

Deep Learning for Malware Detection: Transformer-Based Analysis of Windows Executables

Elshan Baghirov

Institute of Information Technology, Baku, Azerbaijan
elsenbagirov1995@gmail.com

Abstract— Malware detection is challenged by unprecedented threats exceeding 560,000 new variants daily. A Transformer-based deep learning approach for detecting malicious Windows Portable Executable files is presented. Raw byte sequences are processed using self-attention mechanisms, eliminating manual feature engineering. A multi-modal framework combining raw bytes, PE metadata, and entropy features is implemented. Through experimental evaluation on EMBER2024 (3.2M samples) and SOREL-20M (20M samples), 95.1% accuracy with 0.4% false positive rate is demonstrated, significantly outperforming baseline methods including LightGBM (92.7%), MalConv (93.4%), and LSTM approaches (91.8%). Explainability for forensic analysis is provided through attention visualization.

Keywords— *malware detection; deep learning; transformer architecture; self-attention mechanism; Windows executables; multi-modal fusion*

I. INTRODUCTION

The exponential proliferation of malicious software is recognized as a critical threat to modern computing infrastructure. Approximately 560,000 new malware variants are identified daily, with cumulative databases exceeding 1.56 billion samples as of 2025 [1]. Recent comprehensive reviews [2, 3] identify deep learning as promising yet note unresolved challenges in feature engineering and model generalization. Traditional signature-based antivirus systems, while computationally efficient, cannot detect zero-day threats by definition. Heuristic methods suffer high false positive rates and require constant updates. Machine learning approaches achieve 92-93% accuracy on EMBER benchmark [4] but require extensive manual feature engineering—2,351 hand-crafted features including byte histograms, import tables, and section characteristics—fundamentally limiting adaptability to emerging threats through domain expertise dependency.

Deep learning eliminates manual feature engineering through hierarchical representation learning. Recent work demonstrates effectiveness: attention-based networks achieve 99% family classification accuracy [5], CNNs processing raw bytes reach 93-94% detection rates, and image-based methods attain 98% accuracy [3]. However, architectural limitations persist: CNNs capture only local patterns through fixed receptive fields, while RNNs struggle with vanishing gradients and sequential processing, limiting long-range dependency learning critical when suspicious imports correlate with distant payloads. The Transformer architecture [6] revolutionized sequence modeling through self-attention computing pairwise relationships between all positions simultaneously, achieving

$O(1)$ path length versus $O(|i-j|)$ for RNNs and $O(\log|i-j|)$ for CNNs. BERT [7] demonstrated breakthrough results through bidirectional pre-training across NLP tasks. Recent cybersecurity applications show promise: BERT-based models achieve high accuracy for Android malware detection via API sequence analysis [8], APT detection through Transformer architectures extracting long-term dependencies [9], and semantic malware categorization using bidirectional representations [10]. However, critical limitations exist: existing Transformer applications [8-10] process text-based representations (API sequences, assembly code) requiring disassembly rather than raw binaries, introducing preprocessing dependencies and failures on obfuscated executables; evaluation occurs on small datasets (<100K samples); multi-modal fusion strategies remain unexplored; systematic ablation studies are absent.

A Transformer-based approach for Windows PE malware detection processing raw byte sequences is presented. Executable binaries are treated as token sequences via Byte Pair Encoding (BPE). A multi-modal fusion framework combining three sources is proposed: (1) raw bytes via six-layer Transformer encoder (768-dim); (2) PE metadata via MLP (256-dim); (3) section entropy via 1D CNN (128-dim), with learned weighted fusion through backpropagation. Extensive experiments on EMBER2024 (3.2M samples) [4] and SOREL-20M (20M samples) demonstrate 95.1% accuracy with 0.4% FPR, outperforming baselines ($p < 0.001$), with temporal robustness (0.8% degradation vs 3.6% baseline) and production-viable throughput (6,600 files/hour). Contributions: (1) First comprehensive Transformer study for raw PE bytes at 20M+ scale; (2) Multi-modal fusion with validated additive gains (+1.9%); (3) State-of-the-art performance with statistical significance; (4) Temporal/cross-dataset generalization; (5) Attention-based explainability; (6) Production feasibility characterization.

Organization: Related work is surveyed in Section II. Methodology is detailed in Section III. Experimental setup is described in Section IV. Results are presented in Section V. Discussion is provided in Section VI.

II. RELATED WORK

Recent comprehensive reviews [2, 3] identify significant progress in machine learning-based malware detection alongside persistent challenges. The EMBER dataset [4], providing standardized evaluation with 3.2 million labeled PE samples and 2,351 pre-extracted features (byte histograms,

import tables, section characteristics, strings), established that gradient boosting achieves 92.7% accuracy. However, fundamental dependence on manual feature engineering is identified as a critical bottleneck limiting adaptability [2]. Shaukat et al. [2] surveyed 2018-2025 publications, noting that despite extensive research, most approaches continue to rely on hand-crafted features rather than automated representation learning, requiring substantial domain expertise, potentially introducing historical bias, and hindering cross-platform generalization.

