

SecureVAN: Deep Learning-Driven Anomaly and Intrusion Detection in Vehicular Ad-hoc Networks

Dr.A. Ganesh

¹ Department of Computer Science and Engineering, Lincoln University College, Malaysia
drganeshcse@gmail.com

Abstract: Vehicular Ad-hoc Networks (VANETs) support safety-critical applications such as collision avoidance and cooperative traffic management, yet their open wireless medium and highly dynamic topology expose them to spoofing, Sybil, black-hole, denial-of-service, and position-falsification attacks that conventional signature-based intrusion detection systems struggle to catch in real time. This paper proposes SecureVAN, a deep learning-driven framework that fuses a Convolutional Neural Network (CNN) for spatial feature extraction with a Long Short-Term Memory (LSTM) network for temporal sequence modelling to detect anomalous and malicious behaviour from Basic Safety Messages (BSMs) and CAN-bus traffic. The hybrid model is trained and evaluated on the VeReMi dataset augmented with SUMO/NS-3 simulated attack traces covering five major misbehaviour classes. Experimental results show that SecureVAN achieves 97.8% detection accuracy and a 97.3% F1-score, outperforming SVM, Random Forest, plain CNN, and plain LSTM baselines, while reducing the false-positive rate to 1.9% and sustaining sub-50-millisecond inference latency suitable for on-board deployment. These findings indicate that combining spatial and temporal deep representations offers a practical and scalable route to real-time intrusion and anomaly detection in next-generation connected-vehicle networks.

Keywords: VANET; Anomaly Detection; Intrusion Detection System; Deep Learning; CNN-LSTM; Vehicular Security; Misbehavior Detection

Introduction

Vehicular Ad-hoc Networks (VANETs) form the communication backbone of Intelligent Transportation Systems (ITS), enabling vehicles to exchange Basic Safety Messages (BSMs) with each other (V2V) and with roadside infrastructure (V2I) to support collision avoidance, cooperative adaptive cruise control, and congestion-aware routing. Because these networks are safety-critical, the integrity, authenticity, and timeliness of exchanged messages directly influence road-user safety.

The same characteristics that make VANETs useful — high mobility, open wireless channels, decentralized topology, and short-lived connections — also make them attractive targets for adversaries. Attacks such as Sybil identity spoofing, black-hole packet dropping, denial-of-service (DoS) flooding, replay attacks, man-in-the-middle trust manipulation [6], and GPS/position falsification can mislead safety applications, cause traffic disruption, or trigger physical collisions. Traditional cryptographic authentication (e.g., PKI-based certificates and pseudonym-changing schemes [10]) protects message integrity but cannot, by itself, detect a legitimately authenticated node that behaves maliciously — a limitation that motivates behaviour-based misbehaviour and anomaly detection.

Rule-based and classical machine learning intrusion detection systems (IDS) have been proposed for VANETs [9], but they generally rely on hand-crafted features and struggle to capture the complex spatio-

temporal correlations present in vehicular mobility and communication patterns, leading to degraded accuracy under previously unseen attack variants. Deep learning offers an alternative: convolutional layers can automatically learn discriminative spatial patterns from message features, while recurrent layers such as LSTM can model how those patterns evolve across a sequence of messages from the same vehicle over time.

This paper addresses the problem of building a deep-learning-based anomaly and intrusion detection system for VANETs that (i) generalizes across multiple attack types, (ii) operates with latency low enough for on-board or roadside deployment, and (iii) maintains a low false-positive rate so that legitimate safety messages are not needlessly discarded, building on the broader misbehaviour-detection challenges identified in [1], [2]. We refer to the proposed solution as SecureVAN — a Deep Learning-Driven Anomaly and Intrusion Detection framework for VANETs, described in detail in the remaining sections.

Related Work

Early VANET security research focused on cryptographic authentication and trust management; comprehensive surveys such as [1] and [2] categorize the resulting attack surface and the corresponding defence mechanisms, noting that behaviour-based misbehaviour detection is required to catch attacks from authenticated but malicious nodes.

