

Test-time Diverse Reasoning by Riemannian Activation Steering

Ly Tran Ho Khanh^{*1}, Dongxuan Zhu^{*1}, Man-Chung Yue², Viet Anh Nguyen¹,

¹The Chinese University of Hong Kong ²The University of Hong Kong
hokhanhlytran@cuhk.edu.hk, dxzhu@se.cuhk.edu.hk, mc Yue@hku.hk, nguyen@se.cuhk.edu.hk

Abstract

Best-of- N reasoning improves the accuracy of language models in solving complex tasks by sampling multiple candidate solutions and then selecting the best one based on some criteria. A critical bottleneck for this strategy is the output diversity limit, which occurs when the model generates similar outputs despite stochastic sampling, and hence recites the same error. To address this lack of variance in reasoning paths, we propose a novel unsupervised activation steering strategy that simultaneously optimizes the steering vectors for multiple reasoning trajectories at test time. At any synchronization anchor along the batch generation process, we find the steering vectors that maximize the total volume spanned by all possible intervened activation subsets. We demonstrate that these steering vectors can be determined by solving a Riemannian optimization problem over the product of spheres with a log-determinant objective function. We then use a Riemannian block-coordinate descent algorithm with a well-tuned learning rate to obtain a stationary point of the problem, and we apply these steering vectors until the generation process reaches the subsequent synchronization anchor. Empirical evaluations on popular mathematical benchmarks demonstrate that our test-time Riemannian activation steering strategy outperforms vanilla sampling techniques in terms of generative diversity and solution accuracy.

Code — <https://github.com/lythk88/SPREAD>

1 Introduction

Language models (LMs) have revolutionized tasks from code generation (Chen et al. 2021), symbolic reasoning (Wei et al. 2022), to mathematical problem solving (Lewkowycz et al. 2022). In these tasks, the quality of the final solution can improve significantly by exploring multiple plausible reasoning paths and then presenting the best solution aggregated from the information from these paths. This strategy aligns with our problem-solving intuitions, where each path of reasoning explores different techniques, generalizability, and scientific values. A simple and effective method to exploit this explore-then-aggregate idea is Best-of- N sampling (Lightman et al. 2023; Ni et al. 2023), where the model generates N candidates and uses a reward model to select

the most plausible answer. The final accuracy of this Best-of- N strategy is constrained by the diversity of the generated candidates. A straightforward method to encourage exploration is to use a stochastic sampling decoder when generating the next token, leading to a zoo of emerging methods (Fan, Lewis, and Dauphin 2018; Holtzman et al. 2020; Meister et al. 2023). These methods construct a token-level search space with varying access to the internals of an LM, such as logits, next-token distributions, or probability scores.

Yet, stochastic decoding methods frequently suffer from *diversity collapse*, where the outputs of LMs may converge to nearly identical reasoning paths (Yun et al. 2025; Dang et al. 2025). This phenomenon has sparked more aggressive search strategies, such as contrastive search (Su and Collier 2023), balancing the model confidence with the degeneration penalty to avoid repetitiveness. Alternatively, Vijayakumar et al. (2016) maintains multiple diverse hypotheses during beam search by diversity-promoting objectives. Additionally, self-speculative decoding (Zhang et al. 2024) and speculative sampling (Leviathan, Kalman, and Matias 2023) inherently promote diversity through their multi-step prediction mechanisms. Most of these search strategies are computationally intensive, as they require joint consideration of reasoning trajectories distribution (Nagarajan et al. 2025).

Another challenge to promoting reasoning diversity is measuring it. Popular metrics like lexical or semantic diversity may not capture the reasoning diversity well. Lexical diversity measures the number of different meanings or perspectives conveyed among sequences, but is sensitive to text length, rephrasing, or synonyms (Bestgen 2024). Similarly, semantic diversity quantifies the different meanings or perspectives, but it is sensitive to adding or removing details from the text (Han, Kim, and Chang 2022). Computationally, both often invoke extra neural architecture for evaluation, thus adding load to the inference process.

In this paper, we look at the reasoning diversity through another proxy: the diversity of the hidden activations of the generated sequences. Hidden activations are the internal representations computed by a language model at each layer and token position as it processes input. They encode intermediate computations and abstract concepts, acting as the latent “thinking space” where reasoning, memory, and structure are implicitly formed. While there may be strong correlations between the activations and the reasoning paths, there

^{*}These authors contributed equally.

is unfortunately no one-to-one equivalence between them. Thus, admittedly, promoting diversity of the hidden activations does not necessarily lead to diversity in the reasoning paths. However, recent progress suggests that encouraging diversity among neuron activations within the same layer increases the capacity of the model to learn a broader range of features (Laakom et al. 2023). In general, increasing the diversity of hidden activations can reduce estimation error and improve generalization. Moreover, models that facilitate diverse internal activations may be better equipped to represent and synthesize multiple reasoning strategies, as the richer internal space allows for more varied “thought paths” (Naik et al. 2023). Recent results in interpretability indicate that different features or activation clusters can sometimes correspond to different “reasoning circuits” or strategies, especially in LMs (Lindsey et al. 2025). These observations suggest that we should design fast and parameter-efficient mechanisms to promote the diversity of hidden activations during generation, hoping to induce reasoning diversity and improve the accuracy for Best-of- N sampling.

Contributions. We summarize our contributions as follows:

- We propose the **SP**herical intervention for **RE**asoning **D**iversity (**SPREAD**), an unsupervised activation steering method that improves the diversity among reasoning trajectories. At a synchronization anchor, SPREAD extracts the hidden activations from all sequences, then computes the steering vectors that maximize the total volume spanned by all possible subsets of the intervened activations. SPREAD then adds these steering vectors to the respective activations of all subsequent tokens until the next synchronization anchor.
- We show that determining the optimal steering vectors can be reformulated as a manifold optimization problem defined over the product of spheres, where the log-determinant objective function captures the geometric diversity of the intervened activations. We propose using a Riemannian block coordinate descent algorithm, which exploits the product structure of the manifold constraints. We also study the theoretical properties of the optimization problem and prove the convergence guarantee of the algorithm for appropriate step sizes.

Using the steering methods, SPREAD uses readily available hidden activations from the generation process and does not require any additional neural architectures to measure quality or reasoning diversity. Moreover, SPREAD could rely on only one hyperparameter that prescribes the relative radii of the intervention vectors, and it relieves the burden of parameter tuning at inference time.

Our paper unfolds as follows: Section 2 reviews the related works on generative diversity and activation steering, Section 3 presents the mathematical formulation of the SPREAD framework, Section 4 develops the manifold optimization algorithm for computing the optimal steering vectors, and Section 5 empirically illustrates the performance of SPREAD on mathematical reasoning tasks.

Notations. The space of p -dimensional vectors is denoted $\mathbb{R}^p$. For any $x \in \mathbb{R}^p$, $\|x\|_2$ is its Euclidean norm. For a matrix $A \in \mathbb{R}^{p \times N}$, we use $\|A\|_F$ for the Frobenius norm.

We use $\nabla \ell(V)$ and $\nabla^2 \ell(V)$ for the gradient of function ℓ and Hessian matrix of function ℓ with respect to V in the Euclidean sense; while $\text{grad } \ell$ and $\text{Hess } \ell$ are the Riemannian counterparts. We use $\nabla_i \ell(V)$, $\nabla_i^2 \ell(V)$, $\text{grad}_i \ell(V)$ and $\text{Hess}_i \ell(V)$ for the corresponding operator with respect to the i -th block v_i while fixing all other blocks. All proofs are relegated to the appendix.

2 Literature Review

Diversity in generation. Classical approaches, including temperature sampling (Ackley, Hinton, and Sejnowski 1985), top- k sampling (Fan, Lewis, and Dauphin 2018), nucleus sampling (Holtzman et al. 2020), and typical decoding (Meister et al. 2023), promote output diversity by introducing stochasticity into the generation process. Prompt-centric techniques have also been shown to enrich the diversity of reasoning. Li et al. (2023b) focuses on prompt diversity by generating diverse prompts to explore different reasoning paths, while filtering out incorrect answers by a weighted voting scheme. Naik et al. (2023) proposes a self-reflective prompting that leverages the LLM as a guide to design a diverse set of approaches for complex reasoning tasks. Wang et al. (2025) uses multiple adaptive steering vectors for different hallucination types, though maintaining multiple vectors is computationally expensive. Chung et al. (2025) achieves diversity through parameter-efficient prefix tuning, but effectiveness is sensitive to training data quality.

Activation Steering is a lightweight and interpretable method for controlling LMs. This approach injects direction vectors into the residual stream of transformer layers to steer generation toward desired attributes (e.g., truthfulness, sentiment, toxicity) without modifying model weights. *Contrastive steering methods* derive directions by comparing activations between positive and negative examples of desired behaviors. Turner et al. (2023) computes steering vectors by averaging residual stream differences between factual and hallucinatory responses, while Stolfo et al. (2025) contrasts activations with and without specific instructions. *Probe-based steering methods* instead learn to identify relevant concepts through trained classifiers, then extract steering directions from the learned representations (Li et al. 2023a; Zhang et al. 2025). Despite demonstrating effectiveness across domains like toxicity reduction (Zhang et al. 2025), and truthfulness enhancement (Turner et al. 2023), activation steering restricts its broader applicability. Most existing approaches rely on single, fixed direction vectors that constrain model outputs to narrow behavioral modes. The challenge becomes even more pronounced in mathematical reasoning, where defining clear positive and negative exemplars for contrastive learning is inherently difficult because mathematical correctness involves complex, context-dependent logical structures that resist simple binary classification. These domain-specific challenges explain why prior activation steering research has largely avoided mathematical reasoning applications.

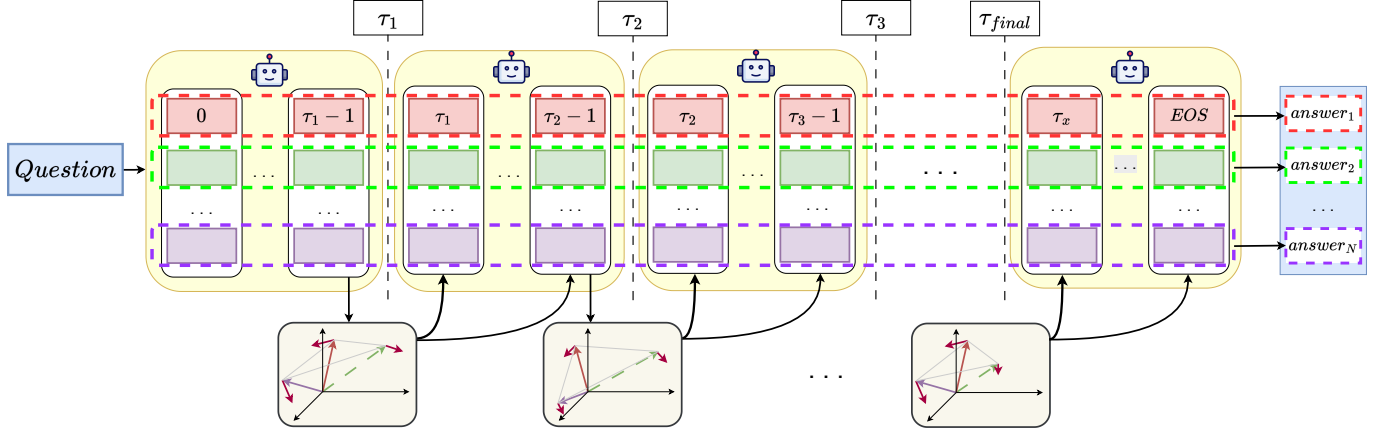


Figure 1: Overview of SPREAD for generating N diverse reasoning answers simultaneously. At each decoding step τ_t , we extract the hidden vectors corresponding to the last token in each path. These hidden vectors serve as inputs to Algorithm 1, where they are projected into a shared activation space to compute N steering vectors. This process is repeated until an end-of-sequence EOS token is generated.

