

SCALE: Constructing Structured Natural Language Comment Trees for Software Vulnerability Detection

Xin-Cheng Wen
Harbin Institute of Technology
Shenzhen, China
xiamenwxc@foxmail.com

Cuiyun Gao*
Harbin Institute of Technology
Shenzhen, China
gaocuiyun@hit.edu.cn

Shuzheng Gao
The Chinese University of Hong Kong
Hong Kong, China
szgao23@cse.cuhk.edu.hk

Yang Xiao
Chinese Academy of Sciences
Beijing, China
xiaoyang@iie.ac.cn

Michael R. Lyu
The Chinese University of Hong Kong
Hong Kong, China
lyu@cse.cuhk.edu.hk

ABSTRACT

Recently, there has been a growing interest in automatic software vulnerability detection. Pre-trained model-based approaches have demonstrated superior performance than other Deep Learning (DL)-based approaches in detecting vulnerabilities. However, the existing pre-trained model-based approaches generally employ code sequences as input during prediction, and may ignore vulnerability-related structural information, as reflected in the following two aspects. First, they tend to fail to infer the semantics of the code statements with complex logic such as those containing multiple operators and pointers. Second, they are hard to comprehend various code execution sequences, which is essential for precise vulnerability detection.

To mitigate the challenges, we propose a Structured Natural Language Comment tree-based vulnerability detection framework based on the pre-trained models, named **SCALE**. The proposed Structured Natural Language Comment Tree (SCT) integrates the semantics of code statements with code execution sequences based on the Abstract Syntax Trees (ASTs). Specifically, SCALE comprises three main modules: (1) *Comment Tree Construction*, which aims at enhancing the model's ability to infer the semantics of code statements by first incorporating Large Language Models (LLMs) for comment generation and then adding the comment node to ASTs. (2) *Structured Natural Language Comment Tree Construction*, which aims at explicitly involving code execution sequence by combining the code syntax templates with the comment tree. (3) *SCT-Enhanced Representation*, which finally incorporates the constructed SCTs for well capturing vulnerability patterns. Experimental results demonstrate that SCALE outperforms the best-performing baseline, including the pre-trained model and LLMs, with improvements of 2.96%, 13.47%, and 3.75% in terms of F1 score on the FFMpeg+Qemu, Reveal, and SVuLD datasets, respectively. Furthermore, SCALE can be applied to different pre-trained models, such as CodeBERT and UniXcoder, yielding the F1 score performance enhancements ranging from 1.37% to 10.87%.

* Corresponding author. The author is also affiliated with Peng Cheng Laboratory and Guangdong Provincial Key Laboratory of Novel Security Intelligence Technologies.

Xin-Cheng Wen, Cuiyun Gao*, Shuzheng Gao, Yang Xiao, and Michael R. Lyu. 2024. SCALE: Constructing Structured Natural Language Comment Trees for Software Vulnerability Detection. In *Proceedings of ACM SIGSOFT International Symposium on Software Testing and Analysis (ISSTA 2024)*. ACM, New York, NY, USA, 13 pages.

1 INTRODUCTION

Nowadays, modern software development is significantly afflicted by the prevalence of source code security vulnerabilities [42]. For instance, according to IBM, the year 2023 witnessed a peak in the average cost of data breaches, reaching US\$ 4.45 million [30]. Over the past decade, there has been a substantial increase in the quantity of identified Common Vulnerabilities and Exposures (CVEs) [2]. Specifically, during the 2022, the identified CVE number reached 25,227 [49], with a 25.1% increase over the number of vulnerabilities detected in 2021. Therefore, it becomes imperative to develop effective methods for detecting these software vulnerabilities.

Conventional Program Analysis (PA)-based vulnerability detection methods heavily rely on user-defined rules and specifications to identify vulnerabilities, such as INFER [16] and CheckMarx [31], which makes them labor-intensive and time-consuming. Recently, DL-based methods have achieved great success as they can reduce the dependence on expert knowledge and offer a broader spectrum of capabilities in detecting various types of software vulnerabilities [13]. Early DL-based methods use Convolutional Neural Networks (CNNs) [61, 62], Recurrent Neural Networks (RNNs) [38, 48] and Graph Neural Networks (GNNs) [10, 65] for supervised training. However, the performance of these models can be largely limited due to the scarcity of vulnerability data [45]. Recently, the advent of pre-trained models has further advanced this field. These models are trained on massive open-source code repositories and have possessed a vast general programming knowledge. The recent studies mainly resort to the pre-training fine-tuning paradigm which injects vulnerability knowledge by further training models on vulnerability datasets. For example, EPVD [63] and SVuLD [43] achieve state-of-the-art performance for vulnerability detection. Although achieving promising performance, the effectiveness of the pre-training techniques is still limited, embodied in two aspects:

(1) *The pre-trained models tend to fail to infer the semantics of the code statements with complex logic such as those containing multiple operators and pointers.* The previous studies [46, 56] have shown that the “reasoning gap” in pre-trained models poses a significant

```

1  int net_init_tap(const Netdev *netdev, ...) {
...
16  if (tap->has_fd) {
17    if (tap->has_ifname || tap->has_script || tap-
>has_downscript || tap->has_vnet_hdr ||...) {
18      error_setg(errp, "ifname=, ...");
19      return -1; }
20    fd = monitor_fd_param(cur_mon, tap->fd, &err);
21    if (fd == -1) {
22      error_propagate(errp, err);
23      return -1; }
...
27    if (err) {
28      error_propagate(errp, err);
29      return -1; }
30  } else if (tap->has_fds) {
31-    char **fds = g_new0(char *, MAX_TAP_QUEUES);
32-    char **vhost_fds = g_new0(char *, MAX_TAP_QUEUES);
+    char **fds;
+    char **vhost_fds;
33    int nfds, nvhosts;
34    if (tap->has_ifname || tap->has_script || tap-
>has_downscript...) {
...
36      return -1; }
+    fds = g_new0(char *, MAX_TAP_QUEUES);
+    vhost_fds = g_new0(char *, MAX_TAP_QUEUES)
...
39    if (tap->has_vhostfds) {... }

```

Figure 1: A code example that presents a memory leak vulnerability in the Qemu project, and is misclassified by UniX-coder as non-vulnerable. The red and green lines represent the patched code segments before and after fixed, respectively.

challenge. It means that these models struggle to infer the logic of the source code, such as containing multiple operators and pointers, where a single inference error may result in the opposite prediction. Different from natural language, due to the strict rules and syntax, source code may need multiple abstract symbols in a statement to express a simple intent. For instance, the code snippet shown in Figure 1 from the Qemu [4] project potentially causes a memory leak vulnerability in the `net_init_tap()` function. In Lines 31-32, the code snippet allocates memory for the arrays of character pointers `fds` and `vhost_fds`. However, if any of these error checks failed after the allocation in Lines 34-36, the function would exit without properly deallocating these arrays, leading to a memory leak vulnerability. In this example, the complex logic of the source code in Line 34, is characterized by the use of multiple operators (`||` and `->`), leading to erroneous predictions. It presents challenges for pre-trained models in detecting such vulnerability patterns.

(2) *The pre-trained models are hard to capture the code execution sequences.* Since pre-trained models primarily utilize code sequences as input, they are hard to comprehend the inherently non-sequential and various code execution sequences of source code, such as commonly-used `if` and `return` statements. For instance, the example shown in Figure 1 contains several `if` statements and `return` statements, in which the specific use of `if` in line 34 and `return` in line 36 directly causes the vulnerability. These statements necessitate multiple conditional `if` statements to trigger a `return -1`. It is challenging for the pre-trained models to predict the vulnerability without well capturing code execution sequences.

To address the limitations above, in this paper, we propose a Structured natural language Comment tree-based vulnAbiLity

detection framework based on the pre-trained models, named **SCALE**. The proposed Structured Natural Language Comment Tree (SCT) integrates the semantics of code statements with code execution sequences based on the Abstract Syntax Trees (ASTs). SCALE has three main modules: (1) *Comment Tree Construction*. SCALE makes the first attempt to incorporate the Large Language Model (LLM)’s general knowledge for comment generation in the vulnerability detection task and further generates the comment tree, which enhances the model’s ability to infer the semantics of code statements. (2) *Structured Natural Language Comment Tree Construction*, which explicitly involves code execution sequence by combining the code syntax templates with the comment tree. (3) *SCT-Enhanced Representation*, which incorporates the constructed SCTs for well capturing vulnerability patterns.

To evaluate SCALE, we use three widely-used datasets in vulnerability detection: FFMpeg+Qemu [65], Reveal [10], and SVuID [43]. We compare SCALE with eleven existing vulnerability detection methods. Experimental results demonstrate that SCALE outperforms the state-of-the-art baseline, with improvements of 2.96%, 13.47%, and 1.17% in terms of F1 score on the three datasets, respectively. Furthermore, SCALE can be applied to different pre-trained models, yielding F1 score performance enhancements ranging from 1.37% to 10.87%. These results underscore the effectiveness of SCALE in enhancing the detection of software vulnerabilities.

