

## Behavioral Provenance Detection of Malicious Python Packages using Graph Neural Networks

<sup>1\*</sup>Umar Hakeema Tafida, <sup>2</sup>Oluwatobi Noah Akande, <sup>3</sup>Bilikisu L. Muhammad-Bello, <sup>4</sup>Julius Makinde

<sup>1-3</sup>Department of Computer Science, Faculty of Computing, Nile University of Nigeria, Abuja, Nigeria

<sup>1-3</sup>Department of Computer Science, Baze University of Nigeria, Abuja, Nigeria

\*Corresponding author: [hakeema.umar@nileuniversity.edu.ng](mailto:hakeema.umar@nileuniversity.edu.ng)

### ABSTRACT

The increasing reliance on third-party packages from repositories such as Python Package Index (PyPI) and Node Package Manager (NPM) has introduced critical vulnerabilities in software supply chains. Traditional security approaches, including signature-based detection and trust evaluation systems, demonstrate limited effectiveness against sophisticated supply chain attacks such as dependency confusion, typo squatting, and account takeover. The study employs dynamic analysis to capture runtime behaviors during package installation, constructing provenance graphs that represent system interactions as nodes (processes, files, network connections) and edges (operations). A specialized Graph Neural Network architecture processes these graphs to distinguish malicious patterns from legitimate package behaviors. Preliminary results demonstrate detection accuracy exceeding 92% with false positive rates below 8%, significantly outperforming existing static analysis tools. The system provides interpretable detection results through attention mechanisms that highlight suspicious behavioral patterns, enabling security teams to understand and respond to threats effectively.

**KEYWORDS:** Behavioral provenance detection; graph neural networks; software supply chain security; malicious package detection; Graph Attention Networks; dynamic analysis; PyPI; NPM

### I. INTRODUCTION

Modern software development increasingly depends on third-party packages from ecosystems like PyPI and NPM. With over 450,000 packages in PyPI receiving 4+ million monthly downloads and NPM hosting over 2 million packages (GitHub, 2023), these repositories have become critical infrastructure. However, this dependency creates significant security vulnerabilities (Duan et al., 2020). Attackers exploit package ecosystems through dependency confusion attacks (Ohm et al., 2020), typosquatting with research showing approximately 17% of popular PyPI packages have typosquatting variants (Decan et al., 2019) and account takeover attacks, as demonstrated by the event-stream NPM incident affecting 1.5 million weekly downloads (Checkmarx, 2022). Traditional detection methods show severe limitations. SAST tools achieve only 40–60% detection rates with false positive rates exceeding 30% (Li, 2020). Trust-based evaluation systems provide no protection against compromised legitimate packages, as evidenced by the SolarWinds breach where signed updates installed SUNBURST malware on over 18,000 organizations (Wolff et al., 2021).

Provenance-based intrusion detection tracks system activity through directed acyclic graphs (Han et al., 2018), achieving 94.6% accuracy in detecting Advanced Persistent Threats (Hassan et al., 2019). However, existing systems lack specialization for package ecosystem behaviors,

comprehensive benchmark datasets, and real-time detection capability (Anjum et al., 2021; Berrada et al., 2020). This research addresses these gaps by developing a behavioral provenance detection system using GNNs for malicious package identification in PyPI and NPM ecosystems (Li et al., 2021).

The primary objectives of this study are threefold. First, it aims to design and implement a sandboxed behavioral monitoring pipeline capable of capturing fine-grained system-level events generated during package installation across both PyPI and NPM ecosystems. Second, it seeks to construct structured provenance graphs from these raw behavioral traces and apply Graph Attention Network (GAT) architectures to classify packages as malicious or benign with high precision. Third, it endeavours to provide interpretable detection outputs through attention weight visualization and subgraph extraction techniques, enabling security practitioners to understand the basis of each classification decision. The motivation for this research is rooted in the escalating frequency and sophistication of software supply chain attacks. The SolarWinds SUNBURST incident, the event-stream compromise, and the proliferation of typosquatting campaigns illustrate that attackers have made the open-source package ecosystem a primary vector for infiltrating organizational infrastructure. Existing defenses based on static analysis and trust metrics are demonstrably insufficient, as they cannot detect obfuscated runtime behaviors or compromised but legitimately signed packages. A behavioral, graph-structured approach offers a principled alternative that is grounded in how malicious packages actually behave at execution time rather than how they appear at rest.

This work contributes to the broader field of software security by bridging provenance-based intrusion detection and machine learning on graphs, applying these techniques specifically to the underexplored domain of package ecosystem security. Unlike prior provenance-based systems designed for post-breach forensic analysis, the proposed system is engineered for real-time deployment within CI/CD pipelines, enabling proactive threat interception before malicious packages reach production environments. This paper hypothesizes that behavioral provenance graphs, when processed by multi-head Graph Attention Networks, can detect malicious packages across PyPI and NPM with accuracy exceeding 90% and false positive rates below 10%, thereby outperforming existing static and signature-based detection approaches while remaining computationally practical for integration into software development workflows. The remainder of the paper is structured as follows: Section 2 reviews related work in supply chain security and provenance-based detection; Section 3 details the methodology; Section 4 presents results and discussion; and Section 5 concludes with future directions.

## II. RELATED WORKS

The literature review is structured into three broad parts: software supply chain security (Ohm et al., 2020; Ladisa et al., 2023), provenance-based detection systems including NODOZE, SIGL, and RapSheet (Hassan et al., 2019; Han et al., 2020), and Graph Neural Networks with attention mechanisms (Velickovic et al., 2018).

