

Trust-aware and Collaborative Routing Protocol for Automatic Fault Diagnosis for IoT Security

by

Sanaz GHORCHIAN

THESIS PRESENTED TO ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
IN PARTIAL FULFILLMENT FOR A MASTER'S DEGREE
WITH THESIS IN ELECTRICAL ENGINEERING
M.A.Sc.

MONTREAL, JULY 13, 2026

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC



Sanaz GHORCHIAN, 2026



This Creative Commons license allows readers to download this work and share it with others as long as the author is credited. The content of this work cannot be modified in any way or used commercially.

BOARD OF EXAMINERS

**THIS THESIS HAS BEEN EVALUATED
BY THE FOLLOWING BOARD OF EXAMINERS**

Mr. Michel Kadoch, Thesis Supervisor
Department of Electrical Engineering, École de technologie supérieure

Mr. Tony Wong, President of the Board of Examiners
Department of Systems Engineering, École de technologie supérieure

Mr. Armin Jabbarzadeh, Member of the Jury
Department of Systems Engineering, École de technologie supérieure

**THIS THESIS WAS PRESENTED AND DEFENDED
IN THE PRESENCE OF A BOARD OF EXAMINERS AND THE PUBLIC
ON JUIN 17, 2026
AT ÉCOLE DE TECHNOLOGIE SUPÉRIEURE**

Protocole de routage collaboratif basé sur la confiance pour le diagnostic automatique des défaillances dans les réseaux distribués appliqués à la sécurité de l'Internet des objets

Sanaz GHORCHIAN

RÉSUMÉ

Cette thèse examine la sécurité des réseaux IoT basés sur RPL en soutenant que le routage et la détection d'intrusion ne doivent pas être considérés comme des problèmes distincts dans des environnements contraints. Dans les réseaux à faible puissance et à pertes élevées, les attaques peuvent se manifester à travers des comportements de routage apparemment légitimes, ce qui rend les hypothèses de détection classiques moins fiables. Pour cette raison, l'étude cherche à déterminer si une stratégie de routage fondée sur la confiance, mise en œuvre sous la forme de TACR-HOF, peut générer des traces comportementales plus utiles pour la détection d'intrusion que celles produites par la fonction objective conventionnelle MRHOF. L'idée centrale est que la politique de routage n'influence pas seulement l'acheminement des paquets, mais qu'elle façonne aussi la qualité des preuves comportementales disponibles pour le système de détection d'intrusion.

Pour répondre à cette question, la thèse développe un cadre comparatif de simulation à l'aide de Contiki-NG et du simulateur Cooja. La même famille de six topologies est testée sous deux schémas de routage, MRHOF et TACR-HOF, tandis que quatre attaques spécifiques à RPL sont injectées : l'attaque du numéro de version, l'attaque de diminution du rang, l'attaque d'inondation DIS et l'attaque de transfert sélectif. Au lieu de s'appuyer sur une capture complète des paquets, l'étude utilise des résumés comportementaux collectés à partir des nœuds et de leurs voisins, puis transformés en jeux de données structurés pour une détection d'intrusion fondée sur l'apprentissage automatique. Cette conception permet de comparer non seulement les performances de routage, mais aussi la manière dont chaque schéma de routage influence la visibilité, l'équilibre et la séparabilité des comportements malveillants dans les données produites.

Les résultats montrent que TACR-HOF fournit une base plus solide pour une détection d'intrusion comportementale que MRHOF. Il génère une empreinte comportementale malveillante plus riche, améliore les performances des classifieurs pour tous les modèles testés et réduit de façon significative l'erreur la plus critique du système de détection, à savoir la classification erronée d'un comportement malveillant comme normal. La détection par type d'attaque s'améliore également pour les quatre classes d'attaques étudiées, en particulier pour les manipulations du plan de contrôle comme les attaques sur la version et le rang. Toutefois, au niveau du routage, les résultats demeurent nuancés : TACR-HOF améliore plusieurs mesures au niveau des scénarios, notamment la livraison des paquets et la réduction des transmissions manquées dans de nombreux cas, mais il ne surpasse pas MRHOF de manière uniforme dans toutes les topologies ni pour toutes les métriques de routage.

Dans l'ensemble, la thèse aboutit à une conclusion équilibrée. Elle ne prétend pas que TACR-HOF soit universellement supérieur pour tous les objectifs de routage, mais elle démontre qu'il

est plus efficace comme approche de routage orientée vers la sécurité. Sa contribution principale consiste à montrer que la modification de la fonction objective de routage peut améliorer la qualité des preuves comportementales fournies à un système de détection d'intrusion, rendant ainsi les attaques plus visibles, plus distinctes et plus faciles à détecter dans les réseaux IoT contraints.

Mots-clés: Internet des objets (IoT), réseaux à faible puissance et à pertes élevées (LLNs), RPL, système de détection d'intrusion (IDS), routage fondé sur la confiance, TACR-HOF, MRHOF, détection comportementale, apprentissage automatique, Contiki-NG, simulateur Cooja, attaques RPL, attaque du numéro de version (VNA), attaque de diminution du rang (DRA), attaque d'inondation DIS (DISA), attaque de transfert sélectif (SFA).

Trust-aware and Collaborative Routing Protocol for Automatic Fault Diagnosis for IoT Security

Sanaz GHORCHIAN

ABSTRACT

This thesis investigates security in RPL-based Internet of Things (IoT) networks by examining the relationship between routing behavior and intrusion detection in constrained, low-power, and lossy environments. In such networks, attacks can be expressed through seemingly legitimate protocol operations, making conventional intrusion detection assumptions less reliable. The study therefore explores whether a trust-aware routing strategy, implemented as TACR-HOF, can generate behavioural traces that are more informative for intrusion detection than those produced by the conventional Minimum Rank with Hysteresis Objective Function (MRHOF).

To address this problem, a comparative simulation framework was developed using Contiki-NG and the Cooja simulator. The same family of six topologies was evaluated under MRHOF and TACR-HOF, while four RPL-specific attacks were introduced: Version Number Attack, Decrease Rank Attack, DIS Flooding Attack, and Selective Forwarding Attack. Rather than relying on full packet capture, the study employed behaviour-based summaries collected from nodes and their neighbours, which were transformed into structured datasets for machine-learning-based intrusion detection. This approach enabled a comparative analysis of both routing behaviour and the quality of the IDS-ready data produced under each routing scheme.

The results show that TACR-HOF provides a stronger foundation for behaviour-based intrusion detection than MRHOF. It produces a richer malicious-behaviour footprint, improves classification performance across all evaluated models, and significantly reduces the most critical IDS error, namely the misclassification of malicious behaviour as normal. Attack-wise recall also improves across all four attack classes, particularly for control-plane attacks involving rank and version manipulation. At the routing level, however, the results remain nuanced: TACR-HOF improves several scenario-level metrics, including packet delivery and missed transmissions in many cases, but it does not outperform MRHOF uniformly in every topology or routing metric.

Overall, the thesis concludes that TACR-HOF should not be considered universally superior for all routing purposes, but it is clearly more effective as a security-oriented routing approach. Its main contribution lies in demonstrating that modifying the routing objective function can improve the quality of behavioural evidence available to an intrusion detection system, making attacks more visible, more distinguishable, and easier to detect in constrained IoT networks.

Keywords: Internet of Things (IoT), RPL, Low-Power and Lossy Networks (LLNs), Intrusion Detection System (IDS), trust-aware routing, collaborative routing, TACR-HOF, MRHOF, behaviour-based intrusion detection, machine learning, Contiki-NG, Cooja simulator, RPL attacks, network security.

TABLE OF CONTENTS

	Page
INTRODUCTION	1
CHAPTER 1 BACKGROUND STUDIES	9
1.1 IoT systems and constrained networking	9
1.2 RPL as the routing foundation for LLNs	10
1.3 Objective functions and parent selection logic	10
1.4 Security threats in RPL-based IoT networks	10
1.5 Intrusion detection perspectives for constrained networks	11
1.6 Trust, collaboration, and security-aware routing	11
1.7 Behavior-based evidence and feature-oriented security analysis	12
1.8 Positioning of the present study	12
1.9 Synthesis and transition	13
1.9.1 Comparative stack of the secured RPL-based IoT environment	13
1.9.2 Comparative background concepts	14
1.9.3 Attack taxonomy in the study context	15
CHAPTER 2 CONCEPTUAL FRAMEWORK AND ANALYTICAL MODEL	17
2.1 Internet of Things environments and low-power and lossy networking	17
2.2 RPL architecture and the logic of DODAG formation	19
2.3 MRHOF routing principles: path cost and hysteresis	21
2.4 Security exposure of RPL and the attack taxonomy used in the thesis	22
2.4.1 Version Number Attack (VNA)	23
2.4.2 Decrease Rank Attack (DRA)	23
2.4.3 DIS Flooding Attack (DISA)	23
2.4.4 Selective Forwarding Attack (SFA)	24
2.5 Intrusion detection theory in constrained IoT networks	25
2.6 Behavioural feature theory for RPL intrusion detection	27
2.7 Classifier and evaluation principles for multiclass RPL intrusion detection	29
2.8 Conceptual integration of trust evidence into parent selection	30
2.9 Positioning TACR-HOF Relative to Existing Secure and Trust-Aware RPL Proposals	35
2.10 Research gap and the conceptual framework of the thesis	37
2.11 Chapter summary	38
CHAPTER 3 METHODOLOGY	40
3.1 Overall research design and workflow	41
3.2 Simulation platform, execution environment, and control variables	42
3.3 Topology design and scenario coverage	44
3.4 Adversarial model and randomized attack scheduling	45
3.5 Data acquisition, behavioural summaries, and feature engineering	47

3.6	Complete feature schema used in the comparative CSV datasets	49
3.7	Routing configurations: MRHOF baseline versus TACR-HOF	51
3.7.1	Trust penalties logic	54
3.8	Reproducibility, internal validity, and methodological limits	55
3.9	Chapter summary	56
CHAPTER 4 SIMULATIONS AND RESULTS		58
4.1	Comparative experiment design and scenario coverage	59
4.2	Dataset composition and class balance	61
4.3	Classifier-level performance comparison	63
4.4	Attack-wise analysis using the Random Forest classifier	65
4.5	Routing-level observations from the raw simulation logs	69
4.6	Synthesis of findings	71
4.7	Chapter conclusion	72
CHAPTER 5 CONCLUSION AND FUTURE WORK		74
5.1	Final answer to the research problem	75
5.2	Quantitative synthesis of the principal findings	76
5.3	Resolution of the research objectives and questions	78
5.4	Scientific, methodological, and practical contributions	80
5.5	Limitations and validity boundaries	81
5.6	Broader implications of the findings	83
5.7	Open work and post-thesis enhancements	84
5.7.1	Protocol and trust-model enhancements	84
5.7.2	IDS and data-analysis enhancements	85
5.7.3	Experimental and validation enhancements	86
5.7.4	Deployment-oriented enhancements	86
5.8	Final chapter conclusion	89
BIBLIOGRAPHY		91
LIST OF REFERENCES		95
APPENDIX I ROUTING OBJECTIVE FUNCTION SOURCE-CODE EXCERPTS		
95		

LIST OF TABLES

	Page
Table 1.1	Core concepts that motivate the integrated routing-and-detection perspective of the thesis. Author's synthesis based on (Bormann et al. 2014; Winter et al. 2012; Gyamfi & Jurcut 2022; Djedjig et al. 2020) 14
Table 1.2	Representative attack categories motivating the comparative security analysis. Author's synthesis based on (Tsao et al. 2015; Almusaylim et al. 2020; Medjek et al. 2021) 15
Table 2.1	Structural differences between conventional IP networks and low-power and lossy IoT networks. Author's synthesis based on (Bormann et al. 2014; Winter et al. 2012; Elrawy et al. 2018; Tsao et al. 2015) 19
Table 2.2	Core RPL concepts and their operational meaning. Author's synthesis based on (Winter et al. 2012; Tsao et al. 2015) 21
Table 2.3	Baseline logic of MRHOF and its security implications. Author's synthesis based on (Gnawali & Levis 2012; Tsao et al. 2015; Muzammal et al. 2022) 22
Table 2.4	Operational comparison of the four RPL attack classes. Author's synthesis based on (Tsao et al. 2015, Almusaylim et al. 2020, Medjek et al. 2021; Mayzaud et al. 2016 25
Table 2.5	IDS paradigms relevant to RPL-based IoT security. Author's synthesis based on (da Costa et al. 2019; Rahman et al. 2025) 27
Table 2.6	Behavioural feature families and their theoretical meaning. Author's synthesis based on (da Costa et al. 2019; Rahman et al. 2025; HUNSR n.d.) 29
Table 2.7	Evaluation metrics used to interpret multiclass intrusion-detection performance. Author's synthesis based on (da Costa et al. 2019; Rahman et al. 2025; Opitz 2024) 30
Table 2.8	Behavioural evidence used by the TACR-HOF concept and its routing consequence. Author's synthesis based on (HUNSR n.d.; Djedjig et al. 2020; Muzammal et al. 2022; Gnawali & Levis 2012) 34
Table 2.9	Comparative positioning of TACR-HOF relative to existing RPL security research. Author's synthesis based on (Winter et al. 2012; Gnawali & Levis 2012; Tsao et al. 2015; Djedjig et al. 2020;

	Muzammal et al. 2022; da Costa et al. 2019; Elrawy et al. 2018; Gyamfi and Jurcut 2022; Rahman et al. 2025; HUNSR n.d.) 36
Table 2.10	Research gaps identified in the literature and the conceptual response adopted by this thesis. Author's synthesis based on (da Costa et al. 2019; Rahman et al. 2025; Muzammal et al. 2022; HUNSR n.d.) 38
Table 3.1	Core simulation and comparative-experiment parameters (HUNSR, n.d.) 43
Table 3.2	Inventory of topology files extracted directly from the uploaded .csc scenarios. Author's extraction from the Cooja .csc topology scenario files used in the experiment 45
Table 3.3	Attack classes and their operational manifestation in the simulation. Author's synthesis based on (Tsao et al. 2015; Verma & Ranga 2020; HUNSR n.d.) 47
Table 3.4	Full feature inventory of the generated MRHOF/TACR datasets (HUNSR, n.d.) 49
Table 3.5	Methodological comparison between the baseline MRHOF and the TACR-HOF routing configuration. Author's synthesis based on (Gnawali & Levis 2012; da Costa et al. 2019; Rahman et al. 2025; Muzammal et al. 2022; HUNSR n.d.) 52
Table 4.1	Topology-level summary of the comparative simulation campaign 61
Table 4.2	Dataset composition and class balance 62
Table 4.3	Overall IDS performance on the MRHOF and TACR datasets 64
Table 4.4	Random Forest class-wise precision, recall, and F1-score 65
Table 4.5	Routing-level log summary 70
Table 5.1	Quantitative synthesis of the principal results 77
Table 5.2	Closure of the thesis objectives and research questions 79
Table 5.3	Principal limitations of the study and their implications 82
Table 5.4	Recommended post-thesis roadmap for enhancement and validation 88

LIST OF FIGURES

	Page
Figure 1.1 Conceptual stack linking IoT constraints, RPL routing, trust-aware decisions, and behavior-based intrusion detection	13
Figure 2.1 Conceptual stack of the secured RPL-based IoT environment adopted in this thesis	18
Figure 2.2 RPL DODAG structure and control-message flow used to explain upward and downward routing behaviour. Author's illustration based on (Winter et al. 2012)	20
Figure 2.3 Attack-to-symptom chain for the four RPL threats analysed in this thesis. Author's illustration based on (Tsao et al. 2015, Almusaylim et al. 2020, Medjek et al. 2021, and Mayzaud et al. 2016)	24
Figure 2.4 Contrasting packet-capture and behaviour-summary IDS collection paradigms in constrained RPL networks	26
Figure 2.5 Conceptual grouping of the behavioural schema later operationalized in the comparative datasets	28
Figure 2.6 Conceptual logic of TACR-HOF: MRHOF path cost is retained and adjusted by behaviour-derived trust penalties	31
Figure 3.1 End-to-end methodological pipeline used to move from Contiki-NG configuration to comparative IDS-ready datasets	41
Figure 3.2 The six fixed Cooja topologies used in the comparative campaign. (Contiki-NG/Cooja GUI)	44
Figure 4.1 Topology layouts used in the comparative experiment. (Contiki-NG/Cooja GUI)	60
Figure 4.2 Class distribution of the generated MRHOF and TACR datasets	62
Figure 4.3 Accuracy comparison for the three evaluated classifiers	64
Figure 4.4 Attack-wise recall comparison for the Random Forest classifier	66
Figure 4.5 Confusion matrix of the Random Forest classifier on the MRHOF dataset	67

Figure 4.6	Confusion matrix of the Random Forest classifier on the TACR dataset	68
Figure 4.7	Packet delivery ratio by topology for MRHOF and TACR	69
Figure 5.1	Integrated synthesis of the thesis conclusions	76
Figure 5.2	Post-thesis enhancement roadmap organized around protocol, IDS, validation, and deployment workstreams	84

LIST OF ABBREVIATIONS

6LoWPAN	IPv6 over Low-Power Wireless Personal Area Networks
CSV	Comma-Separated Values
DAO	Destination Advertisement Object
DIO	DODAG Information Object
DIS	DODAG Information Solicitation
DISA	DIS Flooding Attack
DODAG	Destination-Oriented Directed Acyclic Graph
DRA	Decrease Rank Attack
ETX	Expected Transmission Count
HMAC	Hash-Based Message Authentication Code
IDS	Intrusion Detection System
IETF	Internet Engineering Task Force
IoT	Internet of Things
IP	Internet Protocol
IPv6	Internet Protocol Version 6
KNN	K-Nearest Neighbours
LLN	Low-Power and Lossy Network
LLNs	Low-Power and Lossy Networks
LSM-RPL	Lightweight Security Mode for RPL

MAC	Medium Access Control
MRHOF	Minimum Rank with Hysteresis Objective Function
PDR	Packet Delivery Ratio
RPL	Routing Protocol for Low-Power and Lossy Networks
SFA	Selective Forwarding Attack
TACR-HOF	Trust-Aware Cooperative Routing with Hysteresis Objective Function
TSCH	Time Slotted Channel Hopping
UDGM	Unit Disk Graph Medium
UDP	User Datagram Protocol
VNA	Version Number Attack

INTRODUCTION

Overview

The Internet of Things (IoT) has transformed networked sensing and control from a specialized engineering capability into a pervasive digital infrastructure supporting homes, factories, healthcare systems, transportation corridors, energy grids, and environmental monitoring platforms. In these environments, large numbers of constrained devices cooperate through low-power and lossy wireless links, frequently under strict limits on energy, memory, processing capacity, and communication bandwidth. These constraints create an architectural condition in which routing efficiency, network stability, and security cannot be treated as independent design targets. A routing protocol that optimizes path cost while ignoring adversarial behavior may preserve nominal connectivity yet simultaneously degrade the quality of observable evidence needed by an intrusion detection system. Conversely, a security mechanism that consumes excessive communication or computation resources may compromise the operational sustainability of the network it seeks to protect. This thesis is situated precisely at that intersection. (Atzori et al., 2010; Bormann et al., 2014; Djedjig et al., 2020; Gyamfi & Jurcut, 2022).

Components

The problem domain addressed in this work is organized around several interacting components. The first component is the constrained IoT node, which senses, processes, and forwards data while operating under severe resource limitations. The second component is the routing layer, where the Routing Protocol for Low-Power and Lossy Networks (RPL) organizes nodes into a destination-oriented directed acyclic graph and selects forwarding parents according to an objective function. The third component is the communication behavior layer, which includes both the control-plane exchanges that maintain routing state and the data-plane transmissions that carry sensed information toward the sink. The fourth component is the attack surface, encompassing malicious manipulations that distort rank, trigger unnecessary control traffic,

or disrupt packet forwarding. The fifth component is the intrusion detection perspective, which interprets behavioral deviations as evidence of abnormal or malicious activity. The final component is the trust-aware routing perspective, which uses observations about neighbor behavior to influence path selection and thereby reshape the patterns that an IDS later analyzes. (Winter et al., 2012; Friehtat et al., 2024; Gyamfi & Jurcut, 2022; Djedjig et al., 2020).

Features

Several features define the technological and analytical setting of the study. IoT routing in low-power networks is distributed rather than centrally orchestrated, meaning local decisions accumulate into global traffic patterns. RPL provides loop avoidance and structured upward forwarding, making it suitable for many constrained deployments, but its effectiveness depends strongly on the choice of objective function. Behavior-based intrusion detection emphasizes summary indicators extracted from node activity rather than deep packet inspection, which is advantageous in constrained environments where full traffic capture is impractical. Trust-aware routing extends conventional path optimization by penalizing suspicious or unstable forwarding candidates instead of relying solely on link-quality-centric metrics. The present thesis treats these features not as isolated properties, but as elements of an integrated system in which routing decisions influence security observability and security-relevant behavior in turn influences routing performance. (Winter et al., 2012; Gyamfi & Jurcut, 2022; Djedjig et al., 2020).

Applications

The relevance of this research extends across application domains where IoT devices must remain dependable even under hostile or unstable conditions. In industrial automation, routing disruptions or stealthy forwarding anomalies can distort process monitoring and affect operational safety. In healthcare monitoring, unreliable or manipulated forwarding may delay critical physiological readings. In smart-city infrastructure, control and sensing devices distributed

across large spaces must operate under heterogeneous interference and varying trust assumptions. In environmental and agricultural monitoring, remote deployments often lack continuous human oversight, making early detection of routing anomalies especially valuable. In all such settings, there is practical interest in mechanisms that improve not only packet delivery and topology maintenance, but also the detectability of behavior associated with malicious activity or severe protocol misuse. (Atzori et al., 2010; Djedjig et al., 2020; da Costa et al. (2019)).

Challenges

Despite the maturity of IoT networking research, several challenges remain unresolved. Low-power and lossy links produce fluctuating connectivity, so routing behavior may vary even in the absence of malicious activity. Constrained nodes cannot support heavy cryptographic, monitoring, or analytics burdens without exhausting scarce resources. RPL control traffic is essential for topology maintenance, yet the same control structures can be exploited by adversaries through false version increments, rank manipulation, or excessive solicitation messages. Some attacks operate on the control plane, while others degrade the data plane more subtly, complicating the task of behavioral discrimination. Another challenge lies in evaluation methodology: routing mechanisms and intrusion detection models are often studied separately, making it difficult to determine whether a change in routing policy improves or worsens the quality of evidence available to downstream security analytics. This separation leaves a methodological gap that limits defensible conclusions about cross-layer security effectiveness. (Bormann et al., 2014; Winter et al., 2012; Frieht et al., 2024; Alsukayti & Alreshoodi, 2023; Gyamfi & Jurcut, 2022).

Problem Statement

Existing IoT security studies frequently assume that routing and intrusion detection can be optimized independently. However, in constrained RPL-based networks, the routing objective

function influences which parents are selected, how control and data traffic are distributed, how instability propagates through the topology, and which behavioral signatures become visible in collected logs or derived datasets. Standard objective functions such as Minimum Rank with Hysteresis Objective Function (MRHOF) emphasize path quality and hysteresis, but they do not explicitly account for trust-related evidence when selecting forwarding paths. As a result, they may tolerate behaviors that preserve short-term path cost while obscuring or diluting indicators of malicious activity. The central problem addressed in this thesis is therefore whether the choice of routing objective function materially changes the behavioral evidence available for intrusion detection, and whether a trust-aware cooperative routing strategy can improve the discriminability of attack-related behavior without resorting to unrealistic assumptions about resource abundance or global knowledge. (Winter et al., 2012; Gnawali & Levis, 2012; Djedjig et al., 2020; Gyamfi & Jurcut, 2022).

Research Objectives

The primary objective of the study is to investigate the relationship between routing strategy and intrusion detection quality in RPL-based IoT networks. More specifically, the work aims to compare a baseline routing objective function with a trust-aware alternative, to determine whether trust-sensitive parent selection can generate network behavior that is more informative for attack detection. A second objective is to analyze the response of the network under multiple attack classes that affect both routing control traffic and data forwarding behavior. A third objective is to frame the evaluation in a way that is scientifically defensible, linking routing-level observations to classification-level outcomes rather than presenting security and networking results as disconnected findings. A final objective is to establish an integrated methodological foundation that can support later enhancements in protocol design, feature engineering, and security analytics.