In summary, the major contributions of this paper are summarized as follows:

- (1) To the best of our knowledge, we propose a novel structured natural language comment tree, which enhances the model’s ability to infer code semantics by integrating code comments and code syntax templates.
- (2) We propose SCALE, a novel vulnerability detection framework for well capturing vulnerability patterns by enhancing code representation with the structured natural language comment tree.
- (3) We perform an extensive evaluation on three popular datasets, and the results demonstrate the effectiveness of SCALE in software vulnerability detection.

2 BACKGROUND AND RELATED WORK

We classify the existing vulnerability detection methods into four types: Program analysis-based, Supervised-based, Pre-trained model-based and LLM-based methods, as illustrated in Figure 2.

2.1 Program Analysis Methods

Numerous program analysis methodologies have been proposed and demonstrated their effectiveness in vulnerability detection, such as CheckMarx [31], FlawFinder [60], PCA [34] and RATs [6]. These methodologies typically employ pre-defined rules or patterns to identify improper operations within source code. Figure 2a illustrates an example of the rules proposed by Saber [50], which use the symbolic rules [7, 32] to simulate programming logic in source code. However, the creation of well-defined vulnerability rules or patterns heavily relies on expert knowledge [35, 36], making it challenging to cover all potential cases. Furthermore, the complex programming logic inherent in real-world software projects hinders the

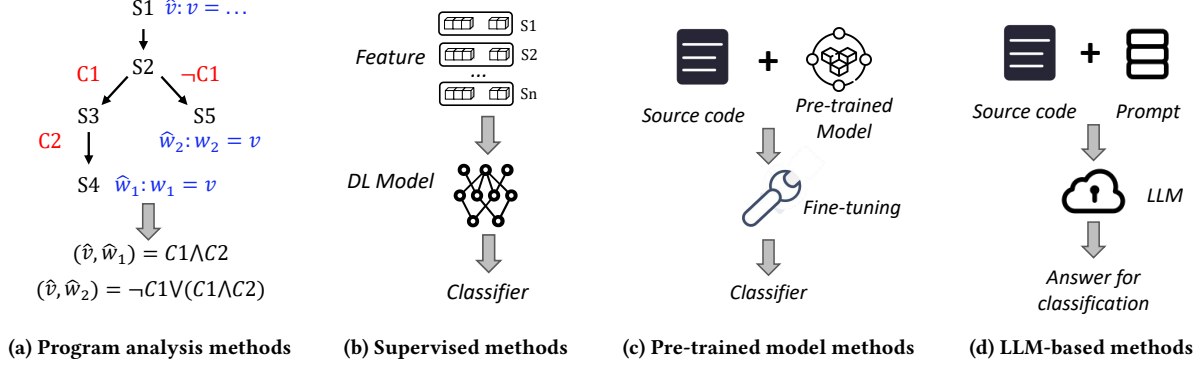


Figure 2: Four types of existing vulnerability detection methods.

manual identification of these rules, undermining the performance of traditional program analysis-based approaches.

2.2 Supervised Methods

With the advent of Deep Learning (DL), numerous supervised methods have been proposed that utilize DL-based models to capture program semantics for the identification of potential software vulnerabilities, which mainly include the sequence-based [38, 39] and graph-based [10, 37, 65] approaches. Figure 2b illustrates the process of these methods. They typically extract features from source code and adopt the DL models for classification. For instance, VulDeePecker [39] and SySeVR [38] use the bidirectional Long Short-Term Memory (LSTM) network for vulnerability detection and extract the code gadgets from the source code.

Then, graph-based methods have gained significant popularity due to their interpretability, which has demonstrated more effectiveness than sequence-based methods. They extract structured representations from the source code, such as AST, CFG, DFG, and Code Property Graphs (CPG) [54]. They then employ Graph Neural Networks (GNNs) [53] models to learn the graph representation for classification. In comparison to traditional program analysis-based approaches, these methods can automatically extract implicit vulnerability patterns from previously vulnerable code, eliminating the need for expert involvement. However, these methods are constrained by semantic embedding and overly complex graph structures.

2.3 Pre-trained Model Methods

While GNN-based methods have demonstrated satisfactory performance, they necessitate large and high-quality data for training [14], which can be challenging to procure in real-world scenarios [59]. An alternative approach involves the use of pre-trained models as the foundational backbone, which are then fine-tuned for the specific downstream task. Figure 2c illustrates the process of pre-trained model-based methods [20, 22, 26, 64]. These methods, which do not require expert involvement or the generation of structured graphs, utilize source code as input. They then simply fine-tune the pre-trained model and only train the classifier for vulnerability detection. All pre-trained code models with encoder components,

such as CodeBERT [18], can be utilized to detect vulnerabilities. CodeT5 [55] and UniXcoder [25] are designed to support both code-related classification and generation tasks. By leveraging the knowledge encapsulated in pre-trained models, these pre-trained model-based approaches achieve the best performance in vulnerability detection. EPVD [63] proposes an execution path selection algorithm and adopts a pre-trained model to learn the path representations. SVulD [43] constructs contrastive paired instances and uses the pre-trained model to learn distinguishing semantic representations.

However, the pre-trained methods generally employ code sequences as input during prediction. They may fail to infer the semantics of the code statement with complex logic and hard to capture the code execution sequence. In this work, we aim to mitigate the issues of pre-trained models for better learning the vulnerability-related structural information.

2.4 LLM-based Methods

In recent years, Large Language Models (LLMs) have gained prominence [8] and widespread adoption in the fields of Natural Language Processing (NLP) [40] and Software Engineering (SE) [28]. Figure 2d illustrates the process of these LLMs [11, 21]. Notable examples include the series of GPT models proposed by OpenAI, such as ChatGPT [11] and GPT-4 [44], as well as the LLaMA models introduced by Meta, including LLaMA [51] and LLaMA2 [52].

Furthermore, beyond these general-purpose LLMs, specific LLMs trained on code repositories have also achieved parameters in the billion-scale range, such as INCODER [19], StarCoder [33], and Code Llama [47]. For example, Code Llama [47], with its substantial 34-billion-parameter model, excels in code generation and completion tasks. These models have consistently demonstrated commendable performance across a spectrum of code-related intelligence tasks [23].

However, these LLMs face significant challenges when applied to software vulnerability detection [21]. It is mainly caused by individual code snippets often containing numerous undefined identifiers and LLMs often lack domain knowledge for vulnerability detection.

3 PROPOSED FRAMEWORK

In this section, we elaborate on the overall architecture of SCALE. As shown in Figure 3, SCALE mainly consists of three modules: comment tree construction, structured natural language comment tree construction, and SCT-enhanced representation, with details as below.

3.1 Comment Tree Construction

The purpose of the comment tree construction phase is to enrich corresponding code snippets with natural language descriptions. It can relieve the difficulties in understanding the semantics of code statements with multiple operators and pointers, and thus help the model infer the code logic more effectively. Inspired by the remarkable achievement of LLMs, we seek to utilize their extensive knowledge base to bridge the reasoning gap in vulnerability detection. Specifically, we leverage ChatGPT to generate comments for each code snippet. Following OpenAI’s gpt-best-practices document [3], we employ the role-based basic prompt to remind ChatGPT [11] of its job (i.e., generate comments) for each provided code snippet and use the “gpt-3.5-turbo-0301” API for the comment generation.

To avoid LLMs from generating unnormalized code, we then normalize the comments to facilitate the subsequent program analysis part. Specifically, we employ some normalization operations including removing all blank lines and replacing multi-line comments enclosed in the form of triple single quotation marks (‘’) to double slashes (//). SCALE then generates Abstract Syntax Trees (ASTs) based on source code by Tree-sitter [5] library for constructing the comment tree. A generated AST can be formulated as a directed graph (N, E) , where N denotes the set of nodes and E represents the set of edges. Each node $n \in N$ within the AST is a representation of a pair (V, T) , with $v \in V$ representing the node’s value and $t \in T$ representing its type. For adding the comment node, we find the first node on the subsequent row and designate this as the target node for the comment. The comment node is then integrated into the AST by taking the parent of the target node as its parent and positioning it as the parent of the target node. For example, as shown in Figure 4 and 5a, we generate the comment for Line 3 and traverse the AST of the source code, where the first node in Line 3 is a “if_statement”. We add the comment node B as a previous sibling node of A and integrate it into AST. This approach ensures that the comments are logically and structurally aligned with the following structured natural language rules.

3.2 Structured Natural Language Comment Tree Construction

In this section, we first present our structured natural language rules for structured natural language comment synthesis. Then we describe our proposed algorithm for building SCTs. The SCT integrates code comments and code syntax templates, and provides rich code execution information for inferring vulnerability patterns.

3.2.1 Structured Natural Language Rules. To integrate execution sequences with the natural language for more effectively identifying vulnerability patterns, we introduce a set of structured natural language rules. Following the guidelines of C++ language reference [1], we select four distinct categories of statements that control how

Table 1: The symbolic rules for constructing the SCT. “[]” indicates the node value to replace and “_” denotes the original AST node type parsed by the Tree-sitter.

