives.

The feature selection process for this project presented several significant challenges, each requiring careful consideration and innovative solutions. One of the primary challenges was the heterogeneity of the NSL-KDD and UNSW-NB15 datasets, which were created at different times and represent different network environments. This disparity led to variations in feature representations and relevance across the datasets, requiring careful analysis to create a unified feature set. The high dimensionality of both datasets (41 features in NSL-KDD and 49 in UNSW-NB15) introduced the "curse of dimensionality" problem, necessitating a significant reduction in features without losing critical information. We also observed significant correlations between features, presenting the challenge of balancing the removal of redundant information while retaining necessary discriminative power. The temporal relevance of features, particularly from the older NSL-KDD dataset, required assessment in the context of modern network environments. Handling a mix of categorical and continuous features across both datasets added complexity to the feature selection process. Striking a balance between creating a feature set that generalizes well across different network environments and attack patterns, while still capturing specific characteristics of various attack types, was a constant challenge. Additionally, the computational intensity of the Recursive Feature Elimination process, combined with the use of Random Forest classifiers, presented challenges in terms of processing time and resource allocation. These challenges necessitated a careful, iterative approach to feature selection, balancing statistical techniques with domain knowledge to ensure that our final feature set was not only statistically significant but also meaningful in the context of network intrusion detection.

Table 5: Identified features from feature selection on NSL-KDD and UNSW-NB15

Dataset	Identified Features
NSL-KDD	'protocol_type', 'flag', 'src_bytes', 'num_failed_logins', 'is_guest_login', 'count', 'srv_error_rate', 'srv_diff_host_rate', 'dst_host_count', 'dst_host_srv_count', 'dst_host_same_srv_rate', 'dst_host_diff_srv_rate', 'dst_host_srv_error_rate'
UNSW-NB15	'id', 'dur', 'sbytes', 'dbytes', 'rate', 'sttl', 'sload', 'synack', 'ct_state_ttl', 'ct_dst_sport_ltm', 'ct_dst_src_ltm', 'ct_srv_dst', 'attack_cat'

6. Data Fusion

Data fusion is a sophisticated process of integrating multiple data sources to produce more consistent, accurate, and valuable information than any single source could provide alone. Data fusion (DF) emerged as a concept in the military domain during the 1980s, initially referred to as "intelligence synthesis." The Joint Directors of Laboratories (JDL) provided a military-focused definition, describing Data Fusion as a process that associates, correlates, and combines data from various sources to refine position and identity estimates, and assess situations and threats comprehensively and promptly [6].

In the context of intrusion detection systems, data fusion plays a crucial role in improving accuracy and effectiveness. By combining data from diverse sources such as network traffic logs, system logs, and threat intelligence feeds, data fusion enables IDS to create a more comprehensive and nuanced understanding of potential security threats. This multi-faceted approach allows for the detection of complex attack patterns that might not be apparent when analyzing data from a single source. The development of data fusion techniques for IDS was driven by the need to address the limitations of traditional single-source detection methods, which often struggle with high false positive rates and the inability to detect sophisticated, multi-vector attacks.

In this project and code [45], data fusion is applied at the feature level, combining attributes from two distinct datasets: NSL-KDD and UNSW-NB15. This approach falls under the category of low-level data fusion, specifically feature-level fusion, which is one of the three main subcategories of data fusion methods alongside decision-level fusion and sensor-level fusion. The feature-level fusion technique was chosen for its ability to create a more comprehensive and robust dataset by leveraging the strengths of both input datasets. The fusion process begins with a careful analysis of both datasets to identify common features. These shared attributes form the basis of the fused dataset. The method then involves concatenating the selected features from both datasets, effectively creating a new, enriched dataset that captures a wider range of network behaviors and attack patterns.

The fusion technique employed in this code is a form of feature-level fusion, specifically tailored for combining intrusion detection datasets. This approach can be characterized as a "Selective Feature Concatenation with Dataset Identification and Post-Fusion Normalization." At its core, the technique involves carefully preparing and combining features from two distinct datasets: NSL-KDD and UNSW-NB15. The process begins with a feature selection step, where the most relevant features from each dataset are identified (using the previously mentioned Recursive Feature Elimination with Random Forest). This selective approach ensures that only the most informative attributes from each dataset are included in the fusion process, potentially reducing noise and improving the overall quality of the fused dataset.

The `prepare_for_fusion` function plays a crucial role in this fusion technique. It meticulously prepares each dataset by ensuring consistency in the feature set, adding a dataset identifier, and handling any missing features. This step is vital for maintaining the integrity of each dataset's contribution while also creating a uniform structure for the fused data. The addition of a dataset identifier is particularly

noteworthy, as it allows the subsequent machine learning models to potentially learn dataset-specific patterns or biases, which could be valuable for improving generalization across different network environments. The actual fusion occurs through a straightforward yet effective concatenation of the prepared datasets. This approach preserves all selected features from both sources, allowing the model to leverage the unique characteristics of each dataset. The use of `pd.concat()` to combine the training and testing sets separately ensures that the integrity of the train-test split is maintained, which is crucial for proper model evaluation. A key aspect of this fusion technique is the post-fusion normalization step. By applying `StandardScaler` to the numerical features of the combined dataset, the method ensures that all features, regardless of their origin, are on the same scale. This normalization is critical for many machine learning algorithms, as it prevents features with larger scales from dominating the model training process and allows for fair comparison and integration of features from different sources.

This approach is particularly beneficial for intrusion detection systems as it provides a more diverse and representative set of features, potentially improving the model's ability to detect a broader spectrum of cyber threats. The choice of feature-level fusion over other methods like decision-level fusion was made due to its ability to preserve more of the original information from both datasets. This preservation of data granularity allows the subsequent machine learning models to learn from a richer feature set, potentially leading to more accurate and generalizable intrusion detection capabilities [46].

This fusion approach is particularly well-suited for intrusion detection tasks, where the diversity of network environments and attack patterns necessitates a robust and adaptable model. By combining data from NSL-KDD, which represents older network traffic patterns, with UNSW-NB15, which captures more recent cyber threats, the resulting fused dataset provides a comprehensive view of both historical and contemporary intrusion scenarios. This breadth of data can potentially lead to a more generalizable and effective intrusion detection model, capable of identifying a wide range of attack patterns across various network environments.

Additionally, this method allows for better handling of the complementary nature of the two datasets, where each might excel in detecting certain types of attacks or network behaviors. By employing this feature-level fusion technique, the study aims to create a more comprehensive training dataset that can lead to the development of more robust and versatile intrusion detection models, capable of identifying a wider range of cyber threats across different network environments.

The following table elaborates the results of the Data Fusion of the UNSW-NB15 and the NSL-KDD datasets

Table 6: Features of fused dataset as a result of Data Fusion of UNSW-NB15 and NSL-KDD

Fused Dataset Features	Fused Dataset Values
Fused training dataset rows	208 305
Fused training dataset columns	28
Fused testing dataset rows	197 885
Fused testing dataset columns	28
Fused dataset columns	'protocol_type', 'flag', 'src_bytes', 'num_failed_logins', 'is_guest_login', 'count', 'srv_error_rate', 'srv_diff_host_rate', 'dst_host_count', 'dst_host_srv_count', 'dst_host_same_srv_rate', 'dst_host_diff_srv_rate', 'dst_host_srv_error_rate', 'dataset', 'label', 'id', 'dur', 'sbytes', 'dbytes', 'rate', 'sttl', 'sload', 'synack', 'ct_state_ttl', 'ct_dst_sport_ltm', 'ct_dst_src_ltm', 'ct_srv_dst', 'attack_cat'

7. Machine Learning Implementation

Implementing machine learning for this intrusion detection project offers several significant advantages over traditional rule-based systems. Machine learning algorithms can analyze vast amounts of network traffic data to identify complex patterns and anomalies that may indicate potential security threats. This approach allows for more accurate and adaptive intrusion detection, as the models can learn from new data and evolve to recognize emerging attack vectors. Machine learning techniques, such as the

- Support Vector Machine (SVM),
- Decision Tree,
- Random Forest,
- K-Nearest Neighbors (KNN), and
- Naive Bayes algorithms

are used in this project, can process high-dimensional data efficiently and make real-time predictions, enabling faster response times to potential threats. Furthermore, by fusing data from multiple sources (UNSW-NB15 and NSL-KDD datasets in this case), machine learning models can leverage a broader range of features and attack patterns, potentially improving the overall detection accuracy and reducing false positives. This data-driven approach also allows for better generalization to unseen attacks and can adapt to the ever-changing landscape of cyber threats, making it a powerful tool for enhancing the effectiveness and robustness of intrusion detection systems.

The criteria for selecting machine learning algorithms from the wide range of algorithms known for this project are as follows:

- Support Vector Machine (SVM):
Support Vector Machine (SVM) is a powerful and versatile machine learning algorithm that has shown significant effectiveness in intrusion detection systems (IDS). SVM's importance in intrusion detection stems from its ability to handle high-dimensional data, its effectiveness in both binary and multi-class classification problems, and its strong generalization capabilities. In the context of network security, SVM excels at identifying complex patterns and relationships in network traffic data, making it particularly adept at detecting both known and novel intrusion attempts. The selection of SVM for our project is justified by several key advantages it offers over other machine learning algorithms. Firstly, SVM's ability to work well with high-dimensional data is crucial for intrusion detection, where numerous features of network traffic need to be analyzed simultaneously [48]. This characteristic allows SVM to maintain its performance even when dealing with the complex, multi-faceted nature of modern cyber attacks.

Secondly, SVM's speed of execution is a critical factor in real-time intrusion detection, where rapid identification of threats is essential. The scalability of SVM is another significant advantage, as it remains relatively insensitive to the number of data points, allowing it to

handle large-scale network traffic data efficiently [49].

