
A Multi-Power Law for Loss Curve Prediction Across Learning Rate Schedules

Kairong Luo 2024310643 Zhixuan Pan 2024311550 Chumeng Jiang 2024316092

Abstract

Training large language models is both resource-intensive and time-consuming, making it crucial to understand the quantitative relationship between model performance and hyperparameters. In this paper, we derive an empirical law that is able to predict the entire pretraining loss curve across various learning rate schedules, including constant, cosine, and step decay schedules. This law takes a multi-power form, combining a power law based on the sum of learning rates and additional power laws to account for a loss reduction effect as learning rate decays. We validate this law extensively on Llama-2 models of varying sizes and demonstrate that, after fitting on a few learning rate schedules, it accurately predicts the loss curves for unseen schedules of different shapes and horizons. Moreover, by minimizing the predicted final pretraining loss across learning rate schedules, we are able to find a schedule that outperforms the widely-used cosine learning rate schedule. Interestingly, this automatically discovered schedule bears some resemblance to the recently proposed Warmup-Stable-Decay (WSD) schedule (Hu et al., 2024) but achieves a slightly lower final loss. We believe these results could offer valuable insights for understanding the dynamics of pretraining and for designing learning rate schedules to improve efficiency.

1 Introduction

Language models can achieve strong performance if pretrained at a very large scale with an appropriate configuration of hyperparameters, such as model width, model depth, number of training steps, and learning rate. However, a full-scale grid search over these hyperparameters is often impossible since one large-scale pretraining run can take weeks or even months.

To reduce the cost of hyperparameter tuning, researchers have proposed various scaling laws that aim to predict the final pretraining loss or downstream performance at scale. These laws usually try to capture an empirical relationship between the final performance and a few key hyperparameters, and use a simple parameterized function to approximate this relationship. A notable example is the Chinchilla scaling law (Hoffmann et al., 2022), which approximates the final pretraining loss as a function of the model size N and the total number of training steps T (or alternatively, the number of training tokens), $\mathcal{L}(N, T) = L_0 + A \cdot T^{-\alpha} + B \cdot N^{-\beta}$. Based on a few experiments with varying N and T , one can fit the parameters L_0, A, B, α, β and use the formula to infer the optimal choice of N and T given a fixed compute budget $C = NT$.

However, existing scaling laws fall short in providing guidance on the choice of *Learning Rate* (LR), which is arguably the most critical hyperparameters in optimization. It is indeed very challenging to incorporate the effect of LR into the scaling laws, as its impact on the training speed and stability is intricate and not yet well understood in a quantitative manner. A qualitative understanding is as follows: a large LR can reduce the training loss quickly, but in the long term, it may cause overshooting and oscillation along sharp directions on the loss landscape. In contrast, a small LR leads to a more stable training process, but at the cost of slowing down the convergence. Practitioners often trade-off between these two extremes by starting training with a large LR and then gradually

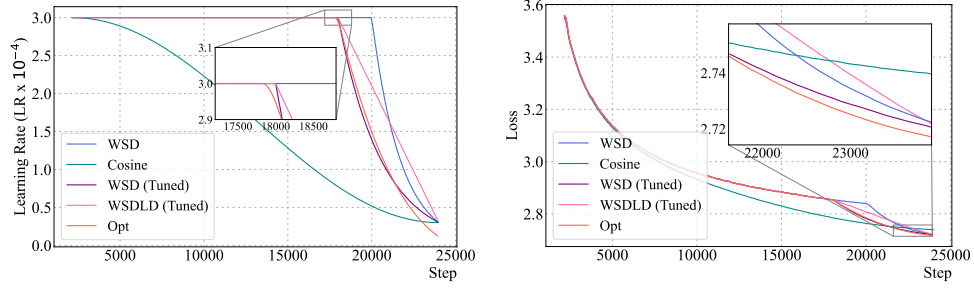


Figure 1: Performance of various schedules, including our optimized LR schedule (Opt), on a 400M Llama-2 (Touvron et al., 2023) model over 12B tokens. Zoom in/out facilitates the readers who are interested in the local details. (a) Our optimal schedule comprises constant and decay stages post-warmup, aligning with WSD (Hu et al., 2024). See Section 4 for details. (b) Our optimized schedule outperforms cosine LR and tuned WSD variants (WSD uses exponential decay; WSDLD uses linear decay).

reducing it over time, following a *Learning Rate schedule* (LR schedule) (Bengio, 2012). These LR schedules sometimes include a warmup phase at the beginning, where the LR is gradually increased from a small value to a large value over a few thousand steps, and only after this warmup phase does the LR start to decay. The most commonly used LR schedule in language model pretraining is the cosine schedule (Loshchilov & Hutter, 2017), which decays the LR following a cosine curve. Other schedules include the cyclic (Smith, 2017), Noam (Vaswani, 2017), and Warmup-Stable-Decay (WSD) schedules (Hu et al., 2024), but there is no consensus on which schedule is the best.

In this paper, we aim to obtain a quantitative understanding of the empirical relationship between the LR schedule and the final training loss in language model pretraining. More specifically, we study the following problem, which we call the *schedule-aware loss curve prediction problem*: *Can we use a simple formula to accurately predict the training loss curve $\mathcal{L}(t)$ ($1 \leq t \leq T$) given a LR schedule $E := \{\eta_1, \eta_2, \dots, \eta_T\}$ for T steps of training?* Following the standard practice in pretraining, we assume that each training step is taking fresh samples from a data stream, thus there is no generalization gap between the training and test loss. We focus on learning rate schedules that decay the LR over time ($\eta_t \leq \eta_{t-1}$) as these schedules are the most common in practice.¹ Moreover, we assume that we have already picked a good initial LR $\eta_{\max}$ that is nearly optimal for short training runs without LR decay. Starting from this initial LR $\eta_1 = \eta_{\max}$, we are interested in predicting the loss curve as the LR decays over time.

Existing scaling laws are insufficient for this problem because they are usually overfitted to one predetermined LR schedule. For example, Hoffmann et al. (2022) fitted the parameters L_0, A, B, α, β in the Chinchilla scaling law $\mathcal{L}(N, T) = L_0 + A \cdot T^{-\alpha} + B \cdot N^{-\beta}$ for training runs that have gone through the entire cosine LR schedule, which makes the law inapplicable to other LR schedules, or even to the same LR schedule with early stopping. Finding a good scaling law for this problem also requires a more sophisticated approach. In contrast to existing laws that make predictions based on two or three hyperparameters, which are easy to visualize and analyze, here we need to predict the loss curve based on the entire LR schedule, which is inherently high-dimensional. A more careful experimental design is thus needed to speculate what a good approximation formula could be.

Our Contribution: Multi-Power Law. In this paper, we propose the following empirical law (1) for schedule-aware loss curve prediction:

$$\mathcal{L}(t) = L_0 + A \cdot S_1(t)^{-\alpha} - \text{LD}(t), \quad \text{where} \quad S_1(t) := \sum_{\tau=1}^t \eta_{\tau}. \quad (1)$$

Here, $L_0 + A \cdot S_1(t)$ can be seen as a naïve extension of the Chinchilla scaling law by replacing the number of steps T with the sum of LR and neglecting the dependence on the model size. This alone can be seen as a crude approximation of the loss curve that linearizes the contribution of the LR at each step, but it is agnostic to the shape of LR decay. The remaining term $\text{LD}(t)$ serves as a correction term that captures how decaying the LR to smaller values leads to a reduction in the loss:

$$\text{LD}(t) := B \sum_{k=2}^t (\eta_{k-1} - \eta_k) \cdot G(\eta_k^{-\gamma} S_k(t)), \quad S_k(t) := \sum_{\tau=k}^t \eta_{\tau}, \quad G(x) := 1 - (Cx + 1)^{-\beta}. \quad (2)$$

More specifically, $\text{LD}(t)$ is the sum of the LR reduction $\eta_{k-1} - \eta_k$ at each step k multiplied by a nonlinear factor. The factor gradually saturates to a constant as the training progresses, and the speed of saturation follows a power law in a scaled sum of LR $\eta_k^{-\gamma} S_k(t)$.

¹If there is a LR warmup phase in the schedule, we focus on the decay phase right after the warmup phase.

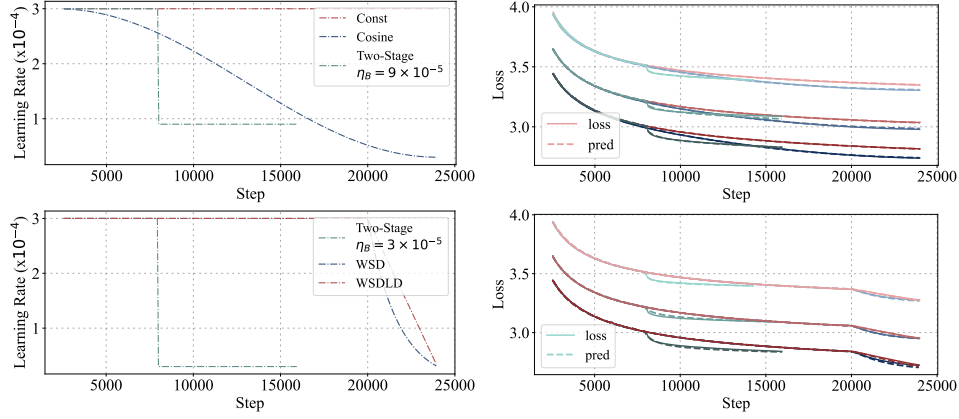


Figure 2: Loss Curves of 25M, 100M, and 400M models from up to down. (a) **Fit on Training Set:** Our multi-power law is reducible to two-stage and Constant LR schedules, and captures Cosine LR decay effects. (b) **Prediction on Test Set:** Our law generalizes to unseen schedules like WSDL and WSD, and handles steep decays in Two-Stage cases.