3 Activation Steering for Diverse Generation under the Optimization Lens

We aim to elicit diverse reasoning paths from an LM at test time without any fine-tuning by sampling multiple sequences per input prompt. Figure 1 shows the core idea of the SPREAD framework, which intervenes directly in the model’s activation space during the autoregressive generation process. Specifically, we simultaneously generate N output sequences and extract the final hidden state vectors after τ tokens, $H = [h_1, \dots, h_N] \in \mathbb{R}^{p \times N}$, for all sequences. We then compute an additive steering vector v_i for each hidden state h_i , yielding a new set of hidden states $H_{new} = H + V$, where $V = [v_1, \dots, v_N]$. The core of our approach is to maximize a geometric measure of diversity for H_{new} (2), or equivalently with respect to the steering vectors V . A natural measure for the dispersion of a set of vectors is the volume of the parallelepiped they span.

Definition 1 (Parallelepiped). Given n vectors $h_1, \dots, h_n \in \mathbb{R}^p$, the parallelepiped is their convex hull:

$$\mathcal{P}(\{h_1, \dots, h_n\}) \triangleq \left\{ h \in \mathbb{R}^p : h = \sum_{i=1}^n \lambda_i h_i, \lambda \in [0, 1]^n \right\}.$$

To encourage robust diversity, we require that not only the entire set of N steered vectors be diverse, but that *any* subset of these vectors also be diverse. This prevents degenerate solutions where, for instance, $N - 1$ vectors are clustered together and only one is pushed far away. Let $\mathbb{I} \in 2^{[N]}$ be any index subset of the N sequences. We aim to maximize the sum of the squared volumes of the parallelepiped associated with all possible subsets:

$$\max_{\forall i: \|v_i\|_2^2 \leq \alpha_i} \sum_{\mathbb{I} \in 2^{[N]}} \text{Volume}(\mathcal{P}(\{h_i + v_i\}_{i \in \mathbb{I}}))^2, \quad (1)$$

where α_i is a magnitude for the intervention vector v_i on the i -th path. The squared norm constraints effectively keep the modified hidden state $h_i + v_i$ within a “trust region”

of moderate radius $\sqrt{\alpha_i}$ around the original state h_i . If α_i is too big, the vector v_i could erase meaningful information stored in h_i , leading to generation collapse. The next proposition gives an explicit form of the objective function in Problem (1) as a log-determinant function.

Proposition 2 (Objective function equivalence). *Problem (1) is equivalent to the following log-determinant optimization problem*

$$\min_{\substack{\forall i: \|v_i\|_2^2 \leq \alpha_i \\ V = [v_1, \dots, v_N]}} \ell(V) \triangleq -\log \det[I + (H + V)^\top (H + V)]. \quad (2)$$

Problem (2) has a convex feasible set, but its objective function ℓ is non-convex, as illustrated in the next example.

Example 3 (Non-convexity of ℓ). *Take*

$$H = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, V_1 = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}, V_2 = \begin{bmatrix} -2 & 0 \\ 0 & -2 \end{bmatrix},$$

and $V_3 = \frac{1}{2}(V_1 + V_2)$. Then $\ell(V_1) = \ell(V_2) = -\log 4$ and $\ell(V_3) = 0$. We find $2\ell(V_3) - (\ell(V_1) + \ell(V_2)) = 2\log 4 > 0$, which implies that ℓ is not convex.

Problem (2) turns out to be a non-convex problem, and it is, in general, NP-hard to find its global optimum (Jin et al. 2021). However, we can show a qualitative result asserting that the optimal steering vectors of Problem (1) will make the norm constraint $\|v_i\|_2^2 \leq \alpha_i$ binding, and we obtain another equivalent problem with equality constraints.

Proposition 4 (Constraint equivalence). *Problem (1) is further equivalent to the following log-determinant optimization problem with equality constraints*

$$\min \{ \ell(V) : \|v_i\|_2^2 = \alpha_i, V = [v_1, \dots, v_N] \}. \quad (3)$$

The equality constraints $\|v_i\|_2^2 = \alpha_i$ are no longer convex. However, the advantage of the reformulation (3) is that we can cast Problem (3) as an optimization problem over a Riemannian manifold. In the next section, we will describe

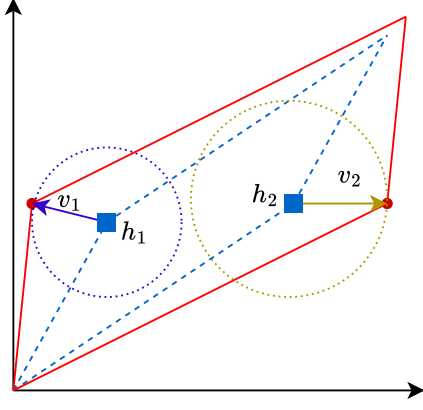


Figure 2: An illustration of the volume maximization intuition behind SPREAD. The hidden vectors h_1 and h_2 (originally blue squares) are pushed toward target positions using the corresponding steering vectors v_1, v_2 , found via Riemannian Block Coordinate Descent (Algorithm 1). After intervention, the new parallelepiped (red) has a larger volume than the original parallelepiped (dashed blue).

the procedure for solving Problem (3) with manifold optimization methods.

Hyperparameters. Problem (1) and its equivalent form (3) requires N radii values $\{\alpha_i\}_{i=1}^N$ as input hyperparameters. We propose to set $\alpha_i = C\|h_i\|_2/p$ for all i , where p is the dimension of h_i , and thus the number of hyperparameters is reduced to only one relative parameter $C > 0$.

4 Manifold Optimization for Steering

This section is to devise an efficient algorithm for solving Problem (3) based on Riemannian optimization. Formally, we let $\mathcal{M}_i$ be the sphere in $\mathbb{R}^p$ of radius $\sqrt{\alpha_i}$:

$$\mathcal{M}_i = \{v_i \in \mathbb{R}^p : \|v_i\|_2^2 = \alpha_i\},$$

and define the product manifold $\mathcal{M} = \mathcal{M}_1 \times \dots \times \mathcal{M}_N$. Problem (3) can be cast as a Riemannian optimization problem over the product manifold $\mathcal{M}$, which is naturally solved using Riemannian optimization algorithms.

4.1 Preliminaries about the Manifold $\mathcal{M}$

Riemannian optimization exploits the geometric structure of the constraint set to perform updates along curved spaces (manifolds) rather than flat Euclidean space. Key tools include (i) the Riemannian gradient, which generalizes classical gradient methods to manifold settings, and (ii) the exponential map, which generalizes the concept of moving in a straight line along a gradient in Euclidean space to moving along a shortest path on a manifold.

At a point on a manifold, a tangent space is a vector space that “touches” the manifold at that point and contains all possible directions in which one can move along the manifold from that point. For an individual sphere, the tangent space at $v_i \in \mathcal{M}_i$ is

$$T_{v_i}\mathcal{M}_i = \{z \in \mathbb{R}^p : z^\top v_i = 0\}.$$

For the product manifold, the tangent space at $V = [v_1, \dots, v_N]$ is

$$T_V\mathcal{M} = \{Z = (z_1, \dots, z_N) : z_i \in T_{v_i}\mathcal{M}_i \forall i\}.$$

The Riemannian gradient of ℓ defined on the manifold $\mathcal{M}$ is the steepest ascent direction that lies within the tangent space of the manifold at a given point. It is computed by first taking the Euclidean gradient of ℓ in the ambient space $\mathbb{R}^{p \times N}$, and then projecting this gradient onto the tangent space of the manifold. Let $G \in \mathbb{R}^{p \times N}$ be the computed Euclidean gradient of ℓ at V , the projection of G onto $T_V\mathcal{M}$ is decomposable into N projections of the columns g_i onto the tangent space of the individual sphere $\mathcal{M}_i$

$$\text{Proj}_{T_{v_i}\mathcal{M}_i}(g_i) = (I - \frac{1}{\alpha_i}v_iv_i^\top)g_i \quad \forall i.$$

The next lemma formalizes the computation of the Riemannian gradient for ℓ .

Lemma 5. For any $V = [v_1, \dots, v_N] \in \mathcal{M}$, the Riemannian gradient of ℓ in the i -th block is given by

$$\text{grad}_i\ell(V) = g_i - \frac{1}{\alpha_i}g_i^\top v_iv_i \in \mathbb{R}^p, \quad (4)$$

where $g_i = -2((H + V)M^{-1})_i \in \mathbb{R}^p$ and $M = I + (H + V)^\top(H + V)$. Furthermore, the Riemannian gradient of ℓ at point V corresponding to the product manifold $\mathcal{M}$ is given by $\text{grad}\ell(V) = [\text{grad}_1\ell(V), \dots, \text{grad}_N\ell(V)]$.

The negative Riemannian gradient at V gives the direction of steepest descent within the tangent space. The exponential map then moves the incumbent solution along this direction, not in a straight line, but along a curved path that fits the spherical surface. This allows us to take meaningful steps while staying on the manifold throughout the iterations. Given the descent direction $d_i = -\text{grad}_i\ell(V)$ and a step size η_i for the i -th block, the exponential map gives (Boumal 2023, Section 10.2)

$$\begin{aligned} \text{Exp}_{v_i}(\eta_id_i) = \\ \cos(\frac{\eta_i\|d_i\|_2}{\sqrt{\alpha_i}})v_i + \sin(\frac{\eta_i\|d_i\|_2}{\sqrt{\alpha_i}})\frac{d_i}{\|d_i\|_2}\sqrt{\alpha_i} \in \mathcal{M}_i. \end{aligned} \quad (5)$$

Given $\text{Exp}_{v_i}(\eta_id_i)$ in (5), the exponential map on the product manifold $\mathcal{M}$ is given by

$$\begin{aligned} \text{Exp}_V([\eta_1d_1, \dots, \eta_Nd_N]) = \\ [\text{Exp}_{v_1}(\eta_1d_1), \dots, \text{Exp}_{v_N}(\eta_Nd_N)] \in \mathcal{M}. \end{aligned} \quad (6)$$

4.2 Block-Coordinate Riemannian Descent

Since our manifold $\mathcal{M}$ decomposes as a product of scaled spheres, its natural block structure allows one global update to be replaced by N cheaper spherical updates. Consequently, Riemannian Block-coordinate Descent (Gutman and Ho-Nguyen 2023) provides an intuitive and efficient method for Problem (3).

In the outer iteration k , Algorithm 1 updates the blocks v_i sequentially for $i = 1, \dots, N$. Fixing all other blocks at their most recent values, it first computes the Euclidean gradient g_i . In Line 5, g_i is then projected onto the tangent space to

Require: Hidden activation vectors $\{h_i\}_{i=1}^N$, magnitudes $\{\alpha_i\}_{i=1}^N$, learning rate $\eta_i > 0$, max iterations K

```

1: Initialize  $v_i^{(0)}$  using (7)
2: for  $k = 1$  to  $K$  do
3:   for  $i = 1$  to  $N$  do
4:     Compute Euclidean gradient in the variable  $v_i$  by  $g_i \leftarrow \nabla_i \ell(v_1^{(k)}, \dots, v_i^{(k-1)}, \dots, v_N^{(k-1)})$ 
5:     Compute the descent direction  $d_i \leftarrow -g_i + \frac{1}{\alpha_i} ((v_i^{(k-1)})^\top g_i) v_i^{(k-1)}$ 
6:     Compute  $v_i^{(k)} \leftarrow \cos(\frac{\eta_i \|d_i\|_2}{\sqrt{\alpha_i}}) v_i^{(k-1)} + \sin(\frac{\eta_i \|d_i\|_2}{\sqrt{\alpha_i}}) \frac{d_i}{\|d_i\|_2} \sqrt{\alpha_i}$ 
7:   end for
8: end for
9: return  $V^{(K)} = [v_1^{(K)}, \dots, v_N^{(K)}]$ 

```

yield the Riemannian descent direction d_i , and then Line 6 moves $v_i^{(k-1)}$ along the geodesic on $\mathcal{M}_i$ via the exponential map (5), thus preserving its feasibility on the sphere.