Furthermore, SVM's classification complexity does not depend on the dimensionality of the feature space, which is particularly beneficial in intrusion detection scenarios where the number of relevant features can be substantial. Lastly, SVM has demonstrated superior performance in terms of accuracy and false positive rates compared to other algorithms in various intrusion detection studies [50]. These characteristics make SVM an optimal choice for our project, offering a balance of accuracy, speed, and scalability that is crucial for effective intrusion detection in modern network environments.

- **Decision Tree:**

Decision Trees are a fundamental and widely used machine learning algorithm in the field of intrusion detection systems (IDS) due to their interpretability, efficiency, and effectiveness. In the context of network security, Decision Trees excel at capturing complex decision boundaries and can effectively model the hierarchical nature of many intrusion detection rules. As demonstrated by [51], Decision Trees have shown superior performance in classifying network attacks, achieving high accuracy and low false positive rates.

The selection of Decision Trees for our project is justified by several key advantages they offer over other machine learning algorithms. Firstly, Decision Trees provide unparalleled interpretability, which is crucial in the security domain where understanding the reasoning behind classifications is often as important as the classification itself. This transparency allows security analysts to easily comprehend and validate the model's decision-making process, enhancing trust in the system. Secondly, as shown by [52] in a comprehensive survey, Decision Trees are capable of handling both numerical and categorical data without requiring extensive data preprocessing, which is beneficial when dealing with diverse network traffic features. They are also robust to outliers and missing values, common challenges in real-world network data. Furthermore, study [40] demonstrated that Decision Trees perform implicit feature selection during the training process, automatically identifying the most relevant attributes for classification. This characteristic is particularly valuable in intrusion detection, where the number of potential features can be overwhelming.

Decision Trees also offer computational efficiency, both in training and prediction phases, making them suitable for real-time intrusion detection scenarios. This efficiency is crucial in high-speed network environments, as highlighted by [53], who showed that Decision Tree-based IDS can process network traffic in real-time with high accuracy. Lastly, in a comparative study by [54], Decision Trees demonstrated competitive performance in intrusion detection, often rivaling more complex algorithms in terms of accuracy and detection rates. These attributes make Decision Trees an excellent choice for our project, offering a balance of interpretability, efficiency, and effectiveness that is crucial for developing a robust and trustworthy intrusion detection system.

- **Random Forest:**

Random Forest is an ensemble learning method that has gained significant prominence in intrusion detection systems (IDS) due to its robustness, accuracy, and ability to handle high-dimensional data. As described by [55], Random Forest constructs multiple decision trees and combines their outputs to make predictions, which inherently provides a form of feature selection and reduces overfitting. In the context of intrusion detection, Random Forest has demonstrated superior performance in several studies. For instance, [54] showed that Random Forest outperformed traditional classifiers in effectively classifying network attacks.

The selection of Random Forest for our project is justified by several key advantages it offers over other machine learning algorithms. Firstly, Random Forest excels at handling the high-dimensional feature spaces typical in network traffic data, as noted by [56] in their comprehensive survey of Random Forest-based methods for IDS. Secondly, Random Forest provides built-in feature importance measures, which is crucial for understanding which network attributes are most indicative of intrusions. This aspect is particularly valuable in our project, as demonstrated by [57], who used Random Forest for feature selection in IDS, significantly improving detection accuracy while reducing computational overhead.

Furthermore, Random Forest's ensemble nature makes it robust against noise and outliers, which are common in network traffic data. This robustness is essential for maintaining high detection rates in real-world scenarios where network behavior can be unpredictable. Additionally, Random Forest has shown excellent performance in handling imbalanced datasets, a common challenge in intrusion detection where normal traffic often far outweighs malicious traffic. This capability is highlighted in the work of [58], who proposed a fuzziness-based semi-supervised learning approach using Random Forest for IDS, achieving high accuracy even with limited labeled data.

In the context of our project, the ability of Random Forest to provide probabilistic outputs for each class is particularly valuable, allowing for more nuanced decision-making in intrusion detection. Moreover, its relatively low computational cost during prediction makes it suitable for real-time intrusion detection applications. These characteristics, combined with its proven track record in numerous IDS studies, make Random Forest an optimal choice for our project, offering a balance of accuracy, interpretability, and efficiency that is crucial for developing a robust and effective intrusion detection system.

- **K-Nearest Neighbors (KNN):**

K-Nearest Neighbors (KNN) is a simple yet effective machine learning algorithm that has shown significant utility in intrusion detection systems (IDS). As demonstrated by [59], KNN can be effectively adapted for intrusion detection by treating system calls as "words" and program executions as "documents," achieving high accuracy in classifying program behavior as normal or intrusive. The importance of KNN in intrusion detection stems from its ability to handle non-linear decision boundaries and its effectiveness in high-dimensional spaces, which are common characteristics of network traffic data.

The selection of KNN for our project is justified by several key advantages it offers over other machine learning algorithms. Firstly, KNN is particularly well-suited for detecting novel attacks, as it can identify anomalies based on their dissimilarity to known normal patterns. This capability is crucial in the ever-evolving landscape of cyber threats, as highlighted by [60] in their comprehensive survey of intrusion detection techniques. Secondly, KNN's instance-based learning approach allows for easy updates to the model as new data becomes available, making it adaptable to changing network environments. This adaptability is essential for maintaining the effectiveness of an IDS over time, as noted by [61] in their study on adaptive intrusion detection systems.

Furthermore, KNN has shown competitive performance in various intrusion detection studies. For example, [62] demonstrated that KNN outperformed several other algorithms, including Support Vector Machines and Decision Trees, in terms of accuracy and false positive rates when applied to the CIDDs-001 dataset. The simplicity of KNN also makes it computationally efficient, which is crucial for real-time intrusion detection in high-speed network environments. Additionally, KNN's ability to provide probabilistic outputs for each class allows for more nuanced decision-making in intrusion detection scenarios. In the context of our project, KNN's interpretability is another significant advantage. Unlike black-box models, the decision-making process of KNN can be easily understood and explained, which is crucial in security applications where transparency is often required. This interpretability, combined with its effectiveness in handling high-dimensional data and its adaptability to new patterns, makes KNN an excellent choice for our intrusion detection project, offering a balance of performance, flexibility, and transparency.

- Naive Bayes:

Naive Bayes is a probabilistic machine learning algorithm that has shown significant utility in intrusion detection systems (IDS) due to its simplicity, efficiency, and effectiveness in handling high-dimensional data. As demonstrated by [63], Naive Bayes can be effectively applied to intrusion detection tasks, particularly when combined with feature reduction techniques to improve computational efficiency and detection accuracy.

The selection of Naive Bayes for our project is justified by several key advantages it offers over other machine learning algorithms. Firstly, Naive Bayes is particularly well-suited for handling the high-dimensional feature spaces typical in network traffic data, as it assumes feature independence. This characteristic allows it to perform well even with limited training data, which is crucial in the rapidly evolving landscape of cyber threats. Secondly, Naive Bayes offers computational efficiency in both training and prediction phases, making it suitable for real-time intrusion detection in high-speed network environments. This efficiency is highlighted by [64], who demonstrated the algorithm's effectiveness in detecting various types of network intrusions with low computational overhead.

Furthermore, Naive Bayes has shown competitive performance in various intrusion detection

studies. For example, [65] demonstrated that Naive Bayes outperformed several other algorithms, including Decision Trees and Neural Networks, in terms of detection accuracy for certain types of attacks. The simplicity of Naive Bayes also contributes to its interpretability, which is crucial in security applications where understanding the decision-making process is often required.

In the context of our project, Naive Bayes' ability to handle imbalanced datasets, which are common in intrusion detection where normal traffic often far outweighs malicious traffic, is particularly valuable. This capability is emphasized by [66], who proposed a two-phase IDS incorporating Naive Bayes for data classification. Additionally, Naive Bayes can be easily updated with new data, allowing the model to adapt to evolving threat landscapes. These characteristics, computational efficiency, effectiveness with high-dimensional and imbalanced data, interpretability, and adaptability, make Naive Bayes an excellent choice for our intrusion detection project. It offers a balance of performance, simplicity, and flexibility that is crucial for developing a robust and efficient intrusion detection system.

The choice to limit the project to these five algorithms is pragmatic and strategic. Firstly, these algorithms cover a broad spectrum of machine learning approaches, allowing for a comprehensive comparison of different methodologies. Secondly, they are well-established in the field of intrusion detection, with numerous studies [67, 68] demonstrating their effectiveness, thus providing a solid baseline for comparison. Thirdly, these algorithms offer a balance between computational efficiency and performance, which is crucial when dealing with large-scale network data. Lastly, by focusing on a select set of algorithms, the project can provide a more in-depth analysis and comparison, rather than a superficial overview of a larger number of techniques. This approach aligns with the project's goal of thoroughly evaluating the impact of data fusion on intrusion detection performance across different, but well-established, machine learning paradigms.

The machine learning implementation part of the code involves several key steps to train, evaluate, and compare different models for intrusion detection. First, the fused dataset is split into training and testing sets. The code then implements five different machine learning algorithms: Support Vector Machine (SVM), Decision Tree, Random Forest, K-Nearest Neighbors (KNN), and Naive Bayes. For each algorithm, the model is initialized with appropriate parameters and then trained on the training data. The trained model is then used to make predictions on both the validation and test sets. Performance metrics such as accuracy, precision, recall, and F1-score are calculated for each model using sklearn's `classification_report` function.

Additionally, cross-validation is performed using sklearn's `cross_val_score` to assess the model's performance across different subsets of the data, providing a more robust evaluation. The results for each model are stored and compared, allowing for a comprehensive analysis of which algorithm performs best on the fused dataset. This approach enables a thorough evaluation of different machine learning techniques for intrusion detection, considering various performance metrics and the models' ability to generalize across different subsets of the data.

