

Article

A Machine Learning-Based Interest Flooding Attack Detection System in Vehicular Named Data Networking

Arif Hussain Magsi ^{1,2} , Syed Agha Hassnain Mohsan ³ , Ghulam Muhammad ^{4,*}  and Suhni Abbasi ²¹ State Key Laboratory of Networking and Switching Technology, Beijing University of Posts and Telecommunications, Beijing 100876, China; ahmagsi@bupt.edu.cn² Information Technology Center, Sindh Agriculture University, Tandojam 70060, Pakistan; suhni.abbasi@sau.edu.pk³ Optical Communications Laboratory, Ocean College, Zhejiang University, Zhoushan 316021, China; hassnainagha@zju.edu.cn⁴ Department of Computer Engineering, College of Computer and Information Sciences, King Saud University, Riyadh 11543, Saudi Arabia

* Correspondence: ghulam@ksu.edu.sa

Abstract: A vehicular ad hoc network (VANET) has significantly improved transportation efficiency with efficient traffic management, driving safety, and delivering emergency messages. However, existing IP-based VANETs encounter numerous challenges, like security, mobility, caching, and routing. To cope with these limitations, named data networking (NDN) has gained significant attention as an alternative solution to TCP/IP in VANET. NDN offers promising features, like intermittent connectivity support, named-based routing, and in-network content caching. Nevertheless, NDN in VANET is vulnerable to a variety of attacks. On top of attacks, an interest flooding attack (IFA) is one of the most critical attacks. The IFA targets intermediate nodes with a storm of unsatisfying interest requests and saturates network resources such as the Pending Interest Table (PIT). Unlike traditional rule-based statistical approaches, this study detects and prevents attacker vehicles by exploiting a machine learning (ML) binary classification system at roadside units (RSUs). In this connection, we employed and compared the accuracy of five (5) ML classifiers: logistic regression (LR), decision tree (DT), K-nearest neighbor (KNN), random forest (RF), and Gaussian naïve Bayes (GNB) on a publicly available dataset implemented on the ndnSIM simulator. The experimental results demonstrate that the RF classifier achieved the highest accuracy (94%) in detecting IFA vehicles. On the other hand, we evaluated an attack prevention system on Python that enables intermediate vehicles to accept or reject interest requests based on the legitimacy of vehicles. Thus, our proposed IFA detection technique contributes to detecting and preventing attacker vehicles from compromising the network resources.

Keywords: vehicular network; named data networking; interest flooding attack; machine learning

Citation: Magsi, A.H.; Mohsan, S.A.H.; Muhammad, G.; Abbasi, S. A Machine Learning-Based Interest Flooding Attack Detection System in Vehicular Named Data Networking. *Electronics* **2023**, *12*, 3870. <https://doi.org/10.3390/electronics12183870>

Academic Editors: Chen Chen and Zhiyuan Ren

Received: 27 July 2023

Revised: 8 September 2023

Accepted: 12 September 2023

Published: 13 September 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

The exponential global surge in the use of conventional vehicles has undoubtedly enhanced individual convenience but has also escalated the risk of accidents [1]. According to World Health Organization (WHO) statistics from 2023, fatalities from road accidents account for 29% of all reported injuries [2]. To address these challenges, the deployment of a vehicular ad hoc network (VANET) [3] has emerged as a promising solution. The primary objective is to reduce road accidents and optimize traffic flow through efficient vehicle-to-everything (V2X) communication [4]. The onboard unit (OBU)-equipped vehicles possess robust communication, storage, and procession capabilities, effectively handling data transmission, storage, and computational tasks. In addition to safety applications (e.g., collision warning messages, emergency information dissemination, traffic conditions, speed limit warnings, and lane change assistance), vehicles can provide infotainment services [5].

Despite numerous features, VANET faces various challenges due to the traditional transmission control protocol/internet protocol (TCP/IP). The existing TCP/IPs encounter issues such as intermittent connectivity, scalability, security, and privacy concerns in the context of VANETs. Specifically, a vehicular environment requires efficient crucial information dissemination and the ability to handle a large number of users within challenging intermittent conditions. In addition to its limitations, TCP/IP operates as a host-centric network, which adds an additional burden and worsens the overall network latency [6]. Consequently, existing IP-based network architecture is inefficient for VANET. Alternatively, named data networking (NDN) [7] has emerged as a promising network architecture of the information-centric network (ICN) [8] for VANET. Instead of relying on IP addresses for establishing connections and data transmission, NDN employs a unique hierarchical naming approach (e.g., /VNDN/infotainment/music/album/video.mp4) that prioritizes content over its host. This content-centric model proves advantageous. One notable feature of NDN is its in-network content-caching mechanism, which enables vehicles to retrieve content from nearby nodes rather than solely relying on the original host. Table 1 compares IP-based communication limitations and NDN-based solutions address those limitations.

Table 1. Limitations of TCP/IP communications and NDN-based solutions.

Limitations of TCP/IP Communication	NDN-Based Solutions
It is host-oriented	It is a content-oriented network interested in content rather than the host.
It relies on IP addresses.	NDN uses unique content names that reduce the dependency on IP addresses.
It is connection-oriented.	NDN is a connectionless network architecture that does not require establishing explicit connections between two ends.
TCP/IP faces intermittent connectivity issues	The in-network content caching and name-based forwarding strategy support intermittent connectivity.
It secures the channel.	NDN secures content with a cryptographic signature rather than the communication channel.
It has limited scalability in large networks	NDN's architecture allows for scalable content retrieval through its distributed caching mechanism, improving performance in large-scale networks.
Lack of inherent support for multi-cast.	NDN inherently supports multi-cast communication, enabling efficient dissemination of content to multiple recipients simultaneously.

NDN is one of the five research endeavors supported by the National Science Foundation (NSF) within its future internet architecture program, encompassing the fundamental principles of information-centric networking (ICN). In 2009, VAN Jacobson initially proposed the concept of a content-centric network [9], which evolved into NDN under the NSF-funded future internet architecture project [10] as a future internet architecture [11]. NDN uses two types of packets: interest and data packets. The content consumer always initiates the interest packet to request specific data, and the data packet contains the content in response to the interest packet. As mentioned below, nodes within NDN are categorized into three categories according to the situation. (1) Content consumer node: It is a content-intensive entity that initiates communication by broadcasting an interest request for specific content. (2) Content producer node: The content producer node matches the requested content and provides it to the content consumer. (3) Intermediate node: The intermediate node in the NDN architecture serves two distinct roles based on the context of the received packet. Firstly, when the requested content name matches the available content

name within the node's storage, it functions as a content producer, directly providing the requested data to the content consumer node. Secondly, if the intermediate node does not have the requested content, it acts as a relay node, forwarding the interest packet to the next hop in the network. This dynamic behavior of intermediate nodes facilitates efficient content retrieval and distribution, contributing to the overall robustness of the NDN network.

In addition, every NDN node contains three data structures:

Content store (CS): The CS allows the NDN node to cache data packets and serve the content consumers without forwarding interest requests to the content producer every time. The CS reduces network congestion and improves content retrieval.

Pending interest table (PIT): The PIT stores unsatisfied interest requests and their interfaces in a table until the interest request is satisfied.

Forward information base (FIB): The FIB is responsible for forwarding unsatisfied interest packets to the next hops. Unlike traditional IP-based routing, NDN's FIB entries are indexed with name prefixes rather than IP addresses, as described in [12]. The entries in FIB contain next-hop information. This feature allows routers to direct interest packets to one or multiple next-hops, depending on the forwarding strategy, enabling efficient multi-path forwarding in the network. Figure 1 shows the NDN communication architecture in vehicular NDN (VNDN).

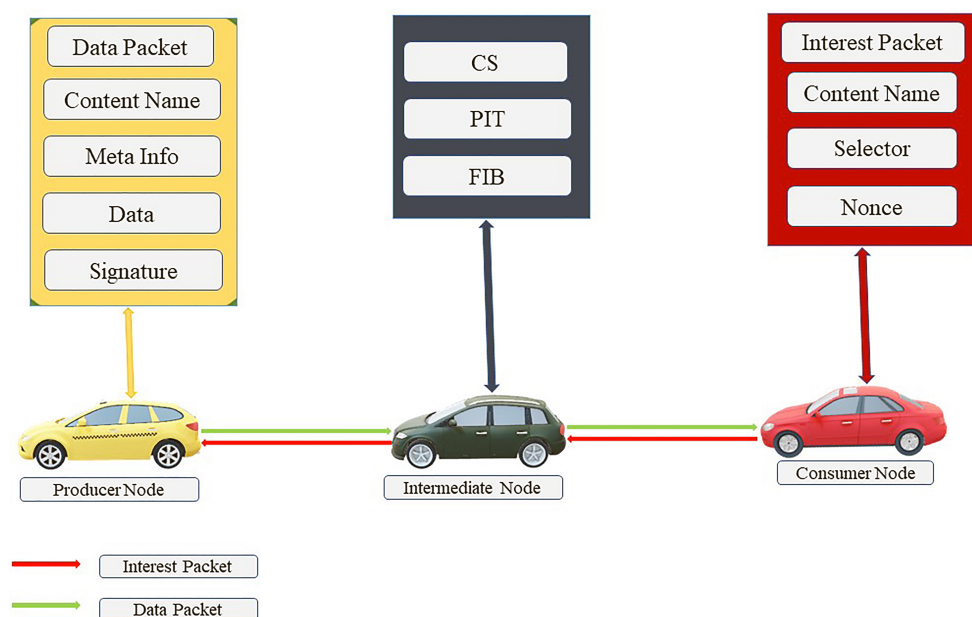


Figure 1. VNDN transmission architecture.

