
Reparameterized LLM Training via Orthogonal Equivalence Transformation

Zeju Qiu¹ Simon Buchholz¹ Tim Z. Xiao¹ Maximilian Dax¹ Bernhard Schölkopf¹ Weiyang Liu^{1,2,*}

¹Max Planck Institute for Intelligent Systems, Tübingen ²The Chinese University of Hong Kong

Abstract

While Large language models (LLMs) are driving the rapid advancement of artificial intelligence, effectively and reliably training these large models remains one of the field’s most significant challenges. To address this challenge, we propose POET, a novel reParameterized training algorithm that uses Orthogonal Equivalence Transformation to optimize neurons. Specifically, POET reparameterizes each neuron with two learnable orthogonal matrices and a fixed random weight matrix. Because of its provable preservation of spectral properties of weight matrices, POET can stably optimize the objective function with improved generalization. We further develop efficient approximations that make POET flexible and scalable for training large-scale neural networks. Extensive experiments validate the effectiveness and scalability of POET in training LLMs.

1 Introduction

Recent years have witnessed the increasing popularity of large language models (LLMs) in various applications, such as mathematical reasoning [13] and program synthesis [3] and decision-making [77]. Current LLMs are typically pre-trained using enormous computational resources on massive datasets containing trillions of tokens, with each training run that can take months to complete. Given such a huge training cost, how to effectively and reliably train them poses significant challenges.

The *de facto* way for training LLMs is to directly optimize weight matrices with the Adam optimizer [37, 55]. While conceptually simple, this direct optimization can be computationally intensive (due to the poor scaling with model size) and requires careful hyperparameter tuning to ensure stable convergence. More importantly, its generalization can remain suboptimal even if the training loss is perfectly minimized [36]. To stabilize training and enhance generalization, various weight regularization methods [4, 10, 12, 47, 49, 79] and weight normalization techniques [28, 38, 39, 50, 52, 54] have been proposed. Most of these methods boil down to improving spectral properties of weight matrices (*i.e.*, singular values) either explicitly or implicitly. Intuitively, the spectral norm of a weight matrix (*i.e.*, the largest singular value) provides an upper bound on how much a matrix can amplify the input vectors, which connects to the generalization properties. In general, smaller spectral norms (*i.e.*, better smoothness) are considered to be associated with stronger generalization, which inspires explicit spectrum control [33, 59, 67, 79]. Theoretical results [6] also suggest that weight matrices with bounded spectrum can provably guarantee generalization. Given the importance of the spectral properties of weight matrices, *what prevents us from controlling them during LLM training?*

- **Inefficacy of spectrum control:** Existing spectrum control methods constrain only the largest singular value, failing to effectively regularizing the full singular value spectrum. Moreover, there is also no guarantee for spectral norm regularization to effectively control the largest singular value.
- **Computational overhead:** Both spectral norm regularization [79] and spectral normalization [59] require computing the largest singular value of weight matrices. Even with power iteration, this still adds a significant overhead to the training process, especially when training large neural networks. Additionally, spectral regularization does not scale efficiently with increasing model size.

*Project lead & Corresponding author

Project page: spherelab.ai/poet

To achieve effective weight spectrum control without the limitations above, we propose POET, a reParameterized training algorithm that uses Orthogonal Equivalence Transformation to indirectly learn weight matrices. Specifically, POET reparameterizes a weight matrix $W \in \mathbb{R}^{m \times n}$ with RW_0P where $W_0 \in \mathbb{R}^{m \times n}$ is a randomly initialized weight matrix, $R \in \mathbb{R}^{m \times m}$ and $P \in \mathbb{R}^{n \times n}$ are two orthogonal matrices. Instead of optimizing weight matrices directly, POET keeps the randomly initialized weight matrix W_0 unchanged during training and learns two orthogonal matrices R, P to transform W_0 . This reparameterization preserves the singular values of weights while allowing flexible optimization of the singular vectors. POET effectively addresses the above limitations:

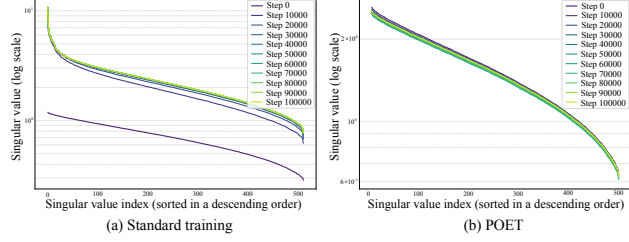


Figure 1: Training dynamics of singular values of the same weight matrix in a LLaMA model. Standard training on the left strictly follows the common practice for training LLMs (direct optimization with AdamW). POET on the right uses the proposed approximation for large-scale LLM training. The slight (almost negligible) singular value changes in POET are due to numerical and approximation error.

- **Strong spectrum control:** Because orthogonal transformations do not change the singular values of weight matrices, POET keeps the weight spectrum the same as the randomly initialized weight matrices (empirically validated by Figure 1 even with approximations). Through the initialization scheme, POET thus directly controls the singular value distribution of its weight matrices. As a result, and in contrast to standard LLM training, POET matrices avoid undesirable large singular values after training (Figure 1 and Appendix I). To further facilitate the POET algorithm, we introduce two new initialization schemes: normalized Gaussian initialization and uniform spectrum initialization, which can ensure the resulting weight matrices have bounded singular values.
- **Efficient approximation:** While a naive implementation of POET can be computationally expensive, its inherent flexibility opens up opportunities for efficient and scalable training. To address the key challenge of optimizing large orthogonal matrices, we introduce two levels of approximations:
 - *Stochastic primitive optimization:* The first-level approximation aims to reduce the number of learnable parameters when optimizing a large orthogonal matrix. To this end, we propose the stochastic primitive optimization (SPO) algorithm. Given a large orthogonal matrix $R \in \mathbb{R}^{m \times m}$, SPO factorizes it into a product of primitive orthogonal matrices, each involving significantly fewer trainable parameters. These primitives are constructed by parameterizing randomly sampled submatrices of the full matrix. This factorization is implemented as a memory-efficient iterative algorithm that sequentially updates one primitive orthogonal matrix at a time. To improve the expressiveness of the sequential factorization, we adopt a merge-then-reinitialize trick, where we merge each learned primitive orthogonal matrix into the weight matrix, and then reinitialize the primitive orthogonal matrix to be identity after every fixed number of iterations.
 - *Approximate orthogonality via Cayley-Neumann parameterization:* The second-level approximation addresses how to maintain orthogonality without introducing significant computational overhead. To achieve this, we develop the Cayley-Neumann parameterization (CNP) which approximates the Cayley orthogonal parameterization [48, 65] with Neumann series. Our merge-then-reinitialize trick can effectively prevent the accumulation of approximation errors.

POET can be viewed as a natural generalization of orthogonal training [48, 51, 65], wherein the model training is done by learning a layer-shared orthogonal transformation for neurons. Orthogonal training preserves the hyperspherical energy [47, 49] within each layer—a quantity that characterizes pairwise neuron relationships on the unit hypersphere. While preserving hyperspherical energy proves effective for many finetuning tasks [51], it limits the flexibility of pretraining. Motivated by this, POET generalizes energy preservation to spectrum preservation and subsumes orthogonal training as its special case. The better flexibility of POET comes from its inductive structures for preserving weight spectrum, rather than more learnable parameters. We empirically validate that POET achieves better pretraining performance than orthogonal training given the same budget of parameters.

To better understand how POET functions, we employ *vector probing* to analyze the learning dynamics of the orthogonal matrices. Vector probing evaluates an orthogonal matrix R using a fixed, randomly generated unit vector v by computing $v^\top Rv$ which corresponds to the cosine similarity between Rv and v . By inspecting the cosine similarities of seven orthogonal matrices throughout training, we

observe that the learning process can be divided into three distinct phases (Figure 2): (1) *conical shell searching*: The cosine starts at 1 (i.e., $\mathbf{R}$ is the identity) and gradually converges to a stable range of $[0.6, 0.65]$, which we observe consistently across all learnable orthogonal matrices.

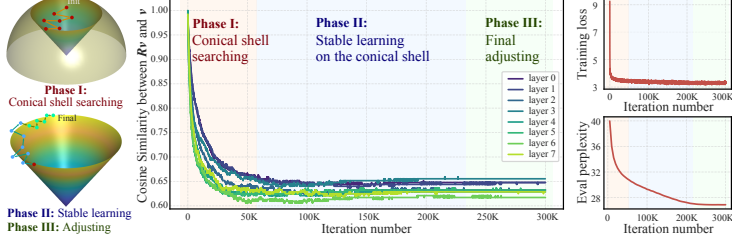


Figure 2: POET’s three learning phases. Left: illustration; Middle: angle; Right: loss and validation.

This suggests that $\mathbf{R}$ transforms $\mathbf{v}$ into a thin conical shell around its original direction. (2) *stable learning on the conical shell*: The cosine remains within this range while the model begins to learn stably. Despite the cosine plateauing, validation perplexity continues to improve almost linearly. (3) *final adjusting*: Learning slows and eventually halts as the learning rate approaches zero. We provide an in-depth discussion and empirical results in Appendix A, G. Our major contributions are:

- We introduce POET, a novel training framework that provably preserves spectral properties of weight matrices through orthogonal equivalence transformation.
- To enhance POET’s scalability, we develop two simple yet effective approximations: stochastic principal submatrix optimization for large orthogonal matrices and the Cayley-Neumann parameterization for efficient representation of orthogonal matrices.
- We empirically validate POET’s training stability and generalization across multiple model scales.

2 From Energy-preserving Training to Spectrum-preserving Training

Orthogonal training [48, 51, 65] is a framework to train neural networks by learning a layer-shared orthogonal transformation for neurons in each layer. For a weight matrix $\mathbf{W} = \{\mathbf{w}_1, \dots, \mathbf{w}_n\} \in \mathbb{R}^{m \times n}$ where $\mathbf{w}_i \in \mathbb{R}^m$ is the i -th neuron, the layer’s forward pass is $\mathbf{y} = \mathbf{W}^\top \mathbf{x}$ with input $\mathbf{x} \in \mathbb{R}^m$ and output $\mathbf{y} \in \mathbb{R}^n$. Unlike standard training, which directly optimizes $\mathbf{W}$, orthogonal training keeps $\mathbf{W}$ fixed at its random initialization $\mathbf{W}_0 = \{\mathbf{w}_1^0, \dots, \mathbf{w}_n^0\}$ and instead learns an orthogonal matrix $\mathbf{R} \in \mathbb{R}^{m \times m}$ to jointly transform all neurons in the layer. The forward pass becomes $\mathbf{y} = (\mathbf{R}\mathbf{W}_0)^\top \mathbf{x}$. The effective weight matrix is $\mathbf{W}_R = \{\mathbf{w}_1^R, \dots, \mathbf{w}_n^R\}$ where $\mathbf{w}_i^R = \mathbf{R}\mathbf{w}_i$. A key property is its *preservation of hyperspherical energy*. With $\hat{\mathbf{w}}_i = \mathbf{w}_i / \|\mathbf{w}_i\|$, orthogonal training ensures

$$\text{HE}(\mathbf{W}_0) := \sum_{i \neq j} \|\hat{\mathbf{w}}_i^0 - \hat{\mathbf{w}}_j^0\|^{-1} = \sum_{i \neq j} \|\mathbf{R}\hat{\mathbf{w}}_i - \mathbf{R}\hat{\mathbf{w}}_j\|^{-1} =: \text{HE}(\mathbf{W}^R), \quad (1)$$

where hyperspherical energy $\text{HE}(\cdot)$ characterizes the uniformity of neurons on a unit hypersphere. Prior work [47–49, 76] has shown that energy-preserving training can effectively improve generalization. Orthogonal finetuning (OFT) [51, 65] also demonstrates that finetuning foundation models while preserving hyperspherical energy achieves a favorable trade-off between efficient adaptation to downstream tasks and retention of pretraining knowledge. While the hyperspherical energy preservation is effective for finetuning, it can be too restrictive for pretraining. To allow greater flexibility in the pretraining phase, we relax the constraint from preserving hyperspherical energy to preserving the singular-value spectrum instead. By inherently maintaining the spectrum, energy-preserving training is a special case of spectrum-preserving training. As a generalization, spectrum-preserving training learns a transformation $\mathcal{T} : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}^{m \times n}$ that preserves the spectrum:

$$\{\sigma_1(\mathcal{T}(\mathbf{W}_0)), \sigma_2(\mathcal{T}(\mathbf{W}_0)), \dots, \sigma_{\min(m,n)}(\mathcal{T}(\mathbf{W}_0))\} = \{\sigma_1(\mathbf{W}_0), \sigma_2(\mathbf{W}_0), \dots, \sigma_{\min(m,n)}(\mathbf{W}_0)\}, \quad (2)$$

where $\sigma_i(\mathbf{W}_0)$ denotes the i -th singular value of $\mathbf{W}_0$ (sorted by descending order with σ_1 being the largest singular value). How we instantiate the transformation $\mathcal{T}$ results in different algorithms. Generally, $\mathcal{T}$ is a spectrum-preserving map, and can be either linear [42] or nonlinear [5]. If we only consider $\mathcal{T}$ to be a linear map, then Theorem 1 can fully characterize the form of $\mathcal{T}$:

Theorem 1 (Simplified results from [42]). *For a linear map $\mathcal{T} : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}^{m \times n}$ ($m \neq n$), if $\sigma_1(\mathcal{T}(\mathbf{W})) = \sigma_1(\mathbf{W})$ always holds for all $\mathbf{W} \in \mathbb{R}^{m \times n}$, then the linear map $\mathcal{T}$ must be of the following form: $\mathcal{T}(\mathbf{W}) = \mathbf{R}\mathbf{W}\mathbf{P}$, for all $\mathbf{W} \in \mathbb{R}^{m \times n}$ where $\mathbf{R} \in \mathbb{R}^{m \times m}$ and $\mathbf{P} \in \mathbb{R}^{n \times n}$ are some fixed elements in orthogonal groups $O(m)$ and $O(n)$, respectively.*

All parameterizations for the linear map $\mathcal{T}$ can be expressed as $\mathcal{T}(\mathbf{W}) = \mathbf{R}\mathbf{W}\mathbf{P}$, where $\mathbf{R}$ and $\mathbf{P}$ are orthogonal matrices. For instance, OFT is an energy-preserving method (a special case of spectrum-preserving training), where the map simplifies to $\mathcal{T}(\mathbf{W}) = \mathbf{R}\mathbf{W}\mathbf{I}$, with $\mathbf{I}$ as the identity.

POET preserves hyperspherical energy under isotropic Gaussian initialization. [48] shows that weight matrices that are initialized by zero-mean isotropic Gaussian distribution (*e.g.*, [17]) are guaranteed to have small hyperspherical energy. Because zero-mean isotropic Gaussian is invariant under orthogonal transformation, RWP also has small energy (see Section 4 and Appendix B).

3 Reparameterized Training via Orthogonal Equivalence Transformation

This section introduces the POET framework, which reparameterizes each neuron as the product of a fixed random weight matrix and two learnable orthogonal matrices applied on both sides. POET serves as a specific implementation of spectrum-preserving training. Inspired by Theorem 1, it parameterizes the spectrum-preserving transformation $\mathcal{T}$ using a left orthogonal matrix that transforms the column space of the weight matrix and a right orthogonal matrix that transforms its row space.

3.1 General Framework

Following the general form of spectrum-preserving linear maps discussed in the last section, POET reparameterizes the neuron as RW_0P , where $W_0 \in \mathbb{R}^{m \times n}$ is a randomly initialized weight matrix that remains fixed during training, and $R \in \mathbb{R}^{m \times m}$, $P \in \mathbb{R}^{n \times n}$ are trainable orthogonal matrices. This reparameterization effectively applies an orthogonal equivalence transformation (OET) to random weight matrices. Specifically, OET is a double-sided transformation, defined as $\text{OET}(W; R, P) = RW_0P$, where the input matrix W is multiplied on the left and on the right by orthogonal matrices R and P , respectively. The forward pass of POET can be thus written as

$$y = W_{RP}^\top x = (RW_0P)^\top x, \quad \text{s.t. } \{R^\top R = RR^\top = I, P^\top P = PP^\top = I\}, \quad (3)$$

where R and P can be merged into a single weight matrix $W_{RP} = RW_0P$ after training. Therefore, the inference speed of POET-trained neural networks is the same as conventionally trained ones.

Spectrum control. POET can be interpreted as learning weight matrices by simultaneously transforming their left singular vectors and right singular vectors while keeping the singular values unchanged. Given the singular value decomposition (SVD) $W_0 = U\Sigma_0V^\top$, the reparameterized neuron weight matrix becomes $W_{RP} = RU\Sigma_0V^\top P$ where both RU and $V^\top P$ are orthogonal matrices. This effectively constitutes an SVD of W_{RP} . It is also straightforward to verify that the spectral properties of W_{RP} remain identical to those of the initial matrix W_0 .

Neuron initialization. Since POET preserves the spectral properties of the initial weight matrix W_0 , the choice of initialization plays a critical role. We consider two common schemes: (1) *standard initialization*, which samples from a zero-mean Gaussian with fixed variance (the default choice for LLaMA models); and (2) *Xavier initialization* [17], which uses a zero-mean Gaussian with variance scaled by the layer dimensions. To facilitate POET, we propose two new initialization schemes. The first method, *uniform-spectrum initialization*, applies SVD to a standard initialization and sets all singular values to 1, balancing spectral properties throughout training. The second, *normalized Gaussian initialization*, normalizes neurons drawn from a zero-mean Gaussian with fixed variance. This is directly inspired by prior work showing that normalized neurons improve convergence [48, 50, 52]. To ensure that the POET-reparameterized network is statistically equivalent to a standard network at initialization, we always initialize both orthogonal matrices as identity matrices.

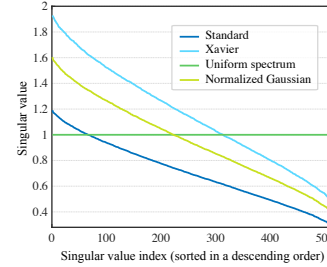


Figure 3: Singular values of a weight matrix of size 512×1376 , randomly generated by different initialization schemes.

3.2 Efficient Approximation to Orthogonality

POET is conceptually simple, requiring only the optimization of two orthogonal matrices. However, these matrices are typically large, and naively optimizing them leads to significant computational challenges. We start by introducing the following efficient approximations.

3.2.1 Stochastic Primitive Optimization

The core idea of SPO is inspired by how QR factorization is performed using Givens rotations and Householder transformations. Both methods construct a large orthogonal matrix R by sequentially applying primitive orthogonal transformations (*e.g.*, Givens rotations or Householder reflections), *i.e.*, $R = \prod_{i=1}^c G_i$, where G_i denotes the i -th primitive orthogonal matrix. While each G_i is of the same size as R , it is parameterized by significantly fewer degrees of freedom. See Figure 4

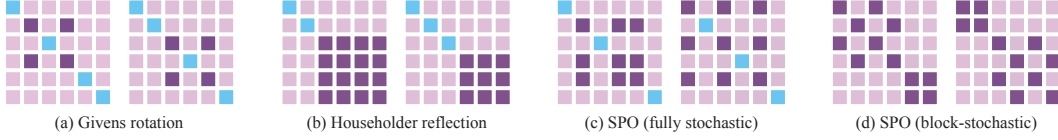


Figure 4: Examples of the primitive orthogonal transformation matrix $\mathbf{G}_i$ in different orthogonalizations (two examples for each method). Note that, blue blocks represent 1, light purple blocks denote 0 and deep purple blocks are the actual orthogonal parameterization to be learned.

for an illustration. Both Givens rotation and Householder reflection use relatively low-capacity parameterizations—for example, each Givens rotation $\mathbf{G}_i$ involves only a single effective parameter—which limits their efficiency in representing the full orthogonal matrix. SPO follows a similar idea of factorizing the original orthogonal matrix into multiple primitive orthogonal matrices. However, unlike Givens and Householder methods, SPO treats the number of effective parameters in each primitive matrix as a tunable hyperparameter and adopts a stochastic sparsity pattern.