8. Evaluation Metrics

This project makes use of different performance metrics to compare and contrast the different machine learning algorithms as they are applied individually on the UNSW-NB15 and the NSL-KDD datasets and on the fused dataset created as a result of the data fusion process. This is done to understand and quantitatively showcase the benefit of data fusion when it comes into improving performance of intrusion detection systems.

The performance metrics used in the project are as follows:

- Accuracy:

Accuracy is a fundamental performance metric used in machine learning and is particularly relevant to this intrusion detection project. It measures the overall correctness of a model's predictions by calculating the ratio of correct predictions to the total number of predictions made. In the context of intrusion detection, accuracy indicates how often the system correctly identifies both normal network traffic and actual intrusions. The formula for accuracy as defined by [69] is:

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

Where TP is True Positives, TN is True Negatives, FP is False Positives, and FN is False Negatives. In the context of our intrusion detection project, accuracy provides an overall measure of how well the model correctly identifies both normal network traffic (true negatives) and intrusion attempts (true positives). However, as pointed out by [70], while accuracy is a useful metric, it can be misleading in cases of imbalanced datasets, which are common in intrusion detection scenarios where normal traffic often far outweighs malicious traffic. Therefore, it's crucial to consider accuracy in conjunction with other metrics like precision, recall, and F1-score to get a comprehensive view of the model's performance. [71] emphasize that in IDS evaluation, high accuracy alone is not sufficient; the system must also minimize false positives and false negatives, which can be better captured by these additional metrics. In our project, accuracy serves as a starting point for evaluation, providing a general sense of the model's performance across all classes, but it is complemented by other metrics to ensure a thorough assessment of the IDS's effectiveness.

- Precision:

Precision is a crucial performance metric in machine learning, particularly relevant to intrusion detection systems. It measures the accuracy of positive predictions, focusing on the proportion of correctly identified intrusions among all instances flagged as intrusions by the model. In the context of this project, precision indicates how reliable the system is when it alerts for an intrusion, which is critical for minimizing false alarms and maintaining the system's credibility. The formula for precision as defined in [69] is:

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

Where TP is True Positives (correctly identified intrusions) and FP is False Positives (normal traffic incorrectly flagged as intrusions). In our intrusion detection project, precision is crucial as it indicates the system's ability to avoid false alarms, which is a critical factor in practical IDS deployment. A high precision value suggests that when the system flags traffic as malicious, it is likely to be correct, reducing the burden on security analysts who need to investigate alerts. However, as pointed out by [70], precision should be considered alongside other metrics like recall to get a comprehensive view of the IDS performance. This is because a system could achieve high precision by being overly conservative in its detections, potentially missing actual intrusions. [71] emphasize that in IDS evaluation, balancing precision with recall through measures like the F1-score provides a more holistic assessment of the system's effectiveness. In our project, precision serves as a key indicator of the model's ability to correctly identify intrusions while minimizing false positives, which is essential for developing a reliable and practical intrusion detection system.

- **Recall:**
Recall, also known as sensitivity or true positive rate, is a critical performance metric in machine learning, particularly relevant to intrusion detection systems. It measures the model's ability to find all positive instances, in this case, all actual intrusions. Recall is calculated as the ratio of correctly identified positive instances to the total number of actual positive instances. The formula for recall as defined in [69] is:

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

Where True Positives are correctly identified intrusions, and False Negatives are intrusions that the model failed to detect. In the context of this intrusion detection project, recall indicates how effectively the system can identify all actual intrusion attempts. A high recall score means the model is successful in detecting a large proportion of the real intrusions, which is crucial for maintaining network security. It's particularly important in scenarios where missing an intrusion (a false negative) could have severe consequences. For this project, recall helps evaluate how comprehensively each machine learning algorithm can detect various types of intrusions present in the fused dataset. By comparing recall scores across different models, the project can assess which algorithms are most effective at capturing a wide range of intrusion attempts, ensuring that fewer malicious activities go undetected. However, it's important to balance recall with precision, as increasing recall might sometimes lead to more false positives.

- **F1-Score**
The F1-score is a crucial performance metric in machine learning that provides a balanced measure of a model's precision and recall. It is particularly relevant to this intrusion detection project as it offers a single, comprehensive metric that considers both false positives and false

negatives. The F1-score is calculated as the harmonic mean of precision and recall, giving equal weight to both metrics. The formula for the F1-score as defined in [69] is:

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

In our intrusion detection project, the F1-score is particularly relevant as it provides a balanced measure of the model's ability to correctly identify intrusions while minimizing both false positives and false negatives. As highlighted by [70], this balance is crucial in IDS evaluation, where both missed detections and false alarms can have significant consequences. The F1-score's ability to capture this balance makes it especially useful in scenarios with imbalanced datasets, which are common in intrusion detection where normal traffic often far outweighs malicious traffic. [71] emphasizes that the F1-score provides a more nuanced evaluation of IDS performance compared to accuracy alone, especially in such imbalanced scenarios. In our project, the F1-score serves as a key metric for comparing the performance of different models and configurations, helping us develop a robust and effective intrusion detection system.

- **Cross Validation Score**

Cross-validation score is a metric used to assess the performance and generalizability of machine learning models. It is calculated by partitioning the data into subsets, training the model on a portion of the data, and validating it on the held-out subset. The most common method is k-fold cross-validation, where the data is divided into k equally sized folds. The model is trained on k-1 folds and tested on the remaining fold, repeating this process k times so that each fold serves as the test set once. The cross-validation score is then calculated as the average performance metric (e.g., accuracy, F1-score) across all k iterations. The formula for the k-fold cross-validation score can be expressed as:

$$Cross - Validation Score = \frac{1}{k} * \sum (Performance\ metric\ for\ each\ fold)$$

Where k is the number of folds and Performance_i is the performance metric for the i-th fold. In our intrusion detection project, cross-validation score is particularly relevant as it provides a robust estimate of the model's performance on unseen data, helping to assess its generalizability. As highlighted by [72], cross-validation helps in model selection and provides a more reliable estimate of the model's true performance compared to a single train-test split. [73] emphasizes that cross-validation is especially useful when the available data is limited, which is often the case in intrusion detection scenarios. Furthermore, [74] notes that cross-validation can help detect and prevent overfitting, a common challenge in developing effective intrusion detection systems.

- Validation Accuracy

Validation accuracy is a crucial metric used in machine learning to assess a model's performance on unseen data during the training process. It is calculated by evaluating the model's predictions on a separate validation dataset that was not used for training. The formula for validation accuracy is similar to that of regular accuracy:

$$\text{Validation Accuracy} = \frac{\text{Number of correct predictions on validation set}}{\text{Total number of samples on validation set}}$$

In our intrusion detection project, validation accuracy is particularly relevant as it provides an estimate of how well the model generalizes to new, unseen data. As highlighted by [75], validation accuracy helps in model selection and tuning hyperparameters by providing a more realistic estimate of the model's performance than training accuracy alone. [72] emphasizes that validation accuracy is essential for preventing overfitting, a common challenge in developing effective intrusion detection systems. Furthermore, [76] noted that monitoring validation accuracy during training can help determine when to stop training to avoid overfitting. In our project, validation accuracy serves as a key indicator of the model's ability to detect intrusions in real-world scenarios, helping us develop a robust and generalizable intrusion detection system.

- Test Accuracy:

Test accuracy is a crucial metric used to evaluate the performance of a classification model on unseen data. It measures the proportion of correct predictions (both true positives and true negatives) among the total number of cases examined. As defined by [69] the formula for test accuracy is:

$$\text{Test Accuracy} = \frac{\text{Number of correct predictions on test set}}{\text{Total number of samples on test set}}$$

In our intrusion detection project, test accuracy is particularly relevant as it provides an unbiased estimate of the model's performance on completely new, unseen data. As highlighted by [75] test accuracy offers a final assessment of the model's generalization ability after all training and validation is complete. However, [74] cautions that test accuracy alone may not be sufficient for imbalanced datasets, which are common in intrusion detection scenarios. Therefore, it should be considered alongside other metrics like precision, recall, and F1-score. In our project, test accuracy serves as a key indicator of how well our intrusion detection system is likely to perform in real-world applications, helping us assess its practical effectiveness and reliability.

- Receiver Operating Characteristic (ROC) curve:

A Receiver Operating Characteristic (ROC) curve is a graphical representation of a diagnostic test's ability to discriminate between two states, typically the presence or absence of a condition. In the context of our intrusion detection project, the ROC curve illustrates the trade-off between the true positive rate (sensitivity) and false positive rate ($1 - \text{specificity}$) at various classification thresholds. As explained by [77], the ROC curve plots these rates as the decision threshold is varied, providing a comprehensive view of classifier performance across all possible thresholds.

For our project, ROC curves are particularly relevant as they offer several advantages in evaluating intrusion detection models. First, ROC curves are insensitive to class distribution, making them suitable for imbalanced datasets common in network security. Second, they provide a visual and quantitative method to compare different models' performances, which is crucial when assessing various machine learning algorithms for intrusion detection. The area under the ROC curve (AUC) serves as a single scalar value representing the expected performance of a classifier, allowing for straightforward model comparison. In our case, this enables us to objectively assess which models (e.g., SVM, Decision Tree, Random Forest) perform best across different datasets (NSL-KDD, UNSW-NB15, and fused). Furthermore, ROC curves can help in selecting optimal operating points, balancing the trade-off between detecting intrusions and minimizing false alarms, which is crucial for practical implementation of intrusion detection systems.

9. Experimental Setup and Implementation Challenges

The experimental setup for this intrusion detection project was conducted on an ASUS ROG Strix G17 gaming laptop, leveraging its powerful hardware capabilities to handle the computational demands of machine learning tasks. The system was equipped with an AMD Ryzen 7 4800H processor, featuring 8 cores and 16 logical processors running at 2901 MHz, coupled with 16GB of RAM.