We call this law of $\mathcal{L}(t)$ the *multi-power scaling law* as it consists of multiple power-law forms. See also Appendix D.3 for the practical version of our law that accounts for the warmup phase. $L_0, A, B, C, \alpha, \beta, \gamma$ are the parameters of the law and can be fitted by running very few pretraining experiments with different LR schedules. We summarize our main contributions as follows:

1. We propose the multi-power law (1) for schedule-aware loss curve prediction, and empirically validate that after fitting the parameters of the law on at most 3 pretraining runs, it can predict the loss curve for unseen LR schedules with remarkable accuracy (see Figure 1). Unlike the Chinchilla scaling law, which relies solely on the final loss of each training run to fit its parameters, our approach utilizes the entire loss curve of each training run to fit the parameters, thus significantly reducing the number of training runs and compute resources needed for accurate predictions (Figure 7). Extensive experiments are presented for various model architectures, sizes, and training horizons (Section 3).
2. Our multi-power law is accurate enough to be used to search for better LR schedules. We show that by minimizing the predicted final loss according to the law, we can obtain an optimized LR schedule that outperforms the standard cosine schedule. Interestingly, the optimized schedule has a similar shape as the recently proposed WSD schedule (Hu et al., 2024), but its shape is optimized so well that it outperforms WSD with grid-searched hyperparameters (Section 4).
3. We use a novel “bottom-up” approach to empirically derive the multi-power law. Starting from two-stage schedules, we conduct a series of ablation studies on LR schedules with increasing complexity, which has helped us to gain strong insights into the empirical relationship between the LR schedule and the loss curve (Section 2).

2 Empirical Derivation of the Multi-Power Law

In this section, we present our empirical derivation of the multi-power law for schedule-aware loss curve prediction. In the first place, we reduce the problem to studying a loss reduction term led by LR decay. Then we take a “bottom-up” approach to study this term for LR schedules with increasing complexity, from two-stage, multi-stage, to general LR decay schedules. For the first two cases, we conduct extensive ablation studies on the behavior of the training loss and derive formulas that can accurately predict the loss reduction term. This has finally inspired us to propose the multi-power law for general cases as a natural unification and generalization of the formulas derived for the two special cases. We will further validate our law with extensive experiments in Section 3.

Background: Learning Rate Schedule. An LR schedule is a sequence $E := \{\eta_1, \dots, \eta_T\}$ that specifies the LR at each step of the training process. In the domain of language model pretraining, the cosine LR schedule (Loshchilov & Hutter, 2017) is the most popular one, which can be expressed as $\eta_t = \frac{1+\alpha}{2}\eta_{\max} + \frac{1-\alpha}{2}\eta_{\max}\cos(\frac{\pi t}{T})$, where $\eta_{\max}$ is the maximum LR and α is usually set to 0.1. The Warmup-Stable-Decay (WSD) schedule (Hu et al., 2024) is a recently proposed LR schedule. This schedule first goes through a warmup phase with W steps, then maintains at a stable LR $\eta_{\max}$ with T_{stable} steps, and finally decays in the form $f(s - T_{\text{stable}})\eta_{\max}$ during stage $T_{\text{stable}} \leq s \leq T_{\text{total}}$. Here $f(x) \in (0, 1)$ can be chosen as linear or exponential decay functions. The visualization of these two LR schedules is in Figure 1(a).

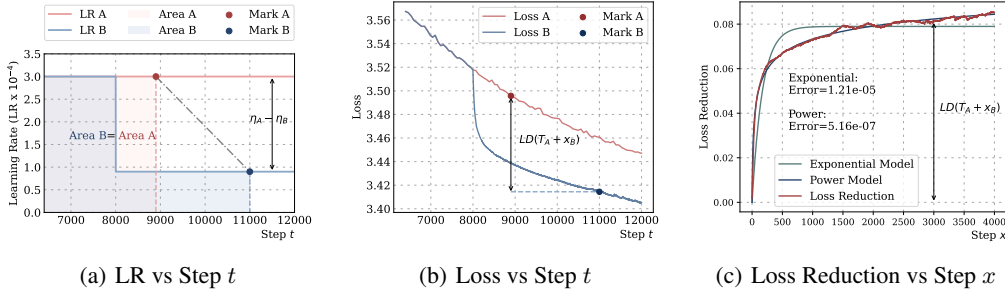


Figure 3: Example of Two-Stage cases: $t_B=11000$, $x_B=3000$, $\eta_B = 9 \times 10^{-5}$, $\eta_A = 3 \times 10^{-4}$, $T_A = 8000$. (a) A and B have the equal LR sums: $x_A = 900$, $t_A = 8900$. (b) Loss Reduction $LD(T_A + x_B) = L_A(t_A) - L_B(t_B)$. (c) Fitting Loss Reduction $\hat{LD}(T_A + x_B)$ with power form results in $0.13(1 - (1 + 0.21x)^{0.15})$; Fitting with exponential form results in $0.0790(1 - e^{-0.01x})$. The shape of loss reduction is closer to a power form instead of exponential.

Background: Warmup Phase. Many LR schedules, such as WSD, contain a warmup phase that increases the LR gradually from 0 to the maximum LR $\eta_{\max}$ over a few steps. Our discussion focuses on the training after the warmup, where the LR is non-increasing in almost all LR schedules. The steps are counted after the warmup phase, i.e., $t = 1$ is the first step after the warmup.

2.1 Our Approach: Learning Rate Sum Matching

Auxiliary Training Process. We first introduce an auxiliary training process to aid our analysis of the loss curve of the actual training process with LR schedule $E := \{\eta_1, \dots, \eta_T\}$. This auxiliary training process is exactly the same as the actual training process for the first K steps, where K is the largest number such that $\eta_1 = \eta_2 = \dots = \eta_K$. Then the auxiliary training process continues training with a constant LR schedule, where the LR is set to η_1 for all the remaining steps. We denote the training loss at step t in this auxiliary process as $\mathcal{L}_{\text{const}}(t)$.

Learning Rate Sum Matching. The multi-power law for approximating the loss curve $\mathcal{L}(t)$ of the actual training process is based on the following decomposition. Define $Z(t)$ as the step in the auxiliary process that has the same sum of LR as the actual training process at step t . Then,

$$\mathcal{L}(t) = \mathcal{L}_{\text{const}}(Z(t)) - \underbrace{(\mathcal{L}_{\text{const}}(Z(t)) - \mathcal{L}(t))}_{=: LD(t)}, \quad \text{where} \quad Z(t) := \frac{1}{\eta_1} \sum_{\tau=1}^t \eta_{\tau}. \quad (3)$$

Here, we first use the training loss at step $Z(t)$ in the auxiliary process, $\mathcal{L}_{\text{const}}(Z(t))$, as an approximation for $\mathcal{L}(t)$, and then write the approximation error term as $LD(t)$. We call $LD(t)$ the *Loss reDuction term* as it is a quantification of the reduction of loss due to learning rate decay.

The rationale behind this approach is that matching the LR sum between the two training processes should result in similar training losses, and thus a more accurate approximation can be obtained by further exploring the loss reduction term $LD(t)$. See Appendix D.2 for more discussion.

Power-Law Ansatz for the Auxiliary Loss. Under constant LR schedules, it is easy to predict the loss curve accurately. Taking inspiration from previous works (Hoffmann et al., 2022; Kaplan et al., 2020), we choose to take a power-law form to approximate the training loss of the auxiliary process, which works reasonably well in our experiments. That is,

$$\mathcal{L}_{\text{const}}(t) \approx \hat{\mathcal{L}}_{\text{const}}(t) := L_0 + \tilde{A} \cdot t^{-\alpha}, \quad (4)$$

where $L_0, \tilde{A}, \alpha$ are parameters. Replacing t with $Z(t) := \frac{1}{\eta_1} \sum_{\tau=1}^t \eta_{\tau} = \frac{1}{\eta_1} S_1(t)$ gives $\hat{\mathcal{L}}_{\text{const}}(Z(t)) = L_0 + \tilde{A} \eta_1^\alpha S_1^{-\alpha}(t)$, where $S_1(t) := \sum_{\tau=1}^t \eta_{\tau}$. Finally, we reparameterize $\tilde{A}$ as $\tilde{A} := A \eta_1^{-\alpha}$, where A is a parameter, and obtain the formula:

$$\hat{\mathcal{L}}_{\text{const}}(Z(t)) = L_0 + A \cdot S_1^{-\alpha}(t). \quad (5)$$

See Figure 15 for an empirical validation of eq. (5) for various constant LR schedules.

Loss Reduction Term. It remains to understand the behavior of the loss reduction term $LD(t)$, which is inherently complex since it depends on each LR used in training. In the rest of the section, we conduct a series of experiments to determine an accurate approximation form for $LD(t)$.

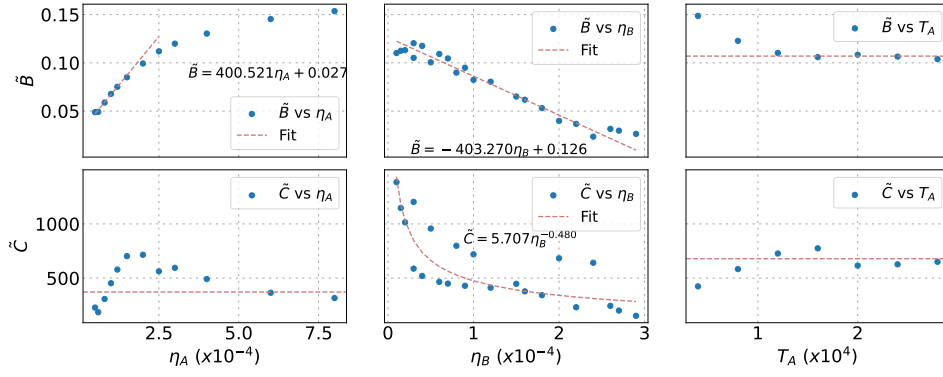


Figure 4: The patterns of $\tilde{B}$, $\tilde{C}$ over η_A , η_B and T_A in Two-Stage cases. For illustration purposes, the second row uses $\tilde{C}$ as the y-axis. $\tilde{B}$ can be approximated to be proportional to $\eta_A - \eta_B$, and $\tilde{C}$ manifests power-law pattern over η_B . The effect of η_A on $\tilde{C}$ and the impact of T_A are unpredictable or negligible, which is approximately ignored in our discussion.