Regardless of the fact that NDN has numerous features, it is highly vulnerable to various attacks, such as interest flooding attacks (IFAs) [13], content poisoning attacks (CPAs) [14], man-in-the-middle attacks [15], and illusion attacks [16]. On top of these attacks, IFA stands out as one of the most prevalent in VNDN. The IFA is a variant of a distributed denial of service (DDoS) attack, where a content consumer initiates IFA in VNDN with a storm of non-existing interest requests. The IFA deliberately depletes resources, including PIT, CS, network bandwidth, and producer resources. This attack is executed by inundating the network with excessive interest packets. By overwhelming the system, the attacker exhausts NDN resources, rendering them inaccessible to legitimate consumers and causing disruption in the network's operation [17]. IFA attackers can consume network resources by employing two distinct techniques: (1) a non-existing interest packet: In this approach, the attacker generates random interest packets that contain invalid requests, such as /VNDN/infotainment/music/5453.txt, where the attacker requests a text file in

the music prefix. These packets refer to content that does not exist in the network. Consequently, intermediate nodes cannot resolve and retain such requests in the PIT. This results in unnecessary resource consumption and potential network congestion. Thus, PIT can be choked with forged interest packets. (2) Valid data request: The attackers target content producers with enormous legitimate interest packets using forged nounce [13]. For example, an attacker initiates an interest packet with /VNDN/infotainment/music/nounce, where nounce is a random value. Using such forged interest packets significantly impacts the producers and network routers by traversing the network resources. Figure 2 visually represents IFA in VNDN, where the attack scenario poses a significant threat to the communication infrastructure of connected intermediate vehicles using a non-existing and valid interest packet.

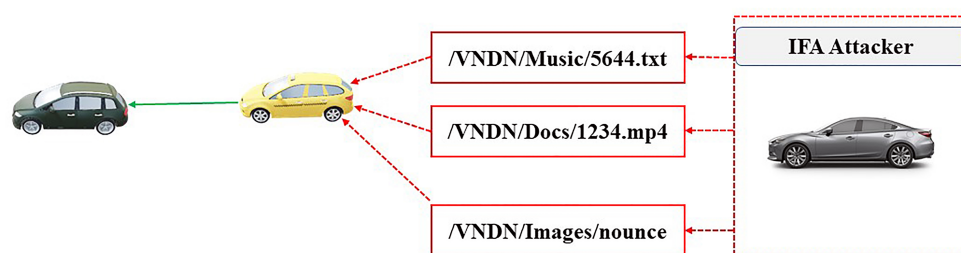


Figure 2. IFA in VNDN.

To address the challenge of IFA in NDN, researchers have explored various approaches, including threshold-based IFA detection [18], statistical-based countermeasures [19], reputation-based IFA detection [20], rating-based approaches [21], and charging/rewarding mechanisms [22]. Although these approaches have contributed significantly to detecting IFA in NDN, they have not provided an efficient solution for accurately detecting and preventing such attacks in the VNDN. On the other hand, ML is gaining momentum in anomaly detection [23] in various fields, including the Internet of Things (IoT) [24], healthcare [25], image processing [26], spam detection [27], unmanned aerial vehicles (UAVs) [28], VANET [29,30], NDN [31] and so on. Specifically, ML has yielded substantial advancements by significantly enhancing capabilities in diverse aspects, including intrusion detection [32], optimal resource allocation [33], offloading strategies, and precise mobility pattern forecasting. Despite its widespread application in various domains, none of the previously mentioned research has explored the use of ML in VNDN. Unlike traditional approaches for detecting IFA in NDN, we are the first to propose an ML-based efficient solution to detect and prevent IFA in VNDN. Considering the challenges and limitations highlighted in the existing literature, the main focus of this research is to propose a resilient network framework that effectively tackles IFA through ML classifiers. To achieve this, we evaluate and propose the most accurate ML classifier for CPA detection. The significant contributions of this research are as follows:

- We propose an ML-based classification technique to identify attackers and legitimate vehicles.
- We evaluate the accuracy of five ML classifiers and propose the most accurate algorithm for IFA detection.
- Based on our ML-based detection results, we propose a simulation-based IFA prevention system in intermediate nodes.

By focusing on the detection and prevention of IFA in VNDN, this research aims to fortify the resilience of vehicular communication systems. Our ML-based approach mitigates the immediate threats posed by IFAs and establishes a foundation for secure VNDN ecosystems. The subsequent sections of this paper are structured as follows:

Section 2 presents a detailed existing work and their limitations in detecting and preventing IFA. Section 3 delves into a comprehensive analysis of the system model, network elements, and proposed ML-based IFA detection and prevention system. We provide IFA detection and prevention results in Section 4 and conclude the paper in Section 5. Finally, Section 6 presents future work.

2. Related Work

The scientific community has seen a marked surge in interest and contributions toward combatting cyber-attacks [34], specifically in the VNDN realm [35]. However, developing and implementing efficient security measures for safeguarding VNDN is in its infancy. Specifically, ML techniques have been explored in detecting attacks [36,37]. Reference [38] discovered IFA in NDN by expressing how a huge number of interest packets can overwhelm the network. Subsequently, numerous research papers have delved into the IFA using several techniques; for example, the authors in [31] proposed an ML-based classification technique for detecting IFA on tree topology (small-scale topology) and Rocketfuel ISP topology (large-scale topology). Another solution for detecting IFA presented a centralized controller-based approach [39]. In this mechanism, a router maintains an unsatisfied interest request threshold system. Based on a predetermined threshold value, a router decides to identify IFA nodes. However, this approach has certain limitations. The metrics used to identify IFA nodes may lead to false detection, and the system might fail to detect IFA, especially in scenarios with significant legitimate traffic or when content producers are unavailable. These shortcomings highlight the need for more advanced and robust detection mechanisms to combat IFA in VNDN effectively. Similarly, in reference [40], the authors presented a threshold-based system for identifying IFA within a local PIT. Instead of relying on a centralized router, the PIT manages a predetermined threshold system in this approach. The threshold-based approach allows the local PIT to assess incoming interest requests autonomously and identify potential IFA scenarios based on the predefined threshold criteria. This solution aims to enhance the efficiency and accuracy of IFA detection within the network by decentralizing the detection process and employing local PIT mechanisms. However, it is essential to consider the trade-offs and limitations of this approach, particularly in terms of scalability and adaptability to various network conditions.

In order to prevent PIT from exhaustion, Wang et al. [41] proposed decoupled legitimate and timeout interest requests. Each router maintains a timeout interest request in this architecture in an m-list. If the prefix is already in the m-list, the router forwards the interest packet outright, avoiding PIT storage. While this approach mitigates some of the impacts of an IFA, it fails to provide a comprehensive solution to thwart such attacks entirely. Additionally, legitimate requests are adversely affected by this approach. In addition, attackers can misuse the router's resources by forging names to flood the m-list, causing the solution to become inefficient. In contrast, few authors have proposed a hypothesis-testing theory-based approach in the literature; for example, the authors in [42] exploited hypothesis-based testing theory for formulating a comprehensive likelihood static hypothesis test theory (SHTT) tailored to address evolving attacks, particularly in incorporating NDN with TCP/IP, which is difficult to address using conventional approaches. Similarly, the authors in [43] tackled IFA by employing a detection approach based on SHTT. The test is free from any reliance on router characteristics or measured values. The framework comprises two main scenarios: (i) When all traffic parameters are known, an optimal test is formulated, and its statistical performance is thoroughly evaluated. (ii) The framework introduces a linear parametric model, which estimates unknown parameters and enables the development of a practical test. However, it is crucial to recognize that the scheme assessment is restricted to a basic binary tree graph with merely eight clients and one adversary. Consequently, assessing the scheme's efficacy under more extensive networks or during distributed attacks presents considerable challenges.

Meanwhile, the authors in [44] presented a Markov-based IFA detection system. This approach involves creating a space vector determined by the fluctuations in the PIT occupancy rate, and the network's state is evaluated using a quantized value. By calculating the Euclidean distance, the system distinguishes between legitimate and malicious interest packets, achieving a high detection rate. However, a notable drawback of this approach is its significant consumption of network resources, particularly when identifying interest packets within a large volume of NDN network traffic.