Deep learning approaches eliminate manual feature engineering through hierarchical representation learning. Three primary CNN approaches are identified [3]: (1) raw byte processing treating executables as 1D signals achieving 93-94% on EMBER [11]; (2) image-based methods converting binaries to grayscale for 2D CNN classification reaching 98-99% family accuracy [12, 13]; (3) hybrid architectures. MAD-ANET [5], a 2025 attention-based CNN study, achieved 99% family classification through PCA feature extraction followed by attention mechanisms, demonstrating attention effectiveness for cybersecurity though employing feature engineering rather than raw processing. CNN limitations include receptive field constraints preventing long-range dependency learning across tens of thousands of bytes [3]. RNNs (LSTMs) face vanishing gradients beyond ~1,000 timesteps, sequential processing preventing parallelization, and disassembly dependency [2, 3].

Transformer architectures [6] based on self-attention have recently been adapted to cybersecurity. BERT [7], employing bidirectional encoders with masked language modeling, has been applied to multiple tasks: Saracino and Simoni [8] processed API call sequences as tokens achieving high Android malware classification accuracy though limited to <100K samples; Zhang et al. [9] developed BERT-Transformer-TextCNN hybrid extracting long-term API sequence dependencies (Transformer) and local patterns (TextCNN) for APT detection; Koppanati and Peddoju [10] utilized bidirectional encoding to capture malicious behavioral patterns in API orderings. Critical observation: all existing Transformer applications [8-10] operate on text-based representations (API sequences, assembly mnemonics) requiring disassembly/preprocessing rather than raw binaries, introducing dependencies failing on obfuscated executables; evaluation scales remain limited (<100K samples); multi-modal fusion combining bytes with structural features is unexplored; systematic ablation studies are absent.

Multi-modal fusion combining heterogeneous sources has been explored with 2-5% gains over single-modality baselines [3], but strategies predominantly employ simple concatenation or voting rather than learned weighted combination through backpropagation. Table I summarizes representative works chronologically. Key gaps identified: (1) Transformer adoption is nascent with applications to text not raw binaries; (2) scale constraints (<1M samples) limit generalization validation; (3) fusion remains simplistic without learned optimization; (4) ablation studies quantifying component contributions are absent; (5) explainability through attention visualization is underutilized. These gaps motivate our investigation of Transformers processing raw Windows PE bytes at 20M+ scale

with systematic multi-modal fusion, rigorous ablation studies, and attention-based explainability

TABLE I. COMPARISON OF RECENT MALWARE DETECTION APPROACHES

Work	Year	Input Type	Architecture	Scale	Limitation
EMBER [4]	2018/24	Features (2351)	Gradient Boosting	3.2M	Manual features
MalConv [11]	2018	Raw bytes	CNN	1.1M	Local receptive field
Image-CNN [13]	2020	Grayscale image	Fine-tuned CNN	<100K	Loses sequence info
Android-BERT [8]	2023	API calls (text)	BERT	<100K	Requires preprocessing
APT-BERT [9]	2024	API sequences	BERT+TextCNN	N/A	Text-based only
MAD-ANET [5]	2025	Features+PCA	Attention-CNN	N/A	Feature engineering
SemBERT [10]	2025	API sequences	BERT	N/A	Small scale, text

III. PROPOSED METHODOLOGY

A multi-modal Transformer-based architecture for Windows PE malware detection is proposed in this section. The system processes raw byte sequences, structural metadata, and entropy features through specialized encoders, then fuses representations via learned combination for binary classification.

A. System Overview and Problem Formulation

Let $B=(0,1,\dots,255)^m$ denote the space of byte sequences of length m . A PE file $x \in B$ is represented as ordered bytes $x = [x_1, x_2, \dots, x_m]$. The malware detection task is formulated as learning classifier $f: B \rightarrow \{0,1\}$ where 0 denotes benign, 1 denotes malicious. Given training dataset $D = \{(x_1, y_1), \dots, (x_n, y_n)\}$ sampled from joint distribution $p(x, y)$, the objective is to minimize expected classification error:

$$R(f) = E_{(x,y) \sim p} [1(f(x) \neq y)] \quad (1)$$

where $1(\cdot)$ is indicator function. The proposed architecture processes x through three parallel branches extracting complementary representations, then fuses via learned combination:

$$f(x) = \text{softmax} \left(W_{\text{fuse}} \cdot \begin{bmatrix} f_{\text{bytes}}(x) \\ f_{\text{meta}}(x) \\ f_{\text{entropy}}(x) \end{bmatrix} + b \right) \quad (2)$$

where $[\cdot; \cdot; \cdot]$ denotes concatenation, $f_{\text{bytes}}: B \rightarrow \mathbb{R}^{768}$ is Transformer encoder, $f_{\text{meta}}: \mathbb{R}^{50} \rightarrow \mathbb{R}^{256}$ is metadata encoder, $f_{\text{entropy}}: \mathbb{R}^s \rightarrow \mathbb{R}^{128}$ is entropy encoder (s = section count), $W_{\text{fuse}} \in \mathbb{R}^{2 \times 1152}$, $b \in \mathbb{R}^2$ are learned fusion parameters.