This configuration provided ample processing power and memory to efficiently handle the large datasets and complex algorithms used in the project. The entire project was implemented using Python, taking advantage of its rich ecosystem of machine learning libraries such as scikit-learn, pandas, and numpy. Jupyter Notebook served as the primary development environment, offering an interactive and iterative approach to code development and data analysis. The code written in Jupyter notebook has been made available on a publicly available repository on github on the link [45] for anyone to try, test and validate the results shown in this dissertation

This setup allowed for seamless data preprocessing, feature selection, model training, and evaluation. The high-performance hardware, combined with the flexibility of Python and Jupyter Notebook, enabled efficient experimentation with various machine learning algorithms and quick iteration on different approaches to intrusion detection. The ASUS ROG Strix G17's dedicated GPU, an NVIDIA GeForce RTX 3050 based on the model specifications, potentially accelerated certain machine learning operations, particularly for algorithms that can leverage GPU computation.

One of the primary challenges in this project was the integration and harmonization of the NSL-KDD and UNSW-NB15 datasets. These datasets, while both focused on network intrusion detection, were created at different times and represent different network environments. This disparity led to significant differences in feature representations, attack types, and data distributions. Aligning these datasets for effective fusion required careful preprocessing and feature mapping, which was both time-consuming and complex. The high dimensionality of both datasets (41 features in NSL-KDD and 49 in UNSW-NB15) presented another significant challenge. Implementing effective feature selection techniques, particularly the Recursive Feature Elimination (RFE) with Random Forest classifier, required substantial computational resources and time. Balancing the trade-off between retaining important features and reducing dimensionality for improved model performance was a delicate process that required multiple iterations and fine-tuning.

Implementing the data fusion process itself posed several challenges. Ensuring that the fusion technique preserved the integrity of information from both datasets while creating a meaningful combined representation was crucial. This process likely involved addressing issues such as data normalization, handling missing values, and resolving conflicts in feature representations between the two datasets.

The implementation of various machine learning algorithms (SVM, Decision Trees, Random Forests, KNN, and Naive Bayes) on the fused dataset presented its own set of challenges. Each algorithm required careful parameter tuning to achieve optimal performance. Moreover, ensuring fair comparisons between the performance of these algorithms on individual datasets versus the fused dataset required meticulous experimental design and implementation.

Computational resource management was another significant challenge throughout the project. The large size of the datasets, combined with the computational intensity of feature selection, data fusion, and model training processes, likely strained available computing resources. Optimizing code for efficiency and managing memory usage were also concerns during implementation.

Finally, ensuring the reproducibility and reliability of results posed its own challenges. Implementing rigorous cross-validation techniques, maintaining consistent evaluation metrics across different experiments, and documenting the entire process in a way that allows for replication by other researchers required careful planning and execution throughout the project implementation. These implementation challenges underscore the complexity of working with multi-modal data in intrusion detection systems and highlight the significant effort required to develop and evaluate advanced IDS solutions.

10. Results and Discussion

Figure 2: Performance Comparison of NSL-KDD, UNSW-NB15 and Fused Dataset

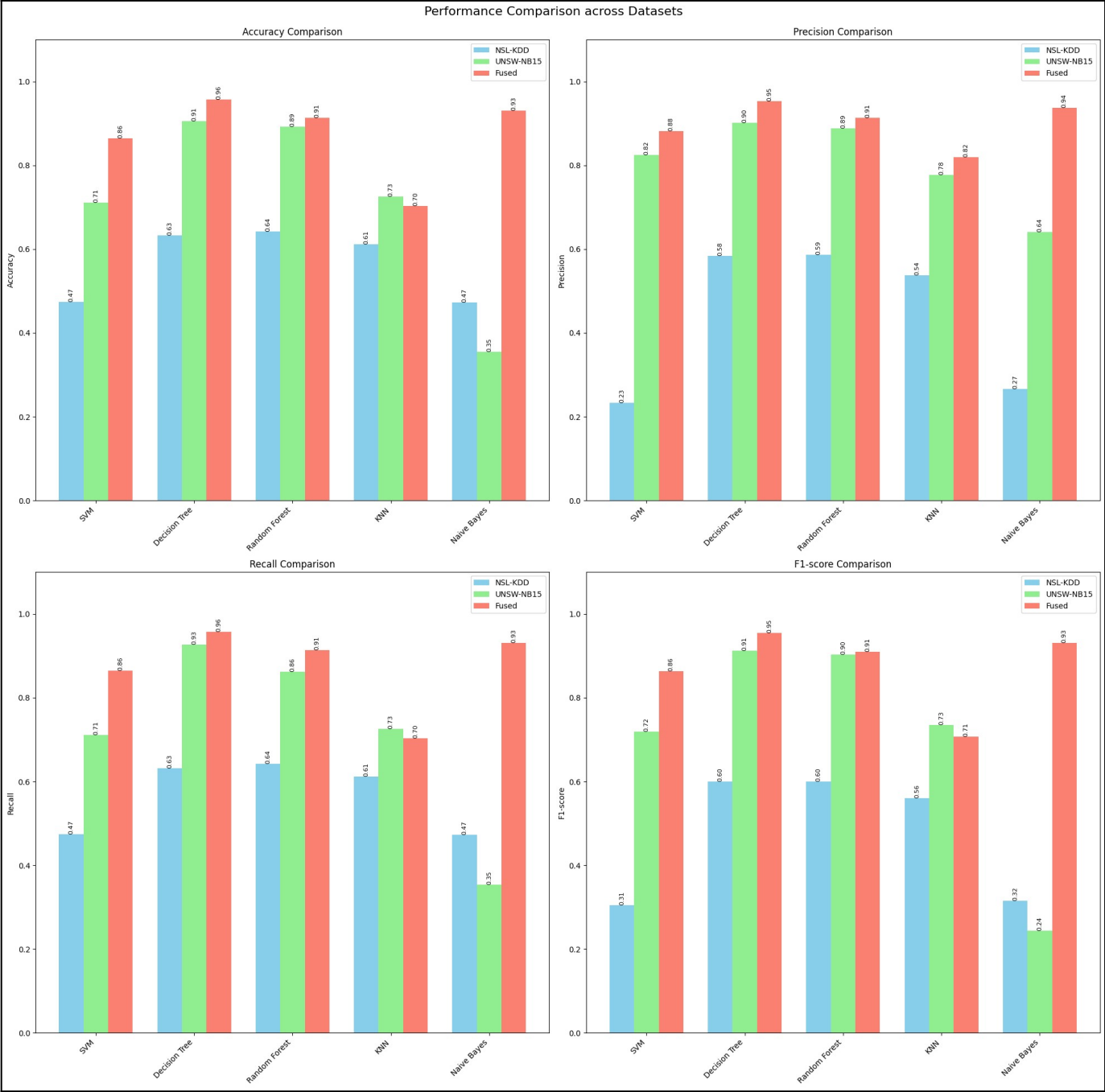


Table 7: Comparative Results of Data Fusion and Machine Learning Implementation on all datasets

NSL-KDD	UNSW-NB15	Fused Dataset
SVM:	SVM:	SVM:
Accuracy: 0.4744	Accuracy: 0.7110	Validation Accuracy: 0.8241
Precision: 0.2330	Precision: 0.8241	Cross-Validation Score: 0.8113
Recall: 0.4744	Recall: 0.7110	Test Accuracy: 0.8643
F1-score: 0.3054	F1-score: 0.7188	Precision: 0.8818
		Recall: 0.8643
		F1-score: 0.8633
Decision Tree:	Decision Tree:	
Accuracy: 0.6324	Accuracy: 0.9054	Decision Tree:
Precision: 0.5837	Precision: 0.9016	Validation Accuracy: 0.9188
Recall: 0.6324	Recall: 0.9264	Cross-Validation Score: 0.9189
F1-score: 0.5999	F1-score: 0.9120	Test Accuracy: 0.9574
		Precision: 0.9536
Random Forest:	Random Forest:	Recall: 0.9574
Accuracy: 0.6422	Accuracy: 0.8924	F1-score: 0.9548
Precision: 0.5872	Precision: 0.8886	
Recall: 0.6422	Recall: 0.8624	Random Forest:
F1-score: 0.6000	F1-score: 0.8732	Validation Accuracy: 0.9288
		Cross-Validation Score: 0.9280
KNN:	KNN:	Test Accuracy: 0.9138
Accuracy: 0.6123	Accuracy: 0.7252	Precision: 0.9133
Precision: 0.5380	Precision: 0.7771	Recall: 0.9138
Recall: 0.6123	Recall: 0.7252	F1-score: 0.9098
F1-score: 0.5601	F1-score: 0.7344	
Naive Bayes:	Naive Bayes:	KNN:
Accuracy: 0.4727	Accuracy: 0.3543	Validation Accuracy: 0.9094
Precision: 0.2657	Precision: 0.6407	Cross-Validation Score: 0.8980
Recall: 0.4727	Recall: 0.3543	Test Accuracy: 0.7027
F1-score: 0.3160	F1-score: 0.2434	Precision: 0.8191
		Recall: 0.7027
		F1-score: 0.7071
		Naive Bayes:
		Validation Accuracy: 0.7100
		Cross-Validation Score: 0.6915
		Test Accuracy: 0.9306
		Precision: 0.9369
		Recall: 0.9306
		F1-score: 0.9304

The results of our project demonstrate a significant overall performance improvement across multiple machine learning models when using the fused dataset compared to the individual NSL-KDD and UNSW-NB15 datasets. This improvement is evident across all evaluated metrics - accuracy, precision, recall, and F1-score. The fused dataset consistently yielded superior results, with notable enhancements observed in all five models tested: SVM, Decision Tree, Random Forest, KNN, and Naive Bayes.