Fully stochastic SPO. The basic idea of fully stochastic SPO is to randomly sample a small submatrix and enforce its orthogonality, allowing it to be easily extended to a full orthogonal matrix by embedding it within an identity matrix—a process similar to Givens or Householder transformations. To represent a large orthogonal matrix $\mathbf{R} \in \mathbb{R}^{m \times m}$, we start by defining c index sets $\mathbf{S}^j = \{s_1^j, \dots, s_b^j\} \subseteq \{1, \dots, m\}$ ($j \in [1, c]$), where each set has cardinality $|\mathbf{S}^j| = b$, a hyperparameter controlling the number of effective parameters of a primitive orthogonal matrix. $\mathbf{S}^j, \forall j$ are randomly sampled from the full indices $\{1, \dots, m\}$. Let $\tilde{\mathbf{G}}_j \in \mathbb{R}^{b \times b}$ be a small orthogonal matrix, and $\mathbf{D}(\mathbf{S}^j) = \{e(s_1^j), \dots, e(s_b^j)\} \in \mathbb{R}^{m \times b}$ be a selection matrix, where $e(k)$ is the standard basis vector with a 1 in the k -th position and 0 elsewhere. The factorization is given by

$$\mathbf{R} = \prod_{i=1}^c \underbrace{(\mathbf{I}_m + \mathbf{D}(\mathbf{S}^i) \cdot (\tilde{\mathbf{G}}_i - \mathbf{I}_b) \cdot \mathbf{D}(\mathbf{S}^i)^\top)}_{\mathbf{G}_i: \text{The } i\text{-th primitive orthogonal matrix}}, \quad \text{s.t. } \tilde{\mathbf{G}}_i^\top \tilde{\mathbf{G}}_i = \tilde{\mathbf{G}}_i \tilde{\mathbf{G}}_i^\top = \mathbf{I}_b, \quad \forall i, \quad (4)$$

where $\mathbf{D}(\mathbf{S}^i) \cdot (\mathbf{A}) \cdot \mathbf{D}(\mathbf{S}^i)^\top$ is a projector that replaces the $b \times b$ sub-block with $\mathbf{A}$. $\mathbf{I}_m$ and $\mathbf{I}_b$ are identity matrices of size $m \times m$ and $b \times b$, respectively. To efficiently parameterize small orthogonal matrices $\tilde{\mathbf{G}}_i$, we can use the CNP introduced in the next section.

Block-stochastic SPO. While fully stochastic SPO is simple, it may fail to transform all neuron dimensions because the identity matrix leaves part of the space unchanged. See the blue blocks in Figure 4(c) as an example. To address this, we propose block-stochastic SPO, which first constructs a block-diagonal orthogonal matrix with small blocks for parameter efficiency, and then applies a random permutation to enhance expressiveness by randomizing the sparsity pattern. Block-stochastic SPO transforms all neuron dimensions simultaneously, as shown in Figure 4(d). Formally we have

$$\mathbf{R} = \prod_{i=1}^c \underbrace{(\Psi_i^\top \cdot \text{Diag}(\tilde{\mathbf{G}}_i^1, \tilde{\mathbf{G}}_i^2, \dots, \tilde{\mathbf{G}}_i^{\lceil \frac{m}{b} \rceil}) \cdot \Psi_i)}_{\mathbf{G}_i: \text{The } i\text{-th primitive orthogonal matrix}}, \quad \text{s.t. } (\tilde{\mathbf{G}}_i^j)^\top \tilde{\mathbf{G}}_i^j = \tilde{\mathbf{G}}_i^j (\tilde{\mathbf{G}}_i^j)^\top = \mathbf{I}_b, \quad \forall i, j, \quad (5)$$

where $\tilde{\mathbf{G}}_i^j \in \mathbb{R}^{b \times b}$ is the j -th block of the block diagonal matrix, and $\Psi_i, \forall i$ are all random permutation matrices. As long as each diagonal block $\tilde{\mathbf{G}}_i^j$ is an orthogonal matrix, both $\mathbf{G}_i$ and $\mathbf{R}$ are also orthogonal matrices. We also use CNP to efficiently parameterize each orthogonal block $\tilde{\mathbf{G}}_i^j$.

The merge-then-reinitialize trick. The factorizations in Equation (4) and (5) offer a simple approach to optimizing large orthogonal matrices by sequentially updating primitive orthogonal matrices. However, storing all previous primitives incurs high GPU memory overhead. To mitigate this, we propose the merge-then-reinitialize trick, where the learned primitive orthogonal matrix can be merged into the weight matrix after every certain number of iterations, and then reinitialized to the identity matrix. After reinitialization, stochastic sampling is repeated to select a new index set (in fully stochastic SPO) or generate a new permutation (in block-stochastic SPO). This trick allows only one primitive matrix to be stored at a time, substantially reducing GPU memory usage.

3.2.2 Cayley-Neumann Parameterization

The classic Cayley parameterization generates an orthogonal matrix $\mathbf{R}$ in the form of $\mathbf{R} = (\mathbf{I} + \mathbf{Q})(\mathbf{I} - \mathbf{Q})^{-1}$ where $\mathbf{Q}$ is a skew-symmetric matrix satisfying $\mathbf{Q} = -\mathbf{Q}^\top$. A minor caveat of this parameterization is that it only produces orthogonal matrices with determinant 1 (*i.e.*, elements of the special orthogonal group), but empirical results in [48, 51, 65] indicate that this constraint does not hurt performance. However, the matrix inverse in the original Cayley parameterization introduces

numerical instability and computational overhead, limiting its scalability to large orthogonal matrices. To address this, we approximate the matrix inverse using a truncated Neumann series:

$$\mathbf{R} = (\mathbf{I} + \mathbf{Q})(\mathbf{I} - \mathbf{Q})^{-1} = (\mathbf{I} + \mathbf{Q}) \cdot \left(\sum_{i=0}^{\infty} \mathbf{Q}^i \right) \approx (\mathbf{I} + \mathbf{Q}) \cdot \left(\sum_{i=0}^k \mathbf{Q}^i \right) = (\mathbf{I} + \mathbf{Q}) \cdot \left(\mathbf{I} + \sum_{i=1}^k \mathbf{Q}^i \right), \quad (6)$$

where a larger number of approximation terms k leads to a smaller approximation error. By avoiding matrix inversion, the training stability of POET is improved; however, this comes with a price—the approximation is valid only when the Neumann series converges in the operator norm. To initialize orthogonal matrices as identity, we set $\mathbf{Q}$ to a zero matrix in CNP, satisfying the convergence condition initially. As the training progresses, however, updates to $\mathbf{Q}$ may cause its operator norm to exceed 1, violating this condition. Fortunately, our merge-then-reinitialize trick mitigates this issue by periodically resetting $\mathbf{Q}$ to a zero matrix, ensuring its operator norm remains small.

3.2.3 Overall Training Algorithm

Step 1: Initialization. We initialize the weight matrices using normalized Gaussian: $\mathbf{W} \leftarrow \mathbf{W}_0$.

Step 2: Orthogonal matrix initialization. For fully stochastic SPO, we randomly sample an index set $\mathcal{S}$, and parameterize $\tilde{\mathbf{G}}_R \in \mathbb{R}^{b \times b}$ and $\tilde{\mathbf{G}}_P \in \mathbb{R}^{b \times b}$ using CNP (Equation (6)). Both matrices are initialized as identity, so $\mathbf{R}$ and $\mathbf{P}$ also start as identity matrices. For block-stochastic SPO, we sample a random permutation matrix Ψ_R, Ψ_P , and parameterize $\{\tilde{\mathbf{G}}_R^1, \dots, \tilde{\mathbf{G}}_R^{\lceil \frac{m}{b} \rceil}\}$ and $\{\tilde{\mathbf{G}}_P^1, \dots, \tilde{\mathbf{G}}_P^{\lceil \frac{n}{b} \rceil}\}$ using CNP. Then we initialize them as the identity, so $\mathbf{R}$ and $\mathbf{P}$ again starts as identity matrices.

Step 3: Efficient orthogonal parameterization. For fully stochastic SPO, we have $\mathbf{R} = \mathbf{I}_m + \mathbf{D}(\mathcal{S})(\tilde{\mathbf{G}}_R - \mathbf{I}_b)\mathbf{D}(\mathcal{S})^\top$ and $\mathbf{P} = \mathbf{I}_n + \mathbf{D}(\mathcal{S})(\tilde{\mathbf{G}}_P - \mathbf{I}_b)\mathbf{D}(\mathcal{S})^\top$. For block-stochastic SPO, we have $\mathbf{R} = \Psi_R^\top \text{Diag}(\tilde{\mathbf{G}}_R^1, \dots, \tilde{\mathbf{G}}_R^{\lceil \frac{m}{b} \rceil}) \Psi_R$ and $\mathbf{P} = \Psi_P^\top \text{Diag}(\tilde{\mathbf{G}}_P^1, \dots, \tilde{\mathbf{G}}_P^{\lceil \frac{n}{b} \rceil}) \Psi_P$.

Step 4: Inner training loop for updating orthogonal matrices. The equivalent weight matrix in the forward pass is $\mathbf{R}\mathbf{W}\mathbf{P}$. Gradients are backpropagated through $\mathbf{R}$ and $\mathbf{P}$ to update $\tilde{\mathbf{G}}_R, \tilde{\mathbf{G}}_P$ (fully stochastic) or $\tilde{\mathbf{G}}_R^i, \tilde{\mathbf{G}}_P^i, \forall i$ (block-stochastic). This inner loop runs for a fixed number of iterations.

Step 5: Merge-then-reinitialize. The learned orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$ are merged into the weight matrix by $\mathbf{W} \leftarrow \mathbf{R}\mathbf{W}\mathbf{P}$. If not terminated, return to **Step 2** for reinitialization.

4 Discussions and Intriguing Insights

Parameter and memory complexity. By introducing a hyperparameter b as the sampling budget, fully stochastic SPO decouples parameter complexity from the size of the weight matrices. With a small b , POET becomes highly parameter-efficient, though at the cost of slower convergence. This offers users a flexible trade-off between efficiency and speed. In contrast, block-stochastic SPO has parameter complexity dependent on the matrix size (*i.e.*, $m + n$), making it more scalable than AdamW, which requires mn trainable parameters. In terms of memory complexity, both POET variants can be much more efficient than AdamW with a suitable sampling budget b . A comparison of parameter and memory complexity is given in Table 1.

Performance under a constant parameter budget. Since POET optimizes two orthogonal matrices $\mathbf{R}, \mathbf{P}$ simultaneously, a natural question arises: *which matrix should receive more parameter budget under a fixed total constraint?* To investigate this, we conduct a controlled experiment where different ratios of trainable parameters are allocated to $\mathbf{R}$ and $\mathbf{P}$ under a fixed total budget. All other settings (*e.g.*, architecture, data) remain unchanged, with full details provided in the Appendix. We use validation perplexity as the evaluation metric. The total parameter budget matches that of fully stochastic POET with $b = \frac{1}{h}m$ for $\mathbf{R}$ and $b = \frac{1}{h}n$ for $\mathbf{P}$, where $h = 8, 4$, and 3 correspond to small, medium, and large budgets, respectively. We explore seven allocation settings: $\mathbf{R} : \mathbf{P} = 1 : 0$ (*i.e.*, orthogonal training [48, 51, 65]), $0.9 : 0.1$, $0.75 : 0.25$, $0.5 : 0.5$ (*i.e.*, standard POET), $0.25 : 0.75$, $0.1 : 0.9$, and $0 : 1$. Results in Figure 5 show that POET with a balanced allocation between $\mathbf{R}$ and $\mathbf{P}$ yields the best performance.

Method	# trainable params	Memory cost
AdamW	mn	$3mn$
GaLore [82]	mn	$mn + mr + 2nr$
POET (FS)	$b(b-1)$	$mn + 3b(b-1)$
POET (BS)	$\frac{1}{2}(m+n)(b-1)$	$mn + \frac{3}{2}(m+n)(b-1)$

Table 1: Comparison to existing methods. Assume $\mathbf{W} \in \mathbb{R}^{m \times n}$ ($m \leq n$), GaLore with rank r and POET with block size b . FS denotes fully stochastic SPO, and BS denotes block-stochastic SPO.

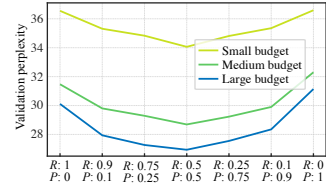


Figure 5: Performance of POET under a constant total parameter budget on $\mathbf{R}, \mathbf{P}$.

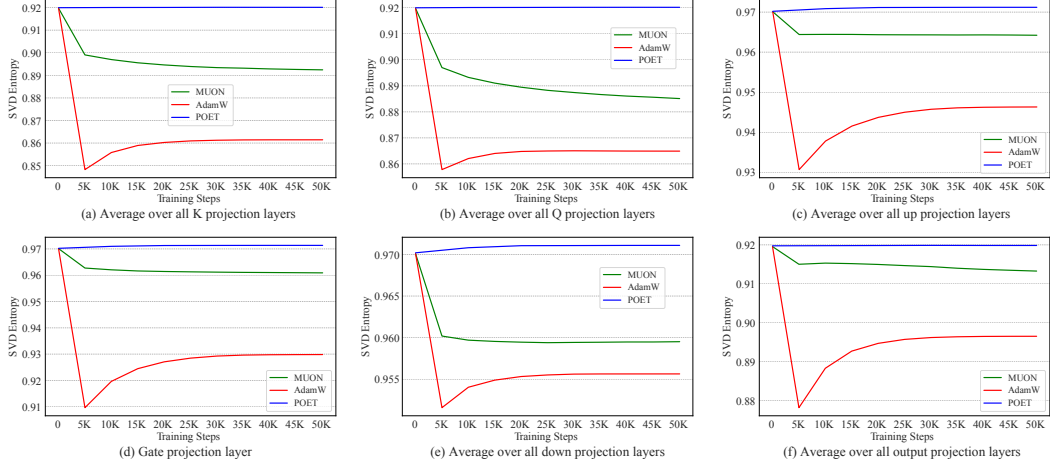


Figure 6: Dynamics comparison of average singular value entropy (singular value diversity) between direct training (AdamW, Muon) and POET.

Guarantees of weight spectrum. For POET with standard and normalized Gaussian initializations, we have proved in Appendix C that the largest and smallest singular values of weights can be bounded. For normalized Gaussian, the bound is only dependent on the row-to-column ratio of the weight matrix. For standard Gaussian, the bound has an extra dependence on the neuron dimension.

Connection to generalization theory. Several generalization results [6, 63, 76] based on bounding the spectral norm of weight matrices. In particular, the spectrally-normalized margin analysis in [6] bounds the misclassification error in terms of a margin-based training loss and a complexity term. The complexity term is proportional to $Q/(\gamma n)$ where γ and n are margin and sample size and Q bounds the spectral complexity. For an L -layer ReLU MLP and maximal width d , Q is bounded by

$$Q = \left(\prod_{i=1}^L \|\mathbf{W}_i\| \right) \left(\sum_{i=1}^L \frac{(\sqrt{d} \|\mathbf{W}_i\|_F)^{2/3}}{\|\mathbf{W}_i\|^{2/3}} \right)^{3/2} \quad (7)$$

where $\|\cdot\|$ and $\|\cdot\|_F$ denote spectral and Frobenius norm respectively. Those norms remain invariant when training the network with POET and at initialization they can be bounded with high probability using standard results from random matrix theory (Appendix C). The scale at initialization is typically chosen such that $\mathbf{W} \in \mathbb{R}^{d \times d}$ satisfies $\|\mathbf{W}\| = O(1)$ and $\|\mathbf{W}\|_F = O(\sqrt{d})$ so that $Q = O_L(d)$.

Approximation properties of SPO. We have seen in Theorem 1 that the factorization $\mathbf{R}\mathbf{W}\mathbf{P}$ with orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$ is the most general spectrum preserving transformation of $\mathbf{W}$. Here we express $\mathbf{R}$ and $\mathbf{P}$ as products of stochastic primitives, but as we state next, this does not reduce representation power when using sufficiently many primitives.

Lemma 1. *If $c \geq \alpha m \ln(m)(m/b)^2$ for some $\alpha > 0$ then with probability at least $1 - m^{-(\alpha-2)}$ over the randomness of the index sets $\mathbf{S}^i$ we can express any orthogonal matrix $\mathbf{R}$ as a product of c primitives $\mathbf{G}_i$ as in Eq. (4). Moreover, the orthogonal matrix $\mathbf{G}_i$ depends only on the sets $\mathbf{S}^j$ and matrices $\mathbf{G}^j$ selected in earlier steps.*

The proof of this lemma can be found in Appendix D. The result extends to Block-stochastic SPO as this is strictly more expressive than fully stochastic SPO. The key idea of the proof is similar to the factorization of orthogonal matrices into a product of Givens rotations. Indeed, by multiplying $\mathbf{R}^\top$ with properly chosen primitive matrices $\mathbf{G}_i$ we can create zeros below the diagonal for one column after another. Note that each $\mathbf{G}_i$ has $b(b-1)/2$ parameters while $\mathbf{R}$ has $m(m-1)/2$ parameters, which implies that generally at least $\Omega((m/b)^2)$ primitives are necessary. In Appendix D we also provide a heuristic that with high probability for $c = O(\ln(m)(m/b)^2)$ every orthogonal matrix can be written as a product of c orthogonal primitives $\mathbf{G}_i$.

Inductive bias. POET-reparameterized neurons result in neural networks that maintain identical architecture and parameter count during inference as conventionally trained networks. While standard training could technically learn equivalent parameters, they consistently fail to do so in practice. This indicates POET provides a unique inductive bias unavailable through standard training. POET also aligns with prior findings in [2, 18] that optimizing factorized matrices yields implicit inductive bias.

Dynamics of singular spectrum. Inspired by [46], we conduct a spectral analysis by comparing the singular spectrum among AdamW, Muon and POET. Following [1, 68], we compute SVD

entropy of the trained Llama-60M model at different iteration. Specifically, the SVD entropy, defined as $H(\sigma) = \frac{1}{\log n} \sum_i \frac{\sigma_i^2}{\sum_j \sigma_j^2} \log \frac{\sigma_i^2}{\sum_j \sigma_j^2}$ measures the diversity of singular values; higher entropy indicates a more uniform and diverse spectrum. [46] attributes the favorable performance of Muon over AdamW to the more diverse spectrum of weight matrices updates. As shown in Figure 6, POET consistently maintains high spectral diversity throughout training, owing to its orthogonal equivalence transformation. Therefore, POET can better explore diverse optimization directions.

POET minimizes hyperspherical energy. While POET generalizes energy-preserving training to spectrum-preserving training, it still ensures low hyperspherical energy when initialized with a zero-mean isotropic Gaussian distribution. This is significant, as POET retains the generalization benefits of minimal hyperspherical energy [47–49]. This property comes from the invariance of zero-mean isotropic Gaussian to orthogonal transformation. [48] shows that for a weight matrix W initialized by zero-mean isotropic Gaussian, its neurons, after being normalized, are uniformly distributed on the unit hypersphere. This property provably leads to a small hyperspherical energy. Since zero-mean isotropic Gaussian is invariant to orthogonal transformation, $OET(W; R, P) = RWP$ does not change the distribution of W , i.e., $RWP =_d W$. We give a derivation in Appendix B. Therefore, the transformed neurons of the new weight matrix RWP are also uniformly distributed over the unit hypersphere after being normalized. This validates that POET provably preserves a small energy.