2.2 Case 1: Two-stage Learning Rate Schedule

To understand the behavior of $LD(t)$, we start with the simplest form of LR decay that consists of two stages: In Stage 1, the LR remains constant at η_A for T_A steps ($\eta_1 = \eta_2 = \dots = \eta_{T_A} = \eta_A$); in Stage 2, the LR suddenly decreases to η_B and the rest of training continues with η_B for T_B steps ($\eta_{T_A+1} = \eta_{T_A+2} = \dots = \eta_{T_A+T_B} = \eta_B$). We call this LR schedule a *two-stage LR schedule*. In this case, the first T_A steps of the auxiliary and actual training processes are the same, and the loss reduction term $LD(t)$ becomes non-zero only in Stage 2.

The Loss Reduction Term Follows a Power Law. In Figure 3, we plot the loss reduction term $LD(t)$ of a two-stage learning rate schedule with $\eta_A = 3 \times 10^{-4}$, $T_A = 8000$, $\eta_B = 9 \times 10^{-5}$. As the number of steps $x := t - T_A$ in Stage 2 increases, $LD(T_A + x)$ monotonically rises from 0 to around 0.09 and eventually saturates. This motivates us to approximate $LD(T_A + x)$ in the form $\tilde{B} \cdot (1 - U(\eta_B x))$, where $\tilde{B}$ is a parameter and $U(s)$ is a function that decreases from 1 to 0 as $s = \eta_B x$ increases from 0 to infinity. The reason we choose $\eta_B x$ instead of x as the argument of U will be clear when we generalize this to multi-stage schedules.

But at what rate should $U(s)$ decrease? After trying different forms of $U(s)$ to fit $LD(T_A + x)$, we find that the power-law form $U(s) = (\tilde{C} \cdot s + 1)^{-\beta}$ for some $\tilde{C}, \beta > 0$ fits most properly, which leads to the following power-law form for the loss reduction term:

$$LD(T_A + x) \approx \widehat{LD}(T_A + x) := \tilde{B}(1 - (\tilde{C} \cdot \eta_B x + 1)^{-\beta}). \quad (6)$$

Figure 3(c) shows that this power law aligns well with the actual loss reduction term $LD(T_A + x)$. In contrast, the exponential form $U(s) = e^{-Bs}$ (so $LD(T_A + x) \approx A(1 - e^{-B\eta_B x})$) struggles to match the slow and steadily increase of $LD(T_A + x)$ when x is large.

We further investigate how to estimate the parameters $\tilde{B}, \tilde{C}, \beta$ in the power law. Our preliminary experiments suggest that the power law fits the loss reduction term very well with a constant β that is independent of η_A, η_B, T_A , so we just set $\beta = 0.4$, which is a constant that works well. Then we conduct experiments to understand how the best parameters $\tilde{B}, \tilde{C}$ to fit $LD(t)$ depend on η_A, η_B, T_A , where we set default values $\eta_A = 3 \times 10^{-4}$, $\eta_B = 3 \times 10^{-5}$, $T_A = 8000$ and change one variable at a time. See Appendix E for experiment details.

$\tilde{B}$ is Linear to LR Reduction. Our first observation is that $\tilde{B} \propto \eta_A - \eta_B$. As shown in the first row of Figure 4, $\tilde{B}$ linearly decreases with η_B and approximately increases linearly with η_A , especially when η_A is not too large. Moreover, the slope of $\tilde{B}$ over η_A and η_B are approximately opposite to each other. This motivates us to hypothesize that $\tilde{B} \propto \eta_A - \eta_B$ and reparameterize $\tilde{B}$ as $\tilde{B} = B(\eta_A - \eta_B)$, where B is a constant.

$\tilde{C}$ Follows a Power Law of η_B . Our second observation is that $\tilde{C}$ follows a power law. As shown in the second row of Figure 4, we find that $\tilde{C}$ is very sensitive to η_B but much less dependent on η_A . We hypothesize that $\tilde{C}$ follows a power law $\tilde{C} \propto \eta_B^{-\gamma}$, and reparameterize $\tilde{C}$ as $\tilde{C} = C\eta_B^{-\gamma}$, where $C > 0$ and $\gamma > 0$ are constants.

LR Reduction Term Depends Less on T_A . We also find that $\tilde{B}$ and $\tilde{C}$ are less sensitive to T_A . As shown in the last column in Figure 4, $\tilde{B}$ and $\tilde{C}$ are relatively stable as T_A varies. This suggests that

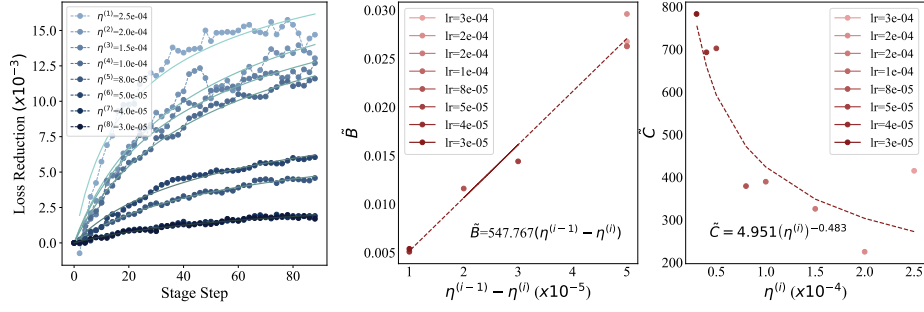


Figure 5: **Left:** The loss reduction $LD_k(t)$ between the adjacent stages of the multi-stage training loss curve still follows the power form. The multi-stage loss curve refers to Figure 14. **Right:** The parameter patterns in the two-stage setting hold in the multi-stage setting approximately. $\tilde{B} \propto \eta^{(i-1)} - \eta^{(i)}$ and $\tilde{C} \propto \eta^i$ keep and the shape of patterns are similar to the patterns in our ablation experiments, as shown in Figure 4.

the loss reduction induced by a single LR decay is approximately independent from the time that LR decays. This insight will be revisited when we generalize the two-stage case to multi-stage cases.

Final Approximation Form. Putting all these together, we have the final approximation form for the loss reduction term in the two-stage schedule:

$$LD(T_1 + x) \approx \widehat{LD}(T_1 + x) := B(\eta_A - \eta_B)(1 - (C\eta_B^{1-\gamma}x + 1)^{-\beta}). \quad (7)$$

2.3 Case 2: Multi-Stage Learning Rate Schedule

We go one step further from two-stage step decay schedule to multi-stage step decay schedules. This class of schedules consists of multiple stages, where the LR decays when a new stage starts but remains constant within each stage. Now, we consider an n -stage LR schedule $E = \{\eta_1, \dots, \eta_T\}$, where the i -th stage lasts from step $t = T_{i-1} + 1$ to $t = T_i$ and uses the LR $\eta_t = \eta^{(i)}$ ($0 = T_0 < T_1 < \dots < T_{n-1} < T_n = T$, $\eta^{(1)} \geq \eta^{(2)} \geq \dots \geq \eta^{(n)}$). See Figure 14 for an example.

Multi-Stage Loss Reduction. To draw insights into the behavior of $LD(t)$ in the multi-stage case, we use the following strategy. Recall that $LD(t)$ is the difference in training losses between the auxiliary and actual training processes at equal LR sums. In addition to these processes, we construct some intermediate processes: for $1 \leq i \leq n$, we define the i -th process to be the same as the actual training process in stages 1 to i but continue to use the learning rate $\eta^{(i)}$ for all stages after i . The first and last processes are the auxiliary and actual training processes themselves.

We again use the trick of LR sum matching: we find the steps of these n processes that have the same LR sum, and then conduct experiments to analyze the loss difference between adjacent processes. Let $\mathcal{L}_i(t)$ be the training loss of the i -th process at step t . For $1 \leq i \leq j \leq n$ and $t \geq T_{j-1}$, we define $Z_{i,j}(t)$ as the step number in the i -th process that has the same LR sum as the j -th process at step t , i.e., $Z_{i,i}(t) := t$ and $Z_{i,j}(t) := T_i + \frac{1}{\eta^{(i)}} \sum_{\tau=T_i+1}^t \eta_\tau$ for $i < j$. Then we have

$$LD(t) = \sum_{k=2}^i LD_k(Z_{k,n}(t)), \quad \text{where} \quad LD_k(t) := \mathcal{L}_{k-1}(Z_{k-1,k}(t)) - \mathcal{L}_k(t). \quad (8)$$

Here, $LD_k(t)$ is the difference between the $(k-1)$ -th and k -th processes at equal LR sums. These two processes are the same for the first $k-1$ stages and diverge only at the beginning of the k -th stage: the former continues to use $\eta^{(k-1)}$ but the latter switches to $\eta^{(k)}$. This is similar to the two-stage case, except that the first $k-1$ stages may not use the same LR.

Interestingly, the power law behavior of the loss reduction term observed in the two-stage case also approximately holds for $LD_k(t)$. As the training enters a new stage $i+1$, a new loss reduction term $LD_i(\cdot)$ is introduced in (8). We observe that $LD_i(T_i + x)$ follows a similar power law behavior as in the two-stage case when x increases. As shown in Figure 5, each $LD_i(T_i + x)$ can be individually approximated by a power law $\tilde{B}(1 - (\tilde{C} \cdot \eta^{(i)}x + 1)^{-\beta})$, and $\tilde{B} \propto (\eta^{(i-1)} - \eta^{(i)})$, $\tilde{C} \propto (\eta^{(i)})^{-\gamma}$.