Categories	Target Nodes	Symbolic Templates
Selection	If	<i>if</i> ([<i>condition</i>]) <i>if-branch</i>
	If-else	<i>if</i> ([<i>condition</i>]) <i>if-branch</i> <i>else</i> [] <i>else-branch</i>
		<i>switch</i> ([<i>condition</i>]) <i>statement</i>
	Switch	
Iteration	While	<i>while</i> ([<i>expression</i>]) <i>statement</i>
	For	<i>for</i> ([<i>init-expression</i> ; <i>condition-expression</i> ; <i>loop-expression</i>]) <i>statement</i>
	Range-based for	<i>for</i> ([<i>for-range-declaration</i> : <i>expression</i>]) <i>statement</i>
Jump	Break	[<i>break</i>];
	Continue	[<i>continue</i>];
	Return	<i>return</i> [<i>expression</i>] ;
	Goto	<i>goto</i> [<i>identifier</i>] ;
Labeled	Case	<i>case</i> [<i>constant-expression</i>]: <i>statement</i>

and in what order programs are manipulated and create structured natural language rules for them, including selection statements, iteration statements, jump statements, and labeled statements. The eleven structured natural language rules for these four categories of statements are shown in Table 1. A structured natural language rule represents a type of statement and corresponding target node types.

Selection Statements: The selection statements encompass “if”, “if-else”, and “switch” statements, which provide a means to conditionally execute sections of code. To construct the structured natural language rules, SCALE starts with finding the target node (e.g., *if* and *switch*) and corresponding comments description. Then it traverses forward on the target node’s child and sibling nodes to confirm the condition branching expression. Finally, these comment descriptions replace the content by utilizing corresponding structured natural language templates. For the example in Figure 5b, SCALE identifies the target node A *if_statement* (highlighted in blue) and extracts the corresponding comment node B (highlighted in grey). It then replaces the child node’s C (highlighted in red) value with the comment node B (highlighted in grey).

Iteration Statements: The iteration statements enable the execution of statements zero or more times, subject to specific termination criteria. We hope to employ natural language rules to

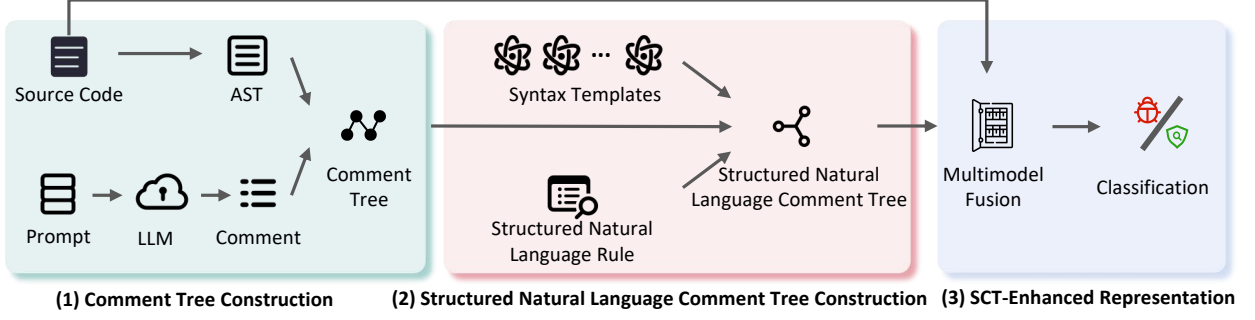


Figure 3: The overview of SCALE.

```

1 int ff_alloc_entries(AVCodecContext *avctx, int count) {...
2   if (avctx->active_thread_type & FF_THREAD_SLICE) {
3-    if (!p->entries) {
4+      if (Check if allocation was successful) {
5-        return AVERROR(ENOMEM); }
6+      return Return an error code if allocation failed;
7-      // Assign count to entries_count of p
8-      p->entries_count = count;
9-      ...
10+     for ([Loop through thread_count number of times]) {
11-       pthread_mutex_init(&p->progress_mutex[i], NULL);
12-       ...}
13 return 0; }

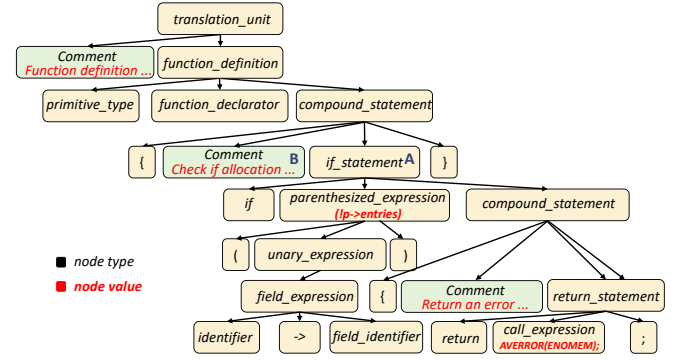
```

Figure 4: An example of structured natural language comment in SCALE. The red and grey lines denote the original source code and structured natural language comment, respectively.

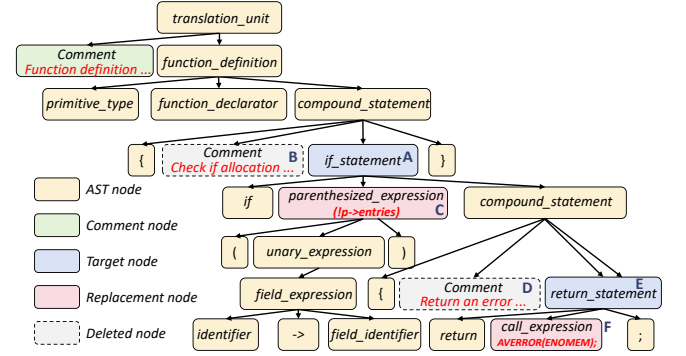
facilitate the comprehension of loop logic and triggering conditions by models.

We design the three structured natural language rules in the second part of Table 1. For the “while” statement, SCALE starts with the target node type of *while*. Then it traverses forward on the target node’s child or sibling nodes to substitute the expression with the corresponding comment. The process of constructing structured natural language rules for the “for” and the “range-based-for” statements closely mirrors that of the “while” statement. For the “for” statement, the replacement encompasses init-expression, condition-expression, and loop-expression, wherein init-expression and loop-expression can contain multiple separated statements. As shown in Figure 4, Line 9 (highlighted in grey) indicates the source code of the “for” statement, and Line 10 (highlighted in red) indicates the generated structured natural language comments.

Jump Statements: The jump statements perform an immediate local transfer of control. These statements share common semantics, such as the “break” statement, thereby altering the program’s execution flow. Despite sharing the same tokens, different “return” statements possess different implicit targets, often rendering them challenging to discern. Consequently, the structured natural language rules for jump statements can be constructed as



(a) An example of the comment tree.



(b) An example of the SCT.

Figure 5: The example of the comment tree and SCT. The black and red font denotes the node type and value, respectively. The yellow, green, blue, red, and gray-shaded nodes denote the original nodes, comment nodes, target nodes, nodes to replace, and deleted nodes, respectively.

intermediaries between the jump statement and the intended inference target, aiding in indicating the code execution sequence. We introduce four categories and the structured natural language templates can be found in the third section of Table 1. For the first

Algorithm 1: Structured Natural Language Comment Tree Construction

Input : The Comment Tree of given function: *func_ct*
Output : Structured Natural Language Comment Tree of the given function, *func_sct*
Initialize: Initialize an action list *ac_list*

```

1 Function Structured Natural Language Comment Tree Construction:
  // Using structured natural language rules to construct
  structured natural language comment tree
2 if node.type  $\in$  ac_list and node.comment exist then
3   if node.type  $\in$  Table 1 and node.parent.type == "comment" then
4     // Processing for child nodes
5     for all child_node  $\in$  node.children is not visited do
6       if child_node.type == "parenthesized_expression" then
7         st_comment = node.comment
8         delete child_node.value
9         child_node.value = st_comment
10        update(func_ct)
11        break
12      end
13    end
14    // Processing for next sibling node
15    if next_sibling_node exists and sibling_node.type ==
    "parenthesized_expression" or "compound_expression" then
16      st_comment = node.comment
17      delete child_node.value
18      child_node.value = st_comment
19      update(func_ct)
20      break
21    end
22  end
23  func_sct = func_ct
24 return func_sct

```

two structured natural language rules of the jump statements, we find the target nodes of *break* and *continue*. Then we combine the initial node value with the corresponding comments to produce the node value. For the *return* and *goto* target nodes, we collect the child node value of *identifier* and *expression*, respectively, where we use the comments to replace the child node values.

In Figure 5b, nodes D to F illustrate an instance of a “return” statement. SCALE identifies the target node E (highlighted in blue) and substitutes the value of node F (highlighted in red) with that of node D (highlighted in grey). This process enriches the “return” statement with relevant in-context information, thereby enhancing the model’s comprehension of the code execution sequence. The generated structured natural language comment is exemplified in lines 5-6 of Figure 4.