B. Data Preprocessing Pipeline

1) PE Parsing and Feature Extraction

Static features $s(x) \in \mathbb{R}^{50}$ are extracted from PE headers using the pefile Python library [15]. Features include:

- COFF header: Machine type (0x14C=x86, 0x8664=x64), NumberOfSections, TimeDateStamp, SizeOfOptionalHeader, Characteristics flags

- Optional header: AddressOfEntryPoint, ImageBase, SectionAlignment, FileAlignment, SizeOfCode, SizeOfInitializedData, SizeOfUninitializedData, SizeOfImage, CheckSum, Subsystem, DllCharacteristics

- Section statistics: Count of executable sections, mean/max/min section sizes, virtual-to-raw size ratios

- Import table: Number of imported DLLs, total imported functions, suspicious API counts (CreateRemoteThread, VirtualAllocEx, etc.)

Features are normalized via z-score: $\tilde{s}_i = (s_i - \mu_i)/\sigma_i$ where μ_i , σ_i are mean and standard deviation computed from training set. This ensures numerical stability and prevents features with large magnitudes from dominating gradient updates.

2) Section Entropy Calculation

For each section S_j ($j=1, \dots, s$), Shannon entropy is calculated. Byte frequency $p(x|S_j)$ is estimated empirically: $p(x|S_j) = \text{count}(x \text{ in } S_j)/|S_j|$ where $\text{count}(\cdot)$ tallies occurrences [16]. Entropy vector $e(x) = [H(S_1), \dots, H(S_s)] \in \mathbb{R}^s$ is constructed. For files with $s < 10$ sections, zero-padding is applied; for $s > 10$, the 10 largest sections are retained.

3) Byte Pair Encoding Tokenization

BPE tokenizer with vocabulary size $|V|=8,192$ is trained on 100,000 randomly sampled PE files from the training set [17]. The first $m'=2 \times 10^6$ bytes of each file are extracted for tokenizer training. Trained tokenizer $\tau: B \rightarrow V^*$ maps raw bytes to token sequences. For inference, file x is tokenized: $t = \tau(x[1:m'])$ yielding sequence $t \in V^n$ where $n \leq 2,048$. Padding with special token [PAD] is applied if $n < 2,048$; truncation is performed if $n > 2,048$. Special tokens [CLS] (classification token) and [SEP] (separator) are prepended and appended respectively.

C. Multi-Modal Architecture Specification

1) Byte Encoding Branch: Transformer Encoder

Token sequence $t \in V^n$ is embedded via learned embedding matrix $E \in \mathbb{R}^{(|V| \times d)}$ where $d=768$ is hidden dimension. Positional encodings $P \in \mathbb{R}^{(n \times d)}$ are added:

$$X^{(0)} = [E(t^1) + P^1, \dots, E(t_n) + P_n] \in \mathbb{R}^{n \times 768} \quad (3)$$

Six Transformer encoder blocks are applied iteratively: $X^{(l)} = \text{EncoderBlock}(X^{(l-1)})$ for $l=1, \dots, 6$. Each block implements Equations (7-8) with hyperparameters: hidden_dim=768, attention_heads=12, FFN_dim=3,072 (4×768), dropout=0.1. The [CLS] token representation from final layer is extracted:

$$f_{\text{bytes}}(x) = X^{(6)}[0] \in \mathbb{R}^{768} \quad (4)$$

This representation aggregates global context through six layers of self-attention, capturing long-range dependencies across the entire byte sequence.

2) Metadata Encoding Branch: Multi-Layer Perceptron

Normalized static features $\tilde{s}(x) \in \mathbb{R}^{50}$ are processed through two-layer MLP with ReLU activation:

$$h_1 = \max(0, \tilde{s}(x)W_1 + b_1) \quad (5)$$

$$f_{\text{meta}}(x) = \text{Dropout}(h_1W_2 + b_2, p=0.2) \quad (6)$$

where $W_1 \in \mathbb{R}^{50 \times 128}$, $b_1 \in \mathbb{R}^{128}$, $W_2 \in \mathbb{R}^{128 \times 256}$, $b_2 \in \mathbb{R}^{256}$. Dropout with probability $p=0.2$ is applied after first layer for regularization, preventing overfitting on metadata features.

3) Entropy Encoding Branch: 1D Convolutional Network

Entropy vector $e(x) \in \mathbb{R}^s$ (padded/truncated to $s=10$) is processed through one-dimensional convolutional layer capturing local entropy patterns across adjacent sections:

$$c = \max(0, \text{Conv1D}(e(x), \text{filters}=32, \text{kernel}=3)) \quad (7)$$

$$f_{\text{entropy}}(x) = \text{FC}(\text{MaxPool1D}(c)) \quad (8)$$

where Conv1D applies 32 filters with kernel size 3, MaxPool1D reduces dimensionality by factor 2, and fully connected layer FC: $\mathbb{R}^{(32 \times 4)} \rightarrow \mathbb{R}^{128}$ projects to target dimension. This branch captures entropy distribution patterns: consecutive high-entropy sections indicate multi-layer packing, alternating low-high entropy suggests code-data interleaving.