Moreover, Xin et al. [45] introduced a cumulative entropy-oriented IFA detection system that monitors abnormal interest requests and identifies malicious prefixes using entropy theory. Similar to our proposed work, few researchers have incorporated ML for IFA detection in NDN. For instance, Azmi et al. [46] proposed a feature selection technique for IFA detection. In this research, the authors employed an information gain and data reduction approach to identify pertinent features from the UNSW-NB 15 dataset, specifically for detecting DDoS attacks. The dataset underwent testing using three distinct classification methods: artificial neural network (ANN), naïve Bayes (NB), and decision table. The experiment's outcomes were thoroughly analyzed, and evaluation metrics, such as true positive (TP), false positive (FP), precision, and accuracy were utilized to categorize the data into two classes: attacks and normal. By employing this methodology, the researchers aimed to develop an effective DDoS attack detection system that relies on relevant characteristics, which can be classified with high precision and accuracy. Similarly, the authors in [47] leveraged ML for detecting IFA. The authors collected a dataset with 12 features implemented in the ndnSIM simulator and performed the accuracy of three ML classifiers, including K-nearest neighbor (KNN), decision tree (DT), and ANN. The results showed that DT outperformed other classifiers with 85.42% accuracy while KNN achieved 81% and ANN yielded 80% accuracy. Although the proposed ML-based IFA detection system is important, it could not achieve high accuracy.

Different from traditional rules-based, threshold-based, and other partial ML-based approaches for detecting IFA, we propose an efficient IFA detection mechanism that exploits five ML classifiers and depicts an accurate algorithm for IFA detection. Additionally, we propose a prevention mechanism to cope with IFA in VNDN. To the best of our knowledge, this study presents the first contribution that leverages ML classifiers for detecting and preventing IFA attacker vehicles in VNDN. Thus, we address the limitations of existing research work by exploiting a novel ML-based classification system in detecting and preventing IFA in VNDN. A comprehensive list of notations used throughout this paper is presented in Table 2.

Table 2. Notations and their descriptions.

Notation	Description
CC_R	Content Consumer Reputation
D_{pkt}	Data Packet
I_{pkt}	Interest Packet
Cnt	Content
C_C	Content Consumer
CC_R^n	Content Consumer New Reputation
$CC_R^n - 1$	Content Consumer Previous Reputation
$AgrCC_R$	Aggregate Content Consumer Reputation

3. System Model

This section is structured into three interconnected and comprehensive subsections, each contributing to the overarching goal of our research. Initially, we present the design of a purpose-built network architecture exclusively for IFA detection within VNDN environments. This architecture is a fundamental framework for efficient data analysis and processing in VNDN. In the second subsection, we leverage ML binary classification techniques to identify attacker and non-attacker vehicles using five ML classifiers. By employing advanced ML techniques, we aim to achieve high accuracy and reliability in identifying potential IFA attacker vehicles amidst the complex data flow within VNDN. Finally, the third subsection introduces our proposed algorithm for IFA prevention, which proactively mitigates the impact of potential intrusions. This algorithm contributes to the

enhanced security and resilience of the VNDN system, safeguarding the network against IFA attacks and fostering a safer and more trustworthy VNDN.

3.1. Proposed Network Architecture

In our proposed network architecture, the intermediate node receives all interest packets, including both satisfied and unsatisfied ones. These packets are then shared with roadside units (RSUs) through a push-based beacon message dissemination system. The interaction between intermediate nodes and RSUs, particularly in terms of communication and data sharing, is elaborated below:

Vehicle to RSU Communication

The existing NDN content propagation is a pull-oriented content retrieval. To extend the scope of current NDN from a pull-based content retrieval to a push-based content propagation, beacon message propagation has been proposed in the literature [48,49], where nodes can broadcast messages to neighboring nodes without considering interest packets. Considering the push-based beacon message dissemination in VNDN, we design a network architecture as depicted in Figure 3, where an intermediate node receives an interest packet from the content consumer and queries RSU about the legitimacy of the content consumer. Based on the reputation of the content consumer, the intermediate vehicles decide to accept or reject the interest request.

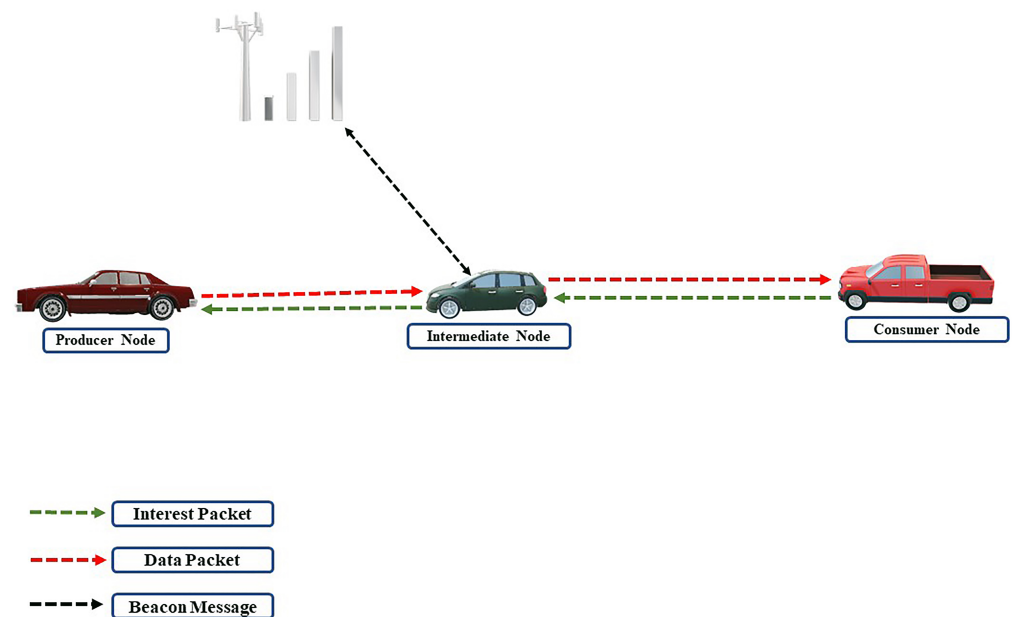


Figure 3. Vehicle to RSU communication.

To classify content consumers as attackers or legitimate vehicles, the RSUs employ an ML algorithm to distinguish them effectively. Figure 4 depicts our network architecture, where a consumer sends an interest packet to request data from the intermediate node. Subsequently, the intermediate node sends the interest packet (satisfied or unsatisfied) to the RSU. The RSU then performs ML classification to classify content consumers as attackers or legitimate vehicles.

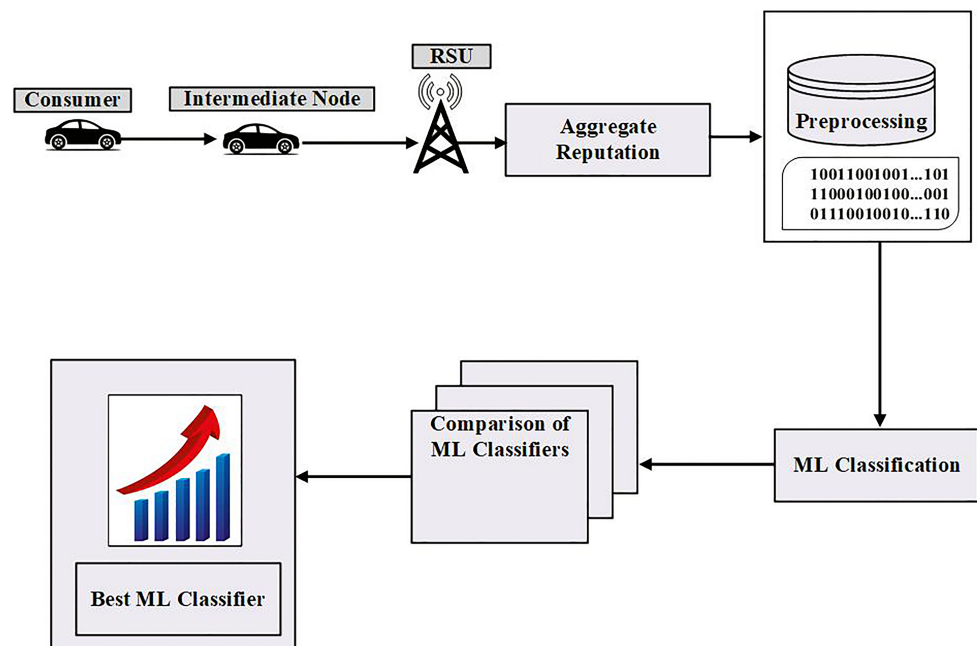


Figure 4. Proposed network architecture.

3.2. ML Classification-Based Attack Detection

This section aims to assess the performance of ML algorithms in distinguishing between legitimate and attacker nodes using a publicly available dataset simulated in ndnSIM. The main objective of this analysis is to leverage ML classification techniques using binary classification on the dataset, differentiating between attackers and legitimate vehicles. Our study primarily focuses on identifying and characterizing the behavior of content consumer vehicles through the utilization of ML algorithms and performance metrics, including accuracy, precision, recall, and F1 score, to evaluate the accuracy of the models.