Labeled Statements: The labeled statements facilitate the direct transfer of program control to a predefined statement. To explain the “case” statement, we also formulate a structured natural language rule. We capture both the target type node *case* and its next sibling node from the comment tree. Subsequently, we employ comments to substitute the value of the *constant – expression*’s node.

3.2.2 SCT Construction Algorithm. Based on the above structured natural language rules, we then introduce the construction process of SCTs. The overall SCT construction process is illustrated in Algorithm 1, with an example shown Figure 5b.

SCALE first traverses the comment tree and constructs the SCT using the above rules. During this traversal, SCALE systematically

identifies and records each visited node including statement, expression, and identifier nodes. When SCALE visits a node that conforms to a structured natural language rule, it captures the node’s values and incorporates the relevant comment following the structured natural language comments rules. According to the parsing rules [5], it involves two conditions: (1) If the target node occurs in the child node, SCALE will traverse all child nodes to add structured natural language comments (Lines 4-12). (2) If the target node occurs in the sibling node, SCALE will check its parent node and only confirm the next sibling node (Lines 13-19). Node values meeting these conditions are then updated. This process will be completed until the comment tree is traversed and return the SCT.

Finally, SCALE flattens the current SCT and generates the structured natural language comment through the AST flatten operator [5]. Figure 4 shows an example of the flattened SCT ultimately obtained for model training.

3.3 SCT-Enhanced Representation

The SCT-enhanced representation module aims to incorporate the constructed SCTs for well-capturing vulnerability patterns. We first utilize UniXcoder as the encoder for source code and structured natural language comments, as UniXcoder [25] is a unified cross-modal (i.e., code, comment, and AST) pre-trained model. SCALE feeds the source code and SCT into the pre-trained model and obtains their representations $h_c \in \mathbb{R}^{l \times n}$ and $h_{ct} \in \mathbb{R}^{l \times n}$, where l and n denote the length of input and the dimension of embedding, respectively. We then propose to fuse the representation of source code and structured natural language comment through Cross-Attention [12]. Specifically, SCALE first maps the representations into:

$$Q = W_Q^{(i)} \cdot \varphi(h_{ct}), K = W_K^{(i)} \cdot \varphi(h_c), V = W_V^{(i)} \cdot \varphi(h_c) \quad (1)$$

where $W_Q^{(i)}$, $W_K^{(i)}$ and $W_V^{(i)}$ are the linear projection of the i -th head. $\varphi(h_{ct})$ and $\varphi(h_c)$ are the intermediate layer’s representation of source code and SCT, respectively. SCALE integrates the information from source code and SCT, then obtains the SCT-enhanced representation as follows:

$$Attention(Q, K, V) = softmax\left(\frac{\|_1^H QK^T}{\sqrt{d}}\right) \cdot V \quad (2)$$

where T denotes the transpose operator and d denotes the embedding size. $\|_1$ and H are the concentrate operator and the head number, respectively.

Finally, SCALE uses the mean operator and leverage a classifier for software vulnerability detection, which can be formulated as follows:

$$M = \sigma(Classifier(Mean(Attention(Q, K, V)))) \quad (3)$$

where σ denotes the sigmoid function. And SCALE trains the classifier [27] by minimizing the Cross-Entropy [15] loss.

4 EXPERIMENTAL SETUP

4.1 Research Questions

In this section, we evaluate the effectiveness of SCALE by comparing it with the state-of-the-art vulnerability detection methods and focus on the following five Research Questions (RQs):

Table 2: Statistics of the datasets.

Dataset	# Total	# Vul.	# Non-vul.	Ratio(%)
FFMPeg+Qemu [65]	22,361	10,067	12,294	45.02
Reveal [10]	18,169	1,664	16,505	9.16
SVuLD [43]	28,730	5,260	23,470	18.31

- RQ1:** How effective is SCALE in vulnerability detection compared with existing approaches?
- RQ2:** How effective is SCALE when applied to other pre-trained models?
- RQ3:** How do different structured natural language comment rules contribute to the performance of SCALE?
- RQ4:** What is the influence of different modules on the detection performance of SCALE?
- RQ5:** What is the influence of hyper-parameters on the performance of SCALE?

4.2 Datasets

To answer the questions above, we choose three widely-used vulnerability datasets, including FFMPeg+Qemu [65], Reveal [10], and SVuLD [43]. The FFMPeg+Qemu dataset originates from two popular C open-source projects. It comprises 22k code snippets, of which 10k have been identified as vulnerable. The Reveal dataset has over 18k code snippets, with around 2k of these exhibiting vulnerabilities. SVuLD [43] is based on Fan et al. [17] and contains both before-fixed and after-fixed code in the training set. The model is required to identify the before-fixed as vulnerable and after-fixed as non-vulnerable simultaneously on this dataset. Table 2 presents the statistics of the experimental datasets.

Following the previous work [41, 43, 59, 65], we split the datasets into disjoint training, validation, and test sets in a ratio of 8:1:1. We use the training set to train the baseline models, use the validation set for selecting best-performance models, and evaluate the performance of SCALE and other baselines in the test set.

4.3 Baselines

To verify the effectiveness of SCALE, we choose the following three types of vulnerability detection approaches as our baselines:

- **Supervised-based methods:** We use Devign [65] and Reveal [10], which are widely adopted as baselines in recent works [9, 37, 59]. They construct the joint graph from the source code and use the GGNN model for function-level vulnerability detection.
- **Pretrained Model-based methods:** We select three popular pre-trained code models, including CodeBERT [18], CodeT5 [55], and UniXcoder [25] to detect vulnerabilities. Besides, we also select recent state-of-the-art methods that further finetune these pre-trained models for vulnerability detection, including EPVD [63], LineVul [20] and SVuLD [43].
- **LLM-based methods:** We also choose Code Llama-7b [47], ChatGPT [11] and GPT3.5-instruct [24] to evaluate the performance of LLMs for vulnerability detection. We follow

the prompts used in [21] to improve their performance on vulnerability detection.

4.4 Evaluation Metrics

Following the previous work [18, 20, 65], we choose the following four metrics to evaluate SCALE’s performance:

- **Accuracy (Acc):** $Acc = \frac{TP+TN}{TP+TN+FN+FP}$. Acc measures the percentage of correctly classified samples out of all samples. TN represents the number of true negatives and $TP + TN + FN + FP$ represents the number of all samples.
- **Precision (Pre):** $Pre = \frac{TP}{TP+FP}$. Pre is the percentage of true vulnerabilities out of all the vulnerabilities that are retrieved. TP and FP denote the number of true positives and false positives, respectively.
- **Recall (Rec):** $Rec = \frac{TP}{TP+FN}$. Rec is the percentage of vulnerabilities that are detected out of all vulnerable code snippets. TP and FN denote the number of true positives and false negatives, respectively.
- **F1 score (F1):** $F1 = 2 \times \frac{Pre \times Rec}{Pre + Rec}$. F1 measures the harmonic mean of Pre and Rec metrics.

4.5 Implementation Details

For all the supervised-based and pre-trained model-based methods except Devign, we directly use the publicly available source code and hyper-parameters released by the authors. For Devign, following the previous work [57, 58], we try our best to reproduce it based on code reproduced by other researchers [10]. For LLM-based methods, we download the Code-llama-7b from HuggingFace [29] and evaluate them on our server. For ChatGPT and GPT-3.5-instruct, we use the OpenAI’s public API “gpt-3.5-turbo-0301” and “gpt-3.5-turbo-instruct” for experiments. And we use the initial parameter provided by OPENAI.

To ensure the fairness of the experiments, we use the same data splitting for all the approaches in all research questions. We use the “gpt-3.5-turbo-0301” API for the comment generation and fine-tune the pre-trained model UniXcoder with a learning rate of $2e-5$. The cross attention head number H is set as 8 and the batch size is set to 32. The influence of these parameters is discussed in Section 5.5. All experiments are conducted on a server with NVIDIA A100-SXM4-40GB GPUs.

5 EXPERIMENTAL RESULTS

5.1 RQ1: Effectiveness of SCALE

To answer RQ1, we compare the SCALE with three types of vulnerability detection baselines. The results are shown in Table 3.

5.1.1 Comparison with Supervised and Pre-trained Model-based Methods. Specifically, we compare SCALE with two supervised-based approaches (i.e., Devign and Reveal) and six pre-trained model-based approaches (i.e., CodeBERT, CodeT5, UniXcoder, EPVD, LineVul, and SVuLD). From the results in Table 3, we observe that SCALE outperforms all the supervised baseline methods on the three datasets in terms of F1 score, by 2.96% for FFMPeg+Qemu, 13.47% for Reveal and 1.17% for SVuLD dataset, respectively. When considering all the four performance metrics in the three datasets

Table 3: Evaluation results of SCALE compared with vulnerability detection baselines on the three datasets. The shaded cells represent the performance of the best methods in each metric. Bold text cells represent the best performance. The results of statistical significance tests are listed in <https://github.com/Xin-Cheng-Wen/Comment4Vul>.

