

SK²DECOMPILE: LLM-BASED TWO-PHASE BINARY DECOMPILE FROM SKELETON TO SKIN

Hanzhuo Tan, Weihao Li, Xiaolong Tian, Siyi Wang, Jiaming Liu, Jing Li, Yuqun Zhang

ABSTRACT

Large Language Models (LLMs) have emerged as a promising approach for binary decompilation. However, the existing LLM-based decompilers are still somewhat limited in effectively presenting a program’s source-level structure with its original identifiers. To mitigate this, we introduce *SK²Decompile*, a novel two-phase approach to decompile from the skeleton (semantic structure) to the skin (identifier) of programs. Specifically, we first apply a Structure Recovery model to translate a program’s binary code to an Intermediate Representation (IR) as deriving the program’s “skeleton”, i.e., preserving control flow and data structures while obfuscating all identifiers with generic placeholders. We also apply reinforcement learning to reward the model for producing program structures that adhere to the syntactic and semantic rules expected by compilers. Second, we apply an Identifier Naming model to produce meaningful identifiers which reflect actual program semantics as deriving the program’s “skin”. We train the Identifier Naming model with a separate reinforcement learning objective that rewards the semantic similarity between its predictions and the reference code. Such a two-phase decompilation process facilitates advancing the correctness and readability of decompilation independently. Our evaluations indicate that *SK²Decompile* significantly outperforms the SOTA baselines, achieving 21.6% average re-executability rate gain over GPT-5-mini on the HumanEval dataset and 29.4% average R2I improvement over Idioms on the GitHub2025 benchmark.

1 INTRODUCTION

Decompilation refers to converting compiled binaries back to high-level source code and has been widely adopted in software security tasks like malware analysis and vulnerability discovery (Brumley et al., 2013; Katz et al., 2018; Wu et al., 2022; Cao et al., 2022; Fu et al., 2019). Ideally, a decompiler ensures both functional correctness and code readability, which can hardly be realized in practice at the same time. For instance, traditional tools like Ghidra (Ghidra, 2024) and IDA (Hex-Rays, 2024) excel at functional correctness but often produce obfuscated, hard-to-read code, while recent Large Language Model (LLM)-based approaches (Hosseini & Dolan-Gavitt, 2022; Armengol-Estap’è et al., 2023; Jiang et al., 2025; Tan et al., 2024; ylfeng et al., 2024; Dramko et al., 2025) generate more readable output but frequently fail to preserve the original program’s functionality (Tan et al., 2024; 2025).

Many research efforts imply the root cause of this trade-off as the intractable complexity of simultaneously inferring control-flow structures, data layouts, and identifiers in a single phase (Lacomis et al., 2019; Xie et al., 2024; Chen et al., 2021b; Console et al., 2023; Patrick-Evans et al., 2020; David et al., 2020; Li et al., 2025). To mitigate this, we introduce *SK²Decompile*, a novel LLM-based decompilation technique that decomposes the binary decompilation task into two phases. In particular, we first derive the program’s skeleton, i.e., its core structure, including control flow and data structure (Aho et al., 2007). Then, we derive the program’s skin, i.e., the meaningful type, variable, and function names reflecting the actual program semantics (Lacomis et al., 2019). Such a two-phase decompilation design allows for tackling the challenges of functionality and readability independently for aggregating their respective effectiveness rather than realizing a trade-off in between. In particular, we design a novel Intermediate Representation (IR) acting as the “skeleton” of the program. This IR essentially refers to the original source code with all identifiers (variable, function, and type names) replaced by generic placeholders (Lachaux et al., 2021) for preserving

structural and functional logic of a program, following the Information Bottleneck principle (Tishby et al., 2000; Tishby & Zaslavsky, 2015). The decompilation process is then split into two sequential phases: *Structure Recovery* where an LLM translates the compiled binary code to our structural IR and *Identifier Naming* where a second LLM enriches the IR by predicting meaningful names reflecting actual program semantics for all placeholders. For Structure Recovery, we first train a sequence-to-sequence model (Cummins et al., 2024; Vaswani et al., 2017) and further tune it with reinforcement learning (RL) (OpenAI, 2023), where the compiler checks the syntax and semantics to provide the reward. A positive reward is generated only if the generated IR successfully compiles, with additional rewards reflecting the correctness of placeholder recovery. For Identifier Naming, we use a separate RL reward. To better capture human-centric readability, this model is not rewarded for exact name match but for the semantic similarity between its output and the reference code (Zhang et al., 2025). In this way, *SK²Decompile* enhances functional correctness and semantic readability simultaneously for LLM-based decompilation.

Our evaluations show that *SK²Decompile* significantly outperforms prior SOTA models on four open-source benchmark suites. To our best knowledge, *SK²Decompile* is the first to approach the average re-executability rate of $\sim 70\%$ on HumanEval (Chen et al., 2021a) and $\sim 60\%$ on MBPP (Austin et al., 2021). It also achieves 21.6% average re-executability rate gain over GPT-5-mini (OpenAI, 2025) on HumanEval and 29.4% average R2I (Eom et al., 2024) improvement over Idioms (Dramko et al., 2025) on the GitHub2025 benchmark (Tan et al., 2025).

The code has been released in anonymous GitHub page ¹. Our main contributions are as follows.

- **Two-phase Decompilation Framework.** We propose the first decompilation framework consisting of two phases: Structure Recovery for advancing the recovery of source-level program structures and Identifier Naming for advancing the recovery of meaningful identifiers reflecting actual program semantics. Each phase trains a model using reinforcement learning with specific rewards respectively.
- **Intermediate Representation (IR).** We propose our IR as the obfuscated source code. This IR satisfies the Information Bottleneck principle by maximizing the compression of the semantics embodied in identifiers while preserving the semantics embodied in the structure of the program, and it is practically simple to generate.
- **Extensive Evaluations.** We perform extensive evaluations on *SK²Decompile* and find out it achieves the optimal performance compared with the studied baselines. For instance, *SK²Decompile* achieves 21.6% average re-executability rate gain over GPT-5-mini on HumanEval and 29.4% average R2I improvement over Idioms on the GitHub2025 benchmark.

2 BACKGROUND

2.1 RELATED WORK

Decompilation, i.e., the reconstruction of source code from binary executables, has long relied on control/data-flow analysis and pattern matching (Brumley et al., 2013; Katz et al., 2018; Wu et al., 2022; Fu et al., 2019). Typically, conventional decompilers like IDA Pro (Hex-Rays, 2024) tend to recover a program’s basic logic, with their generated pseudocode close to low-level assembly code, i.e., their outputs often lack readability and re-executability (Cao et al., 2024; Liu & Wang, 2020).

Motivated by the success of Large Language Models (LLMs) in code-related tasks (Zeng et al., 2022; Wang et al., 2024; Jiang et al., 2024; Su et al., 2024; Wang et al., 2025; Szafraniec et al., 2023), recent research has focused on applying LLMs to refine the pseudocode generated by traditional decompilers (Hu et al., 2024; Wong et al., 2023). Note that as pseudocode is deterministic with the corresponding binary code, we use the terms interchangeably in this paper. Initial efforts, such as LLM4Decompile (Tan et al., 2024), demonstrated that LLMs could effectively learn to translate low-level pseudocode to high-level source code and inspire subsequent studies (yifeng et al., 2024; Feng et al., 2025). Other research focuses on incorporating contextual information. For instance, Idioms (Dramko et al., 2025) enriches the input by incorporating information from adjacent functions in the call graph and attempts to jointly recover user-defined type definitions with the decompiled

¹<https://github.com/albertan017/LLM4Decompile>

code. Recently, D-LIFT (Zou et al., 2025) enhanced the training pipeline by incorporating reinforcement learning, guided by a novel reward function D-SCORE which provides a multi-faceted assessment of code based on accuracy and readability. Despite these advancements, the functional correctness of LLM-based techniques remains a significant challenge, with existing models failing on approximately half of the tasks in the HumanEval-Decompile benchmark (Tan et al., 2024).

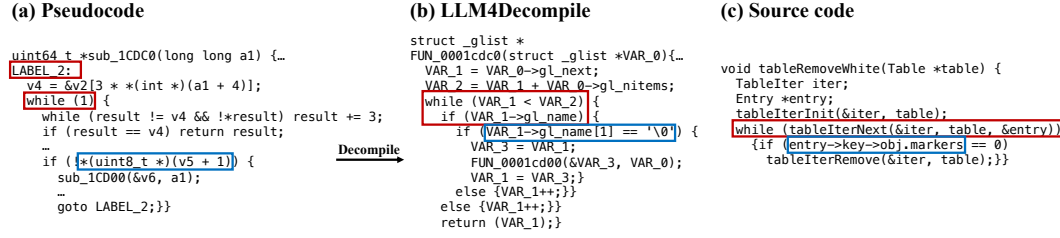


Figure 1: An example with its (a) pseudocode, (b) refinement by LLM4Decompile, and (c) source code. red marks the while loop in different forms, blue represents the data access.

