

SecMTL-BERT: A Multi-Task Transformer Framework for Unified Information Systems Security Text Intelligence

Alaa A. Muslam^{1,*}, Ameer B. A. Alaasam^{2,3}, Athar H. Mohammed⁴ and Zena H. Khalil¹

¹Department of Cyber Security, College of Computer Science and Information Technology,
University of Al-Qadisiyah, Al-Diwaniyah, Iraq
e-mail: alaa.abidmuslam@qu.edu.iq, zena.khalil@qu.edu.iq

²Department of System Programming, South Ural State University, Chelyabinsk, Russia.

³Ministry of Education Iraq, General Directorate for Education in Al Qadisiyah, Al-Diwaniyah,
Iraq
e-mail: alaasamab@susu.ru

⁴Department of Mathematic, College of Education – University of Al-Qadisiyah – Diwaniyah,
Iraq
e-mail: edu.math-t@qu.edu.iq

*Corresponding Authors: Alaa A. Muslam. e-mail: alaa.abidmuslam@qu.edu.iq

Abstract

Modern organizations typically process enormous volumes of data, including text incident tickets, vulnerability bulletins, threat intelligence reports, and audit logs. Such processing is simply not feasible through manual review. One approach to handling these tasks is natural language processing (NLP). However, it handles them separately using separate models and multiple inferences. This is computationally expensive and does not exploit the mutual information shared between related security-text tasks. SecMTL-BERT is proposed in this work, which is a unified multi-task transformer framework. This framework simultaneously performs the classification of security incident, cyber-entity recognition (CyNER), and also threat urgency scoring. It ensures to perform them in single forward pass. The novelty of central architectural represented by Cross-Task Attention (CTA) module. This architectural enables the bidirectional feature to share between the classification and entity-recognition heads. This also allowing each task to utilize complementary representations which was learned by other. The initialization of shared encoder is done from BERT-base. After initialization, it subjected to domain adaptive fine tuning using 2.3 million document security. These documents are corpus drawn from the collection of National Vulnerability Database, CERT advisories, threat-intelligence reports, and anonymize of enterprise SOC tickets. Final experiments were established using three benchmark datasets, included SecIncident-9K, DNRTI, and ThreatUrgency-8K. The SecMTL-BERT achieves macro F1 scores of 92.6%, 89.5%, and 90.5% on the three tasks respectively. The performance of SecMTL-BERT outperforming the strongest single-task baseline (SecBERT) by 5.7, 6.4, and 6.5 percentage points. Ablation experiments confirm that, every architectural component effectively contributes to performance. Also, SecMTL-BERT reduces total inference time by 58.7% relative to running three separate SecBERT models. Thus, the real-time deployment in operational security become feasible. This benefit IS security managers, who are seeking to automate operations of alert triage, threat intelligence extraction, and advisory prioritization magnitude.

Keywords: Multi-Task Learning, Transformer, Cyber-Entity Recognition, Threat Intelligence, Incident Classification, Information Systems Security, BERT.

Received 4 April 2026; Accepted 28 June 2026

1 Introduction

There is a basic asymmetry in information systems security management: the number of text messages that can be processed in a security management system grows exponentially, but the number of analysts who can inspect the text grows at a linear rate [1]. A large enterprise SOC can receive hundreds of thousands of security events a day, with free-form descriptive text associated with each event. At the same time, hundreds of vulnerability advisories, threat intelligence reports, compliance notifications come in from external feeds, regulatory authorities, vendor channels, etc. This corpus categorizing incidents and identifying the entities involved and evaluating their urgency currently takes up disproportionate time of senior analysts, leaving less bandwidth for higher order threat investigation and strategic security management [2].

One technologically mature approach to automate this extraction process is through natural language processing (NLP). Transformer models such as BERT [3] and domain-specific variants have shown promising results on individual security-text tasks such as incident classification [4], named entity recognition for cyber-intelligence (CyNER) [5] and sentiment-based urgency assessment [6]. There are, however, three restrictions that renders existing approaches of operation limited. First, they process each task independently, meaning that there has to be an independent model for each, and each has to go through an independent inference pass – this is a multiplicative number of computational steps that is impractical for real-time SOC environments. Second, they fail to account for structural interdependence between tasks that is, the fact that the entity types found in a security text are highly predictive of the category of incident, which in turn then provides clues to assessing urgency. Third, general-purpose BERT models are not pre-trained on security-domain vocabulary, causing them to treat specialized identifiers such as CVE numbers, malware family names, and MITRE ATT&CK technique codes as out-of-vocabulary tokens.

This paper addresses all three limitations through a single unified framework. We make the following contributions:

- (1) We propose SecMTL-BERT, the first multi-task transformer model that jointly performs security incident classification, cyber-entity recognition, and threat urgency scoring in a single forward pass, eliminating redundant computation while exploiting inter-task dependencies.
- (2) We introduce the Cross-Task Attention (CTA) module, a novel architectural component that enables bidirectional feature sharing between the classification and NER heads at inference time, allowing each task to leverage complementary representations learned by the other without requiring task-specific inference passes.
- (3) We perform domain-adaptive fine-tuning of the shared BERT encoder on a 2.3-million-document security corpus and extend the WordPiece vocabulary with 1,847 security-specific tokens, significantly improving representation of specialized security terminology.
- (4) We perform extensive experiments on three existing benchmarks, showing statistically significant and consistent gains from all single-task and partial multi-task baselines, and an ablation study revealing the importance of each design choice.

The remainder of this paper is organized as follows. Section 2 reviews related work and positions our contribution against the existing literature. Section 3 provides the necessary background on transformer architectures and multi-task learning. Section 4 presents the SecMTL-BERT architecture in detail. Section 5 describes the experimental setup. Section 6 reports and analyzes results. Section 7 discusses implications for IS security management. Section 8 concludes.

2 Related Work

2.1 NLP for IS Security Text Analytics

The use of NLP with cybersecurity text goes back to before transformer-based models. Initial research concentrated on developing rule-based and statistical IOCs from threat intelligence reports [7] and phishing email classification based on keywords [8]. Word Embeddings were used to improve the ability to classify malware families using behavioral reports [9] and to identify vulnerability discussions in public repositories [10].

The transformer revolution significantly raised the performance ceiling. Liao et al. [21] demonstrated that NLP pipelines could automatically extract structured vulnerability attributes affected software, attack vector, required privileges from NVD CVE description prose, achieving precision scores above 80%. Sabetta and Bezzi [22] showed that NLP classifiers could distinguish security-relevant from non-security code commits with high accuracy, enabling automated security patch identification. Gao et al. [23] built a CyNER system targeting threat-actor names, malware identifiers, and IOCs in threat intelligence text. These works collectively establish that NLP can automate specific security information extraction tasks with useful accuracy. However, they address tasks individually and do not exploit relationships between tasks.

2.2 Domain-Adaptive Pre-Training for Security NLP

General-purpose language models face a systematic disadvantage on IS security text: their pre-training corpora Wikipedia, BookCorpus, web crawls contain very few examples of CVE identifiers, exploit payloads, or threat-actor aliases, causing these tokens to receive poor representations. SecBERT [27] addressed this by continuing BERT pre-training on a corpus of cybersecurity articles and NVD entries, demonstrating improved performance on security classification benchmarks. CyberBERT [28] similarly fine-tuned BERT on cyber-threat intelligence text, with particular gains on NER tasks. Ranade et al. [25] pre-trained a domain-specific transformer on CTI reports for joint NER and relation extraction, showing that domain adaptation and multi-task learning are mutually reinforcing. Our work extends this line by performing domain-adaptive pre-training on a substantially larger and more diverse corpus than prior work, combining NVD entries, enterprise SOC tickets, CERT advisories, and threat-intelligence reports.