Dataset	FFMPeg+Qemu				Reveal				SVulD			
Metrics	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1
Devign	56.89	52.50	64.67	57.95	87.49	31.55	36.65	33.91	73.57	9.72	50.31	16.29
Reveal	61.07	55.50	70.70	62.19	81.77	31.55	61.14	41.62	82.58	12.92	40.08	19.31
CodeBERT	62.37	61.55	48.21	54.07	87.51	43.63	56.15	49.10	80.56	14.33	55.32	22.76
CodeT5	63.36	58.65	68.61	63.24	89.53	51.15	54.51	52.78	78.73	14.32	62.36	23.30
UnixCoder	65.19	59.93	59.98	59.96	88.48	47.44	68.44	56.04	77.54	15.11	72.24	24.99
EPVD	63.03	59.32	62.15	60.70	88.87	48.60	63.93	55.22	76.75	14.26	69.58	23.67
LineVul	62.37	61.55	48.21	54.07	87.51	43.63	56.15	49.10	80.57	15.95	64.45	25.58
SVulD	-	-	-	-	-	-	-	-	86.99	21.46	56.84	31.16
Codellama-7B	53.44	46.20	6.37	11.20	83.49	8.23	5.44	6.55	88.40	3.84	5.65	4.57
ChatGPT	53.85	49.72	14.35	22.28	81.77	9.52	8.37	8.91	88.79	14.38	25.81	18.47
GPT-3.5-Instruct	51.31	47.03	42.06	44.41	60.47	9.60	32.22	14.79	65.81	7.26	50.4	12.69
SCALE	66.18	61.88	68.69	65.11	90.02	52.32	78.69	62.85	87.63	22.56	57.03	32.33

(12 combination cases altogether), SCALE achieves the best performance in 10 out of the 12 cases. These results demonstrate that SCALE captures the vulnerability patterns of code more precisely than the supervised and pre-trained model-based baselines. We also notice that the pre-trained model-based methods perform better than supervised-based methods. It can be attributed that pre-trained model-based methods have leveraged general code-related knowledge in the pre-training stage, which enhances the ability to capture vulnerability patterns.

Moreover, experimental results reveal that the detection accuracy on the SVulD dataset is lower compared to the other datasets. It can be attributed to the challenge that current vulnerability detection approaches struggle to differentiate between before- and after-fixed code which often exhibits only small variances. Despite these challenges, SCALE demonstrates improvement in the F1 score of 1.17%, surpassing the SVulD baseline which is specifically designed for this situation.

5.1.2 Comparison with LLM-based Methods. We also compare SCALE with multiple LLM-based approaches, including Code Llama-7b, ChatGPT and GPT-3.5-instruct. In general, we can find that the LLM-based approaches perform worse than SCALE and supervised approaches mentioned before. Specifically, SCALE outperforms the best LLM’s baselines by 9.63%, 170.38%, 83.83%, and 150.53% in terms of accuracy, precision, recall, and F1 score respectively. It indicates that despite the extensive model size and training data, LLM-based approaches still face challenges in vulnerability detection without fine-tuning due to the lack of domain knowledge.

Answer to RQ1: SCALE achieves the best performance in precision and F1 score compared with previous approaches. Specifically, SCALE outperforms the best-performing baseline by 2.96%, 13.47%, and 1.17% in terms of F1 score on the three datasets, respectively.

5.2 RQ2: Effectiveness on Other Pre-trained Models

In this research question, we wonder whether other pre-trained models can have the same effectiveness when equipped with our proposed SCT. To answer this question, we replace UniXcoder with different pre-trained model-based methods including CodeBERT and EPVD to validate the performance. The experimental results are presented in Table 4.

The experimental results show that the structured natural language comment is effective for fine-tuning existing pre-trained models in most cases. We find average improvements on different datasets by 6.67% in CodeBERT, 1.82% in EPVD and 4.88% in UniXcoder, respectively, in terms of F1 score. In particular, incorporating structured natural language comment enables CodeBERT to surpass the performance of the majority of existing baselines on the FFMPeg+Qemu dataset. This indicates that SCT provides the extra explanation in structured natural language ways for intricate logic and leverages the structured natural language rules for involving the code execution sequence, which can be easily learned by existing pre-trained model-based methods. As for EPVD, although it has used multiple syntax control flow paths to capture vulnerability patterns, the structured natural language comment still improves it by alleviating the heavy use of operators and pointers. It improves the

Table 4: The experimental results of applying SCT to the existing pre-trained model-based methods.

Dataset	FFMPeg+Qemu				Reveal				SVulD			
Metrics	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1
CodeBERT#	64.24 (↑+1.87%)	59.07 (↓-2.48%)	72.11 (↑+23.90%)	64.94 (↑+10.87%)	88.78 (↑+1.27%)	48.30 (↑+4.67%)	67.62 (↑+11.47%)	56.41 (↑+7.31%)	80.06 (↓-0.50%)	15.28 (↑+0.95%)	62.74 (↑+7.42%)	24.58 (↑+1.82%)
EPVD#	61.60 (↓-1.43%)	56.40 (↓-2.92%)	72.35 (↑+10.20%)	63.39 (↑+2.69%)	89.49 (↑+0.62%)	63.93 (↑+15.33%)	50.81 (↓-13.12%)	56.62 (↑+1.40%)	76.83 (↑+0.08%)	74.71 (↑+60.45%)	15.04 (↓-54.54%)	25.04 (↑+1.37%)
UniXcoder#	66.18 (↑+0.99%)	61.88 (↑+1.95%)	68.69 (↑+8.71%)	65.11 (↑+5.15%)	89.58 (↑+1.10%)	50.86 (↑+3.42%)	84.84 (↑+16.40%)	63.59 (↑+7.55%)	82.33 (↑+4.79%)	17.14 (↑+2.03%)	62.93 (↓-9.31%)	26.94 (↑+1.95%)

F1 score performance of 2.69%, 1.40%, and 1.37% in three datasets, respectively.

Answer to RQ2: SCT is effective for different pre-trained models, which averagely improves the F1 score by 6.67% in CodeBERT, 1.82% in EPVD, and 4.88% in UniXcoder, respectively.

5.3 RQ3: Effectiveness of Different Structured Natural Language Rules in SCALE

To investigate the effectiveness of each category of rule, we construct four variants of SCALE with four categories of structured natural language rules, including selection, iteration, jump, and labeled statements. All variants only use specific structured natural language rules to construct the SCT.

The performance of four variants and SCALE is presented in Table 5. The experimental results indicate that the four separate structured natural language rules perform worse compared with SCALE in most cases. Specifically, we observe that the variants decrease by 0.63%~1.06% in FFMPeg+Qemu and 0.44%~0.66% in Reveal, 1.41%~6.21% in SVulD, respectively, in terms of accuracy. Compared with other statements, the iteration statements and corresponding structured natural language rules generally contribute more, which outperforms other variants in 6 out of 12 cases. We suppose that it is mainly caused by the structured natural language rules of iteration statements replacing more contents (i.e., init-expression, condition-expression, and loop-expression) and the constructed structured natural language comments can explain the code more precisely.

Besides, we find that leveraging the integration of all structured natural language rules performs better than only using a single rule, which improves the F1 score performance in three datasets by 0.77%, 1.35%, and 1.63%, respectively. This indicates that the integration of all rules can further help the model understand the code’s execution sequence to improve the performance of vulnerability detection.

Answer to RQ3: All structured natural language rules contribute to the performance of SCALE, with an F1 score improvement of 1.41%~6.21%. Leveraging the integration of all

structured natural language rules performs better than using any single rule.

5.4 RQ4: Effectiveness of Different Modules in SCALE

In this section, we explore the impact of different modules of SCALE including Structured Natural Language Comment Tree Construction (SCTC) and SCT-Enhanced Representation (SER) module. The results are shown in Table 6.

5.4.1 Structured Natural Language Comment Tree Construction. To explore the effect of the SCTC module, we deploy one variant (i.e., Without SCTC) by only using the comment tree for SCT-enhanced representation. As shown in Table 6, the SCTC can improve the performance of SCALE on all datasets. Specifically, removing structured natural language comments leads to the F1 score drop of 1.29%, 2.59%, and 2.04% on different datasets. Especially on the imbalanced dataset (i.e., Reveal), SCTC boosts the performance more than on the balanced datasets (i.e., Devign and SVulD), which enhances the four parametric metrics in Reveal by 0.75% for accuracy, 2.32% for precision, 2.87% for recall, and 2.59% for F1 score, respectively. It indicates that structured natural language comment generation has a greater effect on the unbalanced dataset.

5.4.2 SCT-Enhanced Representation. To understand the impact of the SER module, we also deploy a variant (i.e., Without SER) of SCALE without the SER phase. In this variant, we only use structured natural language comments as input without the source code. The performance degradation observed across three datasets consistently demonstrate an average decrease of 3.42% in accuracy and 2.83% in F1 score, respectively. These results indicate the SCT-enhanced representation can better capture vulnerability patterns.