Initialization. To get a feasible initial point, for each i , we can initialize $v_i^{(0)}$ by

$$v_i^{(0)} = \sqrt{\alpha_i} \frac{h_i + \varepsilon_i - \bar{h}}{\|h_i + \varepsilon_i - \bar{h}\|_2} \quad \forall i. \quad (7)$$

where $\varepsilon_i \sim \mathcal{N}(0, \sigma^2 I_d)$ are independent Gaussian noise with small variances and $\bar{h} = \frac{1}{N} \sum_{i=1}^N h_i$ is the centroid of the hidden vectors $\{h_i\}_{i=1}^N$.

4.3 Convergence Analysis

We now study the convergence guarantee of Algorithm 1. The next theorem asserts the convergence of Algorithm 1 for solving Problem (3) with well-chosen step sizes. To this end, let $\bar{\alpha} \triangleq \sum_{i=1}^N \alpha_i$ be the total radii, and for each sequence i , define the quantity

$$L_i \triangleq 2 + 4(\|H\|_F + \sqrt{\bar{\alpha}})^2 + \frac{2}{\sqrt{\alpha_i}}(\|H\|_F + \sqrt{\bar{\alpha}}) \quad \forall i. \quad (8)$$

Theorem 6 (Convergence of Algorithm 1). *Let $V^{(k)} = [v_1^{(k)}, \dots, v_N^{(k)}]$ be the sequence generated from Algorithm 1 with learning rate $\eta_i = 1/L_i$ for $i = 1, \dots, N$. Define*

$$C \triangleq \frac{L_{\min}^2}{4L_{\max}(L_{\min}^2 + L^2 N(N-1))},$$

where $L_{\min} = \min_i L_i$ and $L_{\max} = \max_i L_i$. Then, the following hold:

- We have $\lim_{k \rightarrow \infty} \|\text{grad } \ell(V^{(k)})\|_F = 0$ and

$$\min_{s \leq k} \|\text{grad } \ell(V^{(s)})\|_F \leq \sqrt{\frac{\ell(V^{(0)}) - \ell^*}{Ck}},$$

where ℓ^* is the optimal value of Problem (3).

- Any limit point $V^\infty \in \mathcal{M}$ of $\{V^{(k)}\}_{k \geq 1}$ is a stationary point, i.e., $\text{grad } \ell(V^\infty) = 0$.

Theorem 6 establishes both asymptotic and non-asymptotic convergence results for Algorithm 1. It guarantees that any limit point of the generated sequence is a stationary point of the objective function. The sublinear rate

$O(1/\sqrt{k})$ for the minimum gradient norm bounds the number of iterations needed to achieve a desired accuracy level.

To analyze the convergence of the algorithm, we study the smoothness of the objective function. Firstly, we recite the L -smoothness definition in the Riemannian sense (Gutman and Ho-Nguyen 2023, equation (3)).

Definition 7 (L -smoothness). The function $\ell : \mathcal{M} \rightarrow \mathbb{R}$ is L -smooth if for all $V \in \mathcal{M}$, $U \in T_V \mathcal{M}$ and $Z = \text{Exp}_V(U)$ it holds that

$$\|\Gamma_{V \rightarrow Z}^U \text{grad } \ell(V) - \text{grad } \ell(Z)\|_F \leq L\|U\|_F,$$

where $\Gamma_{V \rightarrow Z}^U : T_V \mathcal{M} \rightarrow T_Z \mathcal{M}$ is the parallel transport operator along the curve $\gamma(t) = \text{Exp}_V(tU)$.

The following lemma shows that the objective function $\ell(V)$ of Problem (3) satisfies Definition 7 with an explicit Lipschitz constant L .

Proposition 8 (Smoothness). *Let $\alpha_{\min} \triangleq \min_i \alpha_i > 0$. The objective function $\ell(V)$ is L -smooth, with constant*

$$L = 2 + 4(\|H\|_F + \sqrt{\bar{\alpha}})^2 + \frac{2}{\sqrt{\alpha_{\min}}}(\|H\|_F + \sqrt{\bar{\alpha}}). \quad (9)$$

We further show that the function ℓ also satisfies the block smoothness, which is defined by restricting Definition 7 to each individual sphere $\mathcal{M}_i$.

Proposition 9 (Block smoothness). *Let $V = [v_1, \dots, v_N] \in \mathcal{M}$, $U = [u_1, \dots, u_N] \in T_V \mathcal{M}$ and $Z = \text{Exp}_V(U)$. For any v_i and $u_i \in T_{v_i} \mathcal{M}_i$, it holds that*

$$\|\Gamma_{v_i \rightarrow z_i}^{u_i} \text{grad}_i \ell(V) - \text{grad}_i \ell(Z)\|_F \leq L_i \|u_i\|_2 \quad (10)$$

with L_i defined in (8). Here $\Gamma_{v_i \rightarrow z_i}^{u_i} : T_{v_i} \mathcal{M}_i \rightarrow T_{z_i} \mathcal{M}_i$ is the parallel transport operator along the curve $\gamma(t) = \text{Exp}_{v_i}(tu_i)$.

Proposition 9 is crucial both for establishing the convergence of Algorithm 1 and for guiding the selection of its step-sizes. Equation (10) together with Boumal (2023, Corollary 10.54) implies that for each i it holds

$$\ell(v_1, \dots, \text{Exp}_{v_i}(u_i), \dots, v_i) \leq \quad (11)$$

$$\ell(V) + \langle \text{grad}_i \ell(V), u_i \rangle + \frac{L_i}{2} \|u_i\|_2^2.$$

Thus, in particular, if one sets $\eta_i = 1/L_i$ then each coordinate-descent update yields a provable decrease in the objective (Gutman and Ho-Nguyen 2023, Lemma 1). This obviates the need for extensive tuning of learning rates η_i .

Model	Temp.	AIME24			MATH500			OlympiadBench		
		SPREAD		Sampling	SPREAD		Sampling	SPREAD		Sampling
		$C = 1$	$C = 10$		$C = 1$	$C = 10$		$C = 1$	$C = 10$	
Qwen2.5-1.5B	1.0	3.3	6.7	0.0	43.2	43.4	42.8	21.5	19.7	19.0
	0.8	6.7	6.7	3.3	51.8	54.2	51.4	25.3	27.0	25.9
	0.6	10.0	3.3	3.3	55.0	54.0	53.4	28.4	28.3	29.3
	0.4	6.7	6.7	6.7	56.2	57.6	55.8	30.8	30.1	30.7
	0.2	10.0	6.7	6.7	55.6	53.8	52.2	28.0	27.1	26.4
Qwen2.5-Math-1.5B-Instruct	1.0	26.7	16.7	20.0	83.8	83.6	84.6	47.6	49.5	48.3
	0.8	26.7	23.3	16.7	85.2	85.0	83.6	49.9	49.2	51.0
	0.6	26.7	26.7	20.0	85.4	84.4	84.6	50.4	51.0	50.8
	0.4	23.3	23.3	23.3	84.6	84.2	84.0	51.7	48.8	51.0
	0.2	20.0	26.7	16.7	82.4	84.4	84.4	52.3	51.1	49.9

Table 1: Pass@ $N \uparrow$ (in %) performance comparison across model variants and sampling temperature on three mathematical reasoning benchmarks. Bold indicates the best method in each dataset.

5 Numerical Experiment

We evaluate SPREAD across three established mathematical reasoning benchmarks: AIME24 (30 problems), MATH500 (500 problems), and OlympiadBench (675 problems). We evaluate using the following metrics:

- *Pass@ N* : The proportion of problems for which at least one of the N generated solutions produces the correct final answer.
- *Solution Diversity*: The float score indicating the overall diversity among N solution
- *Unique Solution Count*: The number of distinct solution approaches among N solutions.

For the latter two metrics, we employ a language-model-as-a-judge paradigm using GPT-4.1-mini. We prompt the model with the question and N generated solutions, requesting a diversity score (float in $[0, 1]$) and a count of unique approaches (integer). The specific prompts and parameters for this evaluation are detailed in the Appendix. Our experiments utilize two model variants: Qwen2.5-1.5B (base) and Qwen2.5-Math-1.5B-Instruct (math-specialized). Steering vectors are applied at layer 28, which corresponds to the final layer in both architectures. The hidden activations have dimension $p = 1536$. The magnitude α_i is chosen as $\alpha_i = C \|h_i\|_2 / p$, where $C \in \{1, 10\}$ is a scaling constant. Steering vectors are computed using Algorithm 1 at token positions $\tau \in \{100, 600, 1100, 1600\}$ with $K = 20$, and the learning rates are calibrated using the input activations following Theorem 6. All experiments are conducted on NVIDIA RTX A5000 24GB hardware using temperature sampling for solution generation with temperature $\in \{0.2, 0.4, 0.6, 0.8, 1.0\}$ and a maximum generation length of 2048 tokens. To ensure reproducibility, we fix the random seed to 42 for all experimental runs. Additional numerical results are presented in the appendix.

5.1 Experiment Results

Table 1 summarizes the performance of SPREAD compared to temperature sampling across various benchmarks and temperature settings. SPREAD with either $C = 1$ or $C = 10$ consistently performs at least as well as vanilla

temperature sampling, and often achieves improvements of several percentage points. Overall, SPREAD decoding (especially with $C = 1$) offers the strongest results under most conditions. Table 2 demonstrates that SPREAD consistently outperforms temperature sampling in generating unique solutions across multiple benchmarks. These results highlight SPREAD’s robustness and effectiveness in promoting correct and diverse reasoning trajectories across various model variants and datasets.

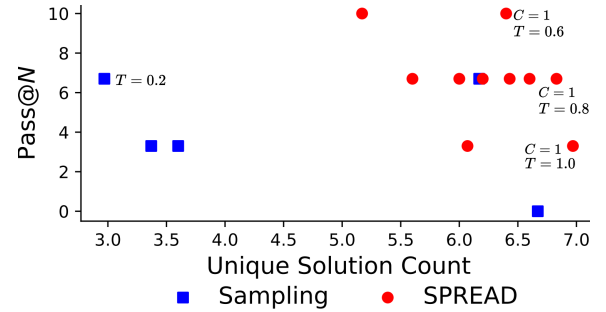


Figure 3: Comparison of SPREAD and sampling methods on AIME24 dataset using Qwen2.5-1.5B model, showing Pass@ N and Unique Solution Count.

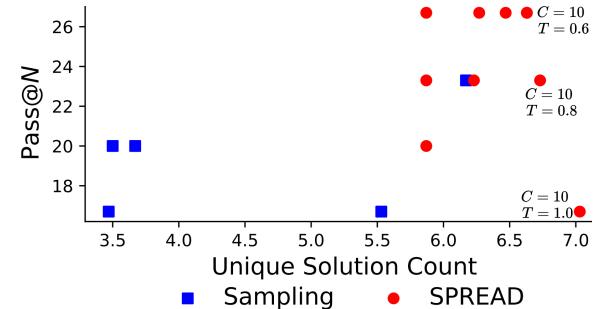


Figure 4: Comparison of SPREAD and sampling methods on AIME24 dataset using Qwen2.5-Math-1.5B-Instruct model, showing Pass@ N and Unique Solution Count.