3.2.1. Dataset Collection

Selecting a suitable dataset for attack detection is one of the most crucial tasks. Considering the relevance and suitability of the dataset, we obtained a simulation-based dataset from a publicly available source [50], implemented using the ndnSIM simulator [51] for IFA detection. The dataset comprises 10 essential features organized in a CSV file, namely InInterests, OutInterests, DropInterests, InData, OutData, InSatisfiedInterest, OutSatisfiedInterest, PITSize, PITSizeInt, and Attack. This dataset uses tree and DFN topology for simulating IFA. Notably, the 'Attack' feature is the target variable, indicating an attacker's presence (designated by the value 1) or a non-attacker scenario (denoted by the value 0). The traces in the dataset contain 24,660 interest requests with attack and legitimate packets.

3.2.2. Data Preprocessing

The preprocessing in ML has a significant role in enhancing the quality and suitability of the data. We involved various steps in data preprocessing, including feature selection, data clearing, removing noise, dataset splitting, and applying cross-validation. Considering the importance of preprocessing, this research initially consolidated all the individual CSV dataset files into a single CSV file. Secondly, we removed missing and duplicate records from the dataset. To ensure the reliability of our proposed model, we randomly shuffled the dataset with 70% as training and 30% as testing and applied a 10-fold cross-validation technique. This approach effectively demonstrated the accuracy of our ML approach in detecting and predicting the legitimacy of vehicles.

3.2.3. Classification

In order to measure the efficiency of various ML algorithms, we evaluated and compared the performance of five ML classifiers on training and testing datasets. During training, these classifiers learned the underlying patterns and relationships between the input features and the corresponding class labels. This learning process enabled the models to learn and understand the dataset's intricacies. The selection of ML classifiers in this study, namely DT, K-nearest neighbor (KNN), random forest (RF), Gaussian naïve Bayes (GNB) and logistic regression (LR), was made with careful consideration of their unique characteristics and their alignment with the objectives of interest flooding attack (IFA) detection in the context of VNDN. These classifiers were selected based on their aptitude for addressing the challenges associated with identifying abnormal patterns in interest propagation, a fundamental requirement for effective IFA detection in VNDN. The rationale for specific classifier selection is mentioned below:

- DT: The DT [52] is a significant method for reaching conclusions based on a set of rules derived from a tree-like structure. We selected DT due to its outstanding capabilities in capturing nonlinear relationships and handling categorical features often utilized in VNDN datasets. Its interpretability offers insights into the decision-making process, aiding in understanding detected attack patterns. The tree comprises two nodes: a decision node and a leaf node. Decision nodes determine the attribute that needs to be selected for further analysis, while leaf nodes represent the final class outcome. The DT employs a top-down approach to provide results. The root is placed at the top of the tree, which acts as the initial decision node. DT uses the information gain technique to select each subsequent DT node, ensuring that each part of the tree selects the most informative attributes. This enables DT to classify and predict results based on the input data characteristics and patterns.
- KNN: The KNN algorithm [53] is the most used ML classifier, popular for its effectiveness in dealing with large datasets. KNN is deemed a suitable ML classifier for recognizing local clustering, which is an essential trait for detecting attack occurrences that might exhibit spatial proximity. It is a simple and flexible classifier that can be applied to regression and classification purposes. The KNN involves categorizing the latest data points by assigning them to the most common class among their K-nearest neighbors in the training set. The KNN then provides the majority class label or the average value of those neighbors. Considering the appropriate value of K is essential and depends on the specific characteristics of the dataset, making KNN a versatile and adaptive choice for various ML scenarios.
- RF: RF [54] combines multiple base models to make predictions. Given the potential noise and outliers in VNDN data, RF's ability to handle such variations becomes crucial. This approach is often called "bootstrapping and aggregation", where the majority vote of the base models on the test data determines the final result. In the RF approach, the data are fed to the base models using row sampling with replacement, a method known as bagging.
- GNB: GNB [55] is a simple yet effective classification method that employs Bayes' theorem for predicting the class of unlabeled data points. We selected GNB for its efficiency in high-dimensional data handling and probabilistic nature, allowing it to capture the likelihood of feature co-occurrences relevant to IFA scenarios. It calculates the prior probabilities of different classes and utilizes this information to make predictions on new, unseen data. One of the key assumptions of GNB is the independence of features, which means that it assumes each feature contributes to the classification independently of other features. This independence assumption simplifies the computation and makes GNB computationally efficient. Due to its simplicity and efficiency, GNB is particularly well-suited for applications with many features and is commonly used in various ML tasks.
- LR: The LR [56] is a statistical method used for predicting the probability of categorical variables, especially in two-class classification problems. It is a well-established binary

classification technique for IFA detection. It utilizes a logistic function to calculate an event's likelihood.

3.2.4. Model Evaluation

Model evaluation plays a pivotal role in assessing the performance and efficacy of an ML model. It assesses the model's ability to generalize to unseen data and accurately predict desired outcomes. Throughout the evaluation process, we utilized several metrics, including accuracy, precision, recall, and the F1 score, to thoroughly assess our model. In particular, during the evaluation phase for classifying "Positive Reputation" samples, we focused on four key performance metrics:

- True positive (TP): represents the count of positive samples correctly classified.
- False positive (FP): indicates the count of samples incorrectly classified as positive.
- True negative (TN): refers to the count of negative samples correctly classified.
- False negative (FN): signifies the count of samples incorrectly classified as negative.

By scrutinizing these metrics, we obtained valuable insights into the model's accuracy in distinguishing between positive and negative instances. This comprehensive evaluation allowed us to identify potential areas for improvement. Below are the corresponding mathematical models for each algorithm:

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

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

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

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

3.3. Attack Prevention System

Upon successful IFA detection using ML techniques, the subsequent crucial step is preventing such attacks. To achieve this, we propose Algorithm 1, which empowers intermediate vehicles to make informed decisions regarding interest requests from content consumers. Leveraging the results obtained from ML classification, Algorithm 1 allows the vehicles to selectively accept or reject these requests based on the legitimacy of the content consumer, thereby mitigating the IFA attack and enhancing network security and reliability.

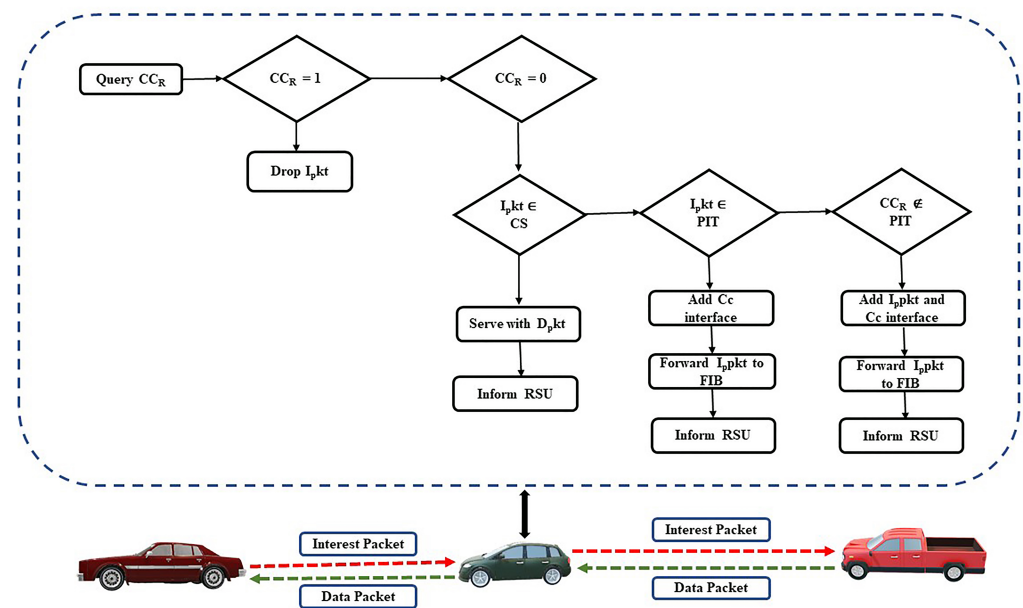
Our proposed algorithm outlines a verification process for incoming I_{pkts} at an intermediate node. When an intermediate node receives an interest request, it queries the content consumer's reputation from RSU. The algorithm takes I_{pkts} as the input variable, representing the interest packet received from the content consumer. If I_{pkts} has a value of 1, the algorithm identifies the consumer as an attacker and discards the interest packet. If I_{pkts} has a value of 0, the algorithm further verifies the interest packet. Additionally, the algorithm shares information about this content transmission with the RSUs. This verification mechanism ensures the network efficiently handles incoming interest packets from content consumers. It promptly discards attacker vehicle interest requests. By employing this mechanism, the network prevents attacker vehicles. Figure 5 depicts our proposed IFA prevention system in V2X.

