

# TRANSFORMERS ARE EFFICIENT COMPILERS, PROVABLY

**Anonymous authors**

Paper under double-blind review

## ABSTRACT

Transformer-based large language models (LLMs) have demonstrated surprisingly robust performance across a wide range of language-related tasks, including programming language understanding and generation. In this paper, we take the first steps towards a formal investigation of using transformers as compilers from an expressive power perspective. To this end, we introduce a representative programming language, **Mini-Husky**, which encapsulates key features of modern C-like languages. We show that if the input code sequence has a bounded depth in both the Abstract Syntax Tree (AST) and type inference (reasonable assumptions based on the *clean code principle*), then the number of parameters required by transformers depends only on the *logarithm of the input sequence length* to handle compilation tasks, such as AST construction, symbol resolution, and type analysis. A significant technical challenge stems from the fact that transformers operate at a low level, where each layer processes the input sequence as raw vectors without explicitly associating them with predefined structure or meaning. In contrast, high-level compiler tasks necessitate managing intricate relationships and structured program information. Our primary technical contribution is the development of a domain-specific language, **Cybertron**, which generates formal proofs of the transformer’s expressive power, scaling to address compiler tasks. We further establish that recurrent neural networks (RNNs) require at least a linear number of parameters relative to the input sequence, leading to an exponential separation between transformers and RNNs. Finally, we empirically validate our theoretical results by comparing transformers and RNNs on compiler tasks within **Mini-Husky**.

## 1 INTRODUCTION

Transformers (Vaswani, 2017) have demonstrated remarkable proficiency across various domains, achieving near-expert performance in solving International Mathematical Olympiad problems (Google Deepmind, 2024) and excelling in complex reasoning tasks in science, coding, and mathematics (OpenAI, 2024a). They also handle routine coding tasks with high precision and have been integrated into code editors to significantly boost programmers’ productivity (cur, 2024; Taelin, 2023a). Despite these advancements, the full extent of their underlying capabilities remains only partially understood.

In this paper, we aim to deepen our understanding of transformers’ abilities to perform compilation tasks. Empirically, transformer-based LLMs have shown rapid progress in code generation and compilation. For example, MetaLL (Cummins et al., 2024) enables LLMs to optimize code by interpreting compiler intermediate representations (IRs), assembly language, and optimization techniques. Gu (2023) highlights the ability of LLMs to generate high-quality test cases for Golang compilers. Surprisingly, Taelin (2023b) demonstrates that models like Sonnet-3.5 can compile legacy code into modern languages like TypeScript, outperforming the now obsolete AgdaJS compiler (Agda Development Team, 2024).

To formally study this problem in a controlled setup, we designed a C-like programming language called **mini-husky**, which encapsulates key features of modern C-like languages such as (Flanagan, 2011) and Rust (Klabnik & Nichols, 2023). We focus on three representative compilation tasks: abstract syntax tree (AST) construction, symbol resolution, and type analysis. The AST is a recursive

structure that represents the input as a tree. From the perspective of programming language design, the AST is considered the true representation of the input, with the textual code serving merely as a convenient interface for human users (Alfred et al., 2007). All syntactic and semantic processing can then be interpreted as specific operations on these trees. Symbol resolution involves verifying the validity of references to entities and flagging errors for undefined symbols. Type analysis encompasses both type inference, which assigns types to variables without explicit annotations, and type checking, which identifies mismatches between actual and expected types.

We demonstrate that under the *clean code principle* (Martin, 2008), transformers can efficiently perform AST construction, symbol resolution, and type analysis, where efficiency means that these tasks can be conducted by transformers with a number of parameters that scale logarithmically with the input code length. To the best of our knowledge, this is the first theoretical demonstration that transformers can function as compilers in a parameter-efficient manner.

We further compare transformers and recurrent neural networks (RNNs). By connecting the type analysis task with the associative recall, we show even under the *clean code principle* (Martin, 2008), RNNs require a memory size that scales linearly with the input sequence length to successfully perform type analysis. Consequently, for type analysis in compilation, transformers can be exponentially more efficient than RNNs. We also empirically validate our theoretical findings by demonstrating the superiority of transformers in the type analysis task.

### Technical Challenges and Our Technique.

Proving that transformers can perform compilation tasks presents several challenges:

- **Transformers operate at too low a level.** Transformers process sequences of floating-point vectors, akin to raw bits in computers, and proving their ability to perform specific tasks is similar to writing specialized parallel machine code. Previous work (Yao et al., 2021) often resorts to graphical illustrations for readability, even for basic tasks.
- **Compilers are exceedingly high-level.** Compilers are among the most complex programming endeavors of our time. Compilation involves numerous sophisticated procedures, some of which are undecidable or computationally expensive, such as code optimization (Alfred et al., 2007) and type analysis (Pierce, 2002). For example, type analysis in complex type systems poses significant challenges, often requiring the development of advanced logical frameworks (Dunfield & Krishnaswami, 2019).

To overcome these challenges, we design a domain-specific language (DSL) called **Cybertron** to serve as the proof vehicle, i.e., a major part of our proof consists of reasoning about type-correct code in Cybertron that represents a transformer. Without using **Cybertron**, writing an equivalent natural language proof would be too complex and intractable. Using code to prove propositions is not new to computer science; it is, in fact, the norm in interactive theorem proving (ITP) (Harrison et al., 2014). ITP focuses on generating computer-verifiable proofs through a combination of human-guided instructions and software automation. For instance, the correctness of the Kepler conjecture (Hales et al., 2017) is verified by the combination of the ITP theorem provers HOL Light (Harrison, 2009) and Isabelle (Paulson, 1994). To the best of our knowledge, we are the first to apply this approach to understanding neural networks.

**Contributions.** We summarize our contributions below:

- **A testbed for compilation tasks:** We introduce **Mini-Husky**, a simple yet representative C-like programming language, designed to formally assess transformers’ capabilities in programming language processing. We anticipate that **Mini-Husky** will become a standard testbed for this purpose.
- **Expressive power theory of transformers for several compilation tasks:** We provide a formal proof that, when the input code sequence has bounded AST depth and inference depth, the number of parameters in transformers only needs to scale logarithmically with the input sequence length to handle compilation tasks such as AST construction, symbol resolution, and type analysis. To the best of our knowledge, this is the first study exploring the power of transformers for these compilation tasks.
- **Transformers vs. RNNs:** Theoretically, we demonstrate a negative result, showing that the number of parameters in RNNs must scale linearly with the input sequence length to perform type

analysis correctly. This result establishes an exponential separation between transformers and RNNs. We further empirically confirm the advantage of transformers for the type analysis task.

- **A Domain-Specific Language for Proofs:** Given the challenges in formal proofs, we design a domain-specific language, **Cybertron**, to serve as a proof vehicle. We believe that **Cybertron**, and the general approach of using DSLs for analysis, can have broader applications in understanding transformers and other architectures.

## 2 RELATED WORK

**Expressive Power of Transformers.** A line of work studies the expressive power of attention-based models. One direction focuses on the universal approximation power (Yun et al., 2019; Bhattamishra et al., 2020b;c; Dehghani et al., 2018; Pérez et al., 2021). More recent works present fine-grained characterizations of the expressive power for certain functions in different settings, sometimes with statistical analyses (Edelman et al., 2022; Elhage et al., 2021; Likhoshesterov et al., 2021; Akyürek et al., 2022; Zhao et al., 2023; Yao et al., 2021; Anil et al., 2022; Barak et al., 2022; Garg et al., 2022; Von Oswald et al., 2022; Bai et al., 2023; Olsson et al., 2022; Akyürek et al., 2022; Li et al., 2023; Hao et al., 2022; Pérez et al., 2019; Strobl, 2023; Chiang et al., 2023; Wei et al., 2022; Wang et al., 2022; Feng et al., 2023; Li et al., 2024; Reddit User, 2013). There are also characterizations of transformers to be as powerful as universal computers if put in a looped context (Giannou et al., 2023). The most related one is Yao et al. (2021) where the authors prove constructively that bounded depth Dyck language can be recognized by encoder-only hard attention transformers, which has similarities to our settings of bounded depth programming language recognized encoder-only hard attention transformers. The major difference is that we introduce concepts and tasks from programming language theory Pierce (2002) to study the semantic powers of transformers.

**Transformers vs. RNN.** It is important to understand the comparative advantages and disadvantages of transformers against RNNs. Empirically, synthetic experiments have shown an advantage of transformers against RNNs for long range tasks (Bhattamishra et al., 2023; Arora et al., 2023). Theoretically, there has been a rich line of work focusing on comparing transformers and RNNs in terms of recognizing formal languages (Bhattamishra et al., 2020a; Hahn, 2019; Merrill et al., 2021), which show that the lack of recursive structure of transformers prevent them from recognizing some formal languages that RNNs can recognize. However, the gap can be mitigated when we consider the bounded length of input or bounded grammar depth (Liu et al., 2022; Yao et al., 2021), which is quite reasonable in practice and is used in this paper. On the other side, prior work (Jelassi et al., 2024; Wen et al., 2024) proves a representation gap between RNNs and Transformers in repeating a long sequence. In summary, it is somehow intuitive that recursive structures with limited memory perform badly at tasks which requires information retrieval. Our paper shows that semantic analysis for programming languages is such a task.

**DSLs for Transformers.** We note that we are not exactly the first one to employ a domain-specific language to understand the expressive powers of transformers. Previously, DSLs with simple typings like RASP (Weiss et al., 2021) were proposed to prove constructively that transformers can do various basic sequence-to-sequence operations. Lindner et al. (2023) writes a compiler that compiles RASP into actual transformers, Friedman et al. (2023) shows that RASP can be learned, and Zhou et al. (2023) uses RASP to prove that simple transformers can perform certain algorithms. The major difference between RASP and our DSL **Cybertron** is that **Cybertron** has a powerful algebraic type system that helps prove complicated operations beyond simple algorithms.

## 3 PRELIMINARIES

The major innovation in the transformer architecture is that it uses self-attention solely without a conjunction with a recurrent network (Vaswani, 2017), which processes input tokens in a distributed manner. This capability enables the model to handle long-range dependencies, a crucial feature for language tasks. We use hard attention and simplified position encoding to simplify our theoretical reasoning.

**Attention.** In practice, attention heads use **soft attention**. Given model dimension  $d_{\text{model}}$ , number of heads  $H$ , and a finite set of token positions Pos, an attention layer with simplified position

encoding is defined as a function  $f_{\text{attn}} : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$  given by

$$\forall p \in \text{Pos}, \quad f_{\text{attn}}(X)_p := W_O \text{Concat} \left( \text{Attn}^{(1)}(X)_p, \dots, \text{Attn}^{(H)}(X)_p \right), \quad (1)$$

where the  $h$ th attention head is defined using soft attention as:  $\text{Attn}^{(h)}(X)_p := \sum_{p' \in \text{Pos}} \alpha_{p,p'}^{(h)} V_{p'}^{(h)}$ .

The attention weights  $\alpha_{p,p'}^{(h)}$  given by:  $\alpha_{p,p'}^{(h)} = \frac{\exp \left( Q_p^{(h)\top} K_{p'}^{(h)} + \lambda^{(h)\top} \Psi_{p'-p} \right)}{\sum_{p'' \in \text{Pos}} \exp \left( Q_p^{(h)\top} K_{p''}^{(h)} + \lambda^{(h)\top} \Psi_{p''-p} \right)}$ , where  $W_O \in \mathbb{R}^{d_{\text{model}} \times d_{\text{model}}}$  are trainable parameters,  $Q_p^{(h)}, K_p^{(h)}, V_p^{(h)} \in \mathbb{R}^{d_{\text{model}}/H}$  are linear transformations of  $X_p$ ,  $\lambda^{(h)} \in \mathbb{R}^2$  depends on the head, and  $\Psi_q = \begin{pmatrix} q \\ 1_{q>0} \end{pmatrix} \in \mathbb{R}^2$  accounts for relative position.

For theoretical convenience, we use hard attention, commonly used in theoretical analysis of transformer (Yao et al., 2021; Hahn, 2019). Hard attention can be viewed as the limit of soft attention when the attention logits become infinitely large. The hard attention head is defined as:

$$\text{Attn}^{(h)}(X)_p := \frac{1}{|S_p|} \sum_{p' \in S_p} V_{p'}^{(h)}, \quad \text{where } S_p = \arg \max_{p' \in \text{Pos}} \left( Q_p^{(h)\top} K_{p'}^{(h)} + \lambda^{(h)\top} \Psi_{p'-p} \right) \quad (2)$$

In other words, hard attention selects the positions  $p'$  that maximize the attention score for each position  $p$ , and averages the corresponding value vectors  $V_{p'}^{(h)}$ .

**Feed-Forward Layer.** Given model dimension  $d_{\text{model}}$ , and a finite set of token positions Pos, a feed-forward layer is a fully connected layer applied independently to each position, defined as a function  $f_{\text{ffn}} : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$  given by

$$\forall p \in \text{Pos}, \quad f_{\text{ffn}}(X)_p = W_2 \sigma_{\text{ReLU}}(W_1 X_p + b_1) + b_2, \quad (3)$$

where  $W_1 \in \mathbb{R}^{d_{\text{ffn}} \times d_{\text{model}}}$  and  $W_2 \in \mathbb{R}^{d_{\text{model}} \times d_{\text{ffn}}}$  are trainable weight matrices,  $b_1 \in \mathbb{R}^{d_{\text{ffn}}}$  and  $b_2 \in \mathbb{R}^{d_{\text{model}}}$  are trainable bias vectors,  $d_{\text{ffn}}$  is the hidden dimension of the feed-forward layer, chosen to be  $2d_{\text{model}}$ , as commonly used in practice,  $\sigma_{\text{ReLU}}$  is the ReLU activation function.

**Encoder-Only Transformer.** Encoder-only transformers consist solely of the encoder stack, making them ideal for tasks like classification, regression, and sequence labeling that do not require sequence generation. Each encoder layer includes a multi-head self-attention mechanism and a feed-forward network, allowing the model to capture complex dependencies and contextual information.

One can define it using the following recurrence,

- The input is given by:  $X^{(0)} = X$ .
- For each layer  $l = 1, 2, \dots, L$ :
  - Compute attention output:  $\hat{X}^{(l)} = X^{(l-1)} + f_{\text{attn}}^{(l)}(X^{(l-1)})$ ,
  - Compute feed-forward output:  $X^{(l)} = \hat{X}^{(l)} + f_{\text{ffn}}^{(l)}(\hat{X}^{(l)})$ .

In the above,  $f_{\text{attn}}^{(l)}$  are the attention layers, and  $f_{\text{ffn}}^{(l)}$  are the feed-forward layers, with the same model dimension  $d_{\text{model}}$ , number of heads  $H$ , and set of token positions Pos. For simplicity, layer normalization is ignored. See Appendix C for full details of transformers and other architectures.

## 4 PROGRAMMING LANGUAGE PROCESSING AND THE TARGET C-LIKE LANGUAGE: MINI-HUSKY

Recently, transformers have expanded to support code analysis and generation (Nijkamp et al., 2023; Chen et al., 2021; Anysphere, 2023). Programming languages offer a cleaner foundation for studying language understanding, as their syntactic and semantic tasks are precisely defined. To formally study the language processing capabilities of transformers, we design **Mini-Husky**, a representative mix of modern C-like languages with strong typing and typical syntactic features. It supports

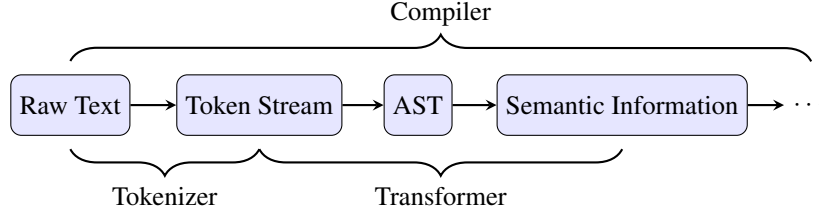


Figure 1: Programming language processing pipeline

user-defined types (e.g., structs, enums) and enforces strict type equality, disallowing implicit conversions. Lexical scoping, including shadowing, ensures proper variable accessibility based on block structures, type inference, and type checking. These features make compiling *Mini-Husky* a representative task to evaluate transformers’ capabilities in syntactic and semantic tasks like symbol resolution and type checking. See Appendix E for the full details of *Mini-Husky*.

The standard pipeline of processing programming languages is shown in Figure 1 (Alfred et al., 2007). The raw text is firstly segmented into parts like literals, identifiers, punctuations, keywords, etc, called token stream, then parsed into a tree-like structure representation generated from the input, finally syntactic and semantic analysis is performed on the tree. *Afterward, an intermediate language program is generated based on the syntactic and semantic analysis, which is further optimized and finally transformed into targeted machine code.* In this paper, to simplify the presentation, we assume the tokenizer has been provided a priori. *Below we describe the programming language processing tasks investigated in the paper.*

**Abstract Syntax Tree Construction.** Abstract Syntax Tree (AST) is a hierarchical, tree-like representation of the syntactic structure of source code in a programming language. Unlike the raw text of the code, the AST abstracts away surface syntax details, capturing the essential elements and their relationships in a structured form. Each node in the AST corresponds to a construct occurring in the source code, such as expressions, statements, or declarations. This representation is central to various stages of language processing, enabling efficient syntax checking, semantic analysis, and code generation. The formal definition of ASTs is standard in the programming language literature but is lengthy, so we defer it to Appendix A.

The AST construction task’s final output is the collection of all AST nodes. We will show transformers can construct AST efficiently.

**Symbol Resolution.** In programming languages, symbols are functions, types, generics, variables, macros, etc. They are defined somewhere and can be used by referring to the corresponding identifier or path in a certain scope. The scope can be within a certain tree of modules, or within a certain curly braced scope within one module. For simplicity, we only consider curly braced scope.

In *Mini-Husky*, the following showcases symbol resolution.

```

1  pub fn f() {
2      fn f1() {}
3
4      let a = 1;
5      let x = a;
6      let a = 2;
7      {
8          let a = 3;
9          { let a = 4; }
10         let y = a;
11     }
12     let z = a;
13 }
14
15 fn g() { f() }
```

The outer function  $f$  is accessible everywhere in the body of function  $g$ . However, the inner function  $f_1$  can only be used inside the body of  $f$  as it is defined within the body. For variables with the same identifier `a`, the first is accessible from line 5, the second is accessible from line 12, the third is accessible from line 10, and the fourth is not accessible from anywhere. Thus  $x = 1, y = 3, z = 2$ .

The output of the **symbol resolution** task is the collection of symbol resolution results on all applicable tokens. More concretely, the output is a sequence of values of type `Option<SymbolResolution>` where `Option<SymbolResolution>` is the type `SymbolResolution` with a null value added for non-applicability and `SymbolResolution` is the type storing the result of the symbol resolution, being either a success with a resolved symbol of type `Symbol` or a failure with an error of type `SymbolResolutionError`. We shall prove that transformers can do symbol resolution and that attention is crucial.

**Type Analysis.** In general, types are essential for conveying the intended usage of the written functions and specifying constraints. As a first exploration of this topic, we try to make the type analysis in **Mini-Husky** as simple as possible yet able to bring out the essential difficulty. The type system consists of four sequential components: (1) *Type definition*, (2) *Type specification*, (3) *Type inference*, and (4) *Type checking*. Due to the page limit, here we only introduce (4) *Type checking* because it is the final step and this is a crucial step which separates transformers and RNNs. See Appendix E.1 for details of (1) *Type definition*, (2) *Type specification*, and (3) *Type inference*.

Type checking ensures that the typed expressions agree with its expectations. For simplicity, we do not allow implicit type conversion, so the agreement means exact equality of types. The arguments of function calls are expected to have types according to the definition of the function. The operand type of field access must be a struct type with a field of the same name. The type of the last expression of the function body or the expr in the return statement must be equal to the return type of the function. For variables defined in the `let` statement, If the types are annotated, the types of the left-hand side and right-hand side should be in agreement.

```
1 // Type Error: the return type is 'i32', yet the last expression is of type 'f32'
2 fn f(a: i32) -> i32 { 1.1 }
3
4 struct A { x: i32 }
5
6 fn g() {
7   // Type Error: 'x' is of type f32 but it's assigned by a value of type 'i32'
8   // Type Error: the first argument of 'f' is expected to be of type 'i32' but gets a
   float literal instead
9   let x: f32 = f(1.1);
10  // Type Error: no field named 'y'
11  let y = A { x: 1 }.y;
12 }
```

The above incorporates typical examples of type disagreements that count as type errors. A compiler should be able to report these errors.

The **type analysis** task's final output is the collection of all type errors. More concretely, the output is a sequence of `Option<TypeError>`, where `Option<TypeError>` denoted the type `TypeError` will a null value added and `TypeError` is the type storing the information of a type error. The position of type errors agrees with the source tokens leading to these errors.

## 5 EXPRESSIVE POWER OF TRANSFORMERS AS EFFICIENT COMPILERS

In this section we discuss main theoretical results about the expressive power of transformers to perform compilation tasks: AST construction, symbol resolution, and type analysis. In Section 5.4, we discuss **Cybertron**, a DSL specifically designed for our proof.

### 5.1 ABSTRACT SYNTAX TREE CONSTRUCTION

We start with a definition that characterizes low-complexity code.

**Definition 1** (code with Bounded AST-Depth). *Let  $\text{MiniHusky}_D$  be the set of token sequences that can be parsed into valid ASTs in **Mini-Husky** with a depth less than  $D$ .*

$D$  in the above definition is small in practice, and a linear dependency on  $D$  is acceptable, but the linear dependency on the length of the token sequence  $L$  is not. The fundamental reason is that the *clean code principle* (Martin, 2008) requires one to write code with as little nested layer as

possible for greater readability. Readability is of the utmost importance because “Programs are meant to be read by humans and only incidentally for computers to execute” (Abelson et al., 1996). This assumption of bounded hierarchical depth is not limited to just programming languages, but is often seen as applicable to natural languages (Frank et al., 2012; Brennan & Hale, 2019; Ding et al., 2017), motivating Yao et al. (2021) to have a similar boundedness assumption. Below is the main result for AST construction using transformers.

**Theorem 1.** *There exists a transformer encoder of model dimension and number of layers being  $O(\log L + D)$  and number of heads being  $O(1)$  that represents a function that maps any token sequence of length  $L$  in  $\text{MiniHusky}_D$  to its abstract syntax tree represented as a sequence.*

We note  $\log L$  is small because 64-bit computers can only process context length at most  $2^{64}$  and  $D$  is small by assumption. Therefore, there exists a transformer with an almost constant number of parameters that is able to process comparatively much longer context length.

*Proof Sketch.* The idea is to construct ASTs in a bottom-up manner with full parallelism. We shall recursively produce the final ASTs in at most  $D$  steps. We shall maintain two values, called `pre_ast` and `ast`. `ast` represents ASTs that have already been allocated, although they might not have been fully initialized. `pre_ast` represents tokens that have yet to form ASTs and new ASTs that have not been fully initialized. For each round, we try to create new ASTs from `pre_ast` and update `ast` and `pre_ast`. For the  $n$ -th round, we provably allocated all ASTs with a depth no more than  $n$ . Then for the  $D$ -th round, all ASTs are properly constructed and allocated. Each round can be represented by a transformer of  $O(1)$  number of heads, model dimension  $O(\log L + D)$ , and  $O(1)$  number of layers. Therefore, the end-to-end process is then representable by a transformer of  $O(1)$  number of heads, model dimension  $O(\log L + D)$ , and  $O(\log L + D)$  number of layers. See full details in Appendix F.  $\square$

## 5.2 SYMBOL RESOLUTION

Next, we show that transformers can effectively perform symbolic resolution as  $\log L$  and  $D$  are almost constant as compared with context length  $L$ .

**Theorem 2.** *There exists a transformer encoder of model dimension and number of layers being  $O(\log L + D)$  and number of heads being  $O(1)$  that represents a function that maps any token sequence of length  $L$  in  $\text{MiniHusky}_D$  to its symbol resolution represented as a sequence of values of type `Option<SymbolResolution>`.*

*Proof Sketch.* First, we need to define the type for scopes. It is represented by a tiny sequence of indices of curly brace block AST that enclose the type/function/variable. We assign the scope by walking through the ASTs in a top-down manner. We not only assign scopes to item definitions, we also: (1) assign scopes to ASTs representing curly brace blocks, with these scopes equal to the scope of block itself, and (2) assign scopes to identifiers waiting to be resolved, with these scopes equal to the maximum possible scope of its resolved definition. The computation process is easily represented in Cybertron, indicating attention is expressive enough for this calculation and it only takes  $O(D)$  number of layers.

After obtaining all the scopes for all items, it takes only one additional layer to obtain the symbolic resolution through attention. As attention is expressed through the dot product of two linear projections  $Q$  and  $K$ , we have to choose the representation of the scope type properly to finish the proof. The full details are in Appendix G.  $\square$

## 5.3 TYPE ANALYSIS

We need an additional definition to characterize the complexity of code for type analysis.

**Definition 2** (code with Bounded AST-Depth and Type-Inference-Depth). *We use  $\text{MiniHuskyAnnotated}_{D,H}$  to denote the subset of  $\text{MiniHusky}_D$  with the depth of type inference no more than  $H$ . The depth of type inference is the number of rounds of computation needed to infer all the types using the type-inference algorithm (described in Appendix E.1).*

In practice,  $H$  is significantly smaller than the context length  $L$  for reasonably written code because it is upper bounded by the number of statements in a function body which is required to be small according to the *clean code principle* (Martin, 2008). Below, we present the main result of using transformers for type analysis. See full details in Appendix H.

**Theorem 3.** *For  $L, D, H \in \mathbb{N}$ , there exists a transformer encoder of model dimension, and number of layers being  $O(\log L + D + H)$  and number of heads being  $O(1)$  that represents a function that maps any token sequence of length  $L$  in  $\text{MiniHuskyAnnotated}_{D,H}$  to its type errors represented as a sequence of values of type `Option<TypeError>`.*

#### 5.4 PROOF VEHICLE: **CYBERTRON**, A DOMAIN-SPECIFIC LANGUAGE

Here we highlight our main proof technique. Proving that transformers can express complex algorithms and software like compilers is a significant challenge due to the inherent differences between how transformers operate and the nature of high-level tasks they are expected to perform. Transformers process input at a low level, where each layer manipulates raw token sequences as vectors without predefined structure or meaning. However, high-level tasks—such as constructing ASTs and performing type and symbol analysis—require handling complex, structured information that depends on long-range relationships and interactions across the input. Bridging the gap between this raw, unstructured processing and the structured, multi-step logic required for these tasks introduces significant difficulty. Compilers, for instance, typically rely on rule-based, step-by-step operations that are abstract and sequential, which transformers must simulate through their attention mechanisms and feedforward layers. The challenge is further compounded by the need to formally prove that transformers can handle such tasks efficiently and accurately, despite operating in a fundamentally different manner. To address these challenges, we propose a domain-specific language (DSL) called **Cybertron**, which allows us to *systematically prove* that transformers are capable of expressing complex algorithms while maintaining sufficient readability.

A key feature of **Cybertron** is its expressive type system, which provides strong correctness guarantees. The type system ensures that every value is strongly typed, making it easier to reason about function composition and ensuring the validity of our proofs. This type system is crucial for managing how transformers represent and manipulate both local and global types—where local types correspond to individual tokens and global types refer to sequences of tokens, encapsulating broader program information.

What transformers output (possibly in the intermediate layers) is a representation in sequences of vector of sequences of values in these types. As types are mathematically interpreted in this paper as a discrete subset of a vector space, **Cybertron** allows us to construct transformers with automatic value validity guarantees if the **Cybertron** code is type-correct.

In **Cybertron**, complex functions are broken down into “atomic” operations through propositions on function compositions and computation graphs (Propositions 11,13,14,2). It is straightforward to prove that these “atomic” operations are representable by transformers, either by feedforward layers or attention layers. For example:

- **Feedforward layers:** boolean operations like AND (Proposition 6), OR (Proposition 7), or NOT (Proposition 5), or operations over option types like `Option::or` (Proposition 9) being applied to each token in a sequence.
- **Attention layers:** operations that require information transmission between tokens such as `nearest_left` and `nearest_right` that collect for each token the nearest left/right non-nil information (Proposition 15).

This approach allows us to break down complex operations into primitive tasks that transformers can simulate. Feedforward layers handle local operations on individual tokens, while attention layers manage long-range dependencies and interactions between tokens, simulating the multi-step reasoning required for higher-level tasks.

**Cybertron**’s expressive type system and function composition framework help bridge the gap between the low-level processing transformers perform and the high-level reasoning necessary for complex tasks like compilation. For full details, including the mathematical foundations of **Cybertron**’s type system and function composition, see Appendix D.

## 6 COMPARISONS BETWEEN TRANSFORMERS AND RNN

Now we compare transformers and RNNs from both theoretical and empirical perspectives.

### 6.1 A LOWER BOUND FOR RNNs FOR TYPE CHECKING

Previously, it has shown that RNN is provably less parameter efficient than transformers for associative recall (Wen et al., 2024). Intuitively speaking, the type checking step covers associative recall. Based on this observation, we obtain the following lower bound for RNNs.