Answer to RQ4: Both the SCTC and SER modules enhance the performance of SCALE. The SCG phase improves F1 score performance on three datasets by 1.29%, 2.59%, and 2.04%, respectively. The SER module improves SCALE by 1.31%, 3.61% and 3.56%, respectively.

Table 5: The performance of SCALE with different structured natural language rules on three different datasets.

Dataset	FFMPeg+Qemu				Reveal				SVulD			
Metrics	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1
Selection	65.55	61.75	65.74	63.68	89.14	49.64	84.02	62.40	80.88	16.22	64.64	25.93
Iteration	65.12	60.23	70.84	65.10	89.18	49.76	84.01	62.50	80.92	16.35	65.21	26.14
Jump	65.45	61.21	67.65	64.27	88.92	49.04	83.61	61.82	80.75	16.22	65.21	25.98
Labeled	65.37	61.06	67.97	64.33	89.23	49.88	82.79	62.25	76.12	13.91	69.58	23.19
SCALE	66.18	61.88	68.69	65.11	89.58	50.86	84.84	63.59	82.33	17.14	62.93	26.94

Table 6: The performance of SCALE in three different datasets when removing the structured natural language comment tree construction module (i.e., Without SCTC) and removing the SCT-enhanced representation module (i.e., Without SER) in four metrics.

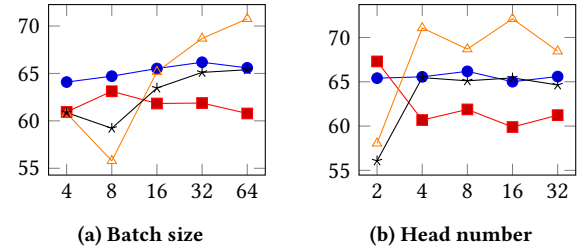
Dataset	FFMPeg+Qemu				Reveal				SVulD			
Metrics	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1	Acc	Pre	Rec	F1
Without SCTC	65.59 (↓0.59%)	61.73 (↓0.15%)	66.06 (↓2.63%)	63.82 (↓1.29%)	89.27 (↓0.75%)	50.00 (↓2.32%)	75.82 (↓2.87%)	60.26 (↓2.59%)	78.32 (↓4.01%)	15.17 (↓1.97%)	69.69 (↑6.76%)	24.99 (↓2.04%)
Without SER	64.28 (↓1.90%)	59.68 (↓2.20%)	68.53 (↓0.16%)	63.80 (↓1.31%)	88.96 (↓1.06%)	49.02 (↓3.30%)	71.72 (↓6.97%)	59.24 (↓3.61%)	75.02 (↓7.31%)	13.90 (↓3.24%)	73.57 (↑10.64%)	23.38 (↓3.56%)
SCALE	66.18	61.88	68.69	65.11	90.02	52.32	78.69	62.85	82.33	17.14	62.93	26.94

5.5 RQ5: Influences of Hyper-paramaters on SCALE

In this section, we study the effect of two hyper-parameters (i.e., batch sizes and head numbers) of SCALE. We assign different values to them and examine their impact on the performance of vulnerability detection. Due to the page limitation, we only present the experimental results of FFMPeg+Qemu in Figure 6.. Experimental results for other datasets are shown in the repository.

5.5.1 Batch Size. Figure 6a shows the performance of SCALE across four metrics with different batch sizes in FFMPeg+Qemu. Specifically, the performance of SCALE increases with the growth of batch size in general, which demonstrates the importance of batch size in improving model generalization. For batch sizes exceeding 32, we observe that the SCALE’s performance exhibits relative stability. Additionally, we observe that batch size seriously influences imbalanced datasets (i.e., Reveal), leading to performance fluctuation by 49.43% ~ 53.60% in precision and 61.07% ~ 88.52% in recall metrics. Conversely, the growth of performance in balanced datasets (i.e., FFMPeg+Qemu and SVulD) is more steady since the similar ratio between negative and positive samples makes the model learning more stable.

5.5.2 Head Number. To evaluate the impacts of the head numbers of SCALE, we vary it from 2 to 32 and present the performance on four metrics in Figure 6b. As can be seen, SCALE achieves the best accuracy performance when the number of attention heads is set as eight. The performance of different head numbers is similar for

**Figure 6: The impact of batch size and head number on the model performance in the FFMPeg+Qemu dataset. The blue, red, orange, and black lines denote the accuracy, precision, recall and F1 score metrics, respectively.**

all datasets. Therefore, we empirically use eight head numbers for all three datasets and achieve 65.11% in FFMPeg+Qemu, 63.59% in Reveal, and 26.94% in SVulD, respectively, in terms of the F1 score.

Answer to RQ5: The hyper-parameter settings can impact the performance of SCALE in the FFMPeg+Qemu, Reveal, and SVUID datasets. SCALE achieves the best performance with the batch size of 32 and head number of 8.

Table 7: The performance of SCALE in FFMpeg+Qemu [65] dataset when using GPT-3.5-Instruct in comment tree construction phase.

Metrics	Acc	Pre	Rec	F1
Comment with GPT3.5-Instruct	63.69	58.23	74.10	65.22
SCALE (Comment with ChatGPT)	66.18	61.88	68.69	65.11

6 DISCUSSION

6.1 Influence of LLM’s Generated Comment

To evaluate the efficacy of the SCT generated by SCALE, we further employ GPT-3.5-instruct to construct SCT. Due to constraints in resources, the experiments are limited to the FFMpeg+Qemu dataset in Table 7. The results indicate that both GPT3.5-instruct and ChatGPT are capable of producing effective comments for the construction of SCTs. Specifically, the SCTs created by GPT3.5-instruct demonstrated an enhancement in Precision and F1 score, showing improvement of 5.49% and 1.98%, respectively, over the best-performing baseline CodeT5, as shown in Table 3. In a comparative analysis with SCTs generated by ChatGPT, the GPT3.5-instruct’s results are comparably effective. They exhibit superiority in recall and F1 score, whereas ChatGPT outperforms in accuracy and precision metrics.

These findings suggest that employing different large language models (LLMs) for comment generation on source code aids in improving the pre-trained model’s comprehension, thereby enhancing the effectiveness of vulnerability detection.

6.2 Threat to Validity

Generalizability on Other Programming Languages. In this paper, we employ the tree-sitter library to parse ASTs for C/C++ programming languages. As a result, our experimental analysis focuses solely on C/C++ datasets, excluding other popular languages such as Java and Python. In future research, we intend to evaluate the efficacy of SCALE in the context of a broader range of programming languages.

Constraints of Domain Knowledge in Structured Natural Language Rules. For the context of C/C++ language reference, SCALE offers eleven structured natural language rules to facilitate the correlation between code and its associated comments. Although it is comprehensive enough in most cases, it still can not cover some statements, such as expression statements. This limitation could potentially introduce biases into the model’s predictions. A potential resolution to this issue could involve more types of AST nodes to broaden the structured natural language rules, thereby enabling a more thorough integration of comment information.

Experiments on the larger pre-trained model. In this experiment, we evaluate SCALE on three pre-trained models. These models are all representative and have shown state-of-the-art performance on benchmarks. However, the size of these models is less than 1B. In the future, we plan to validate the effectiveness of SCALE on larger LLMs such as Code Llama [47].

7 CONCLUSION

In this paper, we propose SCALE, a structured natural language comment tree-based vulnerability detection framework based on the pre-trained models. It mainly comprises a comment tree construction for enhancing the model’s ability to infer the semantics of code statements, a structured natural language comment tree construction module for explicitly involving code execution sequence, and an SCT-enhanced representation for well-capturing vulnerability patterns. Compared with the state-of-the-art methods, the experimental results on three popular datasets validate the effectiveness of SCALE for vulnerability detection.

DATA AVAILABILITY

Our source code and experimental data are available at: <https://github.com/Xin-Cheng-Wen/Comment4Vul>.