Final Approximation Form. Inspired by the above observation and our approximation (7) for the two-stage case, we propose to approximate $LD_k(t)$ with a power law:

$$LD_k(T_{k-1} + x) \approx \widehat{LD}_k(T_{k-1} + x) := B(\eta^{(k-1)} - \eta^{(k)})(1 - (C(\eta^{(k)})^{1-\gamma}x + 1)^{-\beta}). \quad (9)$$

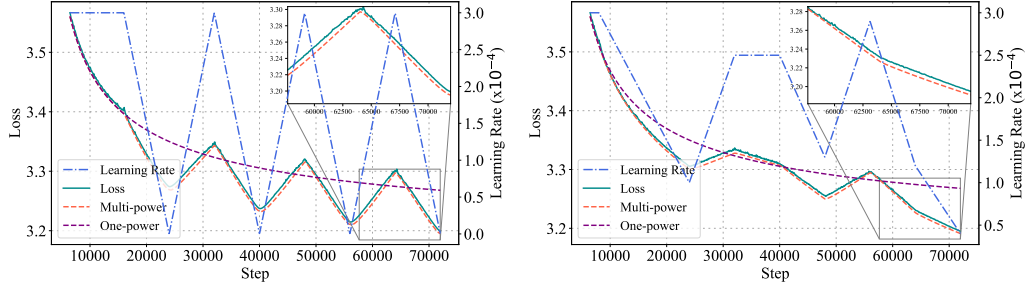


Figure 6: The examples of long-horizon non-monotonic schedules. The one-power line represents the auxiliary process curve. **Left:** The cyclic schedule with 72,000 steps, where each half-cycle spans 8,000 steps, and the first decay begins after 16,000 steps. **Right:** The random-polyline schedule, consisting of piecewise linear interpolation between randomly selected intermediate learning rates in the range of 3×10^{-5} to 3×10^{-4} , with LR milestones occurring at intervals of 8,000 steps.

Plugging this approximation into (8), we can approximate the loss reduction term $\text{LD}(t)$ at step $T_{i-1} < t \leq T_i$ of the actual training process as $\text{LD}(t) \approx \widehat{\text{LD}}(t) := \sum_{k=2}^i \widehat{\text{LD}}_k(Z_{k,n}(t))$, then

$$\text{LD}(t) \approx \widehat{\text{LD}}(t) = \sum_{k=2}^i B(\eta^{(k-1)} - \eta^{(k)})(1 - (C(\eta^{(k)})^{1-\gamma}(Z_{k,n}(t) - T_k) + 1)^{-\beta}).$$

By definition of $Z_k(t)$, we have $Z_{k,n}(t) - T_k = \frac{S_k(t)}{\eta^{(k)}}$, where $S_k(t) := \sum_{\tau=T_k+1}^t \eta_\tau$ is the sum of LRs from the beginning of Stage $k+1$ to step t . We can further simplify the above formula as

$$\text{LD}(t) \approx \widehat{\text{LD}}(t) = \sum_{k=2}^i B(\eta^{(k-1)} - \eta^{(k)})(1 - (C(\eta^{(k)})^{-\gamma} S_k(t) + 1)^{-\beta}). \quad (10)$$

2.4 General Case

Ansatz for the Loss Reduction Term. A general learning rate schedule $E := \{\eta_1, \dots, \eta_T\}$ could be viewed as a T -stage schedule, where the i -th stage uses learning rate η_i for $l_i = 1$ step. This motivates us to make the ansatz that the formula for the loss reduction term $\text{LD}(t)$ in multi-stage schedules (10) can continue to hold even when every stage only lasts for one step.

$$\text{LD}(t) \approx \widehat{\text{LD}}(t) = \sum_{k=2} B(\eta_{k-1} - \eta_k)(1 - (C\eta_k^{-\gamma} S_k(t) + 1)^{-\beta}), \quad (11)$$

where $S_k(t) := \sum_{\tau=k}^t \eta_\tau$ is the sum of LRs from step k to step t , and B, C, γ, β are parameters.

Multi-Power Law. Following the approach of learning rate sum matching in Section 2.1, we first decompose $\mathcal{L}(t)$ as $\mathcal{L}_{\text{const}}(Z(t)) - \text{LD}(t)$ (see (3)). Then we combine the above ansatz for the Loss reduction term with the power-law ansatz for the auxiliary loss, leading to our multi-power law:

$$\mathcal{L}(t) \approx L_0 + A \cdot S_1^{-\alpha}(t) - \sum_{k=2} B(\eta_{k-1} - \eta_k)(1 - (C\eta_k^{-\gamma} S_k(t) + 1)^{-\beta}).$$

See also Appendix D.3 for the practical version of our law that accounts for the warmup phase by slightly changing the $A \cdot S_1^{-\alpha}(t)$ term. See Appendix B.1 for the discussion about the simplification of the multi-power law.

3 Empirical Validation of the Multi-Power Law

The multi-power law (MPL) comes from our speculations from our experiments with special types of LR schedules. Now we present extensive experiments to validate the law for common LR schedules used in practice. Our experiments demonstrate that MPL requires only two or three learning rate (LR) schedules and their corresponding loss curves in the training set to fit the law. The fitted MPL can then predict loss curves for test schedules with different shapes and extended horizons. Details of the experimental setup, fitting approaches, and configurations are provided in Appendix F.

3.1 Results

Generalization to Unseen LR Schedules. The MPL can accurately predict loss curves for LR schedules outside the training set. As illustrated in Figure 2, despite the absence of WSD LR schedules in the training set and the variety of decay functions, MPL successfully predicts their loss curves with high accuracy. Furthermore, MPL can generalize to two-stage schedules with different η_B values from the training set, effectively extrapolating for both continuous and discontinuous cases.

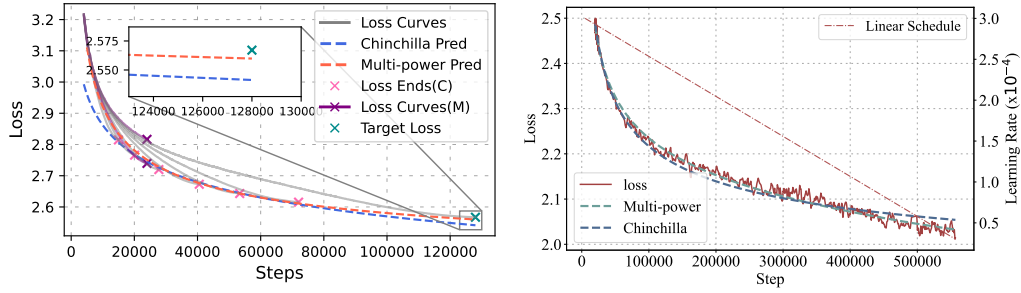


Figure 7: **Left:** Predictions for target loss at 128,000-step for cosine schedule using MPL and CDSL fitting. The CDSL uses the final losses of six cosine losses from 14960 steps to 72000 steps, marked as Loss Ends(C). The MPL uses 24000-step constant and cosine curves, marked as Loss Curves(M). **Right:** Comparison of MPL and CDSL fits on the open-source 7B OLMo curve generated with a linear schedule.

Generalization to Longer Horizons. MPL demonstrates the ability to extrapolate loss curves for horizons exceeding three times the training set length. In our runs, the training set contains approximately 22,000 post-warmup steps, while the test set includes curves with up to 70,000 post-warmup steps. These results validate MPL’s capability to generalize to longer horizons. Notably, the data-to-model ratio for a 25M-parameter model trained over 72,000 steps (36B tokens) is comparable to Llama2 pretraining (70B model, 2T tokens), consistent with trends favoring higher data volumes for fixed model sizes (Dubey et al., 2024).

Generalization to Non-Monotonic Schedules. MPL extends effectively to complex non-monotonic schedules, although derived for monotonic decay schedules. The test set includes challenging cases such as cyclic schedules and the *random-polyline schedule*, where LR values are randomly selected at every 8,000 steps and connected by a polyline. These experiments, conducted on a 25M-parameter model over 72,000 steps, also represent a demanding long-horizon scenario. As shown in Figure 6, MPL accurately predicts these long-horizon non-monotonic schedules, demonstrating its robustness and adaptability.

3.2 Comparison with Baselines

Comparison with Chinchilla Law. While Chinchilla-style data scaling laws, which we abbreviate as CDSLs, are widely utilized (Muennighoff et al., 2023; Hoffmann et al., 2022), MPL offers several distinct advantages: (1) MPL incorporates LR dependency, unlike CDSLs, and (2) MPL predicts the entire loss curve, whereas CDSLs are restricted to final loss predictions. Based on these advantages, the MPL shows higher sample efficiency than the CDSLs. Moreover, we find that two curves of different schedules are enough to fit the MPL with generalizability, as details are discussed in Appendix B.3. As shown in Figure 7, MPL achieves 1/4 error in final loss prediction with 1/4 compute budget compared to CDSL. MPL also shows advantages in the fitting of open-source 7B OLMo (Groeneveld et al., 2024) in Figure 7.

Comparison with Momentum Law. The MPL shows higher accuracy and can apply to the discontinuous schedules compared to the recent Momentum Law (Tissue et al., 2024). The Momentum Law (MTL) (Tissue et al., 2024) incorporates LR annealing effects by modeling loss reduction based on the momentum of LR decay. However, MTL indicates an exponential loss reduction in two-stage LR schedules, which contradicts our observations (see Figure 3). Additionally, as shown in Figure 9, MPL outperforms MTL in predicting loss reduction for WSD schedules with linear LR decay. In the highlighted regions, MPL achieves high accuracy in the decay stage, whereas MTL exhibits substantial error. A summary of prediction results across test sets is provided in Table 1, where MPL consistently outperforms MTL in both average and worst-case scenarios. The details of the MTL and its relation to the MPL can be found in Appendix B.1.

4 The Multi-power Law Induces Better LR Schedules

Due to the high cost of each pretraining run and the curse of dimensionality for LR schedules, it is generally very impossible to tune the LR for every training step. However, in this section, we show that by using the predicted final loss from the MPL, we can optimize the entire LR schedule to significantly reduce the final loss and beat the cosine schedule.

Table 1: Model performance comparison. R^2 , MAE, RMSE, PredE, and WorstE are the coefficient of determination, Mean Absolute Error, Root Mean Square Error, Prediction Error, and Worst-case Error, respectively.

