

(1) Source Engram
(trained on corpus C)

Corpus C
(e.g., 50M tokens)

Source Model A
(e.g., Pythia-160M)

Engram Training
next-token prediction
(unsupervised)

Engram Table E^A

learned & frozen

Source model and Engram table
are frozen at end of Phase 1.

(2) Input-Space Addressing
(tokenizer-agnostic)

Raw tokens from any tokenizer

τ_1 τ_2 τ_3 $\dots$ τ_4

Canonicalization
 $P(\cdot)$: NFKC + lowercase +
accent stripping

Canonical tokens
 c_1 c_2 c_3 $\dots$ c_4

N-gram Hashing ($\varphi_{n,k}$)
 $n = 1 \dots N, k = 1 \dots K$
Memory indices (per head)

(3) Frozen Engram Memory Table
Engram Memory Table E
shared across models frozen

Head 1 ☐ ☐ ☐ $\dots$ ☐

Head K ☐ ☐ ☐ $\dots$ ☐

Slot 0 Slot 1 Slot 2 Slot ($M-1$)

Example: tokenizer-agnostic
canonicalization

Raw text: "The Cafe is open"

Tokenizer A:

The Caf e is open

Tokenizer B:

The Cafe is open

Tokenizer C:

The Ca fe is open

Canonical word units:

[the] [cafe] [is] [open]

2 grams ($n=2$):

(the, cafe) (cafe, is) (is, open)

Identical hash slots
identical memory addresses

(4) Target Model + Trainable Reader (Extractor)

Target Model B

Frozen

Backbone hidden state h_t

Trainable Reader W^B
(learned in Phase 2 only)

Key Projection
 W_K^B

Key Vectors
 h_t

Value Projection
 W_V

Value Vectors
 v_t

Gate (context-aware)
 $\alpha_t = \sigma\left(\frac{h_t^T \text{RMSNorm}(k_t)}{\sqrt{d_B}}\right)$

Lookup
per head

$\alpha_t \cdot v_t$

Augmented
hidden state

Reader variants

Single
Layer

Shared Value
Projection W_V $\rightarrow v_t$ multi-branch (R-branches)

Key Projection
 $W_K^{(1)}$ $\rightarrow k_l^{(1)}$

Gate (context-aware)
 $\alpha_t^{(1)} = \sigma\left(\frac{h_t^T \text{RMSNorm}(k_l^{(1)})}{\sqrt{d_B}}\right) \rightarrow \alpha_t^{(1)} \cdot v_t$

Key Projection
 $W_K^{(2)}$ $\rightarrow k_l^{(2)}$

Gate (context-aware)
 $\alpha_t^{(2)} = \sigma\left(\frac{h_t^T \text{RMSNorm}(k_l^{(2)})}{\sqrt{d_B}}\right) \rightarrow \alpha_t^{(2)} \cdot v_t$

Key Projection
 $W_K^{(R)}$ $\rightarrow k_l^{(R)}$

Gate (context-aware)
 $\alpha_t^{(R)} = \sigma\left(\frac{h_t^T \text{RMSNorm}(k_l^{(R)})}{\sqrt{d_B}}\right) \rightarrow \alpha_t^{(R)} \cdot v_t$

$h_t \leftarrow h_t + o_t$

Average Sum o_t

Multi-layer variant

Layer 1
($h^{(1)}$)

Layer 2
($h^{(2)}$)

Layer 10
($h^{(10)}$)

Layer L
($h^{(L)}$)

Reader

Reader

+

+