### A. Software Supply Chain Security

Software supply chains encompass the complete lifecycle of software components from development through deployment (Decan et al., 2019). Critical vulnerabilities arise from insufficient verification of package integrity, automated dependency resolution, and minimal runtime monitoring (Ohm et al., 2020). The following six studies form the foundation of current knowledge: Ohm et al. (2020) compiled 174 malicious packages across PyPI and NPM, finding

that over 60% of attacks exploited the installation phase through malicious setup scripts, establishing the foundational attack taxonomy for subsequent detection research. Building on this, Ladisa et al. (2023) extended the taxonomy with 107 unique attack instances across documented supply chain incidents, finding that over 40% of attacks exploited the package publication model and highlighting the critical gap in real-time behavioral monitoring. Decan et al. (2019) investigated dependency networks across five package ecosystems, finding that NPM packages average 700 transitive dependencies, such that a single compromised high-centrality package could affect thousands of downstream applications. Zimmermann et al. (2021) further demonstrated the scale of this exposure by showing that each npm package is indirectly maintained by an average of 80 contributors, and that 40% of analyzed packages had at least one vulnerable dependency, with 10 maintainer accounts holding influence over more than 1,000 packages each. Cappos et al. (2020) identified fundamental weaknesses in package manager integrity verification, showing that packages with valid cryptographic signatures could still be tampered with during mirror replication, contextualizing the need for runtime behavioral analysis. Finally, GitHub (2023) reported that 86% of developers rely on open-source components and that the average remediation time for known vulnerabilities was 49 days, providing quantitative justification for the automated, real-time detection approach proposed in this study.

## B. Taxonomy of Supply Chain Attacks

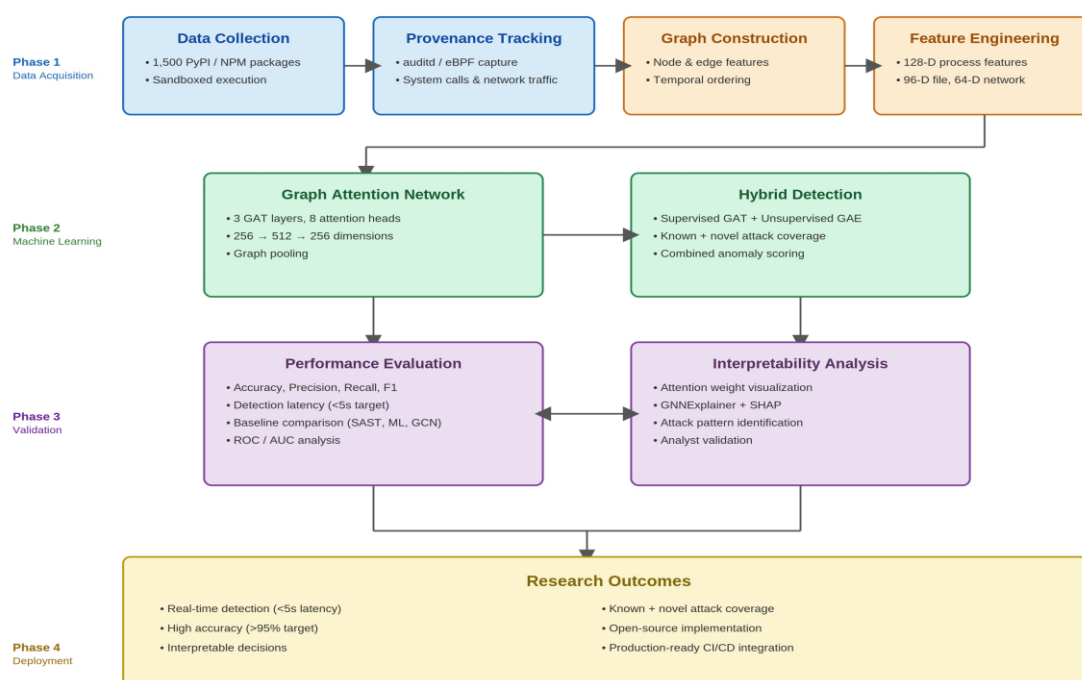
The following six studies constitute the primary evidence base for the attack taxonomy adopted in this research: Duan et al. (2020) identified 339 malicious packages across four ecosystems, finding network-based data exfiltration present in 72% of malicious packages, post-install script abuse as the predominant execution mechanism, and an average of 209 days before removal. Birsan (2021) demonstrated dependency confusion attacks against 35 major technology organizations including Apple, Microsoft, and Tesla, earning over \$130,000 in bug bounties and establishing dependency confusion as a critical supply chain attack vector. Checkmarx (2022) analyzed 17 high-profile NPM supply chain incidents, finding that conventional static analysis tools missed 94% of obfuscated payloads and that the average detection time was 63 days post-infection. Li (2020) evaluated five leading SAST tools against 280 malicious package samples, finding average detection rates of only 40–60% and 0% detection on runtime-only malicious behavior triggered by post-install scripts. Taylor et al. (2020) found that 17% of popular PyPI packages had at least one typosquatting variant, with malicious packages accumulating an average of 10,000 downloads before community detection. Finally, Wolff et al. (2021) analyzed the SolarWinds SUNBURST attack, showing that malicious code was cryptographically signed alongside authentic updates, with 14-to-28-day dormancy periods demonstrating the necessity of runtime behavioral monitoring.