Algorithm 1 Interest packet verification mechanism.**Require:** I_{pkt}

```

1: Query RSU
2: if  $I_{pKT} = 1$  then
3:   Discard  $I_{pkt}$ 
4: end if
5: if  $I_{pkt} = 0$  then
6:   Check CS
7:   if  $Cnt \in CS$  then
8:     Create a  $D_{pkt}$ 
9:     Send  $D_{pkt}$  to the  $C_C$ 
10:    Share information with RSU
11:  else
12:    if  $Cnt \in PIT$  then
13:      Add interface
14:      Remove  $I_{pkt}$ 
15:      Share information with RSU
16:    else
17:      Add entry in the PIT
18:      Forwarded  $I_{pkt}$  to FIB
19:      Share information with RSU
20:    end if
21:  end if
22: end if

```

**Figure 5.** IFA prevention mechanism.**4. Experimental Results and Discussion**

Our experimental evaluation comprises two main components: IFA detection and prevention in VNDN. In the first part, we present the outcomes of IFA detection, where we provide the evaluated results of various ML classifiers. Subsequently, we present the IFA prevention results in the second part. This division allows us to comprehensively analyze our ML-based approach's effectiveness in identifying and mitigating IFA attacks within the VNDN environment.

4.1. ML Evaluation Results

To achieve our objectives, we evaluated various ML algorithms based on the behavior of content consumers. The results of our proposed model, including precision, recall, and F1 score, are presented in Table 3. To further assess the performance of each ML algorithm, we employed precision-recall curves and receiver operating characteristic (ROC) curves for visualization. The precision and recall curves are commonly used for evaluating binary classification performance. On the other hand, the ROC curve illustrates the trade-off between precision and recall values, with a larger area under the curve indicating higher values for both metrics. A high precision value corresponds to a low false positive rate, while a high recall value corresponds to a low false negative rate. Our findings demonstrate that RF performed exceptionally well in detecting IFA attackers, achieving outstanding accuracy in our experiments.

Table 3. ML classifier performance.

ML Classifiers	Precision	Recall	F1 Score
DT	0.85	0.87	0.86
KNN	0.87	0.81	0.84
RF	0.91	0.87	0.89
NB	0.99	0.57	0.72
LR	0.98	0.57	0.72

Visualized Results

To achieve the visual performance of our ML models, we employed visualization techniques using accuracy calculation and ROC analysis. Initially, we utilized precision-recall and ROC curves to evaluate the trade-off between precision and recall in binary classification. Notably, the area under the curve (AUC) revealed exceptional performance in the RF classifier that achieved the highest AUC value. The RF outperforms other ML classifiers in accurately classifying IFA attackers and legitimate vehicles. Figures 6–10 show the precision-recall curve performances for DT, KNN, GNB, and LR classifiers, respectively. Finally, Figure 11 presents the consolidated and comparative evaluation, providing a comprehensive view of the performances of different ML classifiers in tackling IFA detection.

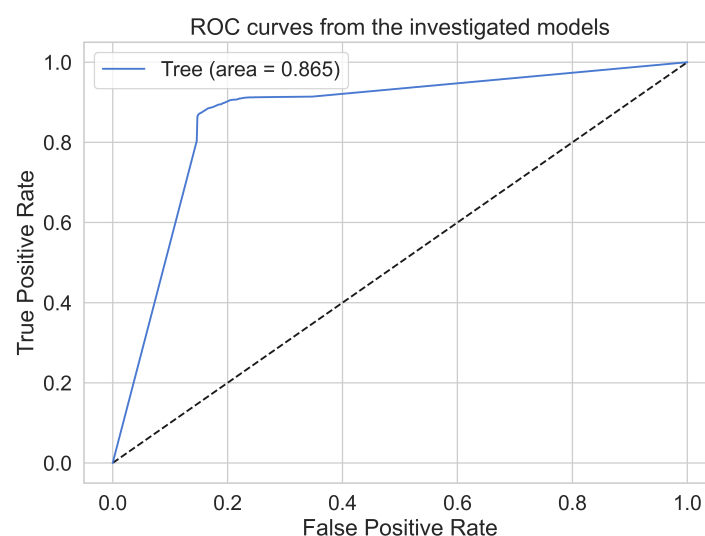


Figure 6. Decision tree accuracy.

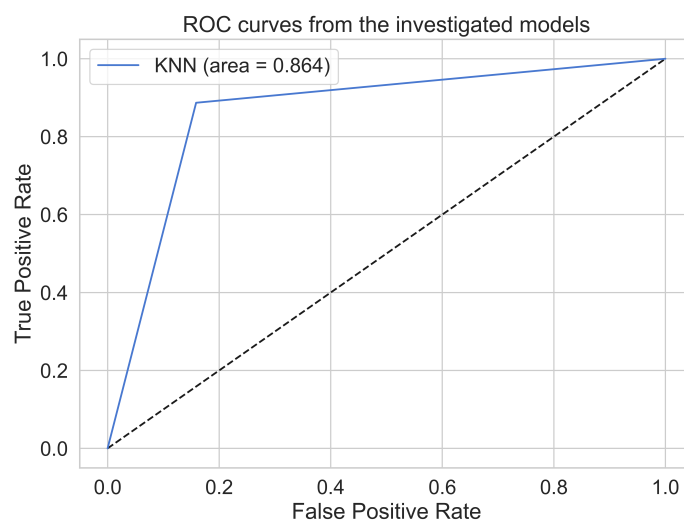


Figure 7. K-nearest neighbor accuracy.

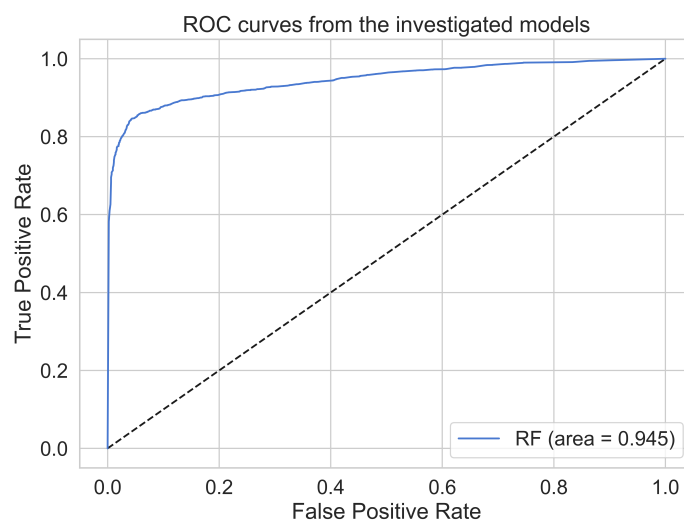


Figure 8. Random forest accuracy.

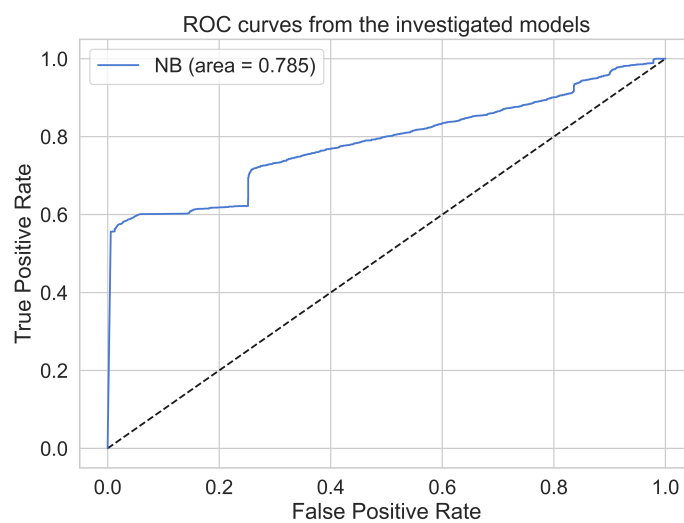


Figure 9. Gaussian naïve Bayes accuracy.

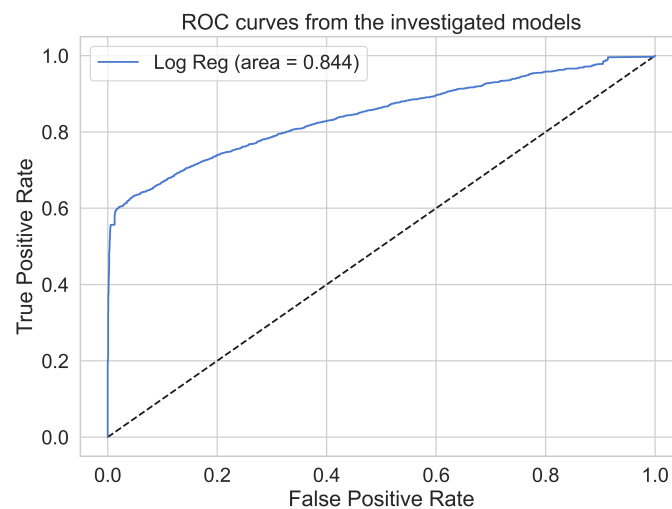


Figure 10. Logistic regression.

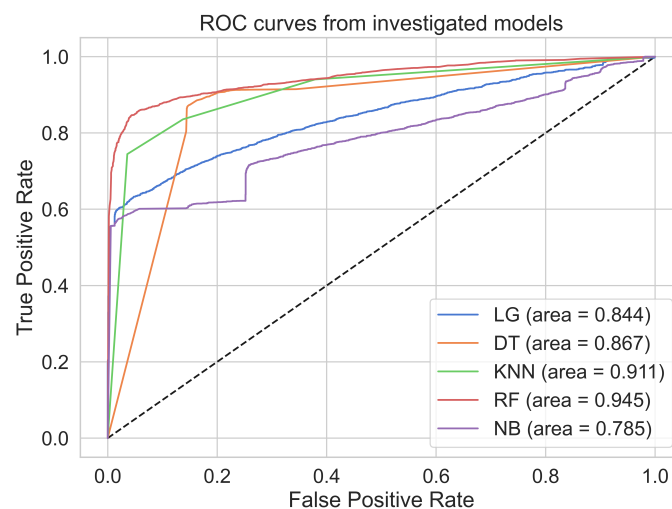


Figure 11. Consolidated accuracy.