Research Questions

To guide the investigation, this study seeks to answer the following research questions:

- Does changing the routing objective function from a conventional metric-oriented scheme to a trust-aware cooperative scheme alter the quality of behavior-level evidence available to an intrusion detection system?
- Are improvements in intrusion detection associated with specific changes in routing and forwarding behavior?
- Does the routing-level cost of trust-aware parent selection remain acceptable within the resource constraints of IoT deployments?

Scope and Limitations

This thesis focuses on low-power and lossy IoT networks that use RPL as the routing foundation and behavior-based analysis as the security evaluation perspective. The work is concerned with comparative analysis between a baseline RPL objective function and a trust-aware cooperative alternative. The study concentrates on attacks that alter routing behavior or forwarding outcomes in ways that can be observed at the behavioral level. The scope does not extend to all possible IoT protocols, all wireless media, or all categories of cyber-physical attacks. It also does not claim that routing adaptation alone is sufficient for complete network security. Rather, the work is bounded to examining how routing policy shapes IDS-relevant evidence under constrained-network assumptions. Limitations therefore include the use of controlled experimental settings, the reliance on selected attack models, and the fact that conclusions are framed with respect to the studied architecture rather than all IoT ecosystems without qualification.

Significance of the Study

The significance of the study lies in its cross-layer viewpoint. Instead of asking only whether a routing protocol forwards packets efficiently, or only whether an IDS classifies attacks accurately, the thesis asks how routing design affects the learnability of malicious behavior. This is

an important distinction because many IoT defenses fail not from the absence of detection algorithms, but from the poor quality of the evidence those algorithms receive. By treating routing as a factor that structures security observability, the thesis contributes a more integrated perspective to IoT security research. It also offers practical value: if trust-aware routing can improve the separability of normal and malicious behavior while preserving acceptable network operation, then routing policy becomes a meaningful design lever for IDS-oriented system hardening (Djedjig et al., 2020; Gyamfi & Jurcut, 2022).

Literature Review Overview

The literature reviewed in this thesis spans four closely related areas: low-power IoT networking, RPL routing and objective functions, intrusion detection in constrained networks, and trust-aware or collaborative security mechanisms. Prior studies establish the importance of RPL as a routing standard for constrained environments, but they also identify its susceptibility to manipulation at the control plane. Research on intrusion detection provides a variety of packet-based, behavior-based, specification-based, and machine-learning-driven approaches, yet many of these assume that the underlying routing behavior is fixed. Studies on trust-aware routing demonstrate that trust metrics can influence path reliability and resilience, but they do not always evaluate the downstream effect on IDS evidence quality using a common comparative framework. The literature therefore motivates a synthesis in which routing, trust assessment, and behavior-based detection are examined together rather than as separate silos (Winter et al., 2012; Tsao et al., 2015; Djedjig et al., 2020; Gyamfi & Jurcut, 2022).

Methodology Overview

The thesis adopts a comparative experimental methodology grounded in repeatable IoT network evaluation. At a high level, the study employs a simulation-based research design to examine how two routing objective functions influence network behavior under normal and adversarial

conditions. The networking environment is implemented using Contiki-NG, while network scenarios are executed through the Cooja simulator. Behavior traces and protocol logs are then processed through Python-based analysis scripts, where security-relevant indicators are organized into an IDS-oriented dataset and evaluated using supervised learning models. At the introductory level, it is sufficient to note that the methodology combines controlled network experimentation, structured data extraction, and comparative analytical modeling. The detailed configuration of topologies, attack timing, logging procedures, feature derivation, and model training is intentionally reserved for the methodology chapter, in accordance with sound thesis structure (HUNSR, n.d.; Oikonomou et al., 2022; Osterlind et al., 2006; da Costa et al., 2019)..

Thesis Structure

The remainder of the thesis is organized to move from conceptual foundations to implementation and evidence. The present introductory section frames the problem, objectives, and research logic of the study. Chapter 1 then presents the background studies needed to understand constrained IoT networking, RPL routing, attack behavior, and intrusion detection perspectives. Chapter 2 develops the conceptual framework and analytical model that connect RPL routing, behaviour-based intrusion detection, and trust-aware parent selection. Chapter 3 explains the methodology, tools, experimental workflow, data engineering pipeline, and analytical procedures used in the research. Chapter 4 presents the results and findings, interpreting both routing-level outcomes and intrusion-detection performance in a defensible comparative manner. Chapter 5 concludes the thesis by summarizing the contributions, discussing limitations, and identifying open work and enhancement directions.

CHAPTER 1

BACKGROUND STUDIES

This chapter reviews the literature needed to position the thesis at the intersection of RPL routing, trust-aware parent selection, and behaviour-based intrusion detection in constrained Internet of Things networks. RPL is widely used for multi-hop communication in Low-Power and Lossy Networks, but it remains vulnerable to topology instability, unreliable forwarding, and routing attacks that can affect both network performance and the behavioural evidence available to an intrusion detection system (Bormann et al., 2014; Winter et al., 2012; Tsao et al., 2015).

Previous studies have proposed improved routing metrics, trust-based parent-selection mechanisms, and intrusion-detection techniques. However, these research directions are often evaluated separately or under limited attack scenarios. This thesis therefore examines whether replacing MRHOF with TACR-HOF produces behavioural patterns that are more informative for attack detection. This chapter establishes the theoretical and research background required to position that comparison within the existing literature, while the experimental design is presented in Chapter 3 (Ghaleb et al., 2019; Djedjig et al., 2020; Gyamfi & Jurcut, 2022; Friehtat et al., 2024).

1.1 IoT systems and constrained networking

IoT systems consist of distributed, resource-constrained devices operating over variable wireless links. Unlike conventional networks, these systems prioritize energy efficiency, limited memory and processing use, and unattended operation rather than high throughput. In Low-Power and Lossy Networks, packet loss, interference, asymmetric links, and topology changes are normal operating conditions. Consequently, routing decisions affect communication reliability, energy consumption, congestion, and the visibility of anomalous behaviour, making routing central to both network operation and security analysis (Atzori et al., 2010; Bormann et al., 2014; Winter et al., 2012).

1.2 RPL as the routing foundation for LLNs

RPL is designed for IPv6-based Low-Power and Lossy Networks and organizes nodes into a Destination-Oriented Directed Acyclic Graph (DODAG) rooted at one or more collection points. Rank represents each node's relative position in the topology and supports loop-free upward forwarding. Parent selection is controlled by an objective function, allowing RPL to adapt to different performance goals. This flexibility is useful, but it also means that network behaviour and security depend strongly on the assumptions embedded in the selected objective function (Winter et al., 2012).

1.3 Objective functions and parent selection logic

An objective function in RPL determines how nodes interpret metrics such as rank, expected transmission count, link quality, or hysteresis when selecting preferred parents. In practical terms, the objective function is the decision rule that converts local observations into routing structure. A conventional metric-oriented objective function, such as MRHOF, seeks stable and efficient paths by minimizing cumulative cost while reducing unnecessary parent oscillations. This design is appropriate for many benign settings, because frequent parent switching can destabilize a constrained network. However, stability-oriented logic may also preserve routes through nodes whose behavior is suspicious but whose short-term path metrics remain superficially acceptable. Trust-aware objective functions respond to this limitation by enriching the selection logic with behavioral evidence, thereby making parent choice sensitive not only to link quality but also to observed forwarding reliability, consistency, or cooperativeness. The theoretical importance of this distinction is that routing becomes behavior-aware rather than exclusively metric-aware (Winter et al., 2012; Gnawali & Levis, 2012; Djedjig et al., 2020).

1.4 Security threats in RPL-based IoT networks

RPL-based IoT networks are exposed to both control-plane and data-plane attacks. Control-plane attacks manipulate routing state through false version numbers, deceptive rank advertisements, or

excessive solicitation messages, while data-plane attacks disrupt packet forwarding by dropping or selectively discarding traffic. The four attacks examined in this thesis represent distinct threat classes: Version Number Attack affects topology consistency, Decrease Rank Attack corrupts parent selection, DIS Flooding increases control overhead, and Selective Forwarding reduces delivery while maintaining partially normal behaviour. Together, they provide a balanced basis for evaluating how routing policy influences observable attack patterns (Tsao et al., 2015; Frieht et al., 2024; Alsukayti & Alreshoodi, 2023).

1.5 Intrusion detection perspectives for constrained networks

Intrusion detection in IoT can be approached through signature matching, specifications of legitimate behavior, anomaly detection, or supervised classification. Each approach offers advantages and trade-offs. Signature-based systems can be effective against known attack patterns but may generalize poorly to evolving behavior. Specification-based methods can be lightweight yet require strong formal assumptions about correct protocol operation. Anomaly detection can surface previously unseen deviations, but it may suffer from false positives in naturally noisy environments. Supervised learning can achieve strong classification performance when meaningful labeled data are available, but its validity depends on the relevance and quality of the feature space. In constrained networks, behavior-based intrusion detection is especially attractive because it can be built from periodic summaries and lightweight indicators rather than continuous payload inspection. This shift from packet content to behavioral evidence is central to the present work, since the thesis asks how routing policy affects the shape of that evidence. (Gyamfi & Jurcut, 2022; da Costa et al., 2019).

1.6 Trust, collaboration, and security-aware routing

Trust-aware routing extends conventional networking logic by introducing a behavioral judgment into forwarding decisions. Rather than treating all candidate parents as equivalent aside from path cost, trust-aware schemes attempt to estimate whether a neighbor behaves in a manner consistent with reliable participation in the routing process. Trust can be inferred from direct

observations, indirect recommendations, forwarding success, control-message consistency, historical stability, or composite evidence derived from several such factors. Collaborative or cooperative trust concepts further assume that security is not solely an end-host concern but a network-structuring concern: the topology itself should adapt away from suspicious participants when evidence is sufficient. From a theoretical perspective, this approach aligns routing with resilience objectives. From an IDS perspective, it may also sharpen the contrast between normal and malicious behavior by isolating unreliable nodes, changing forwarding patterns, and reducing the ambiguity that often exists when attacks are partly masked by conventional metric optimization. (Djedjig et al., 2020; Muzammal et al., 2022).

1.7 Behavior-based evidence and feature-oriented security analysis

Behavior-based evidence summarizes how nodes and their neighbours operate over time through indicators such as control-message counts, forwarding activity, routing inconsistencies, and traffic ratios. This approach is well suited to constrained IoT networks because it avoids continuous packet inspection while still producing structured features for machine-learning analysis. It also supports cross-layer interpretation by linking changes in the feature space to routing instability, control-plane manipulation, or forwarding anomalies. In this thesis, these features form the bridge between routing behaviour and intrusion-detection performance (Gyamfi & Jurcut, 2022; Garcia Ribera et al., 2022).

1.8 Positioning of the present study

This study lies at the intersection of RPL routing, trust-aware networking, and IoT intrusion detection. Although these areas have been studied individually and in partial combination, fewer works examine whether changing the routing objective function improves the behavioural evidence available to an IDS under a common experimental framework. This thesis therefore evaluates whether TACR-HOF produces attack-related behaviour that is more visible, separable, and useful for downstream classification than behaviour generated under MRHOF (Djedjig et al., 2020; Gyamfi & Jurcut, 2022).

1.9 Synthesis and transition

The reviewed literature shows that RPL routing, behaviour-based intrusion detection, and trust-aware parent selection are closely connected in constrained IoT networks. RPL provides the routing structure, the objective function determines parent-selection behaviour, and the resulting traffic patterns shape the evidence available to the IDS. Trust-aware routing extends this process by incorporating behavioural reliability into route selection. The following figure and tables summarize these relationships and the attack categories considered in the thesis, while Chapter 2 develops the corresponding conceptual and analytical framework.

1.9.1 Comparative stack of the secured RPL-based IoT environment

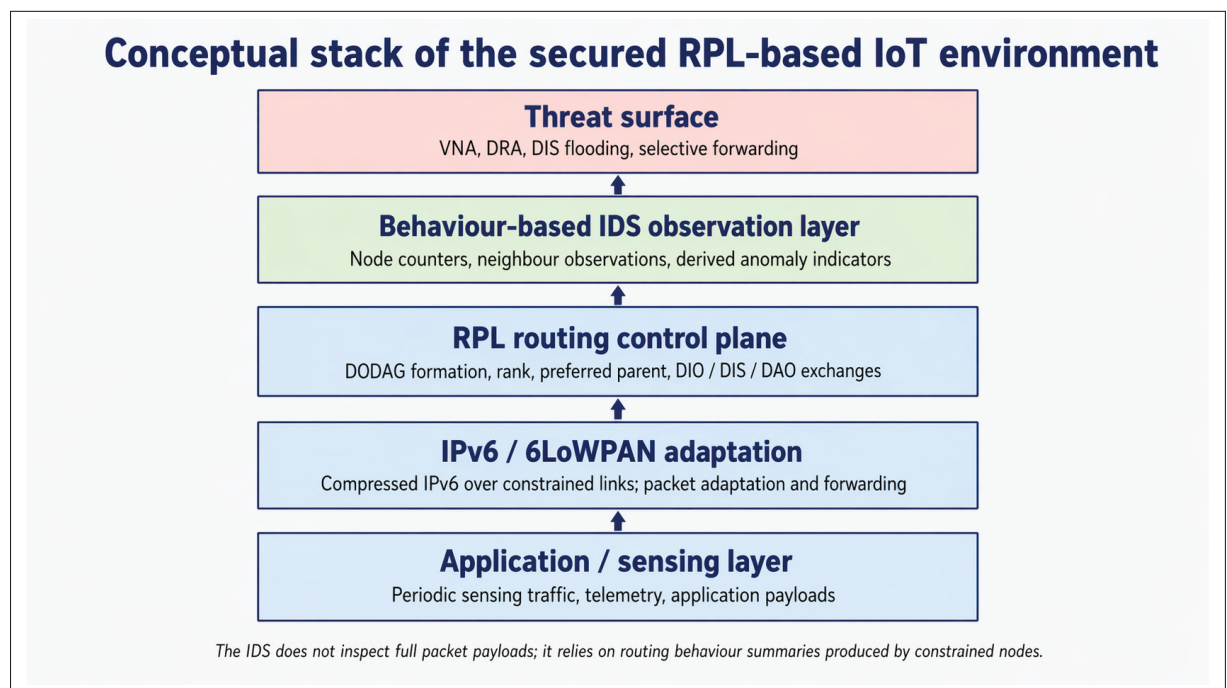


Figure 1.1 Conceptual stack linking IoT constraints, RPL routing, trust-aware decisions, and behavior-based intrusion detection

1.9.2 Comparative background concepts

Table 1.1 Core concepts that motivate the integrated routing-and-detection perspective of the thesis. Author’s synthesis based on (Bormann et al. 2014; Winter et al. 2012; Gyamfi & Jurcut 2022; Djedjig et al. 2020)

Concept	Primary Role	Why it Matters Here	Typical Risk if Ignored
LLN constraints	Define limits on energy, memory, computation, and link reliability	They shape what kinds of routing and security methods are realistically deployable	Overly heavy defenses or analyses become impractical
RPL objective function	Determines parent selection and routing structure	It governs how local metrics become global traffic behavior	Suspicious paths may remain preferred if only path cost is optimized
Behavior-based IDS	Converts protocol activity into structured security evidence	It allows attack detection without continuous payload inspection	Poor feature quality yields weak or misleading classifications
Trust-aware routing	Injects behavioral judgment into routing decisions	It may improve resilience and sharpen security observability	Security remains reactive instead of influencing topology formation

1.9.3 Attack taxonomy in the study context

Table 1.2 Representative attack categories motivating the comparative security analysis. Author’s synthesis based on (Tsao et al. 2015; Almusaylim et al. 2020; Medjek et al. 2021)

Attack Type	Primary Effect on Network Behavior	Analytical Relevance
Version-number manipulation	Disturbs topology consistency and can trigger unnecessary reorganization	Useful for studying control-plane instability and misleading route updates
Rank manipulation	Creates illegitimate parent relationships or abnormal graph positions	Highlights how routing integrity can be corrupted from inside the protocol logic
DIS flooding	Inflates control-message activity and overhead	Reveals whether behavior summaries can distinguish protocol abuse from normal maintenance
Selective forwarding	Drops or withholds data packets while preserving partial normal appearance	Tests the detectability of data-plane degradation and stealthy malicious conduct

CHAPTER 2

CONCEPTUAL FRAMEWORK AND ANALYTICAL MODEL

This chapter develops the conceptual framework and analytical model used to explain and evaluate TACR-HOF. It connects four elements: the operating constraints of low-power and lossy IoT networks, the RPL routing process, behaviour-based intrusion detection, and trust-aware parent selection. The chapter also formalizes the trust score, candidate-rejection rule, and trust-adjusted path cost that distinguish TACR-HOF from MRHOF.

The discussion moves from the general network context to the specific decision model used in this thesis. It first reviews IoT and LLN constraints, RPL and DODAG formation, MRHOF, and the four studied attacks. It then explains the behavioural feature space and multiclass IDS evaluation before presenting the analytical TACR-HOF model and positioning the study relative to existing secure and trust-aware RPL research.

2.1 Internet of Things environments and low-power and lossy networking

IoT environments consist of embedded devices that exchange small amounts of data over constrained wireless links, often with limited energy, memory, processing capacity, and communication bandwidth. In Low-Power and Lossy Networks, connectivity may be unstable, asymmetric, and interference-prone, while nodes frequently depend on multi-hop forwarding to reach the root. Consequently, network performance depends strongly on local link estimation, parent selection, and control-plane stability. A compromised or unreliable forwarding node may therefore affect not only its own traffic but also the traffic of its descendants (Al-Fuqaha et al., 2015; Sobral et al., 2019; Winter et al., 2012).

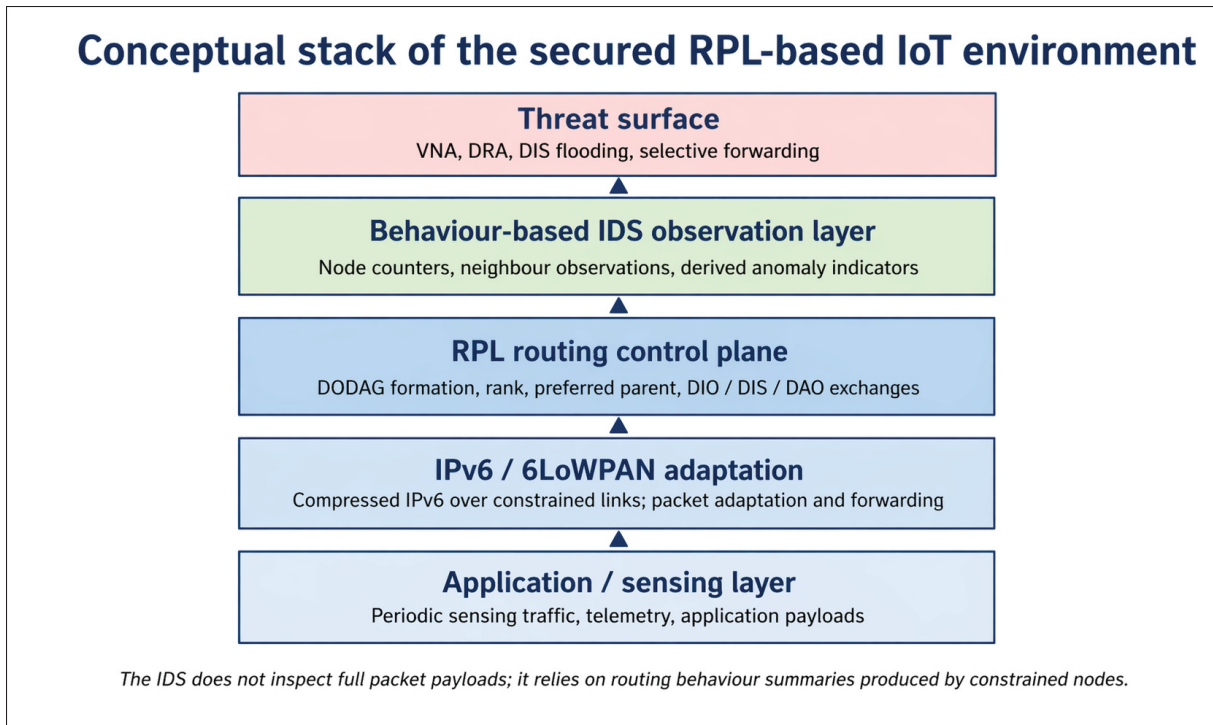


Figure 2.1 Conceptual stack of the secured RPL-based IoT environment adopted in this thesis

Figure 2.1 summarizes the interaction among IoT constraints, RPL routing, behaviour-based intrusion detection, and trust-aware parent selection. Because routing attacks occur within the same protocol layers that generate IDS evidence, routing and security analysis cannot be treated as independent processes. This thesis therefore evaluates them as interacting components of a constrained IoT system (Winter et al., 2012; Tsao et al., 2015; Rahman et al., 2025; Muzammal et al., 2022).

Table 2.1 Structural differences between conventional IP networks and low-power and lossy IoT networks. Author’s synthesis based on (Bormann et al. 2014; Winter et al. 2012; Elrawy et al. 2018; Tsao et al. 2015)

Dimension	Conventional network assumption	LLN / IoT reality	Implication for this thesis
Device capability	Routers and hosts have ample CPU, memory, and storage	Motes operate with tight energy and memory budgets	Security mechanisms must be lightweight and selective
Connectivity	Many paths are stable and high bandwidth	Wireless links are unstable, lossy, and interference-prone	Routing metrics and parent choice materially affect observability
Traffic visibility	Centralized collection is often feasible	Full packet capture across the mesh is expensive and intrusive	Behaviour summaries are more defensible than universal sniffing
Failure domain	Single host compromise often remains local	A bad forwarder can damage upstream and downstream traffic flows	Routing attacks can propagate beyond the malicious node itself

2.2 RPL architecture and the logic of DODAG formation

RPL organizes Low-Power and Lossy Networks as a Destination-Oriented Directed Acyclic Graph (DODAG) rooted at a data-collection point. Each node maintains a rank representing its logical position relative to the root and selects a preferred parent according to the active objective function. Nodes advertise routing information through DIO messages, allowing neighbours to evaluate candidate parents and maintain loop-free upward paths. Because rank reflects routing cost rather than physical distance, manipulated or implausibly attractive values can distort topology formation and traffic flow (Winter et al., 2012; Gnawali & Levis, 2012; Lamaazi & Benamar, 2020).

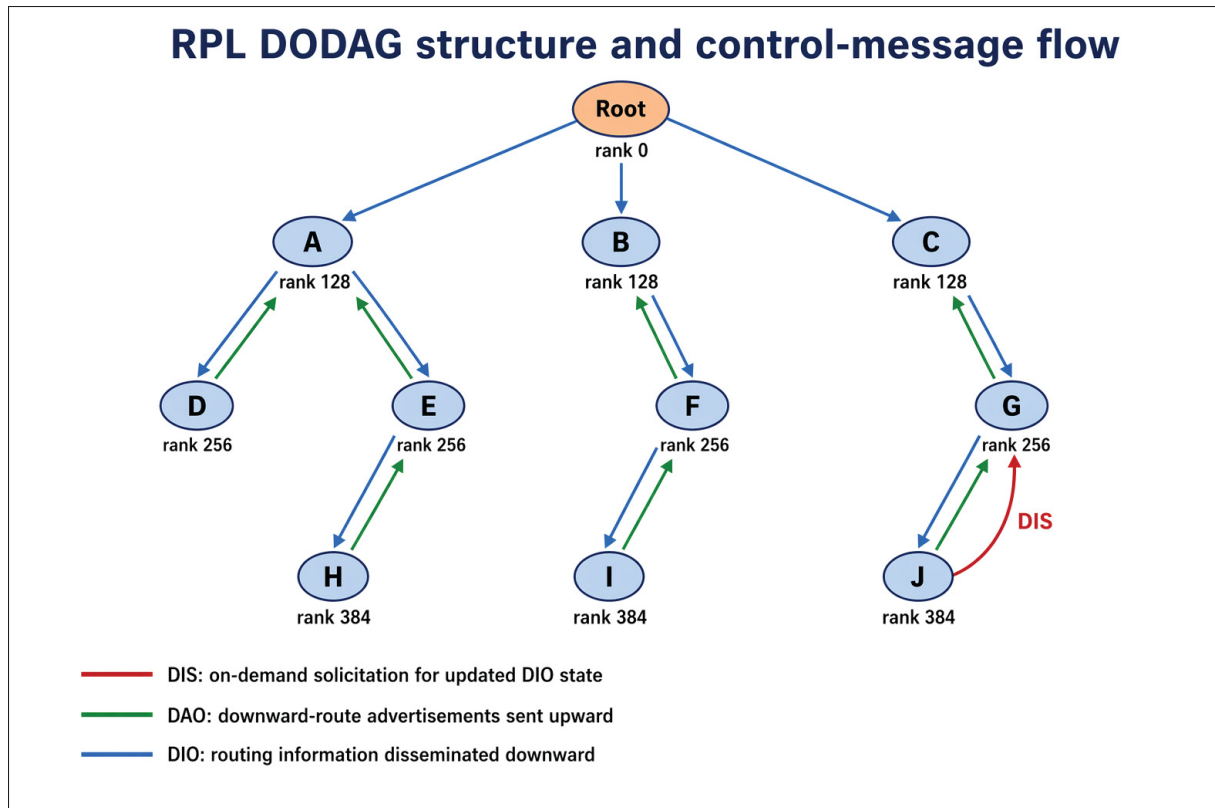


Figure 2.2 RPL DODAG structure and control-message flow used to explain upward and downward routing behaviour. Author's illustration based on (Winter et al. 2012)

The three principal RPL control messages are DIO, DIS, and DAO. DIO messages disseminate routing information and allow neighbouring nodes to learn about the DODAG state. DIS messages are solicitations used when a node needs updated control information.