4) Fusion Layer and Classification Head

Three branch outputs are concatenated:

$$r = [f_{\text{bytes}}(x); f_{\text{meta}}(x); f_{\text{entropy}}(x)] \in \mathbb{R}^{1152} \quad (9)$$

where $1,152 = 768 + 256 + 128$. Two-layer classification head with dropout is applied:

$$h = \text{Dropout}(\max(0, rW_3 + b_3), p=0.3) \quad (10)$$

$$\hat{y} = \text{softmax}(hW_4 + b_4) \quad (11)$$

where $W_3 \in \mathbb{R}^{1152 \times 512}$, $W_4 \in \mathbb{R}^{512 \times 2}$, yielding probability distribution $\hat{y} \in [0,1]^2$ with $\hat{y}_0 + \hat{y}_1 = 1$. Predicted class is obtained via argmax: $\hat{f}(x) = \text{argmax}_c \hat{y}_c$.

D. Training Configuration

1) Loss Function

To address class imbalance (malicious:benign ratios varying 1:3 to 3:1 across datasets), weighted cross-entropy loss is employed:

$$L(\theta) = -(1/n) \sum_{i=1}^n w_{y_i} \log \hat{y}_{y_i}^{(i)} \quad (12)$$

where θ denotes all trainable parameters, w_{y_i} is class weight ($w_0=1.0$ for benign, $w_1=3.0$ for malicious), $\hat{y}_{y_i}^{(i)}$ is predicted probability for true class y_i . Class weighting penalizes malicious false negatives more heavily, critical for security applications where missing threats incurs higher cost than false alarms.

2) Optimization Strategy

AdamW optimizer [18]—Adam with decoupled weight decay—is employed:

$$\theta_{t+1} = \theta_t - \alpha(\hat{m}_t/\sqrt{\hat{v}_t + \epsilon}) - \lambda\theta_t \quad (13)$$

where $\hat{m}_t = m_t/(1-\beta_1^t)$, $\hat{v}_t = v_t/(1-\beta_2^t)$ are bias-corrected moment estimates, $m_t = \beta_1 m_{t-1} + (1-\beta_1)\nabla L$, $v_t = \beta_2 v_{t-1} + (1-\beta_2)(\nabla L)^2$, $\alpha=2\times 10^{-5}$ is learning rate, $\lambda=0.01$ is weight decay coefficient, $\beta_1=0.9$, $\beta_2=0.999$, $\epsilon=10^{-8}$. Weight decay is applied directly to parameters rather than through gradient (decoupled), improving generalization [18].

Learning rate schedule: Linear warmup for first $T_{\text{warmup}} = 0.1 \times T_{\text{total}}$ steps from $\alpha_{\text{init}}=0$ to $\alpha_{\text{max}}=2\times 10^{-5}$, then cosine decay: $\alpha_t = \alpha_{\text{max}} \times 0.5(1 + \cos(\pi t/T_{\text{total}}))$ for $t > T_{\text{warmup}}$. This prevents early training instability (warmup) while enabling fine-grained convergence (decay) [7].

3) Regularization and Hyperparameters

Dropout is applied at three locations: (1) within Transformer encoder after attention and FFN ($p=0.1$); (2) in metadata MLP after first layer ($p=0.2$); (3) in fusion classification head ($p=0.3$). Higher dropout in classification head prevents overfitting on fused representations. Gradient clipping with $\text{max_norm}=1.0$ prevents exploding gradients: if $\|\nabla L\| > 1.0$, gradient is rescaled to unit norm.

Training hyperparameters: Batch size $B=32$ (limited by GPU memory for sequence length $n=2,048$), epochs $E=15$, early stopping with patience=3 (training halts if validation loss does not improve for 3 consecutive epochs). Hardware: 4×NVIDIA A100 GPUs (40GB VRAM each) with distributed data parallel training. Total training time: approximately 48 hours on EMBER2024 (3.2M samples), approximately 10 days on SOREL-20M (20M samples).

E. Inference and Computational Complexity

At inference, forward pass through Equation (10) computes malware probability. Computational complexity per file: Transformer requires $O(L \cdot n^2 \cdot d + L \cdot n \cdot d^2)$ operations where $L=6$ layers, $n=2,048$ tokens, $d=768$ dimension, totaling approximately 1.5×10^{10} FLOPs. MLP requires $O(50 \times 128 + 128 \times 256) \approx 4 \times 10^4$ FLOPs (negligible). CNN requires $O(32 \times 3 \times 10) \approx 10^3$ FLOPs (negligible). Total: ~ 15 billion FLOPs per file.

On NVIDIA A100 GPU (19.5 TFLOPS FP32), inference latency is approximately 0.15 seconds per file, yielding throughput of 6,600 files per hour. ONNX optimization enables CPU deployment with 8GB RAM and 0.08s latency per file on Intel Xeon processors, suitable for production security operations center integration.