In pursuing robust and accurate IFA detection, our study meticulously examined the performance of five distinct classifiers. Through a systematic classification approach, we observed compelling variations in accuracy across these classifiers. As depicted in the consolidated results in Figure 11, the RF classifier emerged as the front-runner, achieving an accuracy of 94%. On the other hand, KNN exhibited a commendable accuracy of 90%. The remaining classifiers exhibited accuracy in the 80s range. These consolidated ROC results collectively shed light on the multifaceted landscape of IFA detection in VNDN.

4.2. IFA Prevention Results

To validate our proposed IFA prevention system in VNDN, we conducted simulations using an Intel Desktop Core i7 CPU operating at 2.6 GHz, with 16 GB of RAM, and running on the Windows 10 operating system. In our simulation, we randomly considered 10 content consumers with attacks (1) and legitimate (0). Our proposed algorithm classified attackers and legitimate information as illustrated in Figure 12, where attackers are identified with red and the legitimate source is identified with green. On the other hand, the existing NDN system could not consider the legitimacy of content consumers and received every interest packet, as shown in Figure 13. Thus, our proposed prevention system identifies the legitimacy of vehicles, followed by an ML-based detection system. Thus, our proposed network architecture classifies the attacker vehicles first and prevents them from violating the PIT and CS of intermediate vehicles.

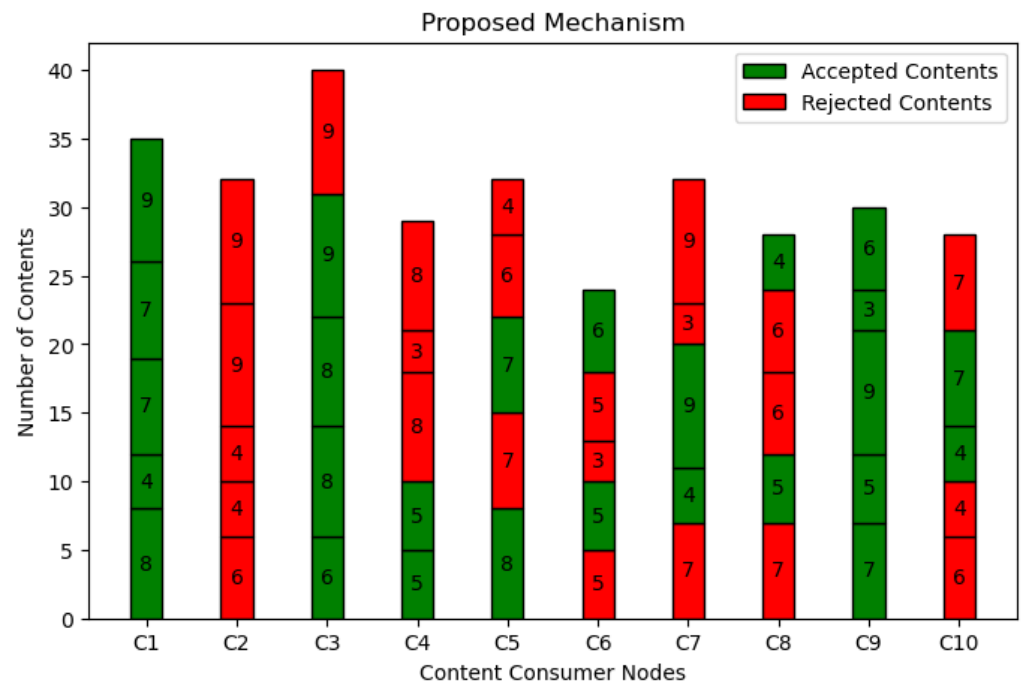


Figure 12. IFA prevention mechanism results.

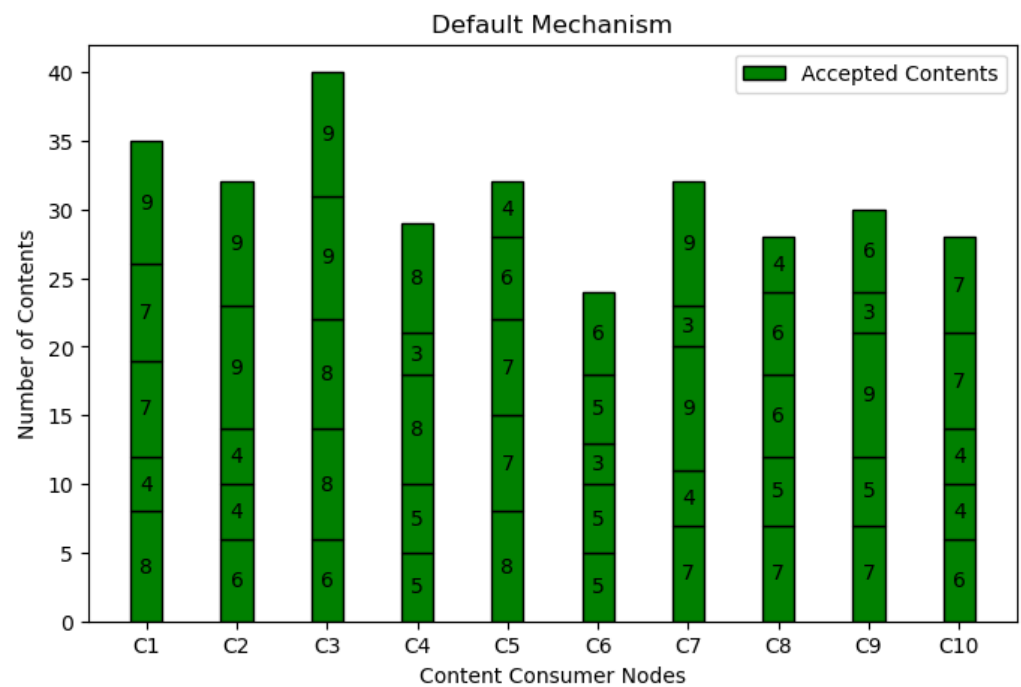


Figure 13. Default NDN mechanism results.

Figure 12 depicts the efficiency of our proposed IFA prevention system. In this validation scenario, both legitimate and malicious interest requests are simulated by content consumers. Leveraging the capabilities of our devised algorithm, the intermediate vehicle undertakes a pivotal role in this endeavor. A distinctive dichotomy emerges in the depicted visualizations: The algorithm adeptly identifies and rejects invalid content requests, visually represented by the prominent red indicators, while conversely accepting valid content consumers through green color. These visual representations underscore the system's capability to discriminate between legitimate and unauthorized interest propagation, reaffirming the effectiveness of our proposed approach. In addition, we measured the efficiency

of the default content acceptance at an intermediate node, as shown in Figure 13, where an intermediate node accepted every interest packet without considering the reputation of the content producer.

4.3. Discussion

The experimental results presented in the previous section provide valuable insights into the effectiveness of our ML-based approach for classifying IFA in VNDN. In this discussion, we delve into the implications of these findings, highlight the strengths and limitations of our approach, and contextualize our results within the broader landscape of vehicular network security. The first component of our evaluation focused on IFA detection, where various ML classifiers were employed to classify content consumers as attackers or legitimate entities. We measured the efficiency of classifiers in terms of precision, recall, and F1 score. These metrics collectively indicate the classifier's ability to identify malicious and benign entities within the network accurately. The high area under the curve (AUC) value obtained by RF reinforces its effectiveness, while the comparative visualization in Figure 11 provides a comprehensive view of the performance differences across different ML classifiers. Notably, RF's accuracy of 94% showcases its superiority in IFA detection compared to other classifiers. The highest accuracy in detecting IFA reflects our algorithm's ability to discern between malicious and benign entities. In contrast, the existing NDN system's limitations in evaluating content consumer legitimacy are illustrated in Figure 13. This highlights the significance of our proposed prevention system in enhancing the network security by identifying attackers and preventing their interference with the PIT and CS of intermediate vehicles.

5. Conclusions

The IFA is one of the most vulnerable attacks in VNDN. The IFA targets intermediate nodes, flooding them with unsatisfying interest requests and saturating network resources, particularly the PIT. This research paper aims to address IFA challenges in VNDN. To cope with IFA in VNDN, the study explored the potential of ML classifiers for detecting the IFA with high accuracy. In this connection, we employed five ML classifiers on a publicly available dataset implemented on the ndnSIM simulator. The study compared their accuracy in identifying IFA vehicles. The findings demonstrated that RF achieved an accuracy of 94% in detecting IFA. This highlights the effectiveness of the proposed IFA detection technique, which empowers VNDN to classify attacker vehicles and proactively prevent them from violating the network. On the other hand, our proposed IFA prevention system classified and prevented attackers with 100% accuracy. The implications of this research are significant, as it contributes to enhancing the security and reliability of vehicular named data networking. Future work may explore further improvements in ML-based detection methods and investigate the application of other ML algorithms to tackle additional security concerns in VNDN. Ultimately, these advancements will foster the development of more resilient and secure VANET architectures, thus promoting the continued growth and success of intelligent transportation systems.