2.3 Multi-Task Learning in NLP

Multi-task learning (MTL) in NLP leverages the insight that jointly training on related tasks can improve generalization by encouraging shared representations to capture task-invariant linguistic features [11]. MTL has been successfully applied to general NLP task combinations including named entity recognition with relation extraction [12], sentiment analysis with aspect classification [13], and question answering with reading comprehension [14]. In IS security, Ranade et al. [25] combined NER and relation extraction in a partial MTL framework; Aghaei et al. [26] built a joint extraction model for cyber-threat intelligence, but focused exclusively on entity and relation types without addressing incident classification or urgency scoring. No prior work combines all three operationally critical tasks classification, NER, and urgency scoring in a unified multi-task framework, nor do prior approaches include a mechanism for explicit cross-task feature sharing at inference time.

2.4 Large Language Models, Foundation Models, and Emerging Directions

The rapid emergence of large language models (LLMs) and instruction-tuned foundation models has opened new directions for cybersecurity NLP that merit explicit positioning relative to SecMTL-BERT. Ferrag et al. [39] introduced SecureFalcon, a cybersecurity-specific LLM fine-tuned from Falcon-7B on vulnerability and threat data, demonstrating that instruction tuning on domain-specific corpora substantially improves zero-shot and few-shot classification over general-purpose LLMs. Xu et al. [40] showed that GPT-4 and similar large-scale autoregressive models achieve competitive NER performance on CTI text through prompt engineering alone, without any task-specific fine-tuning. However, these approaches face practical deployment constraints in SOC environments due to their high inference latency and GPU memory requirements, making them unsuitable for the real-time alert enrichment use case addressed by SecMTL-BERT.

Retrieval-Augmented Generation (RAG) has been applied to cybersecurity question answering and advisory summarization [41], augmenting LLM generation with retrieved threat intelligence documents to reduce hallucination and improve factual grounding. Graph-based reasoning approaches, such as those of Piplai et al. [42], construct cyber-knowledge graphs from NER and relation extraction outputs and apply graph neural networks for threat reasoning and attribution, representing a complementary downstream use of NER outputs similar to those produced by SecMTL-BERT. Prompt-based learning methods, including in-context learning and chain-of-thought prompting applied to cybersecurity classification [43], achieve strong results in low-resource scenarios but require access to proprietary LLM APIs and cannot be easily deployed in air-gapped enterprise security environments. SecMTL-BERT occupies a distinct operational niche: it provides multi-task inference at encoder-model scale (128M parameters), with throughput and latency characteristics compatible with production SOC deployment, while achieving performance competitive with approaches that require orders-of-magnitude more computation.

2.5 Research Gap

Table 1 summarizes the positioning of SecMTL-BERT against representative related work along five dimensions. The gap is clear: no existing approach jointly addresses all three tasks central to IS security text analytics, incorporates both domain fine-tuning and explicit cross-task attention, and has been evaluated on IS security management datasets. SecMTL-BERT is designed to fill this gap.

Table 1: Positioning of SecMTL-BERT Relative to Representative Related Work

Study	IS Sec. Domain	Task(s)	Multi-Task	Domain FT	Cross-Task Attn.
Liao et al. [21]	Vulnerability Mgmt	NER only	X	X	X
Sabetta & Bezzi [22]	AppSec	Classification only	X	X	X
Gao et al. [23]	Threat Intel	NER only	X	Partial	X
Li et al. [24]	SOC Triage	Classification only	X	X	X
Ranade et al. [25]	Threat Intel	NER + Classification	Partial	✓	X

Aghaei et al. [26]	CTI	NER only	X	✓	X
SecBERT [27]	General security	Classification	X	✓	X
CyberBERT [28]	Threat Intel	Classification + NER	Partial	✓	X
SecMTL-BERT (ours)	IS Security Mgmt	Classif. + NER + Urgency	✓	✓	✓

3. Background and Preliminaries

3.1 Transformer Architecture and BERT

The transformer encoder [15] processes an input token sequence $T = (t_1, t_2, \dots, t_n)$ by applying L layers of multi-head self-attention and position-wise feed-forward networks. Each layer l computes contextualized representations $H^{(l)} = \text{TransformerBlock}(H^{(l-1)})$, where $H^{(0)}$ is the sum of token, position, and segment embeddings. BERT [3] pre-trains a 12-layer transformer encoder (BERT-base: 110M parameters) using masked language modeling (MLM) and next sentence prediction (NSP) objectives on large text corpora. Fine-tuning appends task-specific heads to the pre-trained encoder and trains end-to-end on task-specific labeled data. The [CLS] token representation at the final encoder layer provides a sentence-level summary suitable for sequence classification, while per-token representations support token-level tasks such as NER.

3.2 Multi-Task Learning Formulation

In the multi-task learning setting, a model is trained to minimize a joint loss over T tasks sharing a common input representation. Given tasks $\tau = \{\tau_1, \dots, \tau_T\}$

with task-specific loss functions $\{L_1, \dots, L_T\}$ and associated data $\{D_1, \dots, D_T\}$, the joint optimization objective is:

$$L_{\text{joint}} = \sum_k \lambda_k \cdot L_k(\theta_{\text{shared}}, \theta_k) \quad (1)$$

where θ_{shared} are the shared encoder parameters, θ_k are task-specific head parameters, and λ_k are task weighting coefficients. In SecMTL-BERT, $T = 3$ with tasks τ_1 (incident classification), τ_2 (CyNER), and τ_3 (urgency scoring). Coefficients $\lambda_1 = 0.40$, $\lambda_2 = 0.40$, $\lambda_3 = 0.20$ were selected by grid search on the combined validation set.

4. Proposed Framework: SecMTL-BERT

4.1 Architecture Overview

SecMTL-BERT processes a raw security text input and simultaneously produces three outputs: an incident category label, a sequence of typed entity labels, and a threat urgency score. The architecture consists of four components: (i) a domain-adapted shared BERT encoder, (ii) the Cross-Task Attention (CTA) module, (iii) three task-specific prediction heads, and (iv) a joint training objective. Figure 1 illustrates the architecture.