The property that POET preserves singular spectrum while retaining small hyperspherical energy well justifies why POET yields stable training and good generalization. Such a property only holds true when the weight initialization follows zero-mean isotropic Gaussian, which is perfectly satisfied by current weight initialization. It may also justify why zero-mean isotropic Gaussian initialization works better than the uniform spectrum initialization in our experiments. To further validate that POET maintains a small hyperspherical energy, we plot the hyperspherical energy of the Llama-60M model at different iterations in Figure 7. We compare the sum of the hyperspherical energy of all the layers trained by AdamW, Muon or POET. The results again verify POET’s energy-minimizing property. Different from previous orthogonal training [48], POET can not preserve exactly the same hyperspherical energy during training, but it can well minimize hyperspherical energy in an expectation sense. We also find that Muon can better minimize hyperspherical energy than AdamW.

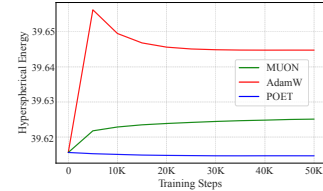


Figure 7: Hyperspherical energy comparison between AdamW, Muon and POET.

5 Experiments and Results

We start by evaluating POET on large-scale LLaMA pretraining, followed by an extensive ablation study to justify our design choices. Detailed settings and additional results are given in Appendices.

5.1 LLM Pretraining using LLaMA Transformers

We perform the pretraining experiments on the Llama transformers of varying sizes (60M, 130M, 350M, 1.3B) for POET. We use the C4 dataset [66], a cleaned web crawl corpus from Common Crawl, widely used for LLM pretraining [29, 56, 82]. For POET-BS, b is the block size of the block-diagonal orthogonal matrix. For POET-FS, $b_{in}=bm$ for R and $b_{out}=bn$

for P . We compare POET against GaLore [82], a low-rank pretraining method, and AdamW, the standard pretraining optimizer. We generally follow the settings in [82]. To better simulate the practical pretraining setting, we significantly increase the number of training tokens for all methods.

Table 2 shows that both POET-FS ($b=1/2$) and POET-BS ($b=256$) consistently outperform both GaLore and AdamW with significantly fewer parameters. For LLaMA-1B, POET-FS ($b=1/2$) yields the best overall performance, achieving a validation perplexity of 13.70, much better than AdamW (14.73) and GaLore (18.33). Block-stochastic POET with $b=256$ achieves the second-best performance

Model (# tokens)	60M (30B)	130M (40B)	350M (40B)	1.3B (50B)
AdamW	26.68 (25.30M)	20.82 (84.93M)	16.78 (302.38M)	14.73 (1.21B)
Galore	29.81 (25.30M)	22.35 (84.93M)	17.99 (302.38M)	18.33 (1.21B)
LoRA _{r=64}	39.70 (4.85M)	32.07 (11.21M)	25.19 (30.28M)	20.55 (59.38M)
POET _{BS, b=64}	29.52 (2.39M)	24.52 (5.52M)	20.29 (14.90M)	18.28 (29.22M)
POET _{BS, b=128}	26.90 (4.81M)	21.86 (11.12M)	18.05 (30.04M)	16.24 (58.91M)
POET _{BS, b=256}	25.29 (9.66M)	19.88 (22.33M)	16.27 (60.32M)	14.56 (118.26M)
POET _{FS, b=1/8}	34.06 (0.53M)	29.67 (1.78M)	24.61 (6.34M)	18.46 (25.39M)
POET _{FS, b=1/4}	28.69 (2.13M)	23.55 (7.13M)	19.42 (25.44M)	17.60 (101.66M)
POET _{FS, b=1/2}	25.37 (8.54M)	19.94 (28.56M)	15.95 (101.86M)	13.70 (406.88M)

Table 2: Comparison of POET with popular pretraining methods using different sizes of LLaMA models. Validation perplexity and the number of trainable parameters are reported.

(14.56), which still surpasses AdamW with only one-tenth of AdamW’s trainable parameters. Similar patterns can be observed for models of smaller sizes. Moreover, we compare the training dynamics between AdamW and POET in Figure 8. The training dynamics of POET is quite different from AdamW. After an initial rapid drop in perplexity, POET improves more slowly than AdamW. As seen in Phase II (Figure 2), this slower but stable progress can lead to better performance in later stages. We attribute this intriguing phenomenon to the unique reparameterization of POET and how we efficiently approximate orthogonality. The exact mechanism behind this phenomenon remains an open question, and understanding it could offer valuable insights into large-scale model training.

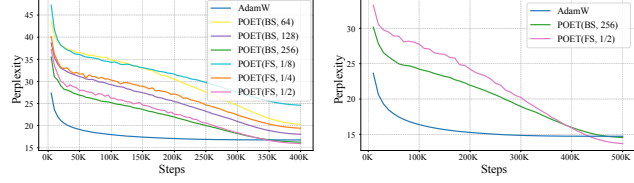


Figure 8: Validation perplexity dynamics on LLaMA-350M and LLaMA-1.3B.

To highlight POET’s non-trivial performance improvement, we increase the training steps (*i.e.*, effectively tokens seen) for AdamW, and find that POET-FS ($b=1/2$) still outperforms AdamW even if AdamW is trained with almost triple the number of tokens. Results are given in Figure 9. In this experiment, the AdamW learning rate was carefully tuned for the full training run, and no training tokens were repeated. Thus, the improvement is non-trivial and cannot be attributed to merely increasing training steps. Interestingly, we also observe from Table 2 that POET’s performance appears strongly correlated with the parameter budget and larger budgets consistently yield better results across model scales. This is particularly important for model scaling law [35]. Another notable observation is that POET significantly outperforms LoRA [26] given a similar parameter budget. For instance, with approximately 30M trainable parameters, POET attains a validation perplexity of 18.05, significantly better than LoRA’s 25.19. We further observe that the block-stochastic variant is more parameter-efficient than the fully stochastic one. On the 130M model, it achieves a validation perplexity of 19.88 with nearly 6M fewer trainable parameters, compared to 19.94 for the fully stochastic variant. We hypothesize that this is due to better coverage of weight parameters. Specifically, the block-stochastic variant ensures all corresponding weights are updated at each step, unlike the more uneven updates in the fully stochastic variant. Experimental details and results on weight update coverage are provided in Appendix H.

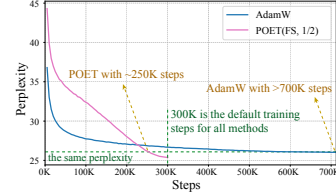


Figure 9: Validation perplexity dynamics of POET (FS, $b=1/2$) and AdamW on Llama-60M. POET outperforms the AdamW trained with almost twice the number of seen tokens.

Pretraining Llama-3B. To demonstrate the scalability of POET, we apply POET to pretrain Llama with 3B parameters while reusing the same hyperparameters from the 1.3B model. This model was trained with only 1/10 of the tokens used for the 1.3B model. Compared to Table 2, the slightly higher final perplexity is attributed to the reduced training data (5B tokens). Table 3 shows that POET maintains the performance advantage over AdamW at the 3B scale, consistent with the conclusion drawn from Table 2 for smaller models. Our findings confirm that the benefits of POET are not limited to smaller models but can extend robustly to larger scales.

Method	Perplexity
AdamW	19.61
POET _{FS, b=1/2}	16.90

Table 3: Validation perplexity of training a LLaMA 3B model on 5 billion tokens.

5.2 LLM Finetuning on Downstream Tasks

To better evaluate models beyond the validation perplexity, we show the results of finetuning the trained model on the GLUE benchmark [75]. This benchmark provides a comprehensive assessment of a model’s

language understanding capabilities through a diverse set of downstream tasks. The performance on GLUE serves as an important indicator of the transferability of the model’s learned representations. Performance was measured using accuracy for MNLI, MRPC, QNLI, QQP, RTE, and SST-2; the Matthews Correlation Coefficient for CoLA; and the Pearson Correlation Coefficient for STS-B. For each task, we finetune the models for 10 epochs using a consistent learning rate of $2e-5$. The evaluation covers several finetuning strategies: full finetuning (Full FT) with AdamW, orthogonal finetuning (OFT) [51, 65], and finetuning with POET. The results in Table 4 show that the POET-pretrained

FT	Model	CoLA	MNLI	MRPC	QNLI	QQP	RTE	SST-2	STS-B
Full FT	AdamW	0.361	0.658	0.696	0.818	0.829	0.534	0.914	0.880
	POET	0.523	0.818	0.824	0.885	0.902	0.661	0.920	0.873
OFT	AdamW	0.388	0.774	0.689	0.842	0.867	0.531	0.915	0.762
	POET	0.437	0.812	0.740	0.855	0.877	0.538	0.924	0.791
POET	AdamW	0.435	0.804	0.806	0.856	0.889	0.653	0.904	0.878
	POET	0.505	0.821	0.826	0.892	0.902	0.682	0.931	0.887

Table 4: Downstream performance on GLUE between AdamW- and POET-pretrained models.

model consistently outperforms the AdamW-pretrained baseline across all tasks and finetuning methods. The results further validate POET’s generalizability. In this setting, OFT uses significantly less parameters than both Full FT and POET, so OFT generally performs worse. However, the comparison between Full FT and POET is fair, since both methods update the entire model.

5.3 Ablation Studies and Empirical Analyses

Initialization schemes. We empirically compare different random initialization schemes for POET, including two commonly used ones (standard Gaussian, Xavier [17]) and two proposed ones (uniform spectrum, normalized Gaussian). Specifically, we use fully stochastic POET with $b=1/2$ to train Llama-60M on 30B tokens and report the validation perplexity in Table 5. Results show that the normalized initialization will lead to the best final performance, and we stick to it as a default choice. Interestingly, uniform spectrum initialization performs poorly. This suggests a trade-off between preserving good weight spectral properties and achieving strong expressiveness. It may limit its expressiveness. Finding the optimal singular value structure for weights remains an important open problem.

Merge-then-reinitialize frequency. The proposed merge-then-reinitialize trick allows POET to train only a small fraction of the large orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$ per iteration, significantly reducing GPU memory usage. However, this trick also introduces a reinitialization frequency hyperparameter T_m , which determines how often the orthogonal matrix is merged and reset to the identity. The index set in POET-FS and the permutation matrix in POET-BS are also resampled at each reinitialization. Therefore, it is quite important to understand how this hyperparameter T_m affects performance. Following the previous initialization experiment, we use POET-FS with $b=1/2$ to train Llama-60M on 30B tokens. We vary the reinitialization frequency from 5 to 1600 and report the validation perplexity in Table 6. Results show that both 200 and 400 perform well. Therefore, we set $T_m = 400$ in all experiments by default.

Neumann series approximation. CNP approximates the matrix inverse using a Neumann series. As the number of Neumann terms directly influences the approximation quality, understanding its impact on model performance is essential. To this end, we evaluate how varying the number of Neumann terms affects performance, using POET-FS with $b = 1/2$ to train LLaMA-130M. Results in Table 7 show that increasing the number of Neumann terms generally improves validation perplexity. However, this also leads to slower training. Moreover, dropping all Neumann terms ($k = 0$) leads to training divergence, highlighting the critical role of maintaining orthogonality. To balance overhead and performance, we find that 3 Neumann terms ($k = 3$) reach a good trade-off.

Additionally, we evaluate the accuracy of the Neumann approximation to understand how the number of Neumann terms affects the orthogonality. The orthogonal approximation error is defined by $e_{\text{orth}} = \|\mathbf{R}\mathbf{R}^T - \mathbf{I}\|_F / \|\mathbf{I}\|_F$. We randomly select a weight matrix and compute the approximation error of corresponding orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$. For clarity, we visualize the error in the initial 1000 training steps in Figure 10. We observe that, with more Neumann terms, the orthogonal approximation error is indeed lower. Moreover, the merge-then-reinitialize trick can periodically reset the error. More results are given in Appendix L.

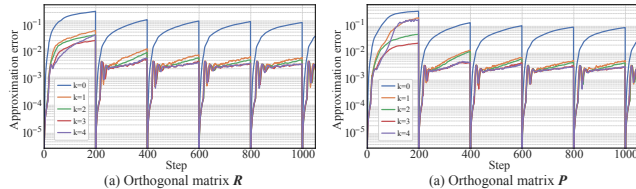


Figure 10: Approximation error of orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$ of a weight matrix.

6 Related Work, Concluding Remarks and Acknowledgement

Related work. Inspired by low-rank adaptation methods such as LoRA [26], a number of recent approaches [11, 20, 27, 30–32, 43–45, 53, 58, 71, 81, 82] have explored low-rank structures to enable efficient pretraining of large language models (LLMs). In parallel, sparsity has also been extensively studied as a means to improve training efficiency in neural networks [9, 14, 15, 25, 72, 78]. Compared to approaches that exploit low-rank structures, relatively few works have explored sparsity for pretraining. Our work broadly aligns with the sparse training paradigm, as POET leverages

Scheme	Perplexity
Standard	26.22
Xavier	25.79
Uni. spectrum	27.29
Normalized	25.37

Table 5: Performance of different initializations.

T_m	Perplexity
5	30.29
25	27.27
50	25.99
200	25.37
400	25.31
1600	25.58

Table 6: Val. perplexity of different T_m .

Scheme	Perplexity
$k = 0$	Not converged
$k = 1$	22.56
$k = 2$	21.54
$k = 3$	20.22
$k = 4$	20.19

Table 7: Number of terms in Neumann series.

sparsely optimized orthogonal matrices to enhance training efficiency. A parallel line of research [34, 46, 60, 69, 80] focuses on developing efficient optimizers for large-scale neural networks. While our work also targets efficient training of large models, it is orthogonal to these efforts, as POET can be integrated with any optimizer. The way POET uses orthogonal matrices to transform neurons may also relate to preconditioned optimizers such as Muon [34], Shampoo [19] and SOAP [73], as well as to the broader field of manifold optimization (e.g., [7]). POET-trained weight matrices remain statistically indistinguishable from randomly initialized ones due to the isotropy of zero-mean independent Gaussian distributions. This yields interesting connections to random neural networks [21, 40, 40, 62, 74], random geometry [22], and random matrix theory [16].

Concluding remarks. This paper introduces POET, a reparameterized training algorithm for large language models. POET models each neuron as the product of two orthogonal matrices and a fixed random weight matrix. By efficiently learning large orthogonal transformations, POET achieves superior generalization while being much more parameter-efficient than existing LLM pretraining methods. Experiments show that POET is broadly applicable to both pretraining and finetuning tasks.

Acknowledgement. The authors would like to sincerely thank Lixin Liu, Han Shi, Gege Gao, Zhen Liu and many colleagues at Max Planck Institute for Intelligent Systems for many helpful suggestions. Additionally, the authors also sincerely thank all the anonymous NeurIPS reviewers for their constructive suggestions that greatly improved the quality of our work.

References

- [1] Orly Alter, Patrick O Brown, and David Botstein. Singular value decomposition for genome-wide expression data processing and modeling. *Proceedings of the National Academy of Sciences*, 2000. 7
- [2] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. In *NeurIPS*, 2019. 7
- [3] Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*, 2021. 1
- [4] Nitin Bansal, Xiaohan Chen, and Zhangyang Wang. Can we gain more from orthogonality regularizations in training deep networks? In *NeurIPS*, 2018. 1
- [5] Line Baribeau and Thomas Ransford. Non-linear spectrum-preserving maps. *Bulletin of the London Mathematical Society*, 32(1):8–14, 2000. 3
- [6] Peter L Bartlett, Dylan J Foster, and Matus J Telgarsky. Spectrally-normalized margin bounds for neural networks. In *NeurIPS*, 2017. 1, 7
- [7] Silvere Bonnabel. Stochastic gradient descent on riemannian manifolds. *IEEE Transactions on Automatic Control*, 58(9):2217–2229, 2013. 11
- [8] Djalil Chafai, Djalil Chafä, Olivier Guédon, Guillaume Lecue, and Alain Pajor. Singular values of random matrices. *Lecture Notes*, 13, 2009. 20
- [9] Beidi Chen, Tri Dao, Eric Winsor, Zhao Song, Atri Rudra, and Christopher Ré. Scatterbrain: Unifying sparse and low-rank attention. In *NeurIPS*, 2021. 10
- [10] Tianlong Chen, Zhenyu Zhang, Yu Cheng, Ahmed Awadallah, and Zhangyang Wang. The principle of diversity: Training stronger vision transformers calls for reducing all levels of redundancy. In *CVPR*, 2022. 1
- [11] Xi Chen, Kaituo Feng, Changsheng Li, Xunhao Lai, Xiangyu Yue, Ye Yuan, and Guoren Wang. Fira: Can we achieve full-rank training of llms under low-rank constraint? *arXiv preprint arXiv:2410.01623*, 2024. 10
- [12] Moustapha Cisse, Piotr Bojanowski, Edouard Grave, Yann Dauphin, and Nicolas Usunier. Parseval networks: Improving robustness to adversarial examples. In *ICML*, 2017. 1
- [13] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021. 1

- [14] Tri Dao, Beidi Chen, Nimit S Sohoni, Arjun Desai, Michael Poli, Jessica Grogan, Alexander Liu, Aniruddh Rao, Atri Rudra, and Christopher Ré. Monarch: Expressive structured matrices for efficient and accurate training. In *ICML*, 2022. 10
- [15] Tri Dao, Albert Gu, Matthew Eichhorn, Atri Rudra, and Christopher Ré. Learning fast algorithms for linear transforms using butterfly factorizations. In *ICML*, 2019. 10
- [16] Alan Edelman and N Raj Rao. Random matrix theory. *Acta numerica*, 14:233–297, 2005. 11
- [17] Xavier Glorot and Yoshua Bengio. Understanding the difficulty of training deep feedforward neural networks. In *AISTATS*, 2010. 4, 10, 19
- [18] Suriya Gunasekar, Blake E Woodworth, Srinadh Bhojanapalli, Behnam Neyshabur, and Nati Srebro. Implicit regularization in matrix factorization. In *NeurIPS*, 2017. 7
- [19] Vineet Gupta, Tomer Koren, and Yoram Singer. Shampoo: Preconditioned stochastic tensor optimization. In *ICML*, 2018. 11
- [20] Andi Han, Jiaxiang Li, Wei Huang, Mingyi Hong, Akiko Takeda, Pratik Jawanpuria, and Bamdev Mishra. Sltrain: a sparse plus low-rank approach for parameter and memory efficient pretraining. *arXiv preprint arXiv:2406.02214*, 2024. 10
- [21] Boris Hanin. Random neural networks in the infinite width limit as gaussian processes. *The Annals of Applied Probability*, 33(6A):4798–4819, 2023. 11
- [22] Boris Hanin and David Rolnick. Complexity of linear regions in deep networks. In *ICML*, 2019. 11
- [23] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *ICCV*, 2015. 19
- [24] Karl Moritz Hermann, Tomáš Kočiský, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. Teaching machines to read and comprehend, 2015. 40
- [25] Torsten Hoefer, Dan Alistarh, Tal Ben-Nun, Nikoli Dryden, and Alexandra Peste. Sparsity in deep learning: Pruning and growth for efficient inference and training in neural networks. *JMLR*, 2021. 10
- [26] Edward J. Hu, yelong shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. LoRA: Low-rank adaptation of large language models. In *ICLR*, 2022. 9, 10
- [27] Jia-Hong Huang, Yixian Shen, Hongyi Zhu, Stevan Rudinac, and Evangelos Kanoulas. Gradient weight-normalized low-rank projection for efficient llm training. In *AAAI*, 2025. 10
- [28] Lei Huang, Xianglong Liu, Bo Lang, Adams Yu, Yongliang Wang, and Bo Li. Orthogonal weight normalization: Solution to optimization over multiple dependent stiefel manifolds in deep neural networks. In *AAAI*, 2018. 1
- [29] Tianjin Huang, Ziquan Zhu, Gaojie Jin, Lu Liu, Zhangyang Wang, and Shiwei Liu. Spam: Spike-aware adam with momentum reset for stable llm training, 2025. 8
- [30] Weihao Huang, Zhenyu Zhang, Yushun Zhang, Zhi-Quan Luo, Ruoyu Sun, and Zhangyang Wang. Galore-mini: Low rank gradient learning with fewer learning rates. In *NeurIPS Workshop on Fine-Tuning in Modern Machine Learning: Principles and Scalability*, 2024. 10
- [31] Minyoung Huh, Brian Cheung, Jeremy Bernstein, Phillip Isola, and Pulkit Agrawal. Training neural networks from scratch with parallel low-rank adapters. *arXiv preprint arXiv:2402.16828*, 2024.
- [32] Ajay Jaiswal, Lu Yin, Zhenyu Zhang, Shiwei Liu, Jiawei Zhao, Yuandong Tian, and Zhangyang Wang. From galore to welore: How low-rank weights non-uniformly emerge from low-rank gradients. *arXiv preprint arXiv:2407.11239*, 2024. 10
- [33] Haoming Jiang, Zhehui Chen, Minshuo Chen, Feng Liu, Dingding Wang, and Tuo Zhao. On computation and generalization of generative adversarial networks under spectrum control. In *ICLR*, 2019. 1
- [34] Keller Jordan, Yuchen Jin, Vlado Boza, Jiacheng You, Franz Cesista, Laker Newhouse, and Jeremy Bernstein. Muon: An optimizer for hidden layers in neural networks, 2024. 11
- [35] Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020. 9