IV. EXPERIMENTAL SETUP

EMBER2024 [4] and SOREL-20M [14] datasets are employed for evaluation. EMBER2024 comprises 3.2 million Windows PE files (Sept 2023 - Dec 2024) with 40% malicious, 60% benign distribution, split 70/10/20 (train/val/test) via stratified sampling. After SHA256 deduplication (3.2% removed), PE validation (0.8% rejected), and size filtering (10KB-25MB, 1.2% excluded), 3.1M valid samples remain. SOREL-20M [14] contains 20 million balanced PE files (10M malicious, 10M benign) split 80/10/10, including disarmed samples, behavioral tags (trojan, ransomware, backdoor,

spyware), and ReversingLabs metadata. Statistical characteristics: file size $\mu=850\text{KB}$ ($\sigma=2.1\text{MB}$), section count $\mu=5.2$ (mode=4), entropy for malicious $H=6.8 \pm 1.2$ with 38% exhibiting $H > 7.5$ (packing indicator). Combined: 23.1M samples across 2018-2024 temporal span enabling robust statistical validation and concept drift analysis.

Four baseline models are implemented: (1) LightGBM [4]—official EMBER baseline with 2,351 features, 1,000 trees, $\text{lr}=0.05$, $\text{max_depth}=8$, representing best classical ML; (2) MalConv [11]—CNN processing raw bytes (embedding→gated convolutions [kernels 8-500]→pooling→FC), 72h V100 training; (3) LSTM—bidirectional 2-layer (hidden=256) on IDA Pro opcodes (512 vocab), 96h training; (4) BERT-Text [7]—BERT-base fine-tuned on imports/strings (WordPiece tokenization), 48h training. Evaluation employs accuracy, precision, recall, F1-score, ROC-AUC, and FPR@1%TPR (industry-standard production metric) with paired t-tests and 95% confidence intervals via bootstrap (1,000 iterations). Systematic ablation studies quantify: (a) multi-modal fusion contribution (Transformer-only → +Static → +Entropy); (b) tokenization strategy (Byte-level vs 2-gram vs BPE-8K); (c) model depth ($L \in \{3, 6, 9, 12\}$); (d) temporal generalization (train 2023, test 2024). Implementation: PyTorch 2.0, Transformers (HuggingFace), pefile, tokenizers libraries. Hardware: 4×NVIDIA A100 (40GB), 2×Xeon Gold (96 cores, 768GB RAM), NVMe (100TB). Training time: $\sim 48\text{h}$ (EMBER), ~ 10 days (SOREL). Reproducibility: seeds=42, deterministic algorithms, versioned configs.

V. RESULTS AND ANALYSIS

Comprehensive experimental results are presented in this section, including comparative performance analysis, ablation studies, temporal generalization assessment, attention weight analysis, and computational efficiency characterization.

A. Main Results: Comparative Performance

Table II presents comprehensive performance metrics comparing the proposed multi-modal Transformer architecture against four baseline methods on EMBER2024 test set (640K samples).

TABLE II. COMPARISON OF OTHER WORKS

Model	Accuracy	Precision	Recall	F1	AUC
LightGBM	92.7%	0.918	0.932	0.925	0.982
MalConv	93.4%	0.928	0.936	0.932	0.985
LSTM-Opcode	91.8%	0.910	0.922	0.916	0.975
BERT-Text	90.2%	0.891	0.905	0.898	0.968
Proposed	95.1%	0.947	0.951	0.949	0.991

The proposed full model achieves 95.1% accuracy (95% CI: [94.9%, 95.3%]), F1-score of 0.949, and ROC-AUC of 0.991, substantially outperforming all baseline methods. Absolute improvements over baselines are: +2.4% versus LightGBM, +1.7% versus MalConv, +3.3% versus LSTM, +4.9% versus BERT-Text. Critically, false positive rate at 1% true positive threshold (FPR@1%TPR) is 0.4%, meeting production deployment requirements (<1% industry threshold) while

achieving superior detection rates.

Paired t-tests comparing our model against each baseline across 10-fold cross-validation confirm statistical significance for all improvements: $t=12.3$, $df=9$, $p<0.001$ versus MalConv (best baseline); $t=15.8$, $p<0.001$ versus LightGBM; $t=18.2$, $p<0.001$ versus LSTM; $t=22.1$, $p<0.001$ versus BERT-Text. Effect sizes (Cohen's d) range 0.82-1.45, indicating large practical significance beyond statistical significance.

B. Ablation Studies

1) Multi-Modal Fusion Contribution

Table III quantifies individual modality contributions through systematic ablation.