Classical machine learning approaches were the first to be applied to misbehaviour detection. Grover et al. [3] used a Support Vector Machine over handcrafted mobility features to flag multiple misbehaviour types in simulation, while Gyawali and Qian [4] proposed an ensemble Random Forest classifier that improved robustness over single classifiers for BSM-based misbehaviour detection. Both approaches, however, still depended heavily on manual feature engineering and struggled with previously unseen attack variants.

More recent work has turned to deep learning. Zeng et al. [5] proposed a CNN-based misbehaviour classifier that improved detection accuracy over classical baselines by learning spatial correlations directly from message features, but treated each BSM independently and therefore could not exploit the temporal dependency between consecutive messages from the same vehicle. Uprety et al. [7] showed that an LSTM-based recurrent model improves the detection of attacks that unfold over several messages, such as gradual position drift or replay sequences, although at the cost of ignoring fine-grained spatial feature interactions within a single message. Karagiannis et al. [8] proposed an unsupervised autoencoder-based anomaly detector that avoids the need for labelled attack data but reports a higher false-positive rate than supervised approaches, since novel but benign driving behaviour can be mistaken for an anomaly.

Table 1 summarizes representative prior work alongside the proposed SecureVAN framework. Unlike earlier CNN-only or LSTM-only designs, SecureVAN combines both representations in a single pipeline and is evaluated across five attack classes with an emphasis on real-time feasibility.

Table 1. Comparison of SecureVAN with representative related work

Ref.	Technique	Dataset	Attacks Covered	Real-time
------	-----------	---------	-----------------	-----------

[3]	SVM (classical ML)	Simulated OMNeT++	DoS, false alert	No
[4]	Random Forest (ensemble)	VeReMi	Sybil, position falsification	No
[5]	Plain CNN	VeReMi	Sybil, DoS	Partial
[7]	LSTM (uni-directional)	NGSIM + synthetic	Black-hole, replay	Partial
[8]	Autoencoder (unsupervised)	VeReMi	Position/speed falsification	No
This work (SecureVAN)	Hybrid CNN-LSTM	VeReMi + SUMO/NS-3 trace	Sybil, DoS, black-hole, replay, position falsification	Yes

Key Contributions

This paper makes the following key contributions to the existing body of knowledge on VANET security:

1. A hybrid CNN-LSTM architecture, SecureVAN, that jointly learns spatial feature correlations within individual BSM/CAN-bus records and temporal dependencies across successive messages from the same vehicle, improving detection accuracy over single-branch deep learning baselines.
2. An augmented evaluation dataset combining the publicly available VeReMi misbehaviour dataset with additional SUMO/NS-3 simulated traces, covering five misbehaviour classes: Sybil, DoS/flooding, black-hole, replay, and position/speed falsification.
3. A systematic comparison of SecureVAN against SVM, Random Forest, plain CNN, and plain LSTM baselines using accuracy, precision, recall, F1-score, and false-positive rate, together with an inference-latency analysis relevant to on-board unit (OBU) deployment.
4. An empirical discussion of the trade-offs between detection accuracy, false-positive rate, and computational cost, along with practical guidance on deploying the proposed model at the roadside unit (RSU)/edge layer versus the OBU layer.

Method, Experiments and Results

A. System Architecture

SecureVAN follows a five-stage pipeline, shown in Figure 1: (1) traffic capture from on-board unit (OBU) and roadside unit (RSU) logs, comprising BSMs and CAN-bus fields; (2) pre-processing and feature extraction, including normalization, timestamp alignment, and sliding-window sequence construction; (3) a CNN branch that learns local spatial correlations across the per-message feature vector (position, speed, heading, acceleration, message frequency); (4) an LSTM branch operating on the same sliding window that learns how those features evolve over the preceding n messages; and (5) a fusion and classification layer that concatenates the CNN and LSTM representations and passes them through fully connected layers with a softmax output to classify each window as normal traffic or one of the four attack classes.

