

A Combinational Model-Based APT Attack Detection Approach

DOI: <https://doi.org/10.54654/isj.v1i24.1078>

Nguyen Thanh Tung*, Nguyen Tien Trong, Nguyen Thi Hien,
Nguyen Quang Hoan, Do Xuan Cho

Abstract— In the context of a world increasingly reliant on digital technology, Advanced Persistent Threats (APT) pose a significant challenge to global cybersecurity. To address this issue, this paper introduces a novel approach called BiLSTM-Attention-GCN (BiLAG), an advanced model combining Bidirectional Long Short-Term Memory (BiLSTM) networks, Attention mechanisms, and Graph Convolutional Networks (GCN). The goal of BiLAG is to provide an effective and accurate method for detecting APT. BiLSTM is employed to capture temporal features related to event sequences, enabling the detection of anomalies over time. The Attention mechanism focuses on the most critical aspects of the dataset, allowing the model to identify hidden signals that indicate potential attacks. Lastly, GCN is utilized to explore complex relationships among network entities, enhancing APT detection by constructing a detailed and precise relational graph. Experimental results demonstrate that BiLAG achieves an accuracy of 99%, with high recall and significantly reduced false positive rates.

Tóm tắt — Trong bối cảnh thế giới ngày càng phụ thuộc vào công nghệ số, các mối đe dọa dai dẳng nâng cao (APT) đang đặt ra một thách thức đáng kể đối với an ninh mạng toàn cầu. Để giải quyết vấn đề đó, bài báo này giới thiệu một phương pháp mới có tên BiLSTM-Attention-GCN (BiLAG), một mô hình tiên tiến kết hợp các mạng nơ-ron bộ nhớ dài ngắn hạn hai chiều (BiLSTM), cơ chế Attention và các Mạng nơ-ron tích chập đồ thị (GCN). Mục tiêu của BiLAG là cung cấp một phương pháp hiệu quả và chính xác để phát hiện các mối đe dọa APT. BiLSTM được sử dụng để nắm bắt các đặc trưng thời gian liên quan đến chuỗi sự kiện, từ đó cho phép phát hiện các điểm bất thường

theo thời gian. Cơ chế Attention tập trung vào những khía cạnh quan trọng nhất của tập dữ liệu, cho phép mô hình xác định các tín hiệu ẩn báo hiệu các cuộc tấn công tiềm ẩn. Cuối cùng, GCN được sử dụng để khám phá các mối quan hệ phức tạp giữa các thực thể trong mạng, qua đó nâng cao khả năng phát hiện APT thông qua việc xây dựng một đồ thị quan hệ chi tiết và chính xác. Kết quả thí nghiệm cho thấy BiLAG đạt được độ chính xác 99%, với độ nhạy cao và tỷ lệ dương tính giả giảm đáng kể.

Keywords— APT attack, BiLSTM model, BiLAG model, GCN, deep learning model.

Từ khóa— Tấn công APT, mô hình BiLSTM, mô hình BiLAG, GCN, mô hình học sâu.

I. INTRODUCTION

A. Overview

APT is a type of complex and prolonged cyberattack aimed at infiltrating and exploiting data from organizations, enterprises, or governments. APT attacks are often carried out by professional hacking groups using sophisticated methods to remain undetected and collect data without exposure. Studies [1, 2, 14] have shown that the number of APT attacks has significantly increased in recent years, with their scale and impact growing exponentially. Detecting APT requires advanced methods that not only rely on specific indicators but also analyze overall network behavior. While machine learning and deep learning have been utilized to analyze network traffic and detect anomalies related to APT [3, 4, 7, 15], standalone deep learning models face limitations in processing large-scale data and recognizing complex relationships among various attack behaviors. To address these challenges, the BiLAG model is proposed, combining three main components: BiLSTM, which identifies

This manuscript was received on January 6, 2025. It was reviewed on March 24, 2025, revised on April 18, 2025 and accepted on May 5, 2025.

* Corresponding author

temporal sequence characteristics and analyzes prolonged abnormal behaviors suitable for APT detection; Attention Mechanism, which focuses on the most prominent abnormal features to enhance the accuracy of distinguishing normal traffic from attack indications; and GCN, which explores graph structures to model relationships between host addresses and detect anomalous behaviors across different parts of the system. This combination allows the BiLAG model to improve detection performance, quickly adapt to new attack patterns, and offer a robust, efficient solution for cybersecurity.

B. Operational Principles of the BiLAG Model

The BiLAG model analyzes and detects APT based on network traffic through the following steps:

Stage 1 Temporal Feature Analysis: BiLSTM and the Attention mechanism are utilized to identify anomalies in network data, prioritizing potential APT flows by aggregating host addresses into feature vectors.

