

Beyond MuP: 4. Preserving Parameter Stability

Su Jianlin

2026-04-24

Through derivations and calculations in previous articles, we can observe that the three stability metrics proposed in the first article *Beyond MuP: 1. Three Characteristics of a Good Model* can generally be divided into *parameter stability* and *increment stability*. In *Beyond MuP: 2. Linear Layers and Steepest Descent* and *Beyond MuP: 3. Special Cases Require Special Treatment*, we demonstrated the process of integrating increment stability with steepest descent to obtain new update rules (optimizers).

However, for parameter stability, we previously only considered initialization. The task of this article is to explore how to maintain parameter stability throughout the entire training process, thus completing the practical application of the theory.

1 Background

Taking *Beyond MuP: 2. Linear Layers and Steepest Descent* as an example, the three stability metrics are:

$$\text{Forward Stability: } \max_{\|\mathbf{x}\|_{RMS}=1} \|\mathbf{x}\mathbf{W}\|_{RMS} = \sqrt{\frac{d_{in}}{d_{out}}} \|\mathbf{W}\|_2 \quad (1)$$

$$\text{Dependency Stability: } \max_{\|\mathbf{x}_1\|_{RMS}=\|\mathbf{x}_2\|_{RMS}=1} \|\mathbf{x}_1\mathbf{W} - \mathbf{x}_2\mathbf{W}\|_{RMS} = 2\sqrt{\frac{d_{in}}{d_{out}}} \|\mathbf{W}\|_2 \quad (2)$$

$$\text{Update Stability: } \max_{\|\mathbf{x}\|_{RMS}=1} \|\mathbf{x}(\mathbf{W} + \Delta\mathbf{W}) - \mathbf{x}\mathbf{W}\|_{RMS} = \sqrt{\frac{d_{in}}{d_{out}}} \|\Delta\mathbf{W}\|_2 \quad (3)$$

where $\mathbf{W} \in \mathbb{R}^{d_{in} \times d_{out}}$ is the parameter of a linear layer. We desire these three metrics to be $\Theta(1)$, which means we want the norm of the parameter and its increment to satisfy $\|\mathbf{W}\|_2 = \Theta(\sqrt{d_{out}/d_{in}})$ and $\|\Delta\mathbf{W}\|_2 = \Theta(\sqrt{d_{out}/d_{in}})$, respectively. In *Beyond MuP: 3. Special Cases Require Special Treatment*, we performed similar calculations for layers such as Embedding and LM Head, with analogous conclusions, though the corresponding norms differ.

We used the increment condition as a stability metric and, based on the principle of steepest descent under “stable acceleration,” derived theoretically optimal update rules. For linear layers, this leads to the Muon optimizer:

$$\operatorname{argmin}_{\|\Delta\mathbf{W}\|_2 \leq \eta \sqrt{\frac{d_{out}}{d_{in}}}} \operatorname{tr}(\mathbf{G}^\top \Delta\mathbf{W}) \quad \Rightarrow \quad \Delta\mathbf{W} = -\eta \sqrt{\frac{d_{out}}{d_{in}}} \operatorname{msign}(\mathbf{G}) \quad (4)$$

For parameter stability, we previously only required that initialization satisfies $\|\mathbf{W}\|_2 = \Theta(\sqrt{d_{out}/d_{in}})$. How to ensure such parameter stability is maintained throughout the entire training process remains unclear.

2 Preliminary Thoughts

How can we ensure that $\mathbf{W}$ maintains $\|\mathbf{W}\|_2 = \Theta(\sqrt{d_{out}/d_{in}})$? More generally, given a parameter ω , which could be a vector, matrix, or even a higher-order tensor, a norm $\|\cdot\|$, typically induced by forward or dependency stability, and a scale τ , how can we ensure that ω maintains $\|\omega\| = \Theta(\tau)$ during training?

A naive idea is to directly enforce $\|\omega\| = \tau$ (where τ can be replaced by a constant multiple, but this does not affect the following discussion). The simplest implementation is to re-scale the norm to τ after each optimization step via normalization (e.g., Hyperball and Nemotron-Flash). Another approach is to replace the original model parameterization with normalized ones, i.e., change $\mathbf{f}(\mathbf{x}; \omega)$ to $\mathbf{f}(\mathbf{x}; \tau\omega/\|\omega\|)$, which theoretically achieves a similar effect.

A more advanced approach combines this with the steepest descent idea to adjust the update rule, as discussed in articles *Steepest Descent on Manifolds: 1. SGD + Hypersphere* and *Steepest Descent on Manifolds: 4. Muon + Spectral Sphere*, and in the paper *Controlled LLM Training on Spectral Sphere*. This approach is more elegant in principle but more complex in practice, often requiring solving a nonlinear equation to obtain the exact update.