Model Size	Method	$R^2 \uparrow$	MAE $\downarrow$	RMSE $\downarrow$	PredE $\downarrow$	WorstE $\downarrow$
25M	Momentum Law	0.9904	0.0047	0.0060	0.0014	0.0047
	Multi-power Law (Ours)	0.9975	0.0039	0.0046	0.0012	0.0040
100M	Momentum Law	0.9959	0.0068	0.0095	0.0022	0.0094
	Multi-power Law (Ours)	0.9982	0.0038	0.0051	0.0013	0.0058
400M	Momentum Law	0.9962	0.0071	0.0094	0.0025	0.0100
	Multi-power Law (Ours)	0.9971	0.0053	0.0070	0.0019	0.0070

4.1 Method

Given that the Multi-Power Law (MPL) provides an accurate estimation of the loss, the final loss prediction by MPL can serve as a surrogate for evaluating schedules. Consider the learning rate (LR) as a T -dimensional vector $\boldsymbol{\eta} = (\eta_1, \dots, \eta_T)$ and the final loss $\mathcal{L}(\boldsymbol{\eta})$ under given hyperparameters. The goal is to identify the optimal LR schedule $\boldsymbol{\eta}^* = \arg \max_{\boldsymbol{\eta}} \mathcal{L}(\boldsymbol{\eta})$. We parameterize the final loss prediction as $\mathcal{L}_{\Theta}(\boldsymbol{\eta})$ using MPL with parameters $\Theta = \{L_0, A, B, C, \alpha, \beta, \gamma\}$. The parameters $\hat{\Theta}$ can be estimated as described in Section 3. Using $\mathcal{L}_{\hat{\Theta}}(\boldsymbol{\eta})$ as a surrogate for $\mathcal{L}(\boldsymbol{\eta})$, we approximate $\boldsymbol{\eta}^*$ by solving:

$$\hat{\boldsymbol{\eta}} = \min_{\boldsymbol{\eta}} \mathcal{L}_{\hat{\Theta}}(\boldsymbol{\eta}) \quad \text{s.t.} \quad 0 \leq \eta_{i+1} \leq \eta_i, \forall 1 \leq i \leq T-1. \quad (12)$$

This process induces an ‘‘optimal’’ schedule $\hat{\boldsymbol{\eta}}$ derived from MPL with parameter $\hat{\Theta}$. We set the initial learning rate η_0 to 3×10^{-4} and assume η_i is monotonically non-increasing based on prior knowledge. The high-dimensional vector $\boldsymbol{\eta}$ is optimized using the Adam optimizer. Additional details are provided in Appendix G.

4.2 Results

Induced LR Schedule Exhibits Stable-Decay Behavior. The induced learning rate schedule follows a Warmup-Stable-Decay (WSD) pattern, comprising two main stages after the warmup phase. It maintains a peak LR for an extended period, followed by a rapid decay to a near-zero LR, as shown in Figure 1 and Figure 19.

Induced LR Schedule Outperforms Cosine Schedule. Figures 1 and 19 compare the induced schedules with the cosine and WSD schedules across models ranging from 25M to 400M. Figure 20 extends this comparison to longer training horizons. The induced schedules consistently outperform the cosine schedule, achieving a margin over 0.02. Notably, no WSD-like schedule is present in the training set, predicting such loss curves an extrapolation by MPL.

Characteristics of the Induced Schedules. The induced schedules provide insights into hyperparameter tuning for WSD schedules. Observations from Figures 1 and 19 highlight the following: (1) Our findings suggest that a lower ending LR—typically below $1/20$ of the peak LR—is more effective in most scenarios, compared to $1/10$ in prior research (Hoffmann et al., 2022; Kaplan et al., 2020). Further details are provided in Appendix G. (2) $f(x) = (1-x)^{-\alpha}$, where $\alpha \approx 1.5$, well captures relationship between normalized steps $\tilde{t}$ and normalized LR $\tilde{\eta}_{\text{avg}}$ in our experiments. This simplified version, referred to as WSD with Sqrt-Cube Decay (WSDSC), is effective across various model sizes and types, as shown in Figures 8 and 10. See Appendix B.2. (3) The induced schedules align closely with the optimal decay steps identified via grid search, as illustrated in Figure 1. See Appendix G.

5 Conclusions and Future Directions

In this paper, we introduce the multi-power law for scheduler-aware loss curve prediction, which can accurately predict loss curves and inspire optimal scheduler derivation. Our findings enhance the understanding of training dynamics in large language models, potentially improving training efficiency. In future work, we will consider refining the law, exploring its underlying mechanisms, and studying the LR relationship with unfixed maximum LR.

References

- Armen Aghajanyan, Lili Yu, Alexis Conneau, Wei-Ning Hsu, Karen Hambardzumyan, Susan Zhang, Stephen Roller, Naman Goyal, Omer Levy, and Luke Zettlemoyer. Scaling laws for generative mixed-modal language models. In *International Conference on Machine Learning*, pp. 265–279. PMLR, 2023.
- Alexander Atanasov, Jacob A Zavatone-Veth, and Cengiz Pehlevan. Scaling and renormalization in high-dimensional regression. *arXiv preprint arXiv:2405.00592*, 2024.
- Yasaman Bahri, Ethan Dyer, Jared Kaplan, Jaehoon Lee, and Utkarsh Sharma. Explaining neural scaling laws. *Proceedings of the National Academy of Sciences*, 121(27):e2311878121, 2024.
- Yoshua Bengio. *Practical Recommendations for Gradient-Based Training of Deep Architectures*, pp. 437–478. Springer Berlin Heidelberg, Berlin, Heidelberg, 2012. ISBN 978-3-642-35289-8. doi: 10.1007/978-3-642-35289-8_26. URL https://doi.org/10.1007/978-3-642-35289-8_26.
- James Bergstra, Dan Yamins, David D Cox, et al. Hyperopt: A python library for optimizing the hyperparameters of machine learning algorithms. *SciPy*, 13:20, 2013.
- Blake Bordelon, Alexander Atanasov, and Cengiz Pehlevan. A dynamical model of neural scaling laws. *arXiv preprint arXiv:2402.01092*, 2024.
- Xiang Cheng, Dong Yin, Peter Bartlett, and Michael Jordan. Stochastic gradient and langevin processes. In *International Conference on Machine Learning*, pp. 1810–1819. PMLR, 2020.
- Zhengxiao Du, Aohan Zeng, Yuxiao Dong, and Jie Tang. Understanding emergent abilities of language models from the loss perspective, 2024. URL <https://arxiv.org/abs/2403.15796>.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- Omer ElKabatz and Nadav Cohen. Continuous vs. discrete optimization of deep neural networks. *Advances in Neural Information Processing Systems*, 34:4947–4960, 2021.
- Jonas Geiping and Tom Goldstein. Cramming: Training a language model on a single gpu in one day. In *International Conference on Machine Learning*, pp. 11117–11143. PMLR, 2023.
- Dirk Groeneveld, Iz Beltagy, Pete Walsh, Akshita Bhagia, Rodney Kinney, Oyvind Tafjord, Ananya Harsh Jha, Hamish Ivison, Ian Magnusson, Yizhong Wang, et al. Olmo: Accelerating the science of language models. *arXiv preprint arXiv:2402.00838*, 2024.
- Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, Tom Hennigan, Eric Noland, Katie Millican, George van den Driessche, Bogdan Damoc, Aurelia Guy, Simon Osindero, Karen Simonyan, Erich Elsen, Jack W. Rae, Oriol Vinyals, and Laurent Sifre. Training compute-optimal large language models, 2022. URL <https://arxiv.org/abs/2203.15556>.
- Shengding Hu, Yuge Tu, Xu Han, Chaoqun He, Ganqu Cui, Xiang Long, Zhi Zheng, Yewei Fang, Yuxiang Huang, Weilin Zhao, et al. Minicpm: Unveiling the potential of small language models with scalable training strategies. *arXiv preprint arXiv:2404.06395*, 2024.
- Peter J Huber. Robust estimation of a location parameter. In *Breakthroughs in statistics: Methodology and distribution*, pp. 492–518. Springer, 1992.
- Frank Hutter, Holger H Hoos, and Kevin Leyton-Brown. Sequential model-based optimization for general algorithm configuration. In *Learning and Intelligent Optimization: 5th International Conference, LION 5, Rome, Italy, January 17-21, 2011. Selected Papers 5*, pp. 507–523. Springer, 2011.
- Marcus Hutter. Learning curve theory. *arXiv preprint arXiv:2102.04074*, 2021.

- Ayush Jain, Andrea Montanari, and Eren Sasoglu. Scaling laws for learning with real and surrogate data. *arXiv preprint arXiv:2402.04376*, 2024.
- Yuchen Jin, Tianyi Zhou, Liangyu Zhao, Yibo Zhu, Chuanxiong Guo, Marco Canini, and Arvind Krishnamurthy. Autolrs: Automatic learning-rate schedule by bayesian optimization on the fly. *arXiv preprint arXiv:2105.10762*, 2021.
- Arlind Kadra, Maciej Janowski, Martin Wistuba, and Josif Grabocka. Scaling laws for hyperparameter optimization. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023. URL <https://openreview.net/forum?id=ghzEUGfRMD>.
- Arlind Kadra, Maciej Janowski, Martin Wistuba, and Josif Grabocka. Scaling laws for hyperparameter optimization. *Advances in Neural Information Processing Systems*, 36, 2024.
- 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, 2020. URL <https://arxiv.org/abs/2001.08361>.
- Aaron Klein, Stefan Falkner, Jost Tobias Springenberg, and Frank Hutter. Learning curve prediction with bayesian neural networks. In *International conference on learning representations*, 2022.
- Lisha Li, Kevin Jamieson, Giulia DeSalvo, Afshin Rostamizadeh, and Ameet Talwalkar. Hyperband: A novel bandit-based approach to hyperparameter optimization. *Journal of Machine Learning Research*, 18(185):1–52, 2018.
- Qianxiao Li, Cheng Tai, and E Weinan. Stochastic modified equations and adaptive stochastic gradient algorithms. In *International Conference on Machine Learning*, pp. 2101–2110. PMLR, 2017.
- Zhiyuan Li and Sanjeev Arora. An exponential learning rate schedule for deep learning. *arXiv preprint arXiv:1910.07454*, 2019.
- Licong Lin, Jingfeng Wu, Sham M Kakade, Peter L Bartlett, and Jason D Lee. Scaling laws in linear regression: Compute, parameters, and data. *arXiv preprint arXiv:2406.08466*, 2024.
- Ilya Loshchilov and Frank Hutter. SGDR: Stochastic gradient descent with warm restarts. In *International Conference on Learning Representations*, 2017. URL <https://openreview.net/forum?id=Skq89Scxx>.
- Chao Ma, Lei Wu, and Weinan E. A qualitative study of the dynamic behavior for adaptive gradient algorithms. In Joan Bruna, Jan Hesthaven, and Lenka Zdeborova (eds.), *Proceedings of the 2nd Mathematical and Scientific Machine Learning Conference*, volume 145 of *Proceedings of Machine Learning Research*, pp. 671–692. PMLR, 16–19 Aug 2022. URL <https://proceedings.mlr.press/v145/ma22a.html>.
- Niklas Muennighoff, Alexander Rush, Boaz Barak, Teven Le Scao, Nouamane Tazi, Aleksandra Piktus, Sampo Pyysalo, Thomas Wolf, and Colin A Raffel. Scaling data-constrained language models. *Advances in Neural Information Processing Systems*, 36:50358–50376, 2023.
- Rui Pan, Haishan Ye, and Tong Zhang. Eigencurve: Optimal learning rate schedule for sgd on quadratic objectives with skewed hessian spectrums. *arXiv preprint arXiv:2110.14109*, 2021.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- Utkarsh Sharma and Jared Kaplan. A neural scaling law from the dimension of the data manifold. *arXiv preprint arXiv:2004.10802*, 2020.
- Leslie N Smith. Cyclical learning rates for training neural networks. In *2017 IEEE winter conference on applications of computer vision (WACV)*, pp. 464–472. IEEE, 2017.
- Jasper Snoek, Hugo Larochelle, and Ryan P Adams. Practical bayesian optimization of machine learning algorithms. *Advances in neural information processing systems*, 25, 2012.