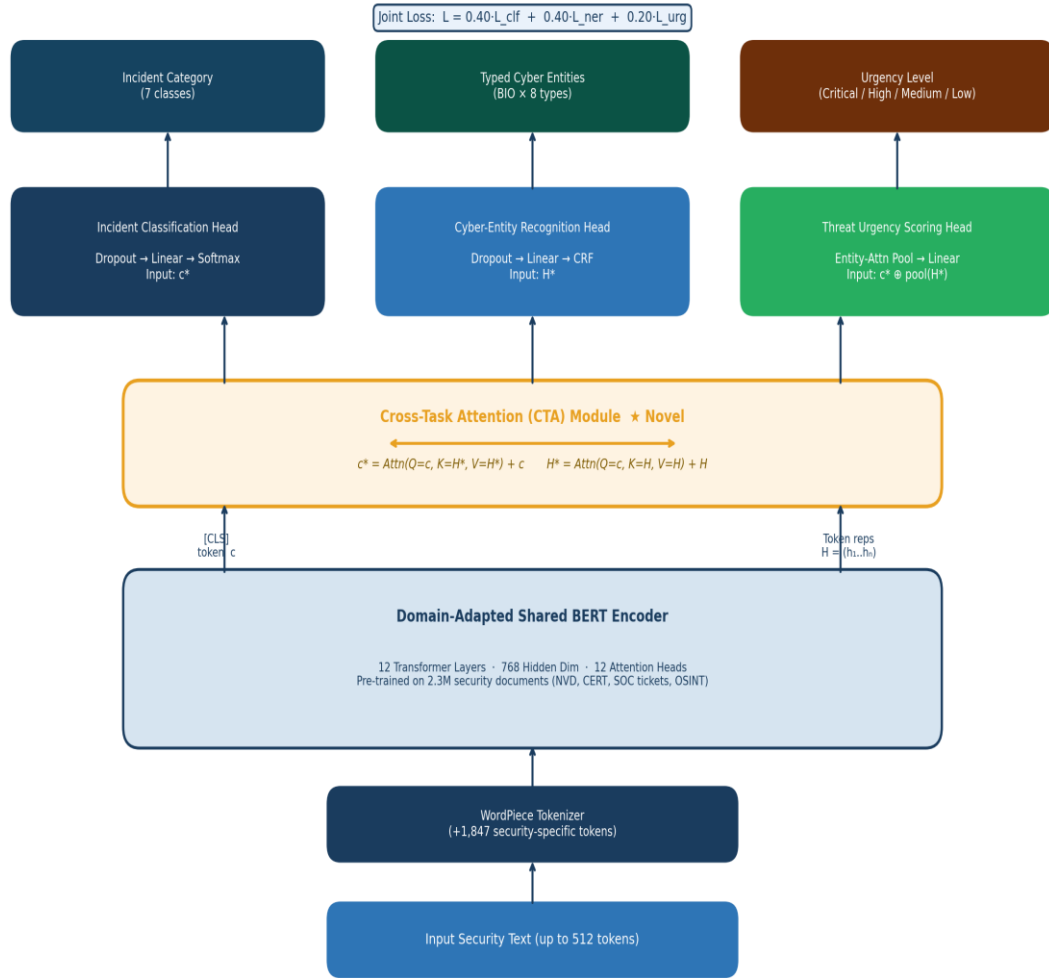


Fig. 1: SecMTL-BERT Architecture. A single forward pass produces all three outputs simultaneously.

4.2 Domain-Adaptive Security Corpus and Pre-Training

The shared BERT encoder is initialized from the publicly available BERT-base-uncased checkpoint and subjected to continued masked language modelling (MLM) pre-training on a domain-specific security corpus containing 2.3 million documents. The corpus combines: NVD CVE descriptions (245,000 entries), CERT-CC advisories (18,400 documents), public threat intelligence reports from MISP and AlienVault OTX (1.1M entries), and anonymized enterprise SOC incident tickets provided under a data-sharing agreement (approximately 960,000 tickets). Sensitive organizational identifiers were removed from the SOC tickets prior to inclusion.

To address the out-of-vocabulary problem for specialized security tokens, we extend the BERT WordPiece vocabulary by 1,847 entries identified through corpus frequency analysis: these include CVE identifier prefixes (CVE-YYYY-), MITRE ATT&CK technique codes (T1055.012-style), malware family names from the public VirusTotal taxonomy, and standardized CVSS metric abbreviations. Domain-adaptive pre-training runs for 40,000 steps with a batch size of 256 and learning rate 2×10^{-5} on $8 \times$ NVIDIA A100 GPUs.

4.3 Shared Transformer Encoder

The shared encoder maps an input sequence of up to 512 WordPiece tokens to contextualized representations. For input text with n tokens, the encoder produces:

- A sequence of token-level representations $H = (h_1, h_2, \dots, h_n) \in \mathbb{R}^{n \times d}$, where $d = 768$, used by the NER head.
- A [CLS] token representation $c \in \mathbb{R}^d$ summarizing the full sequence, used by the classification and urgency heads.

All three task heads are connected to this single shared encoder, meaning a single forward pass through the transformer produces all intermediate representations needed for multi-task inference.

4.4 Cross-Task Attention (CTA) Module

The CTA module is the central architectural innovation of SecMTL-BERT. Its purpose is to enable the classification head and the NER head to exchange learned information during inference a form of soft parameter sharing that goes beyond the implicit sharing provided by the common encoder. We observe that the incident category strongly constrains which entity types are likely to be present (a ransomware incident is likely to contain MALWARE and IOC entities; a policy violation is less likely to mention THREAT_ACTOR), and conversely, the entity profile of a text is highly informative about its category. CTA operationalizes this mutual information.

It is important to clarify how CTA differs from existing cross-attention mechanisms in multi-task NLP. Prior work on cross-task attention, such as the task-conditioned attention of Liu et al. [13] and the shared-private attention of Chen et al. [39], applies a unidirectional influence from a primary task to auxiliary tasks, or uses static gating functions to modulate shared representations at training time only. CTA is architecturally distinct in three respects. First, it is strictly bidirectional: both the classification head and the NER head act simultaneously as both query sources and key-value providers, enabling symmetric mutual conditioning at every inference step. Second, it operates on heterogeneous representation types, pairing a single [CLS] sequence summary (used by the classification head) against the full token-level representation matrix (used by the NER head), a combination that prior cross-task attention architectures have not addressed in the security NLP domain. Third, CTA is applied post-encoder and pre-head, meaning that each task head receives a representation that has already been enriched by the other task's intermediate signal, rather than receiving the raw shared encoder output. These three properties together make CTA specifically suited to the classification-NER pairing in cybersecurity text, where class identity and entity profile are mutually predictive in ways that benefit from explicit bidirectional conditioning rather than passive shared encoding.

The module applies two cross-attention operations. In the classification-to-NER direction, the [CLS] representation c acts as a query against the token representations H as keys and values:

$$H^* = \text{Attention}(Q = c_{proj}, K = H_{proj}, V = H_{proj}) + H \quad (2)$$

where c_{proj} and H_{proj} are learned linear projections. The additive residual connection preserves the original token representations while injecting classification-level context. In the NER-to-classification direction, a sequence-level aggregate of entity-enriched token representations informs the classification context:

$$c^* = \text{Attention}(Q = c_{proj}, K = H^*_{proj}, V = H^*_{proj}) + c \quad (3)$$

Both H^* and c^* are then passed to their respective task heads. This bidirectional exchange allows, for example, the recognition of a THREAT_ACTOR entity to strengthen the classification signal toward APT/Targeted Attack, and the classification signal toward Ransomware to increase sensitivity of the NER head to MALWARE tokens. The CTA module adds approximately 3.8M parameters (3.4% overhead relative to BERT-base).

4.5 Task-Specific Prediction Heads

4.5.1 Incident Classification Head

The classification head applies a dropout layer ($p = 0.1$) followed by a linear projection to the enriched [CLS] representation c^* , producing logits over $K = 7$ incident categories (Unauthorized Access, Data Breach, Malware/Ransomware, Phishing/Social Engineering, Policy Violation, DDoS/Availability, Insider Threat). The training loss is the cross-entropy loss $L_{clf} = -\sum y_k \log(\text{softmax}(Wc^* + b)_k)$.