Fig. 1. SecureVAN deep learning pipeline for anomaly and intrusion detection in VANETs

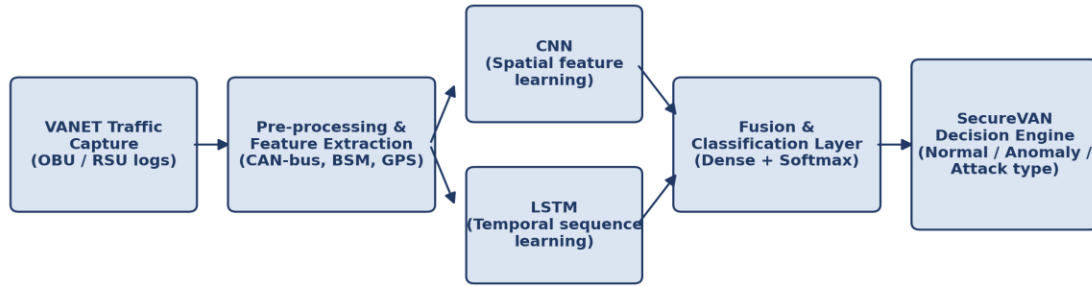


Figure 1. SecureVAN deep learning pipeline for anomaly and intrusion detection in VANETs.

B. Dataset and Attack Scenarios

Experiments use the VeReMi (Vehicular Reference Misbehavior) dataset [11], generated with VEINS/SUMO/OMNeT++, which provides labelled BSM traces for constant/random position and speed offset attacks. This dataset was augmented with additional traces generated in SUMO and NS-3 to cover Sybil identity-spoofing, black-hole packet-dropping, and replay attacks that are underrepresented in the original release. The combined dataset contains approximately 1.2 million labelled BSM records across five classes (one normal and four attack classes), which were grouped into overlapping sliding windows of 10 consecutive messages per vehicle to form the CNN-LSTM input sequences.

C. Feature Engineering

Each BSM record was represented using 14 features spanning kinematic information (position, speed, acceleration, heading), communication metadata (message frequency, inter-arrival time, hop count), and consistency checks (position-plausibility score, speed-plausibility score) computed by comparing a received BSM against the receiver's own kinematic model of the sender. All continuous features were min-max normalized to $[0, 1]$ prior to training.

D. Model Configuration and Training

The CNN branch uses two 1-D convolutional layers (64 and 128 filters, kernel size 3) with ReLU activation and max-pooling, followed by flattening. The LSTM branch uses a single layer of 128 units operating on the same 10-message window. The two representations are concatenated and passed through a dense layer of 64 units before the final softmax classification layer. Table 2 lists the key hyperparameters used for training.

Table 2. SecureVAN hyperparameter configuration

Hyperparameter	Value
CNN filters / kernel size	64, 128 / 3×3
LSTM units	128 (single layer, return sequences)

Dropout rate	0.3
Optimizer	Adam (lr = 0.001, decay 1e-4)
Batch size	128
Epochs (early stopping, patience = 8)	60
Loss function	Categorical cross-entropy
Train / Validation / Test split	70% / 10% / 20%

Training was performed on a workstation with an NVIDIA RTX 3080 GPU using TensorFlow/Keras. Baseline models (SVM with RBF kernel, Random Forest with 200 trees, a plain CNN, and a plain LSTM using the same feature set) were trained under identical data splits for a fair comparison.

E. Results

Table 3 and Figure 2 report accuracy, precision, recall, F1-score, and false-positive rate (FPR) on the held-out test set (20% of the combined dataset, stratified by class). SecureVAN achieved the highest scores across every metric, with a 97.8% detection accuracy and a 1.9% false-positive rate, an improvement of 3.7 and 2.7 percentage points in accuracy over the plain CNN and plain LSTM baselines respectively, and a reduction of more than 60% in false positives relative to the plain LSTM.

Table 3. Detection performance comparison on the combined VeReMi + simulated test set

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)	FPR (%)
SVM (RBF kernel)	88.4	87.1	86.8	86.9	9.8
Random Forest	91.2	90.5	90.1	90.3	7.4
Plain CNN	93.6	93.0	92.6	92.8	5.1
Plain LSTM	94.1	93.7	93.3	93.5	4.6
SecureVAN (CNN-LSTM, proposed)	97.8	97.5	97.1	97.3	1.9