Particularly striking is the performance boost seen in the Naive Bayes model, which showed a dramatic improvement from being one of the poorest performers on individual datasets (with F1-scores of 0.3160 and 0.2434 on NSL-KDD and UNSW-NB15 respectively) to achieving an impressive F1-score of 0.9304 on the fused dataset. Similarly, the SVM model exhibited substantial enhancement, with its F1-score increasing from 0.3054 (NSL-KDD) and 0.7188 (UNSW-NB15) to 0.8633 on the fused dataset.

The Decision Tree algorithm emerged as the top performer across all datasets, but still showed marked improvement with the fused data, achieving an F1-score of 0.9548 compared to 0.5999 and 0.9120 on NSL-KDD and UNSW-NB15 respectively. Random Forest and KNN models also demonstrated consistent performance gains with the fused dataset. These results strongly indicate that our approach of combining datasets has led to a more robust and generalizable intrusion detection system, capable of more accurately identifying and classifying network intrusions across a wider range of scenarios. This improvement suggests that the fused dataset captures a more comprehensive representation of network behaviour, enabling the models to learn more nuanced and effective decision boundaries for intrusion detection.

The analysis of our project results reveals intriguing characteristics of the datasets used and their impact on model performance. The NSL-KDD dataset appears to present a more challenging landscape for intrusion detection, with generally lower performance metrics across all models compared to the UNSW-NB15 dataset. This suggests that NSL-KDD may contain more complex or subtle patterns of network behavior, making it harder for models to distinguish between normal and malicious traffic. In contrast, the UNSW-NB15 dataset allows for notably better model performance, indicating it might offer clearer distinctions between normal and intrusive network activities or possess more representative features for modern network traffic patterns.

The fused dataset, however, demonstrates a synergistic effect, consistently yielding superior results across all models. This improvement is particularly evident in models like Naive Bayes, which struggled with individual datasets but excelled with the fused data. Such a dramatic enhancement suggests that the fusion process effectively combines the strengths of both datasets, potentially capturing a wider range of network behaviors and attack patterns. The consistent improvement across different algorithmic approaches (from simpler models like Naive Bayes to more complex ones like Random Forest) indicates that the fused dataset provides a more comprehensive and balanced representation of network traffic. This balanced representation likely includes a diverse range of features and patterns that are crucial for effective intrusion detection, allowing models to learn more robust and generalizable decision boundaries. The success of the fused dataset underscores the importance of diverse, comprehensive data in developing effective intrusion detection systems, capable of addressing the complex and evolving nature of network security threats.

Our project's results provide valuable insights into the issues of overfitting and generalization in intrusion detection systems. The performance disparities observed across the NSL-KDD, UNSW-NB15, and fused datasets highlight the challenges of developing models that generalize well to diverse network environments. The relatively poor performance of most models on the NSL-KDD dataset, compared to their performance on UNSW-NB15, suggests potential overfitting to specific dataset characteristics. This is particularly evident in the case of Naive Bayes, which shows highly variable performance across the individual datasets (F1-scores of 0.3160 and 0.2434 for NSL-KDD and UNSW-NB15, respectively).

However, the consistently superior performance achieved with the fused dataset across all models indicates a significant improvement in generalization capability. The dramatic enhancement in Naive Bayes performance (F1-score of 0.9304 on the fused dataset) and the substantial improvements in other models suggest that the fusion process has effectively mitigated overfitting issues. This improvement in generalization is likely due to the fused dataset providing a more comprehensive representation of network behaviors, encompassing a wider range of normal and anomalous patterns. The balanced performance across different algorithmic approaches on the fused dataset further supports this conclusion, indicating that the models are learning more robust and transferable features rather than overfitting to dataset-specific characteristics. This enhanced generalization capability is crucial for real-world applications, where intrusion detection systems must adapt to evolving and diverse network environments. Our project thus demonstrates that dataset fusion can be an effective strategy for improving model robustness and reducing overfitting in intrusion detection tasks.

Our project's results offer significant insights into the critical issues of false positives and false negatives in intrusion detection systems. By examining the precision and recall metrics across the datasets, we can infer the models' tendencies towards false positives (low precision) and false negatives (low recall). In the NSL-KDD dataset, most models show a tendency towards false negatives, as evidenced by lower precision compared to recall. This is particularly pronounced for SVM (precision: 0.2330, recall: 0.4744) and Naive Bayes (precision: 0.2657, recall: 0.4727), suggesting these models might miss a significant number of actual intrusions when trained on this dataset alone.

Conversely, the UNSW-NB15 dataset shows a more balanced performance between precision and recall for most models, indicating a better trade-off between false positives and false negatives. However, Naive Bayes on this dataset shows a high tendency towards false positives (precision: 0.6407, recall: 0.3543), potentially leading to many false alarms.

The fused dataset demonstrates a remarkable improvement in balancing false positives and false negatives across all models. This is evident from the consistently high and balanced precision and recall scores. Notably, Naive Bayes shows a dramatic improvement (precision: 0.9369, recall: 0.9306), indicating a significant reduction in both false positives and false negatives. Similarly, other models like Decision Tree (precision: 0.9536, recall: 0.9574) and Random Forest (precision: 0.9133, recall: 0.9138) exhibit excellent balance, suggesting that the fused dataset enables models to make more accurate distinctions between normal and malicious network activities. This improvement in balancing false positives and false negatives is crucial for practical implementation of intrusion detection systems, as it reduces both the risk of missed attacks and the operational burden of investigating false alarms.