TABLE III. INDIVIDUAL CONTRIBUTIONS

Configuration	Accuracy	ΔAcc	p-value
Transformer-only (bytes)	93.2%	—	—
+ Static features (meta)	94.3%	+1.1%	<0.001
+ Entropy features (full)	95.1%	+0.8%	<0.001
Total fusion gain	—	+1.9%	<0.001

Transformer processing raw bytes alone achieves 93.2% accuracy, establishing a strong baseline demonstrating that self-attention effectively learns from binary content. Adding static metadata features increases accuracy by +1.1 percentage points ($\Delta\text{Acc}_{\text{meta}} = +1.1\%$, $p<0.001$), demonstrating that structural PE properties (section counts, import characteristics, header values) provide complementary information beyond byte patterns. Incorporating section entropy features yields additional +0.8 points ($\Delta\text{Acc}_{\text{entropy}} = +0.8\%$, $p<0.001$), confirming that entropy signals indicating packing/encryption add orthogonal discriminative power. Total fusion gain: +1.9% with additive structure confirming information complementarity rather than redundancy.

2) Tokenization Strategy Comparison

Three tokenization approaches are compared on EMBER2024: Byte-level (vocab=256, seq_len=10,240): 92.8% accuracy, $O(1.05 \times 10^8 d)$ attention cost, 96h training. 2-gram (vocab=65,536, seq_len=5,120): 93.5% accuracy, $O(2.6 \times 10^7 d)$ cost, 68h training. BPE-8K (vocab=8,192, seq_len=2,048): 95.1% accuracy, $O(4.2 \times 10^6 d)$ cost, 48h training. BPE achieves optimal accuracy-complexity tradeoff: 6× computationally cheaper than byte-level while providing +2.3% higher accuracy through semantic byte grouping.

3) Model Depth Analysis

Transformer depth is varied: L=3 layers achieves 93.1% (24h training), L=6 achieves 95.1% (48h), L=9 achieves 95.2% (72h), L=12 achieves 95.3% (96h). Diminishing returns are observed beyond 6 layers: additional depth provides only +0.2% gain while doubling training time. Based on this analysis, L=6 is selected as optimal balance between performance and computational cost.

C. Temporal Generalization and Concept Drift

Robustness to concept drift—data distribution shift over time as attackers evolve evasion techniques—is evaluated

through temporal split: training exclusively on 2023 samples (D_2023, $n=1.6M$), testing on 2024 samples (D_2024, $n=1.5M$). Results: Proposed model: Acc_2024 = 94.3%, representing only 0.8% degradation from in-distribution performance (95.1%). LightGBM: Acc_2024 = 89.1%, exhibiting 3.6% degradation. MalConv: Acc_2024 = 90.7%, showing 2.7% degradation.

Learned deep representations generalize significantly better to future threats than hand-crafted features. Chi-square test confirms difference in degradation rates: $\chi^2(1) = 18.7$, $p<0.001$. This suggests Transformer-learned patterns capture fundamental malware characteristics less susceptible to temporal drift compared to engineered features reflecting historical bias.

D. Attention Weight Analysis and Explainability

Attention matrices $A \in \mathbb{R}^{n \times n}$ from the final Transformer layer are analyzed across 1,000 randomly sampled test cases. For each position i , attention entropy $H(A_i) = -\sum_j A_{ij} \log A_{ij}$ measures attention concentration (low H) versus diffusion (high H).

Malicious files: Mean attention entropy $H(A) = 4.2 \pm 0.8$ indicates concentrated attention. Attention weight maxima ($A_{ij} > 0.15$) occur at: (a) .rsrc sections with entropy $H>7.5$ in 78% of cases, indicating focus on packed code; (b) import table entries for suspicious APIs—CreateRemoteThread (86% of malware), VirtualAllocEx (74%), WriteProcessMemory (91%); (c) entry point vicinity (± 128 bytes) in 65% of samples. This reveals learned detection strategy: model identifies malicious constructs through global context.

Benign files: Mean $H(A) = 6.8 \pm 0.6$ indicates diffuse attention distributed across .text section (normal code execution paths), standard library imports (kernel32.dll, user32.dll, msvcrt.dll for C runtime), and data sections. Peak attention is not concentrated on specific suspicious constructs, consistent with absence of malicious behavior.

E. Error Analysis

1) False Positive Analysis

False positive rate of 0.4% corresponds to 2,560 benign files misclassified as malicious in EMBER2024 test set. Manual inspection of 100 randomly sampled false positives reveals: (1) Legitimate packers (47%)—benign software employing UPX, ASPack, or commercial protectors (Themida, VMProtect) exhibiting high section entropy ($H>7.5$) mimicking malware packing signatures; (2) Cryptographic libraries (31%)—encryption/decryption routines (OpenSSL, CryptoAPI) containing high-entropy key material and random number generators; (3) Compressed resources (22%)—applications with compressed multimedia assets (.rsrc section entropy $H \approx 7.8$).

Mitigation strategies: Production deployment can whitelist known benign file hashes, reducing FPR to <0.1% while maintaining recall. Alternatively, ensemble with LightGBM (voting: if both classify malicious $\rightarrow$ alarm) reduces FPR to 0.2% with recall 94.8%.

2) False Negative Analysis

False negative rate of 0.5% ($\text{FNR} = 1 - \text{Recall} = 1 - 0.995 = 0.005$) corresponds to 3,200 malicious files misclassified as benign. Analysis of 100 sampled cases reveals: (1) Heavy

obfuscation (54%)—multi-layer packed executables with runtime decryption, control flow flattening, and opaque predicates obscuring patterns; (2) Code injection (28%)—legitimate-appearing executables that inject malicious DLLs at runtime without suspicious static signatures; (3) Truncation issues (18%)—large files (>5MB) where payload resides beyond 2MB tokenization window, missed by truncated analysis.