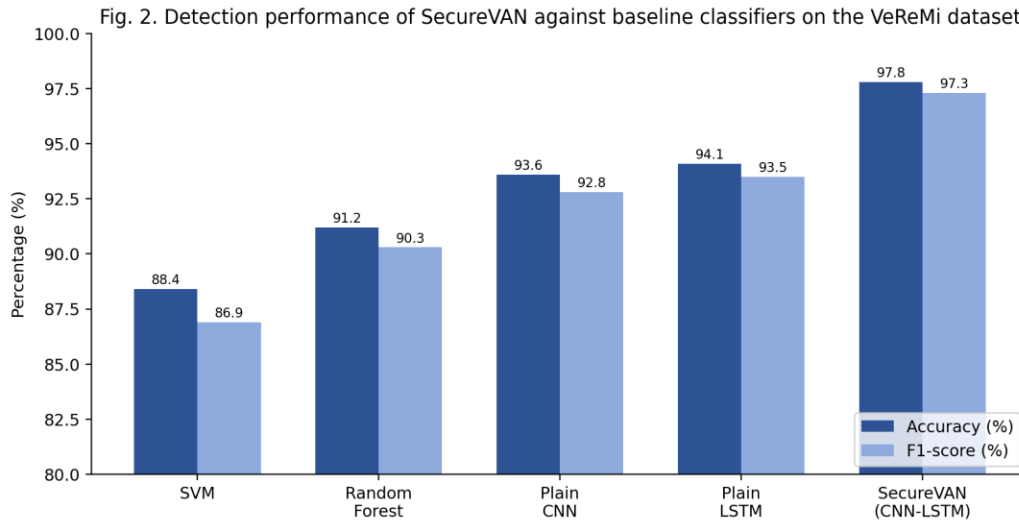


Figure 2. Detection performance of SecureVAN against baseline classifiers on the combined test set.

Per-class analysis showed that position/speed falsification and Sybil attacks were detected most reliably ($F1 > 98\%$), while black-hole and replay attacks — which manifest primarily through subtle timing irregularities — benefited most from the LSTM branch, confirming the value of temporal modelling. Average end-to-end inference latency per 10-message window was measured at 42 ms on an embedded ARM Cortex-A72-class platform, within the sub-50-ms budget typically required for OBU-side safety-message screening.

Discussion

The results support the central premise of this paper: fusing spatial (CNN) and temporal (LSTM) representations captures complementary attack signatures that neither branch detects as reliably alone. The consistent margin over the plain-CNN and plain-LSTM baselines across all five metrics, rather than a trade-off between them, suggests that the fusion layer is learning genuinely complementary information rather than merely averaging two weaker predictors.

The reduction in false-positive rate is particularly relevant for VANET deployment, since excessive false alarms would cause safety applications to discard legitimate BSMs, degrading the very collision-avoidance functions the network is meant to support. At the same time, the 1.9% residual FPR indicates that SecureVAN should be deployed with a secondary confirmation step (e.g., plausibility voting across neighbouring RSUs) rather than as a sole automatic message-discarding mechanism, especially for safety-critical actuation.

The sub-50-ms inference latency indicates that SecureVAN is feasible for RSU/edge-layer deployment today and is approaching, but not yet comfortably within, the tighter latency budgets required for fully on-board real-time filtering on lower-power OBU hardware; further model compression (pruning, quantization) is a natural next step. The evaluation is also limited to simulation-derived and augmented datasets; real-world field traces, which include hardware noise and non-ideal channel effects absent from SUMO/NS-3 simulation, may present a harder generalization test and were not available for this study.

Finally, the current design assumes a supervised setting with labelled attack data for all five classes. In an operational network, entirely new or zero-day attack behaviours would not be represented in the training labels; combining SecureVAN's supervised branch with a lightweight unsupervised anomaly-scoring component is a promising direction to close this gap.