6. Future Work

The proposed research work exploited ML classifiers using a batch learning-based system, whereas deep learning (DL) classifiers can be applied using an incremental learning system, and detailed comparative measures can be considered in the future. Moreover, this study is limited to evaluating the legitimacy of content consumer vehicles, whereas intermediate vehicles can be attackers. Thus, the legitimacy of intermediate vehicles can be determined in future research.

Author Contributions: Conceptualization, A.H.M., G.M. and S.A.; methodology, A.H.M., S.A.H.M., G.M. and S.A.; software, A.H.M., S.A.H.M. and S.A.; validation, S.A.H.M. and G.M.; formal analysis, A.H.M. and S.A.H.M.; investigation, A.H.M., G.M. and S.A.; resources, A.H.M. and G.M.; data curation, S.A.H.M., G.M. and S.A.; writing—original draft, A.H.M., G.M. and S.A.; writing—review and editing, G.M. and S.A.; visualization, A.H.M., S.A.H.M. and S.A.; funding acquisition, G.M. All authors have read and agreed to the published version of the manuscript.

Funding: This research is funded by the Researchers Supporting Project number (RSP2023R34), King Saud University, Riyadh, Saudi Arabia.

Data Availability Statement: Publicly available datasets were analyzed in this study. This data can be found here: <https://github.com/nk10121989/NDN-IFA-Feature-Selection>.

Acknowledgments: The authors acknowledge the Researchers Supporting Project number (RSP2023R34), King Saud University, Riyadh, Saudi Arabia.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Min, H.; Fang, Y.; Wu, X.; Lei, X.; Chen, S.; Teixeira, R.; Zhu, B.; Zhao, X.; Xu, Z. A fault diagnosis framework for autonomous vehicles with sensor self-diagnosis. *Expert Syst. Appl.* **2023**, *224*, 120002. [CrossRef]
2. World Health Organization Statistics. 2023. Available online: <https://www.who.int/data/gho/publications/world-health-statistics> (accessed on 14 July 2023).
3. Shelke, S.; Pundge, A. A Comparative Analysis and Study of Vehicular Ad Hoc Network. In Proceedings of the International Conference on Applications of Machine Intelligence and Data Analytics (ICAMIDA 2022), Aurangabad, India, 22 December 2022; Atlantis Press: Amsterdam, The Netherlands, 2023; pp. 366–381.
4. Chen, S.; Hu, J.; Zhao, L.; Zhao, R.; Fang, J.; Shi, Y.; Xu, H. *Cellular Vehicle-to-Everything (C-V2X)*; Springer: Berlin/Heidelberg, Germany, 2023.
5. Liang, L.; Peng, H.; Li, G.Y.; Shen, X. Vehicular communications: A physical layer perspective. *IEEE Trans. Veh. Technol.* **2017**, *66*, 10647–10659. [CrossRef]
6. Naem, M.A.; Rehman, M.A.U.; Ullah, R.; Kim, B.S. A comparative performance analysis of popularity-based caching strategies in named data networking. *IEEE Access* **2020**, *8*, 50057–50077. [CrossRef]
7. Khelifi, H.; Luo, S.; Nour, B.; Moungla, H.; Faheem, Y.; Hussain, R.; Ksentini, A. Named data networking in vehicular ad hoc networks: State-of-the-art and challenges. *IEEE Commun. Surv. Tutor.* **2019**, *22*, 320–351. [CrossRef]
8. Xylomenos, G.; Ververidis, C.N.; Siris, V.A.; Fotiou, N.; Tsilopoulos, C.; Vasilakos, X.; Katsaros, K.V.; Polyzos, G.C. A survey of information-centric networking research. *IEEE Commun. Surv. Tutor.* **2013**, *16*, 1024–1049. [CrossRef]
9. Jacobson, V.; Smetters, D.K.; Thornton, J.D.; Plass, M.F.; Briggs, N.H.; Braynard, R.L. Networking named content. In Proceedings of the 5th International Conference on Emerging Networking Experiments and Technologies, Rome, Italy, 1–4 December 2009; pp. 1–12.
10. Ambrosin, M.; Compagno, A.; Conti, M.; Ghali, C.; Tsodik, G. Security and privacy analysis of national science foundation future internet architectures. *IEEE Commun. Surv. Tutor.* **2018**, *20*, 1418–1442. [CrossRef]
11. Ahmed, S.H.; Bouk, S.H.; Yaqub, M.A.; Kim, D.; Song, H.; Lloret, J. CODIE: Controlled Data and Interest Evaluation in Vehicular Named Data Networks. *IEEE Trans. Veh. Technol.* **2016**, *65*, 3954–3963. [CrossRef]
12. Song, T.; Yuan, H.; Crowley, P.; Zhang, B. Scalable name-based packet forwarding: From millions to billions. In Proceedings of the 2nd ACM Conference on Information-Centric Networking, San Francisco, CA, USA, 30 September–2 October 2015; pp. 19–28.
13. Benmoussa, A.; Kerrache, C.A.; Lagraa, N.; Mastorakis, S.; Lakas, A.; Tahari, A.E.K. Interest Flooding Attacks in Named Data Networking: Survey of Existing Solutions, Open Issues, Requirements, and Future Directions. *Acm Comput. Surv.* **2022**, *55*, 1–37. [CrossRef]
14. Magsi, A.H.; Yovita, L.V.; Ghulam, A.; Muhammad, G.; Ali, Z. A Content Poisoning Attack Detection and Prevention System in Vehicular Named Data Networking. *Sustainability* **2023**, *15*, 10931. [CrossRef]
15. Al-Shareeda, M.A.; Manickam, S. Man-in-the-middle attacks in mobile ad hoc networks (MANETs): Analysis and evaluation. *Symmetry* **2022**, *14*, 1543. [CrossRef]
16. Lo, N.W.; Tsai, H.C. Illusion attack on vanet applications—a message plausibility problem. In Proceedings of the 2007 IEEE Globecom Workshops, Washington, DC, USA, 26–30 November 2007; pp. 1–8.
17. Kumar, N.; Singh, A.K.; Aleem, A.; Srivastava, S. Security attacks in named data networking: A review and research directions. *J. Comput. Sci. Technol.* **2019**, *34*, 1319–1350. [CrossRef]
18. Pu, C.; Payne, N.; Brown, J. Self-adjusting share-based countermeasure to interest flooding attack in named data networking. In Proceedings of the 2019 International Conference on Internet of Things (iThings) and IEEE Green Computing and Communications (GreenCom) and IEEE Cyber, Physical and Social Computing (CPSCom) and IEEE Smart Data (SmartData), Atlanta, GA, USA, 14–17 July 2019; pp. 142–147.