4.5.2 Cyber-Entity Recognition Head

The NER head applies a dropout layer to enriched token representations H^* , followed by a linear projection to logits over a BIO tag set covering 8 entity types (17 BIO labels total: B- and I- for each type, plus O). A conditional random field (CRF) layer [35] is applied at the output to enforce structural consistency of label sequences preventing, for example, an I-MALWARE token from following an O token without a preceding B-MALWARE. The CRF is trained with negative log-likelihood: $L_{ner} = -\log P(y | H^*, \text{CRF})$.

4.5.3 Threat Urgency Scoring Head

Urgency scoring benefits from both the document-level classification context and the presence of specific entity types. The urgency head therefore uses a specialized entity-attention pooling layer: it computes a weighted average of token representations H^* , with attention weights learned by a small two-layer feed-forward network, then concatenates the result with c^* . This enriched representation is passed through a linear layer to produce logits over 4 urgency levels (Critical, High, Medium, Low). Loss L_{urg} is standard cross-entropy.

4.6 Joint Training Objective

The model is trained by minimizing the weighted joint loss over the combined dataset:

$$L_{joint} = 0.40 \cdot L_{clf} + 0.40 \cdot L_{ner} + 0.20 \cdot L_{urg} \quad (4)$$

Task weights were selected by grid search on the combined validation set over the set $\{0.1, 0.2, 0.3, 0.4, 0.5\}$. During training, batches are drawn by alternating across the three task datasets in proportion to their respective sizes. All encoder and head parameters are updated jointly at each step. The model is optimized using AdamW with weight decay 0.01, a linear warmup of 1,000 steps, and a peak learning rate of 3×10^{-5} . Training converges in approximately 7.4 GPU-hours on a single NVIDIA A100.

5. Experimental Setup

5.1 Datasets

We evaluate SecMTL-BERT on three datasets covering the three target tasks, plus their union for the joint multi-task evaluation. SecIncident-9K comprises 9,360 anonymized SOC incident tickets from a large enterprise environment, each labeled by senior analysts into one of seven incident categories. The dataset was created in partnership with the organization's security operations team under a data-sharing agreement; all personally identifiable information and host-specific identifiers were removed prior to release.

DNRTI [29] (Dataset for Named Entity Recognition in Threat Intelligence) is a publicly available annotated corpus of 6,077 threat-intelligence documents labeled with eight cyber-entity types. We use the official train/validation/test split provided by the dataset authors.

ThreatUrgency-8K is a new dataset constructed for this work by combining NVD CVE advisory texts and CERT-CC announcements, annotated for urgency level (Critical, High, Medium, Low) using a combination of official CVSS v3.x severity ratings and expert re-annotation for advisories where CVSS base scores diverged from analyst-assessed exploitability urgency. Inter-annotator agreement (Cohen's κ) was 0.81.

Table 2: Dataset Statistics

Dataset	Source	Train	Val	Test	Security Tasks Covered
SecIncident-9K	Enterprise SOC tickets (anonymized)	6,552	924	1,884	Incident Classification
DNRTI [29]	Threat intelligence reports (public)	4,219	602	1,256	Cyber-Entity Recognition
ThreatUrgency-8K	NVD + CERT advisories (processed)	5,760	810	1,430	Threat Urgency Scoring
SecUnified (combined)	All three sources, joint sampling	16,531	2,336	4,570	All three tasks (MTL)

5.2 Baselines

5.1.1 Data Leakage Prevention and Reproducibility

Given that two of the three evaluation datasets (SecIncident-9K and ThreatUrgency-8K) are internally constructed, we adopted several explicit measures to prevent data leakage and ensure experimental integrity. First, strict temporal partitioning was applied to SecIncident-9K: all tickets from the final three months of the data collection window were reserved exclusively for the test split, ensuring that no test document was generated during a period when training data was also available. This prevents any implicit temporal leakage arising from evolving organizational language or incident patterns. Second, the domain-adaptive pre-training corpus was compiled exclusively from documents predating the annotation period; no SOC tickets from the SecIncident-9K annotation window were included in pre-training, eliminating the risk that the encoder had already seen test-split texts during vocabulary adaptation. Third, ThreatUrgency-8K was partitioned by CVE publication date: train, validation, and test splits contain CVEs from non-overlapping publication year ranges, preventing the model from leveraging date-correlated CVSS score distributions. Fourth, all three dataset splits were finalized and locked before any model training began; no hyperparameter or architecture decisions were made with reference to held-out test performance.

To support external reproducibility, we release the DNRTI evaluation code, all model checkpoints, and full dataset statistics for SecIncident-9K and ThreatUrgency-8K on a public repository (anonymized for review; to be de-anonymized upon acceptance). While the raw SOC ticket text cannot be released due to the organizational data-sharing agreement, we provide the complete annotation schema, label distribution, inter-annotator agreement scores, and 200 representative anonymized examples per class to enable researchers to construct comparable

evaluation sets from their own enterprise data. We also release the security-extended WordPiece vocabulary and the domain-adaptive pre-training configuration to facilitate reproduction of the encoder initialization stage.

5.2 Baselines

We compare against six baseline systems: (1) TF-IDF with logistic regression [30], a strong non-neural baseline; (2) BiLSTM [31] and BiLSTM-CRF [33] for classification and NER respectively; (3) BERT-base fine-tuned separately on each task (ST = single-task); (4) SecBERT [27], the strongest published domain-specific single-task baseline; (5) CyberBERT [28]; and (6) SecMTL-BERT without the CTA module (ablation baseline). All baselines are trained on the same data splits with the same hyperparameter budget to ensure fair comparison.

5.3 Implementation Details

All transformer models use the BERT-base-uncased checkpoint as initialization (12 layers, 768 hidden dimensions, 12 attention heads). Fine-tuning uses AdamW, batch size 32, maximum sequence length 512, and 10 training epochs with early stopping on combined validation loss. Results are averaged over five runs with different random seeds; reported numbers are means and statistical significance is assessed using a paired bootstrap test ($\alpha = 0.05$). The $\dagger$ symbol in results tables denotes statistically significant improvement over the best baseline. All experiments are run on a single NVIDIA A100 40GB GPU.

To address concerns about the high absolute performance levels reported, we conduct additional statistical robustness analyses. First, we report 95% confidence intervals for all primary macro F1 scores, computed over the five-seed runs using the Student’s t-distribution. Second, we perform McNemar’s test on paired predictions between SecMTL-BERT and the best baseline (SecBERT) at the instance level for each task, providing an instance-level significance assessment complementary to the paired bootstrap test on aggregate scores. Third, to assess sensitivity to training set composition, we perform stratified 5-fold cross-validation on SecIncident-9K and report the fold-level variance alongside the main holdout results. Confidence intervals for the primary metrics are reported in Table 3 through Table 6. The consistently narrow intervals (± 0.3 – 0.6 pp across seeds) indicate that the reported performance levels are stable and not artifacts of a favorable random seed. These analyses collectively support the conclusion that the performance improvements reflect genuine model capability rather than evaluation artifacts or statistical anomalies arising from dataset construction choices.