DAO messages advertise downward-route reachability so that the root can maintain knowledge of reverse paths.

This control exchange is indispensable for network formation, but it also creates an attack surface because malicious manipulation of any of these messages can alter topology perception without necessarily changing application payloads (Winter et al., 2012; Tsao et al., 2015).

Table 2.2 Core RPL concepts and their operational meaning. Author’s synthesis based on (Winter et al. 2012; Tsao et al. 2015)

Concept	Operational role in the network	Why it matters for security analysis
DODAG root	Reference point for route organization and data collection	A natural aggregation point for IDS observations and global state
Rank	Logical distance or path quality indicator under the active objective function	Manipulated rank can attract traffic or destabilize topology
Preferred parent	Neighbour selected as the next hop toward the root	Incorrect parent choice amplifies the impact of malicious behaviour
DIO	Broadcasts routing state to neighbours	Version or rank tampering is usually expressed through DIO inconsistencies
DIS	Requests routing information on demand	Excessive DIS traffic is a signature of solicitation abuse and control-plane stress
DAO	Builds downward routing knowledge	DAO patterns help characterize control-plane balance and stability

2.3 MRHOF routing principles: path cost and hysteresis

In RPL, the objective function determines how candidate parents are compared and how routing paths are formed. MRHOF selects the parent offering the lowest cumulative path cost, commonly using metrics such as Expected Transmission Count (ETX) or rank-based cost. It also applies hysteresis to prevent unnecessary parent changes when competing routes have similar costs, thereby improving stability in variable wireless conditions (Winter et al., 2012; Gnawali & Levis, 2012; Lamaazi & Benamar, 2020).

MRHOF provides an appropriate baseline because it combines efficient path selection with lightweight local decision-making. However, its decisions are primarily metric-driven. A neighbour may remain attractive if its rank or link cost appears favourable, even when its version advertisements, control activity, or forwarding behaviour are suspicious. TACR-HOF addresses

this limitation by retaining the MRHOF path-cost and hysteresis logic while adding behaviour-derived trust penalties to candidate-parent evaluation (Gnawali & Levis, 2012; Muzammal et al., 2022).

Table 2.3 Baseline logic of MRHOF and its security implications. Author’s synthesis based on (Gnawali & Levis 2012; Tsao et al. 2015; Muzammal et al. 2022)

MRHOF component	Routing purpose	Strength under normal conditions	Security limitation
ETX / path cost	Selects paths with lower expected transmission burden	Promotes efficient forwarding and stable routes	Can be misled by a strategically attractive but malicious neighbour
Hysteresis window	Avoids excessive parent oscillation	Improves stability in noisy wireless environments	May delay response when the preferred parent becomes suspicious
Local neighbour comparison	Uses directly observable candidate quality	Computationally lightweight and compatible with LLNs	Does not by itself encode trust, integrity, or forwarding honesty

2.4 Security exposure of RPL and the attack taxonomy used in the thesis

RPL is especially vulnerable to attacks that exploit trust in the routing control plane. A malicious node does not need to compromise cryptographic payload confidentiality to be harmful; it can instead manipulate how routes are formed, maintained, or exercised. Because neighbours infer topology quality from advertised state and observed forwarding, even a modest inconsistency can have path-level consequences. The attacks studied in this thesis were selected because they target this logic directly and because they produce distinguishable but partially overlapping behavioural signatures (Tsao et al., 2015; Almusaylim et al., 2020).

2.4.1 Version Number Attack (VNA)

In a Version Number Attack, the adversary advertises an unjustified increase in the RPL version through DIO messaging. To legitimate neighbours, this can appear as if the DODAG has undergone a valid global update. The effect is a network-wide or neighbourhood-wide repair reaction: nodes reset routing state, exchange fresh control messages, and rebuild their forwarding relationships. The damage is not limited to one packet stream. Instead, VNA injects instability into the protocol's own convergence mechanism (Tsao et al., 2015; Almusaylim et al., 2020).

2.4.2 Decrease Rank Attack (DRA)

In a Decrease Rank Attack, the adversary claims a rank that is artificially low compared with its real topological position. Since lower rank suggests a better path to the root, neighbouring nodes may be attracted to the malicious node as a preferred parent. This changes traffic concentration and can create a basis for further disruption, including loops, congestion, or selective traffic manipulation. The theoretical significance of DRA is that it exploits the semantic meaning of rank itself; it does not simply flood the medium, but corrupts the metric that governs route selection (Tsao et al., 2015; Almusaylim et al., 2020).

2.4.3 DIS Flooding Attack (DISA)

DIS flooding abuses the solicitation mechanism by transmitting excessive DIS messages in order to provoke repeated DIO responses from surrounding nodes. The immediate result is inflated control traffic, unnecessary radio activity, and accelerated consumption of constrained resources. Unlike rank tampering, DIS flooding is not primarily about becoming a preferred parent. Its core effect is to overload the control plane so that legitimate routing maintenance becomes noisy, expensive, and less predictable (Winter et al., 2012; Tsao et al., 2015; Medjek et al., 2021).

2.4.4 Selective Forwarding Attack (SFA)

Selective forwarding is more subtle than direct packet dropping. The malicious node forwards some traffic, but discards selected packets that it should relay onward. In the formulation used in this thesis, the attacker preserves routing control traffic while suppressing data forwarding, thereby attempting to appear like a structurally valid participant in the DODAG while degrading actual delivery. This makes SFA theoretically important: it separates control-plane credibility from data-plane honesty, forcing the IDS to detect discrepancies rather than outright silence (Tsao et al., 2015; Mayzaud et al., 2016)

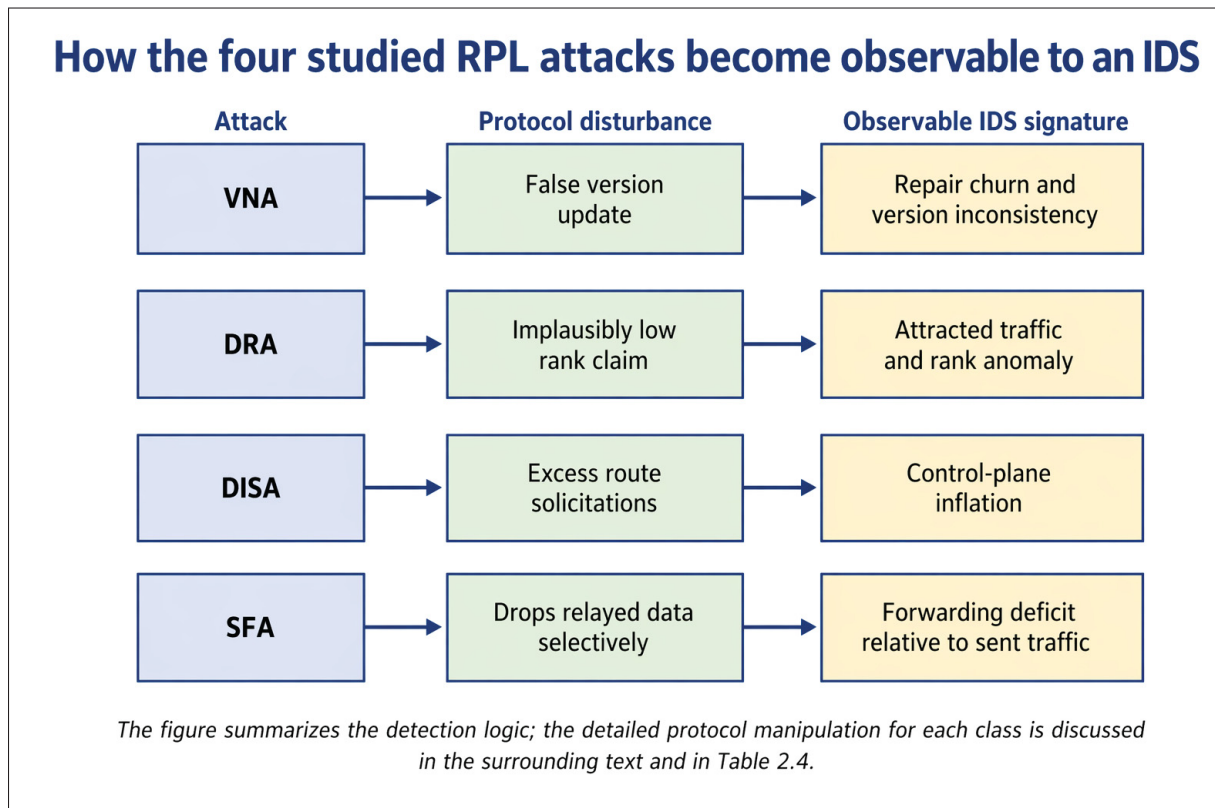


Figure 2.3 Attack-to-symptom chain for the four RPL threats analysed in this thesis. Author's illustration based on (Tsao et al. 2015, Almusaylim et al. 2020, Medjek et al. 2021, and Mayzaud et al. 2016)

Table 2.4 Operational comparison of the four RPL attack classes. Author’s synthesis based on (Tsao et al. 2015, Almusaylim et al. 2020, Medjek et al. 2021; Mayzaud et al. 2016)

Attack	Primary manipulation	Immediate protocol effect	Observable behavioural consequence
VNA	False version increment in DIO state	Unnecessary repair and control churn	Version inconsistency and elevated control ratios
DRA	Advertised rank lower than legitimate value	Traffic attraction and parent misselection	Rank anomalies and altered neighbour reports
DISA	Excess DIS solicitations	Repeated DIO responses from neighbours	Control-plane inflation and temporal instability
SFA	Drops relayed data while preserving protocol participation	Reduced delivery with apparently normal control operation	Forwarding deficits relative to sent traffic

2.5 Intrusion detection theory in constrained IoT networks

Intrusion detection can be organized along several theoretical lines. Signature-based systems compare observations against known attack patterns; specification-based systems encode allowed protocol behaviour; anomaly-based systems learn or infer what normal behaviour looks like and flag significant deviations; and machine-learning systems operationalize classification using labelled or unlabeled data. In constrained IoT environments, the central challenge is not only detection accuracy but also placement, observability, processing cost, and robustness to changing network conditions (da Costa et al., 2019; Rahman et al., 2025).

For RPL networks, the IDS design problem is complicated by multi-hop forwarding and the absence of a single omniscient observation point. A highly distributed IDS may reduce communication burden but increases coordination complexity and demands computation on constrained nodes. A purely centralized packet-capture architecture may offer rich visibility but can be unrealistic because it assumes that most or all traffic can be observed, stored, and

transported to an analysis point without unacceptable overhead (da Costa et al., 2019; Garcia Ribera et al., 2022; Rahman et al., 2025).

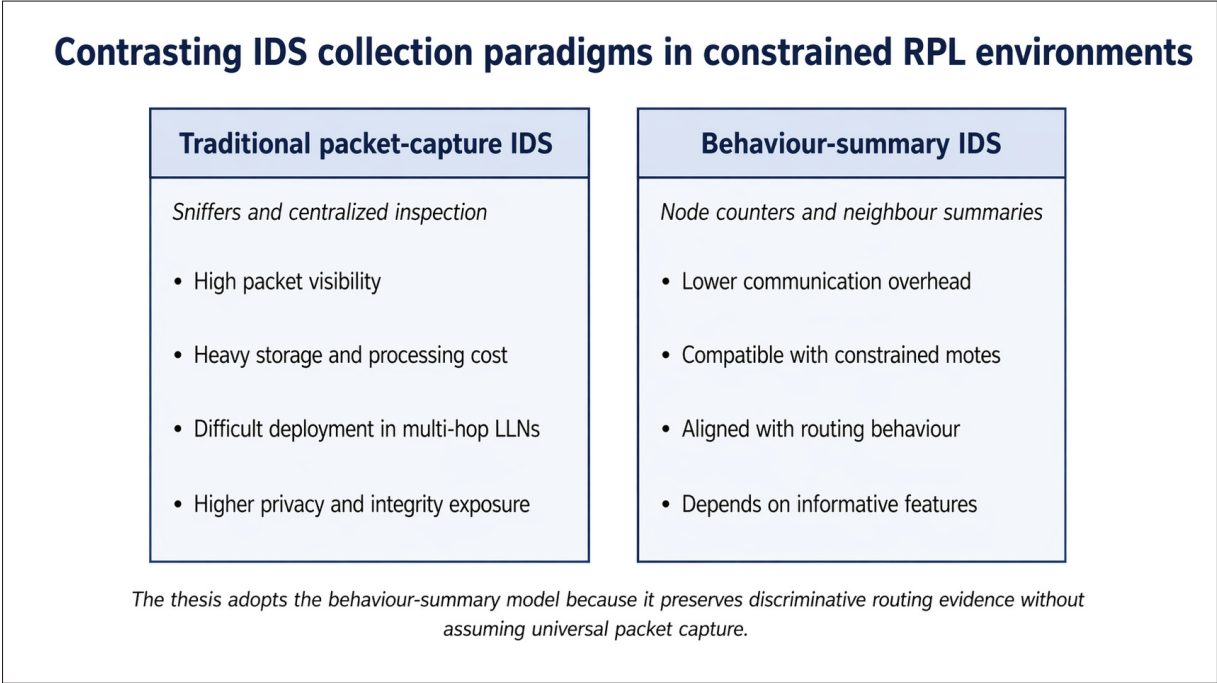


Figure 2.4 Contrasting packet-capture and behaviour-summary IDS collection paradigms in constrained RPL networks

The thesis adopts a behaviour-summary interpretation of anomaly detection. The guiding assumption is that in a resource-constrained mesh, the most defensible IDS is not the one that captures the most bytes, but the one that preserves enough behavioural evidence to discriminate normal from malicious routing activity while respecting LLN constraints. This is why neighbour observations, local counters, and derived ratios become theoretically central. They provide compressed but meaningful evidence about routing integrity (Rahman et al., 2025; da Costa et al., 2019).

Table 2.5 IDS paradigms relevant to RPL-based IoT security. Author’s synthesis based on (da Costa et al. 2019; Rahman et al. 2025)

Paradigm	Core idea	Main advantage	Main limitation in LLNs
Signature-based	Match known patterns	Efficient when attacks are stable and predefined	Weak against novel or slightly modified attacks
Specification-based	Detect deviation from protocol rules	Interpretable and protocol-aware	May be brittle under dynamic network behaviour
Packet-capture anomaly IDS	Learn from raw or parsed traffic traces	High visibility when capture is feasible	Storage, privacy, and deployment burden are high
Behaviour-summary IDS	Use aggregated node and neighbour statistics	Lower overhead and better fit for constrained operation	Success depends on strong feature design and accurate labelling

2.6 Behavioural feature theory for RPL intrusion detection

A behaviour-based IDS does not need to preserve every packet if its summaries retain the routing evidence required to identify malicious activity. In this thesis, the behavioural schema is organized into node-specific variables, neighbour-aggregated observations, and derived anomaly indicators. These categories distinguish what a node reports about itself, what neighbouring nodes observe, and what discrepancies can be calculated between the two perspectives (da Costa et al., 2019; Rahman et al., 2025).

Node-specific variables describe local state and activity, including parent selection, rank, version, and control-message counts. Neighbour-aggregated variables represent externally observed rank, version, forwarding, and control behaviour. Derived indicators convert these observations into discrepancies and ratios, allowing attacks to be identified through abnormal relationships

among variables rather than through individual packet counts alone (Rahman et al., 2025; da Costa et al., 2019; HUNSR, n.d.).

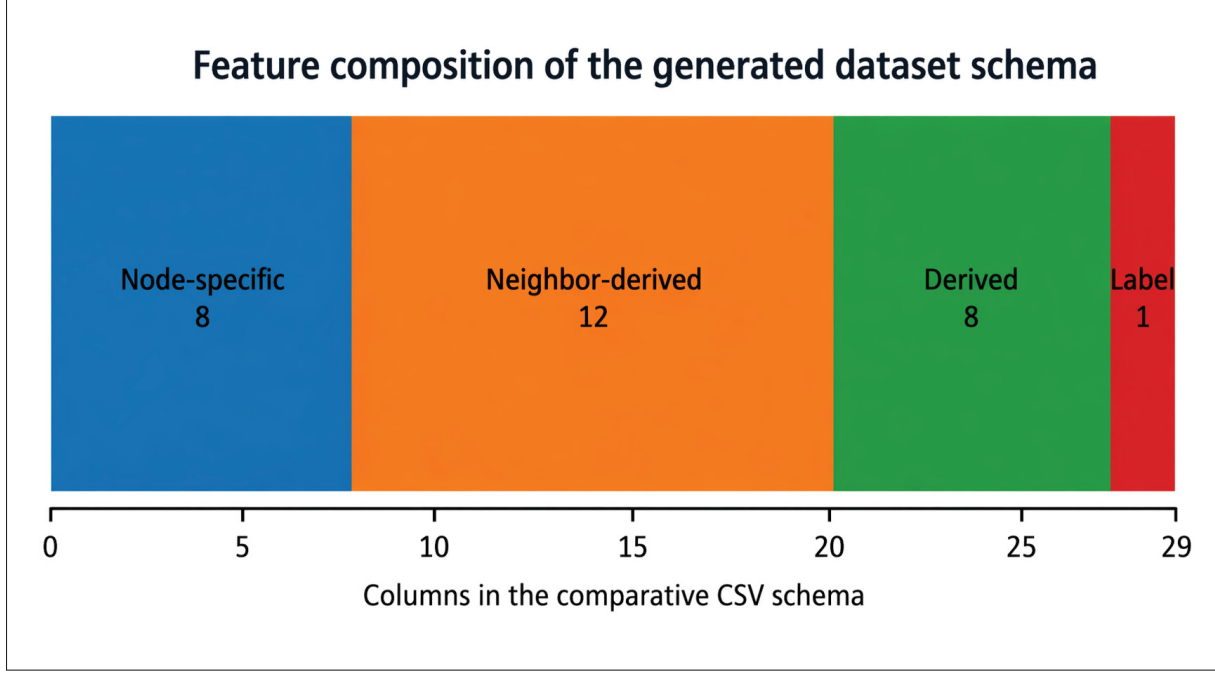


Figure 2.5 Conceptual grouping of the behavioural schema later operationalized in the comparative datasets

The principal derived indicators measure version inconsistency, rank inconsistency, and control-to-data imbalance. Version and rank differences compare a node’s reported state with neighbour observations, while the control-to-data ratio identifies control activity that is disproportionate to useful forwarding. These relationships provide evidence of version manipulation, deceptive rank advertisement, control-plane flooding, and forwarding anomalies (HUNSR, n.d.; da Costa et al., 2019; Rahman et al., 2025).

$$\Delta_{\text{version}} = |V_{\text{node}} - \text{Avg}(V_{\text{neighbours}})| \quad (2.1)$$

$$\Delta_{\text{rank}} = |R_{\text{node}} - \text{Avg}(R_{\text{neighbours}})| \quad (2.2)$$

$$\text{Control-to-data ratio} = \frac{\text{DIS} + \text{DIO} + \text{DAO}}{\text{forwarded data traffic}} \quad (2.3)$$

Source: Author's formulation based on da Costa et al. (2019), Rahman et al. (2025), and HUNSR n.d..

Table 2.6 Behavioural feature families and their theoretical meaning. Author's synthesis based on (da Costa et al. 2019; Rahman et al. 2025; HUNSR n.d.)

Feature family	Representative examples	What the family reveals
Node-specific	time_sec, node_id, parent_id, rpl_ver, rpl_rank, dis_sent, dio_sent, dao_sent	How the node describes its own state and immediate control activity
Neighbour-aggregated	nbr_dis_rcv, nbr_dio_rcv, nbr_rpl_ver_rcv, nbr_rpl_rank_rcv, nbr_fwd_to_me, nbr_fwd_to_others	How surrounding nodes collectively observe the target node
Derived anomaly indicators	diff_rpl_ver, diff_rpl_rank, norm_rank_diff, ctrl_to_data_ratio, rpl_fwd_ratio, total_fwd_ratio	Whether self-reported behaviour and observed behaviour form a plausible routing pattern

2.7 Classifier and evaluation principles for multiclass RPL intrusion detection

The intrusion-detection task in this thesis is multiclass rather than binary. The classifiers must distinguish normal behaviour from Version Number, Decrease Rank, DIS Flooding, and Selective Forwarding attacks. This distinction is important because identifying the attack type provides more useful information for routing defence and security analysis than a general benign-versus-malicious decision (Tsao et al., 2015; da Costa et al., 2019; Rahman et al., 2025).

The labelled behavioural windows are evaluated using Decision Tree, Random Forest, and K-Nearest Neighbours. Decision Tree provides an interpretable rule-based baseline, Random Forest improves robustness through ensemble voting, and K-Nearest Neighbours tests whether attack classes remain separable through local similarity. Using classifiers with different decision

mechanisms helps determine whether improvements under TACR-HOF are consistent rather than specific to one learning model (Quinlan, 1986; da Costa et al., 2019; Rahman et al., 2025).

Because normal observations may outnumber individual attack classes, accuracy alone is insufficient. Precision measures the reliability of class predictions, recall indicates how much of each true class is detected, F1-score balances precision and recall, and the confusion matrix reveals class-specific errors. Together, these metrics provide a defensible assessment of multiclass detection performance (HUNSR, n.d.; Rahman et al., 2025; da Costa et al., 2019).

Table 2.7 Evaluation metrics used to interpret multiclass intrusion-detection performance. Author’s synthesis based on (da Costa et al. 2019; Rahman et al. 2025; Opitz 2024)

Metric	Interpretive meaning	Why it matters in this thesis
Accuracy	Overall proportion of correctly classified instances	Useful as a global indicator but insufficient on its own under class imbalance
Precision	How often a predicted class is actually correct	Important when false alarms would wrongly label benign routing behaviour as malicious
Recall	How much of a true class is successfully retrieved	Critical for showing that minority attack classes are not being missed
F1-score	Harmonic balance between precision and recall	Provides a stronger summary when both missed attacks and false alarms matter
Confusion matrix	Class-by-class distribution of correct and incorrect decisions	Reveals which attacks are separable and which are mutually confusable

2.8 Conceptual integration of trust evidence into parent selection

A behaviour-based IDS can identify suspicious activity, but it does not directly prevent an unreliable neighbour from remaining a preferred parent. TACR-HOF addresses this limitation by incorporating behavioural evidence into parent selection. It retains the MRHOF path-cost

foundation while adding trust penalties, so candidate parents are evaluated according to both routing efficiency and observed reliability (Djedjig et al., 2020; Muzammal et al., 2022).

The contribution is not the first use of trust-aware routing in RPL. Rather, TACR-HOF is evaluated as an MRHOF-compatible extension whose effect is measured through the visibility, separability, and classification of malicious behaviour under a controlled MRHOF-versus-TACR-HOF comparison.