Stage 2 host Network Graph Modeling: The network graph is constructed, where nodes represent hosts and edges represent their relationships. The GCN learns graph features to detect abnormal behaviors.

Stage 3 host Classification: Based on the behavior graph, the model classifies hosts, distinguishing between APT-related hosts and normal ones.

C. Contributions of the Paper

The scientific and practical significance of APT attack detection in this study includes:

- The paper proposes a novel approach combining deep learning and graph analysis to detect APT attacks in network traffic.
- It introduces a method to construct host information from network traffic data, prioritizing suspected APT flows, thereby enhancing detection accuracy.
- The model demonstrates applicability to cybersecurity monitoring systems, enabling organizations to detect threats early and respond swiftly.

II. RELATED WORK

Recent research in APT detection has explored diverse machine learning and deep learning approaches. Arefin et al. ([5]) proposed a hybrid model leveraging modern machine learning techniques, achieving a remarkable accuracy of 96.9%. Their approach significantly outperformed traditional algorithms, including MLP Classifier (94.5%), Gradient Boosting (92.3%), and KNN (76.6%), demonstrating the potential of integrated machine learning strategies in cybersecurity threat detection.

Complementing this, Cho et al. [4] introduced the FIERL model, an innovative approach combining BiLSTM, Attention, and contrastive learning techniques. Their primary objective was to effectively detect APT attacks by extracting anomalous host behaviors and optimizing classification performance. The model demonstrated substantial improvements, outperforming existing methods by over 5% across all evaluation metrics, thus highlighting the effectiveness of advanced deep learning architectures in network security analysis.

Research on APT detection utilizing the NSL-KDD dataset [6] further expanded the methodological landscape by integrating decision trees, Bayesian networks, and deep learning strategies. A particularly noteworthy six-layer deep learning model achieved an impressive accuracy of 98.85% with a remarkably low error rate, underscoring the potential of sophisticated multi-layer neural network architectures in threat detection.

Do Xuan et al. [12] proposed a deep learning method for detecting APTs using network flow analysis, combining BiLSTM and GCN, which outperformed traditional models such as MLP and standalone GCN. Their research is similar to ours in applying deep learning to network data; however, we integrate an Attention mechanism into the BiLSTM-GCN framework and leverage graph representations, enhancing detection accuracy and providing deeper insights into APT behaviors.

Nguyen et al. [13] proposed the MIG model, which integrates MLP, Inference, and GCN to

detect APT through network traffic analysis. In this model, MLP is responsible for aggregating and extracting features from host addresses in network flows, the Inference layer constructs host profiles by grouping and concatenating flows from the same host, and GCN analyzes and reconstructs host features based on behavior extraction from these profiles. Experimental results demonstrate that MIG not only improves APT detection accuracy but also minimizes false alarms. Similar to our research, Nguyen et al. apply deep learning to network traffic analysis for APT detection. However, while MIG focuses on host profiling through Inference, our approach integrates an Attention mechanism into the BiLSTM-GCN framework, enhancing the focus on critical features and providing deeper insights into APT behaviors through graph-based representations.

In a different domain, the BERT-BiGRU-CRF Model [7] presented a novel approach by combining BERT, BiGRU, and CRF techniques to extract APT-related events from web texts. This model achieved a high F1-score, effectively surpassing established models like ERNIE and BERT, and demonstrating the versatility of natural language processing techniques in cybersecurity event extraction.

III. THE BiLAG MODEL FOR APT DETECTION

A. Architecture of the Proposed Model

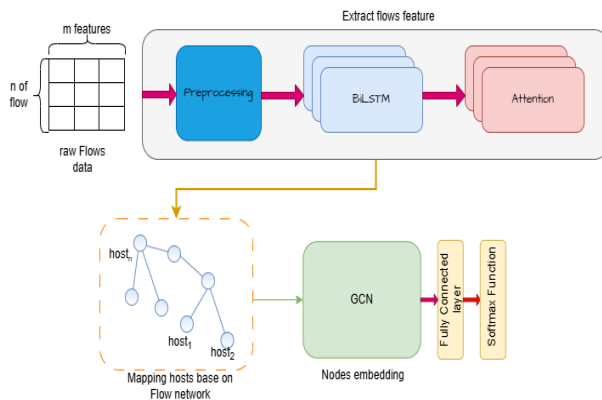


Figure 1. Architecture of the BiLAG Model for APT Detection

Phase 1: Flow Information Aggregation and host Network Construction.

This process includes two main steps:

Step 1: Extracting flow information: This phase focuses on extracting network flow attributes from network traffic. The BiLSTM-Attention model is proposed to aggregate and analyze the behavior of flows in network traffic.

Step 2: Aggregating host information from flows: After the analysis, the model continues to group flows with the same host address. The goal of this step is to build comprehensive information for each host based on the processed flows, serving as a foundation for subsequent host classification tasks.