However, should we really insist on strictly controlling a parameter’s norm to a fixed value? Intuitively, the parameter norm should be determined by the training process itself; at best, we set a prior range for it. Although some work shows that setting an appropriate fixed norm does not harm performance, it still disrupts the original training dynamics and may require extra effort to adapt.

Therefore, the view proposed in this article is that we only need to ensure $\|\omega\| = \mathcal{O}(\tau)$. Specifically, we aim to guarantee $\|\omega\| \leq \tau$ at every step, while leaving the exact value including whether it reaches $\Theta(\tau)$ to be determined by the training algorithm, without further intervention.

3 Post Clip

The next question is: how to achieve $\|\omega\| \leq \tau$? More specifically, suppose the original update rule for ω is

$$\omega_t = \omega_{t-1} - \eta\phi_t \quad (5)$$

How should we modify it to ensure $\|\omega_t\| \leq \tau$? There are many possible methods; for instance, normalization as mentioned earlier is one option. Among them, we aim to select the one that causes the *least disruption* to the optimization process. That is, given ω and a norm $\|\cdot\|$, we aim to minimize the modification to make its norm at most τ . Formally, this is defined as:

$$\lfloor \omega \rfloor_{\|\cdot\| \leq \tau} = \operatorname{argmin}_{\|\tilde{\omega}\| \leq \tau} \|\omega - \tilde{\omega}\|_{RMS} \quad (6)$$

Familiar readers of convex optimization will recognize this as the projection of ω into a ball of radius τ under the specified norm. The key idea is to achieve the goal of bounded norm while minimizing the perturbation to the original parameter ω , thus inducing a specific projection or clipping operation.

The exact computation of $\lfloor \omega \rfloor_{\|\cdot\| \leq \tau}$ depends on the specific norm and will be detailed later. With this operator, one possible scheme is to apply it after each update:

$$\omega_t = \lfloor \omega_{t-1} - \eta\phi_t \rfloor_{\|\cdot\| \leq \tau} \quad (7)$$

We call this scheme “Post Clip,” which is simple and intuitive. However, it may feel “non-smooth”: if initialization starts below τ , the parameter norm increases slowly until reaching τ , at which point clipping “suddenly” activates. While continuous, this is not differentiable at the boundary similar to the $\max(x, 0)$ function.

4 Pre Decay

If this non-smoothness is undesirable, one can imitate weight decay by distributing the penalty across updates. Still starting from update rule (5), suppose $\|\phi_t\| \leq \tau$, then by the triangle inequality:

$$\|\omega_t\| = \|\omega_{t-1} - \eta\phi_t\| \leq \|\omega_{t-1}\| + \eta\tau, \quad (8)$$

meaning the norm could grow by up to $\eta\tau$ per step, accumulating out of control over time.

To prevent this, we can pre-process ω_{t-1} before subtracting $\eta\phi_t$, reducing its norm just enough to counteract the potential increase. Following the weight decay analogy:

$$\omega_t = \lfloor \omega_{t-1} \rfloor_{\|\cdot\| \leq (1-\eta)\|\omega_{t-1}\|} - \eta\phi_t \quad (9)$$

That is, first reduce ω_{t-1} 's norm to $(1 - \eta)$ times its original value, then update. This gives:

$$\|\omega_t\| \leq (1 - \eta)\|\omega_{t-1}\| + \eta\tau \leq \max(\|\omega_{t-1}\|, \tau),$$

and by induction, $\|\omega_t\| \leq \max(\|\omega_0\|, \tau)$. Thus, if initialization satisfies $\|\omega_0\| \leq \tau$, then $\|\omega_t\| \leq \tau$ holds throughout training. This result is independent of the choice of norm; it only relies on the triangle inequality. The minimal-modification operation to reduce norm is precisely the clipping operator in (6), making its use natural here.

We call this ‘‘Pre Decay.’’ Unlike Post Clip, which uses a static threshold τ , Pre Decay uses a dynamic threshold $(1 - \eta)\|\omega_{t-1}\|$ and always triggers a reduction, making the process smoother. Hence, we call it ‘‘decay’’ instead of ‘‘clip,’’ generalizing conventional weight decay.

5 Simple Example

So far, we have established a general framework for constraining parameter norms, with Post Clip and Pre Decay as two schemes, centered on the clipping operator $\lfloor \omega \rfloor_{\|\cdot\| \leq \tau}$ defined in (6). While formally defined, practical implementation depends on the norm.