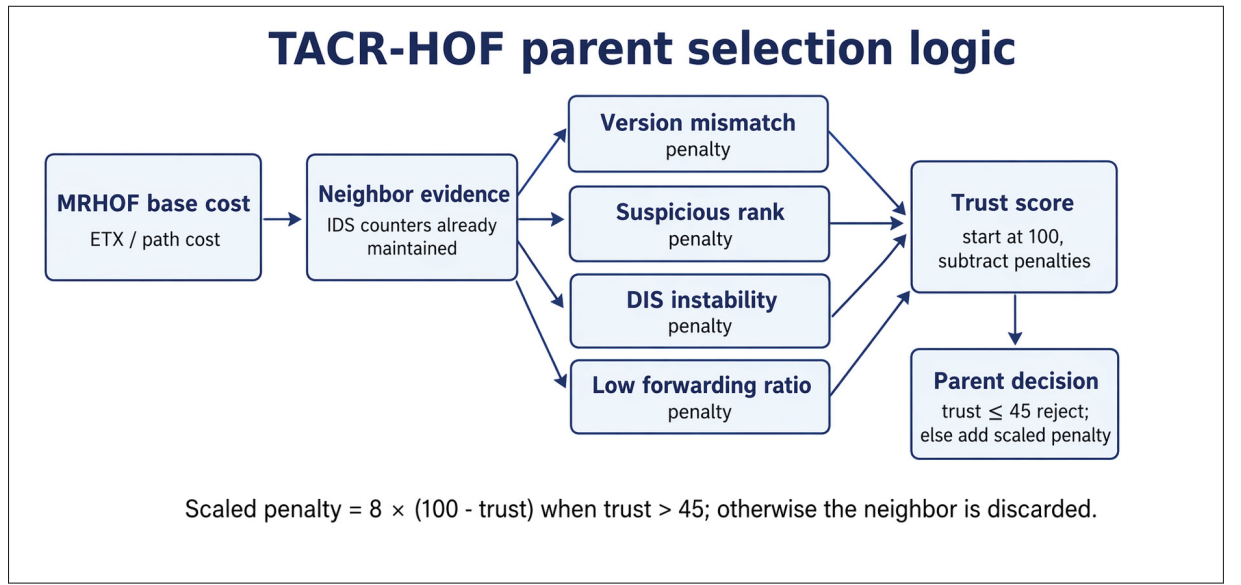


Figure 2.6 Conceptual logic of TACR-HOF: MRHOF path cost is retained and adjusted by behaviour-derived trust penalties

Formal definition of TACR-HOF. Let n denote a candidate parent. TACR-HOF preserves the MRHOF routing basis and augments it with behaviour-derived distrust. The trust-aware path cost is defined as

$$C_{\text{TACR}}(n) = C_{\text{MRHOF}}(n) + P(n),$$

where $C_{\text{MRHOF}}(n)$ is the additive ETX/rank-based path cost inherited from MRHOF, and $P(n)$ is a trust-derived penalty term.

Trust is initialized at 100 and reduced by fixed penalties:

$$T(n) = 100 - p_{\text{ver}}(n) - p_{\text{rank}}(n) - p_{\text{dis}}(n) - p_{\text{fwd}}(n).$$

The penalties used in the implemented prototype are:

$$p_{\text{ver}} = \begin{cases} 35, & \text{if advertised version is inconsistent with local DODAG knowledge} \\ 0, & \text{otherwise} \end{cases}$$

$$p_{\text{rank}} = \begin{cases} 25, & \text{if the neighbour appears implausibly attractive in rank} \\ 0, & \text{otherwise} \end{cases}$$

$$p_{\text{dis}} = \min(20, \text{DIS}_{rx} - \text{DIO}_{rx}) \quad \text{when } \text{DIS}_{rx} > \text{DIO}_{rx} + 3$$

$$p_{\text{fwd}} = \begin{cases} 30, & \text{if forwarding ratio} < 50\% \\ 15, & \text{if } 50\% \leq \text{forwarding ratio} < 70\% \\ 0, & \text{otherwise.} \end{cases}$$

A candidate is rejected when $T(n) \leq 45$. Otherwise, the residual distrust is converted into additive routing cost:

$$P(n) = 8 \times (100 - T(n)).$$

This formulation makes the prototype's logic explicit: TACR-HOF is not a new RPL protocol, but an MRHOF-compatible objective function that modifies parent selection through behaviour-derived penalties.

Algorithm 2.1 summarizes the precise parent-selection rule executed in the uploaded TACR-HOF implementation and makes explicit the relationship between behavioural evidence, trust reduction, and final path cost

Algorithm 2.1 TACR-HOF Candidate Parent Evaluation

Input: Candidate parent n
Output: Adjusted path cost or candidate rejection

```

1  $base\_cost \leftarrow MRHOF\_Path\_Cost(n);$ 
2  $trust \leftarrow 100;$ 
3 if version mismatch then
4    $trust \leftarrow trust - 35;$ 
5 end if
6 if suspiciously attractive rank then
7    $trust \leftarrow trust - 25;$ 
8 end if
9 if  $DIS_{rx} > DIO_{rx} + 3$  then
10   $trust \leftarrow trust - \min(20, DIS_{rx} - DIO_{rx});$ 
11 end if
12 if  $sent\_to\_nbr \geq 8$  then
13    $forwarding\_ratio \leftarrow 100 \times \frac{forwarded\_from\_nbr}{sent\_to\_nbr};$ 
14   if  $forwarding\_ratio < 50$  then
15      $trust \leftarrow trust - 30;$ 
16   end if
17   else if  $forwarding\_ratio < 70$  then
18      $trust \leftarrow trust - 15;$ 
19   end if
20 end if
21 if  $trust \leq 45$  then
22   return  $REJECT(n);$ 
23 end if
24  $trust\_penalty \leftarrow 8 \times (100 - trust);$ 
25 return  $base\_cost + trust\_penalty;$ 

```

TACR-HOF extends MRHOF rather than replacing it. Each candidate parent is first evaluated using the conventional MRHOF path cost. Behavioural evidence related to version consistency, rank plausibility, DIS instability, and forwarding performance then reduces a trust score initialized at 100. Candidates with trust at or below 45 are rejected; otherwise, the remaining distrust is converted into an additive path penalty using $8 \times (100 - \text{trust})$. The adjusted cost is then processed through the existing MRHOF parent-selection and hysteresis logic. Table 2.8 summarizes the behavioural conditions and their routing consequences.

This formulation describes the specific TACR-HOF prototype evaluated in this thesis. The thresholds and penalty values are fixed implementation parameters rather than universally optimal values and would require calibration before deployment under different network conditions.

Table 2.8 Behavioural evidence used by the TACR-HOF concept and its routing consequence. Author’s synthesis based on (HUNSR n.d.; Djedjig et al. 2020; Muzammal et al. 2022; Gnawali & Levis 2012)

Evidence source	Illustrative interpretation	Trust effect in the uploaded prototype	Routing consequence
Version mismatch	Neighbour advertises a version inconsistent with local DODAG knowledge	Subtract 35 trust points	Discourages parents that appear to fabricate global state
Suspicious rank	Neighbour claims a rank that is implausibly attractive	Subtract 25 trust points	Reduces attraction to low-rank deceptive parents
DIS instability	DIS solicitations exceed expected control balance	Penalty up to 20 points	Penalizes control-plane abuse and instability
Poor forwarding ratio	Neighbour receives traffic for forwarding but forwards too little onward	Subtract 30 points if ratio $< 50\%$; 15 points if ratio $< 70\%$	Targets selective forwarding and hidden data-plane dishonesty
Low-trust threshold	Composite trust becomes critically low	Reject if trust ≤ 45 ; otherwise apply $8 \times (100 - \text{trust})$	Either discard the candidate or inflate its effective path cost

The trust mechanism considered in this thesis is deliberately bounded. It uses a fixed set of behaviour-derived indicators and fixed penalty values associated with the implemented TACR-HOF prototype. It does not evaluate learned trust functions, indirect recommendations from multiple nodes, decentralized consensus, adaptive threshold calibration, or adversarial manipulation of trust reports. Consequently, the conclusions apply to the specific TACR-HOF configuration and behavioural schema evaluated in the simulation campaign. Alternative trust architectures remain a separate research direction.

2.9 Positioning TACR-HOF Relative to Existing Secure and Trust-Aware RPL Proposals

Existing RPL security research can be grouped into five related directions: standard RPL and MRHOF design, RPL attack analysis, secure and trust-aware routing, IDS and machine-learning-based detection, and simulation-based evaluation. Standard RPL research establishes rank, DODAG formation, path-cost optimization, and hysteresis, while attack studies examine threats such as version manipulation, rank deception, DIS flooding, and selective forwarding. Trust-aware proposals incorporate reputation or behavioural reliability into parent selection, whereas IDS studies use protocol or behavioural features to classify malicious activity (Winter et al., 2012; Gnawali & Levis, 2012; Tsao et al., 2015; Djedjig et al., 2020; Muzammal et al., 2022; da Costa et al., 2019).

Despite these advances, trust-aware routing is commonly evaluated through routing resilience, packet delivery, or overhead, while IDS research often treats the routing objective function as a fixed background condition. This thesis connects the two perspectives by comparing MRHOF and TACR-HOF using the same topology family, attack classes, behavioural feature schema, and classifier pipeline. The objective is to determine whether trust-aware parent selection improves not only routing resilience but also the quality and separability of IDS evidence (Elrawy et al., 2018; Gyamfi & Jurcut, 2022; Rahman et al., 2025; HUNSR, n.d.).

Table 2.9 Comparative positioning of TACR-HOF relative to existing RPL security research. Author’s synthesis based on (Winter et al. 2012; Gnawali & Levis 2012; Tsao et al. 2015; Djedjig et al. 2020; Muzammal et al. 2022; da Costa et al. 2019; Elrawy et al. 2018; Gyamfi and Jurcut 2022; Rahman et al. 2025; HUNSR n.d.)

Research direction	What prior work mainly does	Remaining gap	Positioning of this thesis
Standard RPL and MRHOF	Defines rank-based routing, ETX/path-cost optimization, and hysteresis-based parent stability.	Does not include behavioural trust or IDS-evidence quality in parent selection.	Uses MRHOF as the baseline for comparison.
Secure RPL and attack mitigation	Studies attacks such as rank manipulation, version abuse, DIS flooding, and selective forwarding.	Often focuses on attack detection or mitigation without testing how objective functions affect IDS evidence.	Evaluates four RPL attack classes under both MRHOF and TACR-HOF.
Trust-aware RPL proposals	Adds trust, reputation, or forwarding reliability to routing decisions.	Usually emphasizes routing resilience more than downstream IDS feature quality.	Positions TACR-HOF as a trust-aware MRHOF extension evaluated through IDS evidence quality.
IDS and ML-based RPL security	Uses behavioural features or classifiers to detect malicious nodes or traffic.	Often treats routing as fixed background infrastructure.	Treats routing policy as a variable that shapes attack visibility and feature separability.
Simulation-based RPL evaluation	Uses Cooja/Contiki-NG to test routing, attacks, and detection under controlled settings.	Simulation does not prove deployment-level robustness.	Frames results as controlled comparative evidence, with hardware validation left for future work.

Table 2.9 clarifies that the thesis does not claim novelty from using RPL, trust, IDS, machine learning, or simulation in isolation. Its contribution lies in combining these elements into a controlled MRHOF-versus-TACR-HOF comparison where the routing objective function becomes the independent variable and IDS evidence quality becomes a central evaluation outcome.

Accordingly, TACR-HOF is positioned in this thesis as a security-oriented routing design whose novelty lies in the comparative evaluation framework around it. The thesis does not claim that trust-aware routing, secure RPL, or IDS-based RPL attack detection are unexplored concepts. Instead, it contributes a controlled MRHOF-versus-TACR-HOF study showing how trust-aware parent-selection logic can reshape the behavioural evidence available to IDS models. This positioning narrows the contribution, but it also makes it more defensible: the thesis is not built on an exaggerated claim of absolute novelty, but on a specific and measurable claim about routing-objective design, IDS evidence quality, and multiclass attack detectability in constrained RPL-based IoT simulations.

2.10 Research gap and the conceptual framework of the thesis

Previous studies have shown that RPL attacks can be detected and that trust or reputation can be incorporated into parent selection. However, fewer studies examine whether the routing objective function itself changes the quality of behavioural evidence available to an intrusion detection system. IDS research often treats routing as a fixed background condition, while trust-aware routing studies commonly emphasize resilience, packet delivery, or overhead rather than downstream classification quality.

This thesis addresses that gap through a controlled comparison between MRHOF and TACR-HOF. It uses behaviour summaries instead of universal packet capture, combines node-reported and neighbour-observed features, and evaluates whether trust-aware parent selection improves the visibility, separability, and classification of four RPL attack classes. The conclusions are limited to the evaluated Contiki-NG/Cooja environment, fixed topologies, periodic UDP traffic, and the

implemented TACR-HOF trust configuration; hardware deployment, mobility, heterogeneous workloads, adaptive trust, and privacy-preserving data collection remain outside the present scope (da Costa et al., 2019; Rahman et al., 2025; Muzammal et al., 2022; HUNSR, n.d.).

Table 2.10 Research gaps identified in the literature and the conceptual response adopted by this thesis. Author’s synthesis based on (da Costa et al. 2019; Rahman et al. 2025; Muzammal et al. 2022; HUNSR n.d.)

Observed gap	Why the gap matters	Conceptual response in this thesis
Dependence on full packet capture	Difficult to justify in multi-hop constrained networks	Use behaviour summaries and root-side feature construction
Weak integration between IDS and routing	Detection may identify a problem without changing route choice	Evaluate a trust-aware objective function beside the MRHOF baseline
Bias from static attack placement or narrow scenarios	Models may memorize nodes instead of learning behaviour	Use randomized attack scheduling and multiple topologies
Overreliance on global accuracy	Minority attack classes may remain poorly detected	Interpret results through precision, recall, F1, and confusion matrices

2.11 Chapter summary

This chapter established the conceptual framework for comparing MRHOF and TACR-HOF. It reviewed the relevant IoT, RPL, attack, intrusion-detection, and behavioural-feature concepts; defined the multiclass evaluation principles; and formalized TACR-HOF through its trust penalties, rejection threshold, and adjusted path-cost model.

The chapter also positioned the study relative to existing secure and trust-aware RPL research and clarified the limits of the implemented trust configuration. Chapter 3 applies this framework through the simulation design, attack scheduling, dataset generation, and comparative evaluation pipeline.

CHAPTER 3

METHODOLOGY

This chapter presents the methodological foundation of the thesis and defines the experimental pipeline used to evaluate intrusion-aware routing in RPL-based IoT networks. The chapter has two roles. First, it explains how the behavioural dataset is generated from Contiki-NG/Cooja simulations through node summaries rather than packet capture. Second, it formalizes the comparative design used later in Chapter 4, where the Minimum Rank with Hysteresis Objective Function (MRHOF) is evaluated against the trust-aware TACR-HOF variant. The aim is not to list tools and parameters, but to justify why the design choices are appropriate, reproducible, and controlled to support conclusions. Methodologically, the study combines protocol-level modification, controlled network simulation, attack scheduling, feature engineering, and downstream IDS evaluation. The uploaded artefacts make this possible at three levels: the base RPL-IDS-Beh simulation environment, the TACR prototype that modifies parent selection inside RPL-lite, and the comparative experiment bundle containing matched folders of logs, figures, and CSV datasets. Accordingly, the present chapter is written from the actual experimental artefacts rather than from an abstract generic workflow (HUNSR, n.d.; Oikonomou et al., 2022).

Methodology at a glance

- The comparative campaign uses six fixed Cooja topologies and two routing schemes, producing twelve principal simulation runs.
- The attack scheduler randomizes target nodes, activation times, and durations while keeping the attack pressure proportional to topology size.
- Behaviour summaries are collected at the node level and converted into a shared CSV schema containing node-specific, neighbor-derived, and derived indicators.
- TACR-HOF preserves MRHOF base path-cost logic but adds trust penalties derived from IDS counters already maintained by the modified stack.
- The methodology is intentionally written with explicit validity notes, including the distinction between fixed topology seeds and randomized attack realizations.

3.1 Overall research design and workflow

At the highest level, the methodology follows a staged experimental design. A Contiki-NG RPL environment is first configured either with the baseline MRHOF objective function or with the TACR-HOF extension. The same family of Cooja topologies is then executed in batch mode. During execution, randomized attacks are injected by the Cooja logger script, while each node continues to participate in ordinary RPL operation and application-layer communication. Nodes periodically expose behavioural summaries that are later converted into structured feature vectors. These vectors are finally used to produce classifier outputs and routing comparisons that are interpreted afterwards (Oikonomou et al., 2022; Osterlind et al., 2006; HUNSR, n.d.).

This separation between network simulation, log extraction, feature construction, and learning-based evaluation is important. It prevents the IDS stage from being conflated with the simulation stage and makes it possible to inspect the intermediate artefacts independently. For thesis-defense purposes, this modularity is a strength: the topology files, the attack scheduler, the raw .testlog files, the generated CSV datasets, and the classifier summary outputs can each be audited without requiring the entire pipeline to be treated as a black box (Osterlind et al., 2006; HUNSR, n.d.).

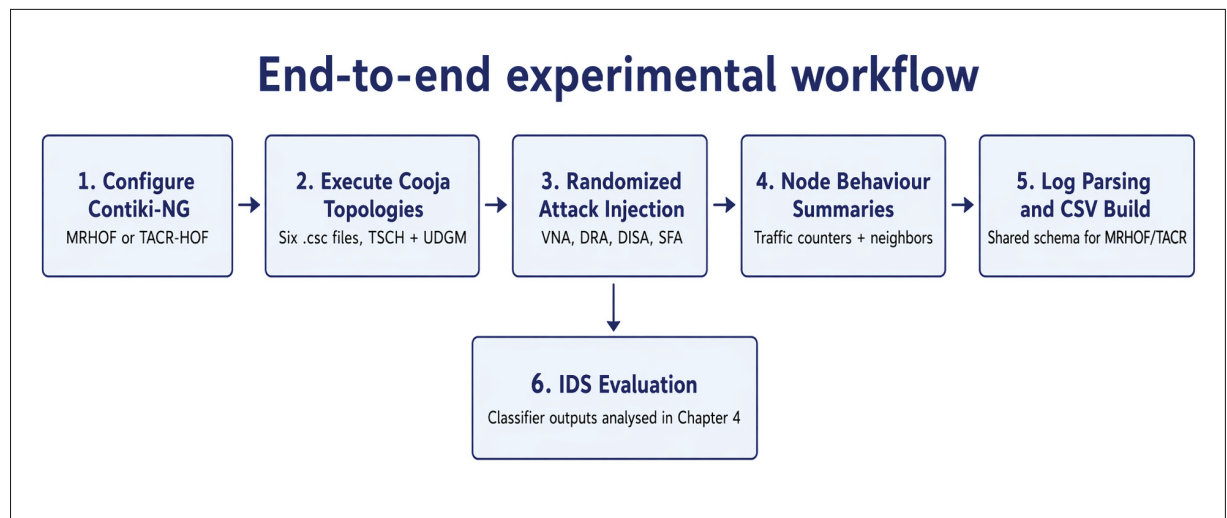


Figure 3.1 End-to-end methodological pipeline used to move from Contiki-NG configuration to comparative IDS-ready datasets

3.2 Simulation platform, execution environment, and control variables

The simulation platform is built on Contiki-NG and the Cooja simulator, using the RPL-lite implementation as the routing layer. The application layer is a lightweight UDP request/response pattern in which non-root motes periodically send data toward the root and the root responds, thereby creating both upward traffic and observable forwarding behavior along the multi-hop path. This is not incidental traffic: it is the mechanism that makes routing anomalies measurable. Without application messages, attacks such as selective forwarding would be substantially harder to characterize because the forwarding path would be only weakly exercised (Oikonomou et al., 2022; Osterlind et al., 2006; Winter et al., 2012; HUNSR, n.d.).

Batch execution is handled by `sequential_executor.py`, which discovers all `.csc` files in the scenario directory and runs them in no-GUI mode. The script also renames the Cooja output from the default `COOJA.testlog` into a topology-specific `.testlog` file after each run, which simplifies later parsing and avoids accidental overwrite. Inside the uploaded topology files, the random seed is fixed at 123456, the radio medium is UDGM, the transmission and interference ranges are 50 m and 100 m respectively, and TSCH is used at the MAC layer. At the application level, the packet-sending interval is 25 s, while IDS summaries are emitted every 480 s. The DIS flooding mechanism, when active, is driven with a 15 s trigger interval (HUNSR, n.d.; Watteyne et al., 2015).

Table 3.1 Core simulation and comparative-experiment parameters (HUNSR, n.d.)

Methodological parameter	Value
Simulator and stack	Contiki-NG with the Cooja simulator and RPL-lite implementation
Comparative routing schemes	Baseline MRHOF and trust-aware TACR-HOF
Topology files	Six fixed .csc scenarios: cooja10-1, cooja10-2, cooja16-1, cooja16-2, cooja50-1, cooja50-2
Total motes per topology	10, 16, and 50 total motes, each instantiated once with a corner root and once with a near-centre root
Random seed in .csc files	123456 in every uploaded topology file
Radio medium	UDGM
Transmission range	50 m
Interference range	100 m
MAC layer	TSCH
Application packet interval	25 s
IDS reporting interval	480 s
DISA trigger interval	15 s when the attack flag is active
Simulation duration	36,000 s (10 h) per run
Execution mode	Headless batch execution through sequential_executor.py, which iterates over all .csc files
Attack scheduling rule	Number of attack windows equals the total number of motes in the topology; activation times and targets are randomized
Nominal attack duration	5,000-10,000 s with a 120 s same-node conflict margin in the scheduler logic
Comparative data products	Six MRHOF logs, six TACR logs, and one CSV dataset per routing scheme with a shared 29-column schema

3.3 Topology design and scenario coverage

The comparative campaign uses six fixed Cooja scenarios: cooja10-1, cooja10-2, cooja16-1, cooja16-2, cooja50-1, and cooja50-2. These names indicate the total number of motes in each topology, including the root. The corresponding networks therefore contain 10, 16, or 50 total motes, with 9, 15, or 49 non-root motes, respectively. This distinction is important because the attack scheduler sets the number of attack windows according to the total mote count.

Each network size is evaluated with two root placements: one near the deployment boundary and one near the centre. Root position affects hop count, parent diversity, traffic concentration, and the influence of compromised forwarding nodes. The smaller topologies provide more interpretable cases, while the 50-mote scenarios evaluate the routing schemes under denser and broader network conditions (Winter et al., 2012; HUNSR, n.d.).

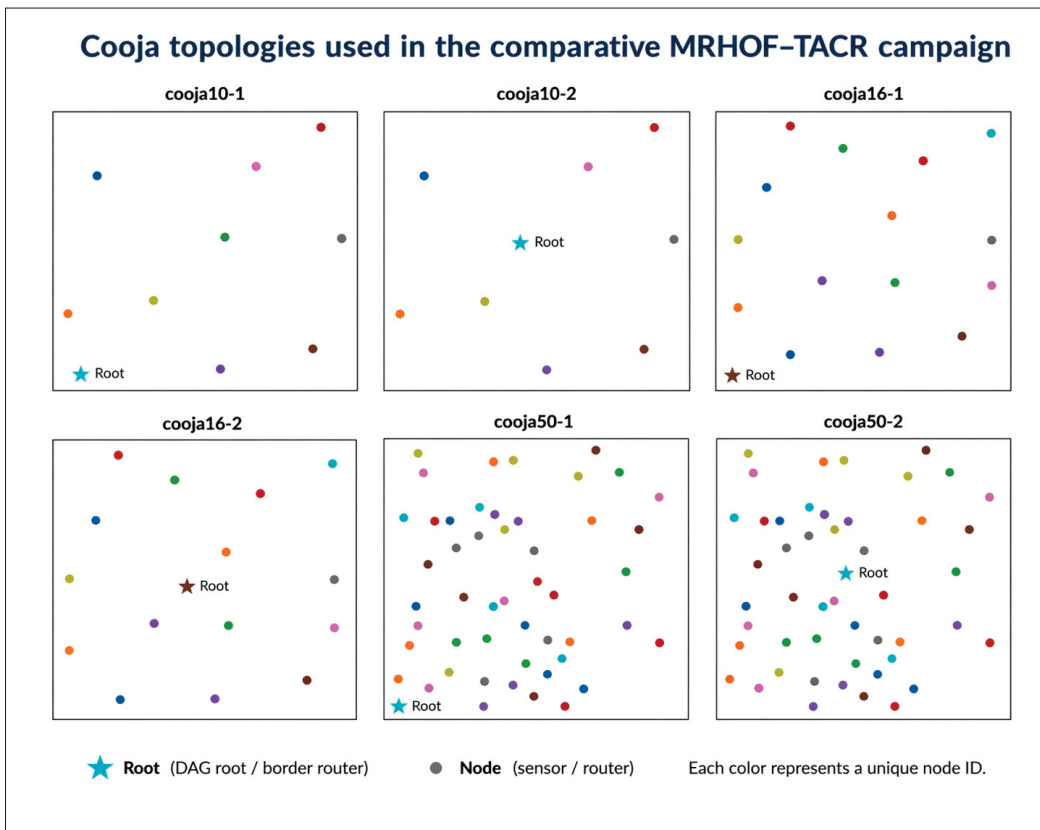


Figure 3.2 The six fixed Cooja topologies used in the comparative campaign. (Contiki-NG/Cooja GUI)

Table 3.2 Inventory of topology files extracted directly from the uploaded .csc scenarios.
Author's extraction from the Cooja .csc topology scenario files used in the experiment

Topology	Total motes	Non-root motes	Root placement	Root coordinates	Seed	TX/INT range (m)
cooja10-1	10	9	Near corner	(9.5, 8.1)	123456	50/100
cooja10-2	10	9	Near center	(56.7, 55.6)	123456	50/100
cooja16-1	16	15	Near corner	(-6.8, 0.6)	123456	50/100
cooja16-2	16	15	Near center	(54.7, 54.0)	123456	50/100
cooja50-1	50	49	Near corner	(4.7, 11.3)	123456	50/100
cooja50-2	50	49	Near center	(89.8, 102.5)	123456	50/100

3.4 Adversarial model and randomized attack scheduling

Four RPL-centric attacks are considered in the methodological design: Version Number Attack (VNA), Decrease Rank Attack (DRA), DIS Flooding Attack (DISA), and Selective Forwarding Attack (SFA). The intent is to cover both control-plane manipulation and data-plane degradation. VNA and DRA distort routing state information advertised by a malicious node, DISA inflates control overhead through repeated solicitations, and SFA targets forwarding reliability while preserving enough plausible behavior to avoid trivial detection (Tsao et al., 2015; Verma & Ranga, 2020).