**Theorem 4.** *For  $L, D, H \in \mathbb{N}$ , for any RNN that represents a function that maps any token sequence of length  $L$  in  $\text{MiniHuskyAnnotated}_{D,H}$  with  $D, H = O(1)$  to its type errors represented as a sequence of values of type `Option<TypeError>`, then its state space size is at least  $\Omega(L)$ .*

Theorem 3 and Theorem 4 give a clear separation between transformers and RNNs in terms of the compilation capability. Specifically, if the input codes satisfy  $D, H \ll L$ , which is typically the case under the *clean code principle* (Martin, 2008), then transformers at most need  $O((\log L + D + H))$  number of parameters, which is significantly smaller what RNNs requires,  $\Omega(L)$ .

### 6.2 EMPIRICAL COMPARISON BETWEEN TRANSFORMERS AND RNNs

We validate our theoretical results by conducting experiments on synthetic data.

**Dataset construction.** The synthetic dataset is parameterized by  $n$  (the number of data pieces),  $f$  (the number of functions in a data piece),  $a$  (the maximum number of arguments of any function),  $c$  (the maximum number of function calls involved in any function),  $d$  (the minimum distance between the declaration and the first call of a function, as well as the minimum distance between its consecutive calls),  $v$  (the probability of using a variable in a function call), and  $e$  (the error rate of using an incorrect type in a function call).

The names of the functions are drawn randomly and uniquely from a list of English words. For each of the arguments of any function, its symbol is randomly drawn from another list of English words and its type is randomly drawn from  $\{\text{Int}, \text{Float}, \text{Bool}\}$ . All the called functions must be declared and not called by at least  $d$  functions ahead of the current one. For each argument of any function call, with probability  $v$ , the argument variable of the enclosing function is used regardless of its type, with probability  $(1 - v)(1 - e)$ , a literal of the correct type is used, and with probability  $(1 - v)e$  an incorrect type literal is used. For integers, the literals are from  $\{0, 1, \dots, 99\}$ ; for floats, the literals are from  $\{0.1, 1.1, \dots, 99.1\}$ ; for booleans, the literals are from  $\{\text{true}, \text{false}\}$ . The training dataset and evaluation dataset use disjoint lists for function names and argument symbols.

Below is a data piece with  $f = 10, a = 5, c = 5, d = 3, v = 0.2, e = 0.5$ :

```
1 fn rename_file ( i : Float , sum : Float ) { }
2 fn parse_data ( list : Int , value : Bool , stack : Float , k : Float , msg : Float ) { }
3 fn parse_json ( position : Bool ) { }
4 fn find_by_id ( error : Float ) { rename_file ( 60.1 , 94.1 ) ; }
5 fn merge ( group : Int , table : Float , error : Bool , count : Int ) { parse_data ( 7 ,
  false , 49.1 , 33.1 , 4.1 ) ; }
6 fn log_info ( val : Bool , m : Bool , xml : Float , path : Float ) { parse_json ( true ) ;
  }
7 fn process ( function : Int , value : Float , keys : Bool ) { find_by_id ( 88.1 ) ;
  rename_file ( value , 40.1 ) ; }
8 fn validate_response ( end : Int , z : Float , max : Bool ) { merge ( 1 , true , 27.1 , 72
  ) ; parse_data ( 11 , 85 , 35.1 , 14.1 , true ) ; }
9 fn print_message ( algorithm : Float ) { parse_json ( 92 ) ; log_info ( true , algorithm ,
  false , 26.1 ) ; }
10 fn print_help ( max : Bool , tree : Int , method : Int , item : Bool ) { process ( 25 , 28
  , false ) ; rename_file ( 48 , 80.1 ) ; }
```

**Model and training.** We use customized BERT models (Devlin et al., 2019) and bidirectional RNN models (Schuster & Paliwal, 1997) in our experiments. To control the model size (i.e., the number of trainable parameters), we adjust only the hidden sizes while keeping other hyperparameters constant. Detailed model specifications can be found in Table 1. For both transformers and RNNs, we use the hyperparameters listed in Table 2 in Appendix J during the training process.

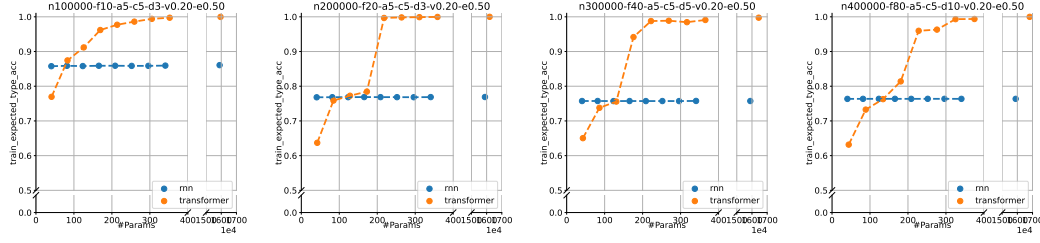


Figure 2: Figures depicting the accuracy of the expected type (see Section 5.3) across different models, measured by their number of trainable parameters, when trained on various datasets. Training accuracies are better indicators of the expressive power of the models (instead of generalizability) than evaluation accuracies. We also report evaluation accuracies in Appendix J.

**Results.** We experimented with multiple combinations of models (Table 1) and datasets (Table 2). For each combination, we conducted independent runs using a fixed set of  $k = 5$  random seeds. When plotting the figures, we took the top  $t = 5$  training/evaluation losses/accuracies from each run and averaged over all the  $k \times t$  values. We plotted separate figures for each dataset and separate sub-figures for each metric. In each sub-figure, the  $x$ -axis represents the number of trainable parameters, and the  $y$ -axis represents the averaged values. Results are shown in Figure 2. They demonstrate that customized BERT models are able to perform better at type checking than bidirectional RNN models when both scale up, corroborating our theories. Other results are in Appendix J.

## 7 CONCLUSION

We demonstrated that transformers can efficiently handle a number of syntactic and semantic analysis tasks in C-like languages, using Cybertron to prove their capacity for tasks like AST generation, symbol resolution, and type analysis. We show a theoretical advantage of transformers over RNNs, particularly in their ability to manage long-range dependencies with logarithmic parameter scaling. In a sense, transformers have the right inductive bias for language tasks. Our experiments confirmed these theoretical insights, showing strong performance on synthetic and real datasets, underscoring the expressiveness and efficiency of transformers in sequence-based learning.

## 8 ACKNOWLEDGEMENT

Xiyu Zhai acknowledges the support of NSF through awards DMS-2031883 and PHY-2019786. Liao Zhang acknowledges the ERC PoC project *FormalWeb3* no. 101156734 and the University of Innsbruck doctoral scholarship *promotion of young talent*.

## REFERENCES

- Cursor: Ai-powered code editor, 2024. URL <https://www.cursor.com/>. Accessed: September 29, 2024.
- Harold Abelson, Gerald Jay Sussman, and with Julie Sussman. *Structure and Interpretation of Computer Programs*. MIT Press/McGraw-Hill, Cambridge, 2nd edition, 1996. ISBN 0-262-01153-0.
- Agda Development Team. *Agda compilers manual v2.6.4.2*, 2024. URL <https://agda.readthedocs.io/en/v2.6.4.2/tools/compilers.html#javascript-backend>.
- Ekin Akyürek, Dale Schuurmans, Jacob Andreas, Tengyu Ma, and Denny Zhou. What learning algorithm is in-context learning? investigations with linear models. *arXiv preprint arXiv:2211.15661*, 2022.

- V Aho Alfred, S Lam Monica, and D Ullman Jeffrey. *Compilers principles, techniques & tools*. pearson Education, 2007.
- Cem Anil, Yuhuai Wu, Anders Andreassen, Aitor Lewkowycz, Vedant Misra, Vinay Ramasesh, Ambrose Slone, Guy Gur-Ari, Ethan Dyer, and Behnam Neyshabur. Exploring length generalization in large language models. *arXiv preprint arXiv:2207.04901*, 2022.
- Anyphere. Cursor, 2023. URL <https://www.cursor.com/features>.
- Simran Arora, Sabri Eyuboglu, Aman Timalsina, Isys Johnson, Michael Poli, James Zou, Atri Rudra, and Christopher R’e. Zoology: Measuring and improving recall in efficient language models. *ArXiv*, abs/2312.04927, 2023. URL <https://api.semanticscholar.org/CorpusID:266149332>.
- Yu Bai, Fan Chen, Huan Wang, Caiming Xiong, and Song Mei. Transformers as statisticians: Provable in-context learning with in-context algorithm selection. *arXiv preprint arXiv:2306.04637*, 2023.
- Boaz Barak, Benjamin Edelman, Surbhi Goel, Sham Kakade, Eran Malach, and Cyril Zhang. Hidden progress in deep learning: Sgd learns parities near the computational limit. *Advances in Neural Information Processing Systems*, 35:21750–21764, 2022.
- S. Bhattamishra, Kabir Ahuja, and Navin Goyal. On the ability and limitations of transformers to recognize formal languages. In *Conference on Empirical Methods in Natural Language Processing*, 2020a. URL <https://api.semanticscholar.org/CorpusID:22225236>.
- S. Bhattamishra, Arkil Patel, Phil Blunsom, and Varun Kanade. Understanding in-context learning in transformers and llms by learning to learn discrete functions. *ArXiv*, abs/2310.03016, 2023. URL <https://api.semanticscholar.org/CorpusID:263620583>.
- S. Bhattamishra, Michael Hahn, Phil Blunsom, and Varun Kanade. Separations in the representational capabilities of transformers and recurrent architectures. *ArXiv*, abs/2406.09347, 2024. URL <https://api.semanticscholar.org/CorpusID:270440803>.
- Satwik Bhattamishra, Kabir Ahuja, and Navin Goyal. On the ability and limitations of transformers to recognize formal languages. *arXiv preprint arXiv:2009.11264*, 2020b.
- Satwik Bhattamishra, Arkil Patel, and Navin Goyal. On the computational power of transformers and its implications in sequence modeling. *arXiv preprint arXiv:2006.09286*, 2020c.
- Jonathan Brennan and John Tracy Hale. Hierarchical structure guides rapid linguistic predictions during naturalistic listening. *PLoS ONE*, 14, 2019. URL <https://api.semanticscholar.org/CorpusID:260538292>.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde De Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, et al. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*, 2021.
- David Chiang, Peter A. Cholak, and Anand Pillay. Tighter bounds on the expressivity of transformer encoders. In *International Conference on Machine Learning*, 2023. URL <https://api.semanticscholar.org/CorpusID:256231094>.
- Chris Cummins, Volker Seeker, Dejan Grubisic, Baptiste Rozière, Jonas Gehring, Gabriele Synnaeve, and Hugh Leather. Meta large language model compiler: Foundation models of compiler optimization. *ArXiv*, abs/2407.02524, 2024. URL <https://api.semanticscholar.org/CorpusID:270924331>.
- Valentin David. *Language Constructs for C++-like languages*. PhD thesis, University of Bergen, 2009.
- Mostafa Dehghani, Stephan Gouws, Oriol Vinyals, Jakob Uszkoreit, and Łukasz Kaiser. Universal transformers. *arXiv preprint arXiv:1807.03819*, 2018.

- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding, 2019. URL <https://arxiv.org/abs/1810.04805>.
- Nai Ding, Lucia Melloni, Xing Tian, and David Poeppel. Rule-based and word-level statistics-based processing of language: insights from neuroscience. *Language, Cognition and Neuroscience*, 32: 570–575, 2017. URL <https://api.semanticscholar.org/CorpusID:46747073>.
- Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*, 2020.
- Jana Dunfield and Neelakantan R Krishnaswami. Sound and complete bidirectional typechecking for higher-rank polymorphism with existentials and indexed types. *Proceedings of the ACM on Programming Languages*, 3(POPL):1–28, 2019.
- Benjamin L Edelman, Surbhi Goel, Sham Kakade, and Cyril Zhang. Inductive biases and variable creation in self-attention mechanisms. In *International Conference on Machine Learning*, pp. 5793–5831. PMLR, 2022.
- N Elhage, N Nanda, C Olsson, T Henighan, N Joseph, B Mann, A Askell, Y Bai, A Chen, T Conerly, et al. A mathematical framework for transformer circuits. *Transformer Circuits Thread*, 2021.
- Husna Farooqui. The curry-howard correspondence. 2021. URL <https://api.semanticscholar.org/CorpusID:244268761>.
- Guhao Feng, Yuntian Gu, Bohang Zhang, Haotian Ye, Di He, and Liwei Wang. Towards revealing the mystery behind chain of thought: a theoretical perspective. *ArXiv*, abs/2305.15408, 2023. URL <https://api.semanticscholar.org/CorpusID:258865989>.
- David Flanagan. *JavaScript: The definitive guide: Activate your web pages*. " O'Reilly Media, Inc.", 2011.
- Bryan Ford. Parsing expression grammars: a recognition-based syntactic foundation. In *Proceedings of the 31st ACM SIGPLAN-SIGACT symposium on Principles of programming languages*, pp. 111–122, 2004.
- S. Frank, Rens Bod, and Morten H. Christiansen. How hierarchical is language use? *Proceedings of the Royal Society B: Biological Sciences*, 279:4522 – 4531, 2012. URL <https://api.semanticscholar.org/CorpusID:11969171>.
- Dan Friedman, Alexander Wettig, and Danqi Chen. Learning transformer programs. *ArXiv*, abs/2306.01128, 2023. URL <https://api.semanticscholar.org/CorpusID:259064324>.
- Shivam Garg, Dimitris Tsipras, Percy S Liang, and Gregory Valiant. What can transformers learn in-context? a case study of simple function classes. *Advances in Neural Information Processing Systems*, 35:30583–30598, 2022.
- Angeliki Giannou, Shashank Rajput, Jy yong Sohn, Kangwook Lee, Jason D. Lee, and Dimitris Papailiopoulos. Looped transformers as programmable computers. *ArXiv*, abs/2301.13196, 2023. URL <https://api.semanticscholar.org/CorpusID:256389656>.
- Google Deepmind. Ai achieves silver-medal standard solving international mathematical olympiad problems, July 2024. URL <https://deepmind.google/discover/blog/ai-solves-imo-problems-at-silver-medal-level/>.
- Qiuhan Gu. Llm-based code generation method for golang compiler testing. *Proceedings of the 31st ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, 2023. URL <https://api.semanticscholar.org/CorpusID:265509921>.

- Michael Hahn. Theoretical limitations of self-attention in neural sequence models. *Transactions of the Association for Computational Linguistics*, 8:156–171, 2019. URL <https://api.semanticscholar.org/CorpusID:189928186>.
- Thomas Hales, Mark Adams, Gertrud Bauer, Tat Dat Dang, John Harrison, Hoang Le Truong, Cezary Kaliszyk, Victor Magron, Sean McLaughlin, Tat Thang Nguyen, et al. A formal proof of the kepler conjecture. In *Forum of mathematics, Pi*, volume 5, pp. e2. Cambridge University Press, 2017.
- Sophie Hao, Dana Angluin, and Roberta Frank. Formal language recognition by hard attention transformers: Perspectives from circuit complexity. *Transactions of the Association for Computational Linguistics*, 10:800–810, 2022. URL <https://api.semanticscholar.org/CorpusID:248177889>.
- John Harrison. Hol light: An overview. In *International Conference on Theorem Proving in Higher Order Logics*, pp. 60–66. Springer, 2009.
- John Harrison, Josef Urban, and Freek Wiedijk. History of interactive theorem proving. In *Handbook of the History of Logic*, volume 9, pp. 135–214. Elsevier, 2014.
- Samy Jelassi, David Brandfonbrener, Sham M. Kakade, and Eran Malach. Repeat after me: Transformers are better than state space models at copying. *ArXiv*, abs/2402.01032, 2024. URL <https://api.semanticscholar.org/CorpusID:267406617>.
- Ralf Jung, Robbert Krebbers, Jacques-Henri Jourdan, Aleš Bizjak, Lars Birkedal, and Derek Dreyer. Iris from the ground up: A modular foundation for higher-order concurrent separation logic. *Journal of Functional Programming*, 28, 2018. URL <https://api.semanticscholar.org/CorpusID:2023423>.
- Steve Klabnik and Carol Nichols. *The Rust programming language*. No Starch Press, 2023.
- Shuai Li, Zhao Song, Yu Xia, Tong Yu, and Tianyi Zhou. The closeness of in-context learning and weight shifting for softmax regression. *arXiv preprint arXiv:2304.13276*, 2023.
- Zhiyuan Li, Hong Liu, Denny Zhou, and Tengyu Ma. Chain of thought empowers transformers to solve inherently serial problems. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=3EWTEy9MTM>.
- Valerii Likhoshervostov, Krzysztof Choromanski, and Adrian Weller. On the expressive power of self-attention matrices. *arXiv preprint arXiv:2106.03764*, 2021.
- David Lindner, János Kramár, Matthew Rahtz, Tom McGrath, and Vladimir Mikulík. Tracr: Compiled transformers as a laboratory for interpretability. *ArXiv*, abs/2301.05062, 2023. URL <https://api.semanticscholar.org/CorpusID:255749093>.
- Bingbin Liu, Jordan T. Ash, Surbhi Goel, Akshay Krishnamurthy, and Cyril Zhang. Transformers learn shortcuts to automata. *ArXiv*, abs/2210.10749, 2022. URL <https://api.semanticscholar.org/CorpusID:252992725>.
- Robert C. Martin. *Clean Code: A Handbook of Agile Software Craftsmanship*. Prentice Hall PTR, USA, 1 edition, 2008. ISBN 0132350882.
- Patrick Massot. Teaching mathematics using lean and controlled natural language. In *International Conference on Interactive Theorem Proving*, 2024. URL <https://api.semanticscholar.org/CorpusID:272330159>.
- The mathlib Community. The lean mathematical library. *Proceedings of the 9th ACM SIGPLAN International Conference on Certified Programs and Proofs*, 2019. URL <https://api.semanticscholar.org/CorpusID:204801213>.
- William Merrill, Ashish Sabharwal, and Noah A. Smith. Saturated transformers are constant-depth threshold circuits. *Transactions of the Association for Computational Linguistics*, 10:843–856, 2021. URL <https://api.semanticscholar.org/CorpusID:248085924>.

- Erik Nijkamp, Hiroaki Hayashi, Caiming Xiong, Silvio Savarese, and Yingbo Zhou. Codegen2: Lessons for training llms on programming and natural languages. *ICLR*, 2023.
- Catherine Olsson, Nelson Elhage, Neel Nanda, Nicholas Joseph, Nova DasSarma, Tom Henighan, Ben Mann, Amanda Askell, Yuntao Bai, Anna Chen, et al. In-context learning and induction heads. *arXiv preprint arXiv:2209.11895*, 2022.
- OpenAI. Openai o1 system card, September 2024a. URL <https://openai.com/index/openai-o1-system-card/>.
- OpenAI. Sora: Creating video from text, February 2024b. URL <https://openai.com/index/sora/>.
- Lawrence C Paulson. *Isabelle: A generic theorem prover*. Springer, 1994.
- Jorge Pérez, Javier Marinkovic, and Pablo Barceló. On the turing completeness of modern neural network architectures. *ArXiv*, abs/1901.03429, 2019. URL <https://api.semanticscholar.org/CorpusID:57825721>.
- Jorge Pérez, Pablo Barceló, and Javier Marinkovic. Attention is turing complete. *The Journal of Machine Learning Research*, 22(1):3463–3497, 2021.
- Benjamin C Pierce. *Types and programming languages*. MIT press, 2002.
- The Univalent Foundations Program. Homotopy type theory: Univalent foundations of mathematics. *arXiv preprint arXiv:1308.0729*, 2013.
- Reddit User. I think the main secret sauce of o1 is the data. [https://www.reddit.com/r/singularity/comments/lfi6yy9/i\\_think\\_the\\_main\\_secret\\_sauce\\_of\\_o1\\_is\\_the\\_data/](https://www.reddit.com/r/singularity/comments/lfi6yy9/i_think_the_main_secret_sauce_of_o1_is_the_data/), 2013. Accessed: 2024-09-28.
- Mike Schuster and Kuldip K Paliwal. Bidirectional recurrent neural networks. *IEEE transactions on Signal Processing*, 45(11):2673–2681, 1997.
- Lena Strobl. Average-hard attention transformers are constant-depth uniform threshold circuits. *ArXiv*, abs/2308.03212, 2023. URL <https://api.semanticscholar.org/CorpusID:260680416>.
- Victor Taelin. Ai and the future of coding. <https://medium.com/jonathans-musings/ai-and-the-future-of-coding-43caad31c3d3>, 2023a. Accessed: 2024-10-01.
- Victor Taelin. Agda to typescript compilation with sonnet-3.5, 2023b. URL <https://x.com/VictorTaelin/status/1837925011187027994>. Accessed: September 29, 2024.
- A Vaswani. Attention is all you need. *Advances in Neural Information Processing Systems*, 2017.
- Johannes Von Oswald, Eyvind Niklasson, Ettore Randazzo, João Sacramento, Alexander Mordvintsev, Andrey Zhmoginov, and Max Vladymyrov. Transformers learn in-context by gradient descent. *arXiv preprint arXiv:2212.07677*, 2022.
- Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Huai hsin Chi, and Denny Zhou. Self-consistency improves chain of thought reasoning in language models. *ArXiv*, abs/2203.11171, 2022. URL <https://api.semanticscholar.org/CorpusID:247595263>.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed Huai hsin Chi, F. Xia, Quoc Le, and Denny Zhou. Chain of thought prompting elicits reasoning in large language models. *ArXiv*, abs/2201.11903, 2022. URL <https://api.semanticscholar.org/CorpusID:246411621>.
- Gail Weiss, Yoav Goldberg, and Eran Yahav. Thinking like transformers. *ArXiv*, abs/2106.06981, 2021. URL <https://api.semanticscholar.org/CorpusID:235421630>.
- Kaiyue Wen, Xingyu Dang, and Kaifeng Lyu. Rnns are not transformers (yet): The key bottleneck on in-context retrieval. *ArXiv*, abs/2402.18510, 2024. URL <https://api.semanticscholar.org/CorpusID:268041425>.

- Shangda Wu, Xu Tan, Zili Wang, Rui Wang, Xiaobing Li, and Maosong Sun. Beyond language models: Byte models are digital world simulators. *ArXiv*, abs/2402.19155, 2024. URL <https://api.semanticscholar.org/CorpusID:268063492>.
- Shunyu Yao, Binghui Peng, Christos H. Papadimitriou, and Karthik Narasimhan. Self-attention networks can process bounded hierarchical languages. In *Annual Meeting of the Association for Computational Linguistics*, 2021. URL <https://api.semanticscholar.org/CorpusID:235166395>.
- Chulhee Yun, Srinadh Bhojanapalli, Ankit Singh Rawat, Sashank J Reddi, and Sanjiv Kumar. Are transformers universal approximators of sequence-to-sequence functions? *arXiv preprint arXiv:1912.10077*, 2019.
- Haoyu Zhao, Abhishek Panigrahi, Rong Ge, and Sanjeev Arora. Do transformers parse while predicting the masked word? *arXiv preprint arXiv:2303.08117*, 2023.
- Hattie Zhou, Arwen Bradley, Etai Littwin, Noam Razin, Omid Saremi, Josh Susskind, Samy Bengio, and Preetum Nakkiran. What algorithms can transformers learn? a study in length generalization. *ArXiv*, abs/2310.16028, 2023. URL <https://api.semanticscholar.org/CorpusID:264439160>.

## A TREE

Trees are one of the most fundamental objects to study in computer science. However, its exact definition differs for different domains. The trees used in “abstract syntax tree” (Section B) is more restrictive than that in mathematics, which we call “typed tree”, so that one can define recursive computation more rigorously.

### A.1 WHAT ARE TREES

Trees in data structures have slightly additional meaning as compared to trees in mathematics. In this paper, all trees are trees in data structures. For clarity, we lay down the precise definition of trees in data structure.

**Definition 3** (Tree). *A tree  $T$  is a set of nodes storing elements such that the nodes have a parent-child relationship that satisfies the following:*

- *If  $T$  is not empty, it has a special node called the **root** that has no parent.*
- *Each node  $v$  of  $T$  other than the root has a unique parent node  $w$ ; each node with parent  $w$  is a child of  $w$ .*

*We denote the nodes of  $T$  as  $N(T)$ .*

**Definition 4** (Recursive Definition of a Tree). *A tree  $T$  is either empty or consists of a node  $r$  (the root) and a possibly empty set of trees whose roots are the children of  $r$ .*

However, the above definition is too permissive. We shall define a typed version as follows:

**Definition 5** (Typed Tree). *A tree type consists of a set of values  $V$  and a set of relationships  $C \subseteq V \times \mathbb{N}$ , and a typed tree under this type is any tree  $T$  such that for each node, a value  $v \in V$  is assigned such that  $(v, n) \in C$  where  $n$  is the number of the children of the node.*

**All trees in this paper are typed.**

**Example 1** (Abstract syntax tree (AST) as Typed Tree). *Consider an AST for a simple arithmetic expression. Let the set of values  $V$  be:*

$$V = \{ \text{num}, \text{add}, \text{sub}, \text{mul}, \text{div} \}$$

*and the set of relationships  $C \subseteq V \times \mathbb{N}$  specify the allowed number of children for each value:*

$$C = \{ (\text{num}, 0), (\text{add}, 2), (\text{sub}, 2), (\text{mul}, 2), (\text{div}, 2) \}$$

*An example AST for the arithmetic expression  $(3 + 5) \times 2$  is the following typed tree:*

- *The root node is labeled **mul** (multiplication), and it has two children.*
  - *The left child is labeled **add** (addition), and it has two children:*
    - \* *The left child of **add** is labeled **num** with the value 3.*
    - \* *The right child of **add** is labeled **num** with the value 5.*
  - *The right child of **mul** is labeled **num** with the value 2.*

*This tree conforms to the typing rules because:*

- ***num** has 0 children,*
- ***add** has 2 children,*
- ***mul** has 2 children,*

*all of which satisfy the relationships in  $C$ .*

## A.2 REPRESENTATIONS OF TREES

It's also important to talk about tree representations. We are studying transformers, and then it's necessary to represent large trees as a sequence, otherwise the model dimension is not large enough to contain the information locally. Let's first talk about the classical **arena pattern** used in system programming for representing trees and we shall slightly adapt it to our use case for studying transformers.

**Arena Pattern.** To represent trees efficiently in memory, especially when trees are frequently modified (such as insertions or deletions of nodes), an arena pattern is often used. The arena pattern provides a way to manage memory allocation for tree structures, allowing for efficient memory usage and avoiding fragmentation. Here's how the arena pattern works in the context of tree representation:

**Definition 6** (Arena Pattern in Tree Representation). *In the arena pattern, a tree is represented by an array (or vector) of nodes, called an **arena**. Each node in the arena contains:*

- *An element or value stored in the node.*
- *References (often indices or pointers) to the node's children and possibly to its parent.*

*The key characteristics of the arena pattern are:*

- *Memory Contiguity: All nodes are stored contiguously in memory within the arena, which allows for efficient traversal and modification operations.*
- *Fixed Capacity: The arena has a fixed or dynamically resizable capacity, and nodes are added sequentially. This avoids the overhead of allocating individual nodes on the heap.*
- *Index-based References: Instead of using pointers, the nodes reference each other using indices within the array, which simplifies memory management and can lead to cache-friendly operations.*
- *Efficient Allocation and Deallocation: Nodes can be efficiently allocated and deallocated within the arena without requiring complex memory management techniques like garbage collection or reference counting.*

The arena pattern is particularly useful in scenarios where the structure of the tree is highly dynamic or when performance is critical. It allows for a simple and efficient way to manage and traverse trees without the typical overhead associated with more traditional pointer-based tree representations.

**Adaptations for Transformers** For transformers, inputs, intermediate values and outputs are all sequences. So the trees are represented as sequences of nodes with node reference representable by token position encoding. Based on the representation, transformers will be able to perform various kinds of recursive tree operations, as we shall see.

## B CONTEXT FREE GRAMMAR

In this section, we lay down the well-known definitions of context free grammar, derivations, and parse trees. To define an abstract syntax tree (AST), one commonly resorts to generation rules, such as context-free grammars (CFG) (Alfred et al., 2007) and parsing expression grammars (PEG) (Ford, 2004). In most cases, just generation rules themselves are not sufficient to define properly a language. Many practical languages like C and C++ cannot be solely described by these rules (David, 2009) so that they can reuse the limited set of special characters on the keyboard. Furthermore, semantic constraints like type correctness are intrinsically contextual and cannot be expressed through CFG or similar rules. However, CFG or other rules provide a valuable construct, the AST. With an AST, one can refine the language definition by putting restrictions on the syntax tree through tree operations. Effectively, a language can be seen as a subset of trees, not as a subset of strings. Semantic analysis like symbol resolution and type checking can be described effectively based on trees. In short, CFG standalone is hardly practical but it provides a useful and clear foundation to build definitions upon.