Phase 2: Modeling Detailed host Behavioral Networks. This phase aims to analyze and extract information based on relationships between hosts through two steps:

Step 1: Constructing the host relationship graph: Data-exchanging hosts are represented as a graph, with nodes as hosts and edges as attributes from the network flow.

Step 2: Extracting relationship information between hosts: GCN is used to exploit the complex relationships between hosts.

Phase 3: Host Classification: The task of this step is to classify hosts into APT-related and normal hosts.

B. Aggregating and Constructing the host Information Network Based on the BiLSTM and Attention Model.

1. Extracting Two-Dimensional Temporal Features of Flows Using BiLSTM.

BiLSTM is an extended RNN architecture that allows learning from both past and future data in a sequence Graves et al. [8]. BiLSTM consists of two layers: Forward LSTM (processing data in a forward direction) and Backward LSTM (processing data in a backward direction). By combining contextual information from both directions, BiLSTM enhances the detection of prolonged abnormal behavior patterns.

BiLSTM generates two hidden states, a forward hidden state h_t^f and a backward hidden state h_t^b , which are then concatenated to form the final hidden state:

$$h_t^{\square} = \text{Concat}(h_t^f, h_t^b) \quad (1)$$

where $h_t^{\square}$ is the hidden state at time t , and "concat" represents the concatenation of the two vectors.

This approach enables BiLSTM to leverage information from both preceding events and nearby future events, providing superior sequence data analysis compared to traditional LSTM.

2. Attention Mechanism Focusing on Abnormalities in Each Flow.

Attention is a crucial mechanism in deep learning, enabling models to focus on essential information in a sequence. Bahdanau et al. [9] introduced Attention in machine translation to identify the most relevant words for translation. Vaswani et al [10] extended it into Self-Attention within the Transformer, allowing efficient processing of long sequences.

In this paper, a self-attention mechanism is proposed to aggregate the output matrix H of BiLSTM, emphasizing key features within the data. The calculations are as follows:

$$q = h_n^{\square} \omega, \quad K = HW^K, \quad V = HW^V \quad (2)$$

$$\text{Attention}(q, K, V) = \text{softmax}\left(\frac{qK^T}{\sqrt{d^k}}\right) V \quad (3)$$

$$h_v^{\square} = \text{Attention}(q, K, V) \quad (4)$$

Here $h_n^{\square}$ represents the final output of the BiLSTM network, ω is a learnable vector, and W^K, W^V are learnable weight matrices. $h_v^{\square}$ is the final output representation after applying the Attention mechanism.

After processing the flows through the BiLSTM-Attention model, feature vectors $h_v^{\square}$ are obtained for each flow. The model then groups flows with the same host address and aggregates their feature vectors into a single feature vector for each host. This feature vector contains comprehensive information, serving as a robust foundation for effective host classification.

C. Modeling the Detailed Behavioral Network of hosts

Each host has relationships with other hosts through data exchange. These relationships are crucial for evaluating the security level of an host. In this study, relationships between hosts are determined based on their communication: if two hosts communicate, they are connected by an edge in the graph. By analyzing communication between pairs of hosts, a graph is constructed where nodes represent hosts and edges depict their communication relationships.

GCN are a deep learning model that extends Convolutional Neural Networks (CNN) to work with graph-structured data, introduced by Kipf and Welling in 2017 [11]. This model is particularly useful for modeling complex relationships between nodes.

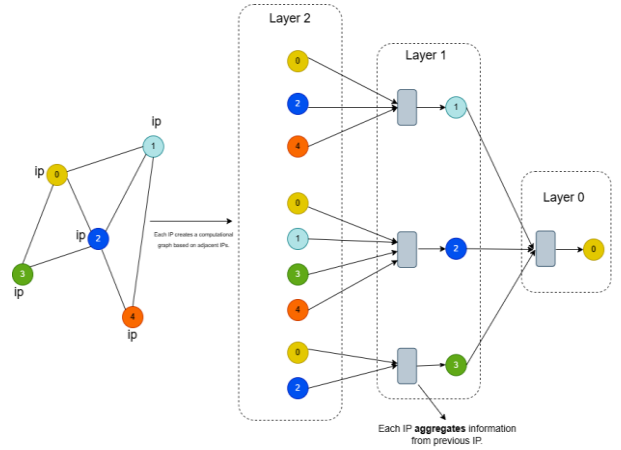


Figure 2. Multi-layer Information Aggregation Process in the GCN for Each Node (host)

In our study, the GCN plays a crucial role in enhancing the detection of APT attacks by effectively processing the graph data of the host information network. Specifically, GCN operates on a graph where nodes represent host addresses, processes, or files, and edges denote relationships between them, such as network connections or interactions between entities. Through graph convolution, GCN aggregates information from neighboring nodes to update the features of each node, constructing detailed representations that reflect both individual behaviors and the relational context within the network. This enables the model to detect complex APT behavioral patterns, such as

propagation through network connections or suspicious interactions, which are challenging to capture using sequential methods like BiLSTM or attention mechanisms alone. Compared to the Graph Attention Network (GAT), which offers greater flexibility by assigning different weights to connections through its attention mechanism, GCN adopts a simpler approach by treating all neighboring nodes with equal weights. This not only reduces computational complexity but also makes it well-suited for graphs with relatively homogeneous structures, as is the case in our research.