Consider a simple example using $\|\cdot\|_{RMS}$. For vectors, this equals the L2 norm; for matrices, it equals the Frobenius norm. We have:

$$\lfloor \omega \rfloor_{\|\cdot\|_{RMS} \leq \tau} = \operatorname{argmin}_{\tilde{\omega} \|\cdot\|_{RMS} \leq \tau} \|\omega - \tilde{\omega}\|_{RMS} = \min \left(1, \frac{\tau}{\|\omega\|_{RMS}} \right) \omega \quad (10)$$

The proof is left to the reader (if stuck, consult Kimi). Substituting $\tau = (1 - \eta)\|\omega\|_{RMS}$ yields:

$$\lfloor \omega \rfloor_{\|\cdot\|_{RMS} \leq (1-\eta)\|\omega\|_{RMS}} = \min \left(1, \frac{(1 - \eta)\|\omega\|_{RMS}}{\|\omega\|_{RMS}} \right) \omega = (1 - \eta)\omega \quad (11)$$

Plugging into (9) gives:

$$\omega_t = (1 - \eta)\omega_{t-1} - \eta\phi_t \quad (12)$$

This is precisely conventional weight decay. Thus, Pre Decay under RMS norm reduces to standard weight decay, the minimal modification scheme to preserve RMS (equivalently, L2 or Frobenius) norm bounds.

6 Singular Clipping

We now turn to the main eventmatrix parameters and Muon. Recall the Muon update:

$$\mathbf{W}_t = \mathbf{W}_{t-1} - \eta\lambda\Phi_t, \quad \Phi_t = \frac{1}{\lambda} \sqrt{\frac{d_{out}}{d_{in}}} \operatorname{msign}(\cdot) \mathbf{G}_t \quad (13)$$

Let $\tau = \frac{1}{\lambda} \sqrt{\frac{d_{out}}{d_{in}}}$, so $\|\Phi_t\|_2 = \tau$. To ensure $\|\mathbf{W}_t\|_2 \leq \tau$, two schemes are:

$$\text{Post Clip: } \mathbf{W}_t = \lfloor \mathbf{W}_{t-1} - \eta\lambda\Phi_t \rfloor_{\|\cdot\|_2 \leq \tau} \quad (14)$$

$$\text{Pre Decay: } \mathbf{W}_t = \lfloor \mathbf{W}_{t-1} \rfloor_{\|\cdot\|_2 \leq (1-\eta\lambda)\|\mathbf{W}_{t-1}\|_2} - \eta\lambda\Phi_t \quad (15)$$

Now compute $\lfloor \mathbf{W} \rfloor_{\|\cdot\|_2 \leq \tau}$, which is:

$$\lfloor \mathbf{W} \rfloor_{\|\cdot\|_2 \leq \tau} = \operatorname{argmin}_{\tilde{\mathbf{W}}_{\|\cdot\|_2 \leq \tau}} \|\mathbf{W} - \tilde{\mathbf{W}}\|_F \quad (16)$$

This problem's optimal solution is known as ‘‘Singular Value Clipping (SVC)’’ or $\text{mclip}(\cdot; \cdot)$

$$\lfloor \mathbf{W} \rfloor_{\|\cdot\|_2 \leq \tau} = \text{mclip}(\cdot; \cdot) \mathbf{W}; \tau = \mathbf{U} \min(\Sigma, \tau) \mathbf{V}^\top \quad (17)$$

where $\mathbf{U}\Sigma\mathbf{V}^\top$ is the SVD of $\mathbf{W}$ and $\min(\Sigma, \tau)$ truncates singular values to at most τ . With this notation:

$$\text{Post Clip: } \mathbf{W}_t = \text{mclip}(\cdot; \cdot) \mathbf{W}_{t-1} - \eta\lambda\Phi_t; \tau \quad (18)$$

$$\text{Pre Decay: } \mathbf{W}_t = \text{mclip}(\cdot; \cdot) \mathbf{W}_{t-1}; (1 - \eta\lambda)\|\mathbf{W}_{t-1}\|_2 - \eta\lambda\Phi_t \quad (19)$$

7 Derivation

We now prove (17). Let $\mathbf{W} = \mathbf{U}\Sigma\mathbf{V}^\top$ be the SVD. Then:

$$\|\mathbf{W} - \tilde{\mathbf{W}}\|_F = \|\mathbf{U}\Sigma\mathbf{V}^\top - \tilde{\mathbf{W}}\|_F \quad (20)$$

$$= \|\mathbf{U}^\top(\mathbf{U}\Sigma\mathbf{V}^\top - \tilde{\mathbf{W}})\mathbf{V}\|_F \quad (21)$$

$$= \|\Sigma - \mathbf{U}^\top \tilde{\mathbf{W}} \mathbf{V}\|_F \quad (22)$$

since orthogonal transformations preserve Frobenius norm. Let $\tilde{\Sigma} = \mathbf{U}^\top \tilde{\mathbf{W}} \mathbf{V}$. The objective becomes:

$$\operatorname{argmin}_{\tilde{\Sigma}_{\|\cdot\|_2 \leq \tau}} \|\Sigma - \tilde{\Sigma}\|_F \quad (23)$$