A context-free grammar (CFG) is defined as a 4-tuple  $G = (V, \Sigma, R, S)$ , where:

- $V$  is a finite set of variables (non-terminal symbols).
- $\Sigma$  is a finite set of terminal symbols, disjoint from  $V$ . Sequences of  $\Sigma$ , i.e., elements of  $\Sigma^*$  are called (literal) strings.
- $R \subset V \times (V \cup \Sigma)^*$  is a finite set of production rules, where each rule is of the form  $A \rightarrow \alpha$ , with  $A \in V$  and  $\alpha \in (V \cup \Sigma)^*$ .
- $S \in V$  is the start symbol.

Given a context-free grammar  $G = (V, \Sigma, R, S)$ , we define derivation as follows:

- A **derivation** is a sequence of steps where, starting from the start symbol  $S$ , each step replaces a non-terminal with the right-hand side of a production rule.
- Formally, we write  $u \Rightarrow v$  if  $u = \alpha A \beta$  and  $v = \alpha \gamma \beta$  for some production  $A \rightarrow \gamma$  in  $R$ , where  $\alpha, \beta \in (V \cup \Sigma)^*$  and  $A \in V$ .
- A **leftmost derivation** is a derivation in which, at each step, the leftmost non-terminal is replaced.
- A **rightmost derivation** is a derivation in which, at each step, the rightmost non-terminal is replaced.
- We denote a **derivation sequence** as  $S \Rightarrow^* w$ , where  $w \in \Sigma^*$  is a string derived from  $S$  in zero or more steps.

A **parse tree** (or **syntax tree**) for a context-free grammar  $G = (V, \Sigma, R, S)$  is a tree that satisfies the following conditions:

- The root of the tree is labeled with the start symbol  $S$ .
- Each leaf of the tree is labeled with a terminal symbol from  $\Sigma$  or the empty string  $\epsilon$ .
- Each internal node of the tree is labeled with a non-terminal symbol from  $V$ .
- If an internal node is labeled with a non-terminal  $A$  and has children labeled with  $X_1, X_2, \dots, X_n$ , then there is a production rule  $A \rightarrow X_1 X_2 \dots X_n$  in  $R$ .
- The yield of the parse tree, which is the concatenation of the labels of the leaves (in left-to-right order), forms a string in  $\Sigma^*$  that is derived from the start symbol  $S$ .

## C NEURAL ARCHITECTURES

In this section, we lay down the precise mathematical definitions of neural architectures we are going to use in our proof.

**Definition 7** (Single-Layer Fully Connected Network with  $4 \times$  Intermediate Space).

*Given model dimension  $d_{model}$ , a single-layer feed-forward network with an intermediate space expanded to 4 times the input dimension is a function from  $\mathbb{R}^{d_{model}}$  to  $\mathbb{R}^{d_{model}}$ , denoted by  $f_{fcn}$  and defined as follows:*

*given  $X \in \mathbb{R}^{d_{model}}$ , weights  $W_1 \in \mathbb{R}^{4d_{model} \times d_{model}}$ ,  $W_2 \in \mathbb{R}^{d_{model} \times 4d_{model}}$ , and biases  $B_1 \in \mathbb{R}^{4d_{model}}$ ,  $B_2 \in \mathbb{R}^{d_{model}}$ , the output  $f_{fcn}(X)$  is computed as:*

$$f_{fcn}(X) = W_2 \sigma_{\text{ReLU}}(W_1 X + B_1) + B_2,$$

*where  $\sigma_{\text{ReLU}} : \mathbb{R}^{4d_{model}} \rightarrow \mathbb{R}^{4d_{model}}$  is the Rectified Linear Unit activation function applied element-wise, defined by:*

$$\sigma_{\text{ReLU}}(z) = (\max(z_1, 0), \max(z_2, 0), \dots, \max(z_{4d_{model}}, 0))^T,$$

*for  $z = (z_1, z_2, \dots, z_{4d_{model}})^T \in \mathbb{R}^{4d_{model}}$ .*

The choice of a  $4\times$  intermediate space is common in practice, often used in Transformer architectures. Interestingly, this empirical choice turns out to have a useful theoretical property: it allows the network to express any affine transformation, as we'll see in the following proposition.

**Proposition 1.** *A single-layer fully connected network with a  $4\times$  intermediate space, as defined previously, can express any affine map from  $\mathbb{R}^{d_{\text{model}}}$  to  $\mathbb{R}^{d_{\text{model}}}$ .*

*Proof.* Let  $f : \mathbb{R}^{d_{\text{model}}} \rightarrow \mathbb{R}^{d_{\text{model}}}$  be any affine map given by  $f(X) = AX + b$ , where  $A \in \mathbb{R}^{d_{\text{model}} \times d_{\text{model}}}$  and  $b \in \mathbb{R}^{d_{\text{model}}}$ . We will construct weights  $W_1 \in \mathbb{R}^{4d_{\text{model}} \times d_{\text{model}}}$ ,  $W_2 \in \mathbb{R}^{d_{\text{model}} \times 4d_{\text{model}}}$  and biases  $B_1 \in \mathbb{R}^{4d_{\text{model}}}$ ,  $B_2 \in \mathbb{R}^{d_{\text{model}}}$  such that  $f_{\text{fcn}}(X) = f(X)$  for all  $X \in \mathbb{R}^{d_{\text{model}}}$ .

Define:

$$W_1 = \begin{pmatrix} I_{d_{\text{model}}} \\ -I_{d_{\text{model}}} \\ 0 \\ 0 \end{pmatrix}, \quad B_1 = 0 \in \mathbb{R}^{4d_{\text{model}}},$$

where  $I_{d_{\text{model}}}$  is the  $d_{\text{model}} \times d_{\text{model}}$  identity matrix, and 0 represents zero matrices of appropriate dimensions. Set:

$$W_2 = (A \quad -A \quad 0 \quad 0), \quad B_2 = b.$$

For any  $X \in \mathbb{R}^{d_{\text{model}}}$ , compute:

$$\begin{aligned} f_{\text{fcn}}(X) &= W_2 \sigma_{\text{ReLU}}(W_1 X + B_1) + B_2 \\ &= (A \quad -A \quad 0 \quad 0) \sigma_{\text{ReLU}} \left( \begin{pmatrix} X \\ -X \\ 0 \\ 0 \end{pmatrix} \right) + b \\ &= (A \quad -A \quad 0 \quad 0) \begin{pmatrix} \sigma_{\text{ReLU}}(X) \\ \sigma_{\text{ReLU}}(-X) \\ 0 \\ 0 \end{pmatrix} + b \\ &= A \sigma_{\text{ReLU}}(X) - A \sigma_{\text{ReLU}}(-X) + b. \end{aligned}$$

Note that  $\sigma_{\text{ReLU}}(X) - \sigma_{\text{ReLU}}(-X) = X$ , we have:

$$f_{\text{fcn}}(X) = AX + b = f(X).$$

Therefore, the network can represent any affine map from  $\mathbb{R}^{d_{\text{model}}}$  to  $\mathbb{R}^{d_{\text{model}}}$ .  $\square$

**Definition 8** (Single-Layer Feed Forward Network with  $4\times$  Intermediate Space). *Given model dimension  $d_{\text{model}}$  and position set Pos, the Transformer Feed Forward Network is a function  $f_{\text{ffn}} : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$  defined as follows:*

*For an input  $X \in \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ , the output  $f_{\text{ffn}}(X)$  is computed by applying the single-layer feed-forward network  $f_{\text{fcn}}$  (as defined previously) independently to each position:*

$$f_{\text{ffn}}(X)_p = f_{\text{fcn}}(X_p) \quad \forall p \in \text{Pos}$$

where  $X_p \in \mathbb{R}^{d_{\text{model}}}$  is the  $p$ -th row of  $X$ , corresponding to the  $p$ -th position in the input sequence.

Next, we define the attention mechanism, which is a key component of the Transformer architecture. This definition presents a hard attention layer with a simplified position encoding. We use hard attention here for theoretical simplicity, as it represents a discrete limit of the more commonly used soft attention mechanism. Hard attention forces the model to make a clear choice about which inputs to focus on, which can simplify analysis and provide clearer insights into the model's behavior. It can be viewed as the limiting case of soft attention as the temperature approaches zero, where the softmax operation becomes increasingly peaked and eventually converges to a one-hot vector.

**Definition 9** (Hard Attention Layer with Simplified Position Encoding). *Given model dimension  $d_{\text{model}}$ , number of heads  $H$ , and number of layers  $L$ , a transformer with simplified position encoding and hard attention is defined to be a function  $f_{\text{attn}} : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$  defined by*

$$\forall p \in \text{Pos}, f_{\text{attn}}(X)_p := W_O \text{Concat} \left( \text{Attn}^{(1)}(X)_p, \dots, \text{Attn}^{(H)}(X)_p \right), \quad (4)$$

where the  $h$ th attention head uses hard attention, defined as:

$$\text{Attn}^{(h)}(X)_p := \frac{1}{|S_p|} \sum_{p' \in S_p} V_{p'}^{(h)}, \quad (5)$$

where

- $W_O \in \mathbb{R}^{d_{\text{model}} \times d_{\text{model}}}$  are trainable parameters;
- $S_p = \arg \max_{p' \in \text{Pos}} \left( Q_p^{(h)\top} K_{p'}^{(h)} + \lambda^{(h)\top} \Psi_{p'-p} \right)$  with  $Q_p^{(h)}, K_p^{(h)}, V_p^{(h)}, \lambda^{(h)}, \Psi_q$  defined by
  - $Q_p^{(h)} = W_Q^{(h)} X_p, K_p^{(h)} = W_K^{(h)} X_p$  are vectors of dimension  $d_{\text{model}}/H$ , with trainable parameters  $W_Q^{(h)}, W_K^{(h)} \in \mathbb{R}^{(d_{\text{model}}/H) \times d_{\text{model}}}$ ;
  - $V_p^{(h)} = W_V^{(h)} X_p$  are vectors of dimension  $d_{\text{model}}/H$ , linear transformations of  $X_p$  with trainable parameters  $W_V^{(h)} \in \mathbb{R}^{(d_{\text{model}}/H) \times d_{\text{model}}}$ ;
  - $\lambda^{(h)} \in \mathbb{R}^2$  are constants depending only on head count  $h$ ;
  - $\Psi_q \in \mathbb{R}^2$  are 2-dimensional vectors depending on relative position  $q$  but not on head count  $h$ . It is explicitly defined as

$$\Psi_q = \begin{pmatrix} q \\ 1_{q>0} \end{pmatrix}. \quad (6)$$

This formulation allows for both past and future masking.

Having defined the basic components, we can now proceed to describe the full Transformer architecture. This definition builds upon the previously introduced concepts, incorporating them into a complete model structure.

**Definition 10** (Transformer). *A **Transformer** is a function  $f_{\text{tf}} : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$  that maps an input sequence to an output sequence through a series of layers, each consisting of a multi-head attention mechanism and a position-wise feed-forward network (MLP).*

Given:

- Input sequence  $X \in \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ , where Pos is the set of positions and  $d_{\text{model}}$  is the model dimension.
- Number of layers  $L$ .
- Number of attention heads  $H$ .

The Transformer computes the output  $Y = X^{(L)}$  through recursive application of attention and feed-forward layers:

- Initialization is given by:
- For each layer  $l = 1, 2, \dots, L$ :
  - Compute attention output:

$$\hat{X}^{(l)} = X^{(l-1)} + f_{\text{attn}}^{(l)} \left( X^{(l-1)} \right)$$

– Compute feed-forward output:

$$X^{(l)} = \hat{X}^{(l)} + f_{ffn}^{(l)}(\hat{X}^{(l)})$$

Here:

- $f_{attn}^{(l)}$  are **hard attention layers with simplified position encoding** as previously defined. It operates on  $X^{(l-1)}$  and produces an output in  $\mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ .
- $f_{ffn}^{(l)}$  are **feed-forward networks with  $4 \times$  intermediate space** as previously defined. It operates position-wise on  $\hat{X}^{(l)}$  and produces an output in  $\mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ .

*Remark 1.* For simplicity, we have omitted the Layer Normalization component typically present in Transformer architectures. This simplification allows us to focus on the core attention and feed-forward mechanisms while maintaining the essential structure of the Transformer.

We use  $\text{Tf}_{H,L}^{d_{\text{model}}}$  to denote the set of transformers of model size  $d_{\text{model}}$ , number of heads  $H$  and number of layers  $L$  as functions from  $\mathbb{R}^{d_{\text{model}}}$  to  $\mathbb{R}^{d_{\text{model}}}$ .

For purpose of proof, we shall also need residual multi-layer perceptron. Functions over local types are first represented by multi-layer perceptron, then by Proposition 2 applications of these functions over sequences can be representable by transformers. Residual multi-layer perceptron can be assembled through composition or computer graph, as we shall see.

Here’s the definition of a residual MLP Network.

**Definition 11** (Residual Multi-Layer Perceptron). A Residual Multi-Layer Perceptron (*ResMLP*) is a function  $f_{\text{resmlp}} : \mathbb{R}^{d_{\text{model}}} \rightarrow \mathbb{R}^{d_{\text{model}}}$  defined recursively by

$$X^{(0)} = X, \quad X^{(l)} = X^{(l-1)} + f_{\text{fcn}}(X^{(l-1)}), \quad l = 1, 2, \dots, L, \quad f_{\text{resmlp}}(X) = X^{(L)}$$

where  $X \in \mathbb{R}^{d_{\text{model}}}$  is the input vector,  $L$  is the total number of layers, and  $f_{\text{fcn}} : \mathbb{R}^{d_{\text{model}}} \rightarrow \mathbb{R}^{d_{\text{model}}}$  is the Single-Layer Fully Connected Network with  $4 \times$  Intermediate Space as previously defined in Definition 7.

We use  $\text{ResMlp}_L^{d_{\text{model}}} \subset \mathbb{R}^{d_{\text{model}}} \mathbb{R}^{d_{\text{model}}}$  to represent the set of residual MLPs with dimension  $d_{\text{model}}$  and  $L$  layers, as defined in Definition 11.

The following proposition is quite basic. It demonstrates that any function representable by a ResMLP can be applied position-wise by a Transformer.

**Proposition 2** (Position-wise ResMLP Application is Representable by Transformers). Let  $f : \mathbb{R}^{d_{\text{model}}} \rightarrow \mathbb{R}^{d_{\text{model}}}$  be a function representable by a Residual Multi-Layer Perceptron (*ResMLP*) as defined in Definition 11. Then the function  $F : \mathbb{R}^{\text{Pos} \times d_{\text{model}}} \rightarrow \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ , defined by applying  $f$  position-wise,

$$F(X)_p = f(X_p), \quad \forall p \in \text{Pos},$$

is representable by a Transformer as defined in Definition 10.

*Proof.* Since  $f$  is representable by a ResMLP with  $L$  layers, it is defined recursively by

$$X^{(0)} = X, \quad X^{(l)} = X^{(l-1)} + f_{\text{fcn}}(X^{(l-1)}) \quad \text{for } l = 1, \dots, L,$$

and

$$f(X) = X^{(L)},$$

where  $f_{\text{fcn}} : \mathbb{R}^{d_{\text{model}}} \rightarrow \mathbb{R}^{d_{\text{model}}}$  is the Single-Layer Fully Connected Network with  $4 \times$  intermediate space (Definition 7).

We construct a Transformer with  $L$  layers such that, for any input sequence  $X \in \mathbb{R}^{\text{Pos} \times d_{\text{model}}}$ , the output  $Y = f_{\text{tf}}(X)$  satisfies  $Y_p = f(X_p)$  for all  $p \in \text{Pos}$ .

To achieve this, we configure the Transformer so that the attention mechanism outputs zero at each layer. This can be done by setting the attention weights to zero, ensuring  $f_{\text{attn}}(X^{(l-1)}) = 0$ . Consequently, the update equations simplify to

$$\hat{X}^{(l)} = X^{(l-1)}.$$

We then set the feed-forward network  $f_{\text{ffn}}$  in the Transformer to have the same weights and biases as  $f_{\text{fcn}}$  in the ResMLP. The Transformer layer update becomes

$$X^{(l)} = \hat{X}^{(l)} + f_{\text{ffn}}(\hat{X}^{(l)}) = X^{(l-1)} + \left(f_{\text{fcn}}\left(X_p^{(l-1)}\right)\right)_{p \in \text{Pos}}.$$

This recursion matches that of the ResMLP applied position-wise to  $X$ . Therefore, after  $L$  layers, the Transformer output satisfies  $f_{\text{tr}}(X)_p = f(X_p)$  for all  $p \in \text{Pos}$ .

□

## D CYBERTRON

### D.1 INTRODUCTION

It’s often difficult to directly prove that transformers or in general other low level forms of computation can express complicated algorithms and even complex software. There are way too many details as compared with typical mathematical proofs in machine learning theory. Hence, we propose the domain specific language Cybertron, where we can systematically prove transformers can express complicated algorithms and complex software with sufficient readability.

(Note: Cybertron is fundamentally different from Mini-Husky! Mini-Husky is the target language that we want transformers to analyze yet Cybertron is the domain specific language we use to prove that transformers can do that.)

RASP (Weiss et al., 2021) is quite close to Cybertron in terms of its design purpose. However, Cybertron is more powerful with advanced algebraic type system, global and local function constructions, etc. These additional mechanisms replace a significant part of the chore in proofs with automatic type checking. Thus, using Cybertron one can argue operations more complicated than simple algorithms can be simulated by transformers.

In the broader perspective of computer science, it’s common to use code to prove things. In fact, in the formal verification community, mathematical proofs are viewed as a special case of a much larger universe of possible proof systems (mathlib Community, 2019; Massot, 2024) and constructive proof using code (Harrison et al., 2014; Farooqui, 2021; Jung et al., 2018) is far more applicable with great soundness to the most general settings. In our case, our code doesn’t serve as the whole proof but as an important part that contains most of the chores. However, it’s totally possible to build a full-fledged formalized proof, despite it might be too costly for a single paper to do.

Essentially, Cybertron works as follows:

1. one builds complicated functions from the composition of smaller functions. We have lemmas that prove that the composed functions are representable by certain architectures given that smaller functions are representable.
2. there is an algebraic type system and every value is strongly typed and immutable, making it highly readable and easy to reason about;
3. there is a distinction between global and local types/functions. Local types are those information hovered over a single token, and global types are sequences of local types, i.e., the collection of information over the whole token stream. One can define a global function by mapping a local function.
4. there are many functions that is defined externally, requiring external explanation that they can be represented by transformers.

It’s implemented as a subset of the Rust programming language that can be understood as computation graphs over sequences. It can be executed for testing purposes and we’ve tested our implementation for a range of inputs and validated its correctness.

## D.2 PHILOSOPHY: SEQUENTIAL REPRESENTATION OF EVERYTHING

Before going through the full details, let’s first talk about the fundamental philosophy behind transformer and Cybertron.

One of the fundamental reasons transformers can be easily adapted across multiple modalities, including NLP and CV, is their sequence-to-sequence operation. Everything can be represented as an arbitrary-length sequence. Texts are sequences of words, images are sequences of image patches, videos are sequences of spacetime patches (OpenAI, 2024b), and even graphs with sparse spatial structures can be represented as sequences of indexed nodes with additional information like parent node indices. Since inputs of various modalities can be cast into vector sequences, transformers can be applied to different domains without modifications to their architecture (Dosovitskiy et al., 2020).

Interestingly, this sequence-based thinking is not new. **We’ve actually been representing everything as sequences since the very early days of computer science.** This has been the foundation of how data is stored and processed in computers. However, sequence representations were traditionally viewed as low-level and sometimes inefficient for practical use, prompting the development of higher-level abstractions for programming. The rise of transformers, with their scalable learning capabilities, encourages us to reconsider the significance of sequence-based representations.

From a systems perspective, viewing everything as a sequence is the foundational approach in computer science. Data in a computer is stored as a continuous stream of bits. Whether it’s text, images, videos, or graphs, this data is represented in computer memory as an ordered sequence of bits. This aligns with how transformers handle different types of input by transforming them into sequences of vectors. Thus, the sequence-based operation of transformers mirrors the sequence-based representation of data in computer memory.

In essence, **if a data structure can be represented in computer memory using  $N$  bits, it can be processed as a sequence of bits of length  $N$ .** This natural sequence representation in memory is consistent with how transformers process data, which makes them particularly flexible across different modalities. For example, recent state-of-the-art approaches Wu et al. (2024) show that transformers can even be trained directly on raw bits of data, further emphasizing this connection.

Moreover, this sequence-based viewpoint offers fresh insights when applied to the domain of programming, particularly in areas such as code generation and analysis. With tools like ChatGPT and Copilot being widely used by developers, the impact of transformers on programming workflows is growing. Understanding the complexity of algorithms and programs expressed in sequence form becomes an interesting area of study, as it reveals new possibilities for how we approach computation.

In comparison to traditional systems like CPUs and human cognition, transformers are highly parallel but shallow in their operation. A transformer processes data in a fixed number of layers, while a CPU executes  $10^9$  cycles per second, and humans may take days to process information like reading a book. Transformers, therefore, represent a fundamentally different computational model that is worth studying further in the context of sequence-based operations.

**Example 2. Image to Sequence:** *In computer memory, an image is typically stored as a continuous block of pixel values, often in row-major order, where each pixel’s value is encoded as bits in a sequence. When a transformer processes an image, it divides the image into patches (e.g.,  $16 \times 16$  pixels), and each patch is flattened into a vector of pixel values. This creates a sequence of patches, where each patch corresponds to a vector. The way transformers represent these patches as a sequence closely aligns with how the image data is sequentially stored in computer memory.*

**Example 3. Video to Sequence:** *A video is stored in computer memory as a sequence of frames, where each frame is essentially an image. In a similar manner to images, these frames are stored as continuous pixel values. Transformers process videos by dividing the frames into spacetime patches, where each patch captures a small region of space over a short segment of time. These spacetime patches are flattened and arranged into a sequence for the transformer to process. The sequential ordering of these patches matches how video frames and pixel data are stored in computer memory.*

**Example 4. Graph to Sequence:** *In computer memory, a graph is typically stored using an adjacency list or adjacency matrix, where nodes and their connections (edges) are stored sequentially in a data structure. Transformers process graphs by representing each node and its features as a vector, and then creating a sequence of these vectors. The sequence may also encode additional*

information, such as the parent-child relationships between nodes. This sequence-based representation of graphs is consistent with how graph data is stored in memory, where nodes and edges are arranged in a structured order.

**Example 5. Text to Sequence:** Text is naturally stored in computer memory as a sequence of characters or words, where each character is encoded as a sequence of bits (such as ASCII or Unicode values). When transformers process text, they convert each word into a word embedding, which is a vector of real numbers. The sequence of word embeddings corresponds to the sequence of characters or words stored in memory. This natural sequential representation of text in both memory and transformers ensures efficient handling of linguistic data.

**Example 6. AST (Abstract Syntax Tree) to Sequence:** In computer memory, an abstract syntax tree (AST) is typically stored as a tree-like structure, where each node represents a component of the program (e.g., operators, variables, or statements). However, this tree can be linearized into a sequence by traversing it in a specific order (e.g., pre-order traversal). When transformers process an AST, they convert it into a sequence of tokens, where each token corresponds to a node in the tree. This sequential representation of the tree in transformers mirrors how the tree is stored as nodes and edges in memory, and how it can be flattened into a linear sequence.

In conclusion, the sequence-based representation in transformers is not just a novel approach for deep learning but is deeply rooted in how data has been stored and processed in computer memory since the early days of computing. This consistency between how data is stored in memory and how transformers process data as sequences is a key factor in their adaptability across different domains.

### D.3 LOCAL AND GLOBAL TYPES

Now we define the type foundation of Cybertron.

Types are fundamental objects for programming language theory. Here we use types to facilitate our proofs. Type signatures contain rich information that help guarantee correctness of the program. Here, we choose a mathematical definition of types that is most convenient for the discussion in this paper. We introduce the notion of “local type”. Roughly speaking, they are types without heap allocation and intended to be represented with  $\mathbb{R}^{d_{\text{model}}}$  over a single token. For more complicated heap-allocated data structures like trees, graphs, etc., we shall represent them by sequences of these “local type”s, which translate directly to vector sequences for transformers.

**Definition 12 (Local Type).** Given a base space  $B$  with at least two elements and a countably infinite identification space  $\Psi$ , a local type  $\mathcal{T}$  over  $B$  is a finite set  $S$  together with an embedding  $\phi$  from  $S$  to  $B^d$  and some fixed  $d \in \mathbb{N}$  and an identification  $\psi \in \Psi$ .

For convenience, define  $\text{Set}(\mathcal{T}) = S$ ,  $d_{\mathcal{T}} = d$  and  $\phi_{\mathcal{T}} = \phi$  and  $\psi_{\mathcal{T}} = \psi$ . And let  $0_B$ , and  $1_B$  be two different elements of  $B$ . And  $B^0 := \{0_B\}$  so that  $|B^i| = |B|^i$  holds for all  $i \in \mathbb{N}$ .

**Remark 2.** We need  $B$  to be at least size 2, so that  $B^d$  can be as large as we want for  $d$  large enough. For typical computer representation, we can take  $B$  to be  $\mathbf{2} = \{0, 1\}$ . For transformers or neural networks in general, we can take  $B$  to be  $\mathbb{R}$  if we ignore precision. If we don’t ignore precision,  $B$  should be some finite set of floating point numbers. Thus, we shall keep the generality of  $B$  throughout our discussion as all of these settings are important.

**Remark 3.** The role of identification  $\psi_{\mathcal{T}} \in \Psi$  is to make two types mathematically different even if they have the same underlying set, encoding dimension, and encoding. Basically we are establishing a specialized type of theory tailored towards the expressive power of transformers upon a foundation of intuitive set theory.

**Example 7 (Finite Set).** In mathematics, we have the finite set denoted by  $[n] = \{0, 1, \dots, n-1\}$ . Here we use a slightly different notation for a type with underlying set  $\llbracket n \rrbracket$  and some encoding.

**Example 8 (Position Encoding).** Position encoding can be viewed as the encoding of a type denoted by  $\text{Pos}(n)$  with the underlying set being  $[n]$  where  $n$  is the context length. Although it has the same underlying set as type  $\llbracket n \rrbracket$ , it is a different type for a different purpose and might have different encoding.

If  $B$  is  $\mathbb{R}$ , then the position encoding can be understood as the encoding of type  $\llbracket L \rrbracket$  where  $L$  is the context length. More explicitly, we have

$$\phi(x) = (e^{iL^{-i/d}x})_{i \in [d/2]}, \quad (7)$$

viewed as a  $d$  dimensional  $\mathbb{R}$ -vector through the natural conversion of  $\mathbb{C}$  to  $\mathbb{R}^2$ , since  $d$  is even.

It's too cumbersome to manually give the underlying set and the encoding. Here we introduce a classical concept from programming language theory [Program \(2013\)](#) that makes it super easy to construct new types and make things fairly readable.

**Definition 13** (Finite Algebraic Data Type, Mathematical Forms). *We define two ways of creating new types by combining existing types:*

1. *Sum type. Given types  $\mathcal{T}_i = (S_i, \phi_i, d_i)$  over base space  $B$  for  $i = 1, \dots, n$ , we define the sum type of  $\mathcal{T}_i$ , denoted by  $\sum_{i=1}^n \mathcal{T}_i$ , as follows,*

- *let  $S = (\{1\} \times S_1) \sqcup \dots \sqcup (\{n\} \times S_n)$ ;*
- *let  $d = d_{[n]} + \max_{i=1}^n d_i$ ;*
- *let  $\phi : S \rightarrow B^d$  be such that*

$$\forall i \in [n], s \in S_i, \phi((i, s)) = \phi_{[n]}(i) \oplus \phi_i(s) \in B^{d_{[n]}+d_i} \subseteq B^d. \quad (8)$$

*Note that  $|S| = \sum_{i=1}^n |S_i|$ , thus the name sum type.*

2. *Product type. Given Local Types  $\mathcal{T}_i = (S_i, \phi_i, d_i)$  over base space  $B$  for  $i = 1, \dots, n$ , we define the product type of  $\mathcal{T}_i$ , denoted by  $\prod_{i=1}^n \mathcal{T}_i$ , as follows,*

- *let  $S = S_1 \times \dots \times S_n$ ;*
- *let  $d = \sum_{i=1}^n d_i$ ;*
- *let  $\phi : S \rightarrow B^d$  be such that*

$$\forall s = (s_1, \dots, s_n) \in S, \phi(s) = \phi_1(s_1) \oplus \dots \oplus \phi_n(s_n) \in B^d. \quad (9)$$

*Note that  $|S| = \prod_{i=1}^n |S_i|$ , thus the name product type.*

Although we can define things and refer to things in terms of mathematical equations, it's sometimes cumbersome to do so. So we shall frequently refer to types using a programming language form, like `CybertronForm` or more complicated things like `Option<T>` a builtin generic type.