2.2 MOTIVATING EXAMPLE

As shown in Figure 1, while LLM4Decompile, a widely-studied LLM-based decompiler, correctly interprets constructs like `while(1)` and `goto LABEL_2` found in the IDA pseudocode and successfully recovers them to a semantically equivalent and more readable `while` loop. However, the decompiler struggles with recovering the program’s data type structure, and its ability to assign meaningful identifier names reflecting actual program semantics remains limited. For instance, domain-specific types such as `Table` and `Entry` are erroneously mapped to a generic struct named `_glist`. This fundamental limitation in the decompiler’s understanding of the data organization leads to the failure in generating meaningful identifiers. Consequently, variables and functions are reduced to generic placeholders like `VAR_1` and `FUN_0001cdc0`, making it even more difficult to understand the original intent of the program. Such deficiencies motivate a two-phase decompilation process for recovering both program structure and meaningful identifiers respectively rather than realizing a trade-off in between, as illustrated in Section 3.

3 SK²DECOMPILE

3.1 OVERVIEW

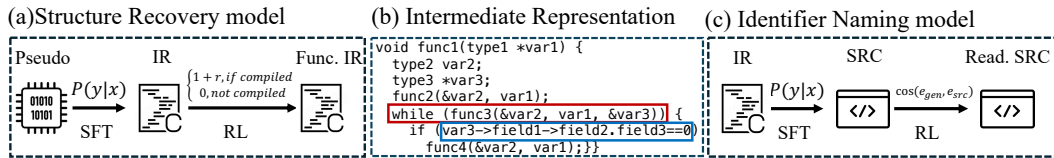


Figure 2: Overview of *SK²Decompile* with two-phase decompilation process (a) Structure Recovery and (c) Identifier Naming. Obfuscated source code (b) serves as the Intermediate Representation (IR) connecting these two phases. For each model, a supervised fine-tuning (SFT) process is performed followed by Reinforcement Learning (RL) with its respective reward.

Figure 2 presents the framework of *SK²Decompile* (**S**keleton-to-**S**kin **D**ecompile) which includes a two-phase decompilation process, i.e., Structure Recovery and Identifier Naming (Section 3.2) which are realized upon the design of the Intermediate Representation (IR, Sections 3.3 and 3.4), with their respective reward functions to advance the correctness and readability of the final decompiled code (Section 3.5).

3.2 TWO-PHASE DECOMPILE PROCESS

We analogize the two-phase decompilation process of *SK²Decompile* to the human body. In particular, Structure Recovery refers to constructing the global code structure, such as loops, conditionals, and data structures, as deriving the program’s “skeleton”. Identifier Naming refers to inferring meaningful names for functions, types, fields, and variables to further reflect actual program semantics.

We formalize *SK²Decompile* from a probabilistic perspective. In particular, the goal of decompilation is to find the most probable source code (s) given a low-level representation, i.e., the pseudocode (u) in our paper. Correspondingly, the decompilation goal can be modeled as maximizing the conditional probability $P(s|u)$. Our core hypothesis is that introducing an intermediate representation (i) can simplify this task. Using the chain rule of probability, we can decompose the probability $P(s|u)$ as $\sum_i P(s|i, u) \cdot P(i|u)$. This decomposition effectively splits the decompilation task into two more manageable sub-tasks, which we illustrate with the example in Figure 2(b) and its corresponding pseudocode and source code presented in Figure 1(a) and Figure 1(c).

Structure Recovery. This phase corresponding to $P(i|u)$ focuses exclusively on translating the syntax and control flow of the low-level pseudocode (Figure 1(a)) to a well-formed, high-level IR (Figure 2(b)). For example, this task includes identifying that a `while(1)` combined with a `goto LABEL_2` in the pseudocode corresponds to a single, conditional `while()` loop structure in the IR. This phase also transforms opaque pointer arithmetic, like `*(uint8_t *) (v5 + 1)`, to a clean, nested structure access `var3->field2->field3`.

Identifier Naming: This phase corresponding to $P(s|i, u)$ takes the recovered IR (Figure 2(b)) and infers meaningful names for variables and functions to produce the final, human-readable source code (Figure 1(c)), e.g., transforming a generic call `var3->field2->field3` to the one with more meaningful, semantic names `entry->key->obj.markers`.

A key insight is that once the clean structured IR is recovered, the original, messy pseudocode provides almost no additional information for the naming task. For instance, after recovering the structure `var3->field2->field3`, the model no longer needs the pointer expression `*(uint8_t *) (v5 + 1)` to infer the correct variable names. This insight allows us to make a Markov assumption (Evans & Rosenthal, 2004), which simplifies the naming probability from $P(s|i, u)$ to $P(s|i)$. This simplification yields our final probabilistic model:

$$P(s|u) \approx \sum_i P(s|i) \cdot P(i|u) \quad (1)$$

By decomposing the problem, we create a focused, two-phase process. First, we solve the complex Structure Recovery challenge ($P(i|u)$), and then perform the Identifier Naming task on a clean, abstract representation ($P(s|i)$). In this way, we reduce the overall complexity for more robust learning and higher-quality decompilation.

3.3 INTERMEDIATE REPRESENTATION

The two-phase decompilation process necessitates an intermediate representation (IR) that serves as a bridge between pseudocode and source code. However, designing the IR presents a fundamental challenge, i.e., it must be simple enough to be reliably recovered from pseudocode, yet rich informative enough to enable accurate source code reconstruction.

This challenge naturally frames our problem as an Information Bottleneck (IB) optimization task (Tishby et al., 2000; Tishby & Zaslavsky, 2015). In particular, for any information flow pseudocode $\rightarrow$ IR $\rightarrow$ source, the intermediate representation acts as a bottleneck that must balance two competing factors, i.e., *compression* and *relevance*. More specifically, compression means that the IR should discard irrelevant details from the pseudocode to make the Structure Recovery phase tractable. Moreover, relevance refers to that the IR must preserve sufficient information to reconstruct the source code in the Identifier Naming phase. Ideally, the IR should be maximally inferable from the pseudocode and structurally close to the target source code. Accordingly, the Information Bottleneck (IB) principle formalizes this trade-off through the objective:

$$\min_{P(i|u)} \mathcal{L}_{IB} = I(u; i) - \beta I(i; s), \quad (2)$$

where $I(u; i)$ measures the mutual information between pseudocode and IR (to be minimized) and $I(i; s)$ measures the mutual information between IR and source code (to be maximized). Formula 2 guides our choice of IR. We thus propose using obfuscated source code (Lachaux et al., 2021), particularly the original source with all identifiers replaced by generic placeholders. Such a representation emerges naturally from the IB objectives. In particular, for the compression objective, the model should distill high-level structural abstractions from the noisy, low-level patterns of the input pseudocode. This process inherently discards irrelevant input details, thus minimizing the mutual information between the input and our IR. Meanwhile, for the relevance objective, the obfuscated code is an ideal structural representation as it can be theoretically recoverable from compiled binary code even when the semantics embodied in original identifiers is lost during compilation. Consequently, this IR preserves the maximum possible relevant information about the source code, thereby maximizing the mutual information between the IR and the source.

Note that the obfuscated code can be automatically generated from source code through identifier obfuscation (Section 3.4), making it practical in real world.