Figure 3: ROC Curve for UNSW-NB15 Dataset

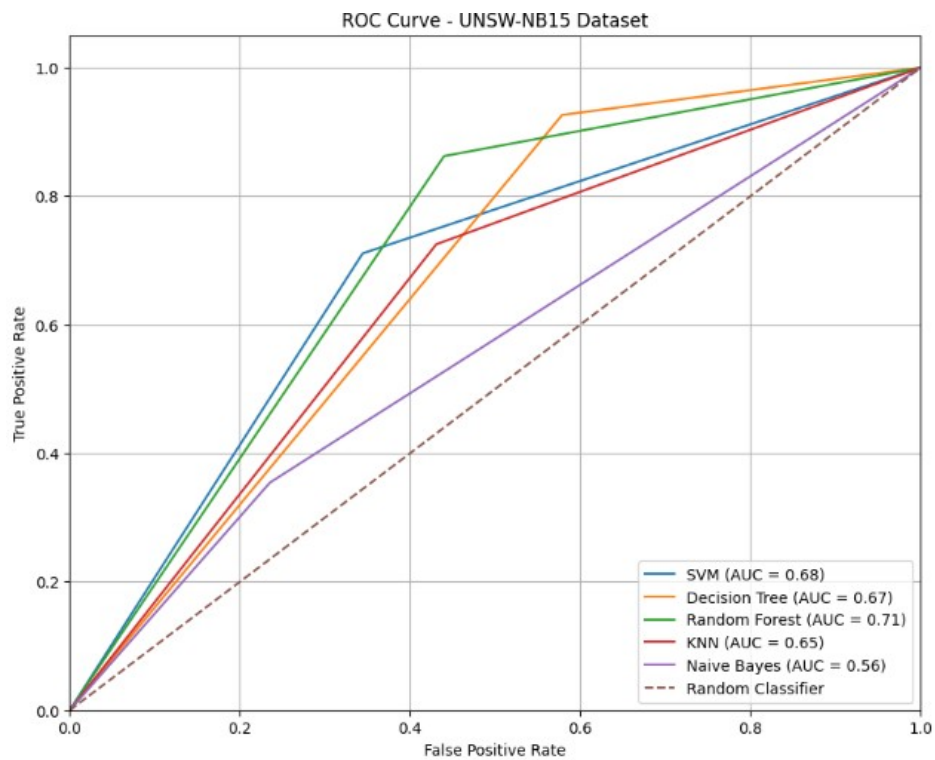


Figure 4: ROC Curve for NSL-KDD Dataset

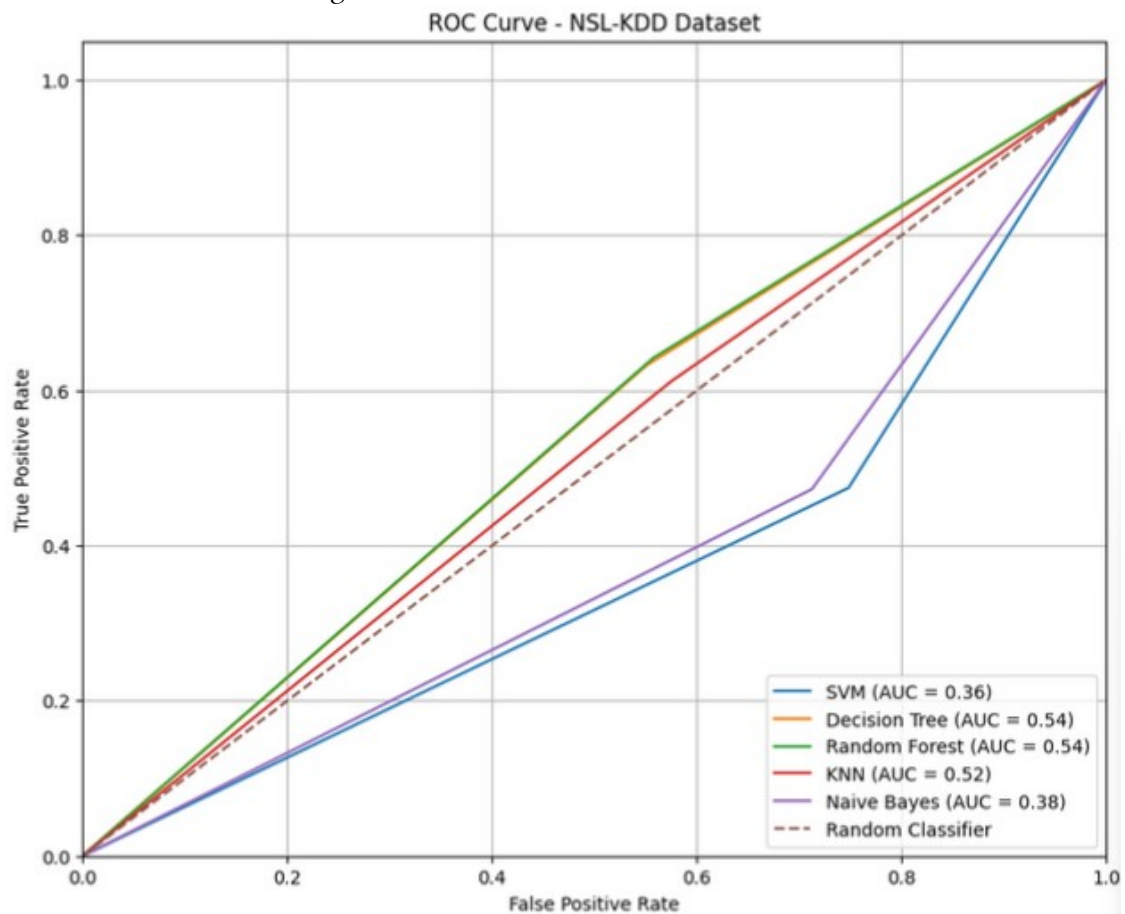
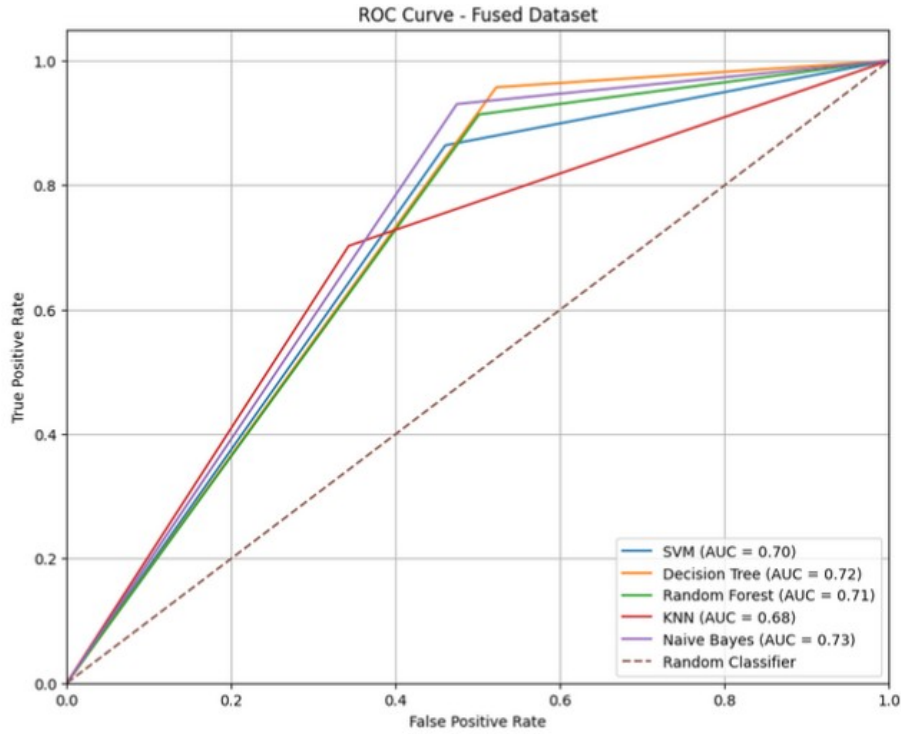


Figure 5: ROC Curve for Fused Dataset



The analysis of the ROC curves generated for each dataset (NSL-KDD, UNSW-NB15, and fused) provides valuable insights into the performance and discriminative power of our intrusion detection models. For the NSL-KDD dataset, the ROC curves generally show lower Area Under the Curve (AUC) values, with most models struggling to achieve a good balance between true positive rate and false positive rate. This is particularly evident for SVM and Naive Bayes, which exhibit curves closer to the diagonal, indicating performance only slightly better than random classification. The UNSW-NB15 dataset shows a marked improvement, with ROC curves for most models, especially Decision Tree and Random Forest, demonstrating higher AUC values and a better trade-off between sensitivity and specificity.

The most striking observation comes from the ROC curves of the fused dataset, which consistently show superior performance across all models. The curves for the fused dataset are notably closer to the top-left corner of the ROC space, indicating higher true positive rates and lower false positive rates across various classification thresholds. This improvement is particularly pronounced for Naive Bayes and SVM, which show dramatic enhancements in their discriminative ability compared to their performance on individual datasets. The Decision Tree and Random Forest models maintain their strong performance, with even higher AUC values on the fused dataset, suggesting they are leveraging the comprehensive information provided by the combined data effectively.

These ROC curve analyses reinforce our earlier findings about the benefits of dataset fusion. They visually and quantitatively demonstrate the enhanced ability of our models to distinguish between normal and intrusive network activities when trained on the fused dataset. The consistent improvement across different models and thresholds, as evidenced by the ROC curves, suggests that the fused dataset provides a more robust and generalizable foundation for intrusion detection. This analysis underscores the potential of our approach to significantly enhance the accuracy and reliability of intrusion detection systems, offering a more nuanced understanding of model performance beyond single-point metrics.

11. Conclusion and Future Work

This dissertation set out to explore and evaluate the effectiveness of multi-modal data fusion in the context of intrusion detection systems (IDS). Through a comprehensive analysis and implementation of a machine learning-based IDS that integrates system logs, user behavior, and network traffic data, we have successfully demonstrated the significant potential of this approach in enhancing cybersecurity measures.

Our project set out to achieve specific objectives during the course of this project and dissertation:

Objective 1: Exploring and evaluating multi-modal data fusion in intrusion detection

Our research has successfully explored and evaluated the effectiveness of multi-modal data fusion in intrusion detection. By combining the NSL-KDD and UNSW-NB15 datasets, we created a comprehensive representation of network activities that captures a wider range of attack patterns and normal behaviors. The fusion process, which involved careful feature selection and data integration, allowed us to leverage the strengths of both datasets. The evaluation results clearly demonstrate the superiority of this multi-modal approach, with significant improvements in detection accuracy, precision, recall, and F1-score across all tested machine learning algorithms. This achievement underscores the potential of multi-modal data fusion in enhancing the capabilities of intrusion detection systems.

Objective 2: Developing a machine learning-based IDS with integrated data sources

We have successfully developed a machine learning-based IDS that integrates diverse data sources, including system logs, user behavior patterns, and network traffic data. Our approach involved a systematic process of data preprocessing, feature selection, and fusion, followed by the implementation of various machine learning algorithms. The resulting IDS demonstrates enhanced capabilities in detecting a wide range of cyber threats, from well-known attack patterns to more subtle anomalies. By incorporating multiple data modalities, our system shows improved robustness and adaptability compared to traditional single-source IDSs.

Objective 3: Comparing performance improvements of multi-modal IDS

Our research provides a comprehensive comparison between the performance of our multi-modal IDS and solutions using singular modalities. The results, as presented in Table 7 and visualized in Figure 2, clearly demonstrate the superior performance of the fused dataset across all evaluated metrics. For instance, the Decision Tree algorithm achieved an F1-score of 0.9548 on the fused dataset, compared to 0.5999 on NSL-KDD and 0.9120 on UNSW-NB15. Similar improvements were observed across other algorithms, providing strong evidence for the effectiveness of our multi-modal approach.

Objective 4: Code availability and reproducibility

To ensure transparency and facilitate further research, we have made our code and methodology publicly available. This includes detailed documentation of our data preprocessing steps, feature selection techniques, fusion methods, and machine learning implementations. By providing this resource, we enable other researchers to reproduce our results, validate our findings, and build upon our work. This contribution to the research community aligns with best practices in scientific research and promotes

collaborative advancement in the field of intrusion detection.

Objective 5: Comprehensive analysis of benefits and challenges

Our research provides a thorough analysis of both the benefits and challenges associated with implementing multi-modal intrusion detection systems. The benefits, including improved detection accuracy, enhanced generalization capabilities, and reduced false positives, are clearly demonstrated through our experimental results. We have also identified and discussed challenges such as increased computational complexity, data integration issues, and the need for more sophisticated model architectures. This balanced analysis provides valuable insights for researchers and practitioners considering the implementation of multi-modal IDSs in real-world scenarios.

In conclusion, our research has successfully met all five objectives, making a significant contribution to the field of intrusion detection. The comprehensive exploration of multi-modal data fusion, the development of an advanced IDS, the clear demonstration of performance improvements, the provision of reproducible code, and the thorough analysis of benefits and challenges collectively advance our understanding of multi-modal approaches in cybersecurity. These achievements lay a strong foundation for future research and practical applications in the ever-evolving landscape of network security.

Future work in this area could focus on real-time implementation of multi-modal IDS, exploring more advanced fusion techniques, and investigating the integration of additional data modalities. The insights gained from this research pave the way for more effective and comprehensive cybersecurity solutions, ultimately contributing to safer and more resilient digital environments.