**Definition 14** (Unit Type). *The unit type is a type with  $S = \{0\}$  and  $\phi : S \rightarrow B^0, 0 \mapsto 0_B$ . In Cybertron, it's denoted as `()`.*

**Definition 15** (Array Type). *Given a type  $\mathcal{T}$ , the array type of  $\mathcal{T}$  with length  $\ell \in \mathbb{N}$  is the type with  $S = S(\mathcal{T})^\ell$ ,  $d = \ell d_{\mathcal{T}}$  and  $\phi : S \rightarrow B^{\ell d_{\mathcal{T}}}, (s_1, \dots, s_\ell) \mapsto \phi_{\mathcal{T}}(s_1) \oplus \dots \oplus \phi_{\mathcal{T}}(s_\ell)$ . It's denoted by  $\mathcal{T}^\ell$ . In Cybertron, it's denoted as `[T;N]`.*

**Definition 16** (Vector Type of Finite Capacity). *Given a type  $\mathcal{T}$ , the vector type of finite capacity of  $\mathcal{T}$  with maximal length  $\ell \in \mathbb{N}$  is the type with  $S = \bigsqcup_{i=1}^\ell \text{Set}(\mathcal{T})^i$ ,  $d = d_{[ \ell ]} + \ell d_{\mathcal{T}}$  and  $\phi : S \rightarrow B^{d_{[ \ell ]} + \ell d_{\mathcal{T}}}, (s_1, \dots, s_i) \mapsto \phi_{[ \ell ]}(i) \oplus \phi_{\mathcal{T}}(s_1) \oplus \dots \oplus \phi_{\mathcal{T}}(s_i) \oplus 0_B \oplus \dots \oplus 0_B$  with just enough number of copies of  $0_B$  such that the dimensionality matches. It's denoted by  $\mathcal{T}^{\leq \ell}$ . In cybertron, it's denoted as `BoundedVec<T,N>`.*

However, it's cumbersome and obtuse to define and operate in mathematical forms only. So we shall give a definition closer to actual programming that is more convenient and easy to read.

**Definition 17** (Finite Algebraic Data Type, the Code Forms). *We define two ways to create new types:*

1. *Enum type. An enum type is the sum type of a finite set of variant types. Each variant type is associated with a different identifier and can be*

- *unit like, a unit type;*
- *struct like, a product of several types, each called a field of the variant, and associated with an identifier;*
- *tuple like, a product of several types, each called a field of the variant, but not associated with an identifier.*

Syntactically, an enum type is specified as follows,

```

1  enum <type-name> {
2      <identifier> {           // 1st variant, struct like
3          <identifier>: <type>, // 1st named field of 1st variant
4          <identifier>: <type>, // 2nd named field of 1st variant
5          ...
6      },
7      <identifier> {           // 2nd variant, struct like
8          <identifier>: <type>, // 1st field of 2nd variant
9          ...
10     },
11     <identifier> (           // 3rd variant, tuple like
12         <type>,              // 1st tuple field of 3rd variant
13         <type>,              // 2nd tuple field of 3rd variant
14         ...
15     ),
16     <identifier>,           // 4th variant, unit like
17     ...
18 }

```

For example,

```

1  enum Expr {
2      Variable(IdentToken),    // 1st variant, tuple like
3      Binary {                 // 2nd variant, struct like
4          lopd: ExprId,
5          opr: BinaryOprToken,
6          ropd: ExprId,
7      },
8      Prefix {                 // 3rd variant, struct like
9          opr: PrefixOprToken,
10         opd: ExprId,
11     },
12     Suffix {                  // 4th variant, struct like
13         opd: ExprId,
14         opr: SuffixOprToken,
15     },
16     Panic,                   // 5th variant, unit like
17 }

```

2. Struct type. A struct type is just the product type of

```

1  struct <type-name> {
2      <identifier>: <type>,
3      <identifier>: <type>,
4      ...
5  }

```

```

1  struct A {
2      a: i32
3  }

```

To show how convenient this is, we can define the very useful option type as follows,

**Definition 18** (Option type). For a local type  $T$ , we can define an option type as

```

1  enum Option<T> {
2      Some (T),
3      None
4  }

```

**Definition 19** (Global Types). Global types are defined to be sequences of local types.

**Example 9** (Representation of Graphs). Graphs can be represented as sequences of its nodes. We can use position index to use as node references.

#### D.4 COMPUTATION GRAPH

For convenience, we shall use computation graph as a vehicle to describe complicated computation processes. Computation graph is close to actual computation process and one can derive an understanding of the computation difficulty from the graph's mathematic properties (width, depth, etc.)

**Definition 20** (Directed Simple Graph). *A directed simple graph  $G$  is a pair  $(V, E)$  where  $V$  is a finite set, and  $E \subseteq V \times V$  is called edges.*

In the following, we shall simplify the "directed simple graph" to just graph.

**Definition 21** (Computation Graph). *A computation graph is an acyclic directed graph  $G = (V, E)$  with additional structures:*

1. *for each vertex  $v \in V$ , there is an associated type, denoted by  $T_v$ ;*
2. *for each vertex  $v \in V$  with a positive number of incoming edges, let  $v_1, \dots, v_n$  be the other vertices for the incoming edges, then there is an associated function  $f_v$  from  $T_{v_1} \times \dots \times T_{v_n}$  to  $T_v$ .*

A computation graph naturally generates a function from **source vertices** to **sink vertices**. Let  $v_1^{\text{in}}, \dots, v_n^{\text{in}}$  be the set of vertices with no incoming edges, and let  $v_1^{\text{out}}, \dots, v_m^{\text{out}}$  be the set of vertices with no outgoing edges. Then we can construct a function from  $T_{v_1^{\text{in}}} \times \dots \times T_{v_n^{\text{in}}}$  to  $T_{v_1^{\text{out}}} \times \dots \times T_{v_m^{\text{out}}}$  in the following obvious manner:

1. let  $(x_1, \dots, x_n) \in T_{v_1^{\text{in}}} \times \dots \times T_{v_n^{\text{in}}}$  be an input;
2. for each  $v_i^{\text{in}}$ , assign it with value  $x_i$ ;
3. for each vertex  $v \in V$  with all its incoming vertices  $v_1, \dots, v_l$  assigned with a value, assign it with the value  $f_v(x_{v_1}, \dots, x_{v_l})$  where  $x_{v_i}$  denotes the value assigned to  $v_i$ ;
4. repeat the process until all vertices are assigned a value, then take  $(x_{v_1^{\text{out}}}, \dots, x_{v_m^{\text{out}}})$  as the output.

Our goal is to make a hypothesis class using the above graph. To control the statistical and computational complexity, we put restrictions on the choice of  $T_v$  and  $f_v$ , as follows:

**Definition 22** (Restricted Computation Graph). *Let  $\mathcal{U}$  be a set of types, and for any  $A, B \in \mathcal{U}$ , there is a set of functions  $\text{Mor}(A, B)$  from  $A$  to  $B$ . We require that  $T_v, T_v^{\text{in}} \in \mathcal{U}$  and  $f_v \in \text{Mor}(T_v^{\text{in}}, T_v)$  where  $T_v^{\text{in}} := \prod_{v'v \in E} T_{v'}$ . We also require that the underlying graph  $G$  satisfies certain conditions (width, depth, etc.)*

**Definition 23** (Restricted Computation Graph Of Sequences). *Let  $\mathcal{U}$  be a universe such that for a set of types  $\mathcal{U}_0$  all types in  $\mathcal{U}$  are of the form  $A^*$  for some type  $A \in \mathcal{U}_0$ , and  $\text{Mor}(A^*, B^*)$  are functions that preserve sequence lengths.*

Given a restriction, the class of functions generated by restricted computation graphs is the central object to study in this paper. We shall use an even more restricted computation graph of sequences. We shall argue about the class of functions formed that

1. it's rich enough to contain many interesting operations including SQL, compiler (type inference, static analysis)
2. it's computationally reasonable, and can be represented by transformers with pragmatic bounds
3. it has a reasonable statistical complexity

As a corollary, our theories suggest that transformers can possibly learn to do many interesting things with reasonable computational and statistical complexity.

To our knowledge, this is the first theoretical paper that gives pragmatic optimistic bounds for the power of transformers in a wide range of meaningful language tasks.

Now we introduce graph-theoretical measures that will play key roles in our new complexity theory.

The most basic one is the following:

**Definition 24** (Depth of Graph). *The depth of a computation graph is defined to be the length of the longest path, denoted by  $d_G$ .*

For convenience, we define the following vertex-wise depth.

**Definition 25** (Depth of Graph Vertex). *The depth of a vertex  $v$  of a computation graph is defined as the length of the longest path with end  $v$ , denoted by  $d_v$ .*

The smaller  $d_G$  is, the more parallel the computation is.

However, we shall discuss a more nuanced measure, containment, as follows:

## D.5 FUNCTIONS OVER LOCAL TYPES

**Definition 26** (Functions over Local Types). *Given Local Types  $\mathcal{T}, \mathcal{R}$ , the functions from  $\mathcal{T}$  to  $\mathcal{R}$  are defined to be just the functions from  $\text{Set}(\mathcal{T})$  to  $\text{Set}(\mathcal{R})$ .*

*Remark 4.* The domains and codomains are all finite sets, so there aren't many constraints we want to enforce here. Basically, these are "discrete" functions.

**Definition 27** (Functions over Algebraic Data Types). *Let  $\mathcal{T}, \mathcal{S}_1, \dots, \mathcal{S}_m, \mathcal{R}$  be Local Types, and suppose that  $\mathcal{T}$  is an algebraic data type, then we can construct functions from  $\mathcal{T} \times \mathcal{S}_1 \times \dots \times \mathcal{S}_m$  to  $\mathcal{R}$  as follows,*

1. *suppose that  $\mathcal{T}$  is the sum type of  $\mathcal{T}_1, \dots, \mathcal{T}_n$ . Then given functions  $f_i : \mathcal{T}_i \times \mathcal{S}_1 \times \dots \times \mathcal{S}_m$  for  $i = 1, \dots, n$ , we can construct a function  $f$ , by letting*

$$f((i, t), s_1, \dots, s_m) = f_i(t, s_1, \dots, s_m), \quad (10)$$

*for each  $t \in \text{Set}(\mathcal{T}_i), s_1 \in \text{Set}(\mathcal{S}_1), \dots, s_m \in \text{Set}(\mathcal{S}_m)$ .*

*(Note that we use the pair  $(i, t)$  because the underlying set of  $\mathcal{T}$  is  $\bigsqcup_{i=1}^n \{i\} \times \text{Set}(\mathcal{T}_i)$ .)*

2. *suppose that  $\mathcal{T}$  is the product type of  $\mathcal{T}_1, \dots, \mathcal{T}_n$ . Then given a function  $f_* : \mathcal{T}_1 \times \dots \times \mathcal{T}_n \times \mathcal{S}_1 \times \dots \times \mathcal{S}_m$  for  $i = 1, \dots, n$ , we can construct a function  $f$ , by letting*

$$f((t_1, \dots, t_n), s_1, \dots, s_m) = f_*(t_1, \dots, t_n, s_1, \dots, s_m), \quad (11)$$

*for each  $t \in \text{Set}(\mathcal{T}_i), s_1 \in \text{Set}(\mathcal{S}_1), \dots, s_m \in \text{Set}(\mathcal{S}_m)$ .*

It is not enough to just mathematically construct. We should also discuss how neural networks can represent these functions. We define the representation of functions over Local Types formally as follows:

**Definition 28** (Representation of Functions over Local Types Using Multi-Layer Perceptions). *Let  $\mathcal{T}, \mathcal{R}$  be Local Types. Given a function  $f$  from  $\mathcal{T}$  to  $\mathcal{R}$ , we say it is representable by MLP of dimension  $d \geq \max\{d_{\mathcal{T}}, d_{\mathcal{R}}\}$  and number of layers  $L$ , if there exists  $\tilde{f} \in \text{ResMlp}_L^d$  such that*

$$\iota_1 \circ \phi_{\mathcal{R}} \circ f = \tilde{f} \circ \iota_2 \circ \phi_{\mathcal{T}}, \quad (12)$$

*where  $\iota_1 : \mathbb{R}^{d_{\mathcal{R}}} \rightarrow \mathbb{R}^d$  and  $\iota_2 : \mathbb{R}^{d_{\mathcal{T}}} \rightarrow \mathbb{R}^d$  are the canonical embeddings by putting zeros to fit the dimensionalities.*

Here are some trivially true facts:

**Proposition 3.** *[Identities are Representable] For any Local Type  $\mathcal{T}$ , the identity map  $\text{Id}_{\mathcal{T}}$  is representable in  $\text{ResMlp}_1^{d_{\mathcal{T}}}$ .*

*Proof.* Just take  $W_0^{(1)} = I_d, W_1^{(1)} = W_2^{(2)} = 0, B_1^{(1)} = B_2^{(2)} = 0$ .  $\square$

**Proposition 4.** *[Equality is Representable] The equality function for any local type  $\mathcal{T}$  is representable in  $\text{ResMlp}_2^{2d}$ , where  $d$  is the encoding dimension of  $\mathcal{T}$ .*

*Proof.* Let  $x, y \in \mathcal{T}$  be the inputs. We encode them as  $\phi_{\mathcal{T}}(x), \phi_{\mathcal{T}}(y) \in \mathbb{R}^d$ . The equality function can be represented as:

$$f_{\text{eq}}(x, y) = \min \left( 1, A \sum_{i=1}^d |\phi_{\mathcal{T}}(x)_i - \phi_{\mathcal{T}}(y)_i| \right),$$

where  $A$  is a large enough positive constant such that the RHS is either 1 or 0.

This can be implemented in two-layer ResMLP with dimension  $2d$ .  $\square$

**Proposition 5.** [Boolean NOT is Representable] The Boolean NOT function is representable in  $\text{ResMlp}_1^1$ .

*Proof.* It’s affine.  $\square$

**Proposition 6.** [Boolean AND is Representable] The Boolean AND function is representable in  $\text{ResMlp}_1^2$ .

*Proof.* Represent each Boolean value as a binary flag within a 1-dimensional vector. Then AND is just taking the minimum. By  $\min(a, b) = b - \sigma_{\text{ReLU}}(b - a)$ , we’re done.  $\square$

**Proposition 7.** [Boolean OR is Representable] The Boolean OR function is representable in  $\text{ResMlp}_1^2$ .

*Proof.* Represent each Boolean value as a binary flag within a 1-dimensional vector. Then OR is just taking the maximum. By  $\max(a, b) = a + \sigma_{\text{ReLU}}(b - a)$ , we’re done.  $\square$

**Proposition 8.** [THEN\_SOME is Representable] The function  $\text{Bool}::\text{then\_some} : \text{Bool} \times T \rightarrow \text{Option } T$  returns *Some t* if the boolean is true and *None* otherwise. This function is representable in  $\text{ResMlp}_1^{d+1}$ .

*Proof.* Encode the boolean as a binary flag in a  $(d + 1)$ -dimensional vector, where the first component indicates the boolean value and the remaining  $d$  components hold the value of type  $T$ . The residual MLP  $f_{\text{resmlp}}$  constructs the output  $\text{Option } T$  by assembling the flag and the value split into positive and negative parts influenced by the flag:

$$f_{\text{resmlp}}(X) = \begin{pmatrix} x_1 \\ \sigma_{\text{ReLU}}(x_{2:d+1} - Ax_1) - \sigma_{\text{ReLU}}(-x_{2:d+1} - Ax_1) \end{pmatrix}.$$

Here,  $A$  is a vector of dimension  $d$  with all entries positive and large enough to ensure proper thresholding. Specifically, each entry of  $A$  should be larger than the maximum absolute value that can be represented in the corresponding dimension of type  $T$ . This ensures that when  $x_1 = 1$ , the subtraction  $x_{2:d+1} - A$  will always be negative, and when  $x_1 = 0$ , it will not affect the value.

When the flag is true ( $x_1 = 1$ ),  $\sigma_{\text{ReLU}}(x_{2:d+1} - A) = 0$  and  $\sigma_{\text{ReLU}}(-x_{2:d+1} - A)$  retains the negated value, resulting in *Some t*. When the flag is false ( $x_1 = 0$ ), both ReLU terms preserve the value, yielding *None*. Thus,  $f_{\text{resmlp}}$  effectively implements  $\text{Bool}::\text{then\_some}$  within a single layer of the MLP.  $\square$

**Proposition 9.** [Option Or is Representable] Let  $T$  be a local type, let  $\text{Option}::\text{or}$  be the function that maps two values  $a, b$  of type  $\text{Option } T$  to a value  $c$  of type  $\text{Option } T$  such that  $c$  is equal to  $a$  when  $a$  is not none, and equal to  $b$  otherwise. Then  $\text{Option}::\text{or}$  is representable in  $\text{ResMlp}_1^{2(d+1)}$ .

*Proof.* Each  $\text{Option } T$  is represented as a  $(d + 1)$ -dimensional vector, where the first component is a binary flag indicating the presence (1 for *Some*, 0 for *None*), and the remaining  $d$  components encode the value. Given inputs  $a, b \in \text{Option } T$ , the residual MLP  $f_{\text{resmlp}}$  processes the concatenated vector

$$X = \begin{pmatrix} a_{\text{flag}} \\ a_{\text{val}} \\ b_{\text{flag}} \\ b_{\text{val}} \end{pmatrix}.$$

The MLP is designed to separate  $b_{\text{val}}$  into positive and negative parts ( $b_+$ ,  $b_-$  respectively) influenced by  $a_{\text{flag}}$ . Specifically, it computes:

$$\begin{aligned} f_{\text{resmlp}}(X) &= a_{\text{val}} + \sigma_{\text{ReLU}}(b_+ - Aa_{\text{flag}}) - \sigma_{\text{ReLU}}(b_- - Aa_{\text{flag}}) \\ &= a_{\text{val}} + \sigma_{\text{ReLU}}(b_{\text{val}} - Aa_{\text{flag}}) - \sigma_{\text{ReLU}}(-b_{\text{val}} - Aa_{\text{flag}}), \end{aligned} \quad (13)$$

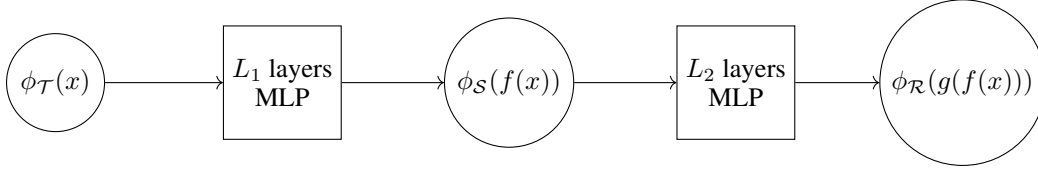


Figure 3: Transformation from  $\phi_{\mathcal{T}}(x)$  to  $\phi_{\mathcal{S}}(f(x))$  to  $\phi_{\mathcal{R}}(g(f(x)))$  with MLP layers.

where  $A$  is a vector with large positive entries that ensures the ReLU activation zeroes out the non-selected parts based on the flag. When  $a_{\text{flag}} = 1$ , the terms involving  $b$  are suppressed, resulting in  $c = a$ . Conversely, when  $a_{\text{flag}} = 0$ , the positive part of  $b$  remains, effectively selecting  $b$ . Thus,  $f_{\text{resmlp}}$  accurately implements the `Option::or` function, demonstrating that it is representable within  $\text{ResMlp}_1^{2(d+1)}$ .  $\square$

**Proposition 10** (Field Access Is Representable in ResMlp). *For algebraic data type, either struct field access, enum discriminator, and variant field access can be represented in  $\text{ResMlp}_1^d$  where  $d$  is the encoding dimension.*

*Proof.* Obvious because these operations are linear.  $\square$

**Proposition 11.** [Composition of Functions Representable in ResMlp] *For local types  $\mathcal{T}, \mathcal{S}, \mathcal{R}$ , with maps  $f : \mathcal{T} \rightarrow \mathcal{S}$  and map  $g : \mathcal{S} \rightarrow \mathcal{R}$  representable in  $\text{ResMlp}_{L_1}^{d_1}$  and  $\text{ResMlp}_{L_2}^{d_2}$  respectively. Then  $g \circ f$  is representable in  $\text{ResMlp}_{L_1+L_2}^{\max\{d_1, d_2\}}$ .*

*Proof.* Obvious by using the first  $L_1$  layers to map from  $\mathcal{T}$  to  $\mathcal{S}$  and using the rest  $L_2$  layers to map from  $\mathcal{S}$  to  $\mathcal{R}$ . The process can be visualized as in Figure 3.  $\square$

**Proposition 12.** [Computation Graph of Functions Representable in ResMlp] *Let  $\mathcal{G}$  be a computation graph, with each vertex  $v$  being of some local type  $\mathcal{T}_v$ , and the construction functions are representable in  $\text{ResMlp}_{L_v}^{d_v}$ . For convenience, if  $v$  is a source vertex,  $d_v$  is defined to be the encoding dimension of  $\mathcal{T}_v$  and  $L_v = 0$ . Then the function induced by the computation graph is representable in  $\text{ResMlp}_{\sum_{v \in \mathcal{G}} d_v}^{\sum_{v \in \mathcal{G}} d_v}$ .*

*Proof.* We construct a global residual multi-layer perceptron (ResMLP) that simulates the computation graph  $\mathcal{G}$  by aggregating and updating the states of all vertices simultaneously. Let  $D = \sum_{v \in \mathcal{G}} d_v$  be the total dimension, where  $d_v$  is the dimension associated with vertex  $v$ . The global ResMLP will have a depth of  $\text{Depth}(\mathcal{G})(\max_v L_v + 1)$ .

Consider the concatenated state vector  $X^{(t)} \in \mathbb{R}^D$ , which is a concatenation of the states of all vertices:

$$X^{(t)} = \left( X_v^{(t)} \right)_{v \in \mathcal{G}},$$

where  $X_v^{(t)} \in \mathbb{R}^{d_v}$  is the state of vertex  $v$  at layer  $t$ .

Initialization occurs at depth zero, corresponding to the source vertices of the computation graph. The state vector  $X^{(0)}$  is set by assigning the input vectors to the source vertices and initializing all other vertices to zero. Formally, if  $V_0$  denotes the set of source vertices, then:

$$X_v^{(0)} = \begin{cases} x_v & \text{if } v \in V_0, \\ 0 & \text{otherwise,} \end{cases}$$

where  $x_v \in \mathbb{R}^{d_v}$  is the input to source vertex  $v$ . Because  $X_v^{(0)}$  is of dimensionality  $d_v$  equal to the encoding dimension, this agrees with our convention for representing functions over local types.

We proceed inductively over the depth levels of the computation graph. For each depth level  $k = 1, 2, \dots, \text{Depth}(\mathcal{G})$ , we perform the following steps in the global ResMLP.

1. Input Aggregation Layer. We apply a linear transformation to gather the outputs from the predecessor vertices of each vertex at depth  $k$  and feed them as inputs to these vertices. Specifically, we define a linear mapping  $W_{\text{agg}}^{(k)} \in \mathbb{R}^{D \times D}$  such that:

$$\tilde{X}^{(t_k)} = W_{\text{agg}}^{(k)} X^{(t_{k-1})},$$

where  $t_{k-1}$  is the layer after processing depth  $k-1$ , and  $\tilde{X}^{(t_k)}$  is the aggregated input for the vertices at depth  $k$ . The matrix  $W_{\text{agg}}^{(k)}$  rearranges and combines the outputs from predecessor vertices to provide the correct inputs to each vertex at depth  $k$ . Specifically, for each vertex  $v$  at depth  $k$ , and for each predecessor  $u$  of  $v$  in the computation graph, the matrix  $W_{\text{agg}}^{(k)}$  contains entries that copy the output of  $u$  into the input positions of  $v$ . All other entries in  $W_{\text{agg}}^{(k)}$  are set to zero or identity as appropriate.

2. Local Computation Layers. For each vertex  $v$  at depth  $k$ , we simulate its local ResMLP of depth  $L_v$ . Since the depths  $L_v$  may vary, we pad the local ResMLPs to have a uniform depth  $L = \max_v L_v$  by adding identity mappings where necessary. The updates for vertex  $v$  are computed as:

$$\begin{aligned} X_v^{(t_k+1)} &= \tilde{X}_v^{(t_k)} + f_{\text{fcn}_v}(\tilde{X}_v^{(t_k)}), \\ X_v^{(t_k+k')} &= X_v^{(t_k+k'-1)} + f_{\text{fcn}_v}(X_v^{(t_k+k'-1)}), \quad \text{for } k' = 2, \dots, L_v, \\ X_v^{(t_k+k')} &= X_v^{(t_k+k'-1)}, \quad \text{for } k' = L_v + 1, \dots, L. \end{aligned}$$

Here,  $f_{\text{fcn}_v}$  denotes the single-layer fully connected network (as per Definition 7) for vertex  $v$ .

3. State Update. After completing the local computations for depth  $k$ , we update the global state vector  $X^{(t_k+L)}$  by concatenating the updated states of all vertices:

$$X^{(t_k+L)} = \left( X_v^{(t_k+L)} \right)_{v \in \mathcal{G}}.$$

The total number of layers added for depth  $k$  is  $L+1$ , accounting for the input aggregation layer and the  $L$  layers simulating the local ResMLPs.

By repeating this process for each depth level  $k = 1, 2, \dots, \text{Depth}(\mathcal{G})$ , we simulate the entire computation graph within a global ResMLP of depth  $\text{Depth}(\mathcal{G})(\max_v L_v + 1)$ .

Lastly, we use the final layer to perform a linear mapping so that the output is in the correct linear representation, clearing out the intermediate values.

Therefore, the function computed by the global ResMLP is equivalent to the function induced by the computation graph  $\mathcal{G}$ , and it is representable in  $\text{ResMlp}_{\text{Depth}(\mathcal{G})(\max_v L_v + 1)}^D$ .  $\square$

*Remark 5.* We only prove things around MLPs here. Later, we shall show that this will imply that the induced map operation over sequences can be represented by transformers.

## D.6 FUNCTIONS OVER GLOBAL TYPES

The task we want transformers to express is too complicated to be cleanly described in one shot. So we introduce the following lemma to significantly simplify things. The lemma shall be useful for our future papers on this topic.

**Proposition 13.** *[Composition of Functions Representable in Transformers] For local types  $\mathcal{T}, \mathcal{S}, \mathcal{R}$ , with maps  $f : \mathcal{T}^* \rightarrow \mathcal{S}^*$  and  $g : \mathcal{S}^* \rightarrow \mathcal{R}^*$  representable in  $\text{Tf}_{H_1, L_1}^{d_1}$  and  $\text{Tf}_{H_2, L_2}^{d_2}$  respectively. Then the composition  $g \circ f$  is representable in  $\text{Tf}_{\max\{H_1, H_2\}, L_1 + L_2}^{\max\{d_1, d_2\}}$ .*

*Proof.* This is basically the same as the proof of Proposition 11.  $\square$

**Proposition 14.** [Computation Graphs of Functions Representable in Transformers] Suppose we have a computation graph  $G = (V, E)$  with types  $\mathcal{T}_v = T_v^*$  together with encoding map  $\psi_v : T_v \rightarrow \mathbb{R}^{d_v}$  and decoding map  $\phi_v : \mathbb{R}^{d_v} \rightarrow T_v$ , satisfying  $\phi_v \circ \psi_v \equiv \text{id}_{T_v}$ , and there exists some positive integer  $d_0$  such that for each  $v \in V$ ,  $f_v$  can be represented in

$$\text{Tf}_{H_v, L_v}^d$$

Let  $f$  be the function generated by the computation graph. Then  $f$  can be represented in  $\text{Tf}_{H, L}^d$  if  $d \geq \sum_v d_v + Hd_0$ ,  $L \geq \frac{|G|}{H} + d_G$  where  $d_G$  is the depth of the graph.

*Remark 6.* This doesn't really cover the above. The bound in Proposition 14 isn't always tight for model dimension when the computation graph is deep and Proposition 13 complements it.

*Proof.* WLOG, assume that  $d = \sum_{v \in V} d_v + Hd_0$ . Then

$$\mathbb{R}^d = \underbrace{\left( \bigoplus_{v \in V} \mathbb{R}^{d_v} \right)}_C \oplus \underbrace{\left( \bigoplus_{h \in [H]} \mathbb{R}^{d_0} \right)}_A. \quad (14)$$

Here  $C$  stands for "cache" used for storing computed values, and  $A$  stands for "active" used for storing intermediate computation results.

Make an order of all the nodes in the graph, say  $V = \{v_1, \dots, v_{|G|}\}$  such that  $\text{Depth}(v_i) \leq \text{Depth}(v_j)$  if  $i \leq j$ .

We now imagine the transformer computation process as gradually evaluating the value of each vertex, starting from  $v_1$  to  $v_{|G|}$ . Every  $\max_v L_v$  layers form a layer group, and after each layer group, at most  $H$  vertices are assigned values. The equation 14 implies that we have enough memory to cache the computed values and intermediate values in small transformers.

Now let this process continue until we compute all the values. It must be finite because after each layer group, at least one of the vertices is computed. But this bound is too loose. We claim the following:

**Claim:** the number of layer groups where less than  $H$  vertices are assigned values is smaller than  $\text{Depth}(G)$ .

**Sketch of Proof of Claim:** for any layer group where less than  $H$  vertices are assigned, all the vertices that aren't assigned after this layer group must have larger depth than any vertices that are assigned values before this layer group, otherwise such a vertex can be evaluated in this layer group. Define the depth of any layer group to be the smallest depth of vertices evaluated in this layer group. Then for any unsatiated layer group, it must have a larger depth than the previous layer group. But depth can only increase  $\text{Depth}(G)$  times, thus there are at most  $\text{Depth}(G)$  unsatiated layer groups.