Algorithm 1 Generation of Intermediate Representation (IR)

Require: Source code C , corresponding pseudocode P

Ensure: Obfuscated source code (IR)

```

1: Analyze  $P$  to extract names that need to be preserved:  $F_P$ 
2: Parse  $C$  into an abstract syntax tree (AST):  $T$ 
3: Initialize rename maps  $R[\cdot] \leftarrow \emptyset$  for func, type, field, var
4: Initialize counters  $cnt[\cdot] \leftarrow 1$  for each identifier type
5: Initialize replacement list  $\mathcal{L} \leftarrow \emptyset$ 
6: Traverse( $node$ ):
7:   ( $id\_type, name$ )  $\leftarrow$  classify  $node$ 
8:   if  $name \notin F_P$  then
9:     if  $name \notin R[id\_type]$  then
10:       $new \leftarrow id\_type \parallel cnt[id\_type]$ 
11:       $R[id\_type][name] \leftarrow new$ 
12:       $cnt[id\_type] \leftarrow cnt[id\_type] + 1$ 
13:     end if
14:     Append replacement ( $start(node), end(node), R[id\_type][name]$ ) to  $\mathcal{L}$ 
15:   end if
16:   for each child  $c$  of  $node$  do
17:     Traverse( $c$ )
18:   end for
19: Obfuscate( $C, \mathcal{L}$ ):
20:   Sort  $\mathcal{L}$  by start position in descending order
21:   Let  $IR$  be a mutable copy of  $C$ 
22:   for each ( $s, e, new$ ) in  $\mathcal{L}$  do
23:     Replace substring  $IR[s : e]$  with  $new$ 
24:   end for
25:   return  $IR$ 
26: return  $IR$ 

```

3.4 IR GENERATION

Algorithm 1 illustrates the process of generating the obfuscated code (IR) from the source code. Specifically, the pseudocode is first analyzed to extract all function and type names that should remain unchanged in the obfuscated code, e.g., the standard type `int` and library function `memcpy`, which are stored in the reserved list F_P (line 1). The source code is then parsed into an abstract syntax tree (AST) to provide precise identifier positions (line 2). For each identifier category, we initialize renaming maps and counters, as well as an empty replacement list (lines 3–5). We then invoke the recursive procedure `TRAVERSE` on the root of the AST (lines 6–18). During traversal, each node is classified to determine its identifier type and name (line 7). If the name does not appear in the reserved list F_P , a new obfuscated name is generated and stored in the renaming map (lines 8–13). A replacement entry containing the start and end offsets together with the new

name is then appended to the replacement list (line 14). The procedure continues recursively on all children of the current node (lines 16-18). After traversal, the replacements are applied (lines 19-25) in OBFUSCATE where the list is sorted in descending order of start position (line 20) so that later modifications do not shift earlier offsets, and all substitutions are performed on the original code (lines 21-24). Finally, the obfuscated code, namely, IR, is returned (line 26).

3.5 ENHANCEMENT WITH REINFORCEMENT LEARNING

To recover the structured IR from pseudocode and the identifier names from the structured IR, we adopt the sequence-to-sequence (S2S) paradigm, which is adopted in many neural machine translation models that aim to predict the output given the input sequence (Vaswani et al., 2017). This paradigm typically minimizes the cross-entropy (CE) loss for the predicted tokens: $y_i: \mathcal{L}_{CE}(\theta) = -\sum_{i=1}^N \log P_{\theta}(y_i | y_{<i}, x)$, i.e., calculating the total loss by summing the negative log probabilities of the model correctly predicting each token in a sequence, given all the preceding tokens.

Such CE loss refers to an aggregation of local, token-level prediction errors, serving as a baseline to train an LLM-based decompiler. However, it lacks syntactic and semantic awareness, including assigning equal penalties for unequal errors. For example, a misplaced semicolon, which breaks compilation, might receive a similar penalty to choosing a semantically equivalent but different variable name. Therefore, only adopting a supervised model for Structure Recovery might generate syntactically plausible code that does not compile, limiting the effectiveness of *SK²Decompile*.

To further enhance Structure Recovery, after the S2S training, we perform reinforcement learning (RL) to align the outputs with compiler’s preference and type constraints such that the generated IR could better represent a compilable and functionally sound program. Specifically, we design the reward made up of two components. First, for each generated IR, we provide the compiler with the header of the ground-truth IR in order to verify its compilability and grant a reward only upon success, for advancing functional correctness. Additionally, we reward the accurate recovery of placeholder identifiers by computing the Jaccard similarity coefficient between the generated (I_{gen}) and ground-truth sets (I_{IR}). The placeholder recovery reward encourages the model to accurately reconstruct the program’s data layout. Formally:

$$r_{\text{placeholder}} = \frac{|I_{\text{gen}} \cap I_{\text{IR}}|}{|I_{\text{gen}} \cup I_{\text{IR}}|}, \quad r_{\text{structure}} = \begin{cases} 0.0, & \text{if } IR \text{ cannot be compiled} \\ 1.0 + r_{\text{placeholder}}, & \text{if } IR \text{ can be compiled} \end{cases} \quad (3)$$

Note that using compiler feedback as a reward is feasible and natural. One possible alternative is to build the reward based on the executing unit tests. However, creating unit tests and replicating execution environments for real-world programs is often prohibitively complex and costly.

Similarly, for Identifier Naming, the CE loss undesirably penalizes cases where identifiers differ superficially but are semantically equivalent, while such differences are negligible from a human’s perspective. To mitigate this, we also perform RL and formulate the corresponding reward as the semantic similarity between the embedded generated code ($\mathbf{e}_{\text{gen}}$) and the reference source code ($\mathbf{e}_{\text{src}}$), measured by the cosine similarity:

$$r_{\text{identifier}} = \cos(\mathbf{e}_{\text{gen}}, \mathbf{e}_{\text{src}}) = \frac{\mathbf{e}_{\text{gen}} \cdot \mathbf{e}_{\text{src}}}{\|\mathbf{e}_{\text{gen}}\| \|\mathbf{e}_{\text{src}}\|} \quad (4)$$

By optimizing for this similarity metric, we encourage the model to generate names that are more semantically aligned with the ground truth, in contrast to the CE loss, which strictly enforces an exact lexical match.

4 EXPERIMENTS

Training Data. We collected our training corpus from the C programs of Exebench (Armengol-Estap’e et al., 2022) and Decompile-Bench (Tan et al., 2025) datasets. We compiled the source files into binaries for the x86 Linux platform using GCC and Clang (Clang, 2025), applying optimization levels -O0 through -O3. To ensure data quality, we normalized the code by removing all comments and applying clang-format to the source code, while formatting the pseudocode to adhere to the R2I standard (Eom et al., 2024). We further employed MinHash-LSH to identify and remove

near-duplicates (Broder, 2000). Following previous reverse engineering practices (Lacomis et al., 2019; Chen et al., 2021b; Xie et al., 2024), we stripped all binaries and used IDA Pro (Hex-Rays, 2024) to generate pseudocode (please refer to Appendix A.1 for stripping examples). This process yielded a comprehensive dataset of approximately 5 million samples, totaling around 2B tokens of pseudocode, 1.5B tokens of IR, and 1.5B tokens of source code.

Evaluation Benchmarks and Metrics. For evaluation, we adopted a set of standard benchmarks widely used in previous studies: HumanEval (Chen et al., 2021a), ExeBench (Armengol-Estap’e et al., 2022), MBPP (Austin et al., 2021), and Github2025 (Tan et al., 2025). These benchmarks were processed using the same compilation pipeline as our training data. To assess the quality of the generated decompiled code, we used three primary metrics, i.e., *R2I* (Eom et al., 2024), *GPT-Judge* (Xu et al., 2025; Liu et al., 2023), and *re-executability rate* (Armengol-Estap’e et al., 2023; Tan et al., 2024). In particular, R2I measures the relative readability of code structure. GPT-Judge uses GPT-5-mini (OpenAI, 2025) to evaluate the Identifier Naming effectiveness of the output, with 1 for poor performance to 5 for excellent performance. For benchmarks that support execution (HumanEval and MBPP), we also measure the re-executability rate (Armengol-Estap’e et al., 2023; Tan et al., 2024), which checks if the decompiled code can be successfully re-compiled and passes the original test cases. For tests on stripped binaries, we restore the original function name in the generated code. Note that Exebench is excluded from the evaluation on re-executability rate because the stripping process disrupts its required execution environment. Detailed definitions for each metric and the prompt of GPT-Judge are provided in Appendix A.2.

Baselines. We compare against GPT-5-mini (OpenAI, 2025), a state-of-the-art commercial model, as well as two leading open-source decompilation models LLM4Decompile (Tan et al., 2024) and Idioms (Dramko et al., 2025). Other LLM-based decompilers, such as Nova (Jiang et al., 2025), Ref-Decomp (Feng et al., 2025), and D-Lift (Zou et al., 2025), were not included in our comparison because they do not provide details about their data preprocessing approaches or do not release their models, hindering fair and reproducible evaluations.