The attack scheduler implemented in the logger scripts is randomized at runtime. For each topology, the script sets `total_attacks` equal to the number of motes in the simulation, then repeatedly samples an attack type, a target node, a start time, and a duration. The duration is drawn from a 5,000-10,000 s band, and a 120 s same-node conflict margin is imposed so that two attack windows do not overlap too closely on the same victim mote. This design is methodologically stronger than attaching fixed attacks to fixed nodes because it reduces the risk

that a classifier merely memorizes node identity, constant rank values, or static attack timing patterns (HUNSR, n.d.).

Parsing the actual comparative logs confirms that the realized attack-window counts match the topology size: 10 windows in each 10-mote run, 16 in each 16-mote run, and 50 in each 50-mote run. The observed duration range in the comparative artefacts remains inside the intended 5,000-10,000 s band for both routing schemes. However, a defensible reading of the experiment must note that randomized scheduling means the MRHOF and TACR runs are matched at the level of topology family and attack-pressure distribution, not necessarily at the level of identical per-node, per-second attack traces unless an external fixed schedule is enforced (HUNSR, n.d.).

Table 3.3 Attack classes and their operational manifestation in the simulation. Author's synthesis based on (Tsao et al. 2015; Verma & Ranga 2020; HUNSR n.d.)

Attack class	Operational behavior in the simulation	Primary protocol surface	Expected behavioural footprint
Version Number Attack (VNA)	The malicious mote advertises an incremented DODAG version to provoke unnecessary global repairs and control churn.	RPL version field / DIO dissemination	Abnormal version discrepancies, elevated control activity, and topology instability.
Decrease Rank Attack (DRA)	The attacker announces an artificially attractive rank so that neighbors route through it.	RPL rank field	Suspicious rank gaps, route attraction, and degraded forwarding behaviour.
DIS Flooding Attack (DISA)	The compromised mote repeatedly emits DIS requests so neighbors respond with fresh DIO messages.	DIS control exchanges	Inflated DIS and DIO activity, excessive control overhead, and energy wastage.
Selective Forwarding Attack (SFA)	The attacker forwards protocol control traffic but drops data packets selectively to remain less visible than a hard blackhole.	Forwarding path / packet relay behaviour	Reduced delivery reliability and forwarding-ratio anomalies.

3.5 Data acquisition, behavioural summaries, and feature engineering

A central design principle of the methodology is that the dataset is behaviour-based rather than packet-capture-based. Instead of streaming every frame to a sniffer infrastructure, each node periodically summarizes what it has observed about itself and about neighboring nodes. These summaries are rich enough to preserve attack signatures while remaining far more

practical for low-power and lossy networks than exhaustive packet collection. The root-side and post-processing stages can then reconstruct a structured view of network behaviour from counters, ranks, versions, and forwarding statistics rather than from raw payload-level traces (Elrawy et al., 2018; Verma & Ranga, 2020; HUNSR, n.d.).

The resulting comparative CSV schema contains 29 columns in total: eight node-specific variables, twelve neighbor-derived variables, eight derived variables, and one ground-truth label column. The node-specific variables describe the local state claimed by the node itself; the neighbor-derived variables encode what surrounding nodes effectively report about that node through observed forwarding and control behavior; and the derived variables convert discrepancies and ratios into more detection-friendly indicators. This layered design is methodologically sound because many RPL attacks are not obvious from any single raw counter in isolation but become visible when local claims are contrasted with externally observed behaviour (HUNSR, n.d.; Elrawy et al., 2018).

Two derived indicators are especially important for understanding the methodological logic. The first is the rank-discrepancy family, which measures how far the node’s own advertised rank deviates from what neighbors imply. The second is the traffic-ratio family, which relates control activity and forwarding activity to data traffic. A DIS flooding attacker, for example, is expected to inflate control traffic disproportionately, while a selective forwarder is expected to distort forwarding ratios. In analytical form, these concepts can be written as:

$$\Delta\text{Rank} = \left| \text{Rank}_{\text{node}} - \text{AvgRank}_{\text{neighbors}} \right| \quad (3.1)$$

$$\text{Control-to-Data Ratio} = \frac{\text{Total Control Packets}}{\text{Total Data Packets}} \quad (3.2)$$

Ground-truth labelling is tied directly to controlled attack activation rather than to post hoc heuristic rules. When the logger turns on a malicious flag for a node, the node records that status

in the periodic behaviour log. This avoids circular labelling based on thresholds the IDS is later expected to learn (HUNSR, n.d.).

3.6 Complete feature schema used in the comparative CSV datasets

Table 3.4 lists the full 29-column schema extracted from the generated comparative CSV files. Using the actual column names from the uploaded artefacts is important because downstream analysis in Chapter 4 is performed on these exact fields rather than on a simplified conceptual subset (HUNSR, n.d.).

Table 3.4 Full feature inventory of the generated MRHOF/TACR datasets (HUNSR, n.d.)

No.	Column	Category	Description
1	time_sec	Node-specific	Timestamp associated with the reporting window.
2	node_id	Node-specific	Identifier of the reporting node.
3	parent_id	Node-specific	Identifier of the currently selected preferred parent.
4	rpl_ver	Node-specific	Current RPL version advertised by the node.
5	rpl_rank	Node-specific	Current rank value maintained by the node.
6	dis_sent	Node-specific	Number of DIS control messages sent.
7	dio_sent	Node-specific	Number of DIO control messages sent.
8	dao_sent	Node-specific	Number of DAO control messages sent.
9	nbr_dis_rcv	Neighbor-derived	Neighbor-observed DIS receptions attributed to the node.
10	nbr_dio_rcv	Neighbor-derived	Neighbor-observed DIO receptions attributed to the node.
11	nbr_dao_ack_rcv	Neighbor-derived	Neighbor-observed DAO acknowledgement receptions.
12	nbr_fwd_to_me	Neighbor-derived	Packets neighbors forwarded onward to this node.
13	nbr_fwd_to_others	Neighbor-derived	Packets neighbors observed the node forwarding to other relays.

Table 3.4 Full feature inventory of the generated MRHOF and TACR datasets (HUNSR, n.d.) (continued)

No.	Column	Category	Description
14	nbr_fwd_bcast	Neighbor-derived	Packets forwarded by the node toward broadcast destinations.
15	nbr_rpl_ctrl	Neighbor-derived	Neighbor-derived RPL control-traffic count associated with the node.
16	nbr_non_rpl_ctrl	Neighbor-derived	Neighbor-derived non-RPL control-traffic count associated with the node.
17	nbr_rpl_ver_rcv	Neighbor-derived	Neighbor-observed RPL version values received in DIO exchanges.
18	nbr_rpl_rank_rcv	Neighbor-derived	Neighbor-observed rank values received from the node.
19	nbr_fwd_rpl	Neighbor-derived	RPL packets sent by the node for onward forwarding.
20	nbr_fwd_non_rpl	Neighbor-derived	Non-RPL packets sent by the node for onward forwarding.
21	diff_rpl_rank	Derived	Absolute discrepancy between the node rank and the rank reflected by neighbors.
22	diff_rpl_ver	Derived	Absolute discrepancy between the node version and neighbor-observed version.
23	norm_rank_diff	Derived	Normalized version of rank discrepancy to reduce scale effects.
24	ctrl_to_data_ratio	Derived	Ratio of control traffic to application/data traffic.
25	non_rpl_to_rpl_ratio	Derived	Ratio of non-RPL control traffic to RPL control traffic.
26	rpl_fwd_ratio	Derived	Ratio of forwarded RPL traffic to data traffic.
27	non_rpl_fwd_ratio	Derived	Ratio of forwarded non-RPL traffic to data traffic.
28	total_fwd_ratio	Derived	Overall ratio of forwarded traffic to data traffic.

Table 3.4 Full feature inventory of the generated MRHOF and TACR datasets (HUNSR, n.d.) (continued)

No.	Column	Category	Description
29	label	Label	Ground-truth class label: 0 normal, 1 VNA, 2 DRA, 3 DISA, 4 SFA.

3.7 Routing configurations: MRHOF baseline versus TACR-HOF

The comparative contribution of this thesis does not stop at dataset generation. It also alters the routing layer itself through TACR-HOF, a trust-aware objective function implemented in `rpl-tachof.c`. From a methodological standpoint, TACR-HOF is designed as a disciplined extension rather than a completely separate routing protocol. It preserves MRHOF path-cost reasoning as the baseline cost term and only then injects penalties derived from IDS-maintained behavioural evidence. This decision is important because it lets the comparison isolate the value of trust-aware parent selection without discarding the path-quality logic already present in RPL-lite (Gnawali & Levis, 2012; HUNSR, n.d.).

The TACR code shows four principal sources of trust reduction. First, a version mismatch penalty is applied when a neighbor advertises a version inconsistent with the current node's view of the DODAG. Second, a suspicious-rank penalty is triggered when the neighbor appears artificially attractive. Third, a DIS-instability penalty is applied when DIS receptions become excessive relative to DIO evidence. Fourth, a forwarding-suspicion penalty is applied when traffic is repeatedly sent to a neighbor for onward forwarding but the node rarely appears to forward it onward. The trust score starts at 100, declines under these conditions, and causes the neighbor to be rejected completely when trust falls to or below 45. Otherwise, the residual distrust is converted into an additive path penalty scaled by the constant 8 (HUNSR, n.d.).

Table 3.5 Methodological comparison between the baseline MRHOF and the TACR-HOF routing configuration. Author’s synthesis based on (Gnawali & Levis 2012; da Costa et al. 2019; Rahman et al. 2025; Muzammal et al. 2022; HUNSR n.d.)

Design aspect	MRHOF baseline	TACR-HOF variant
Objective function	Standard RPL MRHOF path-cost logic based primarily on ETX/rank cost.	Retains MRHOF base cost, then augments it with IDS-derived trust penalties before parent selection.
Activation method	Baseline Contiki-NG configuration.	Enabled through RPL_OCP_TACRHOF and registration of rpl_tacrhof in the RPL-lite stack.
Evidence used for routing	Link-quality and rank/path metrics only.	Version consistency, suspicious rank attraction, DIS instability, and forwarding-ratio evidence from IDS counters.
Neighbor rejection rule	No trust-based rejection.	Neighbors with trust ≤ 45 are discarded as candidate parents.
Penalty scaling	Not applicable.	Penalty added as $8 \times (100 - \text{trust})$ when the trust threshold is not crossed.
Tuning constants visible in project-conf.h	Not applicable beyond baseline RPL defaults.	Max link metric 768, max path cost 32768, rank hysteresis threshold 96, time threshold 120 s, trust-drop threshold 45, penalty scale 8.

One additional methodological nuance must be stated explicitly. The broader repository and reference paper discuss an HMAC-based lightweight security mode (LSM-RPL). However, the uploaded comparative configuration files set CONF_LSM to 0. This means the present MRHOF-versus-TACR experiment isolates routing and IDS-related behavior without activating the optional lightweight security mode. That isolation is acceptable, and arguably desirable, for

the present thesis objective because it prevents security-layer protection from confounding the evaluation of trust-aware path selection itself (HUNSR, n.d.).

```

35 /* ---- TACR-HOF selection ---- */
36 #define RPL_CONF_OF_OCP RPL_OCP_TACRHOF
37 #define RPL_CONF_SUPPORTED_OFS {&rpl_tacrhof, &rpl_mrhof}
38 #define RPL_CONF_MIN_HOPRANKINC 128
39
40 /* ---- TACR-HOF tuning ---- */
41 #define RPL_TACRHOF_CONF_MAX_LINK_METRIC 768
42 #define RPL_TACRHOF_CONF_MAX_PATH_COST 32768
43 #define RPL_TACRHOF_CONF_RANK_THRESHOLD 96
44 #define RPL_TACRHOF_CONF_TIME_THRESHOLD_SEC 120
45 #define RPL_TACRHOF_CONF_TRUST_DROP_THRESHOLD 45
46 #define RPL_TACRHOF_CONF_TRUST_PENALTY_SCALE 8

```

Listing 3.1: TACR-HOF configuration parameters, project-conf.h lines 35–46

Listing 3.1 presents the main configuration parameters used to enable and tune TACR-HOF.

3.7.1 Trust penalties logic

```

124 static uint16_t
125 trust_penalty(rpl_nbr_t *nbr)
126 {
127     #if IDS != 1
128         (void)nbr;
129         return 0;
130     #else
131         int id = node_id_from_nbr(nbr);
132         int trust = 100;
133         unsigned short dis_rx;
134         unsigned short dio_rx;
135         unsigned short forwarded_from_nbr;
136         unsigned short sent_to_nbr;
137         unsigned short adv_version;
138         unsigned short adv_rank;
139
140         if(id < 0 || id >= MAX_NODES) {
141             return 0;
142         }
143
144         dis_rx = nodesinfo[id][1];
145         dio_rx = nodesinfo[id][2];
146         forwarded_from_nbr = nodesinfo[id][7] + nodesinfo[id][8];
147         sent_to_nbr = nodesinfo[id][11] + nodesinfo[id][12];
148         adv_version = nodesinfo[id][9];
149         adv_rank = nodesinfo[id][10];
150
151         /* Version inconsistency penalty */
152         if(currentnodeinfo[2] != 0 && adv_version != 0 && adv_version != currentnodeinfo[2]) {
153             trust -= 35;
154         }
155         /* Suspiciously attractive rank penalty */
156         if(currentnodeinfo[3] != 0 && adv_rank != 0 && adv_rank + curr_instance.min_hoprankinc <
157             currentnodeinfo[3]) {
158             trust -= 25;
159         }
160         /* DIS flooding / instability penalty */
161         if(dis_rx > dio_rx + 3) {
162             trust -= MIN(20, (int)(dis_rx - dio_rx));
163         }
164         /* Forwarding suspicion penalty.
165          * If we keep sending traffic to this neighbor for onward forwarding,
166          * but we rarely overhear it forwarding onward, penalize it.
167          */
168         if(sent_to_nbr >= 8) {
169             int ratio_percent = (int)((100UL * forwarded_from_nbr) / sent_to_nbr);
170             if(ratio_percent < 50) {
171                 trust -= 30;
172             } else if(ratio_percent < 70) {
173                 trust -= 15;
174             }
175         }
176         if(trust < 0) {
177             trust = 0;
178         }
179         if(trust <= TACR_TRUST_DROP_THRESHOLD) {
180             LOG_INFO("reject nbr=%d trust=%d dis=%u dio=%u sent=%u fwd=%u ver=%u rank=%u\n",
181                 id, trust, dis_rx, dio_rx, sent_to_nbr, forwarded_from_nbr,
182                 adv_version, adv_rank);
183             return TACR_MAX_PATH_COST;
184         }
185
186         return (uint16_t)((100 - trust) * TACR_TRUST_PENALTY_SCALE);
187     #endif
188 }
189

```

Listing 3.2: Computation of the TACR-HOF trust penalty for a candidate neighbor.
rpl-tacrhof.c lines 124–192

Listing 3.2 presents the implementation of the TACR-HOF trust-penalty function. The function initializes each candidate neighbor with a trust value of 100 and reduces this value when version inconsistency, suspicious rank advertisement, excessive DIS reception, or weak forwarding behavior is detected. A candidate whose trust value is less than or equal to the configured rejection threshold is assigned the maximum path cost. Otherwise, the remaining trust deficit is converted into an additive routing penalty.

3.8 Reproducibility, internal validity, and methodological limits

From a reproducibility perspective, the methodology has several strengths. The topology files are explicit, the random seed embedded in the .csc scenarios is fixed, the batch runner is available, the routing modification is visible in source code, and the final experiment bundle preserves both raw logs and processed CSV outputs. These artefacts allow the experiment to be repeated or extended without reconstructing the environment from prose alone. In practical terms, a future researcher could swap the objective function, introduce a new attack flag, or extend the feature parser while keeping the rest of the methodological scaffold unchanged (Oikonomou et al., 2022; Osterlind et al., 2006; HUNSR, n.d.).

The internal-validity picture is more mixed, and this is precisely why it should be stated here rather than ignored. Topology structure and simulation parameters are fixed, which is desirable. By contrast, attack scheduling is randomized at runtime, which improves generalization but weakens strict one-to-one trace matching between routing schemes unless identical schedules are externally replayed. Consequently, the comparison is best understood as a controlled comparative campaign under equivalent scenario design rather than as a byte-for-byte replay experiment. That limitation does not invalidate the results; it simply defines the appropriate scope of inference (HUNSR, n.d.).

A second methodological limit is that the study is based on simulation rather than physical nodes. Cooja provides excellent visibility and control, but it cannot capture every hardware-specific radio effect, clock drift nuance, or deployment irregularity found in real IoT installations. A

third limit is that only four attack classes are modeled. They are important RPL attacks, but they do not exhaust the adversarial landscape of low-power routing. Finally, the comparative routing evaluation uses the generated CSV outputs and summary classifier reports already present in the uploaded artefacts. This supports reproducibility, but it also means the thesis is bounded by the exact implementation choices embedded in those artefacts (Osterlind et al., 2006; Tsao et al., 2015; HUNSR, n.d.).

3.9 Chapter summary

In summary, the methodology combines protocol-level simulation, randomized RPL attack injection, behaviour-based dataset construction, and trust-aware routing modification into one coherent experimental framework. The study uses six fixed topologies, twelve principal comparative runs, a shared 29-column feature schema, and a routing-layer comparison between MRHOF and TACR-HOF that is explicit in source code and reproducible from the uploaded files. By grounding the method in actual artefacts rather than generic descriptions, the chapter establishes a defensible bridge to Chapter 4, where the numerical consequences of these design decisions are examined in detail.

CHAPTER 4

SIMULATIONS AND RESULTS

This chapter presents the empirical results of the comparative simulation campaign conducted to evaluate conventional RPL using the Minimum Rank with Hysteresis Objective Function (MRHOF) against the TACR-HOF variant in an intrusion-aware RPL environment. The objective of the chapter is not limited to reporting classifier scores. Instead, the chapter examines three tightly coupled issues: whether TACR-HOF produces behavioural traces that are easier for an intrusion detection system to classify, whether any improvement is consistent across multiple RPL attack classes, and how the routing layer itself behaves once trust-aware penalties are introduced into path selection.

All quantitative results in this chapter were derived from the uploaded comparative experiment outputs: the MRHOF and TACR CSV datasets, the corresponding simulation log files, and the generated confusion-matrix and performance plots. Across both routing schemes, the experiment covers six Cooja topologies and produces a combined total of 20,285 labelled time-window instances. The chapter is therefore written as an evidence-based interpretation of the actual simulation artefacts rather than as a generic restatement of the methodology.

Chapter highlights

- Twelve simulation runs were analysed: six topologies under MRHOF and the same six under TACR-HOF.
- The best IDS result was obtained with Random Forest on the TACR dataset, reaching 99.3% accuracy and 99.3% F1-score.
- Behaviour summaries are collected at the node level and converted into a shared CSV schema containing node-specific, neighbor-derived, and derived indicators.
- Random Forest total classification errors dropped from 71 under MRHOF to 22 under TACR; attack-to-normal misses fell from 46 to 12.

- Routing-level behaviour was mixed rather than uniformly dominant: mean scenario PDR improved and missed transmissions dropped under TACR, but aggregate weighted PDR did not improve in every topology.

4.1 Comparative experiment design and scenario coverage

The comparative campaign reused the same scenario family under two routing configurations. Six topologies were executed for each routing scheme, giving twelve simulation runs overall. The scenario family spans 10-node, 16-node, and 50-node deployments, and each size was instantiated twice: once with the root positioned near the corner and once with the root positioned near the centre. This design choice is methodologically important because root placement directly affects path diversity, hop count, convergence behaviour, and the number of descendants potentially exposed to a compromised forwarding node.

Attack scheduling was dynamic rather than statically attached to a fixed node identity. Inspection of the logs shows that each 10-node run contained 10 scheduled attack windows, each 16-node run contained 16, and each 50-node run contained 50. The observed attack durations remained inside the intended randomised band of approximately 5,000 to 10,000 seconds. In practical terms, this means the adversarial load scaled with network size instead of being artificially constant, which is a stronger experimental design because it keeps the stress level proportionate as the topology grows.

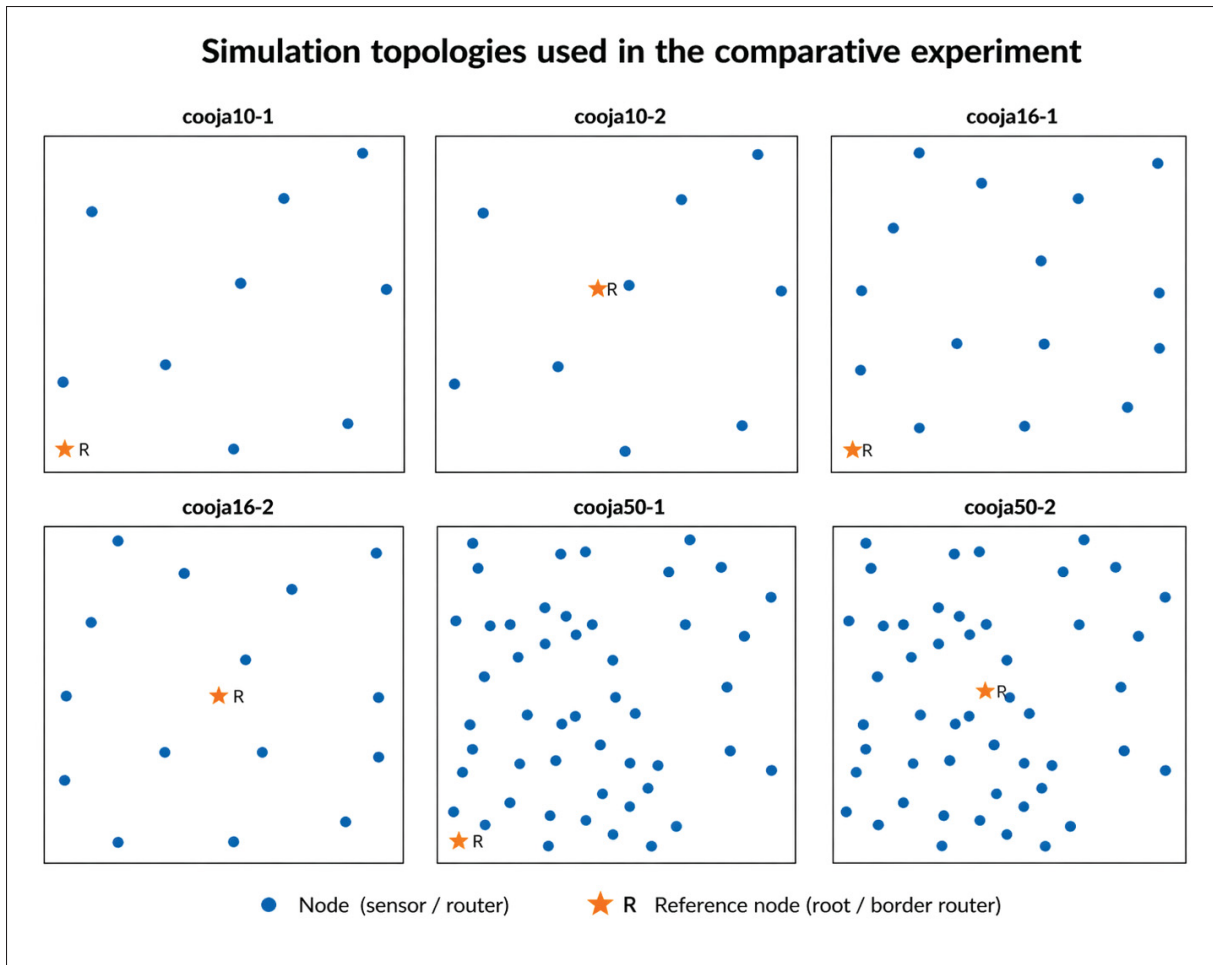


Figure 4.1 Topology layouts used in the comparative experiment. (Contiki-NG/Cooja GUI)

Table 4.1 Topology-level summary of the comparative simulation campaign

Topology	Total nodes	Root placement	Attack windows/run	Observed duration range (s)
cooja10-1	10	near corner	10	5030-9792
cooja10-2	10	near center	10	5438-9911
cooja16-1	16	near corner	16	5200-9917
cooja16-2	16	near center	16	5054-9938
cooja50-1	50	near corner	50	5044-9971
cooja50-2	50	near center	50	5023-9870

Figure 4.1 shows that the six topologies are not trivial perturbations of the same layout. In the corner-root cases, traffic tends to be funnelled through a more directional forwarding structure, whereas in the centre-root cases the forwarding tree has more symmetry and usually offers more parent alternatives. The 50-node layouts are particularly important because they create a denser decision space in which malicious behaviour can interact with route selection more subtly than in the smaller scenarios. From an experimental-rigor standpoint, the scenario coverage is sufficient to support comparative claims about pattern stability. The results discussed later in the chapter therefore reflect repeated behaviour across multiple scales and root placements rather than a single best-case simulation.