**Proof of Claim:** let  $V_1, \dots, V_l$  be the vertices evaluated at each layer group. Note that  $l$  is a different symbol than  $L$  and means that the number of layer groups rather than the number of layers.

For convenience, let  $D_i$  be the minimum of the depths of vertices in  $V_i$ .

Suppose that the  $i$ th layer group is unsatiated, then  $i < l$ . We want to show that  $D_i < D_{i+1}$ . Suppose otherwise, i.e.,  $D_i = D_{i+1}$ . Because the  $i$ th layer group is unsatiated, for any  $v \in V_{i+1}$ ,  $v$  must have dependencies that haven't been evaluated before the  $i$ th layer group. Choose  $v_0 \in V_i, v_1 \in V_{i+1}$  such that  $\text{Depth}(v_0) = \text{Depth}(v_1) = D_i = D_{i+1}$ . Note that any dependency of  $v_1$  must have smaller depths than  $v_0$ , then must have already been evaluated before the  $i$ th layer group. Contradiction!

Now given the claim, we have that for all but at most  $\text{Depth}(G)$  choices of  $i = 1, \dots, l$ , we have  $|V_i| = H$ , then we have

$$|G| = \sum_{i=1}^l |V_i| \geq (l - \text{Depth}(G))H \quad (15)$$

Then  $l \leq \frac{|G|}{H} + \text{Depth}(G)$ .

Then  $L \leq l \cdot \max_{v \in G} L_v = \left( \frac{|G|}{H} + \text{Depth}(G) \right) \max_{v \in G} L_v$ .

□

**Proposition 15.** [Nearest Left/Right] For any local type  $T$ , consider the function that maps a sequence of type  $\text{Option}\langle T \rangle$  to nearest left/right neighbors that are not none. It’s representable in  $\text{Tr}_{1,1}^{d+1}$

*Proof.* There is only one layer and one head needed, so we can omit the layer and head index.

WLOG, we consider the nearest left case.

We just need to make the attention exponential look like this:

$$Q_p^\top K_{p'} + \lambda \Psi_{p'-p} = a_{\text{flag}, p'} - 1_{p'-p > 0}, \quad (16)$$

where  $a_{\text{flag}, p'} \in \{0, 1\}$  indicates whether the value at position  $p'$  is some or none.

We set  $V_{p'}$  to represent the value of type  $\text{Option}\langle T \rangle$ .

For the starter token  $p_0$ , we make it such that

$$Q_p^\top K_{p_0} + \lambda \Psi_{p_0-p} = 1, \quad (17)$$

and

$$V_{p_0} = \mathbf{0}, \quad (18)$$

so that when there are no some to the left, it will give us none. □

**Proposition 16.** [Nearest Two Left/Right] For any local type  $T$ , consider the function that maps a sequence of type  $\text{Option}\langle T \rangle$  to nearest two left/right neighbors that are not none. It’s representable in  $\text{Tr}_{O(1), O(1)}^{O(d)}$  where  $d$  is the encoding dimension of  $T$ .

*Proof.* We can utilize Proposition 15 and 14.

The nearest two left or right is equivalent to first computing the nearest left/right, and then packing them together into one and compute its nearest left/right. The process is represented by a small constant computation graph, then we’re done. □

## D.7 SYNTAX AND SEMANTICS OF CYBERTRON

Having laid the necessary mathematical foundation behind **Cybertron**, we now turn to explaining its surface—its **syntax** and **semantics**. Cybertron serves as a syntax sugar for expressing local and global computation graphs, which are the vehicles used to demonstrate the expressive power of transformers. In Cybertron, computations are divided into two layers: the **local world** and the **global world**. These layers play distinct but complementary roles in constructing computation graphs.

### D.7.1 LOCAL WORLD

The **local world** in Cybertron corresponds to the feed-forward layers of a transformer, focusing on computations over **local types**. Local types represent individual tokens or data points, and computations in this world handle operations on tokens independently of their surrounding context.

**Data Types.** Local types in Cybertron include basic types such as `Bool`, `Idx`, `Pos`, `Fin<n>`, `BoundedVec<T, N>`, etc. These types are essential for building local computation graphs that operate over individual tokens. Compound types, like **structs** and **enums**, can also be defined for more complex token representations. These types serve as the building blocks for the local computation graphs that transform data at the token level.

```

1 struct Node {
2   id: Idx,
3   position: Pos,
4 }
5
6 enum Operation {
7   Add {
8     lhs: Pos,
9     rhs: Pos,
10  },
11   Multiply {
12     factor: Pos,
13  },
14 }

```

**Functions.** Functions in the local world define operations upon information over individual tokens. These operations form nodes in the local computation graphs. For instance, operations like binary or unary expressions, conditionals, and matches on token types are transformed into computation graphs by handling each individual token’s data.

```

1 fn process_ast(ast: AstData) -> Option<Role> {
2   match ast {
3     AstData::LetInit { pattern, initial_value, .. } => {
4       Some(Role::LetStmt { pattern, initial_value })
5     }
6     AstData::Defn { keyword, ident, .. } => {
7       Some(match keyword {
8         DefnKeyword::Struct => Role::StructDefn(ident),
9         DefnKeyword::Enum => Role::EnumDefn(ident),
10        DefnKeyword::Fn => Role::FnDefn(ident),
11      })
12     }
13     _ => None,
14   }
15 }

```

**Control Flow.** In the local world, control flow structures such as `if` and `match` are transformed into computation graphs by treating each branch or arm as an expression that returns an `Option` based on conditions. These `Option` values are then combined using the `Option::or` function. According to **Proposition 9**, `Option::or` maps two `Option<T>` values and returns the first non-`None` value, or the second one otherwise. This allows conditional branches to be represented in computation graphs as sequential option evaluations, where the first matching condition provides the result.

#### D.7.2 GLOBAL WORLD

The **global world** extends beyond individual tokens to sequences of tokens, represented as **global types**. These global types are denoted as `Seq<T>`, where `T` is a local type. The global world represents the full transformer, focusing on operations involving sequences of tokens, including variable definitions, expressions involving variable references, and function calls.

**Variable Definitions.** Variables in the global world are defined using global types, which represent sequences of local tokens. These definitions correspond to nodes in the global computation graph.

**Expressions.** Expressions in the global world consist of references to variables or function calls. Since the global world operates over sequences of tokens, these expressions are translated into sequence-level operations in the computation graph.

**Function Calls.** Function calls are key elements of the global world. They are represented by applying global functions to sequences of tokens. Cybertron provides **map functions** to elevate local functions to global functions by mapping them across sequences. Additionally, **attention methods** like `nearest_left` and `nearest_right` handle dependencies between tokens in the sequence by identifying relationships based on their positions.

```

1 let result = seq_of_values.nearest_left();

```

In the global world, computation graphs are built by composing map functions and attention methods. These graphs, unlike those in the local world, do not include control flow mechanisms.

## D.8 DYCK LANGUAGE

This section demonstrates how the **local world** in Cybertron operates over token-level computations and how the **global world** handles sequence-level operations. We use a Dyck language example to explain the interactions between these two worlds. The example processes a sequence of delimiters (like parentheses) and checks for matching pairs.

**Local World.** In Cybertron, the **local world** operates on individual tokens. Here, the local types are simple, such as `Delimiter` and `PreAst`, which represent information associated with individual tokens. These types allow for token-level operations like comparisons and transformations.

We define a `struct` to represent a delimiter and an `enum` to classify delimiters as either left or right. These definitions reflect local types, as they hold information over a single token.

```
1 // Define a struct 'Delimiter' that wraps a 'u8' value.
2 #[derive(Debug, Clone, Copy, PartialEq, Eq)]
3 pub struct Delimiter(u8);
4
5 // Define an enum 'PreAst' which represents a left or right delimiter.
6 #[derive(Debug, Clone, Copy, PartialEq, Eq)]
7 pub enum PreAst {
8     LeftDelimiter(Delimiter),
9     RightDelimiter(Delimiter),
10 }
```

Here, the local types `Delimiter` and `PreAst` define operations upon individual tokens, representing fundamental units of the computation graph at the local level. The local world is responsible for handling these small, token-level computations independently of the global sequence.

**Global World.** In the **global world**, Cybertron operates on sequences of tokens, treating the collection of local types as a single unit of computation. The global world introduces global types such as `Seq<Option<PreAst>>`, which represents a sequence of optional delimiters. The global world handles sequence-level operations by applying functions like `nearest_left` and `nearest_right` to capture the relationships between tokens in the sequence.

The following function operates on a sequence of `PreAst`, reducing matched pre-asts. The recursive application of `step` gives us the classifier for Dyck language.

```
1 fn step(pre_asts: Seq<Option<PreAst>>) -> Seq<Option<PreAst>> {
2     let pre_asts_nearest_left = pre_asts.nearest_left();
3     let pre_asts_nearest_right = pre_asts.nearest_right();
4     step_aux.apply(pre_asts_nearest_left, pre_asts, pre_asts_nearest_right)
5 }
```

**Local Worlds.** The `step_aux` function matches tokens based on their nearest neighbors within the sequence, eliminating pre-asts if a match is found.

```
1 fn step_aux(
2     pre_ast_nearest_left: Option<(Idx, PreAst)>,
3     pre_ast: Option<PreAst>,
4     pre_ast_nearest_right: Option<(Idx, PreAst)>
5 ) -> Option<PreAst> {
6     match pre_ast? {
7         PreAst::LeftDelimiter(delimiter) => match pre_ast_nearest_right {
8             Some(Idx, PreAst::RightDelimiter(delimiter1)) if delimiter1 == delimiter =>
9                 None,
10                _ => pre_ast,
11             },
12         PreAst::RightDelimiter(delimiter) => match pre_ast_nearest_left {
13             Some(Idx, PreAst::LeftDelimiter(delimiter1)) if delimiter1 == delimiter =>
14                 None,
15                _ => pre_ast,
16             },
17     }
18 }
```

In this example, the global function `step` uses `nearest_left` and `nearest_right` to capture sequence-level dependencies, while the local function `step_aux` uses conditional logic to check for matching pairs of delimiters. The local world handles token-level logic, while the global world coordinates operations across the entire sequence. This separation reflects how Cybertron handles computations at different levels of granularity.

Thus, this example illustrates how Cybertron leverages both the local and global worlds to build comprehensive computation graphs in a convenient, comprehensive yet rigorous manner. The local world performs individual tokenwise operations, and the global world captures relationships between tokens in a sequence, demonstrating how Cybertron enables transformers to express complex computations.

## E MINI-HUSKY DETAILS

Here's the BNF grammar of the Mini-Husky language:

```

<ast> ::= <literal>
        | <ident>
        | <prefix>
        | <binary>
        | <suffix>
        | <delimited>
        | <separated_item>
        | <call>
        | <let_init>
        | <if_stmt>
        | <else_stmt>
        | <defn>

<literal> ::= ...
<ident> ::= ...
<prefix> ::= <prefix_opr> <ast>
<binary> ::= <ast> <binary_opr> <ast>
<suffix> ::= <ast> <suffix_opr>
<delimited> ::= <left_delimiter> <separated_item>* <right_delimiter>
<separated_item> ::= [<ast>] <separator>
<call> ::= <ast> <left_delimiter> <ast>* <right_delimiter>
<let_init> ::= let <ast>
<if_stmt> ::= if <ast> <delimited>
<else_stmt> ::= <if_stmt> else (<delimited> | <else_stmt>)
<defn> ::= <defn_keyword> <ident> <ast>
<prefix_opr> ::= + | - | ! | ...
<binary_opr> ::= + | - | * | / | ...
<suffix_opr> ::= ++ | -- | ...
<left_delimiter> ::= '(' | '[' | '{'
<right_delimiter> ::= ')' | ']' | '}'
<separator> ::= , | ;
<defn_keyword> ::= def | fn | ...

```

Below is a sample piece of codes:

```

1 struct Dog { weight: f32, .. }
2
3 fn see_vet(dog: Dog) -> f32 {
4     assert dog.weight < 100;
5     let mut fee = dog.weight * 10.0;

```

```

6     fee +=100.0;
7     return fee
8 }

```

It should be noted that the above is not the full story. There are additional constraints put on the ASTs. However, these can be easily implemented as tree functions that are easy for transformers to express. As we are focusing on higher level language processing capabilities, we ignore the details here.

Additionally, we need to require that for semantic correctness, we must have proper symbol resolution and type correctness.

## E.1 ADDITIONAL DETAILS ABOUT COMPILER TASKS.

The outputs of the tasks are defined using Cybertron as follows:

- The construction of AST task’s final output is the collection all AST nodes. More concretely, the output is a sequence of `Option<Ast>` with length equal to the input token sequence’s length, where `Option<Ast>` denoted the type `Ast` will a null value added and `Ast` is the type storing the information of a node, including its parent, and its data of type `AstData`. In Cybertron, we define `Ast` and `AstData` explicitly as follows:

```

1  /// Represents a node in an Abstract Syntax Tree (AST).
2  ///
3  /// Each 'Ast' node has a reference to its parent node (if any) and holds
4  /// the associated syntax data (such as expressions, statements, or other
5  /// constructs defined in the 'AstData' enum).
6  pub struct Ast {
7      /// The index of the parent node in the AST, if it exists.
8      ///
9      /// - 'Some(Idx)': The node has a parent, and 'Idx' represents its position.
10     /// - 'None': The node is the root or does not have a parent.
11     pub parent: Option<Idx>,
12     /// The data associated with this AST node.
13     pub data: AstData,
14 }
15
16 /// Enumeration representing different types of Abstract Syntax Tree (AST) nodes
17 pub enum AstData {
18     /// Represents a literal value (e.g., integer, string)
19     Literal(Literal),
20     /// Represents an identifier (e.g., variable name)
21     Ident(Ident),
22     /// Represents a binary expression (e.g., 'x + y', 'a * b')
23     Binary {
24         /// Index of the left operand
25         lopr: Idx,
26         /// Operator in the binary expression (e.g., '+', '*')
27         opr: BinaryOpr,
28         /// Index of the right operand
29         ropd: Idx,
30     },
31     ... // other variants
32 }

```

- The output of the **symbol resolution** task is the collection of symbol resolution results on all applicable tokens. More concretely, the output is a sequence of values of type `Option<SymbolResolution>` where `Option<SymbolResolution>` is the type `SymbolResolution` with a null value added for non-applicability and `SymbolResolution` is the type storing the result of the symbol resolution, being either a success with a resolved symbol of type `Symbol` or a failure with an error of type `SymbolResolutionError`. In Cybertron, we define `SymbolResolution` explicitly as follows:

```

1  // an enum type definition, basically a tagged union type
2  pub enum SymbolResolution {
3      Ok(Symbol), // enum type variant for success with a resolved symbol
4      Err(SymbolResolutionError), // enum type variant for failure with an error
5  }

```

- The **type analysis** task’s final output is the collection of all type errors. More concretely, the output is a sequence of `Option<TypeError>`, where `Option<TypeError>` denoted the type `TypeError` will a null value added and `TypeError` is the type storing the information of a type error. The position of type errors agrees with the source tokens leading to these errors. In Cybertron, we define `TypeError` explicitly as follows:

```
1 // This enum represents various kinds of type errors
2 pub enum TypeError {
3     // This variant indicates a type mismatch
4     // 'expected' is the type that was anticipated
5     // 'actual' is the type that was encountered
6     TypeMismatch { expected: Type, actual: Type },
7 }
```

One can expand the definition to include other kinds of type errors.

- (1) *Type definition*. Types are either identified uniquely by a single identifier like `<identifier>`, or builtin generic types `Option<<identifier>>` or `Vec<<identifier>>`. Users can define custom types without generics like the following (f32 means float32 and i32 means int32 below):

```
1 struct Dog { weight: f32 }
2
3 enum Animal {
4     Dog,
5     Cat,
6 }
```

This part is actually a part of the AST task and type definition is a variant of the `AstData` type:

```
1 /// Enumeration representing different types of Abstract Syntax Tree (AST) nodes
2 pub enum AstData {
3     ...
4     /// Represents a function or variable definition
5     ///
6     /// # defn
7     ///
8     Defn {
9         /// The keyword in the definition (e.g., 'fn', 'enum')
10        keyword: DefnKeyword,
11        /// Index of the identifier in the definition
12        ident_idx: Idx,
13        /// The identifier being defined (e.g., function name, variable name)
14        ident: Ident,
15        /// Index of the content or body of the definition
16        content: Idx,
17    },
18 }
```

- (2) *Type specification*. Each appeared variable has a unique type, either by specification or speculation. All parameters of a function must be specified explicitly by users. Variables defined by let statements might or might not be specified, as follows:

```
1 fn f(a: i32) { // type of 'a' must be specified
2     let x: i32 = a; // type of 'x' specified
3     let y = a; // type of 'y' unspecified
4 }
```

The return type of functions must be specified. The field type of structs and enum variants must be specified. the type of expressions of function calls and field access will be determined correspondingly.

The output of the task is the collection of all type signatures, represented as a sequence of values of type `Option<TypeSignature>` where `TypeSignature` is the type holding the essential information of type specifications. In Cybertron, `TypeSignature` is defined as,

```
1 pub struct TypeSignature {
2     pub key: TypeSignatureKey,
3     pub ty: Type,
4 }
5
6 pub enum TypeSignatureKey {
```

```

7   FnParameter { fn_ident: Ident, rank: Rank },
8   FnOutput { fn_ident: Ident },
9   StructField { ty_ident: Ident, field_ident: Ident },
10  }

```

(3) *Type inference*. As discussed above, not all variables have their types specified.

```

1  fn f() {
2      let x: i32 = 1;
3      let y = x;
4      let z = y;
5  }

```

In the above code, `1` is an ambiguous literal that can be of type `i32`, `i64`, `u32`, `u64`, etc, and the types of `y` and `z` is not specified. However, one easily sees that there exists one and only one choice of the types of `1`, `y`, and `z` such that the whole code is type correct. Utilizing this property, the user can opt out of a significant portion of type specification, achieving static guarantees.

**A Type Inference Algorithm:** For simplicity, we shall prove transformers can implement a simple type inference algorithm: we maintain a table of type assignments for variables. We update the entries of the table by means of reduction, i.e., assuming the whole code is correctly typed and infer more and more unspecified types until we encounter errors or all types are inferred. The process is largely parallel, and we call the number of rounds needed the depth of type inference.

In the above code, the first round, we determine that the type of both `1` and the type of `y` are equal to the type of `x` which is `i32`. But we have no way to determine the type of `z` because the type of `y` is unknown at the first round. In the second round, `z` can be determined to be of type `i32` because the type of `y` is already inferred.

The output of the task is the collection all types inferred, represented as a sequence of values of type `Option<TypeInference>` where `TypeInference` is the type holding the inferred type. In Cybertron, `TypeSignature` is defined as,

```

1  pub struct TypeInference {
2      pub ty: Type,
3  }

```

## F TRANSFORMER AST PROOF

### F.1 HIGH LEVEL OVERVIEW

Here we give the full details of the proof of transformers being able to parse ASTs.

On a high level, we are going to see the parsing of ASTs as an assembly process. First, we immediately get the atomic ones, like identifiers, literals, etc. Then we assembly all composite ASTs with enough precedence until all tokens are consumed. We can prove that at the  $n$ th round, all ASTs with depth no more than  $n$  are already constructed. In the process, we must keep track of the unconsumed tokens and newly constructed ASTs (to be consumed as children for new ASTs in the next round, as we are going bottom up). We use `pre_ast`s to denote all the unconsumed tokens and newly constructed ASTs and use `asts` to denote all the constructed(allocated) ASTs. For correctness guarantees, we give detailed type specifications for tokens, ASTs, and PreASTs as follows.

We define the Token type as follows:

```

1  /// The 'Token' enum represents the various types of tokens that can be
2  /// identified during the lexical analysis phase of a compiler. Each variant
3  /// corresponds to a specific category of token that can be encountered
4  /// in the source code.
5  pub enum Token {
6      /// A literal value, which can be a number, string, or other primitive type.
7      Literal(Literal),
8      /// A reserved keyword in the language, such as 'if', 'else', 'while', etc.
9      Keyword(Keyword),
10     /// An identifier, typically representing variable names, function names,

```

```

11  /// or other user-defined symbols.
12  Ident(Ident),
13  /// An operator, such as '+', '-', '*', '==', etc., representing mathematical
14  /// or logical operations.
15  Opr(Opr),
16  /// A left delimiter, such as '(', '{', '[', used to denote the beginning of
17  /// a block, list, or expression.
18  LeftDelimiter(LeftDelimiter),
19  /// A right delimiter, such as ')', '}', ']', used to denote the end of a
20  /// block, list, or expression.
21  RightDelimiter(RightDelimiter),
22  /// A separator, such as ',', or ';', used to separate elements in a list or
23  /// statements in a block.
24  Separator(Separator),
25  }

```

The type has an encoding dimension  $d_{\text{Token}} = \Theta(\log L)$ , which is large enough to faithfully represent its information.

More specifically, the types `Literal`, `Keyword`, `Ident`, `Opr`, `LeftDelimiter`, `RightDelimiter`, `Separator` are local types assumed to have encoding dimension less than  $d_{\text{Token}}$ . `Keyword`, `Opr`, `LeftDelimiter`, `RightDelimiter`, `Separator` are small, so they can be encoded in a straight-forward manner entirely using  $d_{\text{Token}}$ . However, `Literal` and `Ident` are larger than representable by a limited number of bits because potentially a `Literal` can be a string literal of arbitrary length and an `Ident` can also be of arbitrary length. This can be solved through methods like interning, which gives all literals and identifiers that actually appear in the input distinct encodings. As the context length is  $L$ , the number of different literals/identifiers are bounded by context length and interning needs  $O(d_{\text{Token}}) = O(\log L)$  to work. As far as our theories are concerned, it's totally reasonable to assume that all these types are assumed to have encoding dimension less than  $d_{\text{Token}} = O(\log L)$ .

We define AST type as follows:

```

1  /// Represents a node in an Abstract Syntax Tree (AST).
2  ///
3  /// Each 'Ast' node has a reference to its parent node (if any) and holds
4  /// the associated syntax data (such as expressions, statements, or other
5  /// constructs defined in the 'AstData' enum).
6  pub struct Ast {
7      /// The index of the parent node in the AST, if it exists.
8      ///
9      /// - 'Some(Idx)': The node has a parent, and 'Idx' represents its position.
10     /// - 'None': The node is the root or does not have a parent.
11     pub parent: Option<Idx>,
12     /// The data associated with this AST node.
13     ///
14     /// This field holds the actual syntax information, which is typically
15     /// defined by the 'AstData' enum. This could represent literals, expressions,
16     /// statements, and other constructs in the source language.
17     pub data: AstData,
18 }

```

Note that we intentionally structure the tree by always storing the parent but not necessarily storing all children information. In our assumptions, we only control the depth of ASTs but don't control the number of children. More specifically, a function can have as many statements as possible. To avoid overflowing, we don't store all children information. As we shall see, parent information alone is enough for transformers to perform tree operations.

The `AstData` is the most complicated we define in this paper, as follows:

```

1  /// Enumeration representing different types of Abstract Syntax Tree (AST) nodes
2  pub enum AstData {
3      /// Represents a literal value (e.g., integer, string)
4      Literal(Literal),
5      /// Represents an identifier (e.g., variable name)
6      Ident(Ident),
7      /// Represents a prefix expression (e.g., '!x', '-x')
8      ///
9      /// # exprs
10     ///
11     Prefix {
12         /// Operator in the prefix expression (e.g., '!', '-')
13         opr: PrefixOpr,

```

```

2160 14      /// Operand index of the expression
2161 15      opd: Idx,
2162 16  },
2163 17  /// Represents a binary expression (e.g., 'x + y', 'a * b')
2164 18  Binary {
2165 19      /// Index of the left operand
2166 20      lopd: Idx,
2167 21      /// Operator in the binary expression (e.g., '+', '*')
2168 22      opr: BinaryOpr,
2169 23      /// Index of the right operand
2170 24      ropd: Idx,
2171 25  },
2172 26  /// Represents a suffix expression (e.g., 'x++', 'y--')
2173 27  Suffix {
2174 28      /// Index of the operand
2175 29      opd: Idx,
2176 30      /// Operator in the suffix expression (e.g., '++', '--')
2177 31      opr: SuffixOpr,
2178 32  },
2179 33  /// Represents a delimited expression (e.g., '(x + y)', '{a, b, c}')
2180 34  Delimited {
2181 35      /// Index of the left delimiter in the expression
2182 36      left_delimiter_idx: Idx,
2183 37      /// The left delimiter (e.g., '(', '{')
2184 38      left_delimiter: LeftDelimiter,
2185 39      /// The right delimiter (e.g., ')', '}')
2186 40      right_delimiter: RightDelimiter,
2187 41  },
2188 42  /// Represents an item separated by a separator (e.g., elements in an array or list)
2189 43  SeparatedItem {
2190 44      /// Index of the content, if any
2191 45      content: Option<Idx>,
2192 46      /// The separator (e.g., ',', ';')
2193 47      separator: Separator,
2194 48  },
2195 49  /// Represents a function call or array access (e.g., 'f(...)', 'a[...]')
2196 50  ///
2197 51  /// things like 'f(...)' or 'a[...]')
2198 52  Call {
2199 53      /// Index of the caller (e.g., function or array)
2200 54      caller: Idx,
2201 55      /// The left delimiter of the call (e.g., '(', '[')
2202 56      left_delimiter: LeftDelimiter,
2203 57      /// The right delimiter of the call (e.g., ')', ']')
2204 58      right_delimiter: RightDelimiter,
2205 59      /// Index of the delimited arguments in the call
2206 60      delimited_arguments: Idx,
2207 61  },
2208 62  /// Represents a 'let' statement with an initialization (e.g., 'let x = 5;')
2209 63  ///
2210 64  /// # stmts
2211 65  ///
2212 66  LetInit {
2213 67      /// Index of the expression in the initialization
2214 68      expr: Idx,
2215 69      /// Index of the pattern being initialized
2216 70      pattern: Idx,
2217 71      /// Optional index of the initial value
2218 72      initial_value: Option<Idx>,
2219 73  },
2220 74  /// Represents an 'if' statement
2221 75  If {
2222 76      /// Index of the condition in the 'if' statement
2223 77      condition: Idx,
2224 78      /// Index of the body of the 'if' statement
2225 79      body: Idx,
2226 80  },
2227 81  /// Represents an 'else' statement
2228 82  Else {
2229 83      /// Index of the associated 'if' statement
2230 84      if_stmt: Idx,
2231 85      /// Index of the body of the 'else' statement
2232 86      body: Idx,
2233 87  },
2234 88  /// Represents a function or variable definition
2235 89  ///
2236 90  /// # defn
2237 91  ///
2238 92  Defn {
2239 93      /// The keyword in the definition (e.g., 'fn', 'enum')
2240 94      keyword: DefnKeyword,

```

```

95     /// Index of the identifier in the definition
96     ident_idx: Idx,
97     /// The identifier being defined (e.g., function name, variable name)
98     ident: Ident,
99     /// Index of the content or body of the definition
100    content: Idx,
101  },
102 }

1 /// The 'PreAst' enum represents the intermediate forms of tokens and ASTs that are
2 /// encountered during the parsing phase, before the final AST is constructed.
3 /// Each variant corresponds to a specific type of token or partial
4 /// AST node that contributes to the construction of the final AST.
5 #[derive(Clone, Copy, PartialEq, Eq)]
6 pub enum PreAst {
7     /// A reserved keyword in the language, such as 'if', 'else', 'while', etc.
8     Keyword(Keyword),
9     /// An operator, such as '+', '-', '*', '==', etc., representing mathematical
10    /// or logical operations.
11    Opr(Opr),
12    /// A left delimiter, such as '(', '{', '[', used to denote the beginning of
13    /// a block, list, or expression.
14    LeftDelimiter(LeftDelimiter),
15    /// A right delimiter, such as ')', '}', ']', used to denote the end of a
16    /// block, list, or expression.
17    RightDelimiter(RightDelimiter),
18    /// A partially constructed AST node, representing a more complex structure
19    /// that will be further processed to build the final AST.
20    Ast(AstData),
21    /// A separator, such as ',' or ';', used to separate elements in a list or
22    /// statements in a block.
23    Separator(Separator),
24 }

1 /// this is beyond the scope of Cybertron
2 ///
3 /// rather a general Rust function to integrate for testing
4 pub fn calc_ast_from_input(input: &str, n: usize) -> (Seq<Option<PreAst>>,
5   Seq<Option<Ast>>) {
6     let tokens = tokenize(input);
7     let pre_ast = calc_pre_ast_initial_seq(tokens);
8     let allocated_ast: Seq<Option<Ast>> = tokens.map(|token| token.into());
9     reduce_n_times(pre_ast, allocated_ast, n)
10 }

```

The `reduce` function in Cybertron is designed to progressively refine sequences of pre-abstract syntax trees (pre-ASTs) and allocated abstract syntax trees (ASTs). The function takes two input sequences: `pre_ast`, which is a sequence of optional pre-ASTs, and `allocated_ast`, which is a sequence of optional ASTs. It returns a tuple containing the reduced sequences of pre-ASTs and allocated ASTs.