Configurations. Both the Structure Recovery and Identifier Naming models were initialized from the LLM4Decompile-6.7B checkpoint (Tan et al., 2024). We performed supervised fine-tuning for one epoch using the LLaMA-Factory library (Zheng et al., 2024) with a batch size of 128 and a learning rate of $3e-6$. For the Reinforcement Learning (RL) phase, we leveraged the GRPO (DeepSeek-AI et al., 2025) algorithm in the veRL library (Sheng et al., 2024) and trained on a random subset of 50,000 samples due to computational constraints. The RL reward for code compilability is verified using Psyche-C (Melo, 2025) to generate headers, and the reward for semantic similarity is measured using qwen-embedding-0.6B (Zhang et al., 2025). All experiments were conducted on clusters of NVIDIA H800-80GB GPUs. During inference, we used the vLLM (Kwon et al., 2023) library for accelerated generation and employed greedy decoding to minimize randomness.

4.1 MAIN RESULTS

Table 1: Re-executability results between the studied decompilers.

Re-Executability Rates	HumanEval					MBPP				
	O0	O1	O2	O3	AVG	O0	O1	O2	O3	AVG
GPT-5-mini	67.07	60.78	49.63	49.56	56.75	55.70	49.33	44.13	39.74	47.23
LLM4Decompile	67.07	37.25	33.58	28.32	41.71	61.56	42.42	36.90	31.32	43.05
Idioms	70.73	25.49	12.41	10.62	29.81	54.78	21.58	11.60	8.06	24.01
<i>SK²Decompile</i>	86.59	70.59	61.31	57.52	69.00	69.76	62.33	54.83	51.58	59.63

Table 1 compares the re-executability rates of the studied decompilers on the HumanEval and MBPP datasets across different optimization levels (O0-O3). Notably, *SK²Decompile* achieves the highest performance, surpassing the best-performing baseline GPT-5-mini by 21.6% and 26.3% averagely on each dataset. Specifically, to the best of our knowledge, this is the first model that preserves the functionality of binaries and reaches an average re-executability of $\sim 70\%$ and $\sim 60\%$ of HumanEval and MBPP cases, underscoring the advantage of decomposing decompilation into two sub-tasks.

Table 2: R2I results between the studied decompilers with the compilation optimization levels O0, O3 and the averaged results on -O{0,1,2,3}.

R2I	HumanEval			MBPP			ExeBench			GitHub2025		
	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG
GPT-5-mini	50.52	40.03	42.33	52.13	41.66	44.14	46.91	40.63	45.25	34.82	26.44	28.95
LLM4Decompile	64.23	68.02	68.71	66.10	68.66	69.19	62.64	52.79	57.02	49.65	49.78	47.95
Idioms	68.50	45.43	53.10	72.11	45.55	55.73	66.77	52.17	60.40	63.50	58.21	57.97
<i>SK²Decompile</i>	64.92	72.70	70.27	66.23	74.43	71.66	69.96	73.55	71.50	76.46	74.51	74.99

Table 3: GPT-Judge results between the studied decompilers.

GPT-Judge	HumanEval			MBPP			ExeBench			GitHub2025		
	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG
GPT-5-mini	4.49	4.07	4.23	4.35	3.88	4.08	2.53	2.33	2.37	3.04	2.86	2.87
LLM4Decompile	3.88	3.29	3.42	3.81	3.22	3.41	2.47	2.12	2.22	2.52	2.56	2.62
Idioms	4.30	2.70	3.22	4.07	2.61	3.13	2.46	1.71	2.01	2.51	2.10	2.18
<i>SK²Decompile</i>	4.51	4.05	4.24	4.31	3.95	4.12	2.48	2.47	2.42	3.05	3.02	3.06

Table 2 presents the R2I results of the studied decompilers. *SK²Decompile* consistently outperforms all the baselines. The improvements are particularly significant in the recovery of program structures from real-world binaries. Specifically, on the ExeBench and GitHub2025 datasets, *SK²Decompile* achieves performance gains of 18.4% and 29.4% over the best-performing baseline Idioms.

The effectiveness of Identifier Naming, as evaluated by GPT-Judge, is presented in Table 3, where *SK²Decompile* produces high-quality names on both the HumanEval and MBPP datasets, achieving scores of 4.24 and 4.12 out of 5, respectively. Furthermore, when applied to the real-world datasets, *SK²Decompile* demonstrates an advantage over the existing techniques, outperforming GPT-5-mini by 2.1% and 6.7%.

4.2 ABLATIONS

Table 4: Re-executability results between the *SK²Decompile* variants.

Re-Executability Rates	HumanEval					MBPP				
	O0	O1	O2	O3	AVG	O0	O1	O2	O3	AVG
pseudo-src	73.78	54.25	48.91	42.48	54.86	58.37	48.97	42.77	39.93	47.51
pseudo-ir	78.66	65.35	54.01	52.21	62.56	55.29	49.33	43.22	41.02	47.25
pseudo-ir-rl	87.80	69.28	59.85	58.41	68.84	68.35	59.88	52.25	47.62	57.06
pseudo-ir-src	78.66	66.01	55.47	54.86	63.75	60.95	55.03	48.34	46.89	52.83
pseudo-ir-src-rl	86.59	70.59	61.31	57.52	69.00	69.76	62.45	54.83	51.58	59.63

In the ablation study, we designed a series of *SK²Decompile* variants to indicate the individual effects of its major components, including the task decomposition and the crafted reward (RL) as follows.

- pseudo-src: This model represents a direct, end-to-end approach to decompile from pseudo to source code. It is trained using SFT with the same training data used in *SK²Decompile*.
- pseudo-ir: This model is trained with SFT to convert pseudocode to IR for the evaluation on the effectiveness of Structure Recovery.
- pseudo-ir-src: This model is trained with SFT to convert IR to source code. The output IR from the Structure Recovery phase serves as its input for evaluating the effectiveness of Identifier Naming. Note that the training cost of a direct approach, pseudo-src and decomposed approach, pseudo-ir with pseudo-ir-src, are similar to ensure fair comparison.
- pseudo-ir-rl: Based on pseudo-ir, the model is further tuned with RL on compiler feedback.
- pseudo-ir-src-rl: This model is the complete version of *SK²Decompile* which integrates the decomposed, two-phase framework enhanced with the RL for both phases.

Table 4 presents our ablation study results in terms of the re-executability rates. Noticing that `pseudo-src` establishes a baseline performance with re-executability rates of 54.86% and 47.51% on HumanEval and MBPP dataset, splitting the decompilation process into Structure Recovery and Identifier Naming (`pseudo-ir-src`) could increase corresponding scores to 63.75% and 52.83% respectively. This validates that tackling decompilation as two simpler sub-tasks is indeed a more effective strategy. Notably, even the Structure Recovery model alone (`pseudo-ir`) surpasses the `pseudo-src` baseline on HumanEval, highlighting that an independent phase on program structure recovery is a critical factor. Moreover, `pseudo-ir-rl` achieves dramatic performance gains of 10.0% and 20.8% over the supervised-only model `pseudo-ir`, indicating the benefit of crafted rewards. At last, `pseudo-ir-src-rl` achieves the best performance, demonstrating that each component of *SK²Decompile* is critical and combining them together is essential for optimizing the performance. We observe similar trends for the R2I and GPT-Judge results and present them in Appendix A.3 due to the page limit.

4.3 CASE STUDY

(a) Pseudocode	(b) Source code	(c) <code>pseudo-src</code>
<pre>long long sub_149F00(int a1, long long a2){ signed int v3; long long v4; v3 = (a1 + 7) & 4294967288; if ((unsigned long long)(v3 + *(uint64_t *) (a2 + 8)) > *(uint64_t *) (a2 + 16)) return malloc(a1); v4 = *(uint64_t *) (a2 + 8); *(uint64_t *) (a2 + 8) = v3 + v4; return v4;}</pre>	<pre>char *dtoa_alloc(int i, Stack_alloc *alloc){ char *rv; int aligned_size = MY_ALIGN(i, SIZEOF_CHARP); if (alloc->free+aligned_size<=alloc->end){ rv = alloc->free; alloc->free += aligned_size; } else rv = malloc(i); return rv;}</pre>	<pre>void *_malloc_aligned(int size, void *opaque){ struct malloc_opaque *opaque_ = (struct malloc_opaque *)opaque; size = (size + 7) & ~7; if (opaque->ptr + size > opaque->end) return malloc(size); void *ptr = opaque->ptr; opaque->ptr += size; return ptr;}</pre>
(d) Structure Recovery (<code>pseudo-ir</code>)	(e) Identifier Naming (<code>pseudo-ir-src</code>)	(f) GPT-5-mini
<pre>void *func1(int var1, type1 *var2){ void *var3; var1 = (var1 + 7) & ~7; if (var2->field1+var1<=var2->field2){ var3 = var2->field1; var2->field1 += var1; return var3;} return malloc(var1);}</pre>	<pre>void *_alloc(int size, malloc_state *state) { void *ret; size = (size + 7) & ~7; if (state->avail + size <= state->limit) { ret = state->avail; state->avail += size; return ret;} return malloc(size);}</pre>	<pre>void *arena_alloc(int size, void *arena){ long aligned = (long)(size + 7) & ~7ULL; long cur = *(long *)((char *)arena + 8); long end = *(long *)((char *)arena + 16); if (aligned + cur > end) return malloc(size); *(long *)((char *)arena + 8) = cur+aligned; return (void *)cur;}</pre>

Figure 3: A case study on a memory allocation function with (a) pseudocode, (b) source code, (c) decompilation result from `pseudo-src` in Table 4, (d) Structure Recovery result from `pseudo-ir` in Table 4, (e) Identifier Naming result from `pseudo-ir-src` in Table 4, and (f) decompilation from GPT-5-mini.