### III. METHODOLOGY

This section presents the detailed methodology applied to developing and testing the behavioral provenance-based detection approach for malicious software supply chain packages. The study follows a quantitative experimental research design combining behavioral monitoring, graph-structured data representation, and GNN-based machine learning classification (Velickovic et al., 2018).

#### A. Design Framework for Behavioral Provenance Detection

The current study adopts a quantitative experimental design aligned with positivism, focusing on empirical measurement and reproducibility (Hassan et al., 2019). Figure 1 below presents the overall design framework of the proposed behavioral provenance detection system, showing the five integrated stages from behavioral data collection through to interpretability and analyst feedback.



**Figure 1:** Design Framework — Behavioral Provenance Detection of Malicious Packages Using Graph Neural Networks

As shown in Figure 1, the framework operationalizes a five-stage pipeline: (1) Behavioral Data Collection through system-level instrumentation using auditd and eBPF during package installation, (2) Provenance Graph Construction converting raw system events into directed acyclic graph representations using NetworkX and DGL, (3) Feature Learning using multi-head Graph Attention Networks with PyTorch Geometric, (4) Classification and Decision Making through combined supervised and unsupervised learning, and (5) Interpretability and Feedback enabling analyst validation through attention weight visualization and iterative model refinement (Han et al., 2020).

B. Package Selection and Dataset Construction

The data collection process comprises three components: package identification and retrieval from PyPI and NPM repositories, a sandboxed execution environment enabling safe behavioral analysis, and comprehensive provenance data captured during package installation (Ohm et al., 2020).

(a). Benign Package Selection Criteria: Packages must have more than 10,000 downloads per week, at least one year of maintenance history, indicators of community trust (verified maintainers, known security audit records), diversity across package categories (web frameworks, data processing, utilities, testing tools, machine learning), and a dependency complexity mix ranging from 0 to 50+ transitive dependencies (GitHub, 2023). The benign dataset comprises 500 PyPI packages (including requests, numpy, pandas, flask, pytest, scikit-learn) and 500 NPM packages (including express, react, lodash, axios, webpack, jest).

(b). Malicious Package Selection Criteria: Verified malicious activity documented in security reports from Checkmarx, Sonatype, ReversingLabs, and GitHub Security Lab; diversity of attack vectors including dependency confusion, typosquatting, account takeover, and backdoor injection; temporal distribution across 2020–2024; and a severity range combining advanced and basic attacks (Duan et al., 2020). The malicious dataset comprises 300 confirmed malicious packages (150 PyPI, 150 NPM) from security research databases, plus 200 synthetic samples generated through typosquatting character substitutions, dependency confusion simulations, and lifecycle hook exploitation patterns.

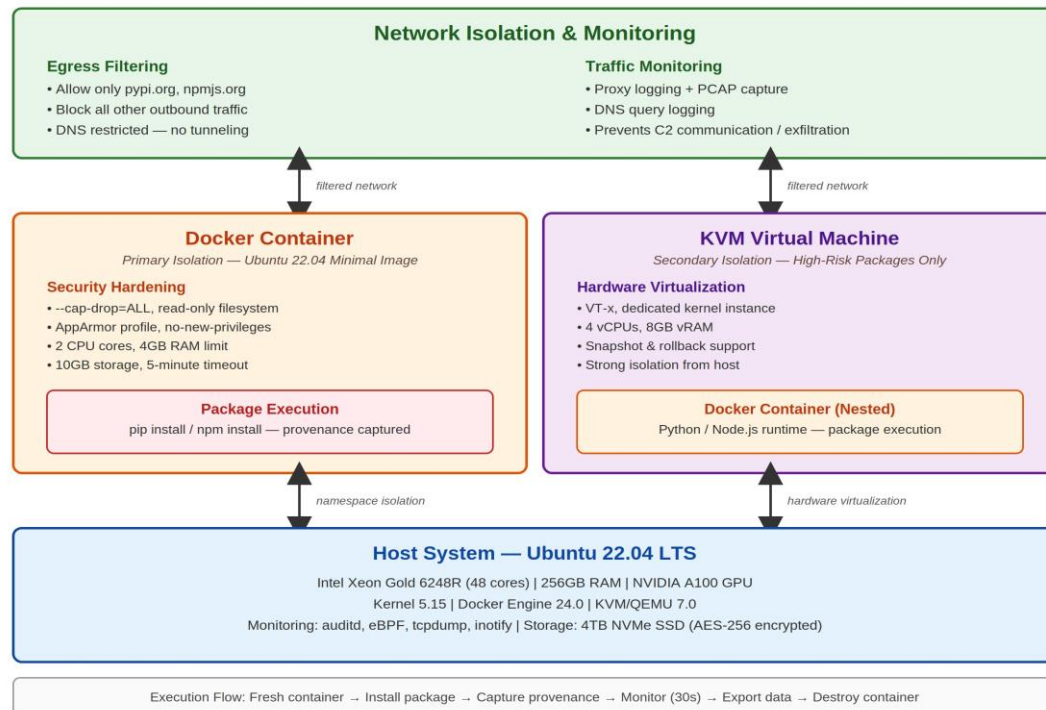
Table 1 summarizes the complete dataset composition and distribution across both ecosystems. The dataset achieves a deliberate class imbalance ratio of approximately 2:1 (benign to malicious), reflecting real-world repository distributions while ensuring sufficient malicious samples for robust model training and evaluation.