4.2 Dataset composition and class balance

Processing the twelve simulation logs produced two ready-to-train datasets with identical schema: 28 behavioural features and one class label. The MRHOF dataset contains 10,087 labelled instances, while the TACR dataset contains 10,198 instances. This near-equal size is advantageous because it reduces the risk that performance differences are merely a consequence of one dataset being much larger than the other. A more consequential difference appears in the class distribution. Under MRHOF, normal traffic dominates 86.59% of all instances, leaving

only 13.41% attack-labelled windows. Under TACR, the normal share falls to 77.82% and the attack share rises to 22.18%. This indicates that TACR changes the observable behavioural profile of the network in a way that exposes a greater number of windows as distinguishably malicious rather than allowing them to remain hidden inside normal-looking traffic patterns. For IDS training and evaluation, that shift is meaningful because it reduces imbalance and increases attack visibility in the feature space.

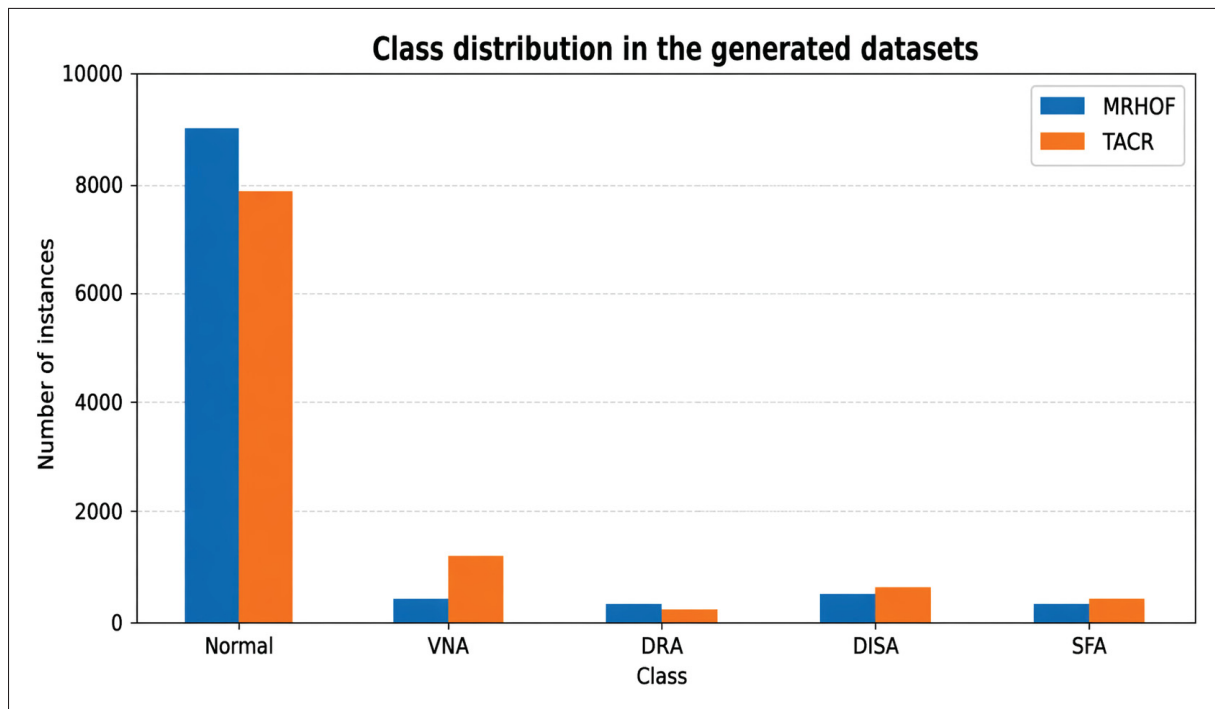


Figure 4.2 Class distribution of the generated MRHOF and TACR datasets

Table 4.2 Dataset composition and class balance

Dataset	Instances	Features	Normal	VNA	DRA	DISA	SFA	Attack Ratio %	Normal Ratio %
MRHOF	10087	28	8734	368	260	473	252	13.41	86.59
TACR	10198	28	7936	1178	184	568	332	22.18	77.82

The largest class shift occurs in Version Number Attack (VNA) instances, which increase from 368 windows under MRHOF to 1,178 under TACR. DISA and SFA also increase moderately,

while DRA windows become somewhat fewer. This pattern should not be interpreted as a simple increase or decrease in attack severity; instead, it suggests that modifying the routing objective function changes which attack manifestations become salient in the summarised behaviour windows. In other words, the routing layer influences what the IDS eventually learns from. The practical implication is that dataset composition is itself part of the result. A trust-aware routing variant does not merely affect delivery paths; it can reshape the empirical distribution of normal and malicious behaviour presented to downstream learning models.

4.3 Classifier-level performance comparison

To test whether the behavioural differences between the two routing schemes translate into improved detection capability, three conventional machine-learning classifiers were evaluated on both datasets: Random Forest, Decision Tree, and K-Nearest Neighbours (KNN). Using more than one learner is important for defensibility, because an improvement observed only under one highly specific model configuration could be dismissed as a modelling artefact rather than a genuine property of the dataset.

The results demonstrate a consistent advantage for TACR across all three classifiers. Random Forest remains the strongest model in both settings, but the TACR-generated dataset raises the performance ceiling further. Random Forest accuracy increases from 97.7% under MRHOF to 99.3% under TACR, with F1-score rising from 97.6% to 99.3%. Decision Tree improves from 96.1% to 98.9%, and KNN improves from 93.7% to 97.8%. The fact that the gain is not restricted to one classifier family strongly supports the interpretation that TACR produces a more separable and internally coherent feature distribution.

Table 4.3 Overall IDS performance on the MRHOF and TACR datasets

Classifier	MRHOF Acc. %	MRHOF F1 %	TACR Acc. %	TACR F1 %	Δ Acc. (pp)	Δ F1 (pp)
Random Forest	97.70	97.60	99.30	99.30	1.60	1.70
Decision Tree	96.10	96.10	98.90	98.90	2.80	2.80
KNN	93.70	93.20	97.80	97.80	4.10	4.60

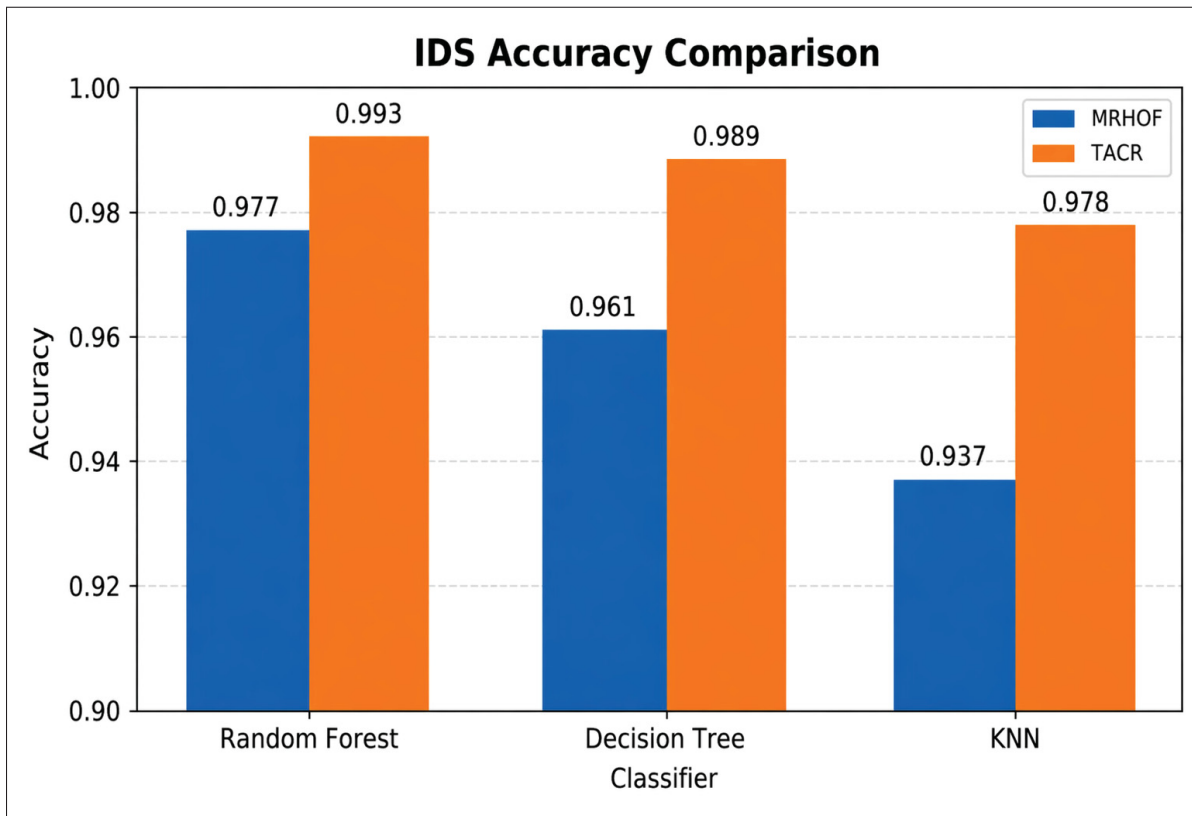


Figure 4.3 Accuracy comparison for the three evaluated classifiers

The magnitude of the improvement is also informative. In absolute terms, TACR improves accuracy by 1.6 percentage points for Random Forest, 2.8 points for Decision Tree, and 4.1 points for KNN. The largest gain in KNN is particularly revealing because KNN depends heavily on local neighbourhood structure in the feature space. A stronger KNN result therefore suggests

that TACR does not simply help a tree-based learner memorise splits; it appears to make normal and malicious windows more locally compact and more clearly separated from one another. For the best-performing Random Forest model, the increase from 97.7% to 99.3% corresponds to an approximately 69.6% reduction in residual classification error. From a thesis-defence perspective, that is a substantial effect size rather than a marginal fluctuation around the same operating point.

4.4 Attack-wise analysis using the Random Forest classifier

Because Random Forest produced the best overall result in both routing conditions, it was selected for detailed class-wise analysis. The confusion matrices show that the dominant error mode under MRHOF is not attack-to-attack confusion; rather, it is the absorption of malicious windows into the Normal class. This is precisely the failure mode that is most problematic in a deployed IDS because it represents missed detections rather than merely incorrect attack naming. Under MRHOF, the Random Forest model produces 71 total errors. Of these, 46 are attack windows misclassified as Normal, and 23 are normal windows incorrectly flagged as attack. Under TACR, total errors drop to only 22, with attack-to-normal misses reduced to 12 and normal-to-attack false alarms reduced to 10. Thus, TACR improves the model in the two ways that matter operationally: it lowers missed detections and it lowers false alarms.

Table 4.4 Random Forest class-wise precision, recall, and F1-score

Label	Prec. MRHOF %	Rec. MRHOF %	F1 MRHOF %	Prec. TACR %	Rec. TACR %	F1 TACR %	Δ Recall (pp)
Normal	98.25	99.12	98.68	99.50	99.58	99.54	0.47
VNA	93.14	88.79	90.91	98.60	99.15	98.88	10.37
DRA	90.91	84.34	87.50	100.00	96.49	98.21	12.15
DISA	95.77	91.28	93.47	98.71	98.71	98.71	7.43
SFA	93.59	87.95	90.68	96.91	94.95	95.92	7.00

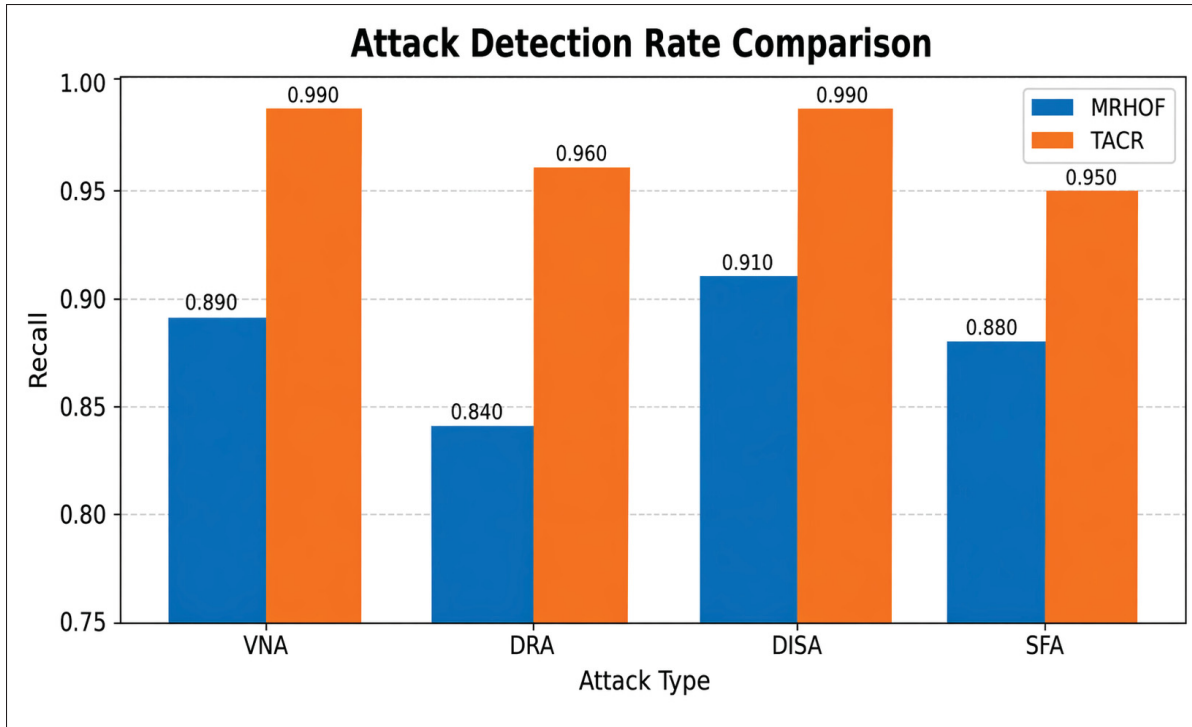


Figure 4.4 Attack-wise recall comparison for the Random Forest classifier

Attack-specific recall improves for every malicious class. VNA recall rises from 88.79% to 99.15%, DRA from 84.34% to 96.49%, DISA from 91.28% to 98.71%, and SFA from 87.95% to 94.95%. The largest improvement is observed for DRA, followed by VNA. This is analytically plausible: rank and version manipulations act directly on the control-plane semantics that influence parent choice and DODAG consistency, so a trust-aware objective function has a stronger chance of making those abnormalities more visible in the derived behavioural features.

DISA was already relatively detectable under MRHOF because excessive DIS activity disturbs control-message statistics in a conspicuous way. Nevertheless, TACR still pushes DISA recall close to the near-perfect range. SFA remains the most challenging malicious class even after improvement. That result is not surprising. Selective forwarding is intentionally stealthier than flood-based or control-field manipulation attacks because the attacker may preserve some control behaviour while selectively discarding only portions of the data plane. The persistence of some

SFA ambiguity therefore strengthens the credibility of the results; it shows that the experiment is detecting a genuinely difficult class rather than an artificially obvious one.

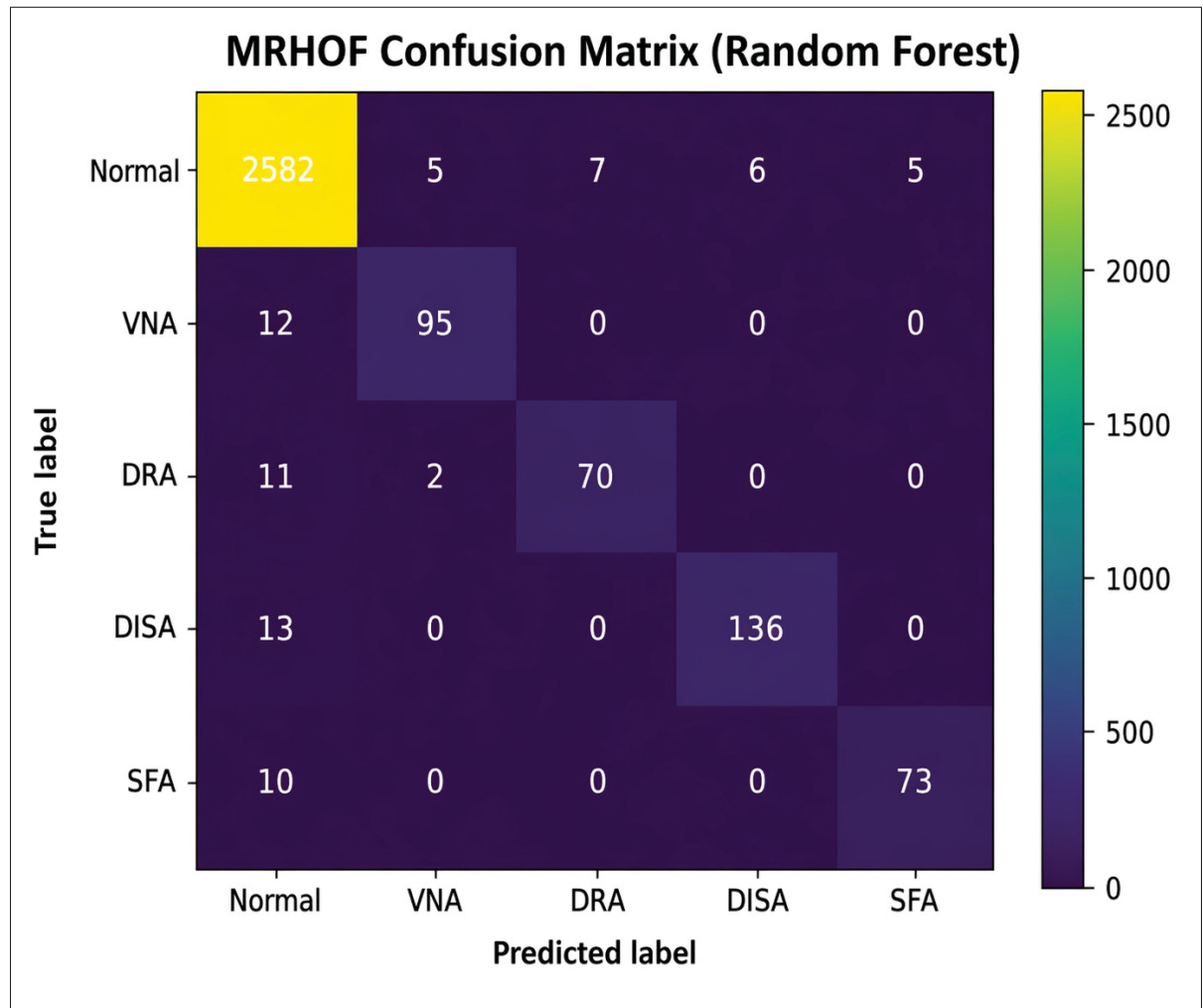


Figure 4.5 Confusion matrix of the Random Forest classifier on the MRHOF dataset

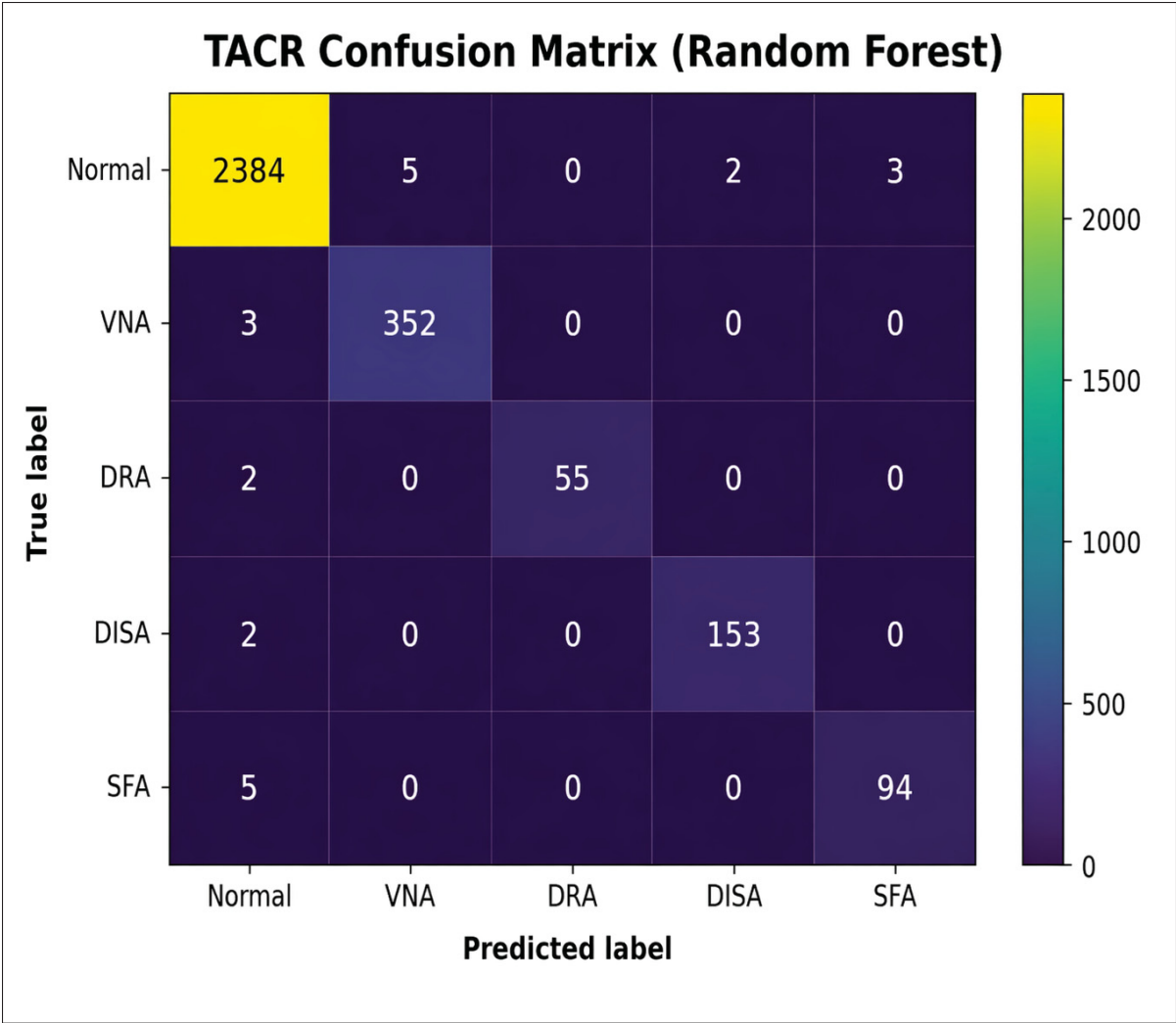


Figure 4.6 Confusion matrix of the Random Forest classifier on the TACR dataset

A visual comparison of Figures 4.5 and 4.6 confirms the numerical interpretation. In the TACR matrix, the diagonal becomes noticeably cleaner and the off-diagonal leakage into the Normal class is sharply reduced. The improvement is especially clear for VNA and DRA, where the TACR matrix retains almost all malicious windows in their correct classes.