The reduction process is carried out in multiple stages, each focusing on different syntactic constructs:

1. `reduce_by_opr` : This step handles reduction by dealing with operators and their precedence. It simplifies expressions involving operations to form more compact ASTs.
2. `reduce_by_delimited` : This step reduces constructs that are delimited, such as those involving parentheses, braces, or other grouping symbols. It ensures that delimited blocks are properly nested and combined in the AST.
3. `reduce_by_call` : In this stage, function or method calls are reduced. This involves identifying and structuring calls within the AST, ensuring correct representation of function invocations.
4. `reduce_by_stmt` : This reduction step addresses statements, ensuring that individual statements are correctly parsed and represented within the AST, such as assignment statements, loops, and conditionals.
5. `reduce_by_defn` : Finally, reduction by definition handles the parsing of definitions, such as variable or function declarations. This step ensures that all definitions are correctly represented within the AST.

By sequentially applying these reduction steps, the `reduce` function progressively transforms the initial sequences into their most refined forms, ready for further syntactic or semantic analysis.

```

1 pub fn reduce(
2     pre_ast: Seq<Option<PreAst>>,
3     allocated_ast: Seq<Option<Ast>>,
4 ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
5     // Reduce ASTs by handling operators and precedence
6     let (pre_ast, allocated_ast) = reduce_by_opr(pre_ast, allocated_ast);
7
8     // Reduce ASTs by handling delimited constructs like parentheses or braces
9     let (pre_ast, allocated_ast) = reduce_by_delimited(pre_ast, allocated_ast);
10
11    // Reduce ASTs by handling function or method calls
12    let (pre_ast, allocated_ast) = reduce_by_call(pre_ast, allocated_ast);
13
14    // Reduce ASTs by handling statements, ensuring proper syntax structure
15    let (pre_ast, allocated_ast) = reduce_by_stmt(pre_ast, allocated_ast);
16
17    // Reduce ASTs by handling definitions, like variables or functions
18    let (pre_ast, allocated_ast) = reduce_by_defn(pre_ast, allocated_ast);
19
20    // Return the final reduced sequences of pre-ASTs and allocated ASTs
21    (pre_ast, allocated_ast)
22 }

```

```

1 pub fn reduce_n_times(
2     mut pre_ast: Seq<Option<PreAst>>,
3     mut allocated_ast: Seq<Option<Ast>>,
4     n: usize,
5 ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
6     for _ in 0..n {
7         let (pre_ast1, allocated_ast1) = reduce(pre_ast, allocated_ast);
8         pre_ast = pre_ast1;
9         allocated_ast = allocated_ast1;
10    }
11    (pre_ast, allocated_ast)
12 }

```

In the above definition, we actually used Rust’s mutable variable semantics. However, it’s straightforward to see that it translates to a computation graph that is a sequential composition of subgraphs with sequential length  $n$ . Because the AST’s depth is bounded by  $D$ , we can just take  $n$  to be  $D$ . Each subgraph is generated from the `reduce` function, then they are all constant graphs constructed by global and local functions, then by Proposition 13.11 and 2 they translate to transformers with  $O(\log L + D)$  depth, model dimension, and number of heads, where  $\log L$  comes from the encoding of types like `Token`.

Below we give full details of the various reduction functions.

As these are implemented as Rust functions, they have been tested against a number of inputs. We don’t guarantee an industry level of correctness, but the key point is well illustrated.

## F.2 OPERATORS

In this section, we lay down the definition of `reduce_by_opr`.

```

1 pub(super) fn reduce_by_opr(
2     pre_ast: Seq<Option<PreAst>>,
3     allocated_ast: Seq<Option<Ast>>,
4 ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
5     let pre_ast_nearest_left2 = pre_ast.nearest_left2();
6     let pre_ast_nearest_right2 = pre_ast.nearest_right2();
7     let new_opr_ast = new_opr_ast.apply(pre_ast_nearest_left2, pre_ast,
8         pre_ast_nearest_right2);
9     let (pre_ast_reduced, new_parents) = reduce_pre_ast_by_opr(pre_ast, new_opr_ast);
10    let pre_ast = update_pre_ast_by_new_ast(pre_ast_reduced, new_opr_ast);
11    let allocated_ast =
12        allocate_ast_and_update_parents(allocated_ast, new_opr_ast, new_parents);
13    (pre_ast, allocated_ast)
14 }

```

```

1 /// a finite function
2 pub(crate) fn new_opr_ast(
3     nearest_left2: Option2<(Idx, PreAst)>,

```

```

2322 4     current: Option<PreAst>,
2323 5     nearest_right2: Option2<(Idx, PreAst)>,
2324 6 ) -> Option<AstData> {
2325 7     let Some(PreAst::Opr(opr)) = current else {
2326 8         return None;
2327 9     };
2328 10    match opr {
2329 11    Opr::Prefix(opr) => {
2330 12        let Some((opd, PreAst::Ast(_))) = nearest_right2.first() else {
2331 13            return None;
2332 14        };
2333 15        if let Some( (_, ast)) = nearest_right2.second() {
2334 16            match ast {
2335 17            PreAst::Keyword(_) => (),
2336 18            PreAst::Opr(right_opr) => match right_opr {
2337 19                Opr::Prefix(_) => (),
2338 20                Opr::Binary(right_opr) => {
2339 21                    // every binary opr in our small language is left associative,
2340 22                    so '<' instead of '<='
2341 23                    if right_opr.precedence() > opr.precedence() {
2342 24                        return None;
2343 25                    }
2344 26                    Opr::Suffix(right_opr) => {
2345 27                        if right_opr.precedence() > opr.precedence() {
2346 28                            return None;
2347 29                        }
2348 30                    }
2349 31                },
2350 32                PreAst::Ast(_) => (),
2351 33                // function call or index takes higher precedence
2352 34                PreAst::LeftDelimiter(_) => return None,
2353 35                PreAst::RightDelimiter(_) => (),
2354 36                PreAst::Separator(_) => (),
2355 37            }
2356 38        };
2357 39        Some(AstData::Prefix { opr, opd })
2358 40    }
2359 41    Opr::Binary(opr) => {
2360 42        let Some((lopd, PreAst::Ast(_))) = nearest_left2.first() else {
2361 43            return None;
2362 44        };
2363 45        let Some((ropd, PreAst::Ast(_))) = nearest_right2.first() else {
2364 46            return None;
2365 47        };
2366 48        if let Some( (_, ast)) = nearest_left2.second() {
2367 49            match ast {
2368 50            PreAst::Keyword(kw) => (),
2369 51            PreAst::Opr(left_opr) => match left_opr {
2370 52                Opr::Prefix(left_opr) => {
2371 53                    if left_opr.precedence() >= opr.precedence() {
2372 54                        return None;
2373 55                    }
2374 56                }
2375 57                Opr::Binary(left_opr) => {
2376 58                    // every binary opr in our small language is left
2377 59                    associative, so '>=' instead of '>'
2378 60                    if left_opr.precedence() >= opr.precedence() {
2379 61                        return None;
2380 62                    }
2381 63                }
2382 64                Opr::Suffix(_) => (), // actually this will be a syntax error
2383 65            },
2384 66            PreAst::Ast(_) => {
2385 67                if opr != BinaryOpr::LightArrow {
2386 68                    return None;
2387 69                }
2388 70            }
2389 71            PreAst::LeftDelimiter(_) => (),
2390 72            PreAst::RightDelimiter(_) => return None,
2391 73            PreAst::Separator(_) => (),
2392 74        }
2393 75    };
2394 76    if let Some( (_, ast)) = nearest_right2.second() {
2395 77        match ast {
2396 78        PreAst::Keyword(kw) => match kw {
2397 79            Keyword::ELSE => return None,
2398 80            _ => (),
2399 81        },
2400 82        PreAst::Opr(right_opr) => match right_opr {
2401 83            Opr::Prefix(_) => (), // actually this will be a syntax error

```

```

2376      83         Opr::Binary(right_opr) => {
2377      84             // every binary opr in our small language is left
2378      85             associative, so '<' instead of '<='
2379      86             if right_opr.precedence() > opr.precedence() {
2380      87                 return None;
2381      88             }
2382      89             Opr::Suffix(right_opr) => {
2383      90                 if right_opr.precedence() >= opr.precedence() {
2384      91                     return None;
2385      92                 }
2386      93             }
2387      94         },
2388      95         // function call or index takes higher precedence
2389      96         PreAst::LeftDelimiter(_) => return None,
2390      97         PreAst::RightDelimiter(_) => (),
2391      98         PreAst::Ast(_) => (),
2392      99         PreAst::Separator(_) => (),
2400     100     }
2401     101     };
2402     102     Some(AstData::Binary { lopd, opr, ropd })
2403     103 }
2404     104 Opr::Suffix(opr) => {
2405     105     let Some((opd, PreAst::Ast(_))) = nearest_left2.first() else {
2406     106         return None;
2407     107     };
2408     108     if let Some( (_, ast)) = nearest_left2.second() {
2409     109         match ast {
2410     110             PreAst::Keyword(_) => (),
2411     111             PreAst::Opr(right_opr) => match right_opr {
2412     112                 Opr::Prefix(right_opr) => {
2413     113                     if right_opr.precedence() > opr.precedence() {
2414     114                         return None;
2415     115                     }
2416     116                 }
2417     117                 Opr::Binary(right_opr) => {
2418     118                     // every binary opr in our small language is left
2419     119                     associative, so '<' instead of '<='
2420     120                     if right_opr.precedence() > opr.precedence() {
2421     121                         return None;
2422     122                     }
2423     123                 }
2424     124                 Opr::Suffix(_) => (),
2425     125             },
2426     126             PreAst::LeftDelimiter(_) => (),
2427     127             PreAst::RightDelimiter(_) => return None,
2428     128             PreAst::Ast(_) => return None,
2429     129             PreAst::Separator(_) => (),
2430     130         }
2431     131     };
2432     132     Some(AstData::Suffix { opr, opd })
2433     133 }
2434     134 }
2435
2436 1 // returns sequence of remaining PreAsts and new parent idxs
2437 2 pub(crate) fn reduce_pre_ast_by_new_ast(
2438 3     pre_ast: Seq<Option<PreAst>>,
2439 4     new_ast: Seq<Option<AstData>>,
2440 5 ) -> (Seq<Option<PreAst>>, Seq<Option<Idx>>) {
2441 6     let new_ast_nearest_left = new_ast.nearest_left();
2442 7     let pre_ast = reduce_pre_ast_by_new_ast.apply(pre_ast, new_ast);
2443 8     let (pre_ast, new_parents) = reduce_pre_ast_by_opr_left
2444 9         .apply_enumerated(new_ast_nearest_left, pre_ast)
2445 10     .decouple();
2446 11     let new_ast_nearest_right = new_ast.nearest_right();
2447 12     reduce_pre_ast_by_opr_right
2448 13         .apply_enumerated(new_ast_nearest_right, pre_ast, new_parents)
2449 14     .decouple()
2450 15 }
2451
2452 1 fn reduce_pre_ast_by_new_ast(pre_ast: Option<PreAst>, new_ast: Option<AstData>) ->
2453 2     Option<PreAst> {
2454 3     if new_ast.is_some() {
2455 4         None
2456 5     } else {
2457 6         pre_ast
2458 7     }
2459 8 }

```

```

2430 1 fn reduce_pre_ast_by_opr_left (
2431 2   idx: Idx,
2432 3   new_ast_nearest_left: Option<(Idx, AstData)>,
2433 4   pre_ast: Option<PreAst>,
2434 5 ) -> (Option<PreAst>, Option<Idx>) {
2435 6   let Some(pre_ast) = pre_ast else {
2436 7     return (None, None);
2437 8   };
2438 9   let Some((new_ast_idx, new_ast_data)) = new_ast_nearest_left else {
2439 10    return (Some(pre_ast), None);
2440 11  };
2441 12  match new_ast_data {
2442 13    AstData::Binary { ropd: opd, .. } | AstData::Prefix { opd, .. } if opd == idx => {
2443 14      (None, Some(new_ast_idx))
2444 15    }
2445 16    _ => (Some(pre_ast), None),
2446 17  }
2447 18 }
2448
2449 1 fn reduce_pre_ast_by_opr_right (
2450 2   idx: Idx,
2451 3   new_ast_nearest_right: Option<(Idx, AstData)>,
2452 4   pre_ast: Option<PreAst>,
2453 5   new_parent: Option<Idx>,
2454 6 ) -> (Option<PreAst>, Option<Idx>) {
2455 7   let Some(pre_ast) = pre_ast else {
2456 8     return (None, new_parent);
2457 9   };
2458 10  if let Some(new_parent) = new_parent {
2459 11    return (None, Some(new_parent));
2460 12  }
2461 13  let Some((new_ast_idx, new_ast_data)) = new_ast_nearest_right else {
2462 14    return (Some(pre_ast), None);
2463 15  };
2464 16  match new_ast_data {
2465 17    AstData::Binary { lopd: opd, .. } | AstData::Suffix { opd, .. } if opd == idx => {
2466 18      (None, Some(new_ast_idx))
2467 19    }
2468 20    _ => (Some(pre_ast), None),
2469 21  }
2470 22 }

```

### F.3 STATEMENTS

In this section, we lay down the definition of `reduce_by_stmt`.

```

2471 1 pub(super) fn reduce_by_stmt (
2472 2   pre_ast_nearest_left2: Option2<(Idx, PreAst)>,
2473 3   pre_ast: Option<PreAst>,
2474 4   allocated_ast: Seq<Option<Ast>>,
2475 5 ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
2476 6   let pre_ast_nearest_left2 = pre_ast.nearest_left2();
2477 7   let pre_ast_nearest_right2 = pre_ast.nearest_right2();
2478 8   let new_stmt_ast =
2479 9     new_stmt_ast.apply(pre_ast_nearest_left2, pre_ast, pre_ast_nearest_right2);
2480 10  let (pre_ast, new_parents) = reduce_pre_ast_by_stmt(pre_ast, new_stmt_ast);
2481 11  let allocated_ast =
2482 12    allocate_ast_and_update_parents(allocated_ast, new_stmt_ast, new_parents);
2483 13  let pre_ast = update_pre_ast_by_new_ast(pre_ast, new_stmt_ast);
2484 14  (pre_ast, allocated_ast)
2485 15 }
2486
2487 1 fn new_stmt_ast (
2488 2   pre_ast_nearest_left2: Option2<(Idx, PreAst)>,
2489 3   pre_ast: Option<PreAst>,
2490 4   pre_ast_nearest_right2: Option2<(Idx, PreAst)>,
2491 5 ) -> Option<AstData> {
2492 6   let PreAst::Keyword(Keyword::Stmt(kw)) = pre_ast? else {
2493 7     return None;
2494 8   };
2495 9   match kw {
2496 10    StmtKeyword::Let => {
2497 11      let Some((idx1, PreAst::Ast(ast))) = pre_ast_nearest_right2.first() else {
2498 12        return None;
2499 13      };
2500 14      if let Some(., pre_ast) = pre_ast_nearest_right2.second() {
2501 15        match pre_ast {
2502 16          PreAst::Keyword(_) => (),

```

```

2484 17         PreAst::Opr(_) | PreAst::LeftDelimiter(_) => return None,
2485 18         PreAst::RightDelimiter(_) => (),
2486 19         PreAst::Ast(_) => return None,
2487 20         PreAst::Separator(separator) => match separator {
2488 21             Separator::Comma => return None,
2489 22             Separator::Semicolon => (),
2490 23         },
2491 24     }
2492 25 }
2493 26 let (pattern, initial_value) = match ast {
2494 27     AstData::Binary {
2495 28         lopd,
2496 29         opr: BinaryOpr::Assign,
2497 30         ropd,
2498 31     } => (lopd, Some(ropd)),
2499 32     AstData::Ident(_)
2500 33     | AstData::Prefix { .. }
2501 34     | AstData::Binary { .. }
2502 35     | AstData::Delimited { .. }
2503 36     | AstData::Call { .. } => (idx1, None),
2504 37     _ => return None,
2505 38 };
2506 39 Some(AstData::LetInit {
2507 40     expr: idx1,
2508 41     pattern,
2509 42     initial_value,
2510 43 })
2511 44 }
2512 45 StmtKeyword::If => {
2513 46     let Some((condition, PreAst::Ast(ast1))) = pre_ast_nearest_right2.first() else
2514 47     {
2515 48         return None;
2516 49     };
2517 50     let Some((
2518 51         body,
2519 52         PreAst::Ast(AstData::Delimited {
2520 53             left_delimiter: LCURL,
2521 54             right_delimiter: RCURL,
2522 55             ..
2523 56         })),
2524 57     ) = pre_ast_nearest_right2.second()
2525 58     else {
2526 59         return None;
2527 60     };
2528 61     Some(AstData::If { condition, body })
2529 62 }
2530 63 StmtKeyword::Else => {
2531 64     let Some((if_stmt, PreAst::Ast(AstData::If { .. }))) =
2532 65     pre_ast_nearest_left2.first()
2533 66     else {
2534 67         return None;
2535 68     };
2536 69     let Some((
2537 70         body,
2538 71         PreAst::Ast(
2539 72             AstData::Delimited {
2540 73                 left_delimiter: LCURL,
2541 74                 right_delimiter: RCURL,
2542 75                 ..
2543 76             }
2544 77         | AstData::If { .. }
2545 78         | AstData::Else { .. },
2546 79     )),
2547 80     ) = pre_ast_nearest_right2.first()
2548 81     else {
2549 82         return None;
2550 83     };
2551 84     if let Some( (_, PreAst::Keyword(Keyword::ELSE))) =
2552 85     pre_ast_nearest_right2.second() {
2553 86         return None;
2554 87     }
2555 88     Some(AstData::Else { if_stmt, body })
2556 89 }
2557 90 }
2558 91 }
2559 92 }
2560 93 }
2561 94 }
2562 95 }
2563 96 }
2564 97 }
2565 98 }
2566 99 }
2567 100 }
2568 101 }
2569 102 }
2570 103 }
2571 104 }
2572 105 }
2573 106 }
2574 107 }
2575 108 }
2576 109 }
2577 110 }
2578 111 }
2579 112 }
2580 113 }
2581 114 }
2582 115 }
2583 116 }
2584 117 }
2585 118 }
2586 119 }
2587 120 }
2588 121 }
2589 122 }
2590 123 }
2591 124 }
2592 125 }
2593 126 }
2594 127 }
2595 128 }
2596 129 }
2597 130 }
2598 131 }
2599 132 }
2600 133 }
2601 134 }
2602 135 }
2603 136 }
2604 137 }
2605 138 }
2606 139 }
2607 140 }
2608 141 }
2609 142 }
2610 143 }
2611 144 }
2612 145 }
2613 146 }
2614 147 }
2615 148 }
2616 149 }
2617 150 }
2618 151 }
2619 152 }
2620 153 }
2621 154 }
2622 155 }
2623 156 }
2624 157 }
2625 158 }
2626 159 }
2627 160 }
2628 161 }
2629 162 }
2630 163 }
2631 164 }
2632 165 }
2633 166 }
2634 167 }
2635 168 }
2636 169 }
2637 170 }
2638 171 }
2639 172 }
2640 173 }
2641 174 }
2642 175 }
2643 176 }
2644 177 }
2645 178 }
2646 179 }
2647 180 }
2648 181 }
2649 182 }
2650 183 }
2651 184 }
2652 185 }
2653 186 }
2654 187 }
2655 188 }
2656 189 }
2657 190 }
2658 191 }
2659 192 }
2660 193 }
2661 194 }
2662 195 }
2663 196 }
2664 197 }
2665 198 }
2666 199 }
2667 200 }
2668 201 }
2669 202 }
2670 203 }
2671 204 }
2672 205 }
2673 206 }
2674 207 }
2675 208 }
2676 209 }
2677 210 }
2678 211 }
2679 212 }
2680 213 }
2681 214 }
2682 215 }
2683 216 }
2684 217 }
2685 218 }
2686 219 }
2687 220 }
2688 221 }
2689 222 }
2690 223 }
2691 224 }
2692 225 }
2693 226 }
2694 227 }
2695 228 }
2696 229 }
2697 230 }
2698 231 }
2699 232 }
2700 233 }
2701 234 }
2702 235 }
2703 236 }
2704 237 }
2705 238 }
2706 239 }
2707 240 }
2708 241 }
2709 242 }
2710 243 }
2711 244 }
2712 245 }
2713 246 }
2714 247 }
2715 248 }
2716 249 }
2717 250 }
2718 251 }
2719 252 }
2720 253 }
2721 254 }
2722 255 }
2723 256 }
2724 257 }
2725 258 }
2726 259 }
2727 260 }
2728 261 }
2729 262 }
2730 263 }
2731 264 }
2732 265 }
2733 266 }
2734 267 }
2735 268 }
2736 269 }
2737 270 }
2738 271 }
2739 272 }
2740 273 }
2741 274 }
2742 275 }
2743 276 }
2744 277 }
2745 278 }
2746 279 }
2747 280 }
2748 281 }
2749 282 }
2750 283 }
2751 284 }
2752 285 }
2753 286 }
2754 287 }
2755 288 }
2756 289 }
2757 290 }
2758 291 }
2759 292 }
2760 293 }
2761 294 }
2762 295 }
2763 296 }
2764 297 }
2765 298 }
2766 299 }
2767 300 }
2768 301 }
2769 302 }
2770 303 }
2771 304 }
2772 305 }
2773 306 }
2774 307 }
2775 308 }
2776 309 }
2777 310 }
2778 311 }
2779 312 }
2780 313 }
2781 314 }
2782 315 }
2783 316 }
2784 317 }
2785 318 }
2786 319 }
2787 320 }
2788 321 }
2789 322 }
2790 323 }
2791 324 }
2792 325 }
2793 326 }
2794 327 }
2795 328 }
2796 329 }
2797 330 }
2798 331 }
2799 332 }
2800 333 }
2801 334 }
2802 335 }
2803 336 }
2804 337 }
2805 338 }
2806 339 }
2807 340 }
2808 341 }
2809 342 }
2810 343 }
2811 344 }
2812 345 }
2813 346 }
2814 347 }
2815 348 }
2816 349 }
2817 350 }
2818 351 }
2819 352 }
2820 353 }
2821 354 }
2822 355 }
2823 356 }
2824 357 }
2825 358 }
2826 359 }
2827 360 }
2828 361 }
2829 362 }
2830 363 }
2831 364 }
2832 365 }
2833 366 }
2834 367 }
2835 368 }
2836 369 }
2837 370 }
2838 371 }
2839 372 }
2840 373 }
2841 374 }
2842 375 }
2843 376 }
2844 377 }
2845 378 }
2846 379 }
2847 380 }
2848 381 }
2849 382 }
2850 383 }
2851 384 }
2852 385 }
2853 386 }
2854 387 }
2855 388 }
2856 389 }
2857 390 }
2858 391 }
2859 392 }
2860 393 }
2861 394 }
2862 395 }
2863 396 }
2864 397 }
2865 398 }
2866 399 }
2867 400 }
2868 401 }
2869 402 }
2870 403 }
2871 404 }
2872 405 }
2873 406 }
2874 407 }
2875 408 }
2876 409 }
2877 410 }
2878 411 }
2879 412 }
2880 413 }
2881 414 }
2882 415 }
2883 416 }
2884 417 }
2885 418 }
2886 419 }
2887 420 }
2888 421 }
2889 422 }
2890 423 }
2891 424 }
2892 425 }
2893 426 }
2894 427 }
2895 428 }
2896 429 }
2897 430 }
2898 431 }
2899 432 }
2900 433 }
2901 434 }
2902 435 }
2903 436 }
2904 437 }
2905 438 }
2906 439 }
2907 440 }
2908 441 }
2909 442 }
2910 443 }
2911 444 }
2912 445 }
2913 446 }
2914 447 }
2915 448 }
2916 449 }
2917 450 }
2918 451 }
2919 452 }
2920 453 }
2921 454 }
2922 455 }
2923 456 }
2924 457 }
2925 458 }
2926 459 }
2927 460 }
2928 461 }
2929 462 }
2930 463 }
2931 464 }
2932 465 }
2933 466 }
2934 467 }
2935 468 }
2936 469 }
2937 470 }
2938 471 }
2939 472 }
2940 473 }
2941 474 }
2942 475 }
2943 476 }
2944 477 }
2945 478 }
2946 479 }
2947 480 }
2948 481 }
2949 482 }
2950 483 }
2951 484 }
2952 485 }
2953 486 }
2954 487 }
2955 488 }
2956 489 }
2957 490 }
2958 491 }
2959 492 }
2960 493 }
2961 494 }
2962 495 }
2963 496 }
2964 497 }
2965 498 }
2966 499 }
2967 500 }
2968 501 }
2969 502 }
2970 503 }
2971 504 }
2972 505 }
2973 506 }
2974 507 }
2975 508 }
2976 509 }
2977 510 }
2978 511 }
2979 512 }
2980 513 }
2981 514 }
2982 515 }
2983 516 }
2984 517 }
2985 518 }
2986 519 }
2987 520 }
2988 521 }
2989 522 }
2990 523 }
2991 524 }
2992 525 }
2993 526 }
2994 527 }
2995 528 }
2996 529 }
2997 530 }
2998 531 }
2999 532 }
3000 533 }
3001 534 }
3002 535 }
3003 536 }
3004 537 }
3005 538 }
3006 539 }
3007 540 }
3008 541 }
3009 542 }
3010 543 }
3011 544 }
3012 545 }
3013 546 }
3014 547 }
3015 548 }
3016 549 }
3017 550 }
3018 551 }
3019 552 }
3020 553 }
3021 554 }
3022 555 }
3023 556 }
3024 557 }
3025 558 }
3026 559 }
3027 560 }
3028 561 }
3029 562 }
3030 563 }
3031 564 }
3032 565 }
3033 566 }
3034 567 }
3035 568 }
3036 569 }
3037 570 }
3038 571 }
3039 572 }
3040 573 }
3041 574 }
3042 575 }
3043 576 }
3044 577 }
3045 578 }
3046 579 }
3047 580 }
3048 581 }
3049 582 }
3050 583 }
3051 584 }
3052 585 }
3053 586 }
3054 587 }
3055 588 }
3056 589 }
3057 590 }
3058 591 }
3059 592 }
3060 593 }
3061 594 }
3062 595 }
3063 596 }
3064 597 }
3065 598 }
3066 599 }
3067 600 }
3068 601 }
3069 602 }
3070 603 }
3071 604 }
3072 605 }
3073 606 }
3074 607 }
3075 608 }
3076 609 }
3077 610 }
3078 611 }
3079 612 }
3080 613 }
3081 614 }
3082 615 }
3083 616 }
3084 617 }
3085 618 }
3086 619 }
3087 620 }
3088 621 }
3089 622 }
3090 623 }
3091 624 }
3092 625 }
3093 626 }
3094 627 }
3095 628 }
3096 629 }
3097 630 }
3098 631 }
3099 632 }
3100 633 }
3101 634 }
3102 635 }
3103 636 }
3104 637 }
3105 638 }
3106 639 }
3107 640 }
3108 641 }
3109 642 }
3110 643 }
3111 644 }
3112 645 }
3113 646 }
3114 647 }
3115 648 }
3116 649 }
3117 650 }
3118 651 }
3119 652 }
3120 653 }
3121 654 }
3122 655 }
3123 656 }
3124 657 }
3125 658 }
3126 659 }
3127 660 }
3128 661 }
3129 662 }
3130 663 }
3131 664 }
3132 665 }
3133 666 }
3134 667 }
3135 668 }
3136 669 }
3137 670 }
3138 671 }
3139 672 }
3140 673 }
3141 674 }
3142 675 }
3143 676 }
3144 677 }
3145 678 }
3146 679 }
3147 680 }
3148 681 }
3149 682 }
3150 683 }
3151 684 }
3152 685 }
3153 686 }
3154 687 }
3155 688 }
3156 689 }
3157 690 }
3158 691 }
3159 692 }
3160 693 }
3161 694 }
3162 695 }
3163 696 }
3164 697 }
3165 698 }
3166 699 }
3167 700 }
3168 701 }
3169 702 }
3170 703 }
3171 704 }
3172 705 }
3173 706 }
3174 707 }
3175 708 }
3176 709 }
3177 710 }
3178 711 }
3179 712 }
3180 713 }
3181 714 }
3182 715 }
3183 716 }
3184 717 }
3185 718 }
3186 719 }
3187 720 }
3188 721 }
3189 722 }
3190 723 }
3191 724 }
3192 725 }
3193 726 }
3194 727 }
3195 728 }
3196 729 }
3197 730 }
3198 731 }
3199 732 }
3200 733 }
3201 734 }
3202 735 }
3203 736 }
3204 737 }
3205 738 }
3206 739 }
3207 740 }
3208 741 }
3209 742 }
3210 743 }
3211 744 }
3212 745 }
3213 746 }
3214 747 }
3215 748 }
3216 749 }
3217 750 }
3218 751 }
3219 752 }
3220 753 }
3221 754 }
3222 755 }
3223 756 }
3224 757 }
3225 758 }
3226 759 }
3227 760 }
3228 761 }
3229 762 }
3230 763 }
3231 764 }
3232 765 }
3233 766 }
3234 767 }
3235 768 }
3236 769 }
3237 770 }
3238 771 }
3239 772 }
3240 773 }
3241 774 }
3242 775 }
3243 776 }
3244 777 }
3245 778 }
3246 779 }
3247 780 }
3248 781 }
3249 782 }
3250 783 }
3251 784 }
3252 785 }
3253 786 }
3254 787 }
3255 788 }
3256 789 }
3257 790 }
3258 791 }
3259 792 }
3260 793 }
3261 794 }
3262 795 }
3263 796 }
3264 797 }
3265 798 }
3266 799 }
3267 800 }
3268 801 }
3269 802 }
3270 803 }
3271 804 }
3272 805 }
3273 806 }
3274 807 }
3275 808 }
3276 809 }
3277 810 }
3278 811 }
3279 812 }
3280 813 }
3281 814 }
3282 815 }
3283 816 }
3284 817 }
3285 818 }
3286 819 }
3287 820 }
3288 821 }
3289 822 }
3290 823 }
3291 824 }
3292 825 }
3293 826 }
3294 827 }
3295 828 }
3296 829 }
3297 830 }
3298 831 }
3299 832 }
3300 833 }
3301 834 }
3302 835 }
3303 836 }
3304 837 }
3305 838 }
3306 839 }
3307 840 }
3308 841 }
3309 842 }
3310 843 }
3311 844 }
3312 845 }
3313 846 }
3314 847 }
3315 848 }
3316 849 }
3317 850 }
3318 851 }
3319 852 }
3320 853 }
3321 854 }
3322 855 }
3323 856 }
3324 857 }
3325 858 }
3326 859 }
3327 860 }
3328 861 }
3329 862 }
3330 863 }
3331 864 }
3332 865 }
3333 866 }
3334 867 }
3335 868 }
3336 869 }
3337 870 }
3338 871 }
3339 872 }
3340 873 }
3341 874 }
3342 875 }
3343 876 }
3344 877 }
3345 878 }
3346 879 }
3347 880 }
3348 881 }
3349 882 }
3350 883 }
3351 884 }
3352 885 }
3353 886 }
3354 887 }
3355 888 }
3356 889 }
3357 890 }
3358 891 }
3359 892 }
3360 893 }
3361 894 }
3362 895 }
3363 896 }
3364 897 }
3365 898 }
3366 899 }
3367 900 }
3368 901 }
3369 902 }
3370 903 }
3371 904 }
3372 905 }
3373 906 }
3374 907 }
3375 908 }
3376 909 }
3377 910 }
3378 911 }
3379 912 }
3380 913 }
3381 914 }
3382 915 }
3383 916 }
3384 917 }
3385 918 }
3386 919 }
3387 920 }
3388 921 }
3389 922 }
3390 923 }
3391 924 }
3392 925 }
3393 926 }
3394 927 }
3395 928 }
3396 929 }
3397 930 }
3398 931 }
3399 932 }
3400 933 }
3401 934 }
3402 935 }
3403 936 }
3404 937 }
3405 938 }
3406 939 }
3407 940 }
3408 941 }
3409 942 }
3410 943 }
3411 944 }
3412 945 }
3413 946 }
3414 947 }
3415 948 }
3416 949 }
3417 950 }
3418 951 }
3419 952 }
3420 953 }
3421 954 }
3422 955 }
3423 956 }
3424 957 }
3425 958 }
3426 959 }
3427 960 }
3428 961 }
3429 962 }
3430 963 }
3431 964 }
3432 965 }
3433 966 }
3434 967 }
3435 968 }
3436 969 }
3437 970 }
3438 971 }
3439 972 }
3440 973 }
3441 974 }
3442 975 }
3443 976 }
3444 977 }
3445 978 }
3446 979 }
3447 980 }
3448 981 }
3449 982 }
3450 983 }
3451 984 }
3452 985 }
3453 986 }
3454 987 }
3455 988 }
3456 989 }
3457 990 }
3458 991 }
3459 992 }
3460 993 }
3461 994 }
3462 995 }
3463 996 }
3464 997 }
3465 998 }
3466 999 }
3467 1000 }