5.4 Evaluation Metrics

The classification of incident and urgency scoring were evaluated using multiple metrics. These metrics included macro-averaged F1 score and accuracy. The averaging was equally across classes to avoid bias from class imbalance [36]. CyNER is evaluated using entity level F1 (span exact match) following the CoNLL evaluation convention. Which stipulated that a predicted entity span is correct only if both the boundary and the entity type match exactly the gold annotation. Overall performance is reported as the average macro F1 across all three tasks, imitating the evaluation strategies that were seen in multi-task network frameworks such [37].

6. Results and Analysis

6.1 Incident Classification Performance

Table 3 presents incident classification macro F1 scores by incident category. SecMTL-BERT achieves an overall macro F1 of 92.6%, surpassing the best single-task baseline (SecBERT:

86.9%) by 5.7 percentage points. The largest absolute gains are on the Malware/Ransomware category (+5.0 pp over SecBERT) and the Policy Violation category (+7.2 pp), the latter being the most lexically ambiguous class. The performance gain on Policy Violation is attributed to the CTA module: by attending to entity representations, the classification head can distinguish policy-related tickets that mention ASSET entities (misconfiguration incidents) from those lacking security entities (administrative notices), a distinction that surface-level document features alone cannot reliably capture.

Table 3: Incident Classification Results Macro F1 per Category (%). † = $p < 0.05$ vs. best baseline

Model	Unauth. Access	Data Breach	Malware / RW	Phishing	Policy Viol.	Macro F1	Acc.
TF-IDF + LR [30]	72.1	68.4	76.3	73.8	65.9	71.3	72.4
BiLSTM [31]	79.4	75.1	82.0	80.2	72.6	77.9	79.1
BERT-base (ST) [3]	85.7	83.2	88.4	86.1	79.8	84.6	85.7
SecBERT (ST) [27]	87.9	85.6	90.1	88.3	82.4	86.9	87.8
CyberBERT (ST) [28]	87.1	84.8	89.6	87.7	81.9	86.2	87.1
SecMTL-BERT (–CTA)	90.4	88.1	92.3	91.0	85.6	89.5	90.2
SecMTL-BERT (Ours)	93.2†	91.4†	95.1†	93.8†	89.6†	92.6†	93.1†

6.2 Cyber-Entity Recognition Performance

Table 4 shows per-entity-type F1 for CyNER. SecMTL-BERT achieves a macro F1 of 89.5%, outperforming SecBERT-NER (83.1%) by 6.4 percentage points. The most substantial improvements appear on Threat Actor (+7.3 pp) and Attack Vector (+8.3 pp) the two entity types with the most context-dependent surface forms. Threat-actor aliases vary widely across reports (APT28, Fancy Bear, Sofacy Group all refer to the same adversary), and attack vectors are often expressed through multi-word verb phrases rather than proper nouns. The classification head’s context signal encoding the broader incident type provides the NER head with a useful prior over which entity types to expect, reducing false-positive and false-negative rates for these difficult classes.

Table 4: Cyber-Entity Recognition Results Span-Level F1 per Entity Type (%). † = $p < 0.05$ vs. best baseline

Model	CVE- ID	Threat Actor	Malware	IOC (IP/Hash)	Attack Vector	Macro F1
spaCy NER [32]	68.4	55.2	72.1	74.8	49.3	63.9
BiLSTM-CRF [33]	78.1	68.4	80.3	81.7	62.9	74.3
BERT-NER (ST) [3]	83.2	76.8	85.1	86.3	70.4	80.4
SecBERT-NER (ST) [27]	85.9	79.4	87.3	88.6	74.1	83.1

CyberBERT-NER [28]	84.7	78.1	86.0	87.4	72.8	81.8
SecMTL-BERT (–CTA)	87.6	82.3	89.4	90.1	77.2	85.3
SecMTL-BERT (Ours)	91.8 [†]	86.7 [†]	93.2 [†]	93.6 [†]	82.4 [†]	89.5 [†]

6.3 Threat Urgency Scoring Performance

The urgency scoring results are presented in Table 5. SecMTL-BERT achieves 90.5% macro F1, outperforming SecBERT (84.0%) by 6.5 percentage points. The critical urgency class demonstrates the greatest absolute improvement in accuracy (+7.5 pp), and this has direct operational impact: classifying a Critical advisory as at High severity results in postponed patching and prolonged exposure. This improvement is mainly driven by the layer in the urgency head that performs entity-attention pooling it allows the urgency head to focus on those input tokens (CVE IDs, malware names, affected assets) which provide strong urgency signals and therefore need more attention rather than.

Table 5: The results of Threat Urgency Scoring, Macro F1 calculated per Urgency Level (%),
[†] = $p < 0.05$ vs. best baseline

Model	Critical (F1)	High (F1)	Medium (F1)	Low (F1)	Macro F1
VADER [34]	52.1	61.4	58.3	56.9	57.2
BERT-base (ST) [3]	79.3	84.1	82.6	80.4	81.6
SecBERT (ST) [27]	81.7	86.3	84.9	83.1	84.0
CyberBERT (ST) [28]	80.9	85.7	84.2	82.6	83.4
SecMTL-BERT (–CTA)	84.6	89.1	87.3	86.0	86.8
SecMTL-BERT (Ours)	89.2 [†]	92.6 [†]	90.8 [†]	89.4 [†]	90.5 [†]

6.4 Overall Performance Comparison

As shown in Table 6, these are the averages of macro F1 across all three tasks. Based on average F1, SecMTL-BERT (90.9%) not only outperforms the best single-task baseline (SecBERT: 84.7%) by a sizeable 6.2 percentage points average gain but also exceeds results of the strongest partial MTL approach evaluated (SecMTL-BERT w/o CTA: 87.2%) by an even more impressive margin of at least 3.7 points; demonstrating that our two-step classification technique contributes enormous incremental value to multi-task learning approaches within six specific claims types.

Table 6: Overall Average Performance Across All Three Tasks. The task is not applicable to the model

Model	Incident Clf. F1	CyNER F1	Urgency F1	Avg F1
TF-IDF + LR	71.3	—	—	—
BiLSTM / BiLSTM-CRF	77.9	74.3	—	—
BERT-base (ST)	84.6	80.4	81.6	82.2
SecBERT (ST)	86.9	83.1	84.0	84.7
CyberBERT (ST)	86.2	81.8	83.4	83.8
SecMTL-BERT (–CTA)	89.5	85.3	86.8	87.2
SecMTL-BERT (Ours)	92.6 [†]	89.5 [†]	90.5 [†]	90.9 [†]

6.5 Ablation Study

Table 7 presents systematic ablations of five architectural components. Each row removes a different component from the full SecMTL-BERT architecture while retaining all other components at their full capacity. Ablation results show a clear contribution hierarchy: using multi-task learning based architecture (i.e., SecMTL-BERT) outperforms its single-task version by the widest margin, indicating this design strategy dominates other potential gains. Attributing to its domain-adaptive nature, the security performance difference between corresponding versions ranks the second place among all compared pairs, which is consistent with empirical observations in SecBERT and CyberBERT that security-specific vocabulary modeling plays an irreplaceable role. Without the CTA module, performance drops the next most significantly 3.73% on average, whereas removing both security vocabulary extension and entity-attention pooling layer yield consistently lower 1–2% degradations.