4.5 Routing-level observations from the raw simulation logs

Superior IDS scores do not automatically imply uniformly superior routing behaviour. For that reason, the raw simulation logs were examined independently using packet-level counters extracted from the six MRHOF and six TACR test logs. The purpose of this additional analysis is to determine whether the trust-aware routing prototype also changes the underlying forwarding dynamics in a consistent way.

When results are averaged per topology, TACR increases mean packet delivery ratio (PDR) from 0.761 to 0.833 and reduces mean missed transmissions from 255.7 to 69.7. TACR outperforms MRHOF in four of the six topology files with respect to PDR: cooja10-1, cooja10-2, cooja16-2, and cooja50-2. However, the advantage is not universal. MRHOF remains better in cooja16-1 and cooja50-1, which indicates that the routing response to trust penalties is topology dependent rather than monotonic.

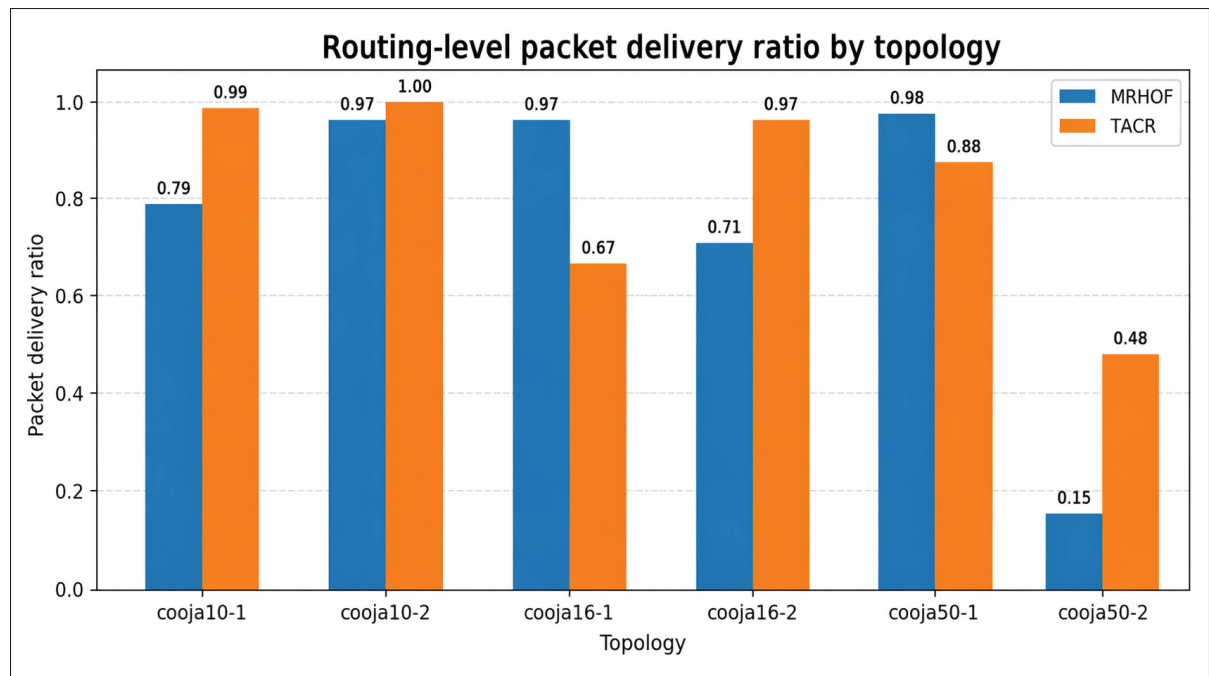


Figure 4.7 Packet delivery ratio by topology for MRHOF and TACR

Table 4.5 Routing-level log summary

Topology	PDR (MRHOF)	PDR (TACR)	MissedTx M	MissedTx T	Trust rejects
cooja10-1	0.794	0.990	268	46	56,476
cooja10-2	0.971	1.000	37	0	112,498
cooja16-1	0.965	0.672	178	5	78,099
cooja16-2	0.707	0.972	85	0	144,535
cooja50-1	0.977	0.884	56	124	127,311
cooja50-2	0.150	0.480	910	243	171,287
Mean	0.761	0.833	255.7	69.7	115,034

A more cautious interpretation emerges when the data are aggregated over total transmitted packets rather than averaged per scenario. Under that weighted view, overall PDR is 0.883 for MRHOF and 0.867 for TACR. The difference arises because the two schemes do not generate the same offered traffic volume in every topology. TACR logs also contain large numbers of trust-rejection events, confirming that the prototype aggressively filters suspicious neighbours. Such filtering can improve attack visibility and break dependence on questionable parents, but it can also reduce forwarding opportunities or temporarily destabilise traffic volume in dense regions.

This mixed routing result is scientifically valuable. It prevents an over-claim. The evidence supports a strong statement that TACR-HOF substantially improves IDS-oriented separability and often reduces retransmission waste, but it does not justify the broader claim that TACR universally improves raw network delivery performance in every topology. Additional tuning of trust thresholds and parent-switching aggressiveness would likely be required before making such a claim.

4.6 Synthesis of findings

Taken together, the results lead to five principal findings. First, TACR-HOF consistently produces behaviour traces that are more learnable for an intrusion detection system than those generated under baseline MRHOF. This conclusion is supported not only by the best model score, but by the fact that all three evaluated classifiers improve on the TACR dataset.

Second, the improvement is attack-wide rather than attack-specific. Every malicious class gains recall under TACR, and the improvement is strongest for the control-plane attacks that most directly manipulate RPL routing semantics. This indicates that the trust-aware routing logic enhances the observability of malicious deviations rather than merely shifting confusion from one attack label to another.

Third, TACR substantially reduces the most dangerous type of IDS error: missed attacks that appear normal. The reduction from 46 attack-to-normal misses to 12 is one of the most persuasive pieces of evidence in the entire chapter because it directly addresses the operational value of the proposed approach.

Fourth, the routing scheme influences dataset composition itself. The TACR dataset is less imbalanced and contains a higher proportion of attack-labelled windows. For future research, this means the routing layer should be treated as an active contributor to dataset generation, not merely as a background transport mechanism.

Fifth, routing-level performance remains nuanced. TACR often improves scenario-level PDR and sharply lowers missed transmissions, yet it also exhibits topology-dependent trade-offs and aggressive neighbour rejection. This nuance increases the credibility of the chapter because it shows that the evaluation distinguishes between detection quality and forwarding efficiency rather than collapsing them into a single undifferentiated claim of superiority.

4.7 Chapter conclusion

The results presented in this chapter demonstrate that the TACR-HOF variant yields a materially stronger foundation for intrusion detection in RPL-based IoT simulations than baseline MRHOF. The strongest result, achieved by Random Forest, reaches 99.3% accuracy and 99.3% F1-score on the TACR dataset, while class-wise analysis shows that the gain extends across VNA, DRA, DISA, and SFA rather than being confined to one attack type. At the same time, the routing-log analysis shows that trust-aware routing introduces measurable trade-offs and remains sensitive to topology, particularly in larger or more stressed layouts.

Accordingly, the central conclusion of this chapter is twofold: TACR-HOF is highly effective as a security-oriented routing modification for producing attack-discriminative behavioural data, but further refinement is still required before it can be claimed as a universally superior routing policy in purely network-performance terms. This dual conclusion is appropriate for a defensible thesis because it is both strongly supported by the data and appropriately bounded by the limits of the observed experiment.

CHAPTER 5

CONCLUSION AND FUTURE WORK

This chapter closes the thesis by integrating the theoretical argument developed in Chapter 2, the experimental methodology established in Chapter 3, and the empirical evidence reported in Chapter 4 into one final, defensible conclusion. Its purpose is not merely to restate earlier results, but to answer the central research problem in a disciplined way, identify the precise value of the TACR-HOF approach, and define the open work required for post-thesis enhancement. A strong conclusion chapter must therefore do three things simultaneously: synthesize what has been proven, state clearly what has not yet been proven, and convert those remaining gaps into a technically coherent future-work agenda.

Viewed in that light, the thesis has established a bounded but meaningful result. The data do not support the simplistic claim that trust-aware routing is universally superior to conventional RPL in every routing circumstance. However, they do strongly support a more important and more defensible conclusion for this research domain: when the evaluation target is security-oriented behavioural discrimination for intrusion detection, TACR-HOF provides a measurably stronger foundation than baseline MRHOF. That conclusion is supported by class-balance shifts, higher classifier accuracy, higher malicious-class recall, and a sharp reduction in the most dangerous IDS error category, namely malicious behaviour that is incorrectly interpreted as normal.

Chapter highlights

- The thesis confirms that TACR-HOF is materially stronger than MRHOF as a routing basis for behaviour-based IDS in RPL-based IoT simulations.
- The strongest empirical evidence is the shift from 97.7% to 99.3% Random Forest accuracy, combined with the reduction of attack-to-normal errors from 46 to 12.
- The routing evidence remains nuanced: mean scenario-level PDR improves under TACR-HOF, but weighted aggregate PDR does not dominate MRHOF in every case.
- The defensible conclusion is therefore a security-oriented superiority claim, not an unconditional routing-optimality claim.

- The most important open work lies in trust calibration, broader validation, temporal IDS modeling, and hardware-grounded deployment assessment.

5.1 Final answer to the research problem

The research problem posed in Chapter 1 asked whether a trust-aware RPL objective function could produce a stronger basis for behaviour-based intrusion detection than conventional MRHOF under constrained IoT conditions. The final answer is yes, but with an important qualification. TACR-HOF improves the discriminative quality of the behavioural traces generated by the network and thereby improves downstream IDS performance in a way that is substantial, repeatable across attack classes, and meaningful for operational security. At the same time, the routing-level evidence shows that this advantage is not equivalent to universal superiority in every traffic scenario or every topology.

This distinction matters for thesis defensibility. A weak conclusion chapter would collapse those two claims and overstate the evidence. A stronger conclusion chapter draws the line precisely. TACR-HOF should be understood as a security-oriented routing enhancement whose main value lies in changing the quality of observed node behaviour, neighbour trust relationships, and attack visibility. The contribution is therefore located at the intersection of routing and detection. The protocol does not merely forward packets differently; it reshapes the empirical conditions under which an IDS must decide whether behaviour is benign or malicious. This positioning is essential because the novelty lies in the MRHOF-versus-TACR-HOF evidence-quality comparison, not in claiming that trust-aware RPL itself is an unexplored concept.

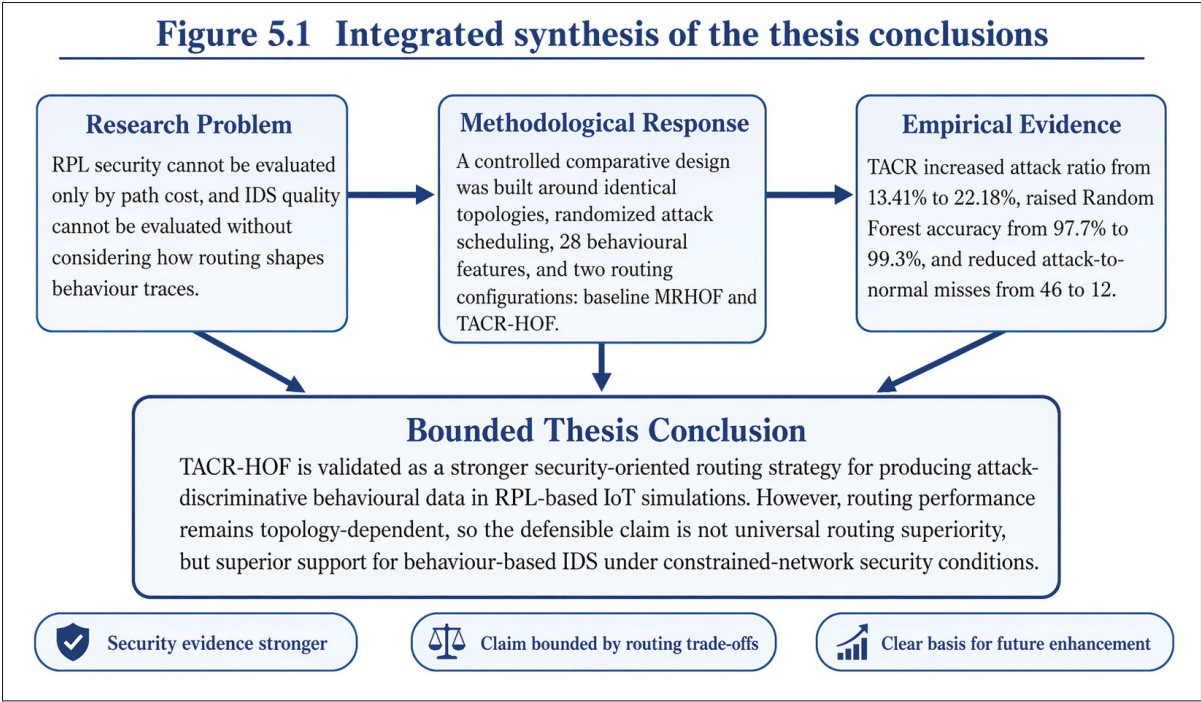


Figure 5.1 Integrated synthesis of the thesis conclusions

5.2 Quantitative synthesis of the principal findings

The numerical evidence accumulated in Chapter 4 can be condensed into a small set of high-value indicators. These indicators are sufficient to justify the final claim because they span dataset composition, classifier-level accuracy, attack-wise recall, routing-level traffic behaviour, and error severity. Table 5.1 brings these dimensions together in one comparative synthesis.

Table 5.1 Quantitative synthesis of the principal results

Dimension	MRHOF	TACR-HOF	Interpretive meaning
Attack-labelled data share	13.41%	22.18%	TACR exposes a richer malicious-behaviour footprint and reduces extreme normal-class dominance.
Best classifier result (Random Forest)	97.7% accuracy; 97.6% F1	99.3% accuracy; 99.3% F1	The strongest overall IDS model becomes more reliable under TACR-generated traces.
Malicious-class recall gains	VNA 88.79%; DRA 84.34%; DISA 91.28%; SFA 87.95%	VNA 99.15%; DRA 96.49%; DISA 98.71%; SFA 94.95%	Improvement is attack-wide rather than confined to a single malicious class.
Attack-to-normal misclassifications	46	12	The most operationally dangerous IDS error is reduced sharply under TACR.
Mean scenario-level PDR	0.761	0.833	On average across topologies, TACR improves successful delivery behaviour.
Weighted overall PDR	0.883	0.867	Aggregate routing performance remains nuanced and prevents a claim of unconditional routing dominance.
Mean missed transmissions	255.7	69.7	TACR materially lowers missed-transmission burden in the average scenario.

The synthesis in Table 5.1 shows why the thesis conclusion is credible. The evidence does not depend on one isolated metric or one unusually favourable topology. Instead, multiple indicators

move in the same direction. The attack classes become more detectable, the strongest classifier improves further, and the most serious false-normal errors fall substantially. The routing metrics are more mixed, but even that mixed result strengthens rather than weakens the credibility of the thesis because it demonstrates disciplined interpretation instead of selective reporting.

5.3 Resolution of the research objectives and questions

A conclusion chapter should also close the logical commitments made in Chapter 1. The thesis introduced a set of objectives and working questions concerning protocol-aware behavioural representation, comparative IDS performance, routing-level implications, and the realism of a bounded claim. Table 5.2 resolves those commitments explicitly.

Table 5.2 Closure of the thesis objectives and research questions

Research objective or question	Evidence base	Final answer	Support level
Can behaviour-based features extracted from RPL logs represent multiclass attack activity in a defensible way?	Chapter 3 defined a 28-feature schema extracted consistently from all comparative runs; Chapter 4 showed separable class behaviour and high multiclass scores.	Yes. The feature space is sufficiently expressive for multiclass IDS evaluation.	Strong
Does TACR-HOF generate a stronger IDS foundation than baseline MRHOF?	Accuracy, F1, and attack recall improved across all three classifiers, with Random Forest reaching 99.3% accuracy under TACR.	Yes. The evidence strongly supports a TACR advantage for IDS-oriented evaluation.	Strong
Does TACR-HOF eliminate routing trade-offs and dominate MRHOF in every network metric?	Scenario-level PDR and missed transmissions often improved, but weighted overall PDR remained slightly lower under TACR and some topologies showed regressions.	No. The advantage is bounded and does not justify universal routing-optimality claims.	Strong negative support
Is a combined routing-and-IDS interpretation more meaningful than evaluating each layer in isolation?	The thesis showed that routing choice altered class balance, attack visibility, and classifier separability, which would be missed by a layer-isolated analysis.	Yes. Integrated evaluation is a key conceptual contribution of the thesis.	Strong

The key point is that every major question posed at the beginning of the thesis can now be answered without ambiguity. The only negative answer in Table 5.2 is also the one that preserves scientific discipline: TACR-HOF should not be portrayed as a universal routing solution. By stating this explicitly, the thesis remains both stronger and more credible.

5.4 Scientific, methodological, and practical contributions

The contribution of the thesis can be read on three levels, but its novelty must be stated precisely. The thesis does not claim that TACR-HOF is the first trust-aware RPL mechanism, nor does it claim that trust-aware routing is universally superior to conventional RPL. Its conceptual contribution is to show that RPL security should be treated as a coupled routing-and-observation problem. The behavioural traces presented to the IDS are not independent of the routing policy that produced them. Therefore, the main contribution is the controlled demonstration that changing the objective function from MRHOF to TACR-HOF can improve the visibility, separability, and learnability of malicious behaviour for behaviour-based intrusion detection.

The second contribution is methodological. The study built a controlled comparative pipeline that begins with fixed Cooja topologies, applies randomized attack scheduling, extracts a consistent 28-feature behavioural schema, and evaluates multiple machine-learning classifiers under identical conditions. This is stronger than an ad hoc experiment because it enables direct attribution: when the detection outcomes change, the explanation can be tied to the routing regime rather than to uncontrolled shifts in data structure.

The third contribution is practical and engineering-oriented. TACR-HOF provides a concrete prototype showing how trust penalties derived from behavioural evidence can be integrated into parent selection without discarding the routing logic of MRHOF altogether. In that sense, the thesis contributes more than an IDS result; it contributes an implementable design direction for security-aware routing in constrained IoT environments.

Taken together, these contributions position the thesis as a bounded but original study of security-oriented routing for RPL-based IoT networks. The originality is not located in any

single isolated element, such as RPL, trust, IDS, machine learning, or simulation. Rather, it lies in the integrated comparative design: the routing objective function is treated as the experimental variable, and IDS evidence quality is treated as a primary evaluation outcome.

5.5 Limitations and validity boundaries

No defensible thesis is complete without a precise statement of its limitations. The purpose of limitations is not to weaken the result, but to define the boundary inside which the result is valid. In this study, those boundaries are clear. The evaluation was performed in simulation, not on physical nodes; the topology family, while varied, was finite; the attack catalogue covered four important RPL-centric attacks rather than the full space of adversarial behaviours; and the IDS stage was evaluated offline using conventional supervised classifiers rather than as an online adaptive detector. These limitations do not invalidate the findings, but they do define the next layer of work required for generalization.

Accordingly, the results should be interpreted as controlled comparative evidence for the TACR-HOF concept, while hardware or hybrid testbed validation remains necessary before making deployment-level claims.

Table 5.3 Principal limitations of the study and their implications

Limitation	Why it matters	Impact on interpretation	Required next step
Simulation-based evaluation	Cooja offers control and observability, but not complete radio-hardware realism.	Results are valid as controlled comparative evidence, not as a final field-deployment guarantee.	Validate on physical or hybrid testbeds.
Finite topology family	Six topologies capture structural diversity but cannot exhaust network-shape variability.	Conclusions are strong across the tested scenario family, but not mathematically universal.	Expand seeds, densities, and placement regimes.
Four attack classes	The studied attacks are important RPL threats, yet they do not represent every malicious strategy or every benign failure mode.	The IDS advantage is proven for the evaluated attack set only.	Add additional attacks and benign-fault controls.
Offline supervised learning	Models were trained and assessed after dataset extraction rather than under live adaptive conditions.	Strong classification evidence does not automatically imply real-time deployment readiness.	Evaluate streaming or online detection.
Incomplete parameter tuning	TACR thresholds and penalties were validated as a working prototype but not exhaustively optimized.	Observed trade-offs may partly reflect calibration choices rather than the conceptual limit of trust-aware routing.	Perform sensitivity analysis and parameter search.

Among these limitations, calibration is especially important. The TACR logs reveal a large number of trust-rejection events, which confirms that the mechanism is active and security-

sensitive. However, an active trust filter can produce both benefits and side effects. When the threshold is too conservative, suspicious neighbours remain attractive longer than they should. When it is too aggressive, forwarding opportunities can be reduced prematurely. Accordingly, one of the most important post-thesis tasks is not to abandon TACR-HOF, but to calibrate it more carefully so that the security gains demonstrated here are retained while routing penalties are reduced.

5.6 Broader implications of the findings

The implications of the thesis extend beyond the specific TACR prototype. The most general implication is that routing policy should be treated as part of the IDS problem formulation in constrained networks. Many IoT security studies implicitly assume that data collection is neutral and that routing merely transports packets. The present thesis shows that this assumption is incomplete. In RPL environments, the objective function shapes which neighbour interactions occur, how often control inconsistencies surface, and how clearly attacks are represented in the resulting behaviour logs (Canbalaban & Şen 2020).

A second implication concerns evaluation practice. Security-oriented routing proposals should not be evaluated only by packet delivery ratio, delay, or energy cost, just as IDS models should not be evaluated only by accuracy without reference to the protocol conditions that produced the data. The stronger approach is joint evaluation. Such a framework allows researchers to determine whether a routing change merely shifts traffic, or whether it actually improves the observability and separability of malicious behaviour.

A third implication is practical. In real constrained deployments, the best defensive solution may not be a single heavyweight IDS placed above the network stack. A more realistic design may instead combine lightweight, protocol-aware evidence collection with routing decisions that gradually penalize suspicious neighbours. The TACR-HOF results do not yet complete that engineering program, but they clearly justify it.

5.7 Open work and post-thesis enhancements

The most constructive use of the present thesis is as a validated baseline for a next-stage research and engineering program. That program should proceed through staged enhancement rather than uncontrolled feature addition. The goal is not to make TACR-HOF more complicated for its own sake, but to preserve the demonstrated attack sensitivity while improving robustness, generality, and deployment realism.