Figure 3 presents a case study on a memory allocation function. A direct decompilation in Figure 3(c) produces non-intuitive code that relies on an explicit type cast from a generic `void * opaque`. It also incorrectly identifies the `free` field as `ptr`. GPT-5-mini in Figure 3(f) fails to reconstruct the data structure, and represents data access with a low-level pointer offset `(char *)arena + 8` instead. In contrast, the Structure Recovery phase of *SK²Decompile* in Figure 3(d) successfully recovers the essential control flow, conditions, and data structures. Building on the clean recovered structure from Figure 3(d), the Identifier Naming phase in Figure 3(e) further enhances readability by assigning meaningful names to identifiers, such as inferring `available` and `state` as semantically appropriate names for the original `free` and `alloc`, leading to a structurally accurate and semantically rich result.

5 CONCLUSION

In this work, we propose *SK²Decompile* which decomposes the binary decompilation task into two phases. First, it recovers the program’s “skeleton”, i.e., its functional structure, using an Intermediate Representation and compiler-guided Reinforcement Learning. Second, it recovers the program’s “skin”, i.e., naming identifiers, with a separate reward on semantic similarity to improve readability. Experimental results show that *SK²Decompile* is the first to achieve the average re-executability rate of approximately 70% on HumanEval and 60% on MBPP datasets. It also achieves a 21.6% average re-executability rate gain over GPT-5-mini on HumanEval and 29.4% average R2I improvement over Idioms on the GitHub2025 benchmark. In conclusion, *SK²Decompile* significantly outperforms the existing techniques in producing functionally correct and human-readable decompilation code.

REPRODUCIBILITY STATEMENT

To ensure full reproducibility of our results, we have made all associated artifacts publicly available in an anonymous GitHub repository. This repository contains the complete source code for our model implementation, training scripts, and evaluation protocols. We also provide the processed testing data, along with scripts for data preparation. For ease of use, pre-trained model weights are also released. The README.md file in the repository offers a step-by-step guide to set up the environment, and replicate the key results presented in this paper.

ETHICS

SK²Decompile was developed under strict ethical guidelines. It is intended for use in legitimate scenarios, such as academic research, debugging, and recovering a company’s own lost source code, where permission is granted or copyright does not apply. To support this, the model was trained exclusively on open-source code from public benchmarks and permissively licensed repositories, e.g., MIT, BSD, and Apache 2.0 (Lozhkov et al., 2024). Notably, commercial software remains well-protected by obfuscation methods that make effective decompilation infeasible (Tan et al., 2024), thus limiting the potential for misuse.

REFERENCES

- Alfred V. Aho, Monica S. Lam, Ravi Sethi, and Jeffrey D. Ullman. Control flow. In *Compilers: Principles, Techniques, and Tools*, chapter 6, pp. 399–408. Addison-Wesley, 2 edition, 2007.
- Jordi Armengol-Estap’e, Jackson Woodruff, Alexander Brauckmann, José Wesley de S. Magalhães, and Michael F. P. O’Boyle. Exebench: an ml-scale dataset of executable c functions. *Proceedings of the 6th ACM SIGPLAN International Symposium on Machine Programming*, 2022. URL <https://api.semanticscholar.org/CorpusID:249536081>.
- Jordi Armengol-Estap’e, Jackson Woodruff, Chris Cummins, and Michael F. P. O’Boyle. Slade: A portable small language model decompiler for optimized assembly. *2024 IEEE/ACM International Symposium on Code Generation and Optimization (CGO)*, pp. 67–80, 2023. URL <https://api.semanticscholar.org/CorpusID:258832373>.
- Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, and Charles Sutton. Program synthesis with large language models, 2021. URL <https://arxiv.org/abs/2108.07732>.
- Andrei Z Broder. Identifying and filtering near-duplicate documents. In *Annual symposium on combinatorial pattern matching*, pp. 1–10. Springer, 2000.
- David Brumley, JongHyup Lee, Edward J. Schwartz, and Maverick Woo. Native x86 decompilation using semantics-preserving structural analysis and iterative control-flow structuring. In Samuel T. King (ed.), *Proceedings of the 22th USENIX Security Symposium, Washington, DC, USA, August 14-16, 2013*, pp. 353–368. USENIX Association, 2013. URL <https://www.usenix.org/conference/usenixsecurity13/technical-sessions/presentation/schwartz>.
- Ying Cao, Ruigang Liang, Kai Chen, and Peiwei Hu. Boosting neural networks to decompile optimized binaries. In *Proceedings of the 38th Annual Computer Security Applications Conference, ACSAC ’22*, pp. 508–518, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450397599. doi: 10.1145/3564625.3567998. URL <https://doi.org/10.1145/3564625.3567998>.
- Ying Cao, Runze Zhang, Ruigang Liang, and Kai Chen. Evaluating the effectiveness of decompilers. In *Proceedings of the 33rd ACM SIGSOFT International Symposium on Software Testing and Analysis, ISSTA 2024*, pp. 491–502, New York, NY, USA, 2024. Association for Computing Machinery. ISBN 9798400706127. doi: 10.1145/3650212.3652144. URL <https://doi.org/10.1145/3650212.3652144>.

Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Pondé de Oliveira Pinto, Jared Kaplan, Harrison Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Joshua Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating large language models trained on code. *CoRR*, abs/2107.03374, 2021a. URL <https://arxiv.org/abs/2107.03374>.

Qibin Chen, Jeremy Lacomis, Edward J. Schwartz, Claire Le Goues, Graham Neubig, and Bogdan Vasilescu. Augmenting decompiler output with learned variable names and types. *ArXiv*, abs/2108.06363, 2021b. URL <https://api.semanticscholar.org/CorpusID:237091166>.

Clang. Clang, 2025. URL <https://clang.llvm.org/>. Accessed: 2025-09-10.

Francesca Console, Giuseppe D’Aquanno, Giuseppe Antonio Di Luna, and Leonardo Querzoni. Binbench: a benchmark for x64 portable operating system interface binary function representations. *PeerJ Computer Science*, 9, 2023. URL <https://api.semanticscholar.org/CorpusID:259029804>.

Chris Cummins, Volker Seeker, Dejan Grubisic, Baptiste Roziere, Jonas Gehring, Gabriel Synnaeve, and Hugh Leather. Meta large language model compiler: Foundation models of compiler optimization. *arXiv preprint arXiv:2407.02524*, 2024.

Yaniv David, Uri Alon, and Eran Yahav. Neural reverse engineering of stripped binaries using augmented control flow graphs. *Proc. ACM Program. Lang.*, 4(OOPSLA), November 2020. doi: 10.1145/3428293. URL <https://doi.org/10.1145/3428293>.

DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojun Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanjuan Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yudian Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu

- Zhang, and Zhen Zhang. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025. URL <https://arxiv.org/abs/2501.12948>.
- Luke Dramko, Claire Le Goues, and Edward J. Schwartz. Idioms: Neural decompilation with joint code and type prediction, 2025. URL <https://arxiv.org/abs/2502.04536>.
- Haeun Eom, Dohee Kim, Sori Lim, Hyungjoon Koo, and Sungjae Hwang. R2i: A relative readability metric for decompiled code. *Proc. ACM Softw. Eng.*, 1(FSE), July 2024. doi: 10.1145/3643744. URL <https://doi.org/10.1145/3643744>.
- Michael J Evans and Jeffrey S Rosenthal. *Probability and statistics: The science of uncertainty*. Macmillan, 2004.
- Yunlong Feng, Bohan Li, Xiaoming Shi, Qingfu Zhu, and Wanxiang Che. Ref decompile: Relabeling and function call enhanced decompile, 2025. URL <https://arxiv.org/abs/2502.12221>.
- Cheng Fu, Huili Chen, Haolan Liu, Xinyun Chen, Yuandong Tian, Farinaz Koushanfar, and Jishen Zhao. Coda: An end-to-end neural program decompiler. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d’Alché-Buc, E. Fox, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL https://proceedings.neurips.cc/paper_files/paper/2019/file/093b60fd0557804c8ba0cbf1453da22f-Paper.pdf.
- Zeyu Gao, Yuxin Cui, Hao Wang, Siliang Qin, Yuanda Wang, Zhang Bolun, and Chao Zhang. DecompileBench: A comprehensive benchmark for evaluating decompilers in real-world scenarios. In Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar (eds.), *Findings of the Association for Computational Linguistics: ACL 2025*, pp. 23250–23267, Vienna, Austria, July 2025. Association for Computational Linguistics. ISBN 979-8-89176-256-5. doi: 10.18653/v1/2025.findings-acl.1194. URL <https://aclanthology.org/2025.findings-acl.1194/>.
- Ghidra. Ghidra software reverse engineering framework, 2024. URL <https://github.com/NationalSecurityAgency/ghidra>.
- Hex-Rays. Ida pro: a cross-platform multi-processor disassembler and debugger, 2024. URL <https://hex-rays.com/ida-pro/>.
- Iman Hosseini and Brendan Dolan-Gavitt. Beyond the C: retargetable decompilation using neural machine translation. *CoRR*, abs/2212.08950, 2022. doi: 10.48550/ARXIV.2212.08950. URL <https://doi.org/10.48550/arXiv.2212.08950>.
- Peiwei Hu, Ruigang Liang, and Kai Chen. Degpt: Optimizing decompiler output with llm. In *Proceedings 2024 Network and Distributed System Security Symposium (2024)*. <https://api.semanticscholar.org/CorpusID>, volume 267622140, 2024.
- Ling Jiang, Junwen An, Huihui Huang, Qiyi Tang, Sen Nie, Shi Wu, and Yuqun Zhang. Binaryai: Binary software composition analysis via intelligent binary source code matching. In *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering, ICSE ’24*, New York, NY, USA, 2024. Association for Computing Machinery. ISBN 9798400702174. doi: 10.1145/3597503.3639100. URL <https://doi.org/10.1145/3597503.3639100>.
- Nan Jiang, Chengxiao Wang, Kevin Liu, Xiangzhe Xu, Lin Tan, Xiangyu Zhang, and Petr Babkin. Nova: Generative language models for assembly code with hierarchical attention and contrastive learning. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=4ytRL3HJrq>.
- Deborah S. Katz, Jason Ruchti, and Eric M. Schulte. Using recurrent neural networks for decompilation. In Rocco Oliveto, Massimiliano Di Penta, and David C. Shepherd (eds.), *25th International Conference on Software Analysis, Evolution and Reengineering, SANER 2018, Campobasso, Italy, March 20-23, 2018*, pp. 346–356. IEEE Computer Society, 2018. doi: 10.1109/SANER.2018.8330222. URL <https://doi.org/10.1109/SANER.2018.8330222>.

- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles*, 2023.
- Marie-Anne Lachaux, Baptiste Roziere, Marc Szafraniec, and Guillaume Lample. Dobf: A de-obfuscation pre-training objective for programming languages. *Advances in Neural Information Processing Systems*, 34:14967–14979, 2021.
- Jeremy Lacomis, Pengcheng Yin, Edward J. Schwartz, Miltiadis Allamanis, Claire Le Goues, Graham Neubig, and Bogdan Vasilescu. Dire: A neural approach to decompiled identifier naming. *2019 34th IEEE/ACM International Conference on Automated Software Engineering (ASE)*, pp. 628–639, 2019. URL <https://api.semanticscholar.org/CorpusID:202676778>.
- Gangyang Li, Xiuwei Shang, Shaoyin Cheng, Junqi Zhang, Li Hu, Xu Zhu, Weiming Zhang, and Nenghai Yu. Beyond the edge of function: Unraveling the patterns of type recovery in binary code, 2025. URL <https://arxiv.org/abs/2503.07243>.
- Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. G-eval: NLG evaluation using gpt-4 with better human alignment. In Houda Bouamor, Juan Pino, and Kalika Bali (eds.), *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 2511–2522, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.153. URL <https://aclanthology.org/2023.emnlp-main.153/>.
- Zhibo Liu and Shuai Wang. How far we have come: testing decompilation correctness of c decompilers. In *Proceedings of the 29th ACM SIGSOFT International Symposium on Software Testing and Analysis*, ISSTA 2020, pp. 475–487, New York, NY, USA, 2020. Association for Computing Machinery. ISBN 9781450380089. doi: 10.1145/3395363.3397370. URL <https://doi.org/10.1145/3395363.3397370>.
- Anton Lozhkov, Raymond Li, Loubna Ben Allal, Federico Cassano, Joel Lamy-Poirier, Nouamane Tazi, Ao Tang, Dmytro Pykhtar, Jiawei Liu, Yuxiang Wei, Tianyang Liu, Max Tian, Denis Kocetkov, Arthur Zucker, Younes Belkada, Zijian Wang, Qian Liu, Dmitry Abulkhanov, Indraneil Paul, Zhuang Li, Wen-Ding Li, Megan Risdal, Jia Li, Jian Zhu, Terry Yue Zhuo, Evgenii Zheltonozhskii, Nii Osa Osa Dade, Wenhao Yu, Lucas Krauß, Naman Jain, Yixuan Su, Xuanli He, Manan Dey, Edoardo Abati, Yekun Chai, Niklas Muennighoff, Xiangru Tang, Muhtasham Oblokulov, Christopher Akiki, Marc Marone, Chenghao Mou, Mayank Mishra, Alex Gu, Binyuan Hui, Tri Dao, Armel Zebaze, Olivier Dehaene, Nicolas Patry, Canwen Xu, Julian McAuley, Han Hu, Torsten Scholak, Sebastien Paquet, Jennifer Robinson, Carolyn Jane Anderson, Nicolas Chapados, Mostofa Patwary, Nima Tajbakhsh, Yacine Jernite, Carlos Muñoz Ferrandis, Lingming Zhang, Sean Hughes, Thomas Wolf, Arjun Guha, Leandro von Werra, and Harm de Vries. Starcoder 2 and the stack v2: The next generation, 2024. URL <https://arxiv.org/abs/2402.19173>.
- Leandro T. C. Melo. A compiler frontend for the c programming language, 2025. URL <https://github.com/ltcmelo/psychec>. Accessed: 2025-09-10.
- OpenAI. GPT-4 technical report. *CoRR*, abs/2303.08774, 2023. doi: 10.48550/ARXIV.2303.08774. URL <https://doi.org/10.48550/arXiv.2303.08774>.
- OpenAI. Gpt-5. Large language model (multimodal) available via OpenAI API, 2025. URL <https://openai.com>. Accessed: 2025-09-25.
- James Patrick-Evans, Lorenzo Cavallaro, and Johannes Kinder. Probabilistic naming of functions in stripped binaries. In *Proceedings of the 36th Annual Computer Security Applications Conference, ACSAC '20*, pp. 373–385, New York, NY, USA, 2020. Association for Computing Machinery. ISBN 9781450388580. doi: 10.1145/3427228.3427265. URL <https://doi.org/10.1145/3427228.3427265>.
- pyparser. Complete c99 parser in pure python, 2025. URL <https://github.com/eliben/pyparser>. Accessed: 2025-09-10.

- Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024.
- Zian Su, Xiangzhe Xu, Ziyang Huang, Kaiyuan Zhang, and Xiangyu Zhang. Source code foundation models are transferable binary analysis knowledge bases. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang (eds.), *Advances in Neural Information Processing Systems*, volume 37, pp. 112624–112655. Curran Associates, Inc., 2024. URL https://proceedings.neurips.cc/paper_files/paper/2024/file/cc83e97320000f4e08cb9e293b12cf7e-Paper-Conference.pdf.
- Marc Szafraniec, Baptiste Roziere, Hugh Leather, Francois Charton, Patrick Labatut, and Gabriel Synnaeve. Code translation with compiler representations, 2023. URL <https://arxiv.org/abs/2207.03578>.
- Hanzhuo Tan, Qi Luo, Jing Li, and Yuqun Zhang. Llm4decompile: Decompiling binary code with large language models. In *Conference on Empirical Methods in Natural Language Processing*, 2024. URL <https://api.semanticscholar.org/CorpusID:268297213>.
- Hanzhuo Tan, Xiaolong Tian, Hanrui Qi, Jiaming Liu, Zuchen Gao, Siyi Wang, Qi Luo, Jing Li, and Yuqun Zhang. Decompile-bench: Million-scale binary-source function pairs for real-world binary decompilation. *arXiv preprint arXiv:2505.12668*, 2025.
- Naftali Tishby and Noga Zaslavsky. Deep learning and the information bottleneck principle. In *2015 IEEE Information Theory Workshop (ITW)*, pp. 1–5. Ieee, 2015.
- Naftali Tishby, Fernando C Pereira, and William Bialek. The information bottleneck method. *arXiv preprint physics/0004057*, 2000.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In Isabelle Guyon, Ulrike von Luxburg, Samy Bengio, Hanna M. Wallach, Rob Fergus, S. V. N. Vishwanathan, and Roman Garnett (eds.), *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pp. 5998–6008, 2017. URL <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.
- Hao Wang, Zeyu Gao, Chao Zhang, Zihan Sha, Mingyang Sun, Yuchen Zhou, Wenyu Zhu, Wenju Sun, Han Qiu, and Xi Xiao. Clap: Learning transferable binary code representations with natural language supervision. In *Proceedings of the 33rd ACM SIGSOFT International Symposium on Software Testing and Analysis, ISSTA 2024*, pp. 503–515, New York, NY, USA, 2024. Association for Computing Machinery. ISBN 9798400706127. doi: 10.1145/3650212.3652145. URL <https://doi.org/10.1145/3650212.3652145>.
- Yongpan Wang, Xin Xu, Xiaojie Zhu, Xiaodong Gu, and Beijun Shen. Salt4decompile: Inferring source-level abstract logic tree for llm-based binary decompilation. *arXiv preprint arXiv:2509.14646*, 2025.
- Wai Kin Wong, Huaijin Wang, Zongjie Li, Zhibo Liu, Shuai Wang, Qiyi Tang, Sen Nie, and Shi Wu. Refining decompiled c code with large language models. *ArXiv*, abs/2310.06530, 2023. URL <https://api.semanticscholar.org/CorpusID:263829777>.
- Ruoyu Wu, Taegyu Kim, Dave (Jing) Tian, Antonio Bianchi, and Dongyan Xu. DnD: A Cross-Architecture deep neural network decompiler. In *31st USENIX Security Symposium (USENIX Security 22)*, pp. 2135–2152, Boston, MA, August 2022. USENIX Association. ISBN 978-1-939133-31-1. URL <https://www.usenix.org/conference/usenixsecurity22/presentation/wu-ruoyu>.
- Danning Xie, Zhuo Zhang, Nan Jiang, Xiangzhe Xu, Lin Tan, and Xiangyu Zhang. Resym: Harnessing llms to recover variable and data structure symbols from stripped binaries. In *Conference on Computer and Communications Security*, 2024. URL <https://api.semanticscholar.org/CorpusID:271540149>.

Jiaqi Xiong, Guoqiang Chen, Kejiang Chen, Han Gao, Shaoyin Cheng, and Weiming Zhang. Hext5: Unified pre-training for stripped binary code information inference. In *Proceedings of the 38th IEEE/ACM International Conference on Automated Software Engineering, ASE '23*, pp. 774–786. IEEE Press, 2024. ISBN 9798350329964. doi: 10.1109/ASE56229.2023.00099. URL <https://doi.org/10.1109/ASE56229.2023.00099>.

Xiangzhe Xu, Zhuo Zhang, Zian Su, Ziyang Huang, Shiwei Feng, Yapeng Ye, Nan Jiang, Danning Xie, Siyuan Cheng, Lin Tan, et al. Unleashing the power of generative model in recovering variable names from stripped binary. In *Proceedings of the Network and Distributed System Security Symposium (NDSS)*, 2025.

ylfeng, Yang Xu, Dechuan Teng, Honglin Mu, Xiao Xu, Libo Qin, Wanxiang Che, and Qingfu Zhu. Self-constructed context decompilation with fined-grained alignment enhancement. In *Conference on Empirical Methods in Natural Language Processing*, 2024. URL <https://api.semanticscholar.org/CorpusID:270710853>.

Zhengran Zeng, Hanzhuo Tan, Haotian Zhang, Jing Li, Yuqun Zhang, and Lingming Zhang. An extensive study on pre-trained models for program understanding and generation. In *Proceedings of the 31st ACM SIGSOFT International Symposium on Software Testing and Analysis, ISSTA 2022*, pp. 39–51, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450393799. doi: 10.1145/3533767.3534390. URL <https://doi.org/10.1145/3533767.3534390>.

Yanzhao Zhang, Mingxin Li, Dingkun Long, Xin Zhang, Huan Lin, Baosong Yang, Pengjun Xie, An Yang, Dayiheng Liu, Junyang Lin, Fei Huang, and Jingren Zhou. Qwen3 embedding: Advancing text embedding and reranking through foundation models. *arXiv preprint arXiv:2506.05176*, 2025.

Yaowei Zheng, Richong Zhang, Junhao Zhang, Yanhan Ye, Zheyao Luo, Zhangchi Feng, and Yongqiang Ma. Llamafactory: Unified efficient fine-tuning of 100+ language models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 3: System Demonstrations)*, Bangkok, Thailand, 2024. Association for Computational Linguistics. URL <http://arxiv.org/abs/2403.13372>.

Muqi Zou, Hongyu Cai, Hongwei Wu, Zion Leonahenahe Basque, Arslan Khan, Berkay Celik, Antonio Bianchi, Dongyan Xu, et al. D-lift: Improving llm-based decompiler backend via code quality-driven fine-tuning. *arXiv preprint arXiv:2506.10125*, 2025.

A APPENDIX

A.1 STRIP

(a) Non-stripped Pseudocode <pre> long long readlock(long long a1, long long a2, unsigned int a3, int a4) { v5 = haslock_search(6v6, a2, a3); if (v5 == -2) { increment_cycle_stats(32LL, v10 == *(uint64_t **)(v13 + 80)); v5 = 0; } if (v5 == -1 && v11) change_victim(v12, 6v6); if (v9) { if (!v11 && !v5 && !(uint32_t *)(v9 + 232)) v5 = -3; cs_unlock(v9); } if (v5 == -1 && v8 != v13) { *(uint8_t *) (v8 + 96) = 1; inline_mysql_cond_broadcast_3(*(uint64_t *) (v8 + 40) + 168LL); cs_unlock(*(uint64_t *) (v8 + 40)); return 0; } return v5; } </pre>	(b) Stripped Pseudocode <pre> long long Sub_FF80(long long a1, long long a2, unsigned int a3, int a4) { v5 = Sub_100390(6v6, a2, a3); if (v5 == -2) { Sub_100630(32LL, v10 == *(uint64_t **)(v13 + 80)); v5 = 0; } if (v5 == -1 && v11) Sub_100670(v12, 6v6); if (v9) { if (!v11 && !v5 && !(uint32_t *)(v9 + 232)) v5 = -3; Sub_FFAD0(v9); } if (v5 == -1 && v8 != v13) { *(uint8_t *) (v8 + 96) = 1; Sub_100170(*(uint64_t *) (v8 + 40) + 168LL); Sub_FF800(*(uint64_t *) (v8 + 40)); return 0; } return v5; } </pre>
---	--

Figure 4: An example with its (a) not striped pseudocode, (b) striped pseudocode

Stripping is the process of removing non-essential information from binary executable files and object files (Patrick-Evans et al., 2020; David et al., 2020; Xiong et al., 2024; Cao et al., 2022). This information, primarily intended for debugging and analysis, is not required for the program’s actual execution. The data typically removed includes *Symbol Tables* and *Debugging Information*. Specifically, symbol tables contain the names and addresses of functions, global variables, and other objects within the program. Debugging information refers to the extra data generated by the compiler

(e.g., with the `-g` flag in GCC) that maps the compiled machine code back to the original source code lines, variable names, and data structures.

Stripping binaries is a **common and standard practice** Lacomis et al. (2019); Chen et al. (2021b); Xie et al. (2024); Xu et al. (2025), particularly for software deployed to production environments, as it ensures size reduction and enhances security. The removal of symbols and debugging information can significantly decrease the size of an executable file. A stripped binary is considerably more difficult for reverse engineers. Without meaningful function and variable names, an attacker or competitor must invest significantly more time and effort to understand the program’s internal workings, business logic, or potential vulnerabilities.