REFERENCES

- [1] 2023. *C++ Language Reference*. <https://learn.microsoft.com/en-us/cpp/cpp/cpp-language-reference?view=msvc-170>
- [2] 2023. *CVE*. <https://cve.mitre.org/>
- [3] 2023. *OpenAI*. <https://platform.openai.com/docs/guides/gptbest-practices>
- [4] 2023. *Qemu*. <https://www.qemu.org/>
- [5] 2023. *Tree-sitter*. <https://tree-sitter.github.io/tree-sitter/>
- [6] [n.d.]. *Rough Audit Tool for Security*. <https://code.google.com/archive/p/rough-auditing-tool-for-security>
- [7] Roberto Baldoni, Emilio Coppa, Daniele Cono D’Elia, Camil Demetrescu, and Irene Finocchi. 2018. A Survey of Symbolic Execution Techniques. *ACM Comput. Surv.* 51, 3 (2018), 50:1–50:39.
- [8] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language Models are Few-Shot Learners. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*.
- [9] Sicong Cao, Xiaobing Sun, Lili Bo, Ying Wei, and Bin Li. 2021. *BGN4VD: Constructing Bidirectional Graph Neural-Network for Vulnerability Detection*. *Inf. Softw. Technol.* 136 (2021), 106576.
- [10] Saikat Chakraborty, Rahul Krishna, Yangruibo Ding, and Baishakhi Ray. 2020. Deep Learning based Vulnerability Detection: Are We There Yet? *CoRR abs/2009.07235* (2020).
- [11] ChatGPT. 2022. ChatGPT. <https://chat.openai.com/>
- [12] Chun-Fu (Richard) Chen, Quanfu Fan, and Rameswar Panda. 2021. CrossViT: Cross-Attention Multi-Scale Vision Transformer for Image Classification. In *2021 IEEE/CVF International Conference on Computer Vision, ICCV 2021, Montreal, QC, Canada, October 10-17, 2021*. IEEE, 347–356.
- [13] Xiao Cheng, Xu Nie, Ningke Li, Haoyu Wang, and Zheng Zheng, and Yulei Sui. 2022. How About Bug-Triggering Paths? Understanding and Characterizing Learning-Based Vulnerability Detectors. *IEEE Transactions on Dependable and Secure Computing*.
- [14] Roland Croft, Muhammad Ali Babar, and M. Mehdi Kholoosi. 2023. Data Quality for Software Vulnerability Datasets. *CoRR abs/2301.05456* (2023).
- [15] Pieter-Tjerk de Boer, Dirk P. Kroese, Shie Mannor, and Reuven Y. Rubinstein. 2005. A Tutorial on the Cross-Entropy Method. *Ann. Oper. Res.* 134, 1 (2005), 19–67.
- [16] Facebook. [n.d.]. *Infer*. <https://fbinfer.com/>
- [17] Jiahao Fan, Yi Li, Shaohua Wang, and Tien N. Nguyen. 2020. A C/C++ Code Vulnerability Dataset with Code Changes and CVE Summaries. In *MSR ’20: 17th International Conference on Mining Software Repositories, Seoul, Republic of Korea, 29-30 June, 2020*, Sunghun Kim, Georgios Gousios, Sarah Nadi, and Joseph Hejderup (Eds.). ACM, 508–512.
- [18] Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xiaocheng Feng, Ming Gong, Linjun Shou, Bing Qin, Ting Liu, Daxin Jiang, and Ming Zhou. 2020. CodeBERT: A Pre-Trained Model for Programming and Natural Languages. In *Findings of the Association for Computational Linguistics: EMNLP 2020, Online Event, 16-20 November 2020 (Findings of ACL, Vol. EMNLP 2020)*, Trevor Cohn, Yulan He, and Yang Liu (Eds.). Association for Computational Linguistics, 1536–1547.
- [19] Daniel Fried, Armen Aghajanyan, Jessy Lin, Sida Wang, Eric Wallace, Freda Shi, Ruiqi Zhong, Wen-tau Yih, Luke Zettlemoyer, and Mike Lewis. 2022. InCoder: A Generative Model for Code Infilling and Synthesis. *CoRR abs/2204.05999* (2022).
- [20] Michael Fu and Chakkrit Tantithamthavorn. 2022. LineVul: A Transformer-based Line-Level Vulnerability Prediction. In *MSR. ACM*, 608–620.