Metric	Model	Temp.	AIME24			MATH500			OlympiadBench		
			SPREAD		Sampling	SPREAD		Sampling	SPREAD		Sampling
			$C = 1$	$C = 10$		$C = 1$	$C = 10$		$C = 1$	$C = 10$	
Unique Solution Count $\uparrow$	Qwen2.5 -1.5B	1.0	6.97	6.60	6.67	3.14	3.21	3.14	6.12	6.14	6.14
		0.8	6.83	6.43	3.6	3.03	3.05	2.97	5.99	5.99	6.04
		0.6	6.40	6.07	3.37	2.91	2.86	2.90	5.57	5.65	5.58
		0.4	6.00	6.20	6.17	2.73	2.74	2.77	5.30	5.30	5.26
		0.2	5.17	5.60	2.97	2.47	2.56	2.59	4.65	4.69	4.69
	Qwen2.5 -Math -1.5B -Instruct	1.0	6.63	7.03	3.67	1.92	1.94	1.93	4.59	4.53	4.57
		0.8	6.63	6.73	3.47	1.87	1.90	1.89	4.28	4.36	4.39
		0.6	6.27	6.47	3.5	1.81	1.84	1.82	4.10	4.16	4.11
		0.4	6.23	5.87	6.17	1.77	1.73	1.78	3.99	3.92	3.89
		0.2	5.87	5.87	5.53	1.70	1.74	1.74	3.67	3.65	3.68
Diversity Score $\uparrow$	Qwen2.5 -1.5B	1.0	0.37	0.33	0.40	0.40	0.39	0.39	0.38	0.40	0.39
		0.8	0.51	0.45	0.37	0.42	0.44	0.41	0.46	0.46	0.46
		0.6	0.52	0.46	0.39	0.41	0.40	0.40	0.46	0.46	0.46
		0.4	0.41	0.47	0.42	0.36	0.36	0.36	0.42	0.43	0.42
		0.2	0.31	0.34	0.33	0.29	0.29	0.28	0.33	0.35	0.35
	Qwen2.5 -Math -1.5B -Instruct	1.0	0.64	0.65	0.63	0.25	0.25	0.25	0.39	0.39	0.38
		0.8	0.60	0.62	0.57	0.23	0.23	0.23	0.36	0.38	0.37
		0.6	0.57	0.55	0.55	0.21	0.22	0.20	0.34	0.34	0.34
		0.4	0.51	0.51	0.51	0.19	0.18	0.18	0.31	0.30	0.31
		0.2	0.46	0.52	0.41	0.15	0.16	0.15	0.26	0.26	0.26

Table 2: Unique Solution Count $\uparrow$ and Diversity Score $\uparrow$ comparison across model variants and sampling temperature on three mathematical reasoning benchmarks. Bold indicates the best method in each dataset.

To strengthen the comparison, we aggregate the metrics and analyze the accuracy-diversity frontiers in Figure 3 and 4 for the AIME24 dataset. We can observe that the performance of SPREAD (red circles) dominates that of vanilla sampling methods with different temperatures (blue squares). The Pareto plots for the MATH and OlympiadBench datasets are presented in the appendix.

5.2 Computational Efficiency

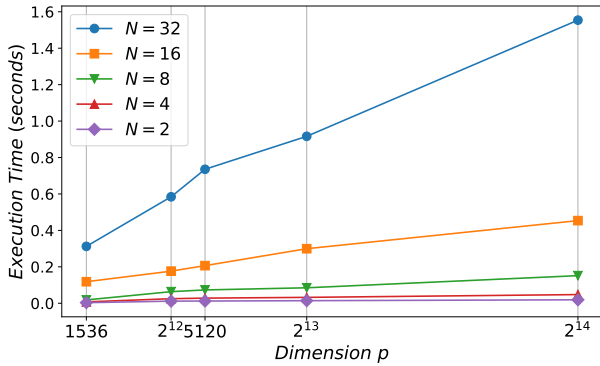


Figure 5: Average running time of Algorithm 1 with varying problem dimension p and number of steering vectors N . Lower execution time is better.

We analyze the execution time of Algorithm 1 for different hidden vector dimensions p and different numbers of sequences N . We generate synthetic activations with p varying from 1536 up to $2^{14} = 16384$, which corresponds to

the hidden size of large-scale models such as LLaMA-3.1-405B. Figure 5 reports the average solution time of Algorithm 1 with $K = 20$ over 30 independent runs. The execution time for $N = 32$ remains under 1.8 seconds even for the largest dimension 16384, showing that our SPREAD method remains highly efficient. These results demonstrate the scalability and practical efficiency of our algorithm for inference-time language model intervention.

6 Conclusions

We introduced SPREAD, a novel unsupervised activation steering framework at test time designed to improve the reasoning diversity of language models. By addressing the well-known challenge of low output variance in best-of- N sampling, SPREAD intervenes meaningful diversity into generation trajectories. Our key innovation lies in reformulating the activation steering problem as a Riemannian optimization over the product of spheres, maximizing the log-determinant volume spanned by intervened representations. This principled formulation enables us to efficiently compute diverse steering directions using Riemannian block-coordinate descent. Empirical evaluations on mathematical reasoning benchmarks such as AIME24, MATH, and OlympiadBench validate the effectiveness of our approach, as SPREAD significantly improves both the diversity and accuracy of generated solutions, outperforming temperature-sampling techniques. These results demonstrate the potential of geometric test-time intervention methods to enhance reasoning in language models.

References

- Ackerman, C. M. 2024. Representation tuning. arXiv:2409.06927.
- Ackley, D. H.; Hinton, G. E.; and Sejnowski, T. J. 1985. A learning algorithm for Boltzmann machines. *Cognitive science*, 9(1): 147–169.
- AI-MO. 2024. AIME24. <https://huggingface.co/datasets/AI-MO/aime-validation-aime>.
- Bestgen, Y. 2024. Measuring lexical diversity in texts: The twofold length problem. *Language Learning*, 74(3): 638–671.
- Boumal, N. 2023. *An Introduction to Optimization on Smooth Manifolds*. Cambridge University Press.
- Chen, M.; Tworek, J.; Jun, H.; Yuan, Q.; de Oliveira Pinto, H. P.; Kaplan, J.; Edwards, H.; Burda, Y.; Joseph, N.; Brockman, G.; Ray, A.; Puri, R.; Krueger, G.; Petrov, M.; Khlaaf, H.; Sastry, G.; Mishkin, P.; Chan, B.; Gray, S.; Ryder, N.; Pavlov, M.; Power, A.; Kaiser, L.; Bavarian, M.; Winter, C.; Tillet, P.; Such, F. P.; Cummings, D.; Plappert, M.; Chantzis, F.; Barnes, E.; Herbert-Voss, A.; Guss, W. H.; Nichol, A.; Paino, A.; Tezak, N.; Tang, J.; Babuschkin, I.; Balaji, S.; Jain, S.; Saunders, W.; Hesse, C.; Carr, A. N.; Leike, J.; Achiam, J.; Misra, V.; Morikawa, E.; Radford, A.; Knight, M.; Brundage, M.; Murati, M.; Mayer, K.; Welinder, P.; McGrew, B.; Amodei, D.; McCandlish, S.; Sutskever, I.; and Zaremba, W. 2021. Evaluating large language models trained on code. arXiv:2107.03374.
- Chung, H.-L.; Hsiao, T.-Y.; Huang, H.-Y.; Cho, C.; Lin, J.-R.; Ziwei, Z.; and Chen, Y.-N. 2025. Revisiting test-time scaling: A Survey and a diversity-aware method for efficient reasoning. arXiv:2506.04611.
- Dang, X.; Baek, C.; Wen, K.; Kolter, J. Z.; and Raghunathan, A. 2025. Weight ensembling improves reasoning in language models. In *Second Conference on Language Modeling*.
- Fan, A.; Lewis, M.; and Dauphin, Y. 2018. Hierarchical neural story generation. In Gurevych, I.; and Miyao, Y., eds., *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 889–898. Melbourne, Australia: Association for Computational Linguistics.
- Gutman, D. H.; and Ho-Nguyen, N. 2023. Coordinate descent without coordinates: Tangent subspace descent on Riemannian manifolds. *Mathematics of Operations Research*, 48(1): 127–159.
- Han, S.; Kim, B.; and Chang, B. 2022. Measuring and improving semantic diversity of dialogue generation. In Goldberg, Y.; Kozareva, Z.; and Zhang, Y., eds., *Findings of the Association for Computational Linguistics: EMNLP 2022*, 934–950. Association for Computational Linguistics.
- He, C.; Luo, R.; Bai, Y.; Hu, S.; Thai, Z. L.; Shen, J.; Hu, J.; Han, X.; Huang, Y.; Zhang, Y.; Liu, J.; Qi, L.; Liu, Z.; and Sun, M. 2024. OlympiadBench: A challenging benchmark for promoting AGI with Olympiad-Level bilingual multimodal scientific problems. arXiv:2402.14008.
- Holtzman, A.; Buys, J.; Du, L.; Forbes, M.; and Choi, Y. 2020. The curious case of neural text degeneration. In *International Conference on Learning Representations*.
- Jin, C.; Netrapalli, P.; Ge, R.; Kakade, S. M.; and Jordan, M. I. 2021. On nonconvex optimization for machine learning: Gradients, stochasticity, and saddle points. *Journal of the ACM (JACM)*, 68(2): 1–29.
- Kulesza, A.; and Taskar, B. 2012. Determinantal point processes for machine learning. *Foundations and Trends® in Machine Learning*, 5(2–3): 123–286.
- Laakom, F.; Raitoharju, J.; Iosifidis, A.; and Gabbouj, M. 2023. WLD-Reg: A data-dependent within-layer diversity regularizer. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(7): 8421–8429.
- Leviathan, Y.; Kalman, M.; and Matias, Y. 2023. Fast inference from transformers via speculative decoding. In *Proceedings of the 40th International Conference on Machine Learning*.
- Lewkowycz, A.; Andreassen, A.; Dohan, D.; Dyer, E.; Michalewski, H.; Ramasesh, V.; Slone, A.; Anil, C.; Schlag, I.; Gutman-Solo, T.; Wu, Y.; Neyshabur, B.; Gur-Ari, G.; and Misra, V. 2022. Solving quantitative reasoning problems with language models. In *Proceedings of the 36th International Conference on Neural Information Processing Systems*.
- Li, K.; Patel, O.; Viégas, F.; Pfister, H.; and Wattenberg, M. 2023a. Inference-time intervention: Eliciting truthful answers from a language model. In Oh, A.; Naumann, T.; Globerson, A.; Saenko, K.; Hardt, M.; and Levine, S., eds., *Advances in Neural Information Processing Systems*, volume 36, 41451–41530.
- Li, Y.; Lin, Z.; Zhang, S.; Fu, Q.; Chen, B.; Lou, J.-G.; and Chen, W. 2023b. Making language models better reasoners with step-aware verifier. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 5315–5333.
- Lightman, H.; Kosaraju, V.; Burda, Y.; Edwards, H.; Baker, B.; Lee, T.; Leike, J.; Schulman, J.; Sutskever, I.; and Cobbe, K. 2023. Let’s verify step by step. In *The Twelfth International Conference on Learning Representations*.
- Lindsey, J.; Gurnee, W.; Ameisen, E.; Chen, B.; Pearce, A.; Turner, N. L.; Citro, C.; Abrahams, D.; Carter, S.; Hosmer, B.; Marcus, J.; Sklar, M.; Templeton, A.; Bricken, T.; McDougall, C.; Cunningham, H.; Henighan, T.; Jermyn, A.; Jones, A.; Persic, A.; Qi, Z.; Thompson, T. B.; Zimmerman, S.; Rivoire, K.; Conerly, T.; Olah, C.; and Batson, J. 2025. On the Biology of a Large Language Model. <https://transformer-circuits.pub/2025/attribution-graphs/biology.html>. Accessed: 2025-08-02.
- Meister, C.; Pimentel, T.; Wiher, G.; and Cotterell, R. 2023. Locally typical sampling. *Transactions of the Association for Computational Linguistics*, 11: 102–121.
- Nagarajan, V.; Wu, C. H.; Ding, C.; and Raghunathan, A. 2025. Roll the dice & look before you leap: Going beyond the creative limits of next-token prediction. In *Forty-second International Conference on Machine Learning*.