- [36] Nitish Shirish Keskar and Richard Socher. Improving generalization performance by switching from adam to sgd. *arXiv preprint arXiv:1712.07628*, 2017. 1
- [37] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *ICLR*, 2015. 1
- [38] Günter Klambauer, Thomas Unterthiner, Andreas Mayr, and Sepp Hochreiter. Self-normalizing neural networks. In *NeurIPS*, 2017. 1
- [39] Hojoon Lee, Youngdo Lee, Takuma Seno, Donghu Kim, Peter Stone, and Jaegul Choo. Hyperspherical normalization for scalable deep reinforcement learning. *arXiv preprint arXiv:2502.15280*, 2025. 1
- [40] Jaehoon Lee, Yasaman Bahri, Roman Novak, Samuel S Schoenholz, Jeffrey Pennington, and Jascha Sohl-Dickstein. Deep neural networks as gaussian processes. *arXiv preprint arXiv:1711.00165*, 2017. 11
- [41] Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Ves Stoyanov, and Luke Zettlemoyer. Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. *arXiv preprint arXiv:1910.13461*, 2019. 40
- [42] Chi-Kwong Li and Nam-Kiu Tsing. Linear operators preserving unitarily invariant norms of matrices. *Linear and Multilinear Algebra*, 26(1-2):119–132, 1990. 3
- [43] Vladislav Lialin, Namrata Shivagunde, Sherin Muckatira, and Anna Rumshisky. Relora: High-rank training through low-rank updates. *arXiv preprint arXiv:2307.05695*, 2023. 10
- [44] Kaizhao Liang, Bo Liu, Lizhang Chen, and Qiang Liu. Memory-efficient llm training with online subspace descent. *arXiv preprint arXiv:2408.12857*, 2024.
- [45] Xutao Liao, Shaohui Li, Yuhui Xu, Zhi Li, Yu Liu, and You He. Galore+: Boosting low-rank adaptation for llms with cross-head projection. *arXiv preprint arXiv:2412.19820*, 2024. 10
- [46] Jingyuan Liu, Jianlin Su, Xingcheng Yao, Zhejun Jiang, Guokun Lai, Yulun Du, Yidao Qin, Weixin Xu, Enzhe Lu, Junjie Yan, et al. Muon is scalable for llm training. *arXiv preprint arXiv:2502.16982*, 2025. 7, 8, 11
- [47] Weiyang Liu, Rongmei Lin, Zhen Liu, Lixin Liu, Zhiding Yu, Bo Dai, and Le Song. Learning towards minimum hyperspherical energy. In *NeurIPS*, 2018. 1, 2, 3, 8
- [48] Weiyang Liu, Rongmei Lin, Zhen Liu, James M Rehg, Liam Paull, Li Xiong, Le Song, and Adrian Weller. Orthogonal over-parameterized training. In *CVPR*, 2021. 2, 3, 4, 5, 6, 8, 19
- [49] Weiyang Liu, Rongmei Lin, Zhen Liu, Li Xiong, Bernhard Schölkopf, and Adrian Weller. Learning with hyperspherical uniformity. In *AISTATS*, 2021. 1, 2, 3, 8, 20
- [50] Weiyang Liu, Zhen Liu, Zhiding Yu, Bo Dai, Rongmei Lin, Yisen Wang, James M Rehg, and Le Song. Decoupled networks. In *CVPR*, 2018. 1, 4
- [51] Weiyang Liu, Zeju Qiu, Yao Feng, Yuliang Xiu, Yuxuan Xue, Longhui Yu, Haiwen Feng, Zhen Liu, Juyeon Heo, Songyou Peng, Yandong Wen, Michael J. Black, Adrian Weller, and Bernhard Schölkopf. Parameter-efficient orthogonal finetuning via butterfly factorization. In *ICLR*, 2024. 2, 3, 5, 6, 9
- [52] Weiyang Liu, Yan-Ming Zhang, Xingguo Li, Zhiding Yu, Bo Dai, Tuo Zhao, and Le Song. Deep hyperspherical learning. In *NIPS*, 2017. 1, 4
- [53] Ziyue Liu, Ruijie Zhang, Zhengyang Wang, Zi Yang, Paul Hovland, Bogdan Nicolae, Franck Cappello, and Zheng Zhang. Cola: Compute-efficient pre-training of llms via low-rank activation. *arXiv preprint arXiv:2502.10940*, 2025. 10
- [54] Ilya Loshchilov, Cheng-Ping Hsieh, Simeng Sun, and Boris Ginsburg. ngpt: Normalized transformer with representation learning on the hypersphere. *arXiv preprint arXiv:2410.01131*, 2024. 1
- [55] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *ICLR*, 2019. 1, 25
- [56] Chao Ma, Wenbo Gong, Meyer Scetbon, and Edward Meeds. Swan: Sgd with normalization and whitening enables stateless llm training, 2025. 8
- [57] Albert W Marshall, Ingram Olkin, and Barry C Arnold. *Inequalities: theory of majorization and its applications*, volume 143. Springer, 1979. 21
- [58] Roy Miles, Pradyumna Reddy, Ismail Elezi, and Jiankang Deng. Velora: Memory efficient training using rank-1 sub-token projections. *arXiv preprint arXiv:2405.17991*, 2024. 10

- [59] Takeru Miyato, Toshiki Kataoka, Masanori Koyama, and Yuichi Yoshida. Spectral normalization for generative adversarial networks. *arXiv preprint arXiv:1802.05957*, 2018. 1
- [60] Zhanfeng Mo, Long-Kai Huang, and Sinno Jialin Pan. Parameter and memory efficient pretraining via low-rank riemannian optimization. In *ICLR*, 2025. 11
- [61] Shashi Narayan, Shay B. Cohen, and Mirella Lapata. Don’t give me the details, just the summary! topic-aware convolutional neural networks for extreme summarization, 2018. 40
- [62] Radford M Neal. Priors for infinite networks. *Bayesian learning for neural networks*, pages 29–53, 1996. 11
- [63] Behnam Neyshabur, Srinadh Bhojanapalli, and Nathan Srebro. A pac-bayesian approach to spectrally-normalized margin bounds for neural networks. In *ICLR*, 2018. 7
- [64] Sean O’Rourke, Van Vu, and Ke Wang. Eigenvectors of random matrices: a survey. *Journal of Combinatorial Theory, Series A*, 144:361–442, 2016. 20
- [65] Zeju Qiu, Weiyang Liu, Haiwen Feng, Yuxuan Xue, Yao Feng, Zhen Liu, Dan Zhang, Adrian Weller, and Bernhard Schölkopf. Controlling text-to-image diffusion by orthogonal finetuning. In *NeurIPS*, 2023. 2, 3, 5, 6, 9
- [66] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of machine learning research*, 21(140):1–67, 2020. 8, 25
- [67] Mihaela Rosca, Theophane Weber, Arthur Gretton, and Shakir Mohamed. A case for new neural network smoothness constraints. *arXiv preprint arXiv:2012.07969*, 2020. 1
- [68] Olivier Roy and Martin Vetterli. The effective rank: A measure of effective dimensionality. In *European signal processing conference*, 2007. 7
- [69] Ishaan Shah, Anthony M Polloreno, Karl Stratos, Philip Monk, Adarsh Chaluvvaraju, Andrew Hojel, Andrew Ma, Anil Thomas, Ashish Tanwer, Darsh J Shah, et al. Practical efficiency of muon for pretraining. *arXiv preprint arXiv:2505.02222*, 2025. 11
- [70] Jack W Silverstein et al. The smallest eigenvalue of a large dimensional wishart matrix. *The Annals of Probability*, 13(4):1364–1368, 1985. 20
- [71] DiJia Su, Andrew Gu, Jane Xu, Yuandong Tian, and Jiawei Zhao. Galore 2: Large-scale llm pre-training by gradient low-rank projection. *arXiv preprint arXiv:2504.20437*, 2025. 10
- [72] Vithursan Thangarasa, Abhay Gupta, William Marshall, Tianda Li, Kevin Leong, Dennis DeCoste, Sean Lie, and Shreyas Saxena. Spdf: Sparse pre-training and dense fine-tuning for large language models. In *UAI*, 2023. 10
- [73] Nikhil Vyas, Depen Morwani, Rosie Zhao, Mujin Kwun, Itai Shapira, David Brandfonbrener, Lucas Janson, and Sham Kakade. Soap: Improving and stabilizing shampoo using adam. *arXiv preprint arXiv:2409.11321*, 2024. 11
- [74] Gilles Wainrib and Jonathan Touboul. Topological and dynamical complexity of random neural networks. *Physical review letters*, 110(11):118101, 2013. 11
- [75] Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel R Bowman. Glue: A multi-task benchmark and analysis platform for natural language understanding. *arXiv preprint arXiv:1804.07461*, 2018. 9
- [76] Bo Xie, Yingyu Liang, and Le Song. Diverse neural network learns true target functions. In *AISTATS*, 2017. 3, 7
- [77] Sherry Yang, Ofir Nachum, Yilun Du, Jason Wei, Pieter Abbeel, and Dale Schuurmans. Foundation models for decision making: Problems, methods, and opportunities. *arXiv preprint arXiv:2303.04129*, 2023. 1
- [78] Xinyu Yang, Jixuan Leng, Geyang Guo, Jiawei Zhao, Ryumei Nakada, Linjun Zhang, Huaxiu Yao, and Beidi Chen. S2ft: Efficient, scalable and generalizable llm fine-tuning by structured sparsity. *arXiv preprint arXiv:2412.06289*, 2024. 10
- [79] Yuichi Yoshida and Takeru Miyato. Spectral norm regularization for improving the generalizability of deep learning. *arXiv preprint arXiv:1705.10941*, 2017. 1

- [80] Yushun Zhang, Congliang Chen, Ziniu Li, Tian Ding, Chenwei Wu, Diederik P Kingma, Yinyu Ye, Zhi-Quan Luo, and Ruoyu Sun. Adam-mini: Use fewer learning rates to gain more. *arXiv preprint arXiv:2406.16793*, 2024. 11
- [81] Zhenyu Zhang, Ajay Jaiswal, Lu Yin, Shiwei Liu, Jiawei Zhao, Yuandong Tian, and Zhangyang Wang. Q-galore: Quantized galore with int4 projection and layer-adaptive low-rank gradients. *arXiv preprint arXiv:2407.08296*, 2024. 10
- [82] Jiawei Zhao, Zhenyu Zhang, Beidi Chen, Zhangyang Wang, Anima Anandkumar, and Yuandong Tian. Galore: Memory-efficient llm training by gradient low-rank projection. In *ICML*, 2024. 6, 8, 10, 25

Appendix

Table of Contents

A	Delving into POET's Three Training Phases	17
A.1	More Details on Vector Probing	17
A.2	Geometric Interpretation of the Trace of Orthogonal Matrices	17
A.3	Empirical Observations	18
B	Minimum Hyperspherical Energy in POET	19
C	Guarantees of Weight Spectrum under POET	20
D	Proofs of Lemma 1	23
E	Experimental Details	25
F	Implementation and CUDA Acceleration	27
G	Results of Vector Probing for R and P	28
H	Weight Update Evenness of Different POET Variants	30
I	Training Dynamics of Singular Values	31
J	Orthogonality Approximation Quality using Neumann Series	37
K	Full Results of Training Dynamics	39
L	More Results of POET as a Finetuning Method	40

A Delving into POET’s Three Training Phases

A.1 More Details on Vector Probing

The three training phases of POET are summarized from the empirical observation of the vector probing results. The idea of vector probing is very straightforward. We generate a constant vector $\mathbf{v}$ that is randomly initialized. Then we let it to be transformed by the learned orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$. Finally, we compute the cosine of their angle: $\mathbf{v}^\top \mathbf{R} \mathbf{v}$ and $\mathbf{v}^\top \mathbf{P} \mathbf{v}$. In this process, the probing vector $\mathbf{v}$ is always fixed. The full results are given in Appendix G.

Beyond a particular constant probing vector, we also consider a set of randomly sampled probing vectors that follow our proposed normalized Gaussian initialization. Specifically, we consider the following expectation:

$$\mathbb{E}_{\mathbf{v} \sim \mathbb{S}^{m-1}} \{\mathbf{v}^\top \mathbf{R} \mathbf{v}\}, \quad (8)$$

where $\mathbf{v}$ is a vector initialized by normalized Gaussian distribution (thus uniformly distributed on a unit hypersphere $\mathbb{S}^{m-1}$). Because $\mathbb{E}\{\mathbf{v} \mathbf{v}^\top\} = \frac{1}{m} \mathbf{I}$, then we have that

$$\mathbb{E}_{\mathbf{v} \sim \mathbb{S}^{m-1}} \{\mathbf{v}^\top \mathbf{R} \mathbf{v}\} = \frac{1}{m} \text{Tr}(\mathbf{R}). \quad (9)$$

where $\text{Tr}(\cdot)$ denotes the matrix trace. Its geometric interpretation is the cosine of the rotation angle between $\mathbf{v}$ and $\mathbf{R} \mathbf{v}$.

Next, we look into the variance of $q(x) = \mathbf{v}^\top \mathbf{R} \mathbf{v}$ (we simplify the expectation over the unit hypersphere to $\mathbb{E}$):

$$\text{Var}(q(x)) = \mathbb{E}\{(\mathbf{v}^\top \mathbf{R} \mathbf{v})^2\} - (\mathbb{E}\{\mathbf{v}^\top \mathbf{R} \mathbf{v}\})^2. \quad (10)$$

First we compute $\mathbb{E}\{(\mathbf{v}^\top \mathbf{R} \mathbf{v})^2\}$:

$$\begin{aligned} \mathbb{E}\{(\mathbf{v}^\top \mathbf{R} \mathbf{v})^2\} &= \frac{\text{Tr}(\mathbf{R})^2 + 2 \left\| \frac{\mathbf{R}^\top + \mathbf{R}}{2} \right\|}{m(m+2)} \\ &= \frac{\text{Tr}(\mathbf{R})^2 + \text{Tr}(\mathbf{R}^2) + m}{m(m+2)} \end{aligned} \quad (11)$$

Then we compute $(\mathbb{E}\{\mathbf{v}^\top \mathbf{R} \mathbf{v}\})^2$:

$$(\mathbb{E}\{\mathbf{v}^\top \mathbf{R} \mathbf{v}\})^2 = \frac{\text{Tr}(\mathbf{R})^2}{m^2}. \quad (12)$$

Finally, we combine pieces and have the final variance:

$$\text{Var}(\mathbf{v}^\top \mathbf{R} \mathbf{v}) = \frac{m + \text{Tr}(\mathbf{R}^2) + \frac{2\text{Tr}(\mathbf{R})^2}{m}}{m(m+2)} \quad (13)$$

which shrinks at the order of $O(1/m)$. Therefore, when the dimension of orthogonal matrices is large, even if we use a fixed random probing vector $\mathbf{v}$, this rotation angle is quite consistent.

A.2 Geometric Interpretation of the Trace of Orthogonal Matrices

Let’s delve deeper into the trace of orthogonal matrices. It generally represents how much a transformation preserves vectors in their original directions. Specifically, the trace indicates how much “alignment” or similarity there is between the original vectors and their images after transformation.

The trace of an orthogonal matrix $\mathbf{R} \in \mathbb{R}^{m \times m}$ can be written as

$$\text{Tr}(\mathbf{R}) = \sum_{i=1}^m \mathbf{e}_i^\top \mathbf{R} \mathbf{e}_i \quad (14)$$

where $\mathbf{e}_i, \forall i$ are unit basis vectors. This expression reveals that the trace measures the sum of inner products between each original direction $\mathbf{e}_i$ and its transformed version $\mathbf{R} \mathbf{e}_i$. Since $\mathbf{e}_i^\top \mathbf{R} \mathbf{e}_i$ can be interpreted as the cosine of the angle between $\mathbf{e}_i$ and $\mathbf{R} \mathbf{e}_i$, the trace thus reflects how much the orthogonal transformation aligns with or deviates from the original coordinate directions.

We also plot the trace of both R and P during the POET training. The results are shown in Figure 13 and Figure 14. After dividing the trace by the orthogonal matrix dimension, we obtain that the result is generally in the range of $[0.6, 0.65]$ after training. This is similar to the results of vector probing. Therefore, we empirically verify the conclusion that the expectation of vector probing results is $\frac{\text{Tr}(R)}{m}$ with a small variance.

A.3 Empirical Observations

The training dynamics of POET presents three geometry-driven phases. We note that these phase changes are based on empirical observation, and further theoretical understanding of this process remains an open problem.

Phase I: conical-shell searching rotates each orthogonal matrix R and P smoothly away from the identity while preserving their singular values, so the cosine similarity between transformed and initial weight vectors falls from 1 to ≈ 0.6 ; this provides a spectrally well-conditioned “cone” in which learning can proceed safely. this phase serves the role of “spectral warm-up”. By plotting the cosine similarity of any one layer, we always see the same graceful slide towards 0.6–0.65, independent of model size, layer type, or whether you train with fully-stochastic or block-stochastic SPO. This phase carves out the thin “shell” in which subsequent learning lives.

Phase II: stable learning on the conical shell occupies the bulk of training: the angles to the initial vectors stay locked in that narrow band, optimization now shears weights *within* the cone, and validation perplexity drops almost linearly because spectra remain frozen and gradients act only on meaningful directions. In this phase, the trace of the orthogonal matrices stay almost as a constant.

Specifically, we hypothesize that the orthogonal transforms have reached a “good” cone; thereafter they mostly shear vectors inside that shell, leaving the angle to the original vector unchanged. The spectrum continues to be exactly that of the random initial matrix, so gradients can no longer distort singular values and instead devote capacity to learning meaningful directions. Because the geometry is stabilized in this phase, the learning of patterns happen in a stable subspace. This stable learning phase takes up 80% of the training time.

Phase III: final adjusting coincides with learning-rate decay; the orthogonal transforms barely move, making only tiny refinements to singular vectors, so additional steps yield diminishing returns. This phase is merely the LR cooldown; weights and spectra are already near their final configuration, so progress naturally slows.

B Minimum Hyperspherical Energy in POET

We start by showing that orthogonal equivalence transformation in POET can provably obtain small hyperspherical energy for its transformed weight matrices. This result holds when the weight matrices are independently initialized by zero-mean isotropic Gaussian distribution. Both Xavier [17] and Kaiming [23] initializations satisfy such a weight initialization condition. Orthogonal equivalence transformation is given by $\text{OET}(\mathbf{W}; \mathbf{R}, \mathbf{P}) = \mathbf{R}\mathbf{W}\mathbf{P}$, where the input matrix $\mathbf{W}$ is multiplied on the left and on the right by orthogonal matrices $\mathbf{R}$ and $\mathbf{P}$, respectively.