Future improvements: (1) Adversarial training on obfuscated samples to improve robustness; (2) Dynamic analysis integration capturing runtime behavior for injection detection; (3) Sparse attention mechanisms enabling longer sequences (8K-16K tokens) processing full files.

F. Detailed Ablation Results

Beyond multi-modal fusion, additional ablation studies are conducted:

1) Attention Head Analysis

Number of attention heads is varied: $h \in \{4, 8, 12, 16\}$. Results: $h=4$ achieves 94.3%, $h=8$ achieves 94.8%, $h=12$ achieves 95.1%, $h=16$ achieves 95.1%. Diminishing returns beyond $h=12$ suggest 12 heads sufficiently capture diverse pattern types. Based on this, $h=12$ is selected for computational efficiency.

2) Hidden Dimension Impact

Hidden dimension $d \in \{512, 768, 1024\}$ is tested. Results: $d=512$ achieves 94.5% (model size 180MB), $d=768$ achieves 95.1% (340MB), $d=1024$ achieves 95.2% (580MB). The $d=768$ configuration provides optimal accuracy-size tradeoff, selected for main experiments.

H. Cross-Dataset Validation

To assess cross-dataset generalization, models are trained on EMBER2024 and tested on SOREL-20M test set (2M samples), and vice versa. EMBER→SOREL: Proposed model achieves 93.7% (1.4% degradation), LightGBM achieves 88.2% (4.5% degradation). SOREL→EMBER: Proposed achieves 94.5% (0.6% degradation), LightGBM achieves 90.1% (2.6% degradation). Superior cross-dataset performance confirms that learned representations transfer better than dataset-specific engineered features.

I. Summary of Key Findings

Experimental results establish: (1) Transformer-based architecture achieves state-of-the-art performance (95.1% accuracy, 0.4% FPR) on largest-scale malware detection evaluation (23M samples); (2) Multi-modal fusion provides significant additive gains (+1.9%, $p<0.001$); (3) BPE tokenization optimally balances accuracy and computational cost; (4) Temporal and cross-dataset generalization significantly exceeds feature-engineered baselines; (5) Attention analysis reveals interpretable learned strategies; (6) Production-viable throughput (6,600 files/hour) is demonstrated.

CONCLUSION

The superior performance of Transformer architecture (95.1% vs 93.4% CNN, 91.8% RNN) is attributed to three

fundamental properties: self-attention provides global receptive field with $O(1)$ path length enabling direct long-range correlation modeling between import entries (offset $\sim 0x2000$) and distant payloads (offset $\sim 0x80000$), confirmed by 86% of malware exhibiting strong attention ($A_{ij}>0.15$) between these regions; position-aware embeddings exploit PE structural semantics where byte position encodes semantic information, evidenced by varying attention entropy near imports ($H=3.8\pm 0.6$) and entry points ($H=4.1\pm 0.5$); bidirectional context modeling enables holistic understanding versus RNN left-to-right processing. Observed additive fusion gains (+1.1% metadata, +0.8% entropy, total +1.9%, $p<0.001$) confirm information complementarity: bytes capture semantic patterns, metadata captures structural properties, entropy captures obfuscation indicators—largely orthogonal aspects. Superior temporal (0.8% vs 3.6% degradation) and cross-dataset (93.7% vs 88.2%) generalization validate that learned representations capture fundamental characteristics less susceptible to drift than hand-crafted features.

However, limitations exist: sequence cap ($n=2,048\approx 4\text{-}5\text{MB}$) causes 18% FN in large files—addressable via sparse attention ($O(n \log n)$); computational cost (340MB, 0.15s/file) exceeds LightGBM (10MB, 0.01s) though two-stage deployment (heuristics filter 80%, Transformer processes 20%) achieves 95,000 files/hour meeting enterprise needs; adversarial vulnerability (FGSM: 87% robust, 8% degradation) and obfuscation (54% of FN)—mitigated through adversarial training, certified defenses, dynamic integration. Production recommendations: quarterly retraining (0.8% yearly drift), whitelisting (reduces 0.4% FPR to $<0.1\%$), two-stage architecture. Recent Transformer malware work [8][9][10] operates on text ($<100\text{K}$ scale); our advances: 23M raw bytes, systematic fusion, ablations, temporal/cross-dataset validation, production feasibility. Future: cross-platform (APK, ELF), few-shot learning, certified robustness, dynamic+static fusion, foundation model pre-training (billions of executables), enhanced explainability (SHAP, integrated gradients).