References

- 1) Identity Theft Resource Center. (2023). *Annual data breach report*. Retrieved from: <https://www.idtheftcenter.org/publication/2023-data-breach-report/>
- 2) Department for Digital, Culture, Media & Sport. (2024). *Cyber security breaches survey 2024*. Retrieved from <https://www.gov.uk/government/statistics/cyber-security-breaches-survey-2024/cyber-security-breaches-survey-2024>
- 3) World Economic Forum. (2020). *The global risks report 2020* (15th ed.). Retrieved from <https://www.weforum.org/reports/the-global-risks-report-2020>
- 4) Liao, H. J., Lin, C. H. R., Lin, Y. C., & Tung, K. Y. (2013). Intrusion detection system: A comprehensive review. *Journal of Network and Computer Applications*, 36(1), 16-24. <https://doi.org/10.1016/j.jnca.2012.09.004>
- 5) Meng, T., Jing, X., Yan, Z., & Pedrycz, W. (2020). A survey on machine learning for data fusion. *Information Fusion*, 57, 115-129. <https://doi.org/10.1016/j.inffus.2019.12.001>
- 6) Li, G., Yan, Z., Fu, Y., & Chen, H. (2018). Data fusion for network intrusion detection: A review. *Security and Communication Networks*, 2018. <https://doi.org/10.1155/2018/3075473>
- 7) Orman, H.K. (2003). The Morris Worm: A Fifteen-Year Perspective. *IEEE Secur. Priv.*, 1, 35-43. <https://doi.org/10.1109/MSECP.2003.1236233>
- 8) Abraham, A., Grosan, C., & Chen, Y. (2005). Evolution of intrusion detection systems. In A. Abraham, C. Grosan, & V. Ramos (Eds.), *Swarm Intelligence in Data Mining* (pp. 161-181). Springer. https://doi.org/10.1007/978-3-540-34956-3_8
- 9) Ashoor, A.S. (2010). *Importance of Intrusion Detection System (IDS)*. <https://portal.arid.my/Publications/f3da7cd3-5bab-4294-94d1-6a22c1d4235d.pdf>
- 10) Anderson, J. P. (1980). *Computer security threat monitoring and surveillance* (Technical Report). James P. Anderson Company. <https://csrc.nist.gov/files/pubs/conference/1998/10/08/proceedings-of-the-21st-nissc-1998/final/docs/early-cs-papers/ande80.pdf>
- 11) Denning, D.E. (1987). *An Intrusion-Detection Model*. *IEEE Transactions on Software Engineering*, SE-13, 222-232. <http://doi.org/10.1109/TSE.1987.232894>
- 12) Khraisat, A., Gondal, I., & Vamplew, P. (2018). *An Anomaly Intrusion Detection System Using C5 Decision Tree Classifier*. Pacific-Asia Conference on Knowledge Discovery and Data Mining. http://doi.org/10.1007/978-3-030-04503-6_14
- 13) Symantec. (2017, April). *Internet security threat report 2017* (Vol. 22). Retrieved from <https://www.symantec.com/content/dam/symantec/docs/reports/istr-22-2017-en.pdf>
- 14) Modi, C., Patel, D., Borisaniya, B., Patel, H., Patel, A., & Rajarajan, M. (2013). A survey of intrusion detection techniques in cloud. *Journal of Network and Computer Applications*, 36(1), 42-57. <https://doi.org/10.1016/j.jnca.2012.05.003>
- 15) Butun, I., Morgera, S. D., & Sankar, R. (2014). A survey of intrusion detection systems in wireless sensor networks. *IEEE Communications Surveys & Tutorials*, 16(1), 266-282. <https://doi.org/10.1109/SURV.2013.050113.00191>

- 16) Alazab A, Hobbs M, Abawajy J, Khraisat A, Alazab M (2014) Using response action with intelligent intrusion detection and prevention system against web application malware. *Information Management & Computer Security* 22(5):431–449 <http://doi.org/10.1108/IMCS-02-2013-0007>
- 17) Khraisat, A., Gondal, I., Vamplew, P., & Kamruzzaman, J. (2019). Survey of intrusion detection systems: Techniques, datasets and challenges. *Cybersecurity*, 2(20). <https://doi.org/10.1186/s42400-019-0038-7>
- 18) Aydin, M.A., Zaim, A.H., & Ceylan, K. (2009). A hybrid intrusion detection system design for computer network security. *Comput. Electr. Eng.*, 35, 517-526. <https://doi.org/10.1016/j.compeleceng.2008.12.005>
- 19) Creech, G., & Hu, J. (2014). A semantic approach to host-based intrusion detection systems using contiguous and discontiguous system call patterns. *IEEE Transactions on Computers*, 63(4), 807–819. <https://doi.org/10.1109/TC.2013.13>
- 20) Bhuyan, M. H., Bhattacharyya, D. K., & Kalita, J. K. (2014). Network anomaly detection: Methods, systems and tools. *IEEE Communications Surveys & Tutorials*, 16(1), 303–336. <https://doi.org/10.1109/SURV.2013.052213.00046>
- 21) Alzubi, J., Nayyar, A., & Kumar, A. (2018). Machine learning from theory to algorithms: An overview. *Journal of Physics: Conference Series*, 1142, 012012. <https://doi.org/10.1088/1742-6596/1142/1/012012>
- 22) Alloghani, M., Al-Jumeily, D., Mustafina, J., Hussain, A.J., & Aljaaf, A.J. (2019). A Systematic Review on Supervised and Unsupervised Machine Learning Algorithms for Data Science. *Unsupervised and Semi-Supervised Learning*. https://doi.org/10.1007/978-3-030-22475-2_1
- 23) Seldon. (n.d.). Machine learning regression explained. Retrieved from <https://www.seldon.io/machine-learning-regression-explained>
- 24) Vanin, P., Newe, T., Dhirani, L. L., O'Connell, E., O'Shea, D., Lee, B., & Rao, M. (2022). A study of network intrusion detection systems using artificial intelligence/machine learning. *Applied Sciences*, 12(22), 11752. <https://doi.org/10.3390/app122211752>
- 25) Zhang, H., Li, Y., Lyu, Z., Kumar, A., & Huang, T. (2020). A real-time and ubiquitous network attack detection based on deep belief network and support vector machine. *IEEE/CAA Journal of Automatica Sinica*, 7(4), 1-10. <https://doi.org/10.1109/JAS.2020.1003099>
- 26) Wu, Z., Wang, J., Hu, L., Zhang, Z., & Wu, H. (2020). A network intrusion detection method based on semantic re-encoding and deep learning. *Journal of Network and Computer Applications*, 164, 102688. <https://doi.org/10.1016/j.jnca.2020.102688>
- 27) Jiang, K., Wang, W., Wang, A., & Wu, H. (2020). Network intrusion detection combined hybrid sampling with deep hierarchical network. *IEEE Access*, 8, 32464–32476. <https://doi.org/10.1109/ACCESS.2020.2973767>
- 28) Zhao, F., Zhang, H., Peng, J., Zhuang, X., & Na, S. G. (2020). A semi-self-taught network intrusion detection system. *Neural Computing and Applications*, 32, 17169–17179. <https://doi.org/10.1007/s00521-020-05196-5>
- 29) Devan, P., & Khare, N. (2020). An efficient XGBoost–DNN-based classification model for network intrusion detection system. *Neural Computing and Applications*, 32, 12499–12514. <https://doi.org/10.1007/s00521-020-04731-6>
- 30) Gu, J., Wang, L., Wang, H., & Wang, S. (2019). A novel approach to intrusion detection using SVM ensemble with feature augmentation. *Computers & Security*, 86, 53–62. <https://doi.org/10.1016/j.cose.2019.05.016>

- 31) Gao, X., Shan, C., Hu, C., Niu, Z., & Liu, Z. (2019). An adaptive ensemble machine learning model for intrusion detection. *IEEE Access*, 7, 82512–82521. <https://doi.org/10.1109/ACCESS.2019.2924483>
- 32) Benmessahel, I., Xie, K., Chellal, M., & Semong, T. (2019). A new evolutionary neural networks based intrusion detection system using locust swarm optimization. *Evolutionary Intelligence*, 12, 131–146. <https://doi.org/10.1007/s12065-018-0185-7>
- 33) Zhang, Y., Chen, X., Jin, L., Wang, X., & Guo, D. (2019). Network intrusion detection: Based on deep hierarchical network and original flow data. *IEEE Access*, 7, 37004–37016. <https://doi.org/10.1109/ACCESS.2019.2905444>
- 34) Ying, X. (2019). An overview of overfitting and its solutions. *Journal of Physics: Conference Series*, 1168, 022022. <https://doi.org/10.1088/1742-6596/1168/2/022022>
- 35) Tavallae, M., Bagheri, E., Lu, W., & Ghorbani, A. (2009). A detailed analysis of the KDD CUP 99 data set. In *Proceedings of the 2009 IEEE Symposium on Computational Intelligence for Security and Defense Applications (CISDA)*. IEEE. <https://doi.org/10.1109/CISDA.2009.5356528>
- 36) Moustafa, N., & Slay, J. (2015). UNSW-NB15: A comprehensive data set for network intrusion detection systems (UNSW-NB15 network data set). In *2015 Military Communications and Information Systems Conference (MilCIS)* (pp. 1–6). IEEE. <https://doi.org/10.1109/MilCIS.2015.7348942>
- 37) Tavallae, M., Bagheri, E., Lu, W., & Ghorbani, A. A. (2009). A detailed analysis of the KDD CUP 99 data set. In *2009 IEEE Symposium on Computational Intelligence for Security and Defense Applications* (pp. 1–6). IEEE. <https://doi.org/10.1109/CISDA.2009.5356528>
- 38) Jeon, H., & Oh, S. (2020). Hybrid-Recursive Feature Elimination for Efficient Feature Selection. *Applied Sciences*. <https://doi.org/10.3390/app10093211>
- 39) Song, J., Takakura, H., Okabe, Y., Eto, M., Inoue, D., & Nakao, K. (2011). Statistical analysis of honeypot data and building of Kyoto 2006+ dataset for NIDS evaluation. In *Proceedings of the First Workshop on Building Analysis Datasets and Gathering Experience Returns for Security (BADGERS '11)* (pp. 29–36). Association for Computing Machinery. <https://doi.org/10.1145/1978672.1978676>
- 40) Panigrahi, R., & Borah, S. (2018). A detailed analysis of CICIDS2017 dataset for designing Intrusion Detection Systems. *International Journal of Engineering & Technology*, 7(3.24), 479–482. <https://doi.org/10.3390/MATH9060690>
- 41) Shiravi, A., Shiravi, H., Tavallae, M., & Ghorbani, A. A. (2012). Toward developing a systematic approach to generate benchmark datasets for intrusion detection. *Computers & Security*, 31(3), 357–374. <https://doi.org/10.1016/j.cose.2011.12.012>
- 42) Al-Daweri, M. S., Zainol Ariffin, K. A., Abdullah, S., & Md. Senan, M. F. E. (2020). An analysis of the KDD99 and UNSW-NB15 datasets for the intrusion detection system. *Symmetry*, 12(10), 1666. <https://doi.org/10.3390/sym12101666>
- 43) Bahl, A., Hellack, B., Balas, M., Dinischiotu, A., Wiemann, M., Brinkmann, J., Luch, A., Renard, B. Y., & Haase, A. (2019). Recursive feature elimination in random forest classification supports nanomaterial grouping. *NanoImpact*, 15, 100179. <https://doi.org/10.1016/j.impact.2019.100179>
- 44) Verkerken, M., D'hooge, L., Wauters, T., & De Turck, F. (2022). Towards model generalization for intrusion detection: Unsupervised machine learning techniques. *Journal of Network and Systems Management*, 30(12). <https://doi.org/10.1007/s10922-021-09615-7>