When $\mathbf{W}$ is initialized by zero-mean isotropic Gaussian distribution, [48] has shown that these random neurons, if normalized, are uniformly distributed on the unit hypersphere. This leads to a provably small hyperspherical energy for the randomly initialized weight matrix. In the following, we will show that after orthogonal equivalence transformation, the weight matrix still maintains a small hyperspherical energy.

We consider a weight matrix $\mathbf{W} \in \mathbb{R}^{m \times n}$ where each entry is *i.i.d.* sampled from a zero-mean Gaussian distribution with variance the same variance σ^2 , *i.e.*, $W_{ij} \sim \mathcal{N}(0, \sigma^2)$. After applying orthogonal equivalence transformation, we have $\mathbf{W}^{\text{new}} = \mathbf{R}\mathbf{W}\mathbf{P}$ where $\mathbf{R}$ and $\mathbf{P}$ are two orthogonal matrices. Then we compute the distribution of $\mathbf{W}^{\text{new}}$. Because linear maps preserve Gaussianity, each entry of $\mathbf{W}^{\text{new}}$ is a finite linear combination of W_{ij} , and hence, $\mathbf{W}^{\text{new}}$ follows a joint Gaussian which can be fully characterized by mean and covariance.

The mean of $\mathbf{W}^{\text{new}}$ is given by

$$\mathbb{E}[\mathbf{W}^{\text{new}}] = \mathbf{R} \cdot \mathbb{E}[\mathbf{W}] \cdot \mathbf{P} = \mathbf{R} \cdot \mathbf{0} \cdot \mathbf{P} = \mathbf{0} = \mathbb{E}[\mathbf{W}]. \quad (15)$$

For its covariance, we consider two generic entries W_{ij}^{new} and $W_{i'j'}^{\text{new}}$:

$$\begin{aligned} W_{ij}^{\text{new}} &= \sum_{k,l} R_{ik} W_{kl}^{\text{new}} P_{lj}, \\ W_{i'j'}^{\text{new}} &= \sum_{u,v} R_{i'u} W_{uv}^{\text{new}} P_{vj'}. \end{aligned} \quad (16)$$

Then we compute the covariance between the two entries:

$$\begin{aligned} \text{Cov}(W_{ij}^{\text{new}}, W_{i'j'}^{\text{new}}) &= \sum_{k,l,u,v} R_{ik} R_{i'u} P_{lj} P_{vj'} \text{Cov}(W_{kl}, W_{uv}) \\ &= \sigma^2 (\mathbf{R}\mathbf{R}^\top)_{ii'} (\mathbf{P}^\top \mathbf{P})_{jj'} \\ &= \sigma^2 \delta_{ii'} \delta_{jj'} \end{aligned} \quad (17)$$

which implies the following results:

- The covariance matrix is a diagonal matrix, so different entries of $\mathbf{W}^{\text{new}}$ are uncorrelated.
- Because $\mathbf{W}^{\text{new}}$ is a joint Gaussian and different entries are uncorrelated, each entry of $\mathbf{W}^{\text{new}}$ is independent.
- Each entry of $\mathbf{W}^{\text{new}}$ has identical variance σ^2 .

To sum up, each entry of $\mathbf{W}^{\text{new}}$ is *i.i.d.* $\mathcal{N}(0, \sigma^2)$, which is identical to the distribution of each entry of $\mathbf{W}$. Because we have $\mathbf{W}^{\text{new}} =_d \mathbf{W}$, we can conclude that, similar to $\mathbf{W}$, $\mathbf{W}_{\text{new}}$ also has provably small hyperspherical energy among neurons.

Despite being extremely simple, we find that this is in fact a significant result. Under zero-mean isotropic Gaussian initialization, spectrum-preserving training and energy-preserving training can be achieved simultaneously. It also partially explains why the proposed normalized Gaussian initialization achieves the best performance.

C Guarantees of Weight Spectrum under POET

For standard Gaussian initialization where each element of the weight matrix $\mathbf{W} \in d \times n$ is sampled with a normal distribution, we have the following standard results [8, 70]:

$$\begin{aligned} \frac{1}{\sqrt{d}} \sigma_{\max}(\mathbf{W}) &\xrightarrow[n \rightarrow \infty]{\text{a.s.}} 1 + \sqrt{\lambda} \\ \frac{1}{\sqrt{d}} \sigma_{\min}(\mathbf{W}) &\xrightarrow[n \rightarrow \infty]{\text{a.s.}} 1 - \sqrt{\lambda} \end{aligned} \quad (18)$$

which gives spectrum guarantees for weight matrices generated by the standard Gaussian initialization.

In the following, we give the spectrum guarantees for the normalized Gaussian initialization. We start by stating the following theorem from [49]:

Theorem 2. *Let $\tilde{\mathbf{v}}_1, \dots, \tilde{\mathbf{v}}_n \in \mathbb{R}^d$ be i.i.d. random vectors where each element follows the Gaussian distribution with mean 0 and variance 1. Then $\mathbf{v}_1 = \frac{\tilde{\mathbf{v}}_1}{\|\tilde{\mathbf{v}}_1\|_2}, \dots, \mathbf{v}_n = \frac{\tilde{\mathbf{v}}_n}{\|\tilde{\mathbf{v}}_n\|_2}$ are uniformly distributed on the unit hypersphere $\mathbb{S}^{d-1}$. If the ratio $\frac{n}{d}$ converges to a constant $\lambda \in (0, 1)$, asymptotically we have for $\mathbf{W} = \{\mathbf{v}_1, \dots, \mathbf{v}_n\} \in \mathbb{R}^{d \times n}$:*

$$\begin{aligned} \lim_{n \rightarrow \infty} \sigma_{\max}(\mathbf{W}) &\leq (\sqrt{d} + \sqrt{\lambda d}) \cdot \left(\max_i \frac{1}{\|\tilde{\mathbf{v}}_i\|_2} \right) \\ \lim_{n \rightarrow \infty} \sigma_{\min}(\mathbf{W}) &\geq (\sqrt{d} - \sqrt{\lambda d}) \cdot \left(\min_i \frac{1}{\|\tilde{\mathbf{v}}_i\|_2} \right) \end{aligned} \quad (19)$$

where $\sigma_{\max}(\cdot)$ and $\sigma_{\min}(\cdot)$ denote the largest and the smallest singular value of a matrix, respectively.

Proof. We first introduce the following lemma as the characterization of a unit vector that is uniformly distributed on the unit hypersphere $\mathbb{S}^{d-1}$.

Lemma 2 ([64]). *Let $\mathbf{v}$ be a random vector that is uniformly distributed on the unit hypersphere $\mathbb{S}^{d-1}$. Then $\mathbf{v}$ has the same distribution as the following:*

$$\left\{ \frac{u_1}{\sqrt{\sum_{i=1}^d u_i^2}}, \frac{u_2}{\sqrt{\sum_{i=1}^d u_i^2}}, \dots, \frac{u_d}{\sqrt{\sum_{i=1}^d u_i^2}} \right\} \quad (20)$$

where $u_1, u_2, \dots, u_d$ are i.i.d. standard normal random variables.

Proof. The lemma follows naturally from the fact that the Gaussian vector $\{u_i\}_{i=1}^d$ is rotationally invariant. $\square$

Then we consider a random matrix $\tilde{\mathbf{W}} = \{\tilde{\mathbf{v}}_1, \dots, \tilde{\mathbf{v}}_n\}$ where $\tilde{\mathbf{v}}_i$ follows the same distribution of $\{u_1, \dots, u_d\}$. Therefore, it is also equivalent to a random matrix with each element distributed normally. For such a matrix $\tilde{\mathbf{W}}$, we have from [70] that

$$\begin{aligned} \lim_{n \rightarrow \infty} \sigma_{\max}(\tilde{\mathbf{W}}) &= \sqrt{d} + \sqrt{\lambda d} \\ \lim_{n \rightarrow \infty} \sigma_{\min}(\tilde{\mathbf{W}}) &= \sqrt{d} - \sqrt{\lambda d} \end{aligned} \quad (21)$$

where $\sigma_{\max}(\cdot)$ and $\sigma_{\min}(\cdot)$ denote the largest and the smallest singular value, respectively.

Then we write the matrix $\mathbf{W}$ as follows:

$$\begin{aligned} \mathbf{W} &= \tilde{\mathbf{W}} \cdot \mathbf{Q} \\ &= \tilde{\mathbf{W}} \cdot \begin{bmatrix} \frac{1}{\|\tilde{\mathbf{v}}_1\|_2} & 0 & \dots & 0 \\ 0 & \frac{1}{\|\tilde{\mathbf{v}}_2\|_2} & \ddots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & \frac{1}{\|\tilde{\mathbf{v}}_n\|_2} \end{bmatrix} \end{aligned} \quad (22)$$

which leads to

$$\begin{aligned}\lim_{n \rightarrow \infty} \sigma_{\max}(\mathbf{W}) &= \lim_{n \rightarrow \infty} \sigma_{\max}(\tilde{\mathbf{W}} \cdot \mathbf{Q}) \\ \lim_{n \rightarrow \infty} \sigma_{\min}(\mathbf{W}) &= \lim_{n \rightarrow \infty} \sigma_{\min}(\tilde{\mathbf{W}} \cdot \mathbf{Q})\end{aligned}\quad (23)$$

We first assume that for a symmetric matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ $\lambda_1(\mathbf{A}) \geq \dots \geq \lambda_n(\mathbf{A})$. Then we introduce the following inequalities for eigenvalues:

Lemma 3 ([57]). *Let $\mathbf{G}, \mathbf{H} \in \mathbb{R}^{n \times n}$ be positive semi-definite symmetric, and let $1 \leq i_1 < \dots < i_k \leq n$. Then we have that*

$$\prod_{t=1}^k \lambda_{i_t}(\mathbf{GH}) \leq \prod_{t=1}^k \lambda_{i_t}(\mathbf{G}) \lambda_t(\mathbf{H}) \quad (24)$$

and

$$\prod_{t=1}^k \lambda_{i_t}(\mathbf{GH}) \geq \prod_{t=1}^k \lambda_{i_t}(\mathbf{G}) \lambda_{n-t+1}(\mathbf{H}) \quad (25)$$

where λ_i denotes the i -th largest eigenvalue.

We first let $1 \leq i_1 < \dots < i_k \leq n$. Because $\tilde{\mathbf{W}} \in \mathbb{R}^{d \times n}$ and $\mathbf{Q} \in \mathbb{R}^{n \times n}$, we have the following:

$$\begin{aligned}\prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}\mathbf{Q}) &= \prod_{t=1}^k \sqrt{\lambda_{i_t}(\tilde{\mathbf{W}}\mathbf{Q}\mathbf{Q}^\top \tilde{\mathbf{W}}^\top)} \\ &= \sqrt{\prod_{t=1}^k \lambda_{i_t}(\tilde{\mathbf{W}}^\top \tilde{\mathbf{W}}\mathbf{Q}\mathbf{Q}^\top)}\end{aligned}\quad (26)$$

by applying Lemma 3 to the above equation, we have that

$$\begin{aligned}\sqrt{\prod_{t=1}^k \lambda_{i_t}(\tilde{\mathbf{W}}^\top \tilde{\mathbf{W}}\mathbf{Q}\mathbf{Q}^\top)} &\geq \sqrt{\prod_{t=1}^k \lambda_{i_t}(\tilde{\mathbf{W}}^\top \tilde{\mathbf{W}}) \lambda_{n-t+1}(\mathbf{Q}\mathbf{Q}^\top)} \\ &= \prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}) \sigma_{n-t+1}(\mathbf{Q})\end{aligned}\quad (27)$$

$$\begin{aligned}\sqrt{\prod_{t=1}^k \lambda_{i_t}(\tilde{\mathbf{W}}^\top \tilde{\mathbf{W}}\mathbf{Q}\mathbf{Q}^\top)} &\leq \sqrt{\prod_{t=1}^k \lambda_{i_t}(\tilde{\mathbf{W}}^\top \tilde{\mathbf{W}}) \lambda_t(\mathbf{Q}\mathbf{Q}^\top)} \\ &= \prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}) \sigma_t(\mathbf{Q})\end{aligned}\quad (28)$$

Therefore, we have that

$$\prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}\mathbf{Q}) \geq \prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}) \sigma_{n-t+1}(\mathbf{Q}) \quad (29)$$

$$\prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}\mathbf{Q}) \leq \prod_{t=1}^k \sigma_{i_t}(\tilde{\mathbf{W}}) \sigma_t(\mathbf{Q}) \quad (30)$$

Suppose we have $k = 1$ and $i_1 = n$, then Eq. (29) gives

$$\sigma_n(\tilde{\mathbf{W}}\mathbf{Q}) \geq \sigma_n(\tilde{\mathbf{W}}) \sigma_n(\mathbf{Q}) \quad (31)$$

Then suppose we have $k = 1$ and $i_1 = 1$, then Eq. (30) gives

$$\sigma_1(\tilde{\mathbf{W}}\mathbf{Q}) \leq \sigma_1(\tilde{\mathbf{W}}) \sigma_1(\mathbf{Q}) \quad (32)$$

Combining the above results with Eq. (21) and Eq. (23), we have that

$$\begin{aligned}
\lim_{n \rightarrow \infty} \sigma_{\max}(\mathbf{W}) &= \lim_{n \rightarrow \infty} \sigma_{\max}(\tilde{\mathbf{W}} \cdot \mathbf{Q}) \leq \lim_{n \rightarrow \infty} (\sigma_{\max}(\tilde{\mathbf{W}}) \cdot \sigma_{\max}(\mathbf{Q})) \\
&= (\sqrt{d} + \sqrt{\lambda d}) \cdot \max_i \frac{1}{\|\tilde{\mathbf{v}}_i\|_2} \\
\lim_{n \rightarrow \infty} \sigma_{\min}(\mathbf{W}) &= \lim_{n \rightarrow \infty} \sigma_{\min}(\tilde{\mathbf{W}} \cdot \mathbf{Q}) \geq \lim_{n \rightarrow \infty} (\sigma_{\min}(\tilde{\mathbf{W}}) \cdot \sigma_{\min}(\mathbf{Q})) \\
&= (\sqrt{d} - \sqrt{\lambda d}) \cdot \min_i \frac{1}{\|\tilde{\mathbf{v}}_i\|_2}
\end{aligned} \tag{33}$$

which concludes the proof. $\square$

Combing with the fact that

$$\lim_{n \rightarrow \infty} \max \frac{\|\mathbf{v}_i\|_2}{\sqrt{d}} = \lim_{n \rightarrow \infty} \min \frac{\|\mathbf{v}_i\|_2}{\sqrt{d}} = 1, \tag{34}$$

we essentially have that

$$\begin{aligned}
\lim_{n \rightarrow \infty} \sigma_{\max}(\mathbf{W}) &\rightarrow 1 + \sqrt{\lambda}, \\
\lim_{n \rightarrow \infty} \sigma_{\min}(\mathbf{W}) &\rightarrow 1 - \sqrt{\lambda}.
\end{aligned} \tag{35}$$

which can be written to the following results:

$ \begin{aligned} \sigma_{\max}(\mathbf{W}) &\xrightarrow[n \rightarrow \infty]{\text{a.s.}} 1 + \sqrt{\lambda} \\ \sigma_{\min}(\mathbf{W}) &\xrightarrow[n \rightarrow \infty]{\text{a.s.}} 1 - \sqrt{\lambda} \end{aligned} \tag{36} $

which shows that under our proposed normalized Gaussian initialization, the maximal and minimal singular values are well bounded by a constant that is only dependent on the size of weight matrix. These results justify the effectiveness of our proposed normalized Gaussian initialization in POET.

D Proofs of Lemma 1

Proof of Lemma 1. We consider an orthogonal matrix $\mathbf{R}$ and orthogonal primitives $\mathbf{G}^i$ corresponding to uniformly random subsets $\mathbf{S}^j \subset [m]$ of size b as explained in the main text (see equation (4)). The main claim we need to prove is that given any vector $\mathbf{v} \in \mathbb{R}^m$ and a set $\mathbf{S} \subset [m]$ with $k \in [m]$ we can find an orthogonal primitive matrix $\mathbf{G}$ corresponding to the set $\mathbf{S}$ such that

$$\begin{aligned} (\mathbf{G}\mathbf{v})_l &= 0 \quad \text{for } l \in \mathbf{S} \text{ with } l > k \\ (\mathbf{G}\mathbf{v})_k &\geq 0 \\ (\mathbf{G}\mathbf{v})_l &= v_l \quad \text{for } l \notin \mathbf{S}. \end{aligned} \tag{37}$$

Moreover, for all $\mathbf{w} \in \mathbb{R}^m$ with $w_i = 0$ for $i \geq k$ the relation

$$\mathbf{G}\mathbf{w} = \mathbf{w} \tag{38}$$

holds. We can assume that the matrix $\mathbf{D}(\mathbf{S}) = \{e(s_1), \dots, e(s_b)\}$ contains the entries s_i in ascending order. Then we write

$$\mathbf{D}(\mathbf{S})^\top \mathbf{v} = \begin{pmatrix} \tilde{v}_1 \\ \tilde{v}_2 \end{pmatrix} \tag{39}$$

where $\tilde{v}_1 \in \mathbb{R}^{b_1}$ corresponds to the entries s_i with $s_i < k$ and $\tilde{v}_2 \in \mathbb{R}^{b_2}$ to the remaining entries, in particular $s_{b_1+1} = k$ because $k \in \mathbf{S}$. It is well known that for every vector $\mathbf{v}$ there is a rotation $\mathbf{Q}$ aligning $\mathbf{v}$ with the first standard basis vector, i.e., such that $\mathbf{Q}\mathbf{v} = \lambda e(1)$ for some $\lambda \geq 0$. Consider such a matrix $\tilde{\mathbf{Q}}$ for the vector $\tilde{v}_2$ and then define the orthogonal matrix

$$\tilde{\mathbf{G}} = \begin{pmatrix} \mathbf{1}_{b_1} & \mathbf{0}_{b_1 \times b_2} \\ \mathbf{0}_{b_2 \times b_1} & \tilde{\mathbf{Q}} \end{pmatrix}. \tag{40}$$

Careful inspection of (4) implies that the last part of (37) is actually true for any $\tilde{\mathbf{G}}$ as the second term has rows with all entries equal to zero for all $l \notin \mathbf{S}$. For the first part we find

$$\mathbf{D}(\mathbf{S})\tilde{\mathbf{G}}\mathbf{D}(\mathbf{S})^\top \mathbf{v} = \mathbf{D}(\mathbf{S})\tilde{\mathbf{G}} \begin{pmatrix} \tilde{v}_1 \\ \tilde{v}_2 \end{pmatrix} = \mathbf{D}(\mathbf{S}) \begin{pmatrix} \tilde{v}_1 \\ \lambda e(1) \end{pmatrix} = \sum_{i \leq b_1} e(s_i)(\tilde{v}_1)_i + \lambda e(k). \tag{41}$$

Here we used $s_{b_1+1} = k$ in the last step. Since in addition

$$((\mathbf{1}_m - \mathbf{D}(\mathbf{S}) \cdot \mathbf{1}_b \cdot \mathbf{D}(\mathbf{S})^\top) \mathbf{v})_l = 0 \tag{42}$$

for all $l \in \mathbf{S}$ we conclude that indeed $(\mathbf{G}\mathbf{v})_l = 0$ for $l \in \mathbf{S}$ and $l > k$, $(\mathbf{G}\mathbf{v})_k \geq 0$. The remaining statement (38) follows from the observation that when decomposing as in (39) we find

$$(\mathbf{D}(\mathbf{S}))^\top \mathbf{w} = \begin{pmatrix} \tilde{w}_1 \\ \mathbf{0}_{b_2} \end{pmatrix} \tag{43}$$

(because $w_i = 0$ for $i \geq k$) and therefore

$$(\tilde{\mathbf{G}} - \mathbf{1}_b)(\mathbf{D}(\mathbf{S}))^\top \mathbf{w} = \mathbf{0}_b \tag{44}$$

by definition of $\tilde{\mathbf{G}}$ and we find $\mathbf{G}\mathbf{w} = \mathbf{w}$.