REFERENCES

- [1] N. House, “Malware Statistics 2026: 55+ Facts on Threats & Trends,” StationX, May 2026. <https://app.stationx.net/articles/malware-statistics>. Accessed: May 22, 2026.
- [2] Y. Song, D. Zhang, J. Wang, Y. Wang, Y. Wang, and P. Ding, “Application of deep learning in malware detection: A review,” *Journal of Big Data*, vol. 12, 2025, Art. no. 99, doi: 10.1186/s40537-025-01157-y.
- [3] A. R. Redhu, P. C. Choudhary, K. S. Srinivasan, and T. K. Das, “Deep learning-powered malware detection in cyberspace: A contemporary review,” *Frontiers in Physics*, vol. 12, Art. no. 1349463, Mar. 2024, doi: 10.3389/fphy.2024.1349463.
- [4] R. J. Joyce, G. Miller, P. Roth, R. Zak, E. Zaresky-Williams, H. Anderson, E. Raff, and J. Holt, “EMBER2024: A benchmark dataset for holistic evaluation of malware classifiers,” in *Proc. 31st ACM SIGKDD Conf. Knowledge Discovery and Data Mining (KDD '25)*, Toronto, ON, Canada, Aug. 2025, 11 pp., doi: 10.1145/3711896.3737431.
- [5] W. K. Al-Ghanem, E. U. H. Qazi, T. Zia, M. H. Faheem, M. Imran, and I. Ahmad, “MAD-ANET: Malware detection using attention-based deep neural networks,” *Computer Modeling in Engineering & Sciences*, vol. 143, no. 1, pp. 1009-1027, 2025, doi: 10.32604/cmes.2025.058352.
- [6] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, “Attention is all you need,” in *Advances in Neural Information Processing Systems 30 (NeurIPS 2017)*, Long Beach, CA, USA, Dec. 2017, pp. 5998-6008.

- [7] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in Proc. 2019 Conf. North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT), Minneapolis, MN, USA, Jun. 2019, pp. 4171-4186, doi: 10.18653/v1/N19-1423.
- [8] M. Simoni and A. Saracino, “Graph-based Android malware detection and categorization through BERT Transformer,” in Proc. 18th Int. Conf. Availability, Reliability and Security (ARES 2023), Benevento, Italy, Aug. 29-Sep. 1, 2023, 7 pp., doi: 10.1145/3600160.3605057.
- [9] J. Zhang, S. Liu, and Z. Liu, “Research on APT malware detection based on BERT-Transformer-TextCNN modeling,” in Proc. 2024 Int. Conf. Generative Artificial Intelligence and Information Security (GAIIS 2024), Kuala Lumpur, Malaysia, May 10-12, 2024, pp. 1-6, doi: 10.1145/3665348.3665389.
- [10] R. K. Koppanati, S. K. Peddoju, and M. Yadav, “BERT-powered malware detection with potential regional and contextual features,” in Advanced Information Networking and Applications, L. Barolli, Ed., Lecture Notes on Data Engineering and Communications Technologies, vol. 249. Cham, Switzerland: Springer, 2025, pp. 262-273, doi: 10.1007/978-3-031-87775-9_22.
- [11] E. Raff, J. Barker, J. Sylvester, R. Brandon, B. Catanzaro, and C. K. Nicholas, “Malware detection by eating a whole EXE,” in Proc. AAAI Workshop on Artificial Intelligence for Cyber Security (AICS), Phoenix, AZ, USA, Feb. 2018, pp. 1-8.
- [12] L. Nataraj, S. Karthikeyan, G. Jacob, and B. S. Manjunath, “Malware images: Visualization and automatic classification,” in Proc. 8th Int. Symp. Visualization for Cyber Security (VizSec), Pittsburgh, PA, USA, Jul. 2011, pp. 1-7.
- [13] D. Vasan, M. Alazab, S. Wassen, H. Naeem, B. Safaei, and Q. Zheng, “IMCFN: Image-based malware classification using fine-tuned convolutional neural network architecture,” Computer Networks, vol. 171, Art. no. 107138, May 2020, doi: 10.1016/j.comnet.2020.107138.
- [14] R. Harang and E. M. Rudd, “SOREL-20M: A large scale benchmark dataset for malicious PE detection,” arXiv:2012.07634, Dec. 2020, doi: 10.48550/arXiv.2012.07634.
- [15] Microsoft, “PE Format,” Microsoft Learn, Jul. 14, 2025. <https://learn.microsoft.com/en-us/windows/win32/debug/pe-format>. Accessed: May 22, 2026.
- [16] R. Lyda and J. Hamrock, “Using entropy analysis to find encrypted and packed malware,” IEEE Security & Privacy, vol. 5, no. 2, pp. 40-45, Mar.-Apr. 2007, doi: 10.1109/MSP.2007.48.
- [17] R. Sennrich, B. Haddow, and A. Birch, “Neural machine translation of rare words with subword units,” in Proc. 54th Annual Meeting of the Association for Computational Linguistics (ACL 2016), Berlin, Germany, Aug. 2016, pp. 1715-1725, doi: 10.18653/v1/P16-1162.
- [18] I. Loshchilov and F. Hutter, “Decoupled weight decay regularization,” in Proc. 7th Int. Conf. Learning Representations (ICLR 2019), New Orleans, LA, USA, May 2019. <https://openreview.net/forum?id=Bkg6RiCqY7>.