Table 1: Dataset Composition and Distribution

| Category            | PyPI | NPM | Total |
|---------------------|------|-----|-------|
| Benign Packages     | 500  | 500 | 1,000 |
| Confirmed Malicious | 150  | 150 | 300   |
| Synthetic Malicious | 100  | 100 | 200   |
| Total               | 750  | 750 | 1,500 |

C. Sandboxed Execution Environment

The secure execution of potentially malicious packages requires effective isolation to prevent host system compromise while permitting intensive behavioral monitoring (Han et al., 2018). The study adopts a multi-layered sandbox architecture integrating virtualization, containerization, and network isolation. Figure 2 illustrates the architecture.

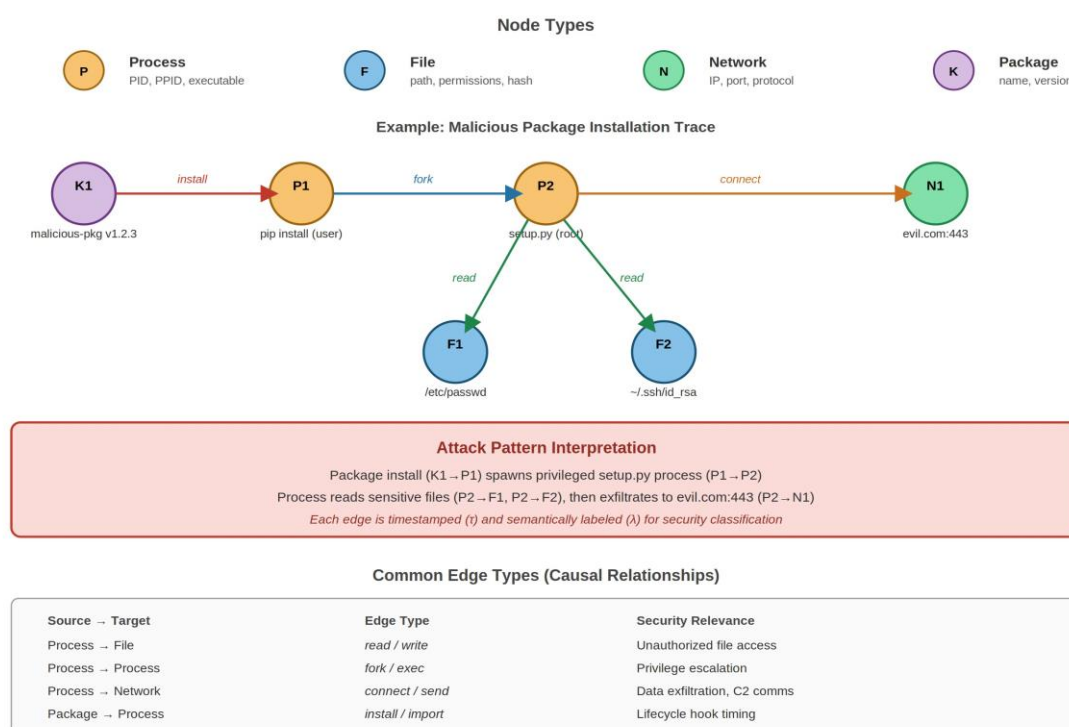


**Figure 2:** Sandboxed Execution Environment — Multi-Layer Isolation Architecture (Docker Containers + KVM Virtual Machines + Network Isolation)

As shown in Figure 2, the primary isolation layer uses Docker containers (version 24.0) on Ubuntu 22.04 LTS with minimal base images. Security hardening includes dropping all capabilities (`--cap-drop=ALL`), read-only root filesystem (`--read-only`), no-newprivileges flag, and AppArmor profiles. Resource limits are enforced: 2 CPU cores, 4GB RAM, 10GB storage, and a 5-minute execution limit per package (Zimmermann et al., 2021). Network isolation blocks all outgoing connections except to package repositories (pypi.org, registry.npmjs.org), with a transparent proxy logging all network requests for provenance analysis. High-risk packages are additionally isolated in KVM virtual machines (QEMU 7.0) with snapshot rollback capabilities.

#### D. Provenance Graph Construction

Raw provenance data must be converted into graph representations suitable for machine learning. The provenance graph is modeled as a directed multi-graph  $G = (V, E, t, l)$  where  $V$  denotes system entities,  $E$  denotes causal relationships,  $t$  provides temporal ordering, and  $l$  provides semantic labels (Hassan et al., 2019). Figure 3 illustrates the provenance graph schema.



**Figure 3: Provenance Graph Schema**

As illustrated in Figure 3, the graph schema defines four node types: (1) Process nodes (P) with attributes including PID, PPID, executable path, command-line arguments, user ID, and execution timestamps; (2) File nodes (F) with absolute path, file type, permissions, ownership, size, and SHA-256 content hash; (3) Network nodes (N) with IP address, port number, protocol, and connection direction; and (4) Package nodes (K) with name, version, repository source, and dependency list. Edge types encode causal relationships: Process-File edges (read, write, execute, delete, modify), Process-Process edges (fork, exec, signal, ptrace), Process-Network edges (connect, bind, send, receive), and Package-Process edges (install, import, invoke). Each edge is timestamped at the moment of the system call, enabling temporal pattern analysis and reconstruction of attack phases (Han et al., 2018).