```

```

6   let new_ast_nearest_right = new_ast.nearest_right();
7   reduce_pre_ast_by_stmt
8     .apply_enumerated(new_ast_nearest_left, new_ast_nearest_right, pre_ast)
9     .decouple()
10  }

1  fn reduce_pre_ast_by_stmt(
2    idx: Idx,
3    new_ast_nearest_left: Option<Idx, AstData>,
4    new_ast_nearest_right: Option<Idx, AstData>,
5    pre_ast: Option<PreAst>,
6  ) -> (Option<PreAst>, Option<Idx>) {
7    if let Some((idx1, ast)) = new_ast_nearest_left {
8      match ast {
9        AstData::LetInit { expr, .. } if expr == idx => (None, Some(idx1)),
10       AstData::If {
11         condition, body, ..
12       } if condition == idx || body == idx => (None, Some(idx1)),
13       AstData::Else { body, .. } if body == idx => (None, Some(idx1)),
14       _ => (pre_ast, None),
15     }
16   } else if let Some((idx1, AstData::Else { if_stmt, .. })) = new_ast_nearest_right
17     && if_stmt == idx
18   {
19     (None, Some(idx1))
20   } else {
21     (pre_ast, None)
22   }
23 }

```

#### F.4 GENERALIZED CALL FORMS

In this section, we lay down the definition of `reduce_by_call`.

```

1  pub(super) fn reduce_by_call(
2    pre_ast: Seq<Option<PreAst>>,
3    allocated_ast: Seq<Option<Ast>>,
4  ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
5    let pre_ast_nearest_left2 = pre_ast.nearest_left2();
6    let pre_ast_nearest_right = pre_ast.nearest_right();
7    let new_call_ast =
8      new_call_ast.apply_enumerated(pre_ast_nearest_left2, pre_ast_nearest_right);
9    let (pre_ast, new_parents) = reduce_pre_ast_by_call(pre_ast, new_call_ast);
10   let allocated_ast =
11     allocate_ast_and_update_parents(allocated_ast, new_call_ast, new_parents);
12   let pre_ast = update_pre_ast_by_new_ast(pre_ast, new_call_ast);
13   (pre_ast, allocated_ast)
14 }

1  fn new_call_ast(
2    idx: Idx,
3    pre_ast_nearest_left2: Option2<Idx, PreAst>,
4    pre_ast_nearest_right: Option<Idx, PreAst>,
5  ) -> Option<AstData> {
6    let (caller, PreAst::Ast(caller_ast)) = pre_ast_nearest_left2.first()? else {
7      return None;
8    };
9    let (
10      delimited_arguments,
11      PreAst::Ast(AstData::Delimited {
12        left_delimiter_idx,
13        left_delimiter,
14        right_delimiter,
15      }),
16    ) = pre_ast_nearest_right?
17  else {
18    return None;
19  };
20  if let Some((_, snd)) = pre_ast_nearest_left2.second() {
21    match snd {
22      PreAst::Keyword(kw) => match kw {
23        Keyword::Defn(kw) => match kw {
24          DefnKeyword::Struct | DefnKeyword::Enum => return None,
25          DefnKeyword::Fn => match left_delimiter.delimiter() {
26            Delimiter::Parenthesis | Delimiter::Box => return None,
27            Delimiter::Curly => (),
28          },

```

```

2592         },
2593         Keyword::Stmt(kw) => match kw {
2594             StmtKeyword::Let => (),
2595             StmtKeyword::If => match left_delimiter.delimiter() {
2596                 Delimiter::Parenthesis | Delimiter::Box => (),
2597                 Delimiter::Curly => return None,
2598             },
2599             StmtKeyword::Else => return None,
2600         },
2601     },
2602     PreAst::Opr(opr) => match opr {
2603         Opr::Prefix(_) | Opr::Binary(_) => match left_delimiter.delimiter() {
2604             Delimiter::Parenthesis | Delimiter::Box => (),
2605             Delimiter::Curly => return None,
2606         },
2607         Opr::Suffix(_) => return None,
2608     },
2609     PreAst::LeftDelimiter(_) => (),
2610     PreAst::RightDelimiter(_) => return None,
2611     PreAst::Ast(snd_ast) => {
2612         if let AstData::Ident(_) = snd_ast
2613             && left_delimiter == LCURL
2614         {
2615             match caller_ast {
2616                 AstData::Binary {
2617                     opr: BinaryOpr::LightArrow,
2618                     ..
2619                 }
2620                 | AstData::Delimited {
2621                     left_delimiter: LPAR,
2622                     right_delimiter: RPAR,
2623                     ..
2624                 } => (),
2625                 _ => return None,
2626             }
2627         } else {
2628             return None;
2629         }
2630     }
2631     PreAst::Separator(_) => (),
2632 }
2633
2634 if left_delimiter_idx != idx {
2635     return None;
2636 }
2637 Some(AstData::Call {
2638     caller,
2639     delimited_arguments,
2640     left_delimiter,
2641     right_delimiter,
2642 })
2643 }
2644
2645 fn reduce_pre_asts_by_call(
2646     pre_asts: Seq<Option<PreAst>>,
2647     new_asts: Seq<Option<AstData>>,
2648 ) -> (Seq<Option<PreAst>>, Seq<Option<Idx>>) {
2649     let new_asts_nearest_left = new_asts.nearest_left();
2650     let new_asts_nearest_right = new_asts.nearest_right();
2651     reduce_pre_ast_by_call
2652         .apply_enumerated(new_asts_nearest_left, new_asts_nearest_right, pre_asts)
2653         .decouple()
2654 }
2655
2656 fn reduce_pre_ast_by_call(
2657     idx: Idx,
2658     new_ast_nearest_left: Option<(Idx, AstData)>,
2659     new_ast_nearest_right: Option<(Idx, AstData)>,
2660     pre_ast: Option<PreAst>,
2661 ) -> (Option<PreAst>, Option<Idx>) {
2662     if let Some((
2663         idx1,
2664         AstData::Call {
2665             delimited_arguments,
2666             ..
2667         },
2668     )) = new_ast_nearest_left
2669         && delimited_arguments == idx
2670     {
2671         (None, Some(idx1))
2672     }

```

```

17   } else if let Some((idx1, AstData::Call { caller, .. })) = new_ast_nearest_right
18     && caller == idx
19   {
20     (None, Some(idx1))
21   } else {
22     (pre_ast, None)
23   }
24 }

```

## F.5 DEFINITIONS

In this section, we lay down the definition of `reduce_by_defn`.

```

1  pub(super) fn reduce_by_defn(
2    pre_asts: Seq<Option<PreAst>>,
3    allocated_asts: Seq<Option<Ast>>,
4  ) -> (Seq<Option<PreAst>>, Seq<Option<Ast>>) {
5    let pre_asts_nearest_left2 = pre_asts.nearest_left2();
6    let pre_asts_nearest_right2 = pre_asts.nearest_right2();
7    let new_defn_asts =
8      new_defn_ast.apply(pre_asts_nearest_left2, pre_asts, pre_asts_nearest_right2);
9    let (pre_asts, new_parents) = reduce_pre_asts_by_defn(pre_asts, new_defn_asts);
10   let allocated_asts =
11     allocate_asts_and_update_parents(allocated_asts, new_defn_asts, new_parents);
12   let pre_asts = update_pre_asts_by_new_asts(pre_asts, new_defn_asts);
13   (pre_asts, allocated_asts)
14 }

```

```

1  fn new_defn_ast(
2    pre_ast_nearest_left2: Option2<Idx, PreAst>,
3    pre_ast: Option<PreAst>,
4    pre_ast_nearest_right2: Option2<Idx, PreAst>,
5  ) -> Option<AstData> {
6    let PreAst::Keyword(Keyword::Defn(keyword)) = pre_ast? else {
7      return None;
8    };
9    {
10     let Some((ident_idx, PreAst::Ast(AstData::Ident(ident)))) =
11       pre_ast_nearest_right2.first()
12     else {
13       return None;
14     };
15     let Some((content, PreAst::Ast(content_ast))) = pre_ast_nearest_right2.second()
16     else {
17       return None;
18     };
19     match keyword {
20       DefnKeyword::Struct => match content_ast {
21         AstData::Delimited { .. } => (),
22         _ => return None,
23       },
24       DefnKeyword::Enum => match content_ast {
25         AstData::Delimited { .. } => (),
26         _ => return None,
27       },
28       DefnKeyword::Fn => match content_ast {
29         AstData::Call { .. } => (),
30         _ => return None,
31       },
32     }
33     Some(AstData::Defn {
34       keyword,
35       ident,
36       ident_idx,
37       content,
38     })
39 }

```

```

1  fn reduce_pre_asts_by_defn(
2    pre_asts: Seq<Option<PreAst>>,
3    new_asts: Seq<Option<AstData>>,
4  ) -> (Seq<Option<PreAst>>, Seq<Option<Idx>>) {
5    let new_asts_nearest_left = new_asts.nearest_left();
6    let new_asts_nearest_right = new_asts.nearest_right();
7    reduce_pre_ast_by_defn
8      .apply_enumerated(new_asts_nearest_left, new_asts_nearest_right, pre_asts)

```

```

9      .decouple()
10 }

1 fn reduce_pre_ast_by_defn(
2   idx: Idx,
3   new_ast_nearest_left: Option<Idx, AstData>,
4   new_ast_nearest_right: Option<Idx, AstData>,
5   pre_ast: Option<PreAst>,
6 ) -> (Option<PreAst>, Option<Idx>) {
7   if let Some((idx1, ast)) = new_ast_nearest_left {
8     match ast {
9       AstData::Defn {
10         keyword,
11         ident_idx,
12         ident,
13         content,
14         ..
15       } if ident_idx == idx || content == idx => (None, Some(idx1)),
16       _ => (pre_ast, None),
17     }
18   } else if let Some((idx1, AstData::Defn { .. })) = new_ast_nearest_right
19     && false
20   {
21     (None, Some(idx1))
22   } else {
23     (pre_ast, None)
24   }
25 }

```

## G TRANSFORMER SYMBOL RESOLUTION PROOF

Here we lay down the code for symbol resolution. The actual process involves many details such as computing ranks (the exact position of an AST node among its siblings), scopes, and roles (a more precise version of AST, computed from its parent recursively), definitions and resolutions.

### G.1 RANKS

```

1 #[derive(Debug, Default, PartialEq, Eq, Clone, Copy)]
2 pub struct Rank(u8);
3
4 impl Rank {
5   fn next(self) -> Self {
6     Self(self.0 + 1)
7   }
8 }
9
10 pub fn calc_ranks(asts: Seq<Option<Ast>>) -> Seq<Option<Rank>> {
11   let counts = asts.count_past_by_attention(asts, |ast, ast1| {
12     let Some(ast) = ast else { return false };
13     let Some(ast1) = ast1 else { return false };
14     ast.parent == ast1.parent
15   });
16   (|c: usize, ast| {
17     ast?;
18     Some(Rank(c.try_into().unwrap()))
19   })
20   .apply(counts, asts)
21 }
22
23 pub fn calc_ranks1(asts: Seq<Option<Ast>>, n: usize) -> Seq<Option<Rank>> {
24   let mut ranks: Seq<Option<Rank>> = asts.map(|_| None);
25   for _ in 0..n {
26     ranks = calc_sibling_indicies_step(asts, ranks);
27   }
28   ranks
29 }
30
31 fn calc_sibling_indicies_step(
32   asts: Seq<Option<Ast>>,
33   ranks: Seq<Option<Rank>>,
34 ) -> Seq<Option<Rank>> {
35   let previous_ranks = ranks.nearest_left_filtered_by_attention(asts, asts, |ast, ast1| {
36     let Some(ast) = ast else { return false };
37     let Some(ast1) = ast1 else { return false };
38     ast.parent == ast1.parent

```

```

39     });
2754
2755     let ranks = (|ast, rank, previous_rank: Option<Option<Rank>>| {
2756         let _ = ast?;
2757         if let Some(rank) = rank {
2758             return Some(rank);
2759         }
2760         let Some(previous_rank) = previous_rank else {
2761             return Some(Default::default());
2762         };
2763         Some(previous_rank?.next())
2764     })
2765     .apply(asts, ranks, previous_ranks);
2766     ranks
2767 }

```

In the above, `count_past_by_attention` that count is representable by transformers by utilizing directly hard attention and the starter token. If the count is  $c$ , we shall get  $c/(c+1)$  from the attention directly.

## G.2 SCOPES

```

1  const D: usize = 8usize;
2
3  pub struct Scope {
4      enclosing_blocks: BoundedVec<Idx, D>,
5  }
6
7  impl Scope {
8      pub fn from_ast(idx: Idx, ast: AstData, parent_scope: Scope) -> Self {
9          match ast {
10             AstData::Delimited {
11                 left_delimiter_idx,
12                 left_delimiter: LCURL,
13                 right_delimiter: RCURL,
14             } => Self {
15                 enclosing_blocks: parent_scope.enclosing_blocks.append(idx),
16             },
17             _ => parent_scope,
18         }
19     }
20
21     pub fn new(idx: Idx) -> Self {
22         Self {
23             enclosing_blocks: todo!(),
24         }
25     }
26
27     pub fn append(self, idx: Idx) -> Self {
28         Self {
29             enclosing_blocks: self.enclosing_blocks.append(idx),
30         }
31     }
32 }
33
34 impl Scope {
35     pub fn contains(self, other: Self) -> bool {
36         let len = self.enclosing_blocks.len();
37         if len > other.enclosing_blocks.len() {
38             return false;
39         }
40         for i in 0..len {
41             if self.enclosing_blocks[i] != other.enclosing_blocks[i] {
42                 return false;
43             }
44         }
45         true
46     }
47 }
48
49 pub fn infer_scopes(asts: Seq<Option<Ast>>, n: usize) -> Seq<Option<Scope>> {
50     let mut scopes = initial_scope.apply_enumerated(asts);
51     for _ in 0..n {
52         let parent_scopes = parent_queries(asts, scopes);
53         scopes = infer_scopes_step(asts, parent_scopes, scopes);
54     }
55     scopes
56 }
57
58 fn initial_scope(idx: Idx, ast: Option<Ast>) -> Option<Scope> {

```

```

2808 59   let ast = ast?;
2809 60   if ast.parent.is_some() {
2810 61       return None;
2811 62   }
2812 63   let scope = Scope::default();
2813 64   Some(Scope::from_ast(idx, ast.data, scope))
2814 65 }
2815 66
2816 67 fn infer_scopes_step(
2817 68     asts: Seq<Option<Ast>>,
2818 69     parent_scopes: Seq<Option<Scope>>,
2819 70     scopes: Seq<Option<Scope>>,
2820 71 ) -> Seq<Option<Scope>> {
2821 72     infer_scope_step.apply_enumerated(asts, parent_scopes, scopes)
2822 73 }
2823 74
2824 75 fn infer_scope_step(
2825 76     idx: Idx,
2826 77     ast: Option<Ast>,
2827 78     parent_scope: Option<Scope>,
2828 79     scope: Option<Scope>,
2829 80 ) -> Option<Scope> {
2830 81     if let Some(scope) = scope {
2831 82         return Some(scope);
2832 83     }
2833 84     Some(Scope::from_ast(idx, ast?.data, parent_scope?))
2834 85 }

```

### G.3 ROLES

```

2830 1 #[derive(Debug, Clone, Copy, PartialEq, Eq)]
2831 2 pub enum Role {
2832 3     LetStmt {
2833 4         pattern: Idx,
2834 5         initial_value: Option<Idx>,
2835 6     },
2836 7     LetStmtInner {
2837 8         pattern: Idx,
2838 9         initial_value: Idx,
2839 10    },
2840 11    LetStmtIdent,
2841 12    LetStmtTypedVariables {
2842 13        variables: Idx,
2843 14        ty: Idx,
2844 15    },
2845 16    StructDefn(Ident),
2846 17    EnumDefn(Ident),
2847 18    FnDefn(Ident),
2848 19    FnDefnCallForm {
2849 20        fn_ident: Ident,
2850 21        scope: Scope,
2851 22    },
2852 23    FnParameters {
2853 24        fn_ident: Ident,
2854 25        has_return_ty: bool,
2855 26        scope: Scope,
2856 27    },
2857 28    FnParametersAndReturnType {
2858 29        fn_ident: Ident,
2859 30        parameters: Idx,
2860 31        scope: Scope,
2861 32        return_ty: Idx,
2862 33    },
2863 34    FnBody(Ident),
2864 35    StructFields(Ident),
2865 36    FnParameter {
2866 37        fn_ident: Ident,
2867 38        rank: Rank,
2868 39        ty: Idx,
2869 40        fn_ident_idx: Idx,
2870 41        scope: Scope,
2871 42    },
2872 43    FnParameterIdent {
2873 44        scope: Scope,
2874 45    },
2875 46    FnParameterSeparated {
2876 47        fn_ident: Ident,
2877 48        rank: Rank,
2878 49        scope: Scope,

```

```

2862   },
2863   FnParameterType {
2864     fn_ident: Ident,
2865     rank: Rank,
2866   },
2867   FnOutputType {
2868     fn_ident: Ident,
2869   },
2870   StructField {
2871     ty_ident: Ident,
2872     field_ident: Ident,
2873     ty_idx: Idx,
2874   },
2875   StructFieldType {
2876     ty_ident: Ident,
2877     field_ident: Ident,
2878   },
2879   TypeArgument,
2880   TypeArguments,
2881   StructFieldSeparated(Ident),
2882   LetStmtVariablesType,
2883   LetStmtVariables,
2884 }
2885
2886 impl Ast {
2887   fn role(self) -> Option<Role> {
2888     match self.data {
2889       AstData::LetInit {
2890         expr,
2891         pattern,
2892         initial_value,
2893       } => Some(Role::LetStmt {
2894         pattern,
2895         initial_value,
2896       }),
2897       AstData::Defn {
2898         keyword,
2899         ident_idx,
2900         ident,
2901         content,
2902       } => Some(match keyword {
2903         DefnKeyword::Struct => Role::StructDefn(ident),
2904         DefnKeyword::Enum => Role::EnumDefn(ident),
2905         DefnKeyword::Fn => Role::FnDefn(ident),
2906       }),
2907       _ => None,
2908     }
2909   }
2910 }
2911
2912 pub fn calc_roles(
2913   asts: Seq<Option<Ast>>,
2914   scopes: Seq<Option<Scope>>,
2915   n: usize,
2916 ) -> Seq<Option<Role>> {
2917   let mut roles: Seq<Option<Role>> = asts.map(|ast| ast?.role());
2918   let ranks = calc_ranks(asts);
2919   for _ in 0..n {
2920     let parent_roles = parent_queries(asts, roles);
2921     roles = calc_roles_step(asts, parent_roles, roles, ranks, scopes);
2922   }
2923   roles
2924 }
2925
2926 fn calc_roles_step(
2927   asts: Seq<Option<Ast>>,
2928   parent_roles: Seq<Option<Role>>,
2929   roles: Seq<Option<Role>>,
2930   ranks: Seq<Option<Rank>>,
2931   scopes: Seq<Option<Scope>>,
2932 ) -> Seq<Option<Role>> {
2933   calc_role_step.apply_enumerated(asts, parent_roles, roles, ranks, scopes)
2934 }
2935
2936 fn calc_role_step(
2937   idx: Idx,
2938   ast: Option<Ast>,
2939   parent_role: Option<Role>,
2940   role: Option<Role>,

```

```

2916 6     rank: Option<Rank>,
2917 7     scope: Option<Scope>,
2918 8 ) -> Option<Role> {
2919 9     if let Some(role) = role {
2920 10         return Some(role);
2921 11     }
2922 12     let ast = ast?;
2923 13     if let Some(role) = ast.role() {
2924 14         return Some(role);
2925 15     }
2926 16     match parent_role? {
2927 17         Role::LetStmt {
2928 18             pattern,
2929 19             initial_value,
2930 20         } => match ast.data {
2931 21             AstData::Ident(ident) if idx == pattern => Some(Role::LetStmtIdent),
2932 22             AstData::Binary {
2933 23                 lopd,
2934 24                 opr: BinaryOpr::Assign,
2935 25                 ropd,
2936 26                 lopd_ident,
2937 27             } if lopd == pattern => Some(Role::LetStmtInner {
2938 28                 pattern,
2939 29                 initial_value: ropd,
2940 30             }),
2941 31             _ => None,
2942 32         },
2943 33         Role::LetStmtInner {
2944 34             pattern,
2945 35             initial_value,
2946 36         } => {
2947 37             if idx == pattern {
2948 38                 match ast.data {
2949 39                     AstData::Ident(ident) => Some(Role::LetStmtIdent),
2950 40                     AstData::Binary {
2951 41                         lopd,
2952 42                         lopd_ident,
2953 43                         opr,
2954 44                         ropd,
2955 45                     } => Some(Role::LetStmtTypedVariables {
2956 46                         variables: lopd,
2957 47                         ty: ropd,
2958 48                     }),
2959 49                     _ => todo!(),
2960 50                 }
2961 51             } else {
2962 52                 None
2963 53             }
2964 54         }
2965 55     }
2966 56     Role::LetStmtIdent => todo!(),
2967 57     Role::FnParameterIdent { scope } => todo!(),
2968 58     Role::StructDefn(ident) => match ast.data {
2969 59         AstData::Literal(_) => todo!(),
2970 60         AstData::Ident(_) => None,
2971 61         AstData::Prefix { opr, opd } => todo!(),
2972 62         AstData::Binary {
2973 63             lopd,
2974 64             opr,
2975 65             ropd,
2976 66             lopd_ident,
2977 67         } => todo!(),
2978 68         AstData::Suffix { opd, opr } => todo!(),
2979 69         AstData::Delimited {
2980 70             left_delimiter_idx,
2981 71             left_delimiter,
2982 72             right_delimiter,
2983 73         } => Some(Role::StructFields(ident)),
2984 74         AstData::SeparatedItem { content, separator } => todo!(),
2985 75         AstData::Call { .. } => todo!(),
2986 76         AstData::LetInit {
2987 77             expr,
2988 78             pattern,
2989 79             initial_value,
2990 80         } => todo!(),
2991 81         AstData::Return { result } => todo!(),
2992 82         AstData::Assert { condition } => todo!(),
2993 83         AstData::If { condition, body } => todo!(),
2994 84         AstData::Else { if_stmt, body } => todo!(),
2995 85         AstData::Defn {
2996 86             keyword,
2997 87             ident_idx,

```

```

2970      87         ident,
2971      88         content,
2972      89     } => todo!(),
2973      90 },
2974      91 Role::EnumDefn(_) => None, // ad hoc
2975      92 Role::FnDefn(fn_ident) => match ast.data {
2976      93     AstData::Literal(_) => todo!(),
2977      94     AstData::Ident(_) => None,
2978      95     AstData::Prefix { opr, opd } => todo!(),
2979      96     AstData::Binary {
2980      97         lopd,
2981      98         opr,
2982      99         ropd,
2983     100         lopd_ident,
2984     101     } => todo!(),
2985     102     AstData::Suffix { opd, opr } => todo!(),
2986     103     AstData::Delimited {
2987     104         left_delimiter_idx,
2988     105         left_delimiter,
2989     106         right_delimiter,
2990     107     } => todo!(),
2991     108     AstData::SeparatedItem { content, separator } => todo!(),
2992     109     AstData::Call {
2993     110         delimited_arguments,
2994     111         ..
2995     112     } => Some(Role::FnDefnCallForm {
2996     113         fn_ident,
2997     114         scope: match scope {
2998     115             Some(scope) => scope.append(delimited_arguments),
2999     116             None => Scope::new(delimited_arguments),
3000     117         },
3001     118     }),
3002     119     AstData::LetInit {
3003     120         expr,
3004     121         pattern,
3005     122         initial_value,
3006     123     } => todo!(),
3007     124     AstData::Return { result } => todo!(),
3008     125     AstData::Assert { condition } => todo!(),
3009     126     AstData::If { condition, body } => todo!(),
3010     127     AstData::Else { if_stmt, body } => todo!(),
3011     128     AstData::Defn {
3012     129         keyword,
3013     130         ident_idx,
3014     131         ident,
3015     132         content,
3016     133     } => todo!(),
3017     134 },
3018     135 Role::FnDefnCallForm { fn_ident, scope } => match ast.data {
3019     136     AstData::Literal(_) => todo!(),
3020     137     AstData::Ident(_) => todo!(),
3021     138     AstData::Prefix { opr, opd } => todo!(),
3022     139     AstData::Binary {
3023     140         lopd,
3024     141         opr,
3025     142         ropd,
3026     143         lopd_ident,
3027     144     } => {
3028     145         if opr == BinaryOpr::LightArrow {
3029     146             Some(Role::FnParametersAndReturnType {
3030     147                 fn_ident,
3031     148                 parameters: lopd,
3032     149                 return_ty: ropd,
3033     150                 scope,
3034     151             })
3035     152         } else {
3036     153             unreachable!()
3037     154         }
3038     155     }
3039     156     AstData::Suffix { opd, opr } => todo!(),
3040     157     AstData::Delimited {
3041     158         left_delimiter_idx,
3042     159         left_delimiter,
3043     160         right_delimiter,
3044     161     } => match left_delimiter.delimiter() {
3045     162         Delimiter::Parenthesis => Some(Role::FnParameters {
3046     163             fn_ident,
3047     164             has_return_ty: false,
3048     165             scope,
3049     166         }),
3050     167         Delimiter::Box => todo!(),