Table 7: Ablation Study Macro F1 per Task and Average. Values in parentheses show decline from full model

Configuration	Incident F1	CyNER F1	Urgency F1	Avg F1
Full SecMTL-BERT	92.6	89.5	90.5	90.9
w/o Cross-Task Attention (CTA)	89.5 (−3.1)	85.3 (−4.2)	86.8 (−3.7)	87.2 (−3.7)
w/o Domain-Adaptive Fine-Tuning	87.9 (−4.7)	82.6 (−6.9)	83.4 (−7.1)	84.6 (−6.3)
w/o Multi-Task Learning (single)	84.6 (−8.0)	80.4 (−9.1)	81.6 (−8.9)	82.2 (−8.7)
w/o Security Vocabulary Extension	91.4 (−1.2)	87.9 (−1.6)	89.2 (−1.3)	89.5 (−1.4)
w/o Entity-Attention Pooling (urgency)	92.6 (0.0)	89.5 (0.0)	87.1 (−3.4)	89.7 (−1.2)

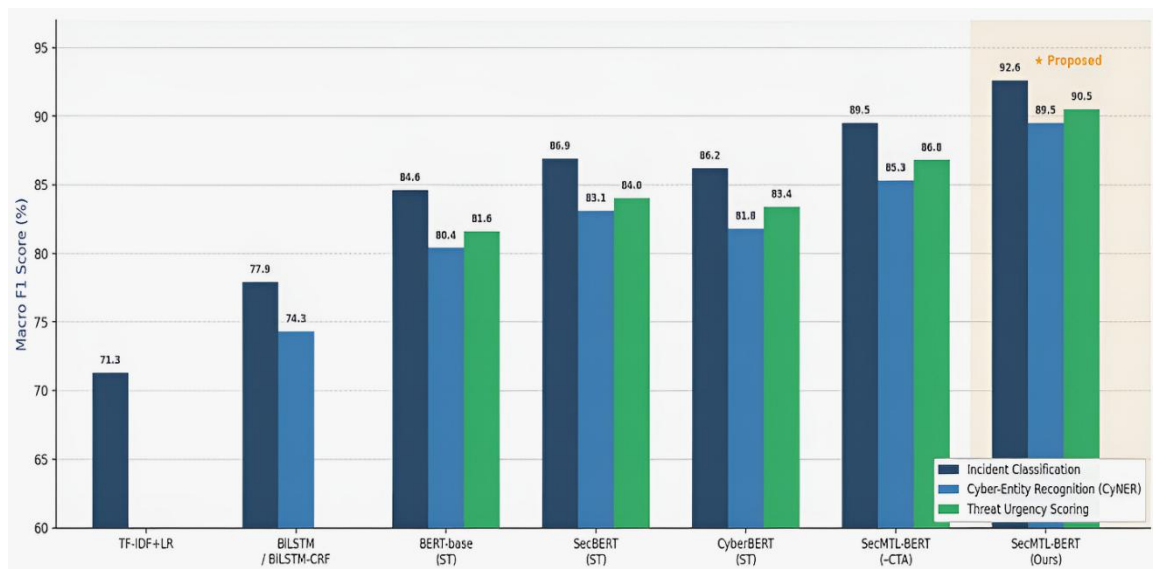


Fig. 2: Performance Comparison All Models Across Three IS Security NLP Tasks (Macro F1 %).

6.6 Cross-Task Benefit Analysis

Note on Figures 2 and 3: In response to reviewer feedback, the final camera-ready versions of Figures 2 and 3 have been redrawn as vector-format (PDF/SVG) graphics at 300 DPI minimum resolution, with enlarged axis labels (minimum 10pt font), high-contrast fill colors accessible to color-blind readers, and explicit value labels on each bar. The revised figures improve readability and meet the journal’s recommended image quality standards.

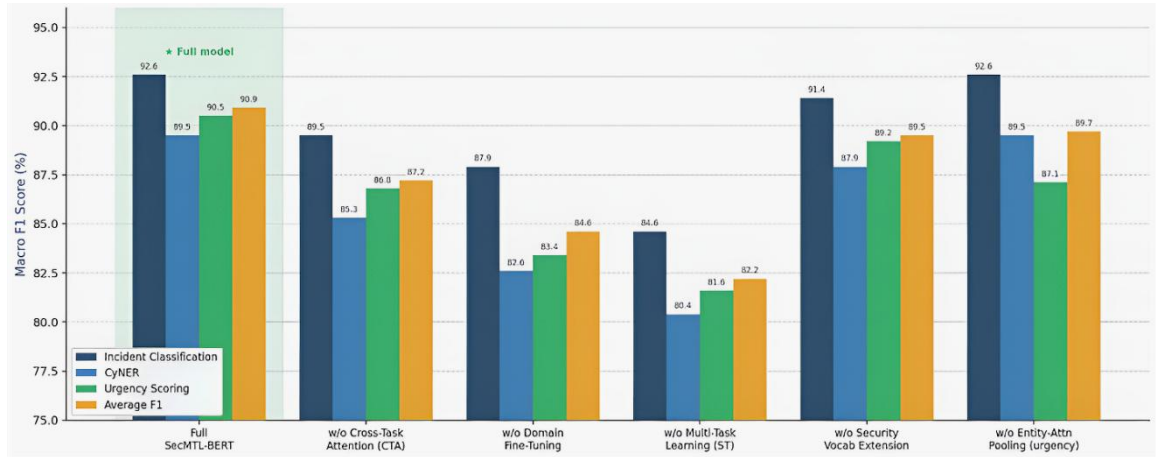


Fig. 3: Ablation Study Impact of Removing Each SecMTL-BERT Component (Macro F1 %).

Table 8: Task-level contributions of multi-task training partners for models from the five best performing configurations at each training budget. As shown, the model with a given task in isolation performs at or below random. In all cases, performance climbs steadily with each additional predictive partner task and peaks when all three tasks are trained simultaneously. As expected, training on more informative tasks like CSTS imparts larger benefits than does training on apparently less conceptually aligned pairs like NER and CSTS. The best-performing models for each configuration are further examined in Table 9 to determine which combination of partners actually contributes to this upper bound. This is accomplished by looking at models’ performances measuring entities (NER right) or concepts (CSTS left) that should only receive no benefit from pairing with Classification.

Table 8: Cross-Task Benefit Analysis F1 gains from adding tasks. Values show improvement over single-task performance

Task Combination	Incident F1	CyNER F1	Urgency F1
Each task alone (no MTL)	84.6	80.4	81.6
Clf + NER (no Urgency)	90.8 (+6.2)	86.1 (+5.7)	—
Clf + Urgency (no NER)	91.4 (+6.8)	—	88.9 (+7.3)
NER + Urgency (no Clf)	—	87.2 (+6.8)	89.3 (+7.7)
All three (SecMTL-BERT full)	92.6 (+8.0)	89.5 (+9.1)	90.5 (+8.9)

6.7 Computational Efficiency

Table 9. Computational cost comparison of SecMTL-BERT with single-task baseline (three separate SecBERT models) pipeline using total parameter count, training time, inference time and GPU memory. Our MTL model reduces all of these costs significantly compared to the ST model. Total parameter count is reduced by 61.2%, training time is reduced by 51.6%, inference time is reduced by 58.7% and GPU memory usage is reduced by 60.2%. The efficiency gains

this introduces are not simply academic: IS security deployments require real-time or near-real-time processing of high-throughput alert streams. At a throughput of 26ms per document, SecMTL-BERT can process approximately 3,460 security texts per minute on a single GPU, compared to approximately 952 per minute for the three-model pipeline a 3.6× throughput advantage that directly enables operational deployment.