- Yunfei Teng, Jing Wang, and Anna Choromanska. Autodrop: Training deep learning models with automatic learning rate drop. *arXiv preprint arXiv:2111.15317*, 2021.
- Howe Tissue, Venus Wang, and Lu Wang. Scaling law with learning rate annealing. *arXiv preprint arXiv:2408.11029*, 2024.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- A Vaswani. Attention is all you need. *Advances in Neural Information Processing Systems*, 2017.
- Zhen Xu, Andrew M Dai, Jonas Kemp, and Luke Metz. Learning an adaptive learning rate schedule. *arXiv preprint arXiv:1909.09712*, 2019.

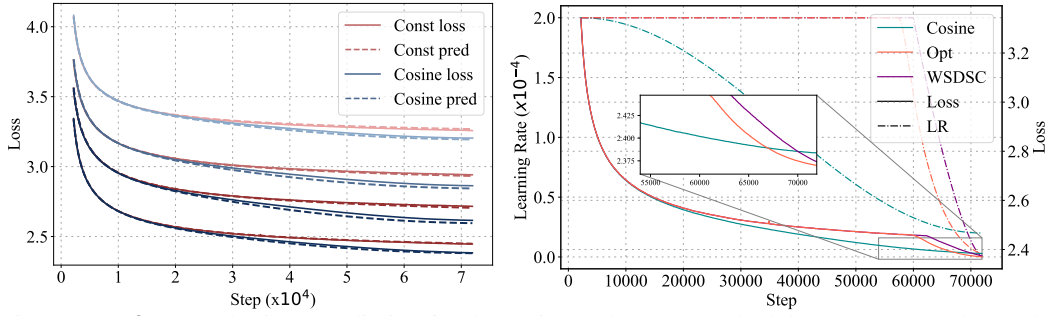


Figure 8: **Left:** Long horizon prediction for the cosine and constant schedules. From up to down, the model sizes range from 25M to 1B. **Right:** The comparison on 1B models between the optimized schedule (Opt), cosine schedule (Cosine), and the simplified optimized schedule (WSDSC, see Section 4.2), a WSD schedule with sqrt-cube decay.

Table 2: Downstream performance comparison for Cosine and induced schedules. Percentage changes ($\uparrow$ or $\downarrow$) indicate relative improvements or regressions compared to the Cosine schedule.

Downstream Dataset	LAMBADA	HellaSwag	PIQA	ARC-E
Cosine Schedule	46.54	37.12	65.13	43.56
Induced Schedule	48.71 ($\uparrow$ 2.17%)	37.74 ($\uparrow$ 0.62%)	65.07 ($\downarrow$ 0.06%)	44.09 ($\uparrow$ 0.53%)

A Discussion

Model Types. We validate the MPL on GPT-2 (Radford et al., 2019) and OLMo (Groeneveld et al., 2024) models to evaluate the generalizability of the MPL across model architectures. In the preceding experiments, we used the Llama2 (Touvron et al., 2023). For experiments on GPT-2, the validation process followed the procedure fit with curves of cosine and constant schedules, described in Section 3. For the 7B OLMo model, we fit the MPL on the open-source training curve, which employs a linear decay schedule, as shown on the right of Figure 7. Our results show that the MPL presents a high prediction accuracy across different model types for both self-run and open-source experiments.

Model Size. We extended the MPL and its induced schedule to a larger scale by training a 1B-parameter model on 144B data tokens. The MPL was fitted over 24,000 steps and successfully predicted loss curves up to 72,000 steps, as shown in Figure 8. We tested the performance of the induced 72,000-step schedule and its simplified version (see Section 4.2) against the widely used cosine schedule. The induced schedule outperformed the cosine schedule, while the simplified version achieved results between the induced and cosine schedules. To further validate the effectiveness of the induced schedules, we compared downstream task performance for models trained using the cosine and induced schedules. As shown in Table 2, the induced schedule led to overall improvements in downstream tasks.

Peak Learning Rate Ablation. We evaluated the applicability of the MPL across different peak learning rates. In previous experiments, the peak learning rate was fixed at 3×10^{-4} . However, as shown in Figure 4, the empirical behavior of two-stage learning rate schedules deviates when the peak learning rate increases. To investigate this, we conducted experiments with peak learning rates of 4×10^{-4} and 6×10^{-4} . The MPL achieved an average R^2 value of 0.9965 for the 4×10^{-4} case and 0.9940 for the 6×10^{-4} case, demonstrating consistently high accuracy.

Batch Size Ablation. We conduct ablation experiments on sequence batch sizes of 64 and 256 over 25M models, apart from 128 in previous experiments. The MPL presents a consistent accuracy with R^2 higher than 0.9970.

Random Seed. We performed an ablation study to examine the impact of random seed variability on curves. We trained a 25M-parameter model for 24,000 steps using the cosine schedules with three random seeds. As shown in Figure 9, the standard deviation of the resulting loss values was approximately less than 0.001, establishing a lower bound for prediction errors. It highlights the prediction accuracy of the MPL discussed in Section 3.

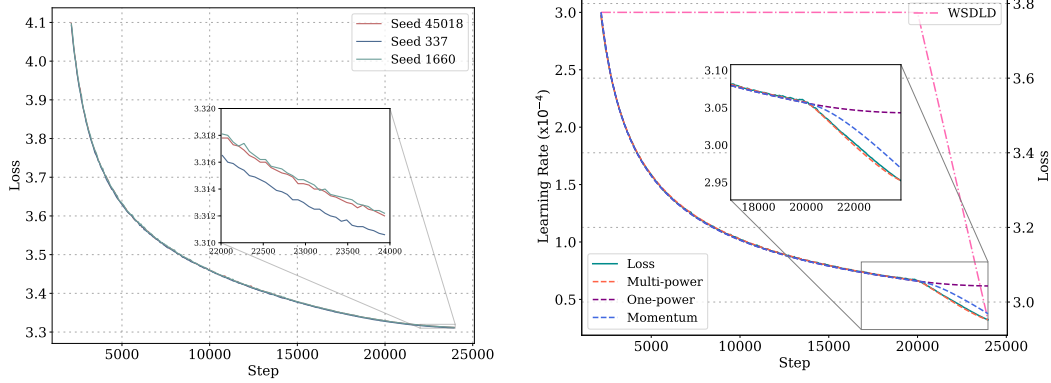


Figure 9: **Left:** The experiments on 25M and 24000 steps of different seeds. The standard variance of final loss is 0.0007 and the max gap is 0.0014. **Right:** Comparison between multi-power law and momentum law. In the decay stage, the multi-power law not only presents higher accuracy to fit the loss curve but also aligns with the curvature of the curve. As a comparison, the momentum law can also fit the loss in the stable stage, but it predicts a counterfactual concave curve in the decay stage.

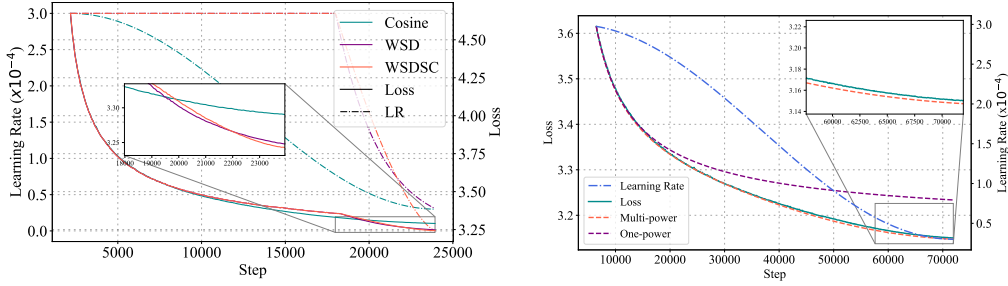


Figure 10: The loss curves of GPT2 models. The multi-power law is fit over 24000-step constant and cosine schedule losses. **Left:** The comparison between the cosine, WSD, and WSDSC (see Section 4.2) schedules; **Right:** Prediction on the 72000-step loss curve of cosine schedule.

B Simplification of Formula, Usage and Fitting.

B.1 Simplification of Formula

We simplify the full multi-power law (MPL; see Equation (1)) at various levels, trading computational complexity for reduced accuracy. Table 3 summarizes the fitting performance of simplified versions and variants of the MPL. The fitting experiments are conducted over 25M models.