In this study, the input to the GCN consists of host graphs constructed following the above principle. The GCN treats hosts as graph nodes and extracts relational attributes between host edges, as illustrated in Figure 2.

The formula for the convolutional operation in each GCN layer is expressed as follows:

$$H^{(l+1)} = \sigma(D^{-1/2} \tilde{A} D^{-1/2} H^{(l)} W^{(l)}) \quad (5)$$

Where:

$H^{(l)}$ is the feature matrix of nodes at layer l .

$\tilde{A} = A + I$ is the adjacency matrix of the graph with added self-loops.

D is the degree matrix normalized for the nodes.

$W^{(l)}$ is the learnable weight matrix in layer l .

σ is the activation function, typically ReLU.

D. Host Classification

To classify APT hosts and normal hosts, the model utilizes two layers: Fully Connected Layers and a Softmax Layer. These layers perform the following functions:

- Fully Connected Layer: Learns the attributes extracted by the GCN layer, similar to a MLP.
- Softmax Layer: Calculates the probability of output labels using the Softmax function.

IV. EXPERIMENTAL EVALUATION

A. Experimental Data

The experimental data was collected from 29 network traffic files in the Malware Capture CTU-13 dataset, consisting of six types of

malware from APT attacks: Andromeda, Cobalt, Cridex, Dridex, Emotet, and Gh0stRAT. Clean data was extracted from the e-Government server of Quang Nam province on July 27, 2019, under the research project KC.01.05/16-20 funded by the Ministry of Science and Technology of Vietnam.

After the network traffic was processed through CICFlowMeter to generate network flows, the data underwent preprocessing. Duplicate flows and flows containing NaN or infinity values were removed. Unnecessary fields such as 'FlowID,' 'SrcIP,' 'SrcPort,' 'DstIP,' 'DstPort,' and 'Timestamp' were also eliminated. The data was normalized to a mean of 0 and a standard deviation of 1. The dataset was split into 80% for training and 20% for evaluation.

TABLE I. COMPOSITION OF THE EXPERIMENTAL DATASET

N^o	Type	Total	APT	BENIGN
1	Flows	1.671.393	871.914 (52.17%)	799.479 (47.83%)
2	Host	1.443	144 (9.98%)	1299 (90.02%)

B. Evaluation Criteria and Scenarios

Four key metrics, as utilized in [12], were employed to evaluate the model's performance. To assess the effectiveness of the proposed model, experiments are conducted to address the following questions:

- QA 1: How does the BiLAG model perform with imbalanced datasets? Fine-tuning BiLAG parameters to find the optimal configuration.
- QA 2: Why choose BiLAG over other models? The following experiments clarify the roles of networks in BiLAG:

Evaluating BiLSTM: Testing models such as CNN-Attention-GCN, RNN-Attention-GCN, and LSTM-Attention-GCN.

Evaluating Attention: Testing BiLSTM-Mean-GCN and BiLSTM-Inference-GCN.

Evaluating GCN: Comparing BiLSTM-Attention-GAT with BiLAG.

- QA 3: How does BiLAG compare to other models such as BiLSTM-GCN [12], CNN-LSTM [4], and MIG [13]?

C. Experimental Results

1. Experiments Addressing QA1

In the experimental dataset (Table 1), the data is imbalanced, with the number of normal hosts being nine times higher than APT hosts. This reflects real-world scenarios in APT detection, as attacks are rare compared to normal traffic. The experimental results of the BiLAG model (Table 2) show that the model was fine-tuned for optimal performance and evaluated for effectiveness against imbalanced datasets.