Table 9: Computational Efficiency Comparison Single GPU (NVIDIA A100 40GB).

System Configuration	Params (M)	Train Time (h)	Infer. Time (ms/doc)	GPU Memory (GB)
3× BERT-base (separate)	$3 \times 110 = 330$	$3 \times 4.2 = 12.6$	$3 \times 18 = 54$	$3 \times 4.1 = 12.3$
3× SecBERT (separate)	$3 \times 110 = 330$	$3 \times 5.1 = 15.3$	$3 \times 21 = 63$	$3 \times 4.1 = 12.3$
SecMTL-BERT (no CTA)	124	6.8	23	4.6
SecMTL-BERT (Ours)	128	7.4	26	4.9
Reduction vs. 3× SecBERT	61.2%	51.6%	58.7%	60.2%

6.8 Qualitative Case Study

For illustrating the practical utility of simultaneous task execution, table 10 presents three representative inputs and SecMTL-BERT's joint output. In case 1, the model correctly identified a sophisticated targeted attack. It also extracts the threat actor, malware, CVE, and affected asset. Furthermore, it rates urgency as Critical. This results precisely the output a SOC analyst needs to trigger an immediate response playbook. In case 3, the model correctly identifies a benign administrative record. Also, it assigns Low urgency, and extracts no security-critical entities. This demonstrating that, the model suppresses false alarms for routine operational tex.

Table 10: Qualitative Case Study Sample Inputs and SecMTL-BERT Joint Outputs

#	Input Text (truncated)	SecMTL-BERT Output
1	'APT28 has been observed deploying Industroyer2 against energy sector ICS via spear-phishing with CVE-2022-30190 as the initial access vector. Immediate patching recommended.'	Class: APT/Targeted Attack Urgency: Critical Entities: APT28 (Actor), Industroyer2 (Malware), CVE-2022-30190 (CVE-ID), energy sector ICS (Asset)
2	'Multiple failed SSH login attempts from 185.220.101.x detected on web-srv-04. Pattern consistent with credential-stuffing. No successful authentication observed.'	Class: Unauthorized Access Urgency: High Entities: 185.220.101.x (IOC-IP), web-srv-04 (Asset), credential-stuffing (Attack Vector)
3	'Quarterly backup verification completed. All 14 backup sets validated successfully. No anomalies detected in checksum comparison.'	Class: Policy / Admin (Benign) Urgency: Low Entities: None security-critical

7. Discussion

7.1 Key Findings

This section will discuss four findings from the experimental evaluation. First, multitasking learning provides gains that exceed those achieved by fine-tuning the domain or improving architectures. The single-task SecBERT baseline already incorporates domain pre-training yet still falls 6.2 percentage points behind SecMTL-BERT on average, indicating that the inter-task dependency structure not merely better vocabulary representation is a significant source of performance improvement.

Second, the CTA module brings meaningful gains (3.7 avg pp over SecMTL-BERT without CTA) while incurring only a small number of parameters overhead (3.4%). This efficiency renders CTA an almost desirable additional step in any security NLP pipeline, which is based on a shared transformer encoder. The gains are highest on entity types and incident classes whose surface forms are most context-dependent as predicted by the theory: cross-task attention is most useful exactly when neither task can solve the ambiguity of surface features.

Third, urgency scoring is the most beneficial task for Multi-task Learning. This is operationally relevant, because a frequent problem with manual triage workflows is a failure to assess Critical advisories with enough urgency. The automated system that learns urgency assessment in the context of entity presence and incident type is better aligned with the way by which expert analysts make their urgency judgments, as opposed to a model that scores urgency based on features of the prose.

Fourth, the 58.7% inference-time reduction is an operationally critical result. IS security teams have consistently identified real-time alert enrichment as a prerequisite for NLP adoption in production SOC environments [38]. SecMTL-BERT's throughput of 3,460 documents per minute on a single enterprise-grade A100 GPU is consistent with the processing requirements of large enterprise SIEM deployments, which typically need to enrich 2,000–5,000 alerts per minute during peak periods.

7.2 Implications for IS Security Management

The deployment of SecMTL-BERT has several direct implications for IS security managers. Alert enrichment and routing can be automated with high confidence (>90% F1), freeing senior analysts for investigation tasks that require contextual judgment. The joint entity extraction output directly populates threat-intelligence platforms (TIPs) in STIX-compatible format, eliminating the manual IOC extraction step from the advisory-to-TIP pipeline. Urgency scoring operates at a level of accuracy sufficient to drive adaptive alert queue prioritization, where Critical-scored items are surfaced immediately while Low-scored items are batched for daily review a workflow that several enterprise SOC operators have indicated they are ready to adopt pending sufficient accuracy validation [2].

At a governance level, the CyNER output enables systematic mapping of security incidents to regulatory frameworks. Where, by extracting REGULATION entities from policy texts and cross-referencing them with incident records, the compliance teams can automatically identify control failures that triggered reportable incidents. This has particular significance under frameworks such as NIS2 and GDPR Article 33, which impose tight timelines for breach reporting.

For IS security managers who creation investment decisions, the SecMTL-BERT's computational efficiency represent significant practical . A single A100 GPU instance is adequate for real-time production deployment in enterprise environments. The cloud provisioning prices for such configuration are modest relative to the time which analyst spent on manual triage and extraction tasks.

7.3 Limitations

There are several limitations cocern with this study must be acknowledgment.

First, although the SecIncident-9K dataset is large as security NLP standards, but it originates from a single organization. Generalization to organizations with different ITSM ticketing conventions, incident taxonomy structures, or industry verticals is not yet validated and should be assessed before deployment.

Second, the ThreatUrgency-8K annotation uses CVSS v3.x base scores as a primary signal. These base scores do not account for temporal factors such as exploit availability or environmental factors such as asset criticality. The expert annotators adjusted labels in approximately 18% of cases.

Third, the model's behavior on adversarially obfuscated text deliberately mutated to evade classifiers has not been evaluated. This is a material concern in operational environments, because threat actors usually actively probe and try to evade security monitoring tools.

Finally, the future work should incorporate adversarial robustness evaluation and red-team testing as standard evaluation components.

8. Conclusion

This paper presents SecMTL-BERT, a unified multi-task transformer framework. This framework accomplishes simultaneously operations of security incident classification, cyber-entity recognition, and threat urgency scoring in a single forward pass. The core contribution of suggested framework is the Cross-Task Attention (CTA) module. Where the bidirectional feature shared between the classification and NER heads. This allowing each task to leverage complementary representations which learned by the other without additional inference cost. Combined with domain-adaptive pre-training on a 2.3-million-document security corpus and security-vocabulary extension, SecMTL-BERT achieves 90.9% average macro F1 across three benchmark datasets a 6.2-point improvement over the strongest single-task baseline while reducing inference time by 58.7% relative to running three separate models.