No Loss Reduction. The necessity of the loss reduction term $LD(t)$ can be assessed by fitting a one-power law (OPL), a simplified MPL where $LD(t) = 0$ or equivalently $B = 0$:

$$\mathcal{L}_{\text{OPL}}(t) = L_0 + A \cdot S_1(t)^{-\alpha}, \quad S_1(t) := \sum_{\tau=1}^t \eta_{\tau}. \quad (13)$$

This formulation approximates the loss curve by linearizing the cumulative learning rate (LR) effects. The fitted results (first row of Table 3) exhibit significant degradation compared to the full MPL, demonstrating the critical role of $LD(t)$.

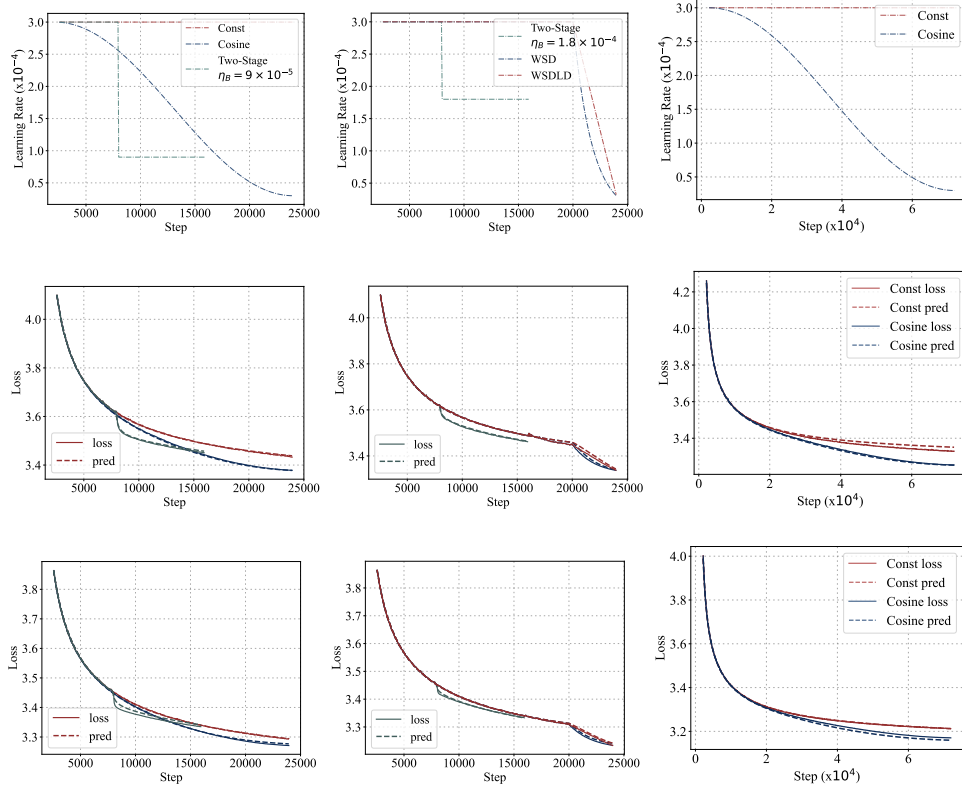


Figure 11: Ablation study on batch sizes. The R^2 values for batch sizes of 64 and 256 are 0.9977 and 0.9973, respectively. **Row One:** Learning rate schedules. **Row Two:** Loss curves for experiments with a sequence batch size of 64. **Row Three:** Loss curves for experiments with a sequence batch size of 256. **Column One:** Training set results. **Column Two:** Test set results, focusing on loss curves with the same horizon as the training set. **Column Three:** Test set results, focusing on loss curves with an extended horizon.

Table 3: Summary of fitting results for simplified laws. Each row corresponds to a specific law, reporting metrics including R^2 , MAE, RMSE, PredE, and WorstE. Higher R^2 values and lower MAE, RMSE, PredE, and WorstE indicate better fitting performance. See Table 1 for metric definitions.

Formula	Differences from MPL	$R^2 \uparrow$	MAE $\downarrow$	RMSE $\downarrow$	PredE $\downarrow$	WorstE $\downarrow$
OPL	$LD(t) = 0$ ($B = 0$)	0.8309	0.0378	0.0412	0.0111	0.0241
LLDL	$G(x) = 1$	0.9797	0.0077	0.0101	0.0023	0.0108
No- γ	$\gamma = 0$	0.9961	0.0046	0.0053	0.0014	0.0041
SPL	$x = t - k$	0.9921	0.0066	0.0075	0.0020	0.0069
MEL	$G(x) = 1 - e^{-Cx}$, $\gamma = 0$	0.9934	0.0044	0.0057	0.0013	0.0047
MTL	$G(x) = 1 - e^{-Cx}$, $x = t - k$	0.9904	0.0047	0.0060	0.0014	0.0047
MPL	(Ours)	0.9975	0.0039	0.0046	0.0012	0.0040

Linear Approximation of Loss Reduction. The loss reduction term $LD(t)$ (defined in Equation (2)) can be simplified by treating the scaling function $G(x)$ as a constant:

$$LD(t) \approx \sum_{k=2}^t B(\eta_{k-1} - \eta_k) = B(\eta_1 - \eta_t).$$

Despite its simplicity, we observe a near-linear relationship between $LD(t)$ and the LR reduction $(\eta_1 - \eta_t)$, regardless of the LR schedule type, as shown in Figure 12. This motivates the Linear Loss

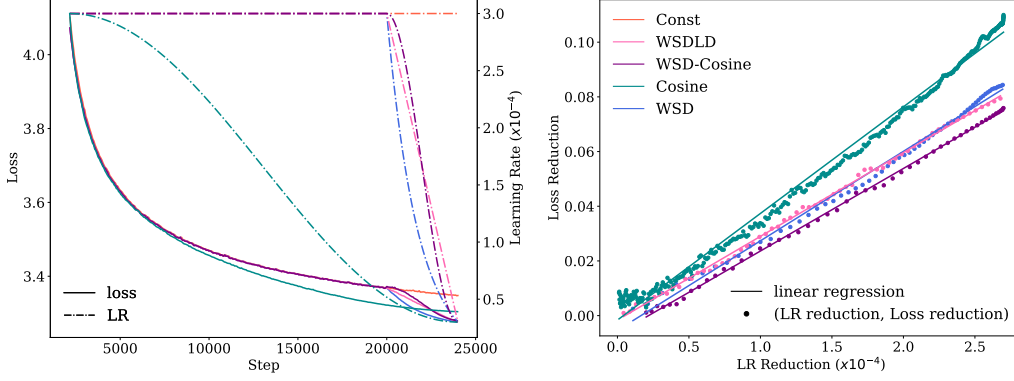


Figure 12: Linear regression between loss reduction and LR reduction over different schedules. The total step number is 24000 and the model size is 25M. WSD-Cosine denotes the WSD schedule with cosine decay function. The decay steps for the WSD schedule and variants are 4000. **Left:** the learning rate schedules and corresponding loss curves. **Right:** the loss reductions over LR reductions for different schedules, as well as their linear regression. The mean R^2 value is 0.9980.

reDUCTION Law (LLDL):

$$\mathcal{L}_{\text{LLDL}}(t) = L_0 + A \cdot S_1(t)^{-\alpha} + B(\eta_1 - \eta_t). \quad (14)$$

As shown in Table 3, LLDL achieves significantly better accuracy than OPL, although it underperforms the full MPL. However, this formulation is unsuitable for optimizing schedules, as its results collapse to trivial solutions.

Loss Reduction Without γ . Next, we simplify $G(x)$ by setting $\gamma = 0$, yielding the No- γ Law:

$$\mathcal{L}_{\text{No-}\gamma} = L_0 + A \cdot S_1(t)^{-\alpha} + B \sum_{k=2}^t (\eta_{k-1} - \eta_k) \cdot G(S_k(t)). \quad (15)$$

Results (third row of Table 3) indicate a slight performance drop, confirming that γ enhances fitting accuracy with minimal additional computational cost. Thus, we retain γ in the final MPL.

Step-Based Approximation. An alternative is to replace $G(\eta_k^{-\gamma} S_k(t))$ with a step-based formulation, $G(t - k)$. This yields the Step Power Law (SPL):

$$\mathcal{L}_{\text{SPL}} = L_0 + A \cdot S_1(t)^{-\alpha} + B \sum_{k=2}^t (\eta_{k-1} - \eta_k) \cdot G(t - k). \quad (16)$$

While simpler, this approximation reduces prediction accuracy and contradicts empirical results, as it implies loss reduction even when LR reaches zero.

Exponential Approximation. Substituting $G(x)$ with an exponential function $G(x) = 1 - e^{-Cx}$ gives the Multi-exponential Law (MEL):

$$\mathcal{L}_{\text{MEL}} = L_0 + A \cdot S_1(t)^{-\alpha} + B \sum_{k=2}^t (\eta_{k-1} - \eta_k) \cdot G(S_k(t)). \quad (17)$$

Results (fifth row of Table 3) show a performance drop compared to the power-based MPL, consistent with observations in Section 2.

Relation to Momentum Law. The concurrently proposed Momentum Law (MTL) is in the form of

$$\mathcal{L}_{\text{MTL}}(t) = L_0 + A \cdot S_1(t)^{-\alpha} + B \cdot S_2, \text{ where } S_1 = \sum_{i=1}^t \eta_i, S_2 = \sum_{i=2}^t \sum_{k=2}^i (\eta_{k-1} - \eta_k) \lambda^{i-k}.$$

$\lambda < 1$ is a hyper-parameter of MTL. It is indeed a variant of MPL since

$$S_2 = \sum_{i=2}^t \sum_{k=2}^i (\eta_{k-1} - \eta_k) \lambda^{i-k} = \sum_{k=2}^t (\eta_{k-1} - \eta_k) \sum_{i=k}^t \lambda^{i-k} = \sum_{k=2}^t (\eta_{k-1} - \eta_k) \left(\frac{1 - \lambda^{t-k+1}}{1 - \lambda} \right).$$

Thus, the Momentum Law (MTL) is a variant of MPL with an exponential step-based approximation:

$$\mathcal{L}_{\text{MTL}}(t) = L_0 + A \cdot S_1(t)^{-\alpha} + B' \cdot G(t - k + 1), \quad G(x) = 1 - e^{-C'x}.$$

Here, $B' = \frac{B}{1-\lambda}$, $C' = -\log \lambda$. While MTL incorporates step-based decay, its performance (last second row of Table 3) lags behind MEL, highlighting the limitations of step-based approximations.