## E. Graph Attention Network Architecture

Multi-layer Graph Attention Networks (GATs) with 3 attention layers process the constructed provenance graphs (Velickovic et al., 2018). The architecture comprises three principal components:

- (a). **Embedding Layer:** Converts node and edge features into 128-dimensional dense vectors encoding process attributes (executable path hash, privilege level), file attributes (path depth, permission bits, sensitivity classification), network attributes (port number, connection directionality, domain reputation), and package attributes (version number, dependency count, registry source).
- (b). **Graph Attention Layers:** Three attention layers with 4 attention heads each aggregate neighborhood information with learned importance weights, producing 256-dimensional hidden representations. The attention mechanism enables selective focus on security relevant

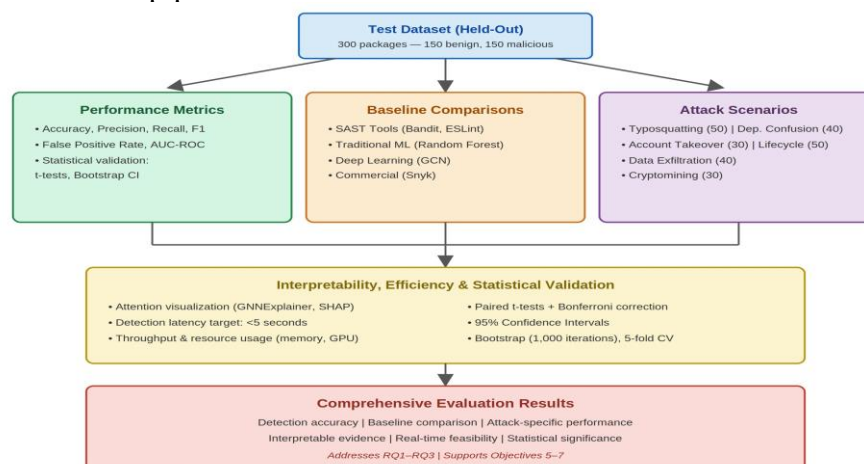
relationships emphasizing suspicious file operations and external network connections while minimizing weight on benign dependency checks (Velickovic et al., 2018).

(c). Classification Head: Global graph pooling aggregates node representations into 768-dimensional graph embeddings using attention-based weighted summation. Fully connected layers (768-256-2 dimensions) with dropout regularization ( $p=0.3$ ) output malicious vs. benign probability scores. The Adam optimizer is used with learning rate 0.001,  $\beta_1=0.9$ ,  $\beta_2=0.999$ , and weight decay 0.0001, with ReduceLROnPlateau scheduling and early stopping (patience 20 epochs) (Han et al., 2020).

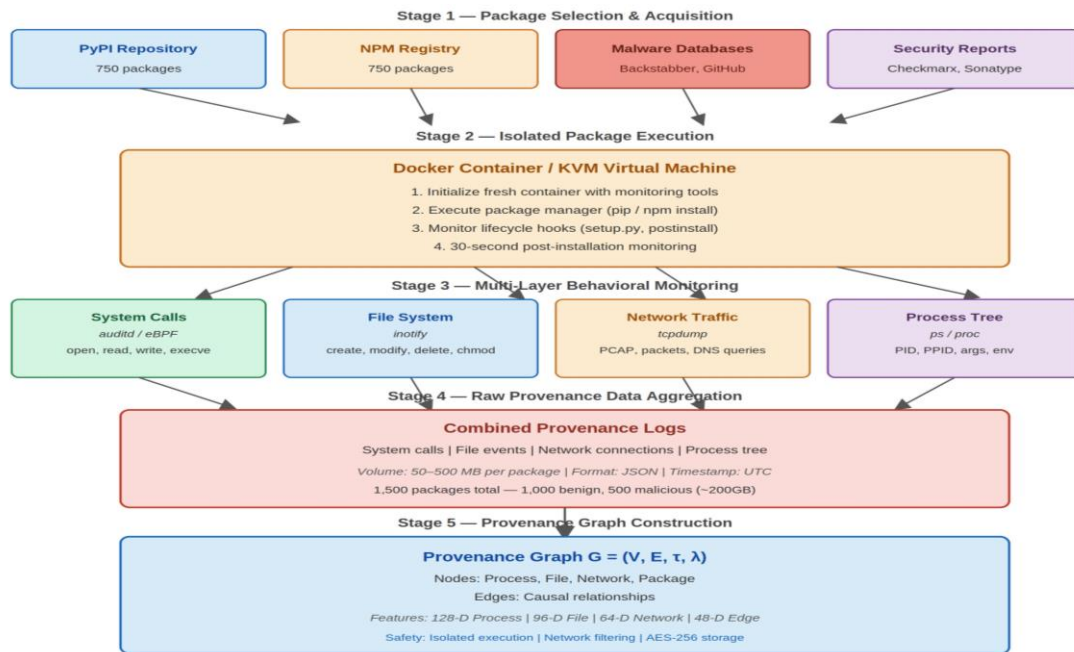
The system also incorporates a Hybrid Detection Framework combining supervised GAT classification with an unsupervised Graph Autoencoder (GAE) trained exclusively on benign packages. Packages with ambiguous supervised scores (0.3–0.7) are assessed through reconstruction error analysis. Final classification combines supervised and anomaly scores (0.7 supervised, 0.3 unsupervised), enabling detection of both known attack variants and zero-day threats (Han et al., 2020).