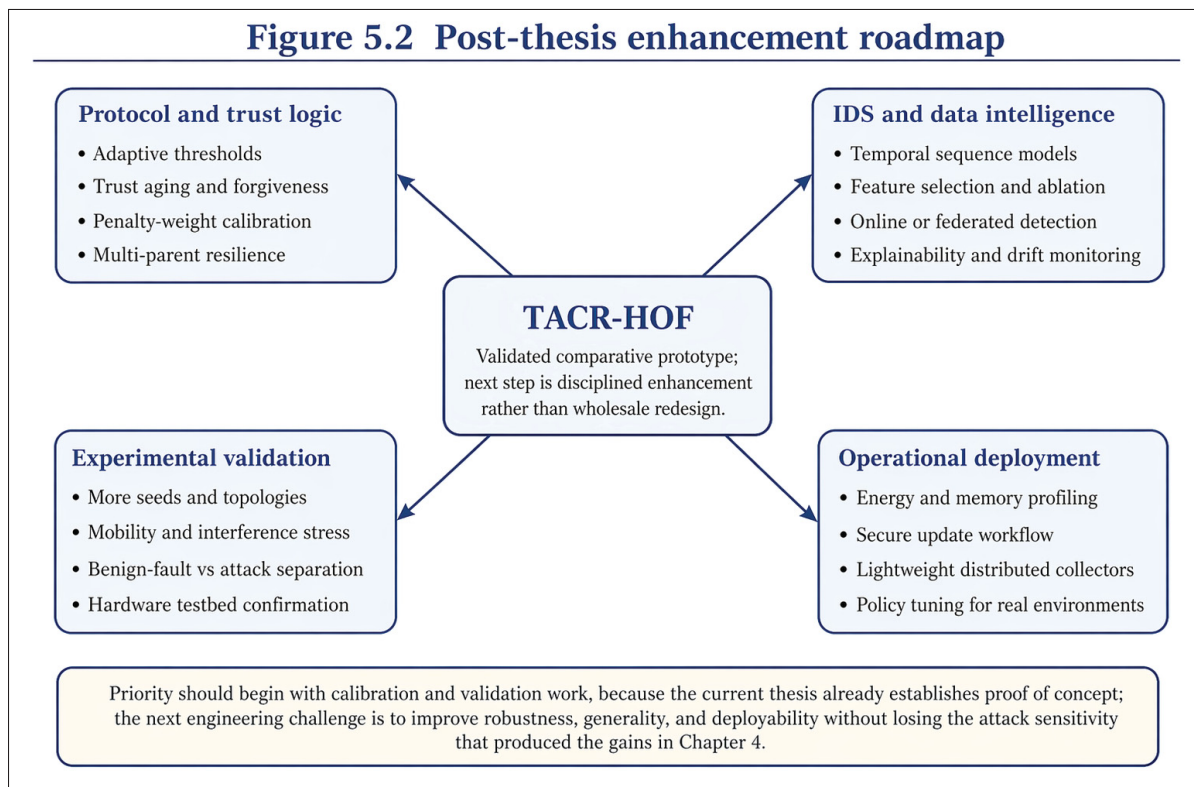


Figure 5.2 Post-thesis enhancement roadmap organized around protocol, IDS, validation, and deployment workstreams

5.7.1 Protocol and trust-model enhancements

The first workstream concerns the trust mechanism itself. A natural enhancement is adaptive thresholding. In the present prototype, trust penalties and rejection thresholds behave as global design choices. A more mature implementation could adapt those thresholds to local congestion,

neighbourhood density, historical link stability, or attack persistence. This would help distinguish transient anomalies from sustained malicious behaviour.

A second enhancement is trust aging and forgiveness. Not all suspicious observations should retain equal weight indefinitely. A neighbour that briefly appears inconsistent and then returns to stable behaviour should not remain permanently disadvantaged. Introducing decay functions, temporal smoothing, or staged recovery logic could reduce unnecessary parent exclusion and stabilize routes without erasing the security memory needed to respond to recurring attacks.

A third enhancement is multi-parent resilience and route diversity control. TACR-HOF currently influences parent choice primarily through trust-penalized cost comparison. Future work could explicitly examine how many trusted backup parents should be maintained, how quickly failover should occur, and whether diversity constraints can reduce overdependence on any one path while still controlling overhead.

5.7.2 IDS and data-analysis enhancements

The IDS layer also presents substantial open work. The present thesis deliberately used conventional classifiers such as Random Forest, Decision Tree, and KNN because they provide interpretable, established baselines for comparative validation. The next step is not necessarily to replace them blindly, but to extend the modeling space. Sequence-aware models, temporal windows, graph-informed representations, and online learning strategies could all be tested to determine whether TACR-generated behaviour remains advantageous under more demanding detection paradigms.

Feature-engineering refinement is another major opportunity. Although the 28-feature schema proved effective, not every feature may contribute equally to attack discrimination. Ablation studies, feature-importance ranking, dimensionality reduction, and robustness testing under noisy logs would clarify which behavioural signals are indispensable and which can be simplified for lower-overhead deployment.

Explainability should also be strengthened. In a field-deployable system, it is not enough to know that a node has been flagged as suspicious; operators and higher-layer control logic should also understand whether that decision arose from version inconsistency, rank manipulation, DIS flooding behaviour, selective forwarding symptoms, or a more complex combination. This is especially relevant if future work integrates routing response and IDS output more tightly.

5.7.3 Experimental and validation enhancements

The current experiment already demonstrates comparative value, but broader validation is essential for generalization. Additional Cooja seeds, denser deployment patterns, different root placements, varying traffic rates, mobility assumptions, and radio-interference conditions should all be introduced systematically. This would show whether the TACR advantage is structurally stable or whether it depends strongly on a narrow subset of network conditions.

Future campaigns should also separate malicious behaviour from benign operational faults more explicitly. In low-power networks, packet loss, congestion, or transient rank changes can arise from non-malicious causes. An advanced study should therefore include benign stress scenarios so that the IDS and the trust mechanism are tested not only on attack discrimination, but also on their ability to avoid over-penalizing ordinary network turbulence.

Finally, hardware-grounded confirmation is highly desirable. A physical testbed, even if smaller than the simulated network family, would help assess timing jitter, memory pressure, radio asymmetry, and energy consequences that cannot be fully represented by simulation alone.

5.7.4 Deployment-oriented enhancements

A post-thesis engineering track should also measure the operational cost of TACR-HOF explicitly. Memory usage, processing overhead, extra control-state maintenance, and energy impact should be profiled. Such measurements are indispensable in constrained environments because a security mechanism that improves detection but overwhelms node resources may not remain viable in practice.

Another practical enhancement is distributed evidence collection. Rather than assuming one centralized log-processing path, future work could explore hierarchical or federated IDS architectures in which nodes or cluster heads maintain partial trust evidence locally and exchange only compact summaries. This would align more naturally with the decentralized character of many IoT deployments.

Finally, deployment readiness requires secure update and policy-governance mechanisms. If trust thresholds, penalty weights, or attack-signature logic are refined after deployment, the system must support authenticated configuration updates and stable rollback procedures. In other words, the path from thesis prototype to deployable solution is not only an algorithmic path, but also an operational one.

For deployment-oriented evolution, the framework should be positioned within a broader risk-management and governance context. Because the current approach relies on behavioural summaries rather than payload inspection, it is already amenable to data-minimizing collection strategies; future deployments should additionally pseudonymize node identifiers, minimize retention windows, and export only compact counters required for routing trust and IDS inference. At the organizational level, TACR-HOF is best interpreted as a technical mechanism that could support a wider security-control and ISMS program, rather than as a standalone claim of compliance. A future deployment study should therefore provide a formal crosswalk from routing hardening, anomaly evidence, logging, and trust response to the relevant organizational control and risk-treatment requirements

Table 5.4 Recommended post-thesis roadmap for enhancement and validation

Phase	Priority	Core actions	Expected benefit	Concrete output
Phase I: Calibration	Highest	Sensitivity analysis for trust thresholds, penalty weights, and trust-aging logic.	Retain IDS gains while reducing avoidable routing penalties.	Parameter-tuned TACR-HOF variant.
Phase II: Generalization	High	Expand topologies, seeds, traffic loads, and benign-fault scenarios.	Determine whether the observed advantage holds across broader conditions.	Extended comparative benchmark.
Phase III: Advanced detection	High	Introduce temporal models, ablation studies, and explainability analysis.	Strengthen detection realism and clarify which features drive decisions.	Next-generation IDS evaluation layer.
Phase IV: Deployment realism	Medium to high	Measure energy, memory, and hardware-testbed behaviour; design distributed collection workflow.	Convert proof of concept into deployment-oriented knowledge.	Prototype deployment package and performance profile.

The ordering proposed in Table 5.4 is deliberate. Calibration and validation should precede radical architectural change. The current thesis already established proof of concept, so the first responsibility of future work is to stabilize and verify that concept across broader conditions. Only after that foundation is secured does it become sensible to move toward more advanced detection architectures or deployment packaging.

5.8 Final chapter conclusion

In final form, the thesis demonstrates that trust-aware routing can be used not merely to alter packet paths, but to create a more security-informative behavioural environment for intrusion detection in RPL-based IoT networks. That is the central intellectual result of the work. TACR-HOF did not simply produce isolated gains on one metric; it changed the data conditions under which attack behaviour was observed, learned, and classified.

Equally important, the thesis has shown how such a claim should be made responsibly. The comparative evidence strongly supports TACR-HOF as a superior security-oriented routing basis relative to MRHOF, yet it also reveals topology-dependent trade-offs that prevent any absolute claim of routing optimality. This bounded conclusion is a strength, not a weakness, because it aligns the argument exactly with the evidence.

Accordingly, the thesis closes with a clear position. The proposed TACR-HOF approach is validated as a meaningful and defensible advance for behaviour-based IDS in constrained RPL environments. The work is sufficiently complete to stand as a graduate-level contribution, and sufficiently open to motivate a substantial next stage of refinement. The future of this line of research is therefore not conceptual rescue, but disciplined enhancement: calibrate the trust logic, generalize the validation regime, extend the IDS layer, and carry the prototype toward deployment-realistic operation.

BIBLIOGRAPHY

- Al-Fuqaha, A., Guizani, M., Mohammadi, M., Aledhari, M. & Ayyash, M. (2015). Internet of Things: A survey on enabling technologies, protocols, and applications. *IEEE Communications Surveys & Tutorials*, 17(4), 2347–2376. doi: 10.1109/COMST.2015.2444095.
- Almusaylim, Z. A., Jhanjhi, N. & Alhumam, A. (2020). Detection and mitigation of RPL rank and version number attacks in the Internet of Things: SRPL-RP. *Sensors*, 20(21), 5997. doi: 10.3390/s20215997.
- Alsukayti, I. S. & Alreshoodi, M. (2023). RPL-based IoT networks under simple and complex routing security attacks: An experimental study. *Applied Sciences*, 13(8), 4878. doi: 10.3390/app13084878.
- Atzori, L., Iera, A. & Morabito, G. (2010). The Internet of Things: A survey. *Computer Networks*, 54(15), 2787–2805. doi: 10.1016/j.comnet.2010.05.010.
- Bormann, C., Ersue, M. & Keränen, A. (2014). Terminology for constrained-node networks. RFC Editor. doi: 10.17487/RFC7228.
- Canbalaban, E. & Şen, S. (2020). A Cross-Layer Intrusion Detection System for RPL-Based Internet of Things. *Ad-Hoc, Mobile, and Wireless Networks*, pp. 214–227.
- da Costa, K. A. P., Papa, J. P., Lisboa, C. O., Munoz, R. & de Albuquerque, V. H. C. (2019). Internet of Things: A survey on machine learning-based intrusion detection approaches. *Computer Networks*, 151, 147–157. doi: 10.1016/j.comnet.2019.01.023.
- Djedjig, N., Tandjaoui, D., Medjek, F. & Romdhani, I. (2020). Trust-aware and cooperative routing protocol for IoT security. *Journal of Information Security and Applications*, 52, 102467. doi: 10.1016/j.jisa.2020.102467.
- Elrawy, M. F., Awad, A. I. & Hamed, H. F. A. (2018). Intrusion detection systems for IoT-based smart environments: A survey. *Journal of Cloud Computing*, 7, 21. doi: 10.1186/s13677-018-0123-6.
- Friehat, N., Anbar, M., Aladaileh, M. A., Hasbullah, I., Shurbaji, T. A., Karuppayah, S. & Almomani, A. (2024). RPL-based attack detection approaches in IoT networks: Review and taxonomy. *Artificial Intelligence Review*, 57, 248. doi: 10.1007/s10462-024-10907-y.
- Garcia Ribera, E., Martinez Alvarez, B., Samuel, C., Ioulianou, P. P. & Vassilakis, V. G. (2022). An Intrusion Detection System for RPL-Based IoT Networks. *Electronics*, 11(23). Consulted at <https://www.mdpi.com/2079-9292/11/23/4041>.

- Gnawali, O. & Levis, P. (2012). The minimum rank with hysteresis objective function. RFC Editor. doi: 10.17487/RFC6719.
- Gyamfi, E. & Jurcut, A. (2022). Intrusion detection in Internet of Things systems: A review on design approaches leveraging multi-access edge computing, machine learning, and datasets. *Sensors*, 22(10), 3744. doi: 10.3390/s22103744.
- HUNSR. RPL-IDS-Behavior-Dataset. Computer software. Retrieved April 22, 2026, Consulted at <https://github.com/HUNSR/RPL-IDS-Behavior-Dataset>.
- Lamaazi, H. & Benamar, N. (2020). A comprehensive survey on enhancements and limitations of the RPL protocol: A focus on the objective function. *Ad Hoc Networks*, 96, 102001. doi: 10.1016/j.adhoc.2019.102001.
- Mayzaud, A., Badonnel, R. & Chrisment, I. (2016). A Taxonomy of Attacks in RPL-based Internet of Things. *Int. J. Netw. Secur.*, 18, 459-473. Consulted at <https://api.semanticscholar.org/CorpusID:10226823>.
- Medjek, F., Tandjaoui, D., Djedjig, N. & Romdhani, I. (2021). Multicast DIS attack mitigation in RPL-based IoT-LLNs. *Journal of Information Security and Applications*, 61, 102939. doi: <https://doi.org/10.1016/j.jisa.2021.102939>.
- Muzammal, S. M., Murugesan, R. K., Jhanjhi, N. Z., Humayun, M., Ibrahim, A. O. & Abdelmaboud, A. (2022). A trust-based model for secure routing against RPL attacks in Internet of Things. *Sensors*, 22(18), 7052. doi: 10.3390/s22187052.
- Oikonomou, G., Duquennoy, S., Elsts, A., Eriksson, J., Tanaka, Y. & Tsiftes, N. (2022). The Contiki-NG open source operating system for next generation IoT devices. *SoftwareX*, 18, 101089. doi: 10.1016/j.softx.2022.101089.
- Opitz, J. (2024). A closer look at classification evaluation metrics and a critical reflection of common evaluation practice. *Transactions of the Association for Computational Linguistics*, 12, 820–836. doi: 10.1162/tacl_a_00675.
- Osterlind, F., Dunkels, A., Eriksson, J., Finne, N. & Voigt, T. (2006). Cross-level sensor network simulation with COOJA. *Proceedings of the 31st IEEE Conference on Local Computer Networks*, pp. 641–648. doi: 10.1109/LCN.2006.322172.
- Quinlan, J. R. (1986). Induction of Decision Trees. *Machine Learning*, 1(1), 81–106. doi: 10.1007/BF00116251.
- Rahman, M. M., Al Shakil, S. & Mustakim, M. R. (2025). A survey on intrusion detection system in IoT networks. *Cyber Security and Applications*, 3, 100082.

doi: 10.1016/j.csa.2024.100082.

- Sobral, J. V. V., Rodrigues, J. J. P. C., Rabêlo, R. A. L., Al-Muhtadi, J. & Korotaev, V. (2019). Routing protocols for low power and lossy networks in Internet of Things applications. *Sensors*, 19(9), 2144. doi: 10.3390/s19092144.
- Tsao, T., Alexander, R., Dohler, M., Daza, V., Lozano, A. & Richardson, M. (2015). A security threat analysis for the routing protocol for low-power and lossy networks. RFC Editor. Consulted at <https://www.rfc-editor.org/rfc/rfc7416.html>.
- Verma, A. & Ranga, V. (2020). Security of RPL based 6LoWPAN networks in the Internet of Things: A review. *IEEE Sensors Journal*, 20(11), 5666–5690. doi: 10.1109/JSEN.2020.2973677.
- Watteyne, T., Palattella, M. R. & Grieco, L. A. (2015). Using IEEE 802.15.4e time-slotted channel hopping (TSCH) in the Internet of Things (IoT): Problem statement. RFC Editor. doi: 10.17487/RFC7554.
- Winter, T., Thubert, P., Brandt, A., Hui, J., Kelsey, R., Levis, P., Pister, K., Struik, R., Vasseur, J. P. & Alexander, R. (2012). RPL: IPv6 routing protocol for low-power and lossy networks. RFC Editor. doi: 10.17487/RFC6550.

APPENDIX I

ROUTING OBJECTIVE FUNCTION SOURCE-CODE EXCERPTS

This appendix presents selected implementation excerpts from the MRHOF and TACR-HOF routing objective functions. The excerpts illustrate the baseline MRHOF parent-evaluation procedure and the integration of the trust penalty into TACR-HOF. The line numbers correspond to the original source files.

1. MRHOF Path-Cost and Parent-Selection Functions

Listing I.1 presents the baseline MRHOF functions used to compute the path cost, derive the rank through a candidate neighbor, verify parent acceptability, apply hysteresis, and select the preferred parent. The excerpt corresponds to lines 130–247 of the `rpl-mrhof.c` source file.

```
130 /*-----*/
131 static uint16_t
132 nbr_path_cost(rpl_nbr_t *nbr)
133 {
134     uint16_t base;
135
136     if(nbr == NULL) {
137         return 0xffff;
138     }
139
140 #if RPL_WITH_MC
141     /* Handle the different MC types */
142     switch(curr_instance.mc.type) {
143         case RPL_DAG_MC_ETX:
144             base = nbr->mc.obj.etx;
145             break;
146         case RPL_DAG_MC_ENERGY:
147             base = nbr->mc.obj.energy.energy_est << 8;
148             break;
149         default:
150             base = nbr->rank;
151             break;
152     }
153 #else /* RPL_WITH_MC */
154     base = nbr->rank;
155 #endif /* RPL_WITH_MC */
156
157     /* path cost upper bound: 0xffff */
158     return MIN((uint32_t)base + link_metric_to_rank(nbr_link_metric(nbr)), 0xffff);
159 }
160 /*-----*/
161 static rpl_rank_t
162 rank_via_nbr(rpl_nbr_t *nbr)
163 {
164     uint16_t min_hoprankinc;
165     uint16_t path_cost;
166
167     if(nbr == NULL) {
168         return RPL_INFINITE_RANK;
169     }
170
171     min_hoprankinc = curr_instance.min_hoprankinc;
172     path_cost = nbr_path_cost(nbr);
173
174     /* Rank lower-bound: nbr rank + min_hoprankinc */
175     return MAX(MIN((uint32_t)nbr->rank + min_hoprankinc, RPL_INFINITE_RANK), path_cost);
176 }
177 /*-----*/
178 static int
179 nbr_has_usable_link(rpl_nbr_t *nbr)
```

```

180 {
181     uint16_t link_metric = nbr_link_metric(nbr);
182     /* Exclude links with too high link metrics */
183     return link_metric <= MAX_LINK_METRIC;
184 }
185 /*-----*/
186 static int
187 nbr_is_acceptable_parent(rpl_nbr_t *nbr)
188 {
189     // hunsr code start
190 #if LSM == 1
191     extern rpl_nbr_t *blacklistnbr;
192     if(blacklistnbr != NULL && nbr == blacklistnbr){
193         LOG_INFO(" hunsr found attack parent  nbr \n");
194         return false;
195     }
196 #endif /* LSM == 1 */
197     // hunsr code End
198
199     uint16_t path_cost = nbr_path_cost(nbr);
200     /* Exclude links with too high link metrics or path cost (RFC6719, 3.2.2) */
201     return nbr_has_usable_link(nbr) && path_cost <= MAX_PATH_COST;
202 }
203 /*-----*/
204 static int
205 within_hysteresis(rpl_nbr_t *nbr)
206 {
207     uint16_t path_cost = nbr_path_cost(nbr);
208     uint16_t parent_path_cost = nbr_path_cost(curr_instance.dag.preferred_parent);
209
210     int within_rank_hysteresis = path_cost + RANK_THRESHOLD > parent_path_cost;
211     int within_time_hysteresis = nbr->better_parent_since == 0
212         || (clock_time() - nbr->better_parent_since) <= TIME_THRESHOLD;
213
214     /* As we want to consider neighbors that are either beyond the rank or time
215        hystereses, return 1 here iff the neighbor is within both hystereses. */
216     return within_rank_hysteresis && within_time_hysteresis;
217 }
218 /*-----*/
219 static rpl_nbr_t *
220 best_parent(rpl_nbr_t *nbr1, rpl_nbr_t *nbr2)
221 {
222     int nbr1_is_acceptable;
223     int nbr2_is_acceptable;
224
225     nbr1_is_acceptable = nbr1 != NULL && nbr_is_acceptable_parent(nbr1);
226     nbr2_is_acceptable = nbr2 != NULL && nbr_is_acceptable_parent(nbr2);
227
228     if(!nbr1_is_acceptable) {
229         return nbr2_is_acceptable ? nbr2 : NULL;
230     }
231     if(!nbr2_is_acceptable) {
232         return nbr1_is_acceptable ? nbr1 : NULL;
233     }
234
235     /* Maintain stability of the preferred parent. Switch only if the gain
236        is greater than RANK_THRESHOLD, or if the neighbor has been better than the
237        current parent for at more than TIME_THRESHOLD. */
238     if(nbr1 == curr_instance.dag.preferred_parent && within_hysteresis(nbr2)) {
239         return nbr1;
240     }
241     if(nbr2 == curr_instance.dag.preferred_parent && within_hysteresis(nbr1)) {
242         return nbr2;
243     }
244
245     return nbr_path_cost(nbr1) < nbr_path_cost(nbr2) ? nbr1 : nbr2;
246 }
247

```

Listing I.1: MRHOF path-cost computation and candidate-parent selection

2. TACR-HOF Trust-Aware Path-Cost Computation

Listing I.2 shows how TACR-HOF augments the baseline MRHOF path cost with the trust penalty. A candidate is assigned an infinite path cost when either the baseline cost is invalid or the trust function returns the configured maximum path cost. The excerpt corresponds to lines 194–207 of the `rpl-tacrhof.c` source file.

```

194 static uint16_t
195 nbr_path_cost(rpl_nbr_t *nbr)
196 {
197     uint16_t base_cost;
198     uint16_t penalty;
199
200     base_cost = mrhof_path_cost(nbr);
201     penalty = trust_penalty(nbr);
202
203     if(base_cost >= 0xffff || penalty >= TACR_MAX_PATH_COST) {
204         return 0xffff;
205     }
206
207     return MIN((uint32_t)base_cost + penalty, 0xffff);
208 }

```

Listing I.2: Integration of the trust penalty into the TACR-HOF path cost

3. TACR-HOF Parent Acceptability and Selection

Listing I.3 presents the TACR-HOF functions used to evaluate candidate-parent acceptability, preserve routing stability through rank and time hysteresis, and select the preferred parent. Because `nbr_path_cost()` already includes the trust penalty, these decisions are based on the combined routing and trust-aware cost. The excerpt corresponds to lines 233–278 of the `rpl-tacrhof.c` source file.

```

233 static int
234 nbr_is_acceptable_parent(rpl_nbr_t *nbr)
235 {
236     #if LSM == 1
237         extern rpl_nbr_t *blacklistnbr;
238         if(blacklistnbr != NULL && nbr == blacklistnbr) {
239             return false;
240         }
241     #endif
242     return nbr_has_usable_link(nbr) && nbr_path_cost(nbr) <= TACR_MAX_PATH_COST;
243 }
244
245 static int
246 within_hysteresis(rpl_nbr_t *nbr)
247 {
248     uint16_t path_cost = nbr_path_cost(nbr);
249     uint16_t parent_path_cost = nbr_path_cost(curr_instance.dag.preferred_parent);
250
251     int within_rank_hysteresis = path_cost + TACR_RANK_THRESHOLD > parent_path_cost;
252     int within_time_hysteresis = nbr->better_parent_since == 0
253         || (clock_time() - nbr->better_parent_since) <= TACR_TIME_THRESHOLD;
254
255     return within_rank_hysteresis && within_time_hysteresis;
256 }
257
258 static rpl_nbr_t *
259 best_parent(rpl_nbr_t *nbr1, rpl_nbr_t *nbr2)
260 {
261     int nbr1_ok = nbr1 != NULL && nbr_is_acceptable_parent(nbr1);
262     int nbr2_ok = nbr2 != NULL && nbr_is_acceptable_parent(nbr2);
263

```

```
264  if(!nbr1_ok) {  
265      return nbr2_ok ? nbr2 : NULL;  
266  }  
267  if(!nbr2_ok) {  
268      return nbr1_ok ? nbr1 : NULL;  
269  }  
270  
271  if(nbr1 == curr_instance.dag.preferred_parent && within_hysteresis(nbr2)) {  
272      return nbr1;  
273  }  
274  if(nbr2 == curr_instance.dag.preferred_parent && within_hysteresis(nbr1)) {  
275      return nbr2;  
276  }  
277  
278  return nbr_path_cost(nbr1) < nbr_path_cost(nbr2) ? nbr1 : nbr2;  
279 }
```

Listing I.3: TACR-HOF parent acceptability, hysteresis, and preferred-parent selection