TABLE II. EXPERIMENTAL RESULTS OF APT DETECTION USING THE BiLAG MODEL

BiLAG			Evaluation of host			
BiLSTM nodes	Attention nodes	GCN nodes				
64 - 128	256	128-128	0.97	0.85	0.79	0.82
128 - 128	256	256-256	0.97	0.86	0.86	0.86
128 - 256	512	256-256	0.98	0.83	0.86	0.89
256 - 256	512	512-512	<u>0.99</u>	<u>0.93</u>	<u>0.93</u>	<u>0.93</u>

We propose the BiLAG model, consisting of two BiLSTM layers, one Attention layer, and two GCN layers, for APT detection. The experimental results (Table 2) indicate that BiLAG achieves an accuracy of 97% to 99%, with the best configuration being BiLSTM 256-256, Attention 512, and GCN 512. Precision and Recall both reached 93%, reflecting a balance between accurate detection and reduced false alarms. Smaller configurations, such as BiLSTM 64-128 and GCN 128, only achieved a Precision of 85% and Recall of 79%, indicating a higher false positive rate. The optimal configuration achieved an F1-score of 93%, demonstrating superior performance in APT detection.

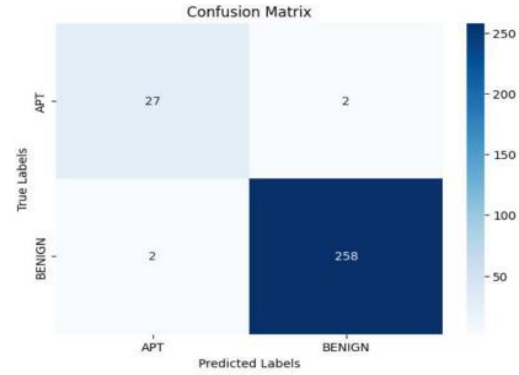


Figure 3. Confusion Matrix of the BiLAG Model on the Evaluation Dataset

Figure 3 shows that the BiLAG model correctly predicted 27 APT hosts, with 2 misclassifications, and 2 errors among 260 normal hosts. Despite the imbalanced dataset (Table 1), BiLAG maintained high performance, thanks to the combination of BiLSTM, Attention, and GCN, which enhanced information extraction and accurate classification.

2. Experiments Addressing QA2

2.1. Evaluating the Effectiveness of the BiLSTM Network

To evaluate the effectiveness of BiLSTM, we replaced it with other popular deep learning networks for information extraction, such as CNN, LSTM, and RNN.

TABLE III. EXPERIMENTAL RESULTS WHEN REPLACING BiLSTM WITH OTHER MODELS

Model	Parameter	Evaluation			
		Acc	Pre	Rec	F1
CNN	[128, 256, 512]	0.96	0.76	0.93	0.84
RNN	[256, 256]	0.94	0.77	0.67	0.71
LSTM	[256, 256]	0.97	0.84	0.87	0.85
BiLSTM	[256, 256]	<u>0.99</u>	<u>0.93</u>	<u>0.93</u>	<u>0.93</u>

Observing Table III clearly shows that BiLSTM outperforms all three traditional architectures (CNN, RNN, and LSTM) across every metric. Specifically, CNN achieves a recall of 0.93 but only a precision of 0.76, resulting in many false positives; RNN suffers from the vanishing gradient problem, as evidenced by its

recall of just 0.67 and low F1-score of 0.71; LSTM partially mitigates this with gating mechanisms raising accuracy to 0.97, precision to 0.84, recall to 0.87, and F1-score to 0.85 yet still falls short of BiLSTM by 2–9 percentage points on each metric. In contrast, with the same [256, 256] parameter configuration, BiLSTM achieves an accuracy of 0.99, precision and recall both at 0.93, and an F1-score of 0.93, demonstrating an optimal balance between sensitivity and specificity. This advantage stems from BiLSTM’s ability to leverage bidirectional sequential context, yielding richer representations for each network event and thus reducing both false positives and false negatives. These results confirm that BiLSTM is a crucial component for significantly enhancing attack detection performance in our system.

2.2. Evaluating the Effectiveness of the Attention Mechanism in the Model

In this scenario, we replaced the Attention mechanism with alternative models such as Mean and Inference to evaluate its impact on performance through experiments using BiLSTM-Mean-GCN and BiLSTM-Inference-GCN.

TABLE IV. EXPERIMENTAL RESULTS WHEN REPLACING ATTENTION WITH ALTERNATIVE MODELS

Model	Evaluation			
	Acc	Pre	Rec	F1
Mean	0.97	0.82	0.90	0.86
Inference	0.97	0.79	0.90	0.84
Attention	<u>0.99</u>	<u>0.93</u>	<u>0.93</u>	<u>0.93</u>

The results in Table 4 show a decline in performance when replacing Attention with Mean or Inference. Specifically, while the accuracy for both models remained at 97%, Precision dropped from 82% (Mean) to 79% (Inference). Recall remained stable at 90%, whereas F1-score decreased from 86% (Mean) to 84% (Inference), with the most significant drop observed in Precision. Compared to BiLAG, the Attention mechanism excelled with an accuracy increase of 2% (99%), Precision improved by 13% (93%), Recall by 3% (93%), and F1-score by 7% (93%), demonstrating significant enhancements in accuracy and balance between Precision and Recall.