19. Wu, Z.; Feng, W.; Yue, M.; Xu, X.; Liu, L. Mitigation measures of collusive interest flooding attacks in named data networking. *Comput. Secur.* **2020**, *97*, 101971. [\[CrossRef\]](#)
20. Khelifi, H.; Luo, S.; Nour, B.; Mounghla, H.; Ahmed, S.H.; Guizani, M. A blockchain-based architecture for secure vehicular Named Data Networks. *Comput. Electr. Eng.* **2020**, *86*, 106715. [\[CrossRef\]](#)
21. Gasti, P.; Tsudik, G.; Uzun, E.; Zhang, L. DoS and DDoS in named data networking. In Proceedings of the 2013 22nd International Conference on Computer Communication and Networks (ICCCN), Nassau, Bahamas, 30 July–2 August 2013; pp. 1–7.
22. Zhang, X.; Li, R. A charging/rewarding mechanism-based interest flooding attack mitigation strategy in NDN. In Proceedings of the 2019 IFIP/IEEE Symposium on Integrated Network and Service Management (IM), Arlington, VA, USA, 8–12 April 2019; pp. 402–407.
23. Apruzzese, G.; Laskov, P.; Montes de Oca, E.; Mallouli, W.; Brdalo Rapa, L.; Grammatopoulos, A.V.; Di Franco, F. The role of machine learning in cybersecurity. *Digit. Threat. Res. Pract.* **2023**, *4*, 1–38. [\[CrossRef\]](#)
24. Alghanmi, N.; Alotaibi, R.; Buhari, S.M. Machine learning approaches for anomaly detection in IoT: An overview and future research directions. *Wirel. Pers. Commun.* **2022**, *122*, 2309–2324. [\[CrossRef\]](#)
25. Alanazi, A. Using machine learning for healthcare challenges and opportunities. *Inform. Med. Unlocked* **2022**, *30*, 100924. [\[CrossRef\]](#)
26. Guan, Z.; Jing, J.; Deng, X.; Xu, M.; Jiang, L.; Zhang, Z.; Li, Y. DeepMIH: Deep invertible network for multiple image hiding. *IEEE Trans. Pattern Anal. Mach. Intell.* **2022**, *45*, 372–390. [\[CrossRef\]](#) [\[PubMed\]](#)
27. Wu, Z.; Cao, J.; Wang, Y.; Wang, Y.; Zhang, L.; Wu, J. hPSD: A hybrid PU-learning-based spammer detection model for product reviews. *IEEE Trans. Cybern.* **2018**, *50*, 1595–1606. [\[CrossRef\]](#)
28. Rasheed, I.; Asif, M.; Ihsan, A.; Khan, W.U.; Ahmed, M.; Rabie, K.M. LSTM-based distributed conditional generative adversarial network for data-driven 5G-enabled maritime UAV communications. *IEEE Trans. Intell. Transp. Syst.* **2022**, *24*, 2431–2446. [\[CrossRef\]](#)
29. Hassan, F.; Yu, J.; Syed, Z.S.; Ahmed, N.; Al Reshan, M.S.; Shaikh, A. Achieving model explainability for intrusion detection in VANETs with LIME. *PeerJ Comput. Sci.* **2023**, *9*, e1440. [\[CrossRef\]](#)
30. Yao, Y.; Zhao, J.; Li, Z.; Cheng, X.; Wu, L. Jamming and Eavesdropping Defense Scheme Based on Deep Reinforcement Learning in Autonomous Vehicle Networks. *IEEE Trans. Inf. Forensics Secur.* **2023**, *18*, 1211–1224. [\[CrossRef\]](#)
31. Liang, H.; Burgess, L.; Liao, W.; Wang, Q.; Yu, W. 16 On Detecting Interest Flooding Attacks in Named Data Networking (NDN)-based IoT Searches. In *AI, Machine Learning and Deep Learning: A Security Perspective*; CRC: Boca Raton, FL, USA, 2023.
32. Hasan, T.; Malik, J.; Bibi, I.; Khan, W.U.; Al-Wesabi, F.N.; Dev, K.; Huang, G. Securing industrial internet of things against botnet attacks using hybrid deep learning approach. *IEEE Trans. Netw. Sci. Eng.* **2022**. [\[CrossRef\]](#)
33. Khan, W.U.; Nguyen, T.N.; Jameel, F.; Jamshed, M.A.; Pervaiz, H.; Javed, M.A.; Jäntti, R. Learning-based resource allocation for backscatter-aided vehicular networks. *IEEE Trans. Intell. Transp. Syst.* **2021**, *23*, 19676–19690. [\[CrossRef\]](#)
34. Li, B.; Zhou, X.; Ning, Z.; Guan, X.; Yiu, K.F.C. Dynamic event-triggered security control for networked control systems with cyber-attacks: A model predictive control approach. *Inf. Sci.* **2022**, *612*, 384–398. [\[CrossRef\]](#)
35. Safwat, M.; Elgammal, A.; AbdAllah, E.G.; Azer, M.A. Survey and taxonomy of information-centric vehicular networking security attacks. *Ad Hoc Netw.* **2022**, *124*, 102696. [\[CrossRef\]](#)
36. Zhang, J.; Peng, S.; Gao, Y.; Zhang, Z.; Hong, Q. APMSA: Adversarial perturbation against model stealing attacks. *IEEE Trans. Inf. Forensics Secur.* **2023**, *18*, 1667–1679. [\[CrossRef\]](#)
37. Han, S.; Ding, H.; Zhao, S.; Ren, S.; Wang, Z.; Lin, J.; Zhou, S. Practical and Robust Federated Learning With Highly Scalable Regression Training. *IEEE Trans. Neural Netw. Learn. Syst.* **2023**. [\[CrossRef\]](#)
38. Compagno, A.; Conti, M.; Gasti, P.; Tsudik, G. NDN interest flooding attacks and countermeasures. In Proceedings of the Annual Computer Security Applications Conference, Orlando, FL, USA, 3–7 December 2012.
39. Salah, H.; Wulfheide, J.; Strufe, T. Coordination supports security: A new defence mechanism against interest flooding in NDN. In Proceedings of the 2015 IEEE 40th Conference on Local Computer Networks (LCN), Clearwater Beach, FL, USA, 26–29 October 2015; pp. 73–81.
40. Salah, H.; Strufe, T. Evaluating and mitigating a collusive version of the interest flooding attack in NDN. In Proceedings of the 2016 IEEE Symposium on Computers and Communication (ISCC), Messina, Italy, 27–30 June 2016; pp. 938–945.
41. Wang, K.; Zhou, H.; Qin, Y.; Chen, J.; Zhang, H. Decoupling malicious interests from pending interest table to mitigate interest flooding attacks. In Proceedings of the 2013 IEEE Globecom Workshops (GC Wkshps), Atlanta, GA, USA, 9–13 December 2013; pp. 963–968.
42. Nguyen, T.; Mai, H.L.; Cogranne, R.; Doyen, G.; Mallouli, W.; Nguyen, L.; El Aoun, M.; De Oca, E.M.; Festor, O. Reliable detection of interest flooding attack in real deployment of named data networking. *IEEE Trans. Inf. Forensics Secur.* **2019**, *14*, 2470–2485. [\[CrossRef\]](#)
43. Nguyen, T.; Cogranne, R.; Doyen, G. An optimal statistical test for robust detection against interest flooding attacks in ccn. In Proceedings of the 2015 IFIP/IEEE International Symposium on Integrated Network Management (IM), Ottawa, ON, Canada, 11–15 May 2015; pp. 252–260.
44. Ding, K.; Liu, Y.; Cho, H.H.; Chao, H.C.; Shih, T.K. Cooperative detection and protection for interest flooding attacks in named data networking. *Int. J. Commun. Syst.* **2016**, *29*, 1968–1980. [\[CrossRef\]](#)

45. Xin, Y.; Li, Y.; Wang, W.; Li, W.; Chen, X. A novel interest flooding attacks detection and countermeasure scheme in NDN. In Proceedings of the 2016 IEEE Global Communications Conference (GLOBECOM), Washington, DC, USA, 4–8 December 2016; pp. 1–7.
46. Azmi, M.A.H.; Foozy, C.F.M.; Sukri, K.A.M.; Abdullah, N.A.; Hamid, I.R.A.; Amnur, H. Feature Selection Approach to Detect DDoS Attack Using Machine Learning Algorithms. *JOIV Int. J. Inform. Vis.* **2021**, *5*, 395–401. [[CrossRef](#)]
47. Subasri, I.; GSR, E.S.; Ramkumar, M. Machine Learning Based Feature Selection for DDoS Detection in Named Data Networking. In Proceedings of the 2022 4th International Conference on Advances in Computing, Communication Control and Networking (ICAC3N), Greater Noida, India, 16–17 December 2022; pp. 305–310.
48. Yaqub, M.A.; Ahmed, S.H.; Kim, D. A detailed simulation study of the push-based protocol for critical data dissemination in vehicular named data networks. In Proceedings of the 2019 International Conference on Networking and Network Applications (NaNA), Daegu, Republic of Korea, 10–13 October 2019; pp. 191–195.
49. Yaqub, M.A.; Ahmed, S.H.; Bouk, S.H.; Kim, D. Enabling critical content dissemination in vehicular named data networks. In Proceedings of the 2018 Conference on Research in Adaptive and Convergent Systems, Honolulu, HI, USA, 9–12 October 2018; pp. 94–99.
50. N.K. NDN-IFA-FeatureSelection. 2018. Available online: <https://github.com/nk10121989/NDN-IFA-FeatureSelection/> (accessed on 22 July 2023).
51. Afanasyev, A.; Moiseenko, I.; Zhang, L. *ndnSIM: NDN Simulator for NS-3*; Technical Report; University of California: Los Angeles, CA, USA, 2012; Volume 4, pp. 1–7.
52. Navada, A.; Ansari, A.N.; Patil, S.; Sonkamble, B.A. Overview of use of decision tree algorithms in machine learning. In Proceedings of the 2011 IEEE Control and System Graduate Research Colloquium, Shah Alam, Malaysia, 27–28 June 2011; pp. 37–42.
53. Kramer, O.; Kramer, O. K-nearest neighbors. In *Dimensionality Reduction with Unsupervised Nearest Neighbors*; Springer: Berlin/Heidelberg, Germany, 2013; pp. 13–23.
54. Liu, Y.; Wang, Y.; Zhang, J. New machine learning algorithm: Random forest. In Proceedings of the Information Computing and Applications: Third International Conference, ICICA 2012, Chengde, China, 14–16 September 2012; Proceedings 3; Springer: Berlin/Heidelberg, Germany, 2012; pp. 246–252.
55. Ontivero-Ortega, M.; Lage-Castellanos, A.; Valente, G.; Goebel, R.; Valdes-Sosa, M. Fast Gaussian Naïve Bayes for searchlight classification analysis. *Neuroimage* **2017**, *163*, 471–479. [[CrossRef](#)] [[PubMed](#)]
56. LaValley, M.P. Logistic regression. *Circulation* **2008**, *117*, 2395–2399. [[CrossRef](#)] [[PubMed](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.