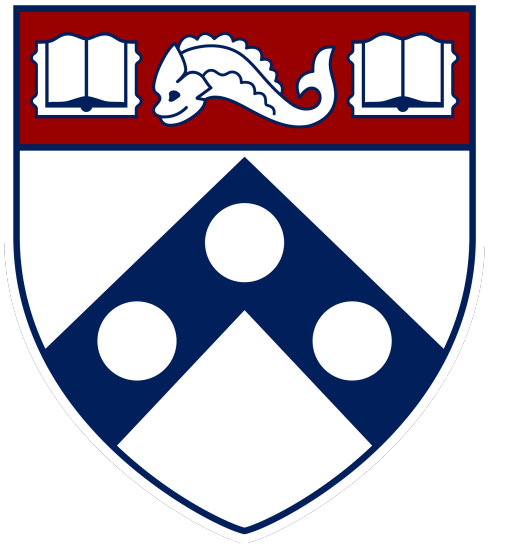


Compiling to transformers

Joey Velez-Ginorio

Department of Computer & Information Science, University of Pennsylvania



Introduction

The representations of transformers are easy to measure but hard to understand; we cannot reliably predict or control behavior from them. To address this, recent work proposed compiling high-level programs to transformers as a tool for studying these representations. By compiling programs to transformers, the algorithms driving their representations are known, and techniques for understanding them can be verified against a ground truth.

However, prior compilers make a *closed program* assumption. They cannot compile programs with free variables. As a consequence, the program must specify the entire algorithm a transformer implements, restricting applicability to hand-constructed transformers rather than ones trained by gradient descent.

We resolve that limitation with Cajal, a language whose programs compile to transformers which are subsequently trained by gradient descent. **We prove its compiler correct, and conduct experiments where we compile a conditional reverse program into the attention head of a transformer, showing that transformers can learn to use compiled programs which specify part of their behavior.** Importantly, this enables verification of interpretability techniques against partial ground truths in transformers trained by gradient descent.

Methods

- Using programming language theory, we provide a formal specification for a compiler from Cajal to differentiable and multilinear maps (Fig. 1).
- Using the compiler, we can embed algorithms into the attention heads of transformers to influence their learning dynamics and behavior (Fig. 3).

$$\begin{aligned}
 & \llbracket - \rrbracket : \text{Program} \rightarrow k\text{-linear map} \\
 & \llbracket x : \tau \vdash x : \tau \rrbracket(\vec{x}) = \vec{x} \\
 & \llbracket \emptyset \vdash * : 1 \rrbracket(0) = 1 \\
 & \llbracket \Delta \vdash \text{inj}_1 e : \tau_1 \oplus \tau_2 \rrbracket(\vec{\sigma}) = (\llbracket e \rrbracket, 0_{[\tau_2]}) \\
 & \llbracket \Delta \vdash \text{inj}_2 e : \tau_1 \oplus \tau_2 \rrbracket(\vec{\sigma}) = (0_{[\tau_1]}, \llbracket e \rrbracket) \\
 & \llbracket \Delta \vdash (e_1, e_2) : \tau_1 \times \tau_2 \rrbracket(\vec{\sigma}) = (\llbracket e_1 \rrbracket, \llbracket e_2 \rrbracket) \\
 & \llbracket \Delta_k \circ \Delta_v \vdash \{e_k : e_v\} : \tau_k \rightarrow \tau_v \rrbracket(\vec{\sigma}) = \vec{x}_q \mapsto \sum_{i=1}^n \langle \llbracket e_k \rrbracket_i, \vec{x}_q \rangle \cdot \llbracket e_v \rrbracket_i \\
 & \llbracket \Delta_1 \circ \Delta_2 \vdash \text{let } x = e_1 \text{ in } e_2 : \tau_2 \rrbracket(\vec{\sigma}) = \llbracket e_2 \rrbracket(\llbracket e_1 \rrbracket) \\
 & \llbracket \Delta_1 \circ \Delta_2 \vdash \text{case } e \text{ of } \{\text{inj}_1 x_1 \Rightarrow e_1; \text{inj}_2 x_2 \Rightarrow e_2\} : \tau_2 \rrbracket(\vec{\sigma}) = \llbracket e_1 \rrbracket(\llbracket e \rrbracket_1) + \llbracket e_2 \rrbracket(\llbracket e \rrbracket_2) \\
 & \llbracket \Delta \vdash \pi_i e : \tau_i \rrbracket(\vec{\sigma}) = \llbracket e \rrbracket_i \\
 & \llbracket \Delta_{kv} \circ \Delta_q \vdash e_{kv}(e_q)_{\mathcal{R}} : \tau_v \rrbracket(\vec{\sigma}) = \llbracket e_{kv} \rrbracket(\llbracket \mathcal{R} \rrbracket(\llbracket e_q \rrbracket)) \\
 & \llbracket \Delta \vdash e_1 \sqcap e_2 : \tau \rrbracket(\vec{\sigma}) = \llbracket e_1 \rrbracket + \llbracket e_2 \rrbracket \\
 & \llbracket \Delta \vdash e_1; e_2 : \tau \rrbracket(\vec{\sigma}) = \llbracket e_1 \rrbracket \cdot \llbracket e_2 \rrbracket
 \end{aligned}$$