2.3. Evaluating the Effectiveness of the GCN in the Model

In this scenario, we replaced GCN with GAT to compare the effectiveness of the two graph structures for APT host classification.

TABLE V. EXPERIMENTAL RESULTS WHEN REPLACING GCN WITH GAT

GAT nodes	Evaluation			
	Acc	Pre	Rec	F1
[128-128]	0.96	0.75	0.90	0.82
[128-256-512]	<u>0.96</u>	<u>0.72</u>	<u>0.93</u>	<u>0.81</u>
[128-256-512-512]	0.94	0.65	0.93	0.77

Table 5 compares results when replacing GCN with GAT across three configurations. The [128-128] configuration of GAT achieved 96% accuracy but lagged behind GCN in Precision (18% lower) and F1-score (11% lower). The [128-256-512] configuration improved Recall to 93% but still underperformed GCN in Precision (21% lower) and F1-score (12% lower). The [128-256-512-512] configuration saw a decrease in accuracy to 94% and F1-score to 77%, underperforming GCN in both Accuracy (5% lower) and F1-score (16% lower). These results indicate that GCN is the optimal choice for APT detection due to its superior capability in capturing global connections compared to GAT.

3. Experiments Addressing QA3.

TABLE VI. EXPERIMENTAL RESULTS OF APT DETECTION WITH DIFFERENT APPROACHES

Model	parameter	Evaluation			
		Acc	Pre	Rec	F1
BiLSTM-GCN[12]	2 BiLSTM - 2 GCN	0.97	0.79	0.90	0.84
CNN-LSTM[4]	4 CNN - 2 LSTM	0.96	0.92	0.65	0.76
MLP-Inference-GCN [13]	3 MLP - Inf-2GCN	<u>0.97</u>	<u>0.89</u>	<u>0.90</u>	<u>0.895</u>
BiLAG	2 BiLSTM - 1 Attention - 2 GCN	<u>0.99</u>	<u>0.93</u>	<u>0.93</u>	<u>0.93</u>

The experimental results reveal differences between network structures for APT detection. BiLSTM-GCN (2 BiLSTM layers, 2 GCN layers) achieved 97% accuracy and 90% recall but only 79% precision, with an F1-score of 84%. CNN-LSTM (4 CNN layers, 2 LSTM layers) achieved the highest precision at 92%, but its recall was the lowest at 65%, leading to an F1-score of 76%. The MIG model (3 MLP layers, 1 inference layer, 2 GCN layers) achieved similar accuracy and recall to BiLSTM-GCN but had a higher precision (89%) and the same F1-score (84%). BiLAG outperformed all models with an accuracy of 99%, precision of 93%, recall of 93%, and an F1-score of 93%.

D. Discussion

The experimental results indicate that BiLAG outperforms BiLSTM-GCN, CNN-LSTM, and MLP-Inference-GCN in APT detection. By integrating BiLSTM, Attention, and GCN, BiLAG provides deep analysis and effectively detects complex APT patterns. The model achieves high accuracy, precision, recall, and F1-score, significantly improving recall and F1-score, reducing false positives, and detecting most APT patterns. BiLAG surpasses BiLSTM-GCN due to the inclusion of the Attention mechanism, which enhances APT recognition, and it outshines CNN-LSTM with its ability to identify hard-to-detect APT patterns. Compared to MLP-Inference-GCN, BiLAG maintains superior accuracy and achieves a better balance between precision and recall.

E. CONCLUSION

This study proposes the BiLAG model for APT detection, achieving superior performance compared to the BiLSTM-GCN, CNN-LSTM, and MLP-Inference-GCN models, as evidenced by higher accuracy, precision, recall, and F1-score. The combination of BiLSTM, Attention, and GCN not only enables more accurate APT detection but also reduces false alarms, demonstrating the model's applicability to other problems such as botnet, DoS, and DDoS detection.

Although the BiLAG model has shown promising results, the study still faces several limitations. Firstly, the current training data is

limited in both scope and diversity of APT samples, which may affect the model's generalizability. Secondly, the process of converting network behaviors into graph structures for applying GCN requires a complex preprocessing pipeline, which may lead to information loss in certain cases. Additionally, the effectiveness of the model on real-world datasets and in complex network environments still needs to be thoroughly validated.

Future work will focus on expanding and diversifying the dataset, optimizing the data preprocessing pipeline, and improving the GCN architecture to enhance the model's ability to learn the latent characteristics of malicious behaviors. Moreover, applying the model to real-world scenarios and evaluating its capability to detect APT variants as well as other abnormal network behaviors is a promising research direction to further increase the system's feasibility and practical applicability.