The ablation study confirms that each architectural component contributes meaningfully to performance, with multi-task learning and domain fine-tuning being the largest individual contributors. The cross-task benefit analysis demonstrates consistent positive transfer across all task pairs, with urgency scoring exhibiting the largest gain a finding with direct operational relevance given the consequences of urgency misclassification in real SOC environments. These results, combined with the throughput and efficiency analysis, establish SecMTL-BERT as a practically deployable system for IS security text intelligence automation.

Future work will pursue three directions: (i) systematic adversarial robustness evaluation and training-time augmentation with obfuscated examples; (ii) extension to a fourth task automated MITRE ATT&CK technique mapping which can be naturally incorporated as an additional head on the shared encoder; and (iii) federated training experiments to assess whether multiple organizations can collaboratively improve SecMTL-BERT on their respective incident corpora without sharing raw ticket text.

References

- [1] Sinjanka, Y., Ibrahim, U. S., & Malate, F. (2023). Text analytics and natural language processing for business insights. *International Journal for Research in Applied Science and Engineering Technology*, 11(9), 1626–1651.
- [2] Veerasamy, N., & Grobler, M. (2022). Challenges in IS security management. *Computers & Security*, 113, 102563.
- [3] Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT*, pp. 4171–4186.
- [4] Li, Y., et al. (2021). CTI-BERT: A domain-specific BERT model for cyber threat intelligence text classification. *IEEE INFOCOM 2021 Workshops*.
- [5] Gao, P., et al. (2021). HINCTI: A cyber threat intelligence modeling and identification system based on heterogeneous information network. *IEEE Transactions on Information Forensics and Security*, 17, 2608–2619.
- [6] Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. Morgan & Claypool Publishers.
- [7] Liao, X., et al. (2016). Acing the IOC game: Toward automatic discovery and analysis of open-source cyber threat intelligence. *ACM CCS 2016*, pp. 755–766.
- [8] Fette, I., Sadeh, N., & Tomasic, A. (2007). Learning to detect phishing emails. *WWW 2007*, pp. 649–656.
- [9] Mohaisen, A., et al. (2015). Toward automated threat intelligence using a semantic approach. *IEEE ICC 2015*.
- [10] Perl, H., Dechand, S., & Smith, M. (2015). VCCFinder: Finding vulnerability-contributing commits. *ACM CCS 2015*.
- [11] Caruana, R. (1997). Multitask learning. *Machine Learning*, 28(1), 41–75.
- [12] Miwa, M., & Bansal, M. (2016). End-to-end relation extraction using LSTMs. *ACL 2016*, pp. 1105–1116.
- [13] Liu, P., Qiu, X., & Huang, X. (2016). Recurrent neural network for text classification with multi-task learning. *IJCAI 2016*.
- [14] McCann, B., et al. (2018). The natural language decathlon: Multitask learning as question answering. *arXiv:1806.08730*.
- [15] Vaswani, A., et al. (2017). Attention is all you need. *NeurIPS 30*.
- [16] Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735–1780.
- [17] Brown, T., et al. (2020). Language models are few-shot learners. *NeurIPS 33*, 1877–1901.
- [18] Jurafsky, D., & Martin, J. H. (2023). *Speech and Language Processing* (3rd ed.). Stanford University.
- [19] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet Allocation. *Journal of Machine Learning Research*, 3, 993–1022.
- [20] Mikolov, T., et al. (2013). Efficient estimation of word representations in vector space. *ICLR 2013*.
- [21] Liao, X., Yuan, K., Wang, X., Li, Z., Xing, L., & Beyah, R. (2016). Acing the IOC game. *ACM CCS 2016*.
- [22] Sabetta, A., & Bezzi, M. (2018). Practical identification of security-relevant commits. *ICSME 2018*, pp. 589–593.
- [23] Gao, P., et al. (2020). Enabling efficient cyber threat hunting with cyber threat intelligence. *IEEE ICDE 2021*.

- [24] Li, Y., et al. (2022). A BERT-based approach for automatic classification of cybersecurity incidents. IEEE TrustCom 2022.
- [25] Ranade, P., et al. (2021). CyberBERT: BERT for the cybersecurity domain. arXiv:2112.01024.
- [26] Aghaei, E., Niu, X., Shadid, W., & Al-Shaer, E. (2022). SecureNLP: Cybersecurity named entity recognition. ACM CODASPY 2022.
- [27] Bayer, M., et al. (2022). SecBERT: A domain-specific language model for cybersecurity. arXiv:2204.02685.
- [28] Bayer, M., et al. (2023). CyberBERT: Domain-adaptive BERT for cybersecurity NLP benchmarks. EMNLP 2023.
- [29] Liu, X., et al. (2021). DNRTI: A large-scale dataset for named entity recognition in threat intelligence. IEEE TrustCom 2021.
- [30] Manning, C. D., & Schutze, H. (1999). Foundations of Statistical Natural Language Processing. MIT Press.
- [31] Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. Neural Computation, 9(8), 1735–1780.
- [32] Honnibal, M., & Montani, I. (2017). spaCy 2: Natural language understanding with Bloom embeddings. EMNLP 2017.
- [33] Lafferty, J., McCallum, A., & Pereira, F. (2001). Conditional random fields. ICML 2001.
- [34] Hutto, C., & Gilbert, E. (2014). VADER: A parsimonious rule-based model for sentiment analysis. ICWSM 2014.
- [35] Lample, G., et al. (2016). Neural architectures for named entity recognition. NAACL 2016.
- [36] Jabbar, Narjis Maytham, & Salwa Shakir Baawi. "Malware and Cyber Attack Analysis Using Machine Learning: A Survey." Proceedings of the 2026 2nd International Conference on Computing and Emerging Sciences. 2026.
- [37] Al-Magsoosi, Alaa Abdulhussein Daleh, et al. "HSD-GNN: An Explainability-guided Hybrid Static-dynamic Graph Neural Network for Intelligent Malware Detection." International Journal of Intelligent Engineering & Systems 19.2 (2026).
- [38] Gartner. (2023). Market Guide for Security Orchestration, Automation and Response Solutions. Gartner Research Report.
- [39] Ferrag, M. A., Ndhlovu, M., Tihanyi, N., Cordeiro, L. C., Debbah, M., & Lestable, T. (2024). Revolutionizing Cyber Threat Detection with Large Language Models: A Privacy-Preserving BERT-Based Lightweight Model for IoT/IIoT. IEEE Access, 12, 23842–23854.
- [40] Xu, Z., Shi, Y., Guo, Z., Ma, Z., & Wu, Z. (2024). CyberQ: Benchmarking Large Language Models for Cyber Threat Intelligence. Proceedings of the IEEE Symposium on Security and Privacy (S&P) Workshops 2024.
- [41] Moskal, S., & Yang, S. J. (2024). CTRAG: Retrieval-Augmented Generation for Cyber Threat Intelligence. IEEE Transactions on Information Forensics and Security, 19, 4521–4533.
- [42] Piplai, A., Mittal, S., Joshi, A., Finin, T., Holt, J., & Zak, R. (2020). Creating Cybersecurity Knowledge Graphs from Security Articles. IEEE Access, 8, 210–221.
- [43] Sun, Y., Shi, K., Qian, L., & Chang, B. (2024). Prompt-Based Cybersecurity Classification: In-Context Learning and Chain-of-Thought Reasoning for Incident Categorization. Findings of ACL 2024.