```

```

3024 168         Delimiter::Curly => Some(Role::FnBody(fn_ident)),
3025 169     },
3026 170     AstData::SeparatedItem { content, separator } => todo!(),
3027 171     AstData::Call { .. } => todo!(),
3028 172     AstData::LetInit {
3029 173         expr,
3030 174         pattern,
3031 175         initial_value,
3032 176     } => todo!(),
3033 177     AstData::Return { result } => todo!(),
3034 178     AstData::Assert { condition } => todo!(),
3035 179     AstData::If { condition, body } => todo!(),
3036 180     AstData::Else { if_stmt, body } => todo!(),
3037 181     AstData::Defn {
3038 182         keyword,
3039 183         ident_idx,
3040 184         ident,
3041 185         content,
3042 186     } => todo!(),
3043 187 },
3044 188 Role::FnParameters {
3045 189     fn_ident, scope, ..
3046 190 } => match ast.data {
3047 191     AstData::Binary {
3048 192         lopd,
3049 193         opr,
3050 194         ropd,
3051 195         lopd_ident,
3052 196     } => {
3053 197         if opr == BinaryOpr::TypeIs {
3054 198             Some(Role::FnParameter {
3055 199                 fn_ident,
3056 200                 fn_ident_idx: lopd,
3057 201                 rank: rank.unwrap(),
3058 202                 ty: ropd,
3059 203                 scope,
3060 204             })
3061 205         } else {
3062 206             unreachable!()
3063 207         }
3064 208     }
3065 209     AstData::SeparatedItem { .. } => Some(Role::FnParameterSeparated {
3066 210         fn_ident,
3067 211         rank: rank.unwrap(),
3068 212         scope,
3069 213     }),
3070 214     _ => unreachable!(),
3071 215 },
3072 216 Role::FnBody(_) => None,
3073 217 Role::StructFields(ty_ident) => match ast.data {
3074 218     AstData::Binary {
3075 219         lopd,
3076 220         opr,
3077 221         ropd,
3078 222         lopd_ident,
3079 223     } => {
3080 224         assert_eq!(opr, BinaryOpr::TypeIs);
3081 225         Some(Role::StructField {
3082 226             ty_ident,
3083 227             field_ident: lopd_ident.unwrap(),
3084 228             ty_idx: ropd,
3085 229         })
3086 230     }
3087 231     AstData::SeparatedItem { content, separator } => {
3088 232         Some(Role::StructFieldSeparated(ty_ident))
3089 233     }
3090 234     _ => None,
3091 235 },
3092 236 Role::FnParameter {
3093 237     fn_ident,
3094 238     fn_ident_idx,
3095 239     rank,
3096 240     ty,
3097 241     scope,
3098 242     ..
3099 243 } => {
3100 244     if idx == ty {
3101 245         Some(Role::FnParameterType { fn_ident, rank })
3102 246     } else if idx == fn_ident_idx {
3103 247         Some(Role::FnParameterIdent { scope })
3104 248     } else {

```

```

3078 249         None
3079 250     }
3080 251 }
3081 252 Role::FnParameterSeparated {
3082 253     fn_ident,
3083 254     rank,
3084 255     scope,
3085 256 } => match ast.data {
3086 257     AstData::Binary {
3087 258         lopd,
3088 259         opr,
3089 260         ropd,
3090 261         lopd_ident,
3091 262     } => {
3092 263         if opr == BinaryOpr::TypeIs {
3093 264             Some(Role::FnParameter {
3094 265                 fn_ident,
3095 266                 fn_ident_idx: lopd,
3096 267                 rank,
3097 268                 ty: ropd,
3098 269                 scope,
3099 270             })
3100 271         } else {
3101 272             unreachable!()
3102 273         }
3103 274     }
3104 275     _ => unreachable!(),
3105 276 },
3106 277 Role::StructField {
3107 278     ty_ident,
3108 279     field_ident,
3109 280     ty_idx,
3110 281 } => {
3111 282     if idx == ty_idx {
3112 283         Some(Role::StructFieldType {
3113 284             ty_ident,
3114 285             field_ident,
3115 286         })
3116 287     } else {
3117 288         None
3118 289     }
3119 290 }
3120 291 Role::StructFieldSeparated(ty_ident) => match ast.data {
3121 292     AstData::Binary {
3122 293         lopd,
3123 294         opr,
3124 295         ropd,
3125 296         lopd_ident,
3126 297     } => {
3127 298         assert_eq!(opr, BinaryOpr::TypeIs);
3128 299         Some(Role::StructField {
3129 300             ty_ident,
3130 301             field_ident: lopd_ident.unwrap(),
3131 302             ty_idx: ropd,
3132 303         })
3133 304     }
3134 305     _ => unreachable!(),
3135 306 },
3136 307 Role::FnParameterType { .. } | Role::StructFieldType { .. } | Role::TypeArgument
3137 308 => {
3138 309     match ast.data {
3139 310         AstData::Delimited {
3140 311             left_delimiter_idx,
3141 312             left_delimiter,
3142 313             right_delimiter,
3143 314         } => Some(Role::TypeArguments),
3144 315         _ => None,
3145 316     }
3146 317 }
3147 318 Role::TypeArguments => match ast.data {
3148 319     AstData::Ident(_) => Some(Role::TypeArgument),
3149 320     AstData::Delimited {
3150 321         left_delimiter_idx,
3151 322         left_delimiter,
3152 323         right_delimiter,
3153 324     } => todo!(),
3154 325     AstData::SeparatedItem { content, separator } => todo!(),
3155 326     AstData::Call {
3156 327         caller,
3157 328         caller_ident,
3158         left_delimiter,

```

```

3132
3133         right_delimiter,
3134         delimited_arguments,
3135     } => todo!(),
3136     _ => None,
3137 },
3138 Role::FnParametersAndReturnType {
3139     fn_ident,
3140     parameters,
3141     return_ty,
3142     scope,
3143 } => {
3144     if idx == parameters {
3145         Some(Role::FnParameters {
3146             fn_ident,
3147             has_return_ty: true,
3148             scope,
3149         })
3150     } else if idx == return_ty {
3151         Some(Role::FnOutputType { fn_ident })
3152     } else {
3153         unreachable!()
3154     }
3155 }
3156 Role::FnOutputType { fn_ident } => todo!(),
3157 Role::LetStmtTypedVariables { variables, ty } => {
3158     if idx == variables {
3159         Some(Role::LetStmtVariables)
3160     } else if idx == ty {
3161         Some(Role::LetStmtVariablesType)
3162     } else {
3163         unreachable!()
3164     }
3165 }
3166 Role::LetStmtVariablesType => todo!(),
3167 Role::LetStmtVariables => todo!(),
3168 }
3169 }

```

#### G.4 DEFNS

```

3161 1 #[derive(Debug, Clone, Copy, PartialEq, Eq)]
3162 2 pub struct SymbolDefn {
3163 3     pub symbol: Symbol,
3164 4     pub scope: Option<Scope>,
3165 5 }
3166
3167 1 pub fn calc_symbol_defns(
3168 2     asts: Seq<Option<Ast>>,
3169 3     scopes: Seq<Option<Scope>>,
3170 4     n: usize,
3171 5 ) -> Seq<Option<SymbolDefn>> {
3172 6     let roles = calc_roles(asts, scopes, n);
3173 7     calc_symbol_defn.apply_enumerated(asts, roles, scopes)
3174 8 }
3175
3176 1 fn calc_symbol_defn(
3177 2     idx: Idx,
3178 3     ast: Option<Ast>,
3179 4     role: Option<Role>,
3180 5     scope: Option<Scope>,
3181 6 ) -> Option<SymbolDefn> {
3182 7     match ast?.data {
3183 8         AstData::Ident(ident) => match role? {
3184 9             Role::LetStmt { .. } => unreachable!(),
3185 10            Role::LetStmtVariables | Role::LetStmtIdent => Some(SymbolDefn {
3186 11                symbol: Symbol {
3187 12                    ident,
3188 13                    source: idx,
3189 14                    data: SymbolData::Variable,
3190 15                },
3191 16                scope,
3192 17            }),
3193 18            Role::FnParameterIdent { scope } => Some(SymbolDefn {
3194 19                symbol: Symbol {
3195 20                    ident,
3196 21                    source: idx,
3197 22                    data: SymbolData::Variable,
3198 23                },

```

```

24         scope: Some(scope),
25     },
26     _ => None,
27 },
28 AstData::Defn {
29     keyword,
30     ident_idx,
31     ident,
32     content,
33 } => Some(SymbolDefn {
34     symbol: Symbol {
35         ident,
36         source: idx,
37         data: SymbolData::Item {
38             kind: keyword.into(),
39         },
40     },
41     scope,
42 },
43 _ => None,
44 }
45 }

```

## G.5 RESOLUTIONS

```

1 pub enum SymbolResolution {
2     Ok(Symbol),
3     Err(SymbolResolutionError),
4 }

```

```

1 pub enum SymbolResolutionError {
2     NotResolved,
3     NotYetDeclared(Symbol),
4 }

```

```

1 pub fn calc_symbol_resolutions(asts: Seq<Option<Ast>>, n: usize) ->
2     Seq<Option<SymbolResolution>> {
3     let scopes = infer_scopes(asts, n);
4     let symbol_defns = calc_symbol_defns(asts, scopes, n);
5     let idents = asts.map(|ast| match ast?.data {
6         AstData::Ident(ident) => Some(ident),
7         _ => None,
8     });
9     let symbols = symbol_defns
10        .map(|symbol_defn| Some(symbol_defn?.symbol))
11        .first_filtered_by_attention(
12            (|ident, scope| (ident, scope)).apply(idents, scopes),
13            symbol_defns,
14            |(ident, scope), symbol_defn| {
15                let Some(ident) = ident else { return false };
16                let Some(symbol_defn) = symbol_defn else {
17                    return false;
18                };
19                if let Some(symbol_defn_scope) = symbol_defn.scope {
20                    if !symbol_defn_scope.contains(scope.unwrap()) {
21                        return false;
22                    }
23                }
24                symbol_defn.symbol.ident == ident
25            },
26        )
27        .map(|s| s.flatten());
28     finalize.apply_enumerated(idents, symbols)
29 }

```

In the above code, we use a somehow complicated attention which we should illustrate why it's representable by transformers. The essence is to prove `symbol_defn_scope.contains(scope.unwrap())` can be represented as part of the inner product in  $Q^T K$ . This can be done by looking closer to what `contains` does. Consider two scopes, `scope1` and `scope2`, which are sequences of bracket ast indices (can be null). The function returns true if the sequence of `scope1` contains the sequence of `scope2` as prefix, which can be achieved by  $\sum_i x_i^T y_i$  where  $x_i, y_i$  are the encoding of  $i$ th ast indices of `scope1` and `scope2` after some transformations (different transformations because the function

is asymmetric) so that  $x_i^\top y_i = 0$  if and only if either  $x_i$  is a None or  $x_i$  represents the same thing as  $y_i$ , and  $x_i^\top y_i < 0$  otherwise. More concretely, if  $x_i$  is a None,  $x_i = 0$  by choice, and equal to  $(1, u_i)$  otherwise where  $u_i$  corresponds to the encoding of the  $i$ th ast index of `scope1`; if  $y_i$  is a None,  $y_i = 0$  by choice, and equal to  $(-1, v_i)$  otherwise where  $A > 0$  and  $v_i$  corresponds to the encoding of the  $i$ th ast index of `scope2`. We should choose the encoding  $u_i, v_i$  such that  $u_i^\top v_i = 1$  if and only if they encode the same index, which is obviously easy enough.

```

1 fn finalize(idx: Idx, ident: Option<Ident>, symbol: Option<Symbol>) ->
  Option<SymbolResolution> {
2   let _ = ident?;
3   let Some(symbol) = symbol else {
4     return Some(SymbolResolution::Err(SymbolResolutionError::NotResolved));
5   };
6   match symbol.data {
7     SymbolData::Item { .. } => (),
8     SymbolData::Variable => {
9       if idx < symbol.source {
10        return Some(SymbolResolution::Err(
11          SymbolResolutionError::NotYetDeclared(symbol),
12        ));
13      }
14    }
15  }
16  Some(SymbolResolution::Ok(symbol))
17 }

```

## H TRANSFORMER TYPE CHECKING PROOF

Here we lay down the code for type analysis. It should be noted that we didn't completely implement all the details. Things like struct fields, enum variant fields are left out. However, we already cover the essential mechanism of type analysis, making it sufficient for proof purposes.

### H.1 TYPE SIGNATURES

```

1 #[deri
2 ve(Debug, PartialEq, Eq, Clone, Copy)]
3 pub struct TypeSignature {
4   pub key: TypeSignatureKey,
5   pub ty: Type,
6 }
7
8
9 #[derive(Debug, PartialEq, Eq, Clone, Copy)]
10 pub enum TypeSignatureKey {
11   FnParameter { fn_ident: Ident, rank: Rank },
12   FnOutput { fn_ident: Ident },
13   StructField { ty_ident: Ident, field_ident: Ident },
14 }
15
16 pub(super) fn calc_ty_signatures(
17   asts: Seq<Option<Ast>>,
18   roles: Seq<Option<Role>>,
19   ty_terms: Seq<Option<Type>>,
20 ) -> Seq<Option<TypeSignature>> {
21   calc_ty_signature.apply(roles, ty_terms)
22 }
23
24 fn calc_ty_signature(role: Option<Role>, ty_term: Option<Type>) -> Option<TypeSignature> {
25   let key = match role? {
26     Role::FnParameterType { fn_ident, rank } => {
27       TypeSignatureKey::FnParameter { fn_ident, rank }
28     }
29     Role::StructFieldType {
30       ty_ident,
31       field_ident,
32     } => TypeSignatureKey::StructField {
33       ty_ident,
34       field_ident,
35     },
36     Role::FnOutputType { fn_ident } => TypeSignatureKey::FnOutput { fn_ident },
37     Role::FnParameters {
38       fn_ident,
39     }

```

```

16         has_return_ty: false,
17         scope,
18     } => {
19         let key = TypeSignatureKey::FnOutput { fn_ident };
20         let ty = Type::new_ident(Ident::new("unit"));
21         return Some(TypeSignature { key, ty });
22     }
23     _ => return None,
24 };
25 // put it here!
26 let ty = ty_term?;
27 Some(TypeSignature { key, ty })
28 }

```

## H.2 TYPE INFERENCE

```

1 pub struct TypeInference {
2     pub ty: Type,
3 }
4
5 pub fn calc_ty_inferences(
6     asts: Seq<Option<Ast>>,
7     symbol_resolutions: Seq<Option<SymbolResolution>>,
8     roles: Seq<Option<Role>>,
9     ty_terms: Seq<Option<Type>>,
10    ty_signatures: Seq<Option<TypeSignature>>,
11    n: usize,
12 ) -> Seq<Option<TypeInference>> {
13     let mut ty_inferences = infer_tys_initial(asts, ty_signatures);
14     let mut ty_designations =
15         calc_initial_ty_designations(asts, roles, symbol_resolutions, ty_inferences,
16                                     ty_terms);
17     for _ in 0..n {
18         ty_inferences |= infer_tys_step(asts, symbol_resolutions, ty_inferences,
19                                       ty_designations);
20         ty_designations |= calc_ty_designations_step(roles, symbol_resolutions,
21                                                     ty_inferences);
22     }
23     ty_inferences
24 }
25
26 fn infer_tys_initial(
27     asts: Seq<Option<Ast>>,
28     ty_signatures: Seq<Option<TypeSignature>>,
29 ) -> Seq<Option<TypeInference>> {
30     inference_literal_tys(asts).or(infer_fn_call_tys(asts, ty_signatures))
31 }
32
33 fn inference_literal_tys(asts: Seq<Option<Ast>>) -> Seq<Option<TypeInference>> {
34     asts.map(|ast| match ast?.data {
35         AstData::Literal(lit) => match lit {
36             Literal::Int(_) => Some(TypeInference {
37                 ty: Type::new_ident(Ident::new("Int")),
38             }),
39             Literal::Float(_) => Some(TypeInference {
40                 ty: Type::new_ident(Ident::new("Float")),
41             }),
42         },
43         _ => None,
44     })
45 }
46
47 fn infer_fn_call_tys(
48     asts: Seq<Option<Ast>>,
49     ty_signatures: Seq<Option<TypeSignature>>,
50 ) -> Seq<Option<TypeInference>> {
51     ty_signatures
52         .first_filtered_by_attention(asts, ty_signatures, |ast, ty_signature| {
53             let Some(ast) = ast else { return false };
54             let Some(TypeSignature {
55                 key: TypeSignatureKey::FnOutput { fn_ident },
56                 ..
57             }) = ty_signature
58             else {
59                 return false;
60             };
61             match ast.data {

```

```

16         AstData::Call {
17             caller,
18             caller_ident,
19             left_delimiter,
20             right_delimiter,
21             delimited_arguments,
22         } if caller_ident == Some(fn_ident) => true,
23         _ => false,
24     }
25 }
26 .map(|ty_inference| {
27     Some(TypeInference {
28         ty: ty_inference?.ty,
29     })
30 })
31 }

```

### H.3 TYPE EXPECTATIONS

```

1 pub struct TypeExpectation {
2     pub ty: Type,
3     pub source: TypeExpectationSource,
4 }

```

```

1 pub enum TypeExpectationSource {
2     CallArgument { caller_ident: Ident, rank: Rank },
3 }

```

```

1 pub fn calc_ty_expectations(
2     asts: Seq<Option<Ast>>,
3     ranks: Seq<Option<Rank>>,
4     ty_signatures: Seq<Option<TypeSignature>>,
5 ) -> Seq<Option<TypeExpectation>> {
6     let parent_ast = asts.index(asts.map(|ast| ast?.parent)).map(Option::flatten);
7     let grandparent_ast = asts
8         .index(parent_ast.map(|parent_ast| parent_ast?.parent))
9         .map(Option::flatten);
10    let ty_expectation_sources = calc_ty_expectation_source.apply(grandparent_ast, ranks);
11    let retrieved_ty_signatures = ty_signatures
12        .first_filtered_by_attention(
13            ty_expectation_sources,
14            ty_signatures,
15            |ty_expectation_source, ty_signature| {
16                let Some(type_expectation_source) = ty_expectation_source else {
17                    return false;
18                };
19                let Some(type_signature) = ty_signature else {
20                    return false;
21                };
22                match (type_expectation_source, type_signature.key()) {
23                    (
24                        TypeExpectationSource::CallArgument {
25                            caller_ident,
26                            rank: rank0,
27                        },
28                        TypeSignatureKey::FnParameter {
29                            fn_ident,
30                            rank: rank1,
31                        },
32                    ) if caller_ident == fn_ident && rank0 == rank1 => true,
33                    _ => false,
34                }
35            },
36        )
37        .map(Option::flatten);
38    (|ty_expectation_source: Option<TypeExpectationSource>,
39     retrieved_ty_signature: Option<TypeSignature>| {
40        Some(TypeExpectation {
41            ty: retrieved_ty_signature?.ty(),
42            source: ty_expectation_source?,
43        })
44    })
45    .apply(ty_expectation_sources, retrieved_ty_signatures)
46 }

```

```

1 fn calc_ty_expectation_source(
2     grandparent_ast: Option<Ast>,

```

```

3   rank: Option<Rank>,
4 ) -> Option<TypeExpectationSource> {
5   let grandparent_ast = grandparent_ast?;
6   let rank = rank?;
7   match grandparent_ast.data {
8     AstData::Call {
9       caller,
10      caller_ident: Some(caller_ident),
11      left_delimiter,
12      right_delimiter,
13      delimited_arguments,
14    } => Some(TypeExpectationSource::CallArgument { caller_ident, rank }),
15    _ => None,
16  }
17 }

```

#### H.4 TYPE ERRORS

```

1 pub enum TypeError {
2   TypeMismatch { expected: Type, actual: Type },
3 }

1 pub fn calc_ty_errors(
2   ty_inferences: Seq<Option<TypeInference>>,
3   ty_expectations: Seq<Option<TypeExpectation>>,
4 ) -> Seq<Option<TypeError>> {
5   calc_ty_error.apply(ty_inferences, ty_expectations)
6 }

1 fn calc_ty_error(
2   ty_inference: Option<TypeInference>,
3   ty_expectation: Option<TypeExpectation>,
4 ) -> Option<TypeError> {
5   let ty_inference = ty_inference?;
6   let ty_expectation = ty_expectation?;
7   if ty_inference.ty == ty_expectation.ty {
8     None
9   } else {
10    Some(TypeError::TypeMismatch {
11      expected: ty_expectation.ty,
12      actual: ty_inference.ty,
13    })
14  }
15 }

```

## I LOWER BOUNDS

```

1 struct <ty-ident-1> {}
2 struct <ty-ident-2> {}
3 struct <ty-ident-3> {}
4 struct <ty-ident-4> {}
5
6 fn <f-ident-1>(a: <arg-ty-ident-1>) {}
7 fn <f-ident-2>(a: <arg-ty-ident-2>) {}
8 fn <f-ident-3>(a: <arg-ty-ident-3>) {}
9 fn <f-ident-4>(a: <arg-ty-ident-4>) {}
10
11 fn g() {
12   let x: <ty-ident> = ...;
13   <f-ident>(x);
14 }

```

### I.1 LOWER BOUNDS FOR RNN: EASY BOUNDS DUE TO MEMORY

*Proof of Theorem 4.* Our proof resonates with the proof of Theorem 4.6 in [Wen et al. \(2024\)](#) and Theorem 8 in [Bhattachamishra et al. \(2024\)](#). For  $L, D, H \in \mathbb{N}$ , suppose that  $D$  makes  $\text{MiniHuskyAnnotated}_{D,H}$  to be nontrivial, i.e., one can define functions with one parameter and use function calls. Simple calculations shows we can choose  $D = 7$  and  $H = 1$ . If a RNN represents a function maps any token sequence of length  $L$  in  $\text{MiniHuskyAnnotated}_{D,H}$  to its type

errors represented as a sequence of values of type `Option<TypeError>`, then the memory right before type checking must store all previous type signatures, the number of which can be as many as  $\Omega(L)$  in the worst case. Assuming proper numerical discretization, the memorization of these type signatures would require the memory size to be  $\Omega(L)$  in the worst case.  $\square$

## J ADDITIONAL EXPERIMENT DETAILS

### J.1 SETUPS

Model details are shown in Table 1, and other hyperparameters are shown in Table 2.

Table 1: Model specification

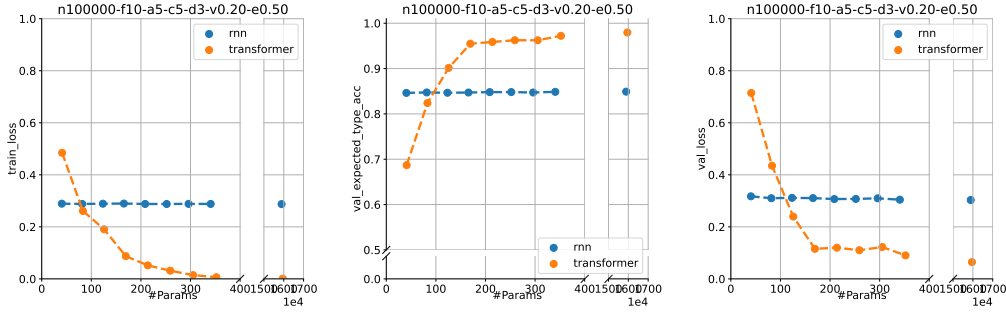
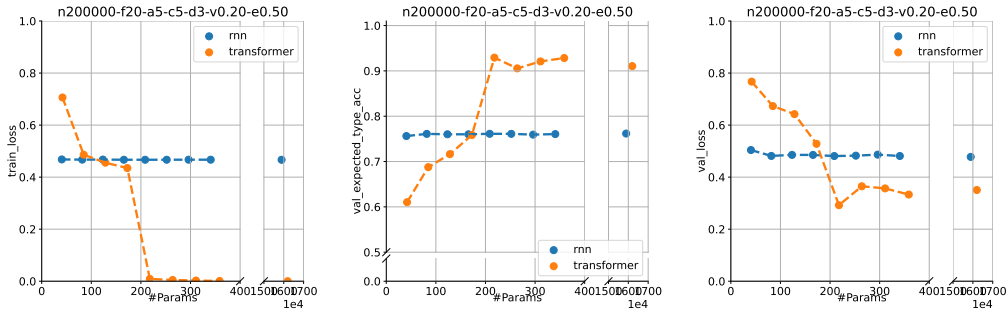
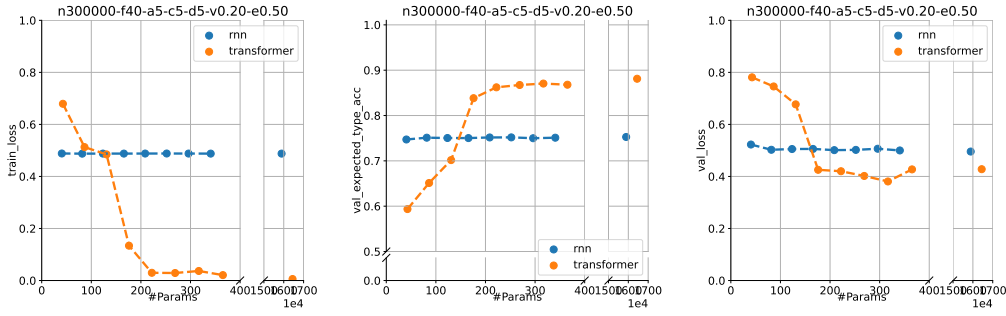
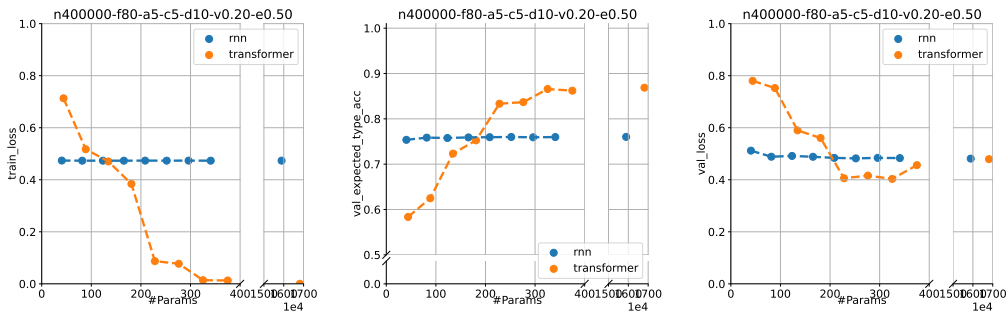
| Specification             | Value                                      |
|---------------------------|--------------------------------------------|
| Transformer               |                                            |
| - Hidden size ( $d_h$ )   | $\{8k \mid 1 \leq k \leq 8\} \cup \{240\}$ |
| - Num attention heads     | 1                                          |
| - Num hidden layers       | 8                                          |
| - Intermediate size       | $2d_h$                                     |
| - Max position embeddings | $\leq 2048$                                |
| RNN                       |                                            |
| - Hidden size             | $\{8k \mid 1 \leq k \leq 8\} \cup \{256\}$ |
| - Num layers              | 8                                          |

Table 2: Hyperparameters of experiments

| Hyperparameter   | Value                                                                      |
|------------------|----------------------------------------------------------------------------|
| Dataset          |                                                                            |
| - $(n, f, d)$    | $\{(100000, 10, 3), (200000, 20, 5), (300000, 40, 10), (400000, 80, 20)\}$ |
| - $(a, c, v, e)$ | $(5, 5, 0.2, 0.5)$                                                         |
| Number of epochs | 80                                                                         |
| Train batch size | 512                                                                        |
| Optimizer        | Adam                                                                       |
| LR scheduler     | Linear warmup-decay                                                        |
| - Warmup min lr  | $1 \times 10^{-5}$                                                         |
| - Warmup max lr  | $1 \times 10^{-3}$                                                         |
| - Warmup steps   | 990                                                                        |

### J.2 ADDITIONAL RESULTS

Figures 4,5,6,7 include other metrics (train loss, accuracies for expected type in validation set, and validation loss) in the experiments. Note that for the expressive power of the models, training accuracies are better indicators.

Figure 4: Figures for the dataset with  $(f, a, c, d, v, e) = (10, 5, 5, 3, 0.2, 0.5)$ .Figure 5: Figures for the dataset with  $(f, a, c, d, v, e) = (20, 5, 5, 3, 0.2, 0.5)$ .Figure 6: Figures for the dataset with  $(f, a, c, d, v, e) = (40, 5, 5, 5, 0.2, 0.5)$ .Figure 7: Figures for the dataset with  $(f, a, c, d, v, e) = (80, 5, 5, 10, 0.2, 0.5)$ .