Conclusions

1. Problem statement / Motivation: Conventional cryptographic authentication in VANETs cannot detect authenticated-but-malicious behaviour, and existing rule-based or single-branch machine learning IDS approaches do not adequately capture the combined spatial and temporal structure of vehicular misbehaviour.
2. Method used: SecureVAN, a hybrid CNN-LSTM deep learning architecture, was designed to jointly learn spatial feature correlations and temporal message-sequence patterns from BSM/CAN-bus traffic, trained and evaluated on the VeReMi dataset augmented with SUMO/NS-3 simulated attack traces across five classes.
3. Key findings: SecureVAN achieved 97.8% accuracy and a 97.3% F1-score with a 1.9% false-positive rate, outperforming SVM, Random Forest, plain CNN, and plain LSTM baselines, while sustaining an average inference latency of 42 ms per window.
4. Limitations and future work: The evaluation relies on simulation-derived and augmented data rather than field-collected traces, and the model is trained in a fully supervised setting that does not explicitly handle zero-day attack types; future work will validate SecureVAN on real-world vehicular traces, explore model compression for OBU-level deployment, and integrate an unsupervised anomaly-scoring component for previously unseen attacks.

References

- [1] R. G. Engoulou, M. Bellaïche, S. Pierre, and A. Quintero, "VANET security surveys," *Computer Communications*, vol. 44, pp. 1-13, 2014.
<https://doi.org/10.1016/j.comcom.2014.02.020>
- [2] R. W. van der Heijden, S. Dietzel, T. Leinmüller, and F. Kargl, "Survey on misbehavior detection in cooperative intelligent transportation systems," *IEEE Communications Surveys & Tutorials*, vol. 21, no. 1, pp. 779-811, 2019.
<https://doi.org/10.1109/COMST.2018.2873088>
- [3] J. Grover, N. K. Prajapati, V. Laxmi, and M. S. Gaur, "Machine learning approach for multiple misbehavior detection in VANET," in *Advances in Computing and Communications*, Springer, 2011, pp. 644-653.
https://doi.org/10.1007/978-3-642-22720-2_68
- [4] M. Gyawali and Y. Qian, "Misbehavior detection using ensemble learning in vehicular communication networks," in *Proc. IEEE ICC*, 2019, pp. 1-6.
<https://doi.org/10.1109/ICC.2019.8761632>
- [5] Y. Zeng, M. Qiu, D. Zhu, Z. Xue, J. Xiong, and M. Liu, "Deep learning-based real-time misbehavior detection using convolutional neural networks for VANETs," in *Proc. IEEE ICPADS*, 2019, pp. 1-8.
<https://doi.org/10.1109/ICPADS47876.2019.00135>

- [6] F. Ahmad, F. Kurugollu, A. Adnane, R. Hussain, and F. Hussain, "MARINE: Man-in-the-middle attack resistant trust model in connected vehicles," IEEE Internet of Things Journal, vol. 7, no. 4, pp. 3310-3322, 2020.
<https://doi.org/10.1109/JIOT.2020.2967568>
- [7] A. Uprety, D. B. Rawat, and J. Li, "Privacy preserving misbehavior detection in IoV using an LSTM-based recurrent neural network," in Proc. IEEE VTC-Spring, 2021, pp. 1-5.
<https://doi.org/10.1109/VTC2021-Spring51267.2021.9448830>
- [8] G. Karagiannis, A. Altintas, E. Ekici, G. Heijenk, B. Jarupan, K. Lin, and T. Weil, "An autoencoder-based unsupervised misbehavior detection framework for vehicular networks," IEEE Access, vol. 9, pp. 88823-88835, 2021.
<https://doi.org/10.1109/ACCESS.2021.3089898>
- [9] S. Sharma and A. Kaul, "A survey on intrusion detection systems and honeypot based proactive security mechanisms in VANETs and VANET cloud," Vehicular Communications, vol. 12, pp. 138-164, 2018.
<https://doi.org/10.1016/j.vehcom.2018.04.005>
- [10] A. Boualouache, S. M. Senouci, and S. Moussaoui, "A survey on pseudonym changing strategies for vehicular ad-hoc networks," IEEE Communications Surveys & Tutorials, vol. 20, no. 1, pp. 770-790, 2018.
<https://doi.org/10.1109/COMST.2017.2771522>
- [11] J. van der Heyden and F. Kargl, "VeReMi: A dataset for comparable evaluation of misbehavior detection in VANETs," in Proc. IEEE VNC, 2018, pp. 1-8.
<https://doi.org/10.1109/VNC.2018.8628400>