REFERENCES

- [1] A. Alshamrani, S. Myneni, A. Chowdhary and D. Huang, "A survey on advanced persistent threats: Techniques, solutions, challenges, and research opportunities", *IEEE Communications Surveys & Tutorials*, vol. 21, no. 2, pp. 1851-1877, 2019.
- [2] Y. Wang, H. Liu, Z. Li, Z. Su and J. Li, "Combating Advanced Persistent Threats: Challenges and Solutions", *IEEE Network*, vol. 33, no. 6, pp. 324 – 333, 2024. <https://doi.org/10.1109/MNET.2024.3389734>
- [3] D. X. Cho, D. T. Huong and D. Duong, "New approach for APT malware detection on the workstation based on process profile", *Journal of Intelligent & Fuzzy Systems*, vol. 43, no. 4, pp. 4815-4834, 2022.
- [4] D. X. Cho and N. H. Cuong, "A novel approach for APT attack detection based on feature intelligent extraction and representation learning", *Plos one*, vol. 19, no. 6, 2024
- [5] S. Arefin, M. Chowdhury, R. Parvez, T. Ahmed, A. S. Abrar and F. Sumaiya, "Understanding APT detection using Machine learning algorithms: Is superior accuracy a thing?", *In 2024 IEEE International Conference on Electro Information Technology (eIT)*, pp. 532-537, May. 2024.
- [6] J. H. Joloudari, M. Haderbadi, A. Mashmool, M. GhasemiGol, S. S. Band and A. Mosavi, "Early

- detection of the advanced persistent threat attack using performance analysis of deep learning”, *IEEE Access*, vol. 8, pp. 186125-186137, 2020.
- [7] G. Xiang, C. Shi and Y. Zhang, “An APT event extraction method based on BERT-BiGRU-CRF for APT attack detection”, *Electronics*, vol. 12, no. 15, pp. 3349, 2023.
- [8] A. Graves and J. Schmidhuber, “Framewise phoneme classification with bidirectional LSTM and other neural network architectures”, *Neural networks*, vol. 18, pp. 602-610, 2005.
- [9] D. Bahdanau, “Neural machine translation by jointly learning to align and translate”, *arXiv preprint arXiv:1409.0473*, 2014.
- [10] A. Vaswani, “Attention is all you need”, *Advances in Neural Information Processing Systems*, 2017.
- [11] T. N. Kipf and M. Welling, “Semi-supervised classification with graph convolutional networks”, *arXiv preprint arXiv:1609.02907*, 2016.
- [12] D. X. Cho, D. M. Hoang and N. H. Dinh, “APT attack detection based on flow network analysis techniques using deep learning”, *Journal of Intelligent & Fuzzy Systems*, vol. 39, no. 3, pp. 4785-4801, 2020.
- [13] N. D. Hoa, D. X. Cho, N. H. Cuong and N. T. Long, “A new framework for APT attack detection based on network traffic”, *Journal of Intelligent & Fuzzy Systems*, vol. 44, no. 3, pp. 3459-3474, 2023.
- [14] D. X. Cho, D.T. Huong and N. Toan, “A novel intelligent cognitive computing-based APT malware detection for Endpoint systems”, *Journal of Intelligent & Fuzzy Systems*, vol. 43, no. 3, pp. 3527-3547, 2022. doi:10.3233/JIFS-220233.
- [15] N. H. Cuong, D. X. Cho, V. T. Long, N. D. Dat and T. Q. Anh, “A Novel Approach for APT Detection Based on Ensemble Learning Model”, *Statistical Analysis and Data Mining: The ASA Data Science Journal*, vol. 18, no. 1, 2025. <https://doi.org/10.1002/sam.70005>.

ABOUT THE AUTHOR



Nguyen Thanh Tung

Workplace: National Institute of Digital Technology and Digital Transformation, Ministry of Science and Technology, Vietnam

Email: tung_nt@mst.gov.vn

Education: Graduated with a Bachelor's and Master's in Information

Technology and Information Systems from the Posts and Telecommunications Institute of Technology, Vietnam, in 2014 and 2021.

Recent research direction: Detection of anomalous behavior of cyber attacks and targeted attacks (APT) based on Artificial Intelligence (AI) technologies.

Tên tác giả: **Nguyễn Thanh Tùng**

Cơ quan công tác: Viện Công nghệ số và Chuyển đổi số Quốc gia, Bộ Khoa học và Công nghệ

Email: tung_nt@mst.gov.vn

Quá trình đào tạo: Tốt nghiệp đại học và thạc sĩ ngành Công nghệ thông tin và Hệ thống thông tin tại Học viện Công nghệ Bưu chính Viễn Thông, Việt Nam vào các năm 2014 và 2021.