- Naik, R.; Chandrasekaran, V.; Yüksekönül, M.; Palangi, H.; and Nushi, B. 2023. Diversity of thought improves reasoning abilities of large language models. *CoRR*, abs/2310.07088.
- Ni, A.; Iyer, S.; Radev, D.; Stoyanov, V.; Yih, W.-t.; Wang, S.; and Lin, X. V. 2023. Lever: Learning to verify language-to-code generation with execution. In *International Conference on Machine Learning*, 26106–26128. PMLR.
- Petersen, K. B.; Pedersen, M. S.; et al. 2008. The Matrix Cookbook. *Technical University of Denmark*, 7(15): 510.
- Pyle, H. R. 1962. Non-square determinants and multilinear vectors. *Mathematics Magazine*, 35(2): 65–69.
- Stolfo, A.; Balachandran, V.; Yousefi, S.; Horvitz, E.; and Nushi, B. 2025. Improving instruction-following in language models through activation steering. In *The Thirteenth International Conference on Learning Representations*.
- Su, Y.; and Collier, N. 2023. Contrastive search is what you need for neural text Generation. *Transactions on Machine Learning Research*.
- Turner, A. M.; Thiergart, L.; Leech, G.; Udell, D.; Vazquez, J. J.; Mini, U.; and MacDiarmid, M. 2023. Steering language models with activation engineering. *arXiv preprint arXiv:2308.10248*.
- Turner, A. M.; Thiergart, L.; Leech, G.; Udell, D.; Vazquez, J. J.; Mini, U.; and MacDiarmid, M. 2024. Steering language models with activation engineering. *arXiv:2308.10248*.
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A. N.; Kaiser, Ł.; and Polosukhin, I. 2017. Attention is all you need. *Advances in neural information processing systems*, 30.
- Vijayakumar, A. K.; Cogswell, M.; Selvaraju, R. R.; Sun, Q.; Lee, S.; Crandall, D. J.; and Batra, D. 2016. Diverse beam search: decoding diverse solutions from neural sequence models. *CoRR*.
- Wang, T.; Jiao, X.; Zhu, Y.; Chen, Z.; He, Y.; Chu, X.; Gao, J.; Wang, Y.; and Ma, L. 2025. Adaptive activation steering: A tuning-free LLM truthfulness improvement method for diverse hallucinations categories. In *Proceedings of the ACM on Web Conference 2025*, 2562–2578. Association for Computing Machinery.
- Wei, J.; Wang, X.; Schuurmans, D.; Bosma, M.; brian ichter; Xia, F.; Chi, E. H.; Le, Q. V.; and Zhou, D. 2022. Chain of thought prompting elicits reasoning in large language models. In Oh, A. H.; Agarwal, A.; Belgrave, D.; and Cho, K., eds., *Advances in Neural Information Processing Systems*.
- Yun, L.; An, C.; Wang, Z.; Peng, L.; and Shang, J. 2025. The price of format: Diversity collapse in LLMs. *arXiv:2505.18949*.
- Zhang, H.; Wang, X.; Li, C.; Ao, X.; and He, Q. 2025. Controlling large language models through concept activation vectors. *Proceedings of the AAAI Conference on Artificial Intelligence*, 39(24): 25851–25859.
- Zhang, J.; Wang, J.; Li, H.; Shou, L.; Chen, K.; Chen, G.; and Mehrotra, S. 2024. Draft& Verify: Lossless large language model acceleration via self-speculative decoding. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 11263–11282. Association for Computational Linguistics.

This is the appendix for the paper “Test-time Diverse Reasoning by Riemannian Activation Steering”. Appendix A collects the proofs of all theoretical results in the main paper. Appendix B provides further details about the experimental settings and additional results.

A Proofs of Main Results

This section contains the proofs of all technical results presented in the main paper. For a matrix $A \in \mathbb{R}^{p \times N}$, $\|A\|_2$ denotes its spectral norm, i.e., the square root of the largest eigenvalue of the matrix $A^\top A$.

Proof of Proposition 2. The volume of the parallelepiped is computed from the determinant of the Gram matrix :

$$\text{Volume}(\mathcal{P}(\{h_i + v_i\}_{i \in \mathbb{I}}))^2 = \det((H + V)_{\mathbb{I}}^\top (H + V)_{\mathbb{I}})$$

where $(H + V)_{\mathbb{I}}$ is the submatrix of $(H + V)$ restricted to the columns indexed by $\mathbb{I}$ (Pyle 1962). By Kulesza and Taskar (2012, theorem 2.1), we find

$$\sum_{\mathbb{I} \subseteq [N]} \det((H + V)_{\mathbb{I}}^\top (H + V)_{\mathbb{I}}) = \det(I + (H + V)^\top (H + V))$$

Consequently, problem (1) is equivalent to

$$\max_{\forall i: \|v_i\|_2^2 \leq \alpha_i} \det[I + (H + V)^\top (H + V)]. \quad (12)$$

Since the logarithm function is a monotonically increasing function, we obtain the desired result. $\square$

Proof of Proposition 4. Introducing dual variables $\lambda_i \geq 0$ for each of the N constraints of problem (2), then the Lagrangian function for problem (2) is

$$\mathcal{L}(V, \lambda) = \ell(V) + \sum_i \lambda_i (\|v_i\|_2^2 - \alpha_i).$$

For each $i = 1, \dots, N$, we define the constraint function $c_i(V) = \|v_i\|_2^2 - \alpha_i$. We have the gradient

$$\nabla_i c_i(V) = (0, \dots, 0, \underbrace{2v_i}_{i\text{-th block}}, 0, \dots, 0). \quad (13)$$

Let $V^* = [v_1^*, \dots, v_N^*]$ be an optimal solution of problem (2). Consider the active set at V^*

$$\mathcal{A} = \{i : c_i(V^*) = \alpha_i\} = \{i : \|v_i^*\|_2^2 = \alpha_i\}.$$

Since $\alpha_i > 0$, from (13) we know that $\{\nabla c_i(V^*), i \in \mathcal{A}\}$ are linearly independent. Thus, the KKT point (V^*, λ^*) of problem (2) satisfies

$$\nabla_i \mathcal{L}(V^*, \lambda^*) = 0 \quad \forall i \quad (14a)$$

$$\|v_i^*\|_2^2 - \alpha_i \leq 0, \quad \lambda_i^* \geq 0 \quad \forall i \quad (14b)$$

$$\lambda_i^* (\|v_i^*\|_2^2 - \alpha_i) = 0 \quad \forall i. \quad (14c)$$

The gradient of $\ell(V)$ with respect to each vector v_i is computed as

$$\nabla_i \ell(V) = \frac{d\ell}{dV} e_i = -2(H + V)[I + (H + V)^\top (H + V)]^{-1} e_i.$$

Thus, the gradient of the Lagrangian with respect to each vector v_i is

$$\nabla_i \mathcal{L}(V^*, \lambda^*) = -2PQe_i - 2\lambda_i^* v_i,$$

where $P = H + V^*$ and $Q = (I + P^\top P)^{-1}$.

In the following, we prove that any KKT point should satisfy $\lambda^* > 0$, and hence by (14c), it implies that $\|v_i^*\|_2^2 - \alpha_i = 0$ for all i . To this end, suppose without any loss of generality that $\lambda_1^* = 0$. Then the stationary condition (14a) gives

$$PQe_1 = 0,$$

which means that

$$Pq_1 = 0, \quad (15)$$

where q_1 is the first column of Q . From the definition of the inverse matrix, it holds that

$$(I + P^\top P)Q = I.$$

Equating the first column of both sides gives

$$(I + P^\top P)q_1 = e_1. \quad (16)$$

Equation (15) implies that $P^\top Pq_1 = 0$, thus, we have $q_1 = e_1$. Consequentially, equation (15) implies that $Pe_1 = 0$, which means that the first column of $P = H + V^*$ is 0, that is, $h_1 + v_1^* = 0$. Now we construct a new matrix $\tilde{V}$ as $\tilde{V} = [\tilde{v}, v_2^*, \dots, v_N^*]$, where $\tilde{v}$ is any vector with norm $\sqrt{\alpha_1}$ but not equal to $-h_1$. One can see that $\tilde{V}$ is feasible for problem (2). Using the Sylvester's determinant identity $\det(I + AB) = \det(I + BA)$, we have

$$\begin{aligned} \ell(\tilde{V}) &= -\log \det[I + (H + \tilde{V})^\top (H + \tilde{V})] \\ &= -\log \det[I + (H + \tilde{V})(H + \tilde{V})^\top] \\ &= -\log \det[I + \sum_{i=2}^N (h_i + v_i^*)(h_i + v_i^*)^\top + (h_1 + \tilde{v})(h_1 + \tilde{v})^\top] \\ &= -\log \det[I + PP^\top + (h_1 + \tilde{v})(h_1 + \tilde{v})^\top] \quad (h_1 + v_1^* = 0) \\ &= -\log \det(I + PP^\top) - \log \det[I + (I + PP^\top)^{-1}(h_1 + \tilde{v})(h_1 + \tilde{v})^\top] \\ &= \ell(V^*) - \log[1 + (h_1 + \tilde{v})^\top (I + PP^\top)^{-1}(h_1 + \tilde{v})]. \end{aligned}$$

Since $I + PP^\top$ is positive definite and $\tilde{v} \neq -h_1$, we have $(h_1 + \tilde{v})^\top (I + PP^\top)^{-1}(h_1 + \tilde{v}) > 0$, which implies that

$$\ell(\tilde{V}) = \ell(V^*) - \log[1 + (h_1 + \tilde{v})^\top (I + PP^\top)^{-1}(h_1 + \tilde{v})] < \ell(V^*).$$

Hence, we get a contradiction about the primal-dual optimality of (V^*, λ^*) . Thus we can conclude that $\lambda_i^* > 0$ for all $i = 1, \dots, N$. Combined with condition (14c), it holds that $\|v_i^*\|^2 = \alpha_i$ for all i . The proof is completed. $\square$

Proof of Lemma 5. The Euclidean gradient of ℓ is given by (Petersen, Pedersen et al. 2008)

$$\nabla \ell(V) = -2(H + V)M^{-1} \in \mathbb{R}^{p \times N}, \quad (17)$$

where $M = I + (H + V)^\top (H + V) \in \mathbb{R}^{N \times N}$. For a column vector v_i , the gradient of ℓ in this vector is the k -th column of the matrix $-2(H + V)M^{-1}$, i.e.,

$$g_i \triangleq \nabla_i \ell(V) = -2((H + V)M^{-1})_i \in \mathbb{R}^p. \quad (18)$$

The tangent space of the scaled sphere $\mathcal{M}_i$ is given by

$$T_{v_i} \mathcal{M}_i = \{z \in \mathbb{R}^p : z^\top v_i = 0\},$$

whose orthogonal projection is

$$\text{Proj}_{T_{v_i} \mathcal{M}_i}(u) = (I - \frac{1}{\alpha_i} v_i v_i^\top)u, \quad u \in \mathbb{R}^p.$$

Therefore, the Riemannian gradient of ℓ at any point $v_i \in \mathcal{M}_i$ is given by

$$\text{grad}_i \ell(V) = \text{Proj}_{T_{v_i} \mathcal{M}_i}(g_i) = g_i - \frac{1}{\alpha_i} g_i^\top v_i v_i \in \mathbb{R}^p. \quad (19)$$

Since $\mathcal{M}$ is a product manifold, according to Boumal (2023, Exercise 3.67), the Riemannian gradient of ℓ at any point $V \in \mathcal{M}$ is

$$\text{grad} \ell(V) = [\text{grad}_1 \ell(V), \dots, \text{grad}_N \ell(V)].$$

This completes the proof. $\square$

The proof of Proposition 8 relies on the following Lemma 10 and Proposition 11.