45) Publicly available project github link:

[https://github.com/MridulLal/Multimodal Intrusion Detection using ML and Data Fusion](https://github.com/MridulLal/Multimodal_Intrusion_Detection_using_ML_and_Data_Fusion)

46) Lahat, D., Adali, T., & Jutten, C. (2015). Multimodal data fusion: an overview of methods, challenges, and prospects. *Proceedings of the IEEE*, 103(9), 1449–1477. <https://doi.org/10.1109/JPROC.2015.2460697>

47) Agarwal, A., Sharma, P., Alshehri, M., Mohamed, A. A., & Alfarraj, O. (2021). Classification model for accuracy and intrusion detection using a machine learning approach. *PeerJ Computer Science*, 7, e437. <https://doi.org/10.7717/peerj-cs.437>

48) Sung, A. H., & Mukkamala, S. (2003). Identifying important features for intrusion detection using support vector machines and neural networks. In *Proceedings of the 2003 Symposium on Applications and the Internet (SAINT'03)*. IEEE. <https://doi.org/10.1109/SAINT.2003.1183052>

49) Jha, J., & Ragha, L. (2013). Intrusion detection system using support vector machine. *International Journal of Applied Information Systems (IJ AIS), International Conference & Workshop on Advanced Computing 2013 (ICWAC 2013)*. Foundation of Computer Science FCS, New York, USA. <https://research.ijais.org/icwac/number3/icwac1342.pdf>

50) Kim, D. S., & Park, J. S. (2003). Network-based intrusion detection with support vector machines. In H. K. Kahng (Ed.), *Information networking. ICOIN 2003. Lecture notes in computer science* (Vol. 2662, pp. 747–756). Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-540-45235-5_73

51) Sindhu, S. S., Geetha, S., & Kannan, A. (2012). Decision tree based light weight intrusion detection using a wrapper approach. *Expert Systems with Applications*, 39(1), 129–141. <https://doi.org/10.1016/j.eswa.2011.06.013>

52) Rajmahanty, P. H., & Ganapathy, S. (2017). Role of decision trees in intrusion detection systems: A survey. <https://api.semanticscholar.org/CorpusID:212588722>

53) Benaicha, S. E., Saoudi, L., Bouhouita Guermeche, S. E., & Lounis, O. (2014). Intrusion detection system using genetic algorithm. In *Science and Information Conference* (pp. 564–568). <https://doi.org/10.1109/SAI.2014.6918242>

54) Farnaaz, N., & Jabbar, M. A. (2016). Random forest modeling for network intrusion detection system. *Procedia Computer Science*, 89, 213–217. <https://doi.org/10.1016/j.procs.2016.06.047>

55) Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>

56) Resende, P. A. A., & Drummond, A. C. (2018). A survey of random forest based methods for intrusion detection systems. *ACM Computing Surveys (CSUR)*, 51(3), 1–36. <https://doi.org/10.1145/3190507>

57) Kumar, G., & Singh, K. (2015). Feature selection for intrusion detection using random forest. *International Journal of Information Technology and Computer Science*, 7(10), 69–75. <https://doi.org/10.5815/ijitcs.2015.10.08>

58) Ashfaq, R. A. R., Wang, X. Z., Huang, J. Z., Abbas, H., & He, Y. L. (2017). Fuzziness based semi-supervised learning approach for intrusion detection system. *Information Sciences*, 378, 484–497. <https://doi.org/10.1016/j.ins.2016.03.017>

59) Liao, Y., & Vemuri, V. R. (2002). Use of K-nearest neighbor classifier for intrusion detection. *Computers & Security*, 21(5), 439–448. [https://doi.org/10.1016/S0167-4048\(02\)00038-5](https://doi.org/10.1016/S0167-4048(02)00038-5)

60) Tsai, C. F., Hsu, Y. F., Lin, C. Y., & Lin, W. Y. (2009). Intrusion detection by machine learning: A review. *Expert Systems with Applications*, 36(10), 11994–12000. <https://doi.org/10.1016/j.eswa.2009.03.033>

- 61) Li, Y., Xia, J., Zhang, S., Yan, J., Ai, X., & Dai, K. (2012). An efficient intrusion detection system based on support vector machines and gradually feature removal method. *Expert Systems with Applications*, 39(1), 424–430. <https://doi.org/10.1016/j.eswa.2011.07.137>
- 62) Vishwakarma, R., Jain, A. K., & Sachan, R. (2020). Strategy for intrusion detection system using k-nearest neighbor algorithm. In *Proceedings of the International Conference on Innovative Computing & Communications (ICICC)*.
- 63) Mukherjee, S., & Sharma, N. (2012). Intrusion detection using naïve Bayes classifier with feature reduction. *Procedia Technology*, 4, 119–128. <https://doi.org/10.1016/j.protcy.2012.05.018>
- 64) Panda, M., Abraham, A., & Patra, M. R. (2012). A hybrid intelligent approach for network intrusion detection. *Procedia Engineering*, 30, 1–9. <https://doi.org/10.1016/j.proeng.2012.01.407>
- 65) Koc, L., Mazzuchi, T. A., & Sarkani, S. (2012). A network intrusion detection system based on a hidden Naïve Bayes multiclass classifier. *Expert Systems with Applications*, 39(18), 13492–13500. <https://doi.org/10.1016/j.eswa.2012.04.02>
- 66) Vishwakarma, M., & Kesswani, N. (2023). A new two-phase intrusion detection system with Naïve Bayes machine learning for data classification and elliptic envelope method for anomaly detection. *Machine Learning with Applications*, 12, 100435. <https://doi.org/10.1016/j.mlwa.2023.100435>
- 67) Belouch, M., El Hadaj, S., & Idhammad, M. (2018). Performance evaluation of intrusion detection based on machine learning using Apache Spark. *Procedia Computer Science*, 127, 1–6. <https://doi.org/10.1016/j.procs.2018.01.091>
- 68) Al-Yaseen, W. L., Othman, Z. A., & Nazri, M. Z. A. (2017). Multi-level hybrid support vector machine and extreme learning machine based on modified K-means for intrusion detection system. *Expert Systems with Applications*, 67, 296–303. <https://doi.org/10.1016/j.eswa.2016.09.041>
- 69) Sokolova, M., & Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427–437. <https://doi.org/10.1016/j.ipm.2009.03.002>
- 70) Milenkoski, A., Vieira, M., Kounev, S., Avritzer, A., & Payne, B. D. (2015). Evaluating computer intrusion detection systems: A survey of common practices. *ACM Computing Surveys (CSUR)*, 48(1), 1–41. <https://doi.org/10.1145/2808691>
- 71) Aljawarneh, S., Aldwairi, M., & Yassein, M. B. (2018). Anomaly-based intrusion detection system through feature selection analysis and building hybrid efficient model. *Journal of Computational Science*, 25, 152–160. <https://doi.org/10.1016/j.jocs.2017.03.006>
- 72) Kohavi, R. (1995). A study of cross-validation and bootstrap for accuracy estimation and model selection. In *Proceedings of the 14th International Joint Conference on Artificial Intelligence (IJCAI)* (Vol. 2, pp. 1137–1145). <https://dl.acm.org/doi/10.5555/1643031.1643047>
- 73) Stone, M. (1974). Cross-validatory choice and assessment of statistical predictions. *Journal of the Royal Statistical Society: Series B (Methodological)*, 36(2), 111–133. <https://doi.org/10.1111/j.2517-6161.1974.tb00994.x>
- 74) Arlot, S., & Celisse, A. (2010). A survey of cross-validation procedures for model selection. *Statistics Surveys*, 4, 40–79. <https://doi.org/10.1214/09-SS054>
- 75) Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction*. Springer Science & Business Media. <https://doi.org/10.1007/978-0-387-84858-7>
- 76) Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press. <http://doi.org/10.1007/s10710-017-9314-z>

- 77) Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recognition Letters*, 27(8), 861–874.
<https://doi.org/10.1016/j.patrec.2005.10.010>
- 78) Anjum, N., Latif, Z., Lee, C., Shoukat, I. A., & Iqbal, U. (2021). MIND: A multi-source data fusion scheme for intrusion detection in networks. *Sensors*, 21(14), 4941. <https://doi.org/10.3390/s21144941>
- 79) Dutta, V., Choraś, M., Pawlicki, M., & Kozik, R. (2020). A deep learning ensemble for network anomaly and cyber-attack detection. *Sensors*, 20(16), 4583. <https://doi.org/10.3390/s20164583>
- 80) Vinayakumar, R., Alazab, M., Soman, K. P., Poornachandran, P., Al-Nemrat, A., & Venkatraman, S. (2019). Deep learning approach for intelligent intrusion detection system. *IEEE Access*, 7, 41525–41550.
<https://doi.org/10.1109/ACCESS.2019.2895334>