Hướng nghiên cứu hiện nay: Nghiên cứu về phát hiện hành vi bất thường của tấn công mạng và tấn công có chủ đích (APT) trên nền tảng của Trí tuệ nhân tạo (AI).



Nguyen Tien Trong

Workplace: Faculty of Information Technology, Posts and Telecommunications Institute of Technology, Vietnam

Email: trongbg2692004@gmail.com

Education: Currently pursuing a bachelor's degree in Information

Technology at the Posts and Telecommunications Institute of Technology since 2022.

Recent research direction: Detection of anomalous behavior of cyber attacks and targeted attacks (APT) based on Artificial Intelligence (AI) technologies.

Tên tác giả: **Nguyễn Tiến Trọng**

Cơ quan công tác: Khoa Công nghệ thông tin, Học viện Công nghệ Bưu chính Viễn thông

Email: trongbg2692004@gmail.com

Quá trình đào tạo: Đang học đại học ngành Công nghệ thông tin tại Học viện Công nghệ Bưu chính Viễn thông từ năm 2022 đến nay.

Hướng nghiên cứu hiện nay: Nghiên cứu về phát hiện hành vi bất thường của tấn công mạng và tấn công có chủ đích (APT) trên nền tảng của Trí tuệ nhân tạo (AI).



Nguyen Thi Hien

Workplace: Faculty of Information Technology, Posts and Telecommunications Institute of Technology, Vietnam

Email: hiennt.b22cn289@stu.ptit.edu.vn

Education: Currently pursuing a bachelor's degree in Information Technology at the Posts and Telecommunications Institute of Technology since 2022.

Recent research direction: Detection of anomalous behavior of cyber attacks and targeted attacks (APT) based on Artificial Intelligence (AI) technologies.

Tên tác giả: **Nguyễn Thị Hiền**

Cơ quan công tác: Khoa Công nghệ thông tin, Học viện Công nghệ Bưu chính Viễn thông

Email: hiennt.b22cn289@stu.ptit.edu.vn

Quá trình đào tạo: Đang học đại học ngành Công nghệ thông tin tại Học viện Công nghệ Bưu chính Viễn thông từ năm 2022 đến nay.

Hướng nghiên cứu hiện nay: Nghiên cứu về phát hiện hành vi bất thường của tấn công mạng và tấn công có chủ đích (APT) trên nền tảng của Trí tuệ nhân tạo (AI).

Email: chodx@ptit.edu.vn

Quá trình đào tạo: Tốt nghiệp đại học, thạc sĩ và tiến sĩ ngành Khoa học máy tính và Cơ sở máy tính tại Đại học Tổng Hợp Kỹ Thuật Điện Saint- Peterburg, Liên bang Nga vào các năm 2008, 2010 và 2014.

Hướng nghiên cứu hiện nay: Nghiên cứu về phát hiện hành vi bất thường của tấn công mạng và tấn công có chủ đích (APT) trên nền tảng của Trí tuệ nhân tạo (AI).



Nguyen Quang Hoan

Workplace: National Institute of Digital Posts and Telecommunications Institute of Technology, Ministry of Science and Technology, Vietnam

Email: quanghoanptit@gmail.com

Education: Graduated with a Bachelor's and PhD in Computer Engineering, Automation, and Information Systems from Moscow, former Soviet Union, and the Institute of Information Technology, National Academy of Science and Technology, in 1973 and 1986 respectively.

Recent research direction: Information Systems and Artificial Neural Networks.

Tên tác giả: **Nguyễn Quang Hoan**

Cơ quan công tác: Học viện Công nghệ Bưu chính Viễn thông, Bộ Khoa học và Công nghệ

Email: quanghoanptit@gmail.com

Quá trình đào tạo: Tốt nghiệp đại học, và tiến sĩ ngành Kỹ thuật Máy tính, Tự động và Hệ thống thông tin tại Moskva, Liên xô cũ và Viện Công Nghệ Thông tin, Viện Hàn lâm Khoa học, Công nghệ Quốc gia vào các năm 1973 và 1986.

Hướng nghiên cứu hiện nay: Hệ thống Thông tin và Mạng nơ ron nhân tạo.



Do Xuan Cho

Workplace: Faculty of Information Security, Posts and Telecommunications Institute of Technology, Vietnam

Email: chodx@ptit.edu.vn

Education: He received his BSc, MSc and PhD degrees in Computer science and computer facilities from the Saint Petersburg Electrotechnical University, Russia. In 2008, 2010 and 2014 respectively.

Recent research direction: Detection of anomalous behavior of cyber attacks and targeted attacks (APT) based on Artificial Intelligence (AI) technologies.

Tên tác giả: **Đỗ Xuân Chợ**

Cơ quan công tác: Khoa An toàn thông tin, Học viện Công nghệ Bưu chính Viễn thông