Lemma 10 (Upper bound of Hessian). *Let $\bar{\alpha} \triangleq \sum_{i=1}^N \alpha_i$. For any $V \in \mathcal{M}$, it holds that*

$$\sup_{\|U\|_F=1} \|\nabla^2 \ell(V)[U]\|_F \leq 2 + 4 \left(\|H\|_F + \sqrt{\bar{\alpha}} \right)^2.$$

Proof of Lemma 10. For any $V \in \mathcal{M}$ and $U \in \mathbb{R}^{p \times N}$ with $\|U\|_F = 1$, the Hessian of ℓ at V evaluated in the direction U is (Petersen, Pedersen et al. 2008)

$$\nabla^2 \ell(V)[U] = -2UM^{-1} + 2(H + V)M^{-1}(U^\top (H + V) + (H + V)^\top U)M^{-1},$$

where $M = I + (H + V)^\top (H + V)$. Thus we have

$$\|\nabla^2 \ell(V)[U]\|_F \leq 2\|UM^{-1}\|_F + 2\|(H + V)M^{-1}(U^\top (H + V) + (H + V)^\top U)M^{-1}\|_F. \quad (20)$$

Since $M = I + (H + V)^\top (H + V)$ and $(H + V)^\top (H + V)$ is always positive semidefinite, the smallest eigenvalue of M must be equal or greater than 1. Thus, the largest eigenvalue of M^{-1} must be no greater than 1, which means $\|M^{-1}\|_2 \leq 1$. Then, we have

$$\|UM^{-1}\|_F \leq \|U\|_F \|M^{-1}\|_2 \leq \|U\|_F = 1. \quad (21)$$

For the term $\|(H + V)M^{-1}(U^\top (H + V) + (H + V)^\top U)M^{-1}\|_F$, it holds that

$$\begin{aligned} & \|(H + V)M^{-1}(U^\top (H + V) + (H + V)^\top U)M^{-1}\|_F \\ & \leq \|H + V\|_F \|M^{-1}\|_2 \|U^\top (H + V) + (H + V)^\top U\|_2 \|M^{-1}\|_2 \\ & \leq \|H + V\|_F \|U^\top (H + V) + (H + V)^\top U\|_2 \\ & \leq \|H + V\|_F (\|U^\top (H + V)\|_2 + \|(H + V)^\top U\|_2) \\ & \leq \|H + V\|_F \times 2\|U\|_2 \|H + V\|_2 \\ & \leq 2\|H + V\|_F^2 \end{aligned} \quad (\|\cdot\|_2 \leq \|\cdot\|_F). \quad (22)$$

Defining $\bar{\alpha} = \sum_i \alpha_i$, according to the triangle inequality, we have

$$\|H + V\|_F \leq \|H\|_F + \|V\|_F = \|H\|_F + \sqrt{\bar{\alpha}},$$

which, together with (22), implies that

$$\|(H + V)M^{-1}(U^\top (H + V) + (H + V)^\top U)M^{-1}\|_F \leq 2(\|H\|_F + \sqrt{\bar{\alpha}})^2. \quad (23)$$

Combining (20), (21) and (23), we have

$$\|\nabla^2 \ell(V)[U]\|_F \leq 2 + 4(\|H\|_F + \sqrt{\bar{\alpha}})^2,$$

which holds for all $U \in \mathbb{R}^{p \times N}$ with $\|U\|_F = 1$. This completes the proof. $\square$

Proposition 11 (Formula for Hess $\ell(V)$). *For all $V \in \mathcal{M}$ and $U \in T_V \mathcal{M}$, it holds that*

$$\text{Hess } \ell(V)[U] = \text{Proj}_{T_V \mathcal{M}}(\nabla^2 \ell(V)[U]) + U\Lambda, \quad (24)$$

where $\Lambda \in \mathbb{R}^{N \times N}$ is a diagonal matrix with entries $\Lambda_{ii} = -\frac{1}{\alpha_i} g_i^\top v_i$.

Proof of Proposition 11. Boumal (2023, Corollary 5.16) shows that

$$\text{Hess } \ell(V)[U] = \text{Proj}_{T_V \mathcal{M}}(\text{DY}(V)[U]), \quad (25)$$

where Y is any smooth extension of $\text{grad} \ell$ and we can pick $Y(V) = \text{grad} \ell(V)$. Here $\text{DY}(V)$ is a linear map defined as

$$\text{DY}(V)[U] = \lim_{t \rightarrow 0} \frac{Y(V + U) - Y(V)}{t}.$$

Since $\mathcal{M}$ is a product manifold, we can compute $\text{Proj}_{T_V \mathcal{M}}$ block-wise, that is,

$$\text{Hess } \ell(V)[U] = \text{Proj}_{T_V \mathcal{M}}(\text{DY}(V)[U]) = [\text{Proj}_{T_{v_1} \mathcal{M}_1}((\text{DY}(V)[U])_1), \dots, \text{Proj}_{T_{v_N} \mathcal{M}_N}((\text{DY}(V)[U])_N)], \quad (26)$$

where $(\text{DY}(V)[U])_i$ is the i -th block of $\text{DY}(V)[U]$. Noting that $Y(V) = \text{grad } \ell(V) = [\text{grad}_1 \ell(V), \dots, \text{grad}_N \ell(V)]$, it holds that

$$\begin{aligned} (\text{DY}(V)[U])_i &= \text{Dgrad}_i \ell(V)[U] \\ &= \text{D}(g_i - \frac{1}{\alpha_i} g_i^\top v_i v_i)(V)[U] \\ &= \text{D}g_i(V)[U] - \frac{1}{\alpha_i} \text{D}(g_i^\top v_i v_i)(V)[U], \end{aligned} \quad (27)$$

where g_i is defined in (18) and here we take it as a function of V . By definition one can see that

$$\text{D}v_i(V)[u_i] = u_i. \quad (28)$$

Noting that $g_i(V) = \nabla_i \ell(V)$ is the Euclidean gradient of ℓ in the i -th block, we have

$$[\text{D}g_1(V)[U], \dots, \text{D}g_N(V)[U]] = \text{D}[g_1, \dots, g_N](V)[U] = \text{D}\nabla \ell(V)[U] = \nabla^2 \ell(V)[U],$$

which implies that $Dg_i(V)[U]$ is the i -th block of $\nabla^2 \ell(V)[U]$, i.e.,

$$Dg_i(V)[U] = (\nabla^2 \ell(V)[U])_i. \quad (29)$$

Also we have

$$\begin{aligned} D(g_i^\top v_i v_i)(V)[U] &= D(g_i^\top v_i)(V)[U] \cdot v_i + g_i^\top v_i \cdot D(v_i)(V)[U] \\ &= [D(g_i^\top)(V)[U]v_i + g_i^\top D(v_i)(V)[U]] \cdot v_i + g_i^\top v_i \cdot D(v_i)(V)[U] \\ &= [(\nabla^2 \ell(V)[U])_i^\top v_i + g_i^\top u_i] v_i + g_i^\top v_i u_i. \end{aligned} \quad (30)$$

Substituting (28), (29) and (30) into (27) gives us

$$\begin{aligned} (DY(V)[U])_i &= Dg_i(V)[U] - \frac{1}{\alpha_i} D(g_i^\top v_i v_i)(V)[U] \\ &= (\nabla^2 \ell(V)[U])_i - \frac{1}{\alpha_i} [(\nabla^2 \ell(V)[U])_i^\top v_i + g_i^\top u_i] v_i - \frac{1}{\alpha_i} g_i^\top v_i u_i. \end{aligned} \quad (31)$$

Since $\frac{1}{\alpha_i} [(\nabla^2 \ell(V)[U])_i^\top v_i + g_i^\top u_i] v_i$ is a vector along v_i , its projection into $T_{v_i} \mathcal{M}_i$ is 0. Since $\frac{1}{\alpha_i} g_i^\top v_i u_i$ is a vector along $u_i \in T_{v_i} \mathcal{M}_i$, its projection into $T_{v_i} \mathcal{M}_i$ is itself. Hence from (31) we have

$$\text{Proj}_{T_{v_i} \mathcal{M}_i}((DY(V)[U])_i) = \text{Proj}_{T_{v_i} \mathcal{M}_i}((\nabla^2 \ell(V)[U])_i) - \frac{1}{\alpha_i} g_i^\top v_i u_i, \quad (32)$$

which together with (26) implies that

$$\begin{aligned} \text{Hess } \ell(V)[U] &= [\text{Proj}_{T_{v_1} \mathcal{M}_1}((DY(V)[U])_1), \dots, \text{Proj}_{T_{v_N} \mathcal{M}_N}((DY(V)[U])_N)] \\ &= [\text{Proj}_{T_{v_1} \mathcal{M}_1}((\nabla^2 \ell(V)[U])_1) - \frac{1}{\alpha_1} g_1^\top v_1 u_1, \dots, \text{Proj}_{T_{v_N} \mathcal{M}_N}((\nabla^2 \ell(V)[U])_N) - \frac{1}{\alpha_N} g_N^\top v_N u_N] \\ &= [\text{Proj}_{T_{v_1} \mathcal{M}_1}((\nabla^2 \ell(V)[U])_1), \dots, \text{Proj}_{T_{v_N} \mathcal{M}_N}((\nabla^2 \ell(V)[U])_N)] + U\Lambda \\ &= \text{Proj}_{T_V \mathcal{M}}(\nabla^2 \ell(V)[U]) + U\Lambda, \end{aligned}$$

where $\Lambda \in \mathbb{R}^{N \times N}$ is a diagonal matrix with entries $\Lambda_{ii} = -\frac{1}{\alpha_i} g_i^\top v_i$. $\square$

Now we are ready to prove Proposition 8.

Proof of Proposition 8. For any $U \in T_V \mathcal{M}$ with $\|U\|_F = 1$, Proposition 11 gives the Riemannian Hessian of $\ell(V)$ as

$$\text{Hess } \ell(V)[U] = \text{Proj}_{T_V \mathcal{M}}(\nabla^2 \ell(V)[U]) + U\Lambda, \quad (33)$$

where $\Lambda \in \mathbb{R}^{N \times N}$ is a diagonal matrix with entries $\Lambda_{ii} = -\frac{1}{\alpha_i} g_i^\top v_i$. Since $\text{Proj}_{T_V \mathcal{M}}$ is a projection and together with Lemma 10 we have

$$\|\text{Proj}_{T_V \mathcal{M}}(\nabla^2 \ell(V)[U])\|_F \leq \|\nabla^2 \ell(V)[U]\|_F \leq 2 + 4 \left(\|H\|_F + \sqrt{\alpha} \right)^2. \quad (34)$$

For the term $U\Lambda$ we have

$$\begin{aligned} \|U\Lambda\|_F^2 &= \sum_{i=1}^N \left\| \frac{1}{\alpha_i} g_i^\top v_i u_i \right\|_2^2 \\ &= \sum_{i=1}^N \frac{1}{\alpha_i^2} (g_i^\top v_i)^2 \|u_i\|_2^2 \\ &\leq \sum_{i=1}^N \frac{1}{\alpha_i^2} \|g_i\|_2^2 \|v_i\|_2^2 \|u_i\|_2^2 \\ &= \sum_{i=1}^N \frac{1}{\alpha_i} \|g_i\|_2^2 \|u_i\|_2^2 \quad (\|v_i\|_2^2 = \alpha_i) \\ &\leq \frac{1}{\alpha_{\min}} \|\nabla \ell(V)\|_F^2 \sum_{i=1}^N \|u_i\|_2^2 \quad (\|g_i\|_2^2 \leq \|\nabla \ell(V)\|_F^2) \\ &= \frac{1}{\alpha_{\min}} \|\nabla \ell(V)\|_F^2, \quad \left(\sum_{i=1}^N \|u_i\|_2^2 = \|U\|_F^2 = 1 \right) \end{aligned}$$

where $\alpha_{\min} = \min_i \alpha_i > 0$. Now, we bound $\|\nabla \ell(V)\|_F$ as

$$\|\nabla \ell(V)\|_F = \|-2(H + V)M^{-1}\|_F \leq 2\|H + V\|_F \|M^{-1}\|_2 \leq 2(\|H\|_F + \sqrt{\bar{\alpha}}), \quad (35)$$

which gives that

$$\|U\Lambda\|_F \leq \frac{1}{\sqrt{\alpha_{\min}}} \|\nabla \ell(V)\|_F \leq \frac{2}{\sqrt{\alpha_{\min}}} (\|H\|_F + \sqrt{\bar{\alpha}}). \quad (36)$$