The rest of the proof is straightforward by induction combined with a simple coin collector problem. For the rest of the proof it is convenient to reverse the indices, i.e., to consider products $\mathbf{G}_c \cdot \dots \cdot \mathbf{G}_1$. Assume that we have chosen $\mathbf{G}_i$ for $i \leq c_k$ and some $c_k \in \mathbb{N}$ such that the product

$$\mathbf{P}^k = \mathbf{G}_{c_k} \cdot \dots \cdot \mathbf{G}_1 \cdot \mathbf{R}^\top \tag{45}$$

satisfies $\mathbf{P}_{l',k'}^k = 0$ for all $k' < k$ and $l' > k'$ and $\mathbf{P}_{k',k'}^k \geq 0$ for $k' < k$. Let $c_{k+1} \geq c_k + \alpha(m/b)^2 \ln(m)$. Then, we can bound for any $l > k$ the probability that there is no $c_k < j \leq c_{k+1}$ such that $\{k, l\} \subset \mathbf{S}^j$ using that $\mathbf{S}^j$ follows a uniform i.i.d. distribution by

$$\mathbb{P}(\nexists c_k < j \leq c_{k+1} : k, l \in \mathbf{S}^j) \leq \left(1 - \frac{b^2}{m^2}\right)^{c_{k+1} - c_k} \leq \exp\left(-\frac{b^2}{m^2} \cdot \alpha \frac{m^2}{b^2} \ln(m)\right) = m^{-\alpha}. \tag{46}$$

The union bound implies that with probability at least $1 - m^{-\alpha+1}$ there is for all $l > k$ a $c_k < j \leq c_{k+1}$ such that $\{k, l\} \subset S^j$. If this holds we set G_j for $c_k < j \leq c_{k+1}$ as constructed above if $k \in S^j$ and $G_j = \mathbf{1}_m$ otherwise. This then ensures that

$$P^{k+1} = G_{c_{k+1}} \cdot \dots \cdot G_1 \cdot R^\top \quad (47)$$

satisfies $P_{l',k'}^{k+1} = 0$ for $k' \leq k$ and $l' > k'$. For $k' < k$ this follows from (38) and for $k' = k$ from (37). We conclude by the union bound that P^m is an upper triangular matrix with non-negative diagonal entries with probability at least $1 - mm^{-\alpha+1} = 1 - m^{-(\alpha-2)}$. But we also know that P^m is orthogonal and therefore satisfies $P^m = \mathbf{1}_m$ and we thus find

$$G_{c_m} \cdot \dots \cdot G_1 = R. \quad (48)$$

□

Next we give a heuristic that actually $O(\ln(m)m^2/b^2)$ terms are sufficient to express every orthogonal map as a product of stochastic primitives. For fixed c we consider the map

$$\Phi : O(b)^c \rightarrow O(m) \quad \Phi(\tilde{G}_1, \dots, \tilde{G}_c) = \prod_{j=1}^c G_j. \quad (49)$$

If $c \geq \alpha \ln(m)m^2/b^2$ we have that with probability at least $1 - m^{-(\alpha-2)}$ for all $k, l \in [m]$ there is $j \leq c$ such that $k, l \in S^j$. Assume that this is the case. Recall that the tangent space of $O(k)$ at the identity is the space of skew-symmetric matrices. Consider a tangent vector $(X_1, \dots, X_c)$ with $X_i \in \text{Skew}(k)$. Then

$$D\Phi(\mathbf{1}_b, \dots, \mathbf{1}_b)(X_1, \dots, X_c) = \sum_{j=1}^c D(S^j) \cdot X_j \cdot D(S^j)^\top. \quad (50)$$

This is a surjective map on $\text{Skew}(m)$ under the condition that for all $k, l \in [m]$ there is $j \leq c$ such that $k, l \in S^j$. We can therefore conclude that the image of Φ contains a neighbourhood of the identity. Moreover, since Φ is a polynomial map, $D\Phi$ is surjective everywhere except for a variety of codimension one. While this is not sufficient to conclude that the image of Φ is $O(d)$ or dense in $O(d)$ it provides some indication that this is the case.

E Experimental Details

Parameter	Llama 60M	Llama 130M	Llama 350M	Llama 1.3B
Hidden dimension	512	768	1024	2048
Intermediate dimension	1280	2048	2816	5376
Number of attention heads	8	12	16	32
Number of hidden layers	8	12	24	24

Table 8: Model architectures for different Llama variants.

Model	Spec.	# GPU	lr (base)	lr (POET)	training steps	batch size	grad acc.
Llama 60M	$b = 1/2$	1	1e-2	1e-3	300,000	256	2
	$b = 1/4$	1	1e-2	2e-3	300,000	256	2
	$b = 1/8$	1	1e-2	4e-3	300,000	256	2
Llama 130M	$b = 1/2$	1	5e-3	1e-3	400,000	128	2
	$b = 1/4$	1	5e-3	2e-3	400,000	128	2
	$b = 1/8$	1	5e-3	4e-3	400,000	128	2
Llama 350M	$b = 1/2$	4	5e-3	1e-3	400,000	128	1
	$b = 1/4$	4	5e-3	2e-3	400,000	128	1
	$b = 1/8$	4	5e-3	4e-3	400,000	128	1
Llama 1.3B	$b = 1/2$	8	1e-3	1e-3	500,000	64	1
	$b = 1/4$	8	1e-3	2e-3	500,000	64	1
	$b = 1/8$	8	1e-3	4e-3	500,000	64	1

Table 9: Hyper-parameter setup of POET-FS.

This section outlines our experimental setup, including the codebase, datasets, and computational resources used.

Code framework. Our method is implemented on top of the codebase from [82]¹ (Apache 2.0 license), which we also use to reproduce the AdamW and GaLore baselines. We will release our code for reproducing all training results prior to publication.

Training details. We employed the AdamW optimizer [55] for all our training runs. The specific hyperparameters used for each experiment are detailed in the Table 9 and Table 10 referenced below. We use the cosine learning rate scheduler with the minimum learning ratio of 0.01. We use the number of warmup steps of 0, weight decay of 0.01 and gradient clipping of 0.1. For the AdamW baseline, we report results for the optimal learning rate from $[1 \times 10^{-2}, 5 \times 10^{-3}, 1 \times 10^{-3}, 5 \times 10^{-4}, 1 \times 10^{-4}, 5 \times 10^{-5}, 1 \times 10^{-5}]$. After each merge-then-reinitialize step, we additionally increase the gradient clipping for 10 training steps to improve training stability.

Model architecture. Our work utilized the **Hugging Face Transformers**² code base to construct the Llama model for pretraining, which is under the **Apache 2.0** license. The specific layer setups for the different scaled Llama models are summarized in Table 8. Note, the intermediate dimension of the Feed-Forward Network (FFN) has been slightly modified for the POET-BS, compared to the configs in [82], because the linear layer dimensions have to be divisible by the POET-BS block size b .

Dataset. We use the *Colossal Clean Crawled Corpus* (C4) dataset [66] for pretraining. The C4 data is a large-scale, meticulously cleaned version of Common Crawl’s web crawl corpus. It was

¹<https://github.com/jiaweizzhao/GaLore>

²<https://github.com/huggingface/transformers>

Model	Spec.	# GPU	lr (base)	lr (POET)	training steps	batch size	grad acc.
Llama 60M	$b = 256$	1	1e-2	1e-3	300,000	256	2
	$b = 128$	1	1e-2	2e-3	300,000	256	2
	$b = 64$	1	1e-2	4e-3	300,000	256	2
Llama 130M	$b = 256$	1	5e-3	1e-3	400,000	256	2
	$b = 128$	1	5e-3	2e-3	400,000	256	2
	$b = 64$	1	5e-3	4e-3	400,000	256	2
Llama 350M	$b = 256$	4	5e-3	1e-3	400,000	128	1
	$b = 128$	4	5e-3	2e-3	400,000	128	1
	$b = 64$	4	5e-3	4e-3	400,000	128	1
Llama 1.3B	$b = 256$	8	1e-3	1e-3	500,000	64	1
	$b = 128$	8	1e-3	2e-3	500,000	64	1
	$b = 64$	8	1e-3	4e-3	500,000	64	1

Table 10: Hyper-parameter setup of POET-BS.

originally introduced for training the Text-to-Text Transfer Transformer (T5) model and has since become a standard pre-training dataset for testing training algorithms for pre-training large language models. The dataset is released under the **ODC-BY** license.

Compute Resources. All the training tasks are performed on a **NVIDIA HGX H100 8-GPU System** node with 80GB memory each. Depending on the model scale, we train on 1, 4 or 8 GPUs.

F Implementation and CUDA Acceleration

To enable efficient POET training, we implement the Cayley–Neumann parameterization. To reduce memory usage, we leverage the structure of the skew-symmetric matrix $\mathbf{Q} \in \mathbb{R}^{n \times n}$, where the diagonal entries are zero ($Q_{ii} = 0$) and off-diagonal elements satisfy $Q_{ij} = -Q_{ji}$. This structure allows us to store only the upper triangular part of $\mathbf{Q}$ as a vector, reducing the number of trainable parameters from n^2 to $n(n - 1)/2$. During the forward pass, $\mathbf{Q}$ is reconstructed on-the-fly using a specialized CUDA kernel, significantly accelerating this process. In addition, the Neumann approximation removes the need for costly and numerically unstable matrix inversion, offering further computational gains. Overall, training a 1.3B LLaMA model on a single H100 8-GPU node yields a 3.8× speedup over the baseline (*i.e.*, native implementation). Table 11 summarizes the contribution of each component to the overall training time.

Design	Speed-Up
Neumann approximation	1.5×
Skew-symmetric CUDA kernel	1.3×
Total	3.8×

Table 11: Method design and clock time speed-up.

G Results of Vector Probing for R and P

In this ablation study, we perform vector probing on the orthogonal matrices $R \in \mathbb{R}^{m \times m}$, $P \in \mathbb{R}^{n \times n}$ for all linear layers for all blocks of a 60M Llama model trained with POET-FS. The cosine similarity results are reported in Figure 11 and Figure 12, and the trace results are reported in Figure 13 and Figure 14. Since we want to understand the learning dynamics of the orthogonal matrices, we employ $b = 1$ with POET learning rate of 5×10^{-4} to eliminate the need for resampling and reinitialization of the orthogonal matrices. Interestingly, we observe this three-phased learning dynamics across different types of linear layers and different-depth transformer blocks.

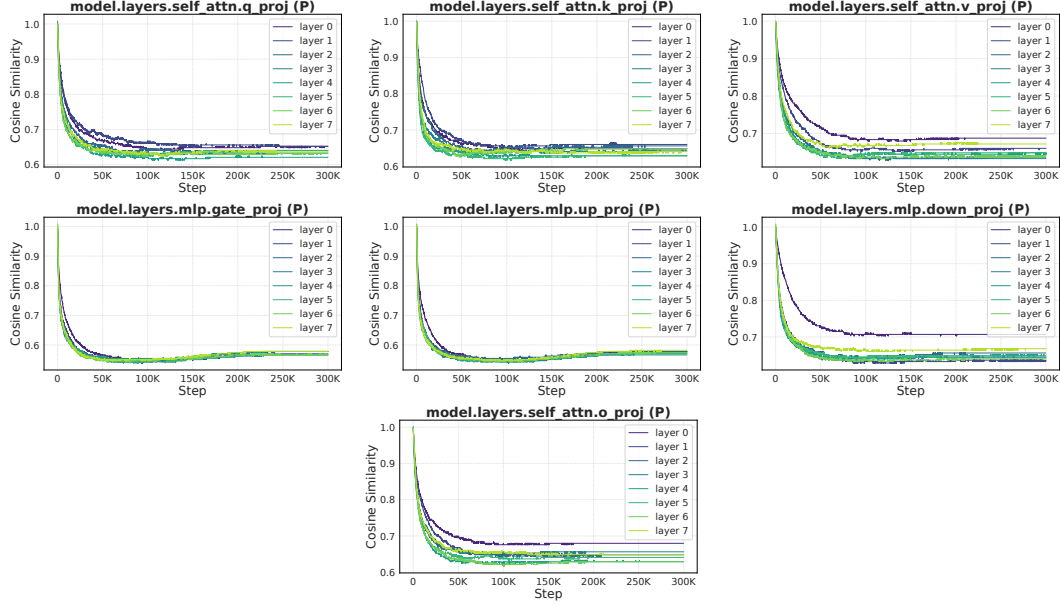


Figure 11: Cosine similarity for vector probing of P across the self-attention components (query, key, value, and output projections) and feed-forward network components (up-, down-, and gate-projections) in all transformer blocks of a POET-trained Llama 60M model.

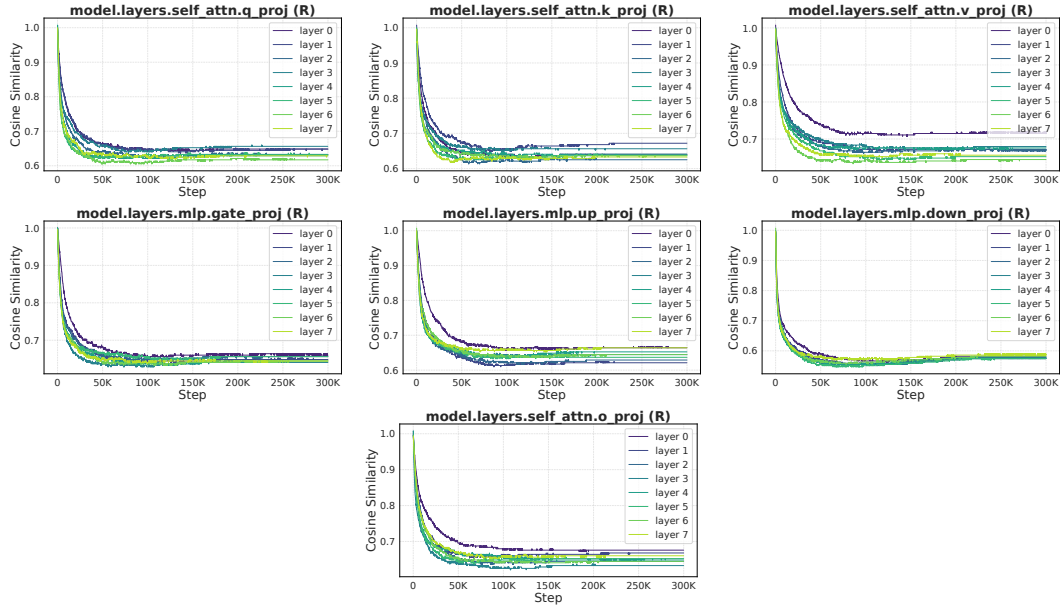


Figure 12: Cosine similarity for vector probing of R across the self-attention components (query, key, value, and output projections) and feed-forward network components (up-, down-, and gate-projections) from all transformer blocks of a POET-trained Llama 60M model.

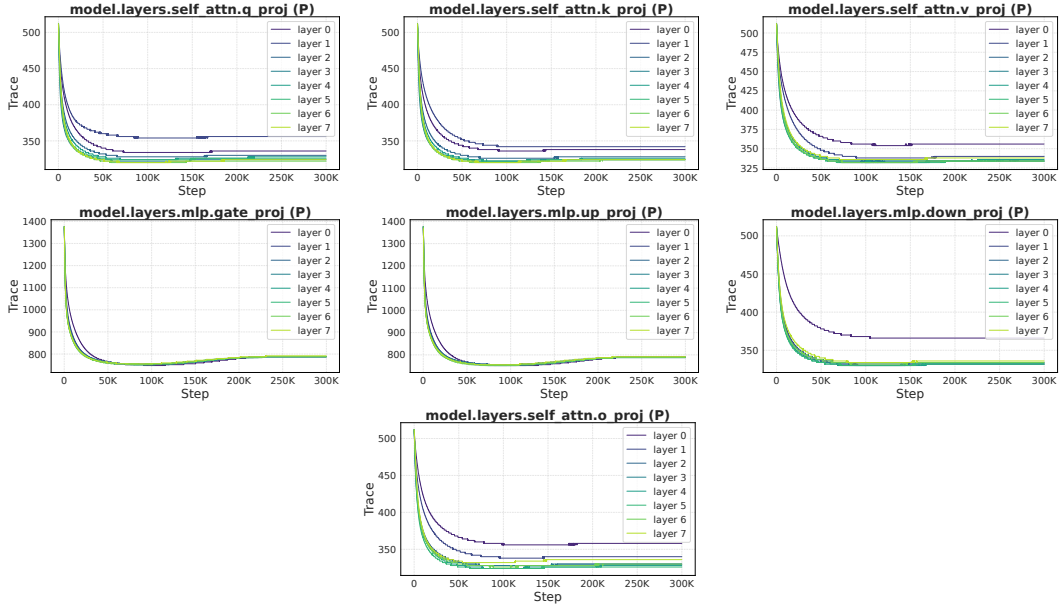


Figure 13: Trace of P across the self-attention components (query, key, value, and output projections) and feed-forward network components (up-, down-, and gate-projections) from all transformer blocks of a POET-trained Llama 60M model.

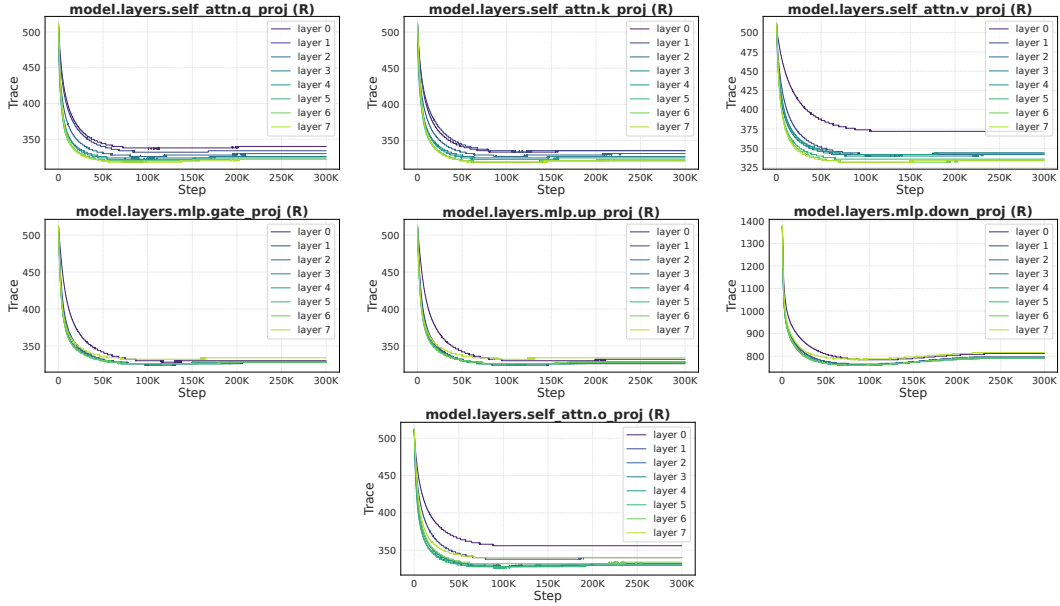


Figure 14: Trace of R across the self-attention components (query, key, value, and output projections) and feed-forward network components (up-, down-, and gate-projections) from all transformer blocks of a POET-trained Llama 60M model.