Fig. 1. Formal specification of the compiler

- We exploit *linear type theory* to define a programming language whose programs compile to multilinear maps.
- The compiler takes as input a program and returns a multilinear map, as shown below.

$$\llbracket x : \text{Bool}, y : \text{Bool} \vdash \text{if } x \text{ then } y \text{ else } y : \text{Bool} \rrbracket \left(\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \right) = \begin{bmatrix} x_1 y_1 + x_2 y_1 \\ x_1 y_2 + x_2 y_2 \end{bmatrix}$$

Fig. 2. Example output of the compiler

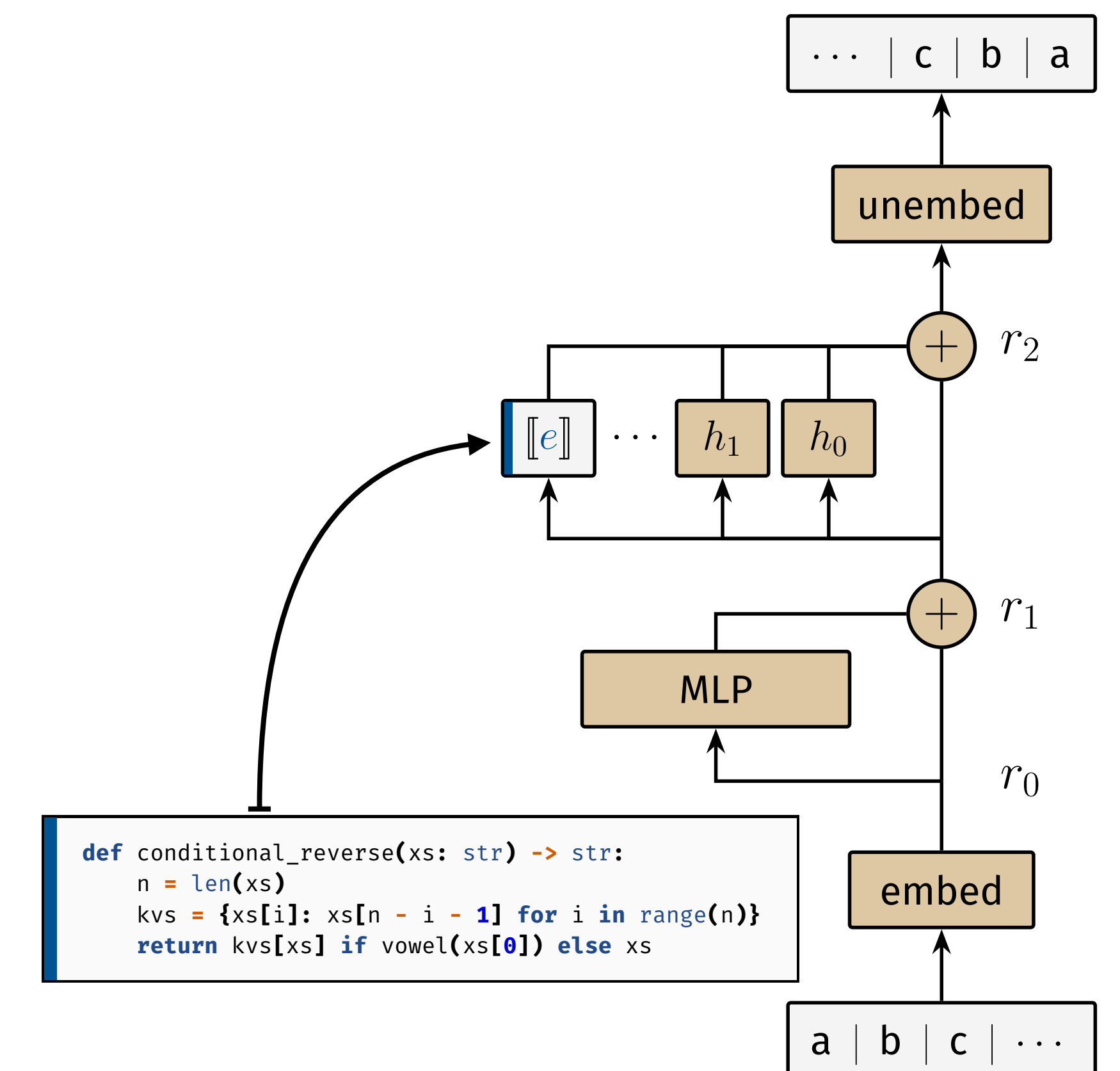


Fig. 3. Transformer with compiled attention head

Future Work

- There are additional constraints needed on the compiler to ensure the multilinear maps we target are compatible with the constraints on functions leveraged in large-scale pretraining of LLMs.
- We believe these constraints are akin to the constraints by an *ordered type system*. We are working on a variant of Cajal with such a system, aiming to compile algorithms into LLMs with billions of parameters.
- We are additionally interested in what happens when we compile generically useful algorithms like *fold* or *scan* into an LLM. The hope is that, as with simpler algorithms, these generic algorithms act as generic inductive biases for learning.