Equations (33), (34) and (36) together implies that for any $U \in T_V \mathcal{M}$ with $\|U\|_F = 1$, we have

$$\|\text{Hess } \ell(V)[U]\|_F \leq \|\text{Proj}_{T_V \mathcal{M}}(\nabla^2 \ell(V)[U])\|_F + \|U\Lambda\|_F \leq 2 + 4 \left(\|H\|_F + \sqrt{\bar{\alpha}} \right)^2 + \frac{2}{\sqrt{\alpha_{\min}}} (\|H\|_F + \sqrt{\bar{\alpha}}). \quad (37)$$

According to Boumal (2023, Corollary 10.47), Equation (37) implies that $\ell(V)$ is L -smooth with the postulated constant. This completes the proof. $\square$

Proof of Proposition 9. The proof is similar to the proof of Proposition 8, and here we focus on the Riemannian Hessian of ℓ about the component v_i instead of V . The Riemannian Hessian of $\ell(V)$ with respect to v_i is computes as

$$\text{Hess}_i \ell(V)[u_i] = \text{Proj}_{T_{v_i} \mathcal{M}_i}(\nabla_i^2 \ell(V)[u_i]) - \frac{1}{\alpha_i} g_i^\top v_i u_i,$$

which implies that

$$\|\text{Hess}_i \ell(V)[u_i]\|_2 \leq \|\text{Proj}_{T_{v_i} \mathcal{M}_i}(\nabla_i^2 \ell(V)[u_i])\|_2 + \left\| \frac{1}{\alpha_i} g_i^\top v_i u_i \right\|_2 \leq \|\nabla_i^2 \ell(V)[u_i]\|_2 + \left\| \frac{1}{\alpha_i} g_i^\top v_i u_i \right\|_2.$$

Following the similar arguments in the proof of Lemma 10, we have

$$\sup_{\|u_i\|_2=1} \|\nabla_i^2 \ell(V)[u_i]\|_2 \leq 2 + 4(\|H\|_F + \bar{\alpha})^2. \quad (38)$$

Following the similar arguments in the proof of Proposition 8, for any $u_i \in T_{v_i} \mathcal{M}_i$ with $\|u_i\|_2 = 1$ we have

$$\begin{aligned} \left\| \frac{1}{\alpha_i} g_i^\top v_i u_i \right\|_2 &\leq \frac{1}{\alpha_i} \|g_i\|_2 \|v_i\|_2 \|u_i\|_2 \\ &\leq \frac{1}{\sqrt{\alpha_i}} \|g_i\|_2 && (\text{since } \|v_i\| = \sqrt{\alpha_i}, \|u_i\|_2 = 1) \\ &\leq \frac{1}{\sqrt{\alpha_i}} \|\nabla \ell(V)\|_F \\ &\leq \frac{2}{\sqrt{\alpha_i}} (\|H\|_F + \sqrt{\bar{\alpha}}), \end{aligned}$$

where the last inequality comes from (35). Thus, for any $u_i \in T_{v_i} \mathcal{M}_i$ with $\|u_i\|_2 = 1$ it holds that

$$\begin{aligned} \|\text{Hess}_i \ell(V)[u_i]\|_2 &\leq \|\nabla_i^2 \ell(V)[u_i]\|_2 + \left\| \frac{1}{\alpha_i} g_i^\top v_i u_i \right\|_2 \\ &\leq 2 + 4(\|H\|_F + \sqrt{\bar{\alpha}})^2 + \frac{2}{\sqrt{\alpha_i}} (\|H\|_F + \sqrt{\bar{\alpha}}) = L_i. \end{aligned}$$

Combine the above bound with Boumal (2023, Proposition 10.47) completes the proof. $\square$

We are now ready to prove Theorem 6.

Proof of Theorem 6. Gutman and Ho-Nguyen (2023, Example 3) shows that the selection rule of block in Algorithm 1 satisfies the $(1, \infty)$ -norm condition. Proposition 8 and Proposition 9 show that our objective function ℓ is both L -smooth and block-wise smooth. Equation (11) further demonstrates that the iterates of Algorithm 1 satisfies Gutman and Ho-Nguyen (2023, Assumption 1). Then the convergence of Algorithm 1 is a direct consequence of Gutman and Ho-Nguyen (2023, Theorem 3). $\square$

B Experimental Details

B.1 Experiment Settings

Here, we rigorously introduce where the activation steering is applied in transformer-based LMs. Transformer architectures (Vaswani et al. 2017) have become the foundation of modern LMs, achieving remarkable performance across diverse natural language processing tasks. We consider an LM with M heads per layer, exhibiting the following per-block information flow for the layer l :

$$x^{(l+1)} = \text{FFN}(\text{MHA}(x^{(l)})) = \text{FFN}\left(\bigoplus_{j=1}^M W_j^o(\text{Attn}_j(x^{(l)}))\right),$$

where $x^{(l)}$ is the input of the specific layer l , with W_j^o denoting the output projection matrix and Attn_j denoting the single-head attention transformation for each head $j = 1, \dots, M$. Here, FFN, MHA denote the feed-forward and multi-head attention, respectively. The direct sum operator $\bigoplus$ concatenates the outputs from all attention heads before applying the feed-forward transformation. Contemporary steering methodologies operate by introducing additive perturbations to the residual stream activations, as demonstrated in prior works (Turner et al. 2024; Ackerman 2024).

$$\tilde{x}^{(l+1)} = x^{(l+1)} + v^{(l+1)}, \quad (39)$$

where $v^{(l+1)}$ is the steering vector. The computation of these steering vectors follows the procedure outlined in Algorithm 1.

B.2 Detail Evaluation for Language-Model-as-a-Judge

We provide below the prompts to evaluate the Diversity Score and the Unique Solution Count using OpenAI GPT4.1-mini.

Diversity Score prompt

```
1 You are given a math reasoning problem and a list of different responses (solutions)
   generated by a model.
2 Your task is to assign a float score from 0 to 1.0 that reflects the **diversity of
   reasoning and approaches** among the responses.
3 Consider differences in:
4 - Mathematical strategies
5 - Solution steps
6 - Structural approach
7 - Logical of reasoning
8
9 ### Problem:
10 {problem}
11
12 ### Responses:
13 {responses}
14
15 ### Instruction:
16 Output **only** a float number between 0 and 1.0 (inclusive), rounded to two decimal
   places.
17 Do **not** include any explanation, symbols, or text - only the score.
18
19 Example output: '0.75'
```

Unique Solution Count prompt

```
1 You are given a math reasoning problem and a list of responses generated by a model.
2
3 Your task is to count how many **unique** responses there are.
4
5 Two responses are considered different if they use different:
6 - Mathematical strategies
7 - Solution steps
8 - Logical approach
9 - Structure of reasoning
10
11 ### Problem:
12 {problem}
13
14 ### Responses:
```

```
15 {responses}
16
17 ### Instruction:
18 Output only the number of unique responses as an integer.
19 Do not include any explanation, text, or symbols - just the number.
20
21 Example output: '3'
```

Temperature setting for evaluation. All evaluations using the language-model-as-a-judge were performed with **GPT-4.1 mini** at `temperature = 0.0`, ensuring deterministic and consistent scoring across all prompts.

B.3 Detailed Results with Variance: Unique Solution Count and Diversity Score

Table 3 reports the mean and standard deviation of Unique Solution Count and Diversity Score across benchmarks and sampling temperatures. The results demonstrate that our framework can generate more unique solutions and yield superior diversity scores across most of the experimental conditions.