H Weight Update Evenness of Different POET Variants

To understand the higher parameter efficiency of POET-BS compared to POET-FS, we employ a toy example to visualize their different weight update mechanisms by counting the total number of updates for each element of the weight matrix. The visualization results are given in Figure 15 and Figure 16. Specifically, in this experiment, a 64×64 matrix was randomly initialized and trained for 100 steps under various POET-BS and POET-FS configurations. The merge-then-reinitialize trick is performed at each iteration, and the same set of weight elements was effectively updated between two successive merge-then-reinitialize operations. For each weight element, we compute its total number of update in these 100 steps.

Given 100 training steps and updates from both R and P , each element of the weight matrix can be updated at most 200 times. This target is consistently achieved by POET-BS, and it is also agnostic to the block size. All POET-BS variants can enable the maximal number of updates for each weight element to be 200. In contrast, POET-FS results in significantly fewer updates per weight element, with updates also unevenly distributed. This unevenness arises from stochasticity, causing certain weights to be updated more frequently than others. While this is less problematic at large iteration counts, it can introduce unexpected training difficulties in earlier stages.

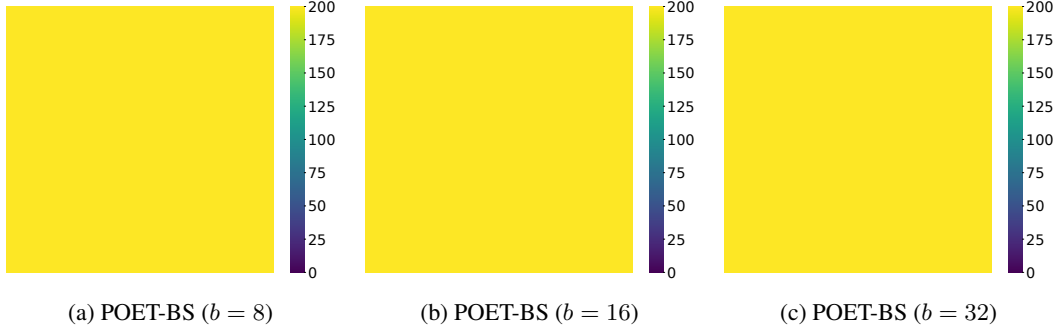


Figure 15: Visualization of the weight update mechanism of POET-BS after 100 steps of update and $T_m = 1$.

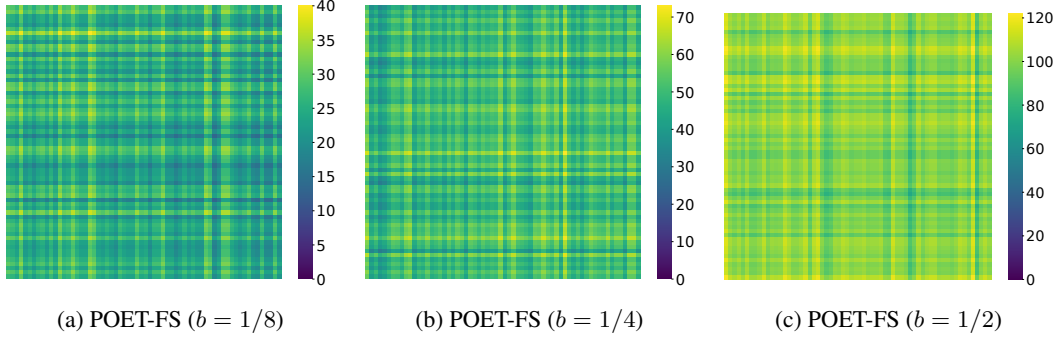


Figure 16: Visualization of the weight update mechanism of POET-FS after 100 steps of update and $T_m = 1$.

I Training Dynamics of Singular Values

We conduct an ablation study to compare the training dynamics of singular values of weight matrices between **AdamW** and **POET**. The results of AdamW are given in Figure 17, Figure 18 and Figure 19. The results of POET are given in Figure 20, Figure 21 and Figure 22. A 60M LLaMA model was trained for 50,000 iterations with an effective batch size of 512, using both AdamW and POET-FS ($b = 1/2$). The model was evaluated every 5,000 steps, and the singular value dynamics are computed by performing singular value decomposition on the weight matrices. For POET, a merge-then-reinitialize step was applied before each evaluation. Training is finished at 50,000 steps, as the spectral norm of the AdamW-trained model plateaued at this point.

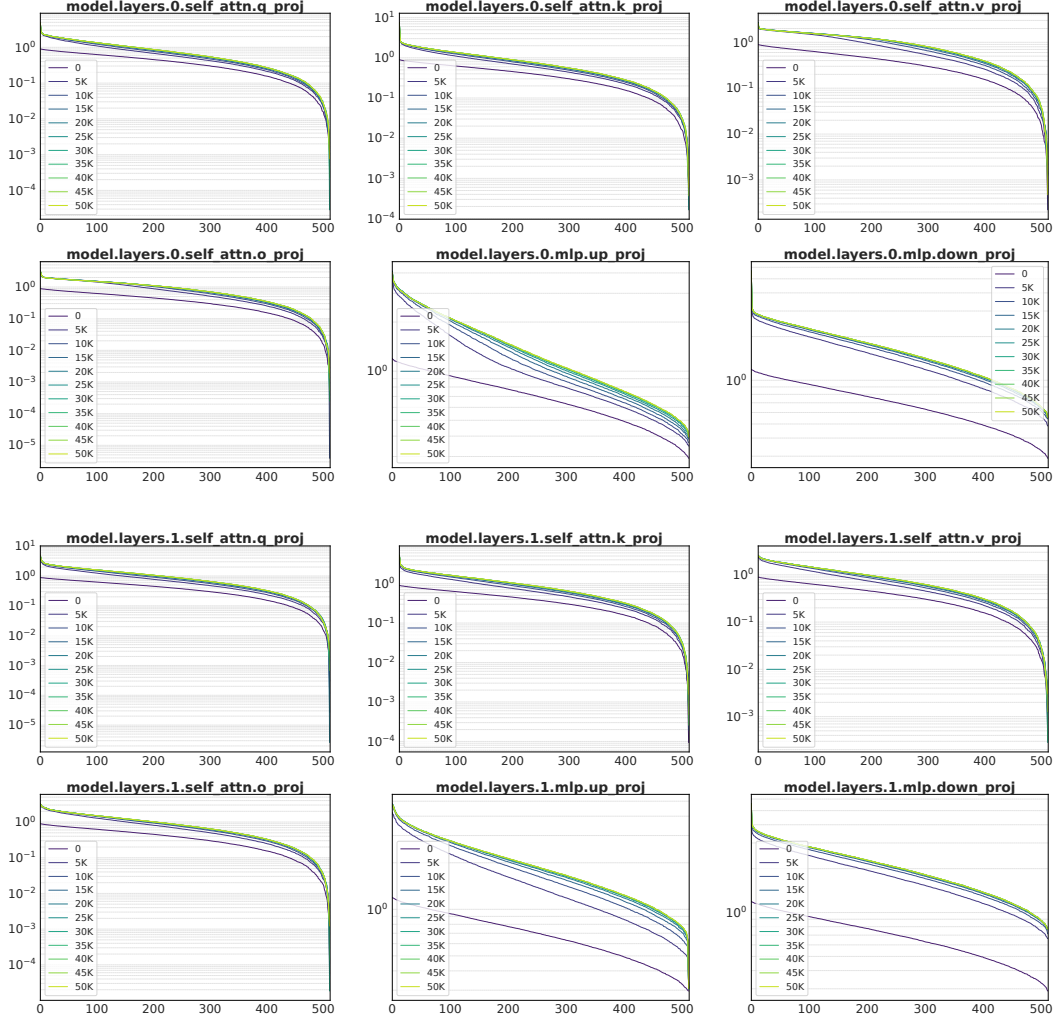


Figure 17: Training dynamics of the singular values of weight matrices within Blocks 0–1 (the i -th row represents Block i) of a 60M Llama Transformer trained with **AdamW**.

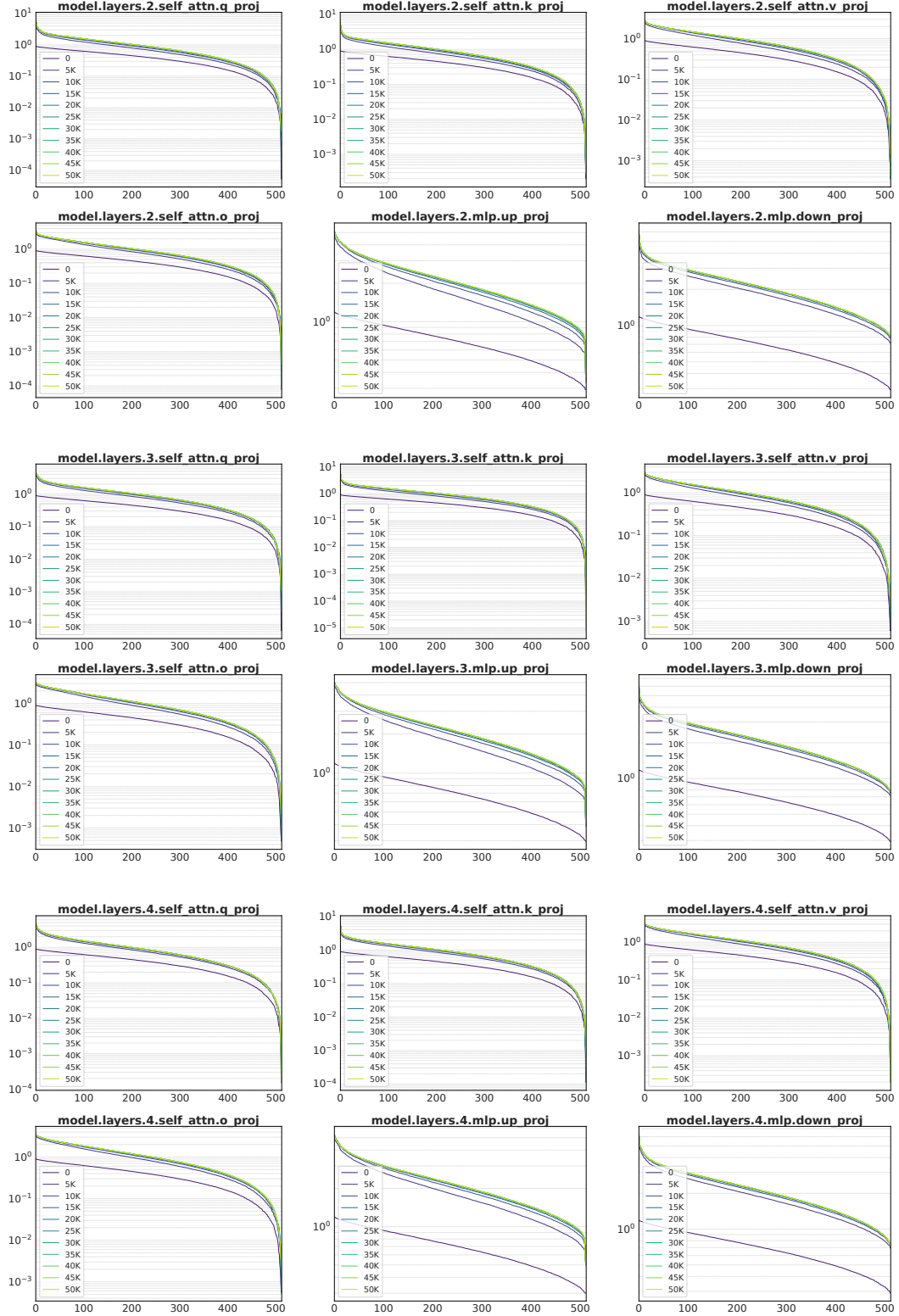


Figure 18: Training dynamics of the singular values of weight matrices within Blocks 2–4 (the i -th row represents Block i) of a 60M Llama Transformer trained with AdamW.

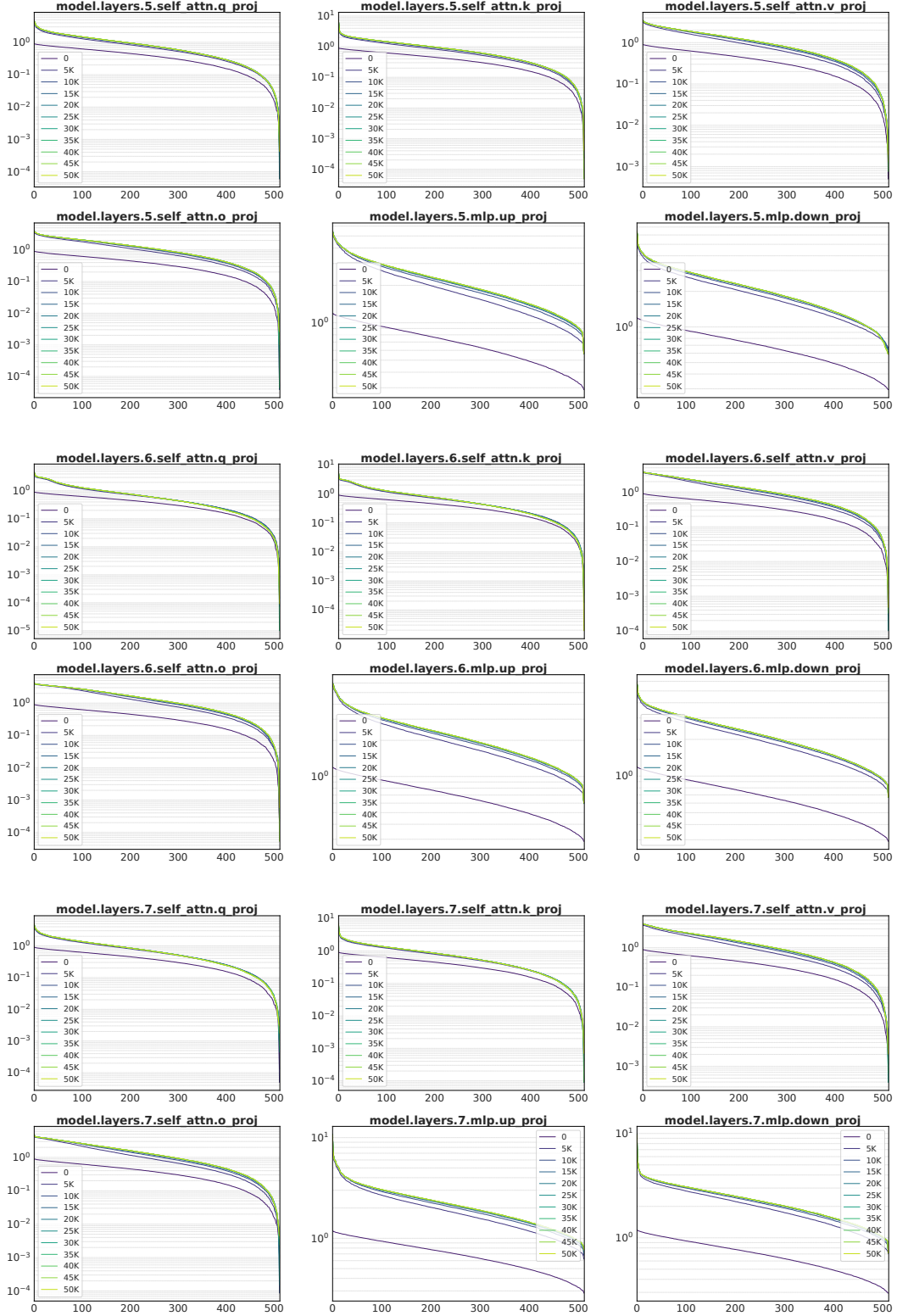


Figure 19: Training dynamics of the singular values of weight matrices within Blocks 5–7 (the i -th row represents Block i) of a 60M Llama Transformer trained with AdamW.

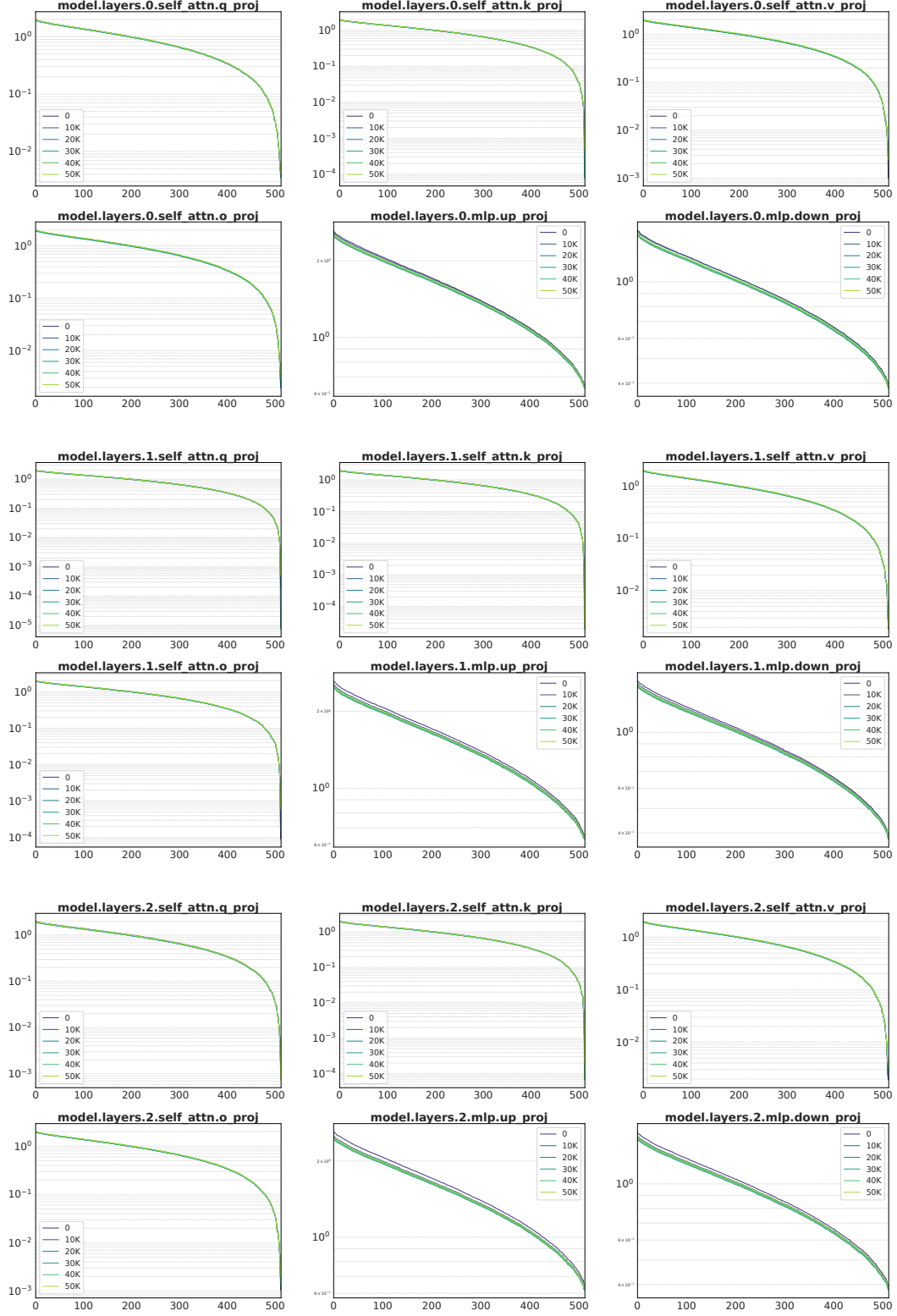


Figure 20: This plot illustrates the singular value training dynamics for individual weight matrices within Blocks 0-2 of a 60M Llama transformer model trained with **POET**. For each block, the dynamics are shown for the self-attention components (query, key, value, and output projections) and the feed-forward network components (up-projection, down-projection, and gate-projection).