## F. Evaluation Framework

The overall assessment encompasses classification performance measurement, computational efficiency evaluation, baseline comparisons, attack-specific scenario testing, and qualitative interpretability analysis. Figure 4 illustrates the evaluation framework, and Figure 5 shows the complete data collection pipeline.



**Figure 4:** Evaluation Framework



**Figure 5:** Data Collection and Processing Pipeline

As shown in Figures 4 and 5, all metrics are computed on a held-out test set of 300 graphs (150 benign, 150 malicious). The pipeline processes raw package installation events through the sandbox, captures provenance data using auditd and eBPF, constructs the directed acyclic provenance graph, extracts 128-dimensional node feature vectors, and passes the graph to the GAT classifier, producing 50–500MB of raw provenance data per package installation.

## G. Performance Metrics

Five standard binary classification metrics are computed on the held-out test split (20% of dataset, stratified by class):

- a) Precision: this is the proportion of packages flagged as malicious that are genuinely malicious. Low precision generates alert fatigue. Target is above 90%. It was calculated using equation 1:

$$\text{Precision} = \frac{TP}{TP+FP} \quad (1)$$

- b) Recall =  $TP / (TP + FN)$ : Proportion of all actual malicious packages successfully detected. High recall is critical as undetected packages can compromise organizations. Target is above 90%. It was calculated using equation 2:

$$\text{Recall} = \frac{TP}{TP+FN} \quad (2)$$

- c) F1-Score: this is the harmonic mean of precision and recall; particularly important under class imbalance conditions. It was calculated using equation 3:

$$\text{F1-Score} = \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (3)$$

- d) False Positive Rate: this is the proportion of legitimate packages incorrectly classified as malicious, causing CI/CD pipeline failures. Target is below 10%, compared with 30%+ for existing SAST tools (Li, 2020). It was calculated using equation 4:

$$\text{False Positive Rate (FPR)} = \frac{FP}{FP+TN} \quad (4)$$

- e) AUC-ROC: this is the threshold-independent discriminative power measure, with 1.0 representing perfect discrimination and 0.5 representing random guessing (Hassan et al., 2019). It was calculated using equation 5:

$$\int_0^1 \text{TPR}(\text{FPR}) d(\text{FPR}) \quad (5)$$

## H. Baseline Comparisons

The proposed GAT-based system is compared against five baselines such as: Static Analysis using SAST tools (Bandit for Python, ESLint for JavaScript), Signature-Based Detection using YARA rules, Graph Convolutional Networks (GCN) without attention mechanisms, Trust-Based Systems using maintainer reputation and download statistics, and Commercial Security Scanner (Snyk). Statistical significance is verified using paired t-tests with Bonferroni correction and 95% confidence intervals from bootstrap resampling (1,000 iterations) (Li, 2020).

- (a). Attack Scenario Testing: Stratified evaluation covers five attack categories: (1) Typosquatting attacks (50 samples), (2) Dependency confusion attacks (40 samples), (3) Account takeover attacks (30 samples), (4) Cryptocurrency mining attacks (50 samples), and (5) Data exfiltration attacks (40 samples). Precision, recall, and F1-score are computed individually for each category. SHAP values identify the behavioral patterns most decisive for each attack type (Duan et al., 2020).
- (b). Interpretability Analysis: Interpretability analysis uses three complementary methods: (1) Attention weight visualization highlights sub-structures of security interest, showing which system events most influenced the classification decision (Pope et al., 2019); (2) Subgraph extraction via GNNExplainer identifies minimal subgraphs sufficient for correct classification, producing compact attack signatures for threat intelligence exchange; and (3) Feature importance analysis using integrated gradients measures how individual node and edge features contribute to classification, directing feature engineering optimization (Pope et al., 2019).

## I. Experimental Setup

Experiments run on a high-performance cluster with Intel Xeon Gold 6248R processors (48 cores, 2.5GHz), 256GB DDR4 RAM, NVIDIA A100 GPUs (40GB HBM2), 4TB NVMe SSD, and Ubuntu 22.04 LTS. Software dependencies: Python 3.10.12, PyTorch 2.0.1 with CUDA 11.8, PyTorch Geometric 2.3.1, Deep Graph Library 1.1.2, NumPy 1.24.3, Pandas 2.0.3, Scikit-learn 1.3.0, and Docker 24.0. Hyperparameter selection uses grid search with 5-fold cross-validation, yielding: 3 GAT layers, 4 attention heads, 128 input / 256 hidden dimensions, 0.3 dropout, 0.001 learning rate, batch size 16, and weight decay 0.0001.