- [21] Michael Fu, Chakkrit Tantithamthavorn, Van Nguyen, and Trung Le. 2023. Chat-GPT for Vulnerability Detection, Classification, and Repair: How Far Are We? *CoRR* abs/2310.09810 (2023).
- [22] Shuzheng Gao, Wenxin Mao, Cuiyun Gao, Li Li, Xing Hu, Xin Xia, and Michael R. Lyu. 2024. Learning in the Wild: Towards Leveraging Unlabeled Data for Effectively Tuning Pre-trained Code Models. *CoRR* abs/2401.01060 (2024).
- [23] Shuzheng Gao, Xin-Cheng Wen, Cuiyun Gao, Wenxuan Wang, Hongyu Zhang, and Michael R. Lyu. 2023. What Makes Good In-Context Demonstrations for Code Intelligence Tasks with LLMs?. In *38th IEEE/ACM International Conference on Automated Software Engineering, ASE 2023, Luxembourg, September 11–15, 2023*. IEEE, 761–773.
- [24] GPT3.5-instruct. 2022. GPT3.5-instruct. <https://platform.openai.com/docs/models>.
- [25] Daya Guo, Shuai Lu, Nan Duan, Yanlin Wang, Ming Zhou, and Jian Yin. 2022. UniXcoder: Unified Cross-Modal Pre-training for Code Representation. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2022, Dublin, Ireland, May 22–27, 2022*, Smaranda Muresan, Preslav Nakov, and Aline Villavicencio (Eds.). Association for Computational Linguistics, 7212–7225.
- [26] Daya Guo, Shuo Ren, Shuai Lu, Zhangyin Feng, Duyu Tang, Shujie Liu, Long Zhou, Nan Duan, Alexey Svyatkovskiy, Shengyu Fu, Michele Tufano, Shao Kun Deng, Colin B. Clement, Dawn Drain, Neel Sundaresan, Jian Yin, Daxin Jiang, and Ming Zhou. 2021. GraphCodeBERT: Pre-training Code Representations with Data Flow. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3–7, 2021*. OpenReview.net.
- [27] Tong He, Weilin Huang, Yu Qiao, and Jian Yao. 2016. Text-Attentional Convolutional Neural Network for Scene Text Detection. *IEEE Trans. Image Process.* 25, 6 (2016), 2529–2541.
- [28] Xinyi Hou, Yanjie Zhao, Yue Liu, Zhou Yang, Kailong Wang, Li Li, Xiapu Luo, David Lo, John C. Grundy, and Haoyu Wang. 2023. Large Language Models for Software Engineering: A Systematic Literature Review. *CoRR* abs/2308.10620 (2023).
- [29] Huggingface hub. 2023. HuggingFace. <https://huggingface.com/>.
- [30] IBM. 2023. Cost of a Data Breach Report 2023. <https://www.ibm.com/reports/data-breach>
- [31] Israel. [n.d.]. Checkmarx. <https://www.checkmarx.com/>.
- [32] Hongzhe Li, Taebeom Kim, Munkhbayer Bat-Erdene, and Heejo Lee. 2013. Software Vulnerability Detection Using Backward Trace Analysis and Symbolic Execution. In *2013 International Conference on Availability, Reliability and Security, ARES 2013, Regensburg, Germany, September 2–6, 2013*. IEEE Computer Society, 446–454.
- [33] Raymond Li, Loubna Ben Allal, Yangtian Zi, Niklas Muennighoff, Denis Kocetkov, Chenghao Mou, Marc Marone, Christopher Akiki, Jia Li, Jenny Chim, Qian Liu, Evgenii Zheltonozhskii, Terry Yue Zhuo, Thomas Wang, Olivier Dehaene, Mishig Davaadorj, Joel Lamy-Poirier, João Monteiro, Oleh Shliazhko, Nicolas Gontier, Nicholas Meade, Armel Zebaze, Ming-Ho Yee, Logesh Kumar Umapathi, Jian Zhu, Benjamin Lipkin, Muhtasham Oblokulov, Zhiruo Wang, Rudra Murthy V, Jason Stillerman, Siva Sankalp Patel, Dmitry Abulkhanov, Marco Zocca, Manan Dey, Zhihan Zhang, Nour Moustafa-Fahmy, Urvashi Bhattacharyya, Wenhao Yu, Swayam Singh, Sasha Luccioni, Paulo Villegas, Maxim Kunakov, Fedor Zhdanov, Manuel Romero, Tony Lee, Nadav Timor, Jennifer Ding, Claire Schlesinger, Hailey Schoelkopf, Jan Ebert, Tri Dao, Mayank Mishra, Alex Gu, Jennifer Robinson, Carolyn Jane Anderson, Brendan Dolan-Gavitt, Danish Contractor, Siva Reddy, Daniel Fried, Dmitry Bahdanau, Yacine Jernite, Carlos Muñoz Ferrandis, Sean Hughes, Thomas Wolf, Arjun Guha, Leandro von Werra, and Harm de Vries. 2023. StarCoder: may the source be with you! *CoRR* abs/2305.06161 (2023).
- [34] Wen Li, Haipeng Cai, Yulei Sui, and David Manz. 2020. PCA: memory leak detection using partial call-path analysis. In *ESEC/FSE '20: 28th ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering, Virtual Event, USA, November 8–13, 2020*, Prem Devanbu, Myra B. Cohen, and Thomas Zimmermann (Eds.). ACM, 1621–1625.
- [35] Yue Li, Tian Tan, Anders Møller, and Yannis Smaragdakis. 2018. Precision-guided context sensitivity for pointer analysis. *Proc. ACM Program. Lang.* 2, OOPSLA (2018), 141:1–141:29.
- [36] Yue Li, Tian Tan, Anders Møller, and Yannis Smaragdakis. 2020. A Principled Approach to Selective Context Sensitivity for Pointer Analysis. *ACM Trans. Program. Lang. Syst.* 42, 2 (2020), 10:1–10:40.
- [37] Yi Li, Shaohua Wang, and Tien N. Nguyen. 2021. Vulnerability detection with fine-grained interpretations. In *ESEC/SIGSOFT FSE*. ACM, 292–303.
- [38] Zhen Li, Deqing Zou, Shouhuai Xu, Hai Jin, Yawei Zhu, and Zhaoxuan Chen. 2022. SySeVR: A Framework for Using Deep Learning to Detect Software Vulnerabilities. *IEEE Trans. Dependable Secur. Comput.* 19, 4 (2022), 2244–2258.
- [39] Zhen Li, Deqing Zou, Shouhuai Xu, Xinyu Ou, Hai Jin, Sujuan Wang, Zhijun Deng, and Yuyi Zhong. 2018. VulDeePecker: A Deep Learning-Based System for Vulnerability Detection. In *25th Annual Network and Distributed System Security Symposium, NDSS 2018, San Diego, California, USA, February 18–21, 2018*. The Internet Society.
- [40] Jiachang Liu, Dinghan Shen, Yizhe Zhang, Bill Dolan, Lawrence Carin, and Weizhu Chen. 2022. What Makes Good In-Context Examples for GPT-3?. In *Proceedings of Deep Learning Inside Out: The 3rd Workshop on Knowledge Extraction and Integration for Deep Learning Architectures, DeeLIO@ACL 2022, Dublin, Ireland and Online, May 27, 2022*. Association for Computational Linguistics, 100–114.
- [41] Shuai Lu, Daya Guo, Shuo Ren, Junjie Huang, Alexey Svyatkovskiy, Ambrosio Blanco, Colin B. Clement, Dawn Drain, Daxin Jiang, Duyu Tang, Ge Li, Lidong Zhou, Linjun Shou, Long Zhou, Michele Tufano, Ming Gong, Ming Zhou, Nan Duan, Neel Sundaresan, Shao Kun Deng, Shengyu Fu, and Shujie Liu. 2021. CodeXGLUE: A Machine Learning Benchmark Dataset for Code Understanding and Generation. In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021*, Joaquin Vanschoren and Sai-Kit Yeung (Eds.).
- [42] Gary McGraw and Bruce Potter. 2004. Software Security Testing. *IEEE Secur. Priv.* 2, 5 (2004), 81–85.
- [43] Chao Ni, Xin Yin, Kaiwen Yang, Dehai Zhao, Zhenchang Xing, and Xin Xia. 2023. Distinguishing Look-Alike Innocent and Vulnerable Code by Subtle Semantic Representation Learning and Explanation. *CoRR* abs/2308.11237 (2023).
- [44] OpenAI. 2023. GPT-4 Technical Report. *CoRR* abs/2303.08774 (2023).
- [45] Yun Peng, Chaozheng Wang, Wenxuan Wang, Cuiyun Gao, and Michael R. Lyu. 2023. Generative Type Inference for Python. *CoRR* abs/2307.09163 (2023).
- [46] Ben Prystawski and Noah D. Goodman. 2023. Why think step-by-step? Reasoning emerges from the locality of experience. *CoRR* abs/2304.03843 (2023).
- [47] Baptiste Rozière, Jonas Gehring, Fabian Gloeckle, Sten Sootla, Itai Gat, Xi-aoqing Ellen Tan, Yossi Adi, Jingyu Liu, Tal Remez, Jérémy Rapin, Artyom Kozhevnikov, Ivan Evtimov, Joanna Bitton, Manish Bhatt, Cristian Canton-Ferrer, Aaron Grattafiori, Wenhan Xiong, Alexandre Défossez, Jade Copet, Faisal Azhar, Hugo Touvron, Louis Martin, Nicolas Usunier, Thomas Scialom, and Gabriel Synaëve. 2023. Code Llama: Open Foundation Models for Code. *CoRR* abs/2308.12950 (2023).
- [48] Rebecca L. Russell, Louis Y. Kim, Lei H. Hamilton, Tomo Lazovich, Jacob Harer, Onur Ozdemir, Paul M. Ellingwood, and Marc W. McConley. 2018. Automated Vulnerability Detection in Source Code Using Deep Representation Learning. In *ICMLA*. IEEE, 757–762.
- [49] Statista. 2023. Common IT vulnerabilities and exposures worldwide 2009–2023. <https://www.statista.com/statistics/>
- [50] Yulei Sui, Ding Ye, and Jingling Xue. 2012. Static memory leak detection using full-sparse value-flow analysis. In *International Symposium on Software Testing and Analysis, ISSTA 2012, Minneapolis, MN, USA, July 15–20, 2012*, Mats Per Erik Heimdahl and Zhendong Su (Eds.). ACM, 254–264.
- [51] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurélien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023. LLaMA: Open and Efficient Foundation Language Models. *CoRR* abs/2302.13971 (2023).
- [52] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajwal Bhargava, Shriti Bhoale, Dan Bikel, Lukas Blecher, Cristian Canton-Ferrer, Moya Chen, Guillem Cucurull, David Esiohu, Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu, Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurélien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. 2023. Llama 2: Open Foundation and Fine-Tuned Chat Models. *CoRR* abs/2307.09288 (2023).
- [53] Lilapati Waikhom and Ripon Patgiri. 2023. A survey of graph neural networks in various learning paradigms: methods, applications, and challenges. *Artif. Intell. Rev.* 56, 7 (2023), 6295–6364.
- [54] Xiaomeng Wang, Tao Zhang, Runpu Wu, Wei Xin, and Changyu Hou. 2018. CPGVA: Code Property Graph based Vulnerability Analysis by Deep Learning. In *10th International Conference on Advanced Infocomm Technology, ICAIT 2018, Stockholm, Sweden, August 12–15, 2018*. IEEE, 184–188.
- [55] Yue Wang, Weishi Wang, Shafiq R. Joty, and Steven C. H. Hoi. 2021. CodeT5: Identifier-aware Unified Pre-trained Encoder-Decoder Models for Code Understanding and Generation. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP 2021, Virtual Event / Punta Cana, Dominican Republic, 7–11 November, 2021*, Marie-Francine Moens, Xuanjing Huang, Lucia Specia, and Scott Wen-tau Yih (Eds.). Association for Computational Linguistics, 8696–8708.
- [56] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. 2022. Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. In *NeurIPS*.
- [57] Xin-Cheng Wen, Yupan Chen, Cuiyun Gao, Hongyu Zhang, Jie M. Zhang, and Qing Liao. 2023. Vulnerability Detection with Graph Simplification and Enhanced Graph Representation Learning. In *45th IEEE/ACM International Conference on*

- Software Engineering, ICSE 2023, Melbourne, Australia, May 14-20, 2023*. IEEE, 2275–2286.
- [58] Xin-Cheng Wen, Cuiyun Gao, Jiaxin Ye, Zhihong Tian, Yan Jia, and Xuan Wang. 2022. Meta-Path Based Attentional Graph Learning Model for Vulnerability Detection. *CoRR* abs/2212.14274 (2022).
- [59] Xin-Cheng Wen, Xincheng Wang, Cuiyun Gao, Shaohua Wang, Yang Liu, and Zhaoquan Gu. 2023. When Less is Enough: Positive and Unlabeled Learning Model for Vulnerability Detection. (2023), 345–357.
- [60] David A. Wheeler. [n.d.]. Flawfinder. <https://dwheeler.com/flawfinder/>
- [61] Fang Wu, Jigang Wang, Jiqiang Liu, and Wei Wang. 2017. Vulnerability detection with deep learning. In *2017 3rd IEEE international conference on computer and communications (ICCC)*. IEEE, 1298–1302.
- [62] Yueming Wu, Deqing Zou, Shihan Dou, Wei Yang, Duo Xu, and Hai Jin. 2022. VulCNN: An Image-inspired Scalable Vulnerability Detection System. In *44th IEEE/ACM 44th International Conference on Software Engineering, ICSE 2022, Pittsburgh, PA, USA, May 25-27, 2022*. ACM, 2365–2376.
- [63] Junwei Zhang, Zhongxin Liu, Xing Hu, Xin Xia, and Shanping Li. 2023. Vulnerability Detection by Learning From Syntax-Based Execution Paths of Code. *IEEE Trans. Software Eng.* 49, 8 (2023), 4196–4212.
- [64] Qinkai Zheng, Xiao Xia, Xu Zou, Yuxiao Dong, Shan Wang, Yufei Xue, Zihan Wang, Lei Shen, Andi Wang, Yang Li, et al. 2023. CodeGeeX: A Pre-Trained Model for Code Generation with Multilingual Evaluations on HumanEval-X. *CoRR* abs/2303.17568 (2023).
- [65] Yaqin Zhou, Shangqing Liu, Jing Kai Siow, Xiaoning Du, and Yang Liu. 2019. Devign: Effective Vulnerability Identification by Learning Comprehensive Program Semantics via Graph Neural Networks. In *Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems 2019, NeurIPS 2019*. 10197–10207.