Since $\|\tilde{\Sigma}\|_2 \leq \tau$ and the minimum of $\|\Sigma - \tilde{\Sigma}\|_F$ occurs when off-diagonal elements of $\tilde{\Sigma}$ are zero and diagonal elements are $\tilde{\sigma}_i = \min(\sigma_i, \tau)$, we obtain:

$$\tilde{\Sigma}^* = \min(\Sigma, \tau) \quad (24)$$

Thus, $\tilde{\mathbf{W}}^* = \mathbf{U} \min(\Sigma, \tau) \mathbf{V}^\top$, proving (17).

8 Main Component Clipping

How to efficiently compute $\text{mclip}(\cdot; \cdot)$? Performing SVD every step is expensive. As discussed in *Computing mclip via msign (Part I)* and *Computing mclip via msign (Part II)*, one can use $\text{msign}(\cdot)$ but it requires multiple passes. An alternative uses iterative top-singular-value estimation (SVD1). Since $\text{mclip}(\cdot; \tau)$ clips all singular values above τ to τ , repeatedly clipping the top singular value achieves the same effect. Assuming training smoothness, one clipping per step suffices. Thus:

$$\text{Post Clip: } \mathbf{W}_t = \tilde{\mathbf{W}}_t - \max(\sigma_1 - \tau, 0) \mathbf{u}_1 \mathbf{v}_1^\top, \quad \sigma_1, \mathbf{u}_1, \mathbf{v}_1 = \text{SVD1}(\tilde{\mathbf{W}}_t), \quad \tilde{\mathbf{W}}_t = \mathbf{W}_{t-1} - \eta\Phi_t \quad (25)$$

$$\text{Pre Decay: } \mathbf{W}_t = \mathbf{W}_{t-1} - \lambda\eta\sigma_1 \mathbf{u}_1 \mathbf{v}_1^\top - \eta\Phi_t, \quad \sigma_1, \mathbf{u}_1, \mathbf{v}_1 = \text{SVD1}(\mathbf{W}_{t-1}) \quad (26)$$

The Pre Decay version matches the spectral weight decay introduced in *From Spectral Gradient to New Weight Decay*.

9 Other Details

For greater accuracy, approximate k largest singular values and clip up to k per step. This requires QR during power iteration (see *Streamed Power Iteration for Muon: Part I*).

Other norms appear in *Beyond MuP: 3*, e.g., row/column RMS for Embedding/LM Head and ℓ_∞ norm for gamma in RMS Norm. Their clipping operators are straightforward: - Embedding: clip each row’s RMS to $\leq \tau$ - LM Head: same for columns - gamma: clip element-wise to $[-\tau, \tau]$, i.e., $\max(\min(., \gamma); \tau, -\tau)$

These are intuitive and left as exercises.

10 Necessary Justification

Some may ask: why not use ordinary weight decay as in *Training Deep Learning Models with Norm-Constrained LMOs*:

$$\mathbf{W}_t = (1 - \eta\lambda)\mathbf{W}_{t-1} - \eta\sqrt{\frac{d_{out}}{d_{in}}} \text{msign}(\cdot)\mathbf{G}_t \quad (27)$$

to limit spectral norm?

The issue is over-intervention. The clipping operator (6) minimizes perturbation. Multiplying by $1 - \eta\lambda$ reduces spectral norm but is not minimal; hence, it over-intervenes. Consequences: - Small λ : insufficient control (τ too large) - Large λ : too much damping

For example, to bound spectral norm by 5 with $d_{in} = d_{out}$, $\lambda = 0.2$ gives excessive 0.2 weight decay (typical is 0.01).

We stress “guarantee”: small models may stay below 5 with decay 0.01, but large models can reach theoretical limits (here 100). Maintaining a tight theoretical bound is essential for safety this is the “stable” in “stable acceleration”. The clipping operator (6) provides the lightest intervention that guarantees bounded norms, thus minimizing performance loss.

11 Summary

This article, based on minimal-perturbation principles, proposes a general framework for maintaining parameter stability during training, including Post Clip and Pre Decay. Under spectral norm, they induce Singular Value Clipping and Spectral Weight Decay. These operations aim to bound key parameter norms while minimizing interference with training dynamics.