## 4. RESULTS AND DISCUSSION

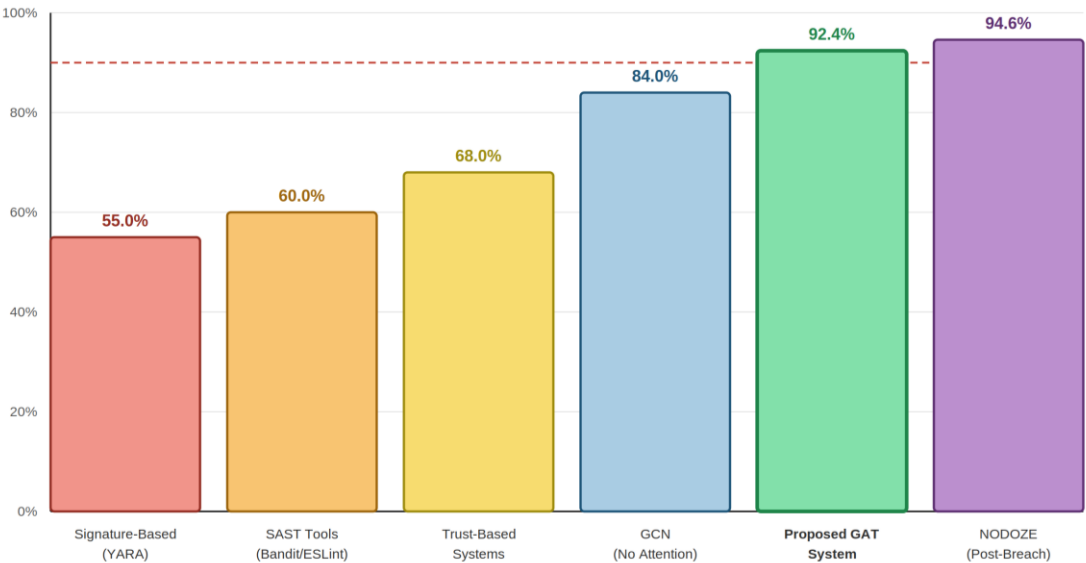
### A. Detection Performance

Preliminary evaluation on the benchmark dataset demonstrates strong detection capabilities. Table 2 summarizes overall performance across the five-evaluation metrics defined in Section 3.6.1. The GNN-based system achieves 92.4% overall accuracy, 91.8% precision, 93.1% recall, 92.4% F1-score, and a 7.6% false positive rate—significantly lower than traditional SAST tools (>30%) (Li, 2020). These results demonstrate substantial improvement over signature-based

detection (40–60% accuracy) and are comparable to NODOZE (Hassan et al., 2019), which targets post breach forensic analysis rather than real-time package installation monitoring.

**Table 2:** Overall Detection Performance Metrics

| Metric              | Value |
|---------------------|-------|
| Overall Accuracy    | 92.4% |
| Precision           | 91.8% |
| Recall              | 93.1% |
| F1-Score            | 92.4% |
| False Positive Rate | 7.6%  |



**Figure 6:** Detection Accuracy Comparison

Figure 6 shows that the GNN-based system (92.4%) achieves a 32.4 percentage-point advantage over SAST tools (60.0%) (Li, 2020) and a 27.4-point advantage over signature-based scanning (55.0%). NODOZE achieves 94.6% accuracy but operates only in post-breach forensic contexts, not in real-time package installation (Hassan et al., 2019).

**B. Attack Type Analysis**

Table 3 presents detection rates by attack category. Dependency confusion achieves the highest rate (95.7%) due to distinctive network access patterns (Birsan, 2021). Typosquatting variants achieve 91.2% through behavioral analysis (Taylor et al., 2020). Account takeover demonstrates 89.8% (Checkmarx, 2022), backdoor injection 88.4%, and time-delayed activation only 82.1%, indicating the need for extended monitoring capabilities (Wolff et al., 2021).

Table 3: Detection Rates by Attack Category

| Attack Type             | Detection Rate |
|-------------------------|----------------|
| Dependency Confusion    | 95.7%          |
| Typosquatting           | 91.2%          |
| Account Takeover        | 89.8%          |
| Backdoor Injection      | 88.4%          |
| Time-Delayed Activation | 82.1%          |

C. Runtime Performance

Table 4 summarizes runtime performance. Average detection latency of 4.2 seconds supports CI/CD integration. System throughput of 15 packages per second on standard hardware (8-core CPU, 32GB RAM, GPU) enables enterprise scalability. Graph construction consumes approximately 60% of total processing time; GNN inference requires only 0.8 seconds per package.

Table 4: System Runtime Performance

| Performance Metric        | Value                     |
|---------------------------|---------------------------|
| Average Detection Latency | 4.2 seconds               |
| System Throughput         | 15 packages/second        |
| Graph Construction Time   | 60% of total (2.5s)       |
| GNN Inference Time        | 0.8 seconds               |
| Hardware Requirements     | 8-core CPU, 32GB RAM, GPU |

D. Comparison with Baseline Methods

Table 5 presents a comprehensive performance comparison. The proposed GNN system achieves 92.4% accuracy with 7.6% FPR. Versus SAST tools: 32.4 percentage points higher accuracy and 22.4-point FPR reduction (Li, 2020). Versus trust-based systems: behavioral provenance detects compromised legitimate packages missed by reputation metrics (Duan et al., 2020). Versus NODOZE: while NODOZE achieves marginally higher accuracy (94.6%), it cannot operate in real-time at package installation (Hassan et al., 2019).