Metric	Model	T	AIME24			MATH500			OlympiadBench		
			SPREAD		Sampling	SPREAD		Sampling	SPREAD		Sampling
			$C = 1$	$C = 10$		$C = 1$	$C = 10$		$C = 1$	$C = 10$	
Unique Solution Count $\uparrow$	Qwen2.5-1.5B	1.0	6.97 ($\pm$ 0.91)	6.60 ($\pm$ 1.02)	6.67 ($\pm$ 0.99)	3.14 ($\pm$ 0.74)	3.21 ($\pm$ 0.68)	3.14 ($\pm$ 0.69)	6.12 ($\pm$ 1.24)	6.14 ($\pm$ 1.23)	6.14 ($\pm$ 1.61)
		0.8	6.83 ($\pm$ 1.00)	6.43 ($\pm$ 1.20)	3.60 ($\pm$ 0.39)	3.03 ($\pm$ 0.72)	3.05 ($\pm$ 0.70)	2.97 ($\pm$ 0.75)	5.99 ($\pm$ 1.22)	5.99 ($\pm$ 1.36)	6.04 ($\pm$ 1.53)
		0.6	6.40 ($\pm$ 1.11)	6.07 ($\pm$ 1.15)	3.37 ($\pm$ 0.37)	2.91 ($\pm$ 0.74)	2.86 ($\pm$ 0.78)	2.90 ($\pm$ 0.77)	5.57 ($\pm$ 1.33)	5.65 ($\pm$ 1.26)	5.58 ($\pm$ 1.67)
		0.4	6.00 ($\pm$ 1.10)	6.20 ($\pm$ 1.05)	6.17 ($\pm$ 0.90)	2.73 ($\pm$ 0.81)	2.74 ($\pm$ 0.80)	2.77 ($\pm$ 0.79)	5.30 ($\pm$ 1.36)	5.30 ($\pm$ 1.40)	5.26 ($\pm$ 1.99)
		0.2	5.17 ($\pm$ 1.19)	5.60 ($\pm$ 1.11)	2.97 ($\pm$ 0.65)	2.47 ($\pm$ 0.78)	2.56 ($\pm$ 0.82)	2.59 ($\pm$ 0.83)	4.65 ($\pm$ 1.48)	4.69 ($\pm$ 1.47)	4.69 ($\pm$ 2.21)
	Qwen2.5-Math 1.5B- Instruct	1.0	6.63 ($\pm$ 0.86)	7.03 ($\pm$ 1.07)	3.67 ($\pm$ 0.60)	1.92 ($\pm$ 1.07)	1.94 ($\pm$ 1.10)	1.93 ($\pm$ 1.06)	4.59 ($\pm$ 2.34)	4.53 ($\pm$ 2.35)	4.57 ($\pm$ 2.34)
		0.8	6.63 ($\pm$ 1.04)	6.73 ($\pm$ 1.50)	3.47 ($\pm$ 0.67)	1.87 ($\pm$ 1.05)	1.90 ($\pm$ 1.07)	1.89 ($\pm$ 1.10)	4.28 ($\pm$ 2.31)	4.36 ($\pm$ 2.32)	4.39 ($\pm$ 2.31)
		0.6	6.27 ($\pm$ 1.28)	6.47 ($\pm$ 1.37)	3.50 ($\pm$ 0.62)	1.81 ($\pm$ 1.01)	1.84 ($\pm$ 1.02)	1.82 ($\pm$ 1.00)	4.10 ($\pm$ 2.24)	4.16 ($\pm$ 2.26)	4.11 ($\pm$ 2.24)
		0.4	6.23 ($\pm$ 1.24)	5.87 ($\pm$ 1.13)	6.17 ($\pm$ 1.57)	1.77 ($\pm$ 0.95)	1.73 ($\pm$ 0.93)	1.78 ($\pm$ 0.98)	3.99 ($\pm$ 2.19)	3.92 ($\pm$ 2.19)	3.89 ($\pm$ 2.14)
		0.2	5.87 ($\pm$ 1.45)	5.87 ($\pm$ 1.28)	5.53 ($\pm$ 1.78)	1.70 ($\pm$ 0.89)	1.74 ($\pm$ 0.92)	1.74 ($\pm$ 0.93)	3.67 ($\pm$ 2.04)	3.65 ($\pm$ 2.04)	3.68 ($\pm$ 2.06)
	Diversity Score $\uparrow$	Qwen2.5-1.5B	1.0	0.37 ($\pm$ 0.08)	0.33 ($\pm$ 0.06)	0.40 ($\pm$ 0.09)	0.40 ($\pm$ 0.05)	0.39 ($\pm$ 0.05)	0.39 ($\pm$ 0.05)	0.38 ($\pm$ 0.06)	0.40 ($\pm$ 0.06)
			0.8	0.51 ($\pm$ 0.09)	0.45 ($\pm$ 0.08)	0.37 ($\pm$ 0.05)	0.42 ($\pm$ 0.05)	0.44 ($\pm$ 0.06)	0.41 ($\pm$ 0.05)	0.46 ($\pm$ 0.06)	0.46 ($\pm$ 0.06)
			0.6	0.52 ($\pm$ 0.08)	0.46 ($\pm$ 0.06)	0.39 ($\pm$ 0.07)	0.41 ($\pm$ 0.06)	0.40 ($\pm$ 0.06)	0.40 ($\pm$ 0.06)	0.46 ($\pm$ 0.06)	0.46 ($\pm$ 0.06)
			0.4	0.41 ($\pm$ 0.05)	0.47 ($\pm$ 0.08)	0.42 ($\pm$ 0.05)	0.36 ($\pm$ 0.05)	0.36 ($\pm$ 0.06)	0.36 ($\pm$ 0.06)	0.42 ($\pm$ 0.05)	0.43 ($\pm$ 0.05)
			0.2	0.31 ($\pm$ 0.02)	0.34 ($\pm$ 0.02)	0.33 ($\pm$ 0.05)	0.29 ($\pm$ 0.05)	0.29 ($\pm$ 0.05)	0.28 ($\pm$ 0.05)	0.33 ($\pm$ 0.05)	0.35 ($\pm$ 0.05)
		Qwen2.5-Math 1.5B- Instruct	1.0	0.64 ($\pm$ 0.04)	0.65 ($\pm$ 0.05)	0.63 ($\pm$ 0.05)	0.25 ($\pm$ 0.06)	0.25 ($\pm$ 0.06)	0.25 ($\pm$ 0.06)	0.39 ($\pm$ 0.08)	0.39 ($\pm$ 0.08)
			0.8	0.60 ($\pm$ 0.06)	0.62 ($\pm$ 0.06)	0.57 ($\pm$ 0.05)	0.23 ($\pm$ 0.05)	0.23 ($\pm$ 0.05)	0.23 ($\pm$ 0.06)	0.36 ($\pm$ 0.08)	0.38 ($\pm$ 0.08)
			0.6	0.57 ($\pm$ 0.05)	0.55 ($\pm$ 0.06)	0.55 ($\pm$ 0.05)	0.21 ($\pm$ 0.05)	0.22 ($\pm$ 0.05)	0.20 ($\pm$ 0.05)	0.34 ($\pm$ 0.07)	0.34 ($\pm$ 0.08)
			0.4	0.51 ($\pm$ 0.05)	0.51 ($\pm$ 0.07)	0.51 ($\pm$ 0.06)	0.19 ($\pm$ 0.04)	0.18 ($\pm$ 0.04)	0.18 ($\pm$ 0.04)	0.31 ($\pm$ 0.07)	0.30 ($\pm$ 0.07)
			0.2	0.46 ($\pm$ 0.06)	0.52 ($\pm$ 0.08)	0.41 ($\pm$ 0.07)	0.15 ($\pm$ 0.04)	0.16 ($\pm$ 0.04)	0.15 ($\pm$ 0.04)	0.26 ($\pm$ 0.05)	0.26 ($\pm$ 0.06)

Table 3: Mean and variance value of Unique Solution Count $\uparrow$ and Diversity Score $\uparrow$ comparison across model variants and sampling temperature (T) on three mathematical reasoning benchmarks, including AIME24 (AI-MO 2024), MATH500 (Lightman et al. 2023) and OlympiadBench (He et al. 2024). Bold indicates the best method in each dataset.

B.4 Pass@N Across Steering Layers

We set the random seed to 42 for all experiments. We conduct experiments to study the effects of changing the steering layers on the performance of SPREAD. For simplicity, we use the AIME24 dataset because it has only 30 samples. We screen the steering layers from 1 to 28 using Qwen2.5-Math-1.5B-Instruct. Figure 6 gives the results measured by Pass@8.

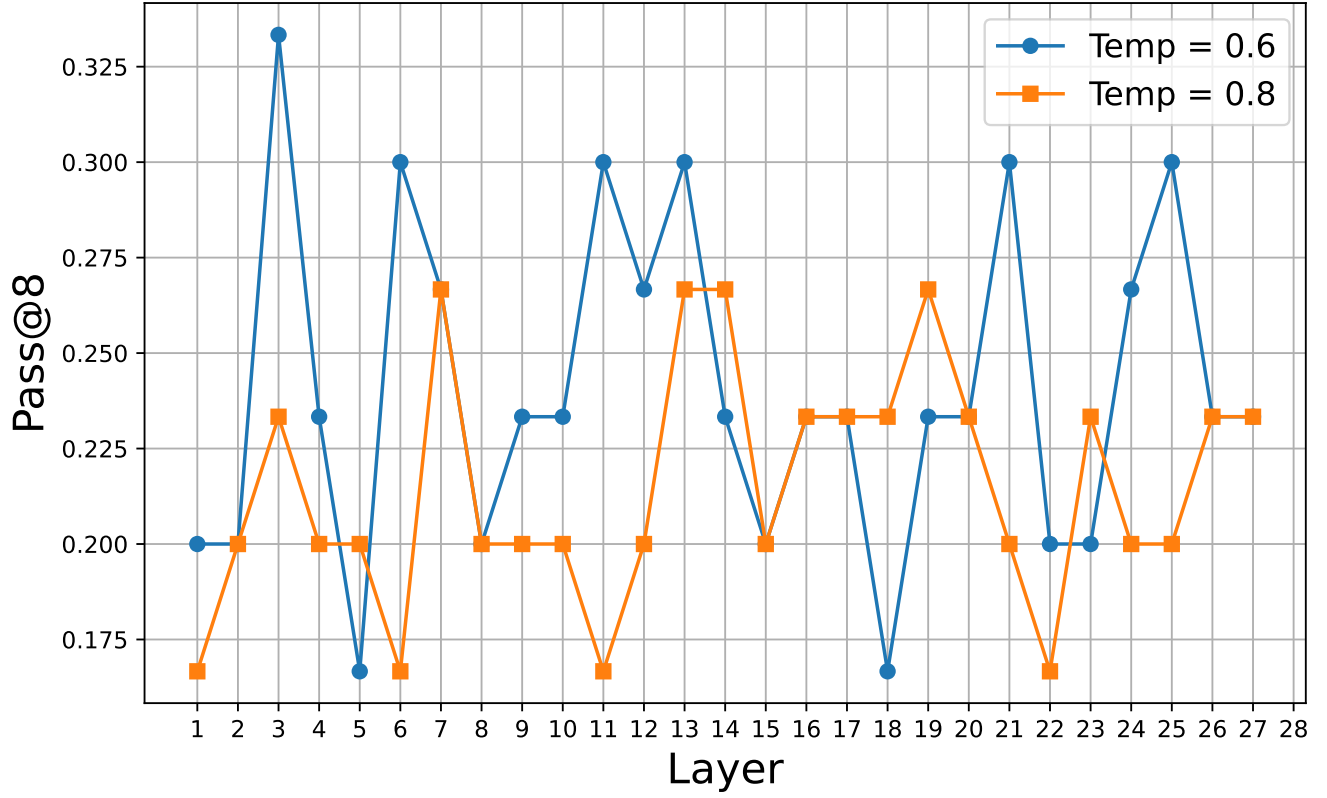


Figure 6: Layer-wise comparison of Pass@N $\uparrow$ (in %) for Qwen2.5-Math-1.5B-Instruct across 28 layers and different temperatures on the AIME24 benchmark.

B.5 Hypothesis Testing

In Section 3, we postulate that a larger volume of the parallelepiped spanned by activations h_{im} will induce diversity in the final answers. We now propose a statistical test to validate this hypothesis.

For each question $q_i, i = 1, \dots, M$, we simultaneously generate $N = 16$ output sequences. At each decoding step τ , we extract the final hidden state vectors of those N paths and form the matrix $H_i^\tau = [h_{i1}^\tau, \dots, h_{iN}^\tau]$ and use Algorithm 1 to compute steer vectors $V_i^\tau = [v_{i1}^\tau, \dots, v_{iN}^\tau]$. For each question q_i , we draw 10 subsamples without replacement of size $\tilde{N} = 8$ from N responses, yielding 10 sub-sampled paths per question. Let $k \in \{1, \dots, 10\}$ index these.

For each sub-sampled path k of question q_i , define the observed diversity

$$U_{ik} = \# \text{ the unique solutions among } \tilde{N} \text{ responses,}$$

and the cumulative ‘‘volume’’ proxy

$$Z_{ik} = \sum_{\tau=1}^{\tau_{final}} \log \det[I + (H_{ik}^\tau + V_{ik}^\tau)^\top (H_{ik}^\tau + V_{ik}^\tau)],$$

where H_{ik}^τ and V_{ik}^τ are activations and steering vectors corresponding to sub-sampled path k of question q_i .

Since U_{ik} is a count between 1 and $\tilde{N}$, we model U_{ik} as a Binomial outcome with $\tilde{N}$ trials and success probability p_{ik} so that

$$U_{ik} \sim \text{Binomial}(\tilde{N}, p_{ik}),$$

where p_{ik} links with Z_{ik} with a logistic regression:

$$p_{ik} = \frac{\exp(\beta_{0,i} + \beta Z_{ik})}{1 + \exp(\beta_{0,i} + \beta Z_{ik})} = \frac{1}{1 + \exp(-\beta_{0,i} - \beta Z_{ik})},$$

where $\beta_{0,i} \in \mathbb{R}$ are question-specific intercepts and $\beta \in \mathbb{R}$ is the common slope parameter. We are interested in testing whether the predictor Z is positively associated with the outcome U :

$$\mathcal{H}_0 : \beta \leq 0 \quad \text{versus} \quad \mathcal{H}_A : \beta > 0.$$

We estimate $\{\beta_{0,i}\}_{i=1}^M$ and β by maximizing the binomial log-likelihood:

$$\max_{\{\beta_{0,i}\}_{i=1}^M, \beta} \sum_{i=1}^M \sum_{k=1}^{10} U_{ik}(\beta_{0,i} + \beta Z_{ik}) - \tilde{N} \log(1 + e^{\beta_{0,i} + \beta Z_{ik}}).$$

Inference is based on a Wald statistic with a cluster-robust (sandwich) variance estimator, clustering at the question level to account for within-question dependence and heteroskedasticity:

$$z_{\text{Wald}} = \frac{\hat{\beta}}{\widehat{\text{SE}}(\hat{\beta})}.$$

Applying this procedure yields $\hat{\beta} = 0.88$, clustered $\widehat{\text{SE}}(\hat{\beta}) = 0.24$, so $z_{\text{Wald}} = 3.6$ and a one-sided $p = 0.001$. Thus, we reject H_0 at conventional levels. Substantively, a one-unit increase in Z multiplies the odds of a response being unique by $\exp(\hat{\beta}) \approx 2.4$, indicating a strong positive association between the proposed volume measure and answer diversity.