The pseudocode snippets in Figure 4 offer an illustration of stripping on a program. In particular, the non-stripped pseudocode in Figure 4(a) is significantly more readable to a human analyst. It features descriptive function names such as `deadlock`, `deadlock_search`, `increment_cycle_stats`, `change_victim`, and `rc_unlock`. These names provide immediate insight into the potential purpose of the code, suggesting it is part of a system designed to detect and handle deadlocks in a database context, possibly related to MySQL as indicated by `inline_mysql_cond_broadcast_3`. On the other hand, the stripped pseudocode in Figure 4(b) is obfuscated. The meaningful function names have been replaced with generic, tool-generated placeholders like `sub_FFB80`, `sub_100390`, `sub_100630`, and `sub_FFAD0`. These names are derived from the memory addresses of the functions and offer no clues about their functionality. An analyst examining this code would face a much steeper challenge in deciphering the program’s logic and intent.

A.2 METRICS

The Relative Readability Index (R2I) (Eom et al., 2024) is a quantitative metric for evaluating and comparing the readability of decompiled C code, producing a normalized score between 0 and 1. It functions by constructing an Abstract Syntax Tree (AST) for each output, extracting predefined features, and calculating a weighted score. However, the original R2I implementation introduces a significant bias. It discards an entire data sample if any single decompiler’s output fails to be parsed by the `pycparser` (pycparser, 2025) library. This is problematic because `pycparser` often fails on code containing user-defined types and functions, skewing the evaluation towards simpler programs. To create a more robust and unbiased metric, we modified the process. First, we use `pschec` (Melo, 2025) to generate headers, improving the likelihood of successful parsing. More importantly, if a specific output still fails to parse, we assign it a score of 0 instead of discarding the entire sample. This allows us to evaluate the other parsable outputs for that program, ensuring a more comprehensive and fair assessment.

Re-executability is a widely adopted metric in decompilation that evaluates the functional equivalence between an original source function and its decompiled output (Tan et al., 2024; Armengol-Estap’e et al., 2023; ylfeng et al., 2024; Jiang et al., 2025; Feng et al., 2025). Ideally, this means the decompiled function should produce the same output as the original function for every conceivable input. However, since testing every input is impossible, we use a practical approach. We run a set of predefined unit tests on both the original code and the decompiled code. If the outputs match for every single test case, we consider the decompilation successful and “re-executable”. This same concept is often called I/O accuracy or pass rate (Armengol-Estap’e et al., 2023; Jiang et al., 2025).

We leverage GPT-Judge to assess the Identifier Naming effectiveness for the decompilers. GPT-Judge has become increasingly adopted for evaluating LLM-based decompilers (Tan et al., 2024; 2025; Gao et al., 2025). In particular, we use GPT-5-mini (OpenAI, 2025) as an automated evaluator, which is prompted to perform a comparative analysis of the decompiled output and the original source code, specifically focusing on the quality of the recovered identifiers. It provides a rating on a 5-point scale, with 1 for poor performance to 5 for excellent performance. The exact prompt used in our evaluation is detailed in Figure 5.

A.3 ADDITIONAL RESULTS

We present ablation results for structural readability (R2I) and identifier quality (GPT-Judge) in Table 5 and Table 6, respectively. The R2I scores in Table 5 exhibit a consistent trend with our

You are an expert reverse engineering analyst tasked with evaluating LLM decompiler performance. You will receive source code and its decompiled version, then assess the readability of the decompiler's output. For each criterion, provide:

1. An integer score from 1 (very poor) to 5 (excellent)
2. A concise 1-2 sentence rationale

****Input Format:****

1. Original Function [SRC]
2. Decompiled Function [DSRC]

****Scoring guidance:****

****1 – Very Poor****

- 1.1 Function, variable, and field names are meaningless (e.g., `func1`, `var1`, `field\_4`).
- 1.2 Names do not reflect their semantic roles (e.g., a counter named `ptr2`).
- 1.3 Types are missing or collapsed into raw pointers/integers, with no sign of higher-level structures.
- 1.4 Access patterns are opaque (e.g., complex pointer arithmetic instead of `arr[i]` or `obj.field`).

****2 – Poor****

- 2.1 Some identifiers exist, but remain generic and uninformative.
- 2.2 Type information is partially present, but arrays, structs, or objects are poorly reconstructed.
- 2.3 Code is slightly more readable than raw disassembly, yet the correspondence to source-level abstractions is weak.

****3 – Fair****

- 3.1 Function and variable names are somewhat descriptive, though often inconsistent or too generic.
- 3.2 Basic type recovery exists: arrays, pointers, and simple structs are recognizable.
- 3.3 Field and array access are partly reconstructed, but may still fall back to pointer arithmetic in places.
- 3.4 Readability is acceptable, but requires effort to interpret correctly.

****4 – Good****

- 4.1 Names are meaningful, semantically relevant, and generally consistent with their roles.
- 4.2 Structs, arrays, and object types are restored in a way close to natural source code.
- 4.3 Field and array access is mostly clean and human-readable (`obj.field`, `arr[i]`).
- 4.4 Overall readability is high, though not fully equivalent to carefully written source code.

****5 – Excellent****

- 5.1 Function, variable, and type names are clear, natural, and semantically accurate.
- 5.2 Type recovery is faithful, with well-structured classes, structs, and arrays that match typical source-level abstractions.
- 5.3 Field access and indexing are intuitive and entirely free of unnecessary pointer arithmetic.
- 5.4 The recovered code feels almost indistinguishable from human-written source, with excellent overall readability.

****Output Format:****

Provide only a valid JSON object with exactly these two fields:

```
```json
{
 "Code Readability Assessment": {
 "score": <int>,
 "rationale": "<string>"
 }
}
```

Output only the JSON object without additional commentary.

Figure 5: GPT-Judge prompt for a qualitative assessment of Identifier Naming effectiveness.

Table 5: R2I results between the  $SK^2Decompile$  variants. Note that since R2I evaluates decompiled code in a relative context quantitatively (Eom et al., 2024), its values can vary significantly for the same decompiler when compared with different baselines.

R2I	HumanEval			MBPP			ExeBench			GitHub2025		
	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG
pseudo-src	54.62	57.77	56.47	56.53	54.43	55.83	59.11	49.34	55.15	54.41	51.71	53.17
pseudo-ir	51.85	60.25	56.39	54.81	55.73	55.26	55.86	53.96	55.18	58.31	55.30	56.46
pseudo-ir-rl	55.51	60.76	57.53	54.49	58.39	57.36	59.68	60.82	60.92	59.24	56.95	57.15
psuedo-ir-src	53.58	59.52	57.10	55.22	56.30	55.80	56.04	55.50	55.73	59.36	56.06	57.33
psuedo-ir-src-rl	56.41	59.01	57.49	54.68	59.28	57.75	59.50	60.59	61.06	59.70	57.56	57.73

re-executability findings (Table 4), further indicating that both task decomposition and specialized rewards improve structural recovery. Table 6 presents a similar trend for Identifier Naming. As expected, the Structure Recovery models score poorly on this metric since they are explicitly designed not to restore original names. Overall, the results altogether validate the effectiveness of our decomposed approach.

Table 6: GPT-Judge results between the  $SK^2Decompile$  variants

GPT-judge	HumanEval			MBPP			ExeBench			GitHub2025		
	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG	O0	O3	AVG
pseudo-src	4.45	3.80	4.05	4.23	3.89	4.03	2.66	2.30	2.37	3.08	2.89	3.00
pseudo-ir	2.88	2.69	2.74	2.78	2.64	2.72	1.96	1.73	1.75	2.42	2.23	2.34
pseudo-ir-rl	2.93	2.69	2.79	2.80	2.67	2.73	1.97	1.73	1.77	2.43	2.32	2.35
pseudo-ir-src	4.48	3.99	4.16	4.26	3.94	4.09	2.47	2.45	2.38	3.02	2.96	3.03
pseudo-ir-src-rl	4.51	4.05	4.24	4.31	3.95	4.12	2.48	2.47	2.42	3.05	3.02	3.06

#### A.4 THE USE OF LARGE LANGUAGE MODELS

During the preparation of this manuscript, we utilized a Large Language Model (LLM) only for assistance with writing. The LLM’s role was strictly limited to proofreading, correcting grammatical errors, and improving the clarity and readability of the text. The core research ideas, methodologies, and conclusions presented in this paper were conceived and developed entirely by the authors.