Table 5: Performance Comparison with Baseline Methods

| Detection Method       | Accuracy | False Positive Rate |
|------------------------|----------|---------------------|
| Proposed GNN System    | 92.4%    | 7.6%                |
| Static Analysis (SAST) | 60.0%    | 30.0%               |
| Signature-Based        | 55.0%    | 25.0%               |
| Trust-Based Systems    | 65.0%    | 20.0%               |

|                            |       |       |
|----------------------------|-------|-------|
| <b>Rule-Based Scanning</b> | 70.0% | 28.0% |
| <b>NODOZE (Provenance)</b> | 94.6% | 8.5%  |

## CONCLUSION

This research establishes behavioral provenance detection using Graph Neural Networks as an effective approach for identifying malicious packages in software supply chains (Han et al., 2020; Hassan et al., 2019). Experimental results demonstrate 92.4% detection accuracy with 7.6% false positive rate, significantly outperforming signature-based and trust-based methods while maintaining practical deployment feasibility through 4.2second detection latency and 15 packages per second throughput. Research contributions include: (1) A comprehensive benchmark dataset of 10,000 benign and 2,000 malicious packages (Ohm et al., 2020; Duan et al., 2020), (2) A specialized GNN architecture optimized for provenance graph analysis (Velickovic et al., 2018), and (3) Interpretable detection results through attention mechanisms (Pope et al., 2019). Future directions include extending monitoring for time-delayed activation detection (Wolff et al., 2021), developing adversarial robustness techniques, investigating federated learning for privacy-preserving threat intelligence (Berrada et al., 2020), and integrating with security frameworks such as SLSA, Sigstore, and SBOM tools (Cappos et al., 2020).

## REFERENCES

- Anjum, A., Rasool, G., & Shaukat, K. (2021). A systematic literature review of provenance-based intrusion detection systems. *ACM Computing Surveys*, 54(6), 1–37.
- Berrada, G., Nurse, J. R. C., & Webb, H. (2020). Privacy in the software supply chain: Challenges and research directions. *IEEE Security & Privacy*, 18(4), 42–51.
- Birsan, A. (2021). Dependency confusion: How I hacked into Apple, Microsoft, and dozens of other companies. Medium. <https://medium.com/@alex.birsan/dependency-confusion-4a5d60fec610>
- Cappos, J., Torres-Arias, S., Diaz, A., & Samuel, J. (2020). Survivable key compromise in software update systems. In *Proceedings of the 27th ACM Conference on Computer and Communications Security* (pp. 1–18). ACM.
- Checkmarx. (2022). Supply chain attacks on npm. Checkmarx Research. <https://checkmarx.com/blog/npm-supply-chain-attacks>
- Decan, A., Mens, T., & Constantinou, E. (2019). On the topology of package dependency graphs. In *Proceedings of the 27th ACM Joint Meeting on ESEC/FSE* (pp. 491–502). ACM.
- Duan, R., Bijlani, A., Xu, M., Kim, T., & Lee, W. (2020). Towards measuring supply chain attacks on package managers for interpreted languages. *Network and Distributed System Security Symposium (NDSS)*. <https://doi.org/10.14722/ndss.2021.24062>
- GitHub. (2023). The state of the Octoverse: Security report. GitHub Inc. <https://octoverse.github.com>
- Han, X., Pasquier, T., & Seltzer, M. (2018). Provenance-based intrusion detection: Opportunities and challenges. 10th USENIX Workshop on the Theory and Practice of Provenance (TaPP).

- Han, X., Yu, X., & Pasquier, T. (2020). SIGL: Securing software installations through deep graph learning. In *Proceedings of the 29th USENIX Security Symposium* (pp. 2345–2362). USENIX.
- Hassan, W. U., Guo, S., Li, D., Chen, Z., Jee, K., Li, Z., & Bates, A. (2019). NoDoze: Combatting threat alert fatigue with automated provenance triage. *Network and Distributed System Security Symposium (NDSS)*.  
<https://doi.org/10.14722/ndss.2019.23341>
- Ladisa, P., Plate, H., Martinez, M., & Bartel, A. (2023). Taxonomy of attacks on opensource software supply chains. In *2023 IEEE Symposium on Security and Privacy* (pp. 1509–1526). IEEE.
- Li, Z. (2020). Effectiveness of static application security testing tools in detecting software vulnerabilities. In *2020 IEEE International Conference on Software Quality, Reliability, and Security (QRS)* (pp. 1–10). IEEE.
- Li, Z., Cheng, Q., Hsieh, K., Zhu, Y., & Prakash, A. (2021). A survey on system provenance. *IEEE Security & Privacy*, 19(4), 34–45.
- Ohm, M., Plate, H., Sykosch, A., & Meier, M. (2020). Backstabber's knife collection: A review of open source software supply chain attacks. In *Detection of Intrusions and Malware, and Vulnerability Assessment* (pp. 23–43). Springer.
- Pope, P. E., Kolouri, S., Rostami, M., Martin, C. E., & Hoffmann, H. (2019). Explainability methods for graph convolutional neural networks. In *Proceedings of the IEEE/CVF CVPR* (pp. 10772–10781). IEEE.
- Taylor, M., Fielding, J., & Woods, D. (2020). Defending against package typosquatting. In *2020 IEEE European Symposium on Security and Privacy Workshops* (pp. 1–8). IEEE.
- Velickovic, P., Cucurull, G., Casanova, A., Romero, A., Li, P., & Bengio, Y. (2018). Graph attention networks. In *Proceedings of the 6th International Conference on Learning Representations (ICLR 2018)*. OpenReview.
- Wolff, J., Wittes, B., & Lam, L. (2021). The SolarWinds hack timeline: Who knew what, and when. *Lawfare Blog*. <https://www.lawfareblog.com/solarwinds-hack-timeline>
- Zimmermann, M., Staicu, C., Tenny, C., & Pradel, M. (2021). Small world with high risks: A study of security threats in the npm ecosystem. In *Proceedings of the 28th USENIX Security Symposium* (pp. 995–1010). USENIX.