B.2 Approximation of Decay Function

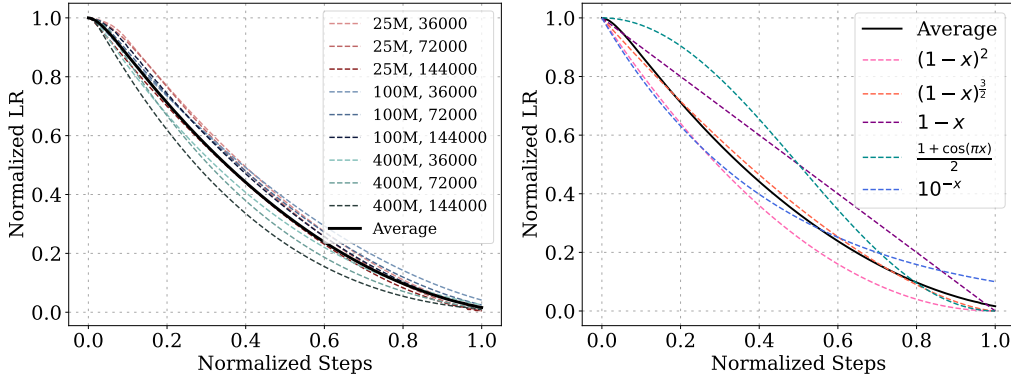


Figure 13: Approximation of decay functions.

To facilitate the use of the optimized LR schedules and find out their decaying trends, we try to approximate the decay function of the WSD-like optimized schedules. Given the optimized schedules $S = \{t, \eta\}_{N,T}$, where N and T denote model sizes and training steps, we compute the normalized LR and steps as follows:

$$\tilde{\eta} = \frac{\eta}{\eta_{\max}}, \quad \tilde{t} = \frac{t - t_{\min}}{t_{\max} - t_{\min}}.$$

Then we average across N and T as shown in the left of Figure 13, computing an averaged LR schedule $\tilde{\eta}_{\text{avg}}$. Then we utilize a symbolic regression approach to search for the approximate function form of the decay function. We get that $f(x) = (1 - x)^{-\alpha}$ best capture the relation between $\tilde{t}$ and $\tilde{\eta}_{\text{avg}}$. In our experiments, $\alpha = \frac{3}{2}$ fits well and we show the average schedule against different candidates function form in the right of Figure 13. We use the WSD with the decay function f to train a 25M model with 24000 steps and 4000 decay steps. The result with a final loss of 3.274 slightly outperforms the WSD with exponential function (Hu et al., 2024) with a final loss of 3.276, but can not match the directly optimized schedule, which reaches below 3.269.

B.3 The Selection of Training Set

We conduct ablation experiments over the loss curves in the training sets, including two-stage, cosine, and constant LR schedules. We remove one of them and keep the other two as training sets. Then we fit over these subsets of the full training sets in the same approach with the multi-power law. The runs are over 25M models. The resulting coefficients are shown in Table 4 and the resulting test metrics are shown in Table 5. The test metrics are measured over the full test sets, including different schedule types and horizons. There are some observations as follows:

- As shown in Table 4, the coefficients of different fittings are consistent overall, while there are a few parameters that vary, like C . We conjecture there are some correlations between C and other parameters like γ . A refined form of multi-power law would be expected in future work.

- As shown in Table 5, the multi-power law shows its robustness over the training sets, with comparable performances between the full-set fitting and the subset-fitting results.

Items	A	B	C	α	β	γ	L_0
Full	0.507	446.4	2.070	0.531	0.406	0.522	3.1
Cosine + 2-stage	0.5272	455.0	6.276	0.5032	0.3622	0.4172	3.147
Constant + 2-stage	0.5279	457.0	7.569	0.5067	0.3613	0.4002	3.149
Constant + Cosine	0.5292	477.4	0.854	0.5041	0.3189	0.6256	3.146

Table 4: Parameters for Different Fittings. “Full” denotes the fitting with the full training set, all three loss curves.

Model	$R^2 \uparrow$	MAE $\downarrow$	RMSE $\downarrow$	PredE $\downarrow$	WorstE $\downarrow$
Full	0.9975	0.0039	0.0046	0.0012	0.0040
Cosine + 2-stage	0.9971	0.0040	0.0046	0.0012	0.0048
Constant + 2-stage	0.9976	0.0037	0.0045	0.0011	0.0039
Constant + Cosine	0.9993	0.0020	0.0031	0.0006	0.0060

Table 5: Performance Metrics for Different Fittings.

C Related Work

Optimal Learning Rate Schedule. Designing an effective learning rate schedule for deep learning has been a prominent research focus. Smith (2017) proposed a cyclical learning rate schedule. Loshchilov & Hutter (2017), inspired by warm restarts, introduced the cosine learning rate schedule, demonstrating its superiority across multiple experimental settings. From a theoretical perspective, Li & Arora (2019) introduced an exponential decay learning rate schedule based on the equivalence of weight decay. Xu et al. (2019) utilized reinforcement learning algorithms to learn a learning rate schedule adaptively. Pan et al. (2021) proposed an eigenvalue-dependent step schedule by incorporating the eigenvalue distribution of the objective function’s Hessian matrix into the design of the learning rate scheduler. Geiping & Goldstein (2023) experimentally compared the performance differences of various learning rate schedules, concluding that quick annealing of the schedule aids in performance improvement. Recently, Hu et al. (2024) introduced a three-phase learning rate schedule with warm-up, stable, and decay phases, showcasing its superior performance across multiple datasets.

However, some of these papers focus on heuristically designing high-performance learning rate schedules without a comprehensible, principled approach to optimize the schedule. Some of the others try to optimize schedules within a function subspace. The effectiveness of the resulting function may be restricted by the subspace. Our paper seeks to open the door to a principled and comprehensible path of the optimal learning rate schedule design.

Scaling Laws. Scaling laws have arguably been the driving force behind the development of large language models. Initially proposed by Kaplan et al. (2020) and further developed by Hoffmann et al. (2022), Kadra et al. (2023), Aghajanyan et al. (2023) and Muennighoff et al. (2023), among others, most scaling laws adopt a power law form. However, due to the lack of dependence on the learning rate, these laws typically predict only the final loss of a training process, lacking guidance for the full training curve. This is because only the final loss bears a full LR decay while the LR decays at the intermediate steps are not sufficient. Typically, they need more than 10 training curves to obtain the scaling law of the final losses for one particular schedule type, the Cosine schedule practically (Hoffmann et al., 2022; Muennighoff et al., 2023). As a comparison, we could fit the LR-dependent multi-power law applicable across different LR schedule types within only 2-3 loss curves.

Several explanations for the power law form of scaling laws have been proposed, ranging from the perspective of data manifolds (Sharma & Kaplan, 2020) to the power law distribution of eigenvalues

in the loss landscape (Lin et al., 2024). While our paper does not delve into the discussion about the model dimension scaling, we discuss the scaling along the data dimension along with the LR dimension. We believe it offers new perspectives and a novel starting point for theoretical investigations.

Hyperparameters Optimization. Hyperparameter optimization (HO) has long been a focal point of research within the machine learning community. For learning rate schedules (LR schedule), early works primarily employed Bayesian optimization-based approaches (Hutter et al., 2011; Snoek et al., 2012; Bergstra et al., 2013) or bandit-based solutions (Li et al., 2018) to tune hyperparameters. However, these works typically parameterized LR schedule as a learnable constant or a family of functions with learnable parameters, without fully exploring the potential of LR schedule. While this form of parameterization offers theoretical and experimental convenience, it often lacks interpretability. Furthermore, methods proposed in Teng et al. (2021); Jin et al. (2021) aim to adjust LR schedule during training automatically, but these approaches cannot identify the optimal LR schedule before training begins, and they fail to fully generalize across different datasets, underutilizing the scaling law information followed by the model. In contrast, Klein et al. (2022) selects hyperparameters based on differences in learning curves for various hyperparameters, while Kadra et al. (2024) recognizes the power law phenomenon and develops HP methods based on power law. However, we propose a more robust scaling law than the power law specifically for the LR schedule dimension and present a comprehensive framework for optimizing LR schedule.

Theory in Scaling Law. Although there are numerous experimental studies on scaling laws, our understanding of the theoretical explanation and origins of scaling laws remains very limited. Sharma & Kaplan (2020) demonstrated that the exponent of the power law is related to the intrinsic dimension of the data in a specific regression task. Hutter (2021) examined a binary classification toy problem, deriving a scaling law with respect to data dimensionality for this problem. Jain et al. (2024) investigated scaling laws in the context of data selection. Bahri et al. (2024) assumed a power-law spectrum on the covariates, obtaining a scaling law with respect to data and model dimensions in the setting of least squares loss. Bordelon et al. (2024) considered scaling laws in regression problems under gradient flow. Atanasov et al. (2024) and Lin et al. (2024) discussed the formation of scaling laws in high-dimensional linear regression problems. Notably, our theoretical analysis is the first to provide a loss prediction throughout the training process from the perspective of the learning rate schedule, formally resembling the multi-power law observed in our experiments.

D Discussions of Multi-Power Law Derivation (Section 2)

D.1 Chinchilla Data Scaling Laws for the Final Loss Prediction.

Considering a fixed hyper-parameter setting, including the batch size, learning rate schedule, model type, previous work Muennighoff et al. (2023); Hoffmann et al. (2022) mainly follows the Chinchilla Scaling Laws to extrapolate model size N and data quantity D : $\mathcal{L}(N, D) = L_0 + A \cdot D^{-\alpha} + B \cdot N^{-\beta}$. According to the form of law, the data scaling is roughly independent of the model scaling. Thus, we could focus on the data scaling and refer to $\mathcal{L}(T) = L_0 + A \cdot T^{-\alpha}$ as the Chinchilla Data Scaling Laws (CDSLs), where T denotes the