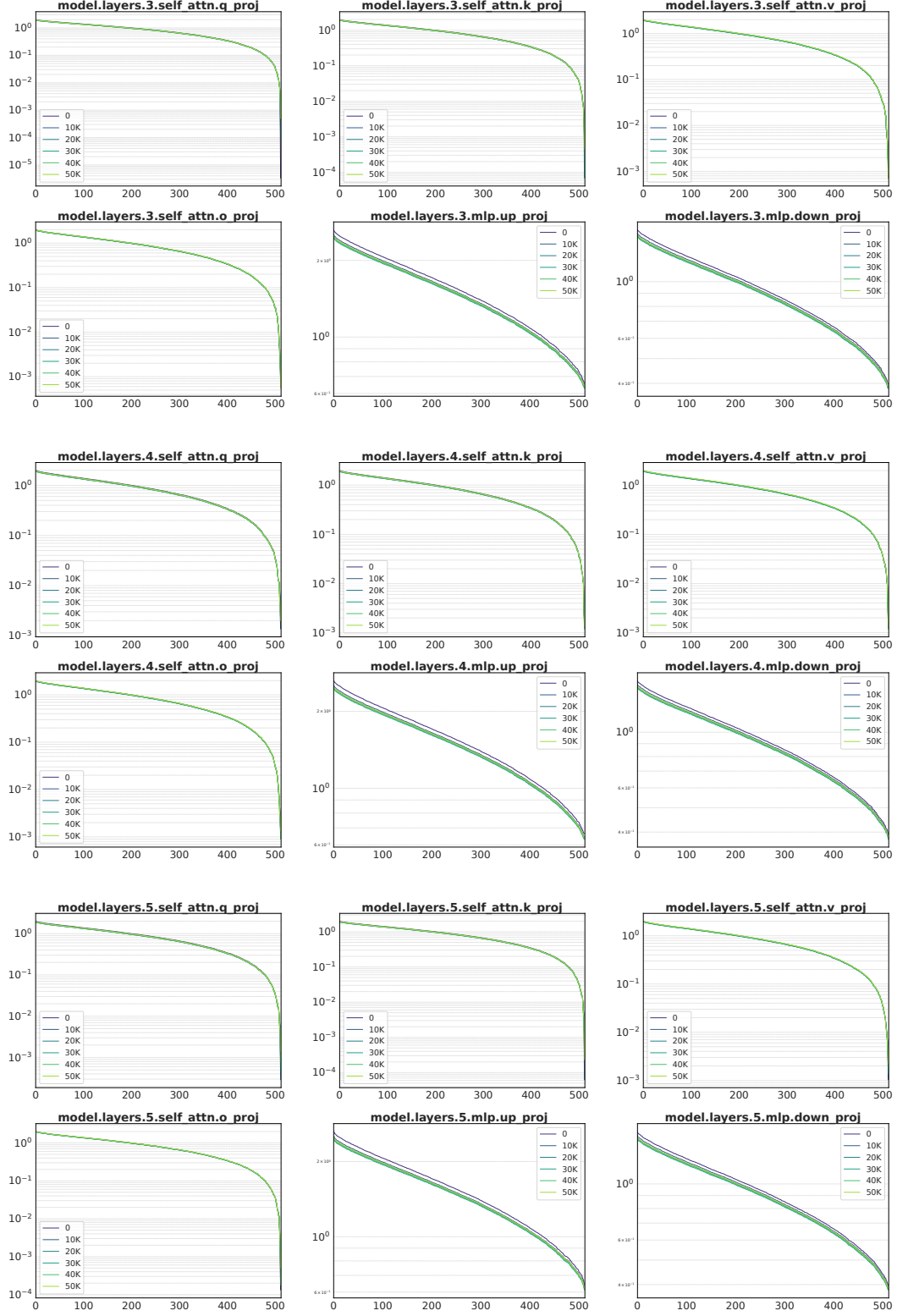


Figure 21: This plot illustrates the singular value training dynamics for individual weight matrices within Blocks 3-5 of a 60M Llama transformer model trained with **POET**. For each block, the dynamics are shown for the self-attention components (query, key, value, and output projections) and the feed-forward network components (up-projection, down-projection, and gate-projection).

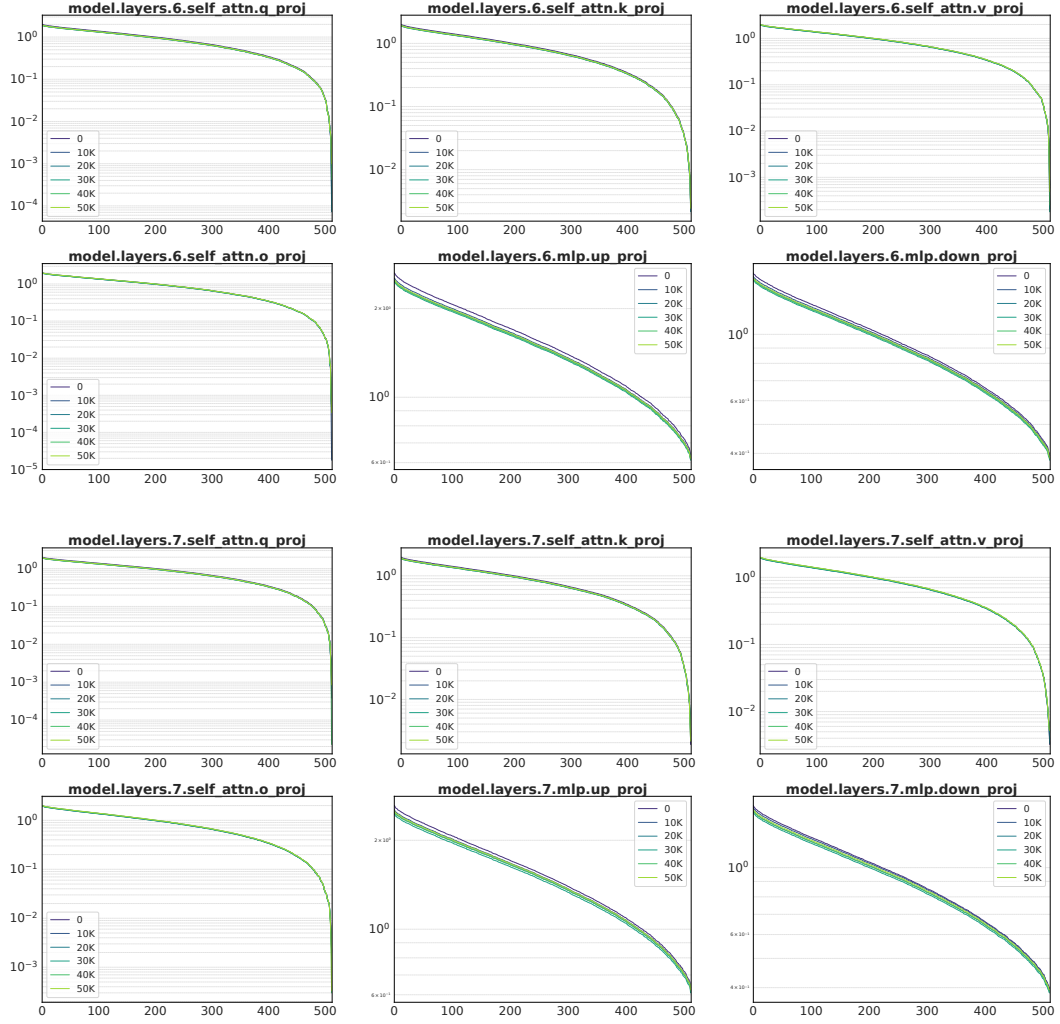


Figure 22: This plot illustrates the singular value training dynamics for individual weight matrices within Blocks 6-7 of a 60M Llama transformer model trained with **POET**. For each block, the dynamics are shown for the self-attention components (query, key, value, and output projections) and the feed-forward network components (up-projection, down-projection, and gate-projection).

J Orthogonality Approximation Quality using Neumann Series

In this ablation study, we evaluate the approximation error of the orthogonal matrices $R \in \mathbb{R}^{m \times m}$ and $P \in \mathbb{R}^{n \times n}$ across all linear layers in Block 0 of a 130M LLaMA model trained with POET-FS ($b = 1/2$) for 10,000 steps. Figure 23 and Figure 24 show the approximation error over the first 1,000 steps. Since the error difference between $k = 3$ and $k = 4$ was negligible, we used $k = 3$ for better computational efficiency. Empirically, while $k = 1$ or $k = 2$ suffices for smaller LLaMA models, larger k values are needed to avoid training divergence caused by exploding gradients due to approximation error.

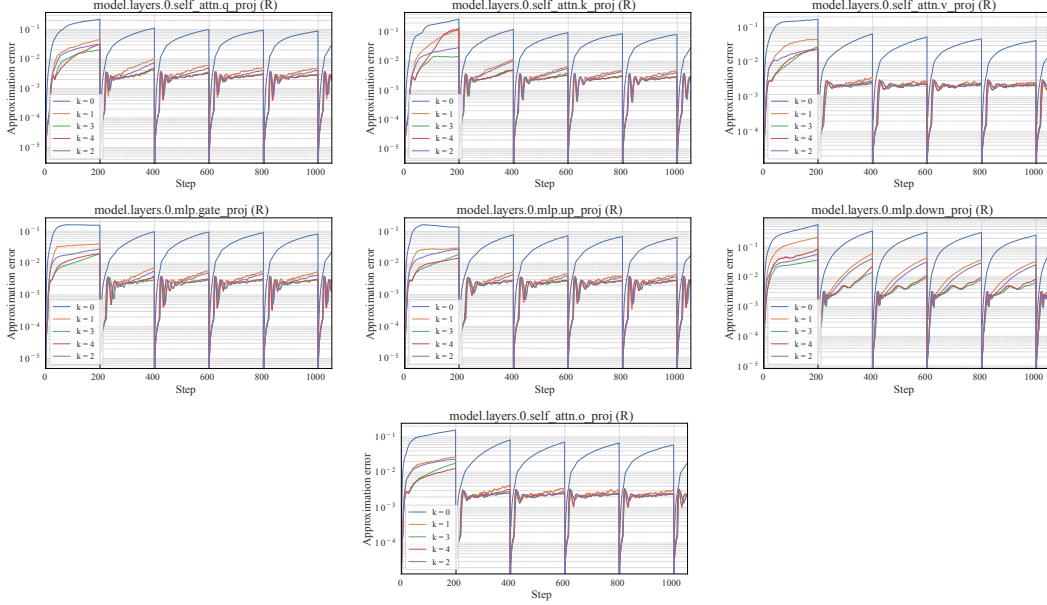


Figure 23: For the transformer block 0, we show approximation error of orthogonal matrix R for the self-attention components (query, key, value, and output projections) and the feed-forward network components (up-projection, down-projection, and gate-projection).

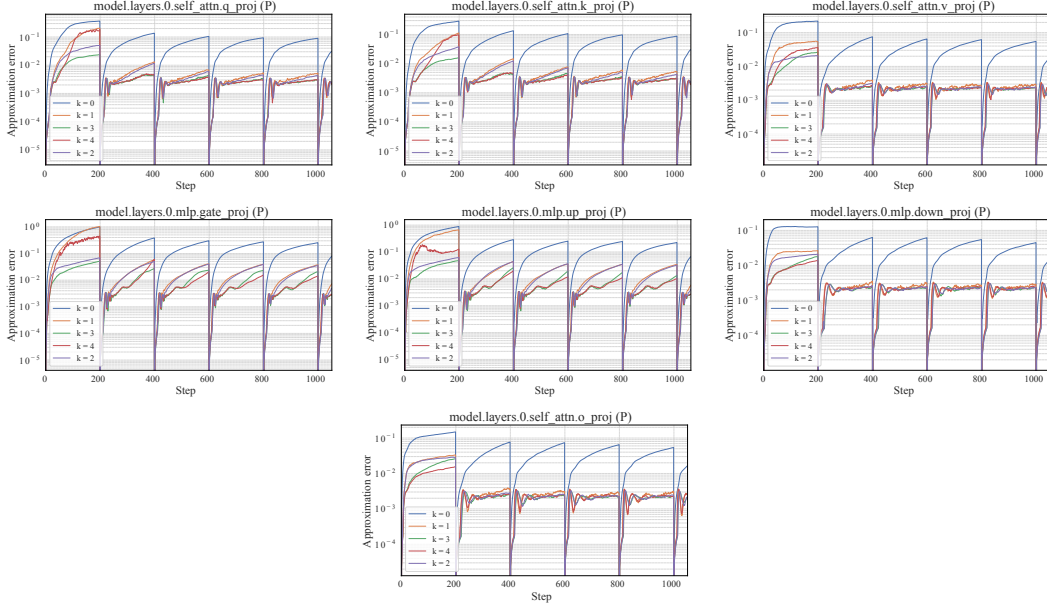


Figure 24: For the transformer block 0, we show approximation error of orthogonal matrix P for the self-attention components (query, key, value, and output projections) and the feed-forward network components (up-projection, down-projection, and gate-projection).

Additionally, Figure 25 shows the orthogonality approximation error of Neumann series with different k over the first 10,000 training steps, illustrating how it decreases as training progresses. We observe a general downward trend in approximation error, indicating improved approximation over time. The results also suggest that using too few Neumann series terms can lead to training divergence in POET.

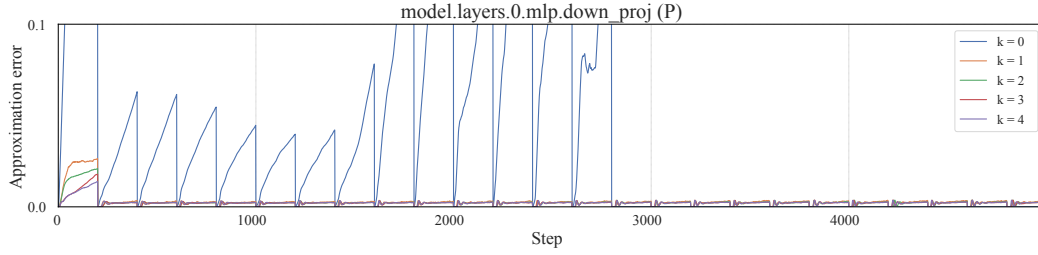


Figure 25: The approximation error of orthogonal matrix P in a randomly selected down-projection layer after training 10000 steps.

K Full Results of Training Dynamics

We provide the full training dynamics of different POET variants under Llama 60M, Llama 130M, Llama 350M and Llama 1.3B in Figure 26. This figure is essentially an extended result of Figure 8. One can observe that the training dynamics of POET is quite different from AdamW, and more importantly, POET consistently yields better parameter-efficiency and generalization.

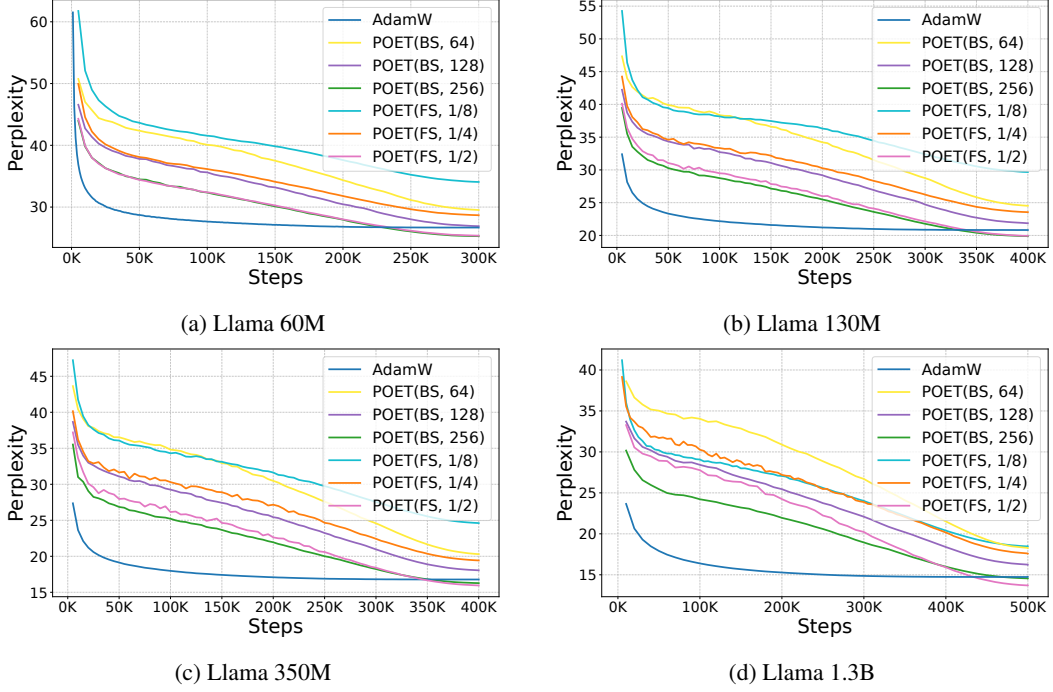


Figure 26: Validation perplexity during the training of the Llama-based transformer with 60M, 130M, 350M and 1.3B parameters.

L More Results of POET as a Finetuning Method

To demonstrate the applicability of POET to general finetuning tasks, we apply it to finetune a BART-large model [41] on the NLP task of text summarization. Specifically, we evaluate POET on the XSum [61] and CNN/DailyMail [24] datasets, reporting ROUGE-1/2/L scores in Table 12. We note that both LoRA and OFT are designed solely for parameter-efficient finetuning and are not applicable to pretraining. Our goal here is to demonstrate that POET is also effective as a finetuning method. For consistency, we use the same configuration as in the pretraining setup, resulting in a higher parameter count. Experimental results show that POET not only supports finetuning effectively but also outperforms both full-model finetuning and parameter-efficient methods.

Method	# Params	XSum	CNN/DailyMail
LoRA ($r=32$)	17.30M	43.38 / 20.20 / 35.25	43.17 / 20.31 / 29.72
OFT ($b=64$)	8.52M	44.12 / 20.96 / 36.01	44.08 / 21.02 / 30.68
Full FT	406.29M	45.14 / 22.27 / 37.25	44.16 / 21.28 / 40.90
POET (FS, $b=1/2$)	144.57M	45.23 / 22.41 / 37.28	44.27 / 21.29 / 41.02

Table 12: Finetuning BART-large on XSum and CNN/DailyMail for text summarization. We report ROUGE-1/2/L results (higher is better).

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and introduction reflect our work's goal, scope, idea and contribution.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We have included analysis in the paper to discuss our method's limitations.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Yes, we provide a detailed derivation of the theory and reference all theoretical results our method is based on.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: Yes, all methods are described in details to ensure the reproducibility of the results. We will also release the code and data.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: Yes, the code and data will be released to ensure full reproducibility.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We have described the training and testing in details both in the main paper and also in the Appendix section to ensure reproducibility and ease the understanding.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We have described the experiments in detail in the Appendix sections, ensuring reproducibility and soundness of our results.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have described the experiments in the main paper, more focusing on its setting and its results and have included the specific type of compute we used to develop our methods. We have included detailed descriptions in our Appendix section.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Yes, our research conforms with the ethics guidelines of NeurIPS.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: Our method mainly focuses on methods and its theoretical implications and other interesting findings. Our work does not facilitate down-stream applications that will have a negative impact on the society.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.

- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We mainly work on methods, we will release the code but will not release any data or trained model.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We will credit the owners of the code, data and models in our Appendix section, where we describe the implementation more thoroughly. We already included the citations for the assets we used in the main paper. Here are an overview of the most important licenses: Bart (License: apache-2.0), Llama 2 (LLAMA 2 COMMUNITY LICENSE AGREEMENT), Code (apache-2.0), C4 dataset (ODC-BY).

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.

- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: Our project does not release any new assets (dataset, model weights etc.).

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: Our project does not involve any crowd sourcing or involve human experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Our project does not involve any crowd sourcing or involve human experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: We did not use LLMs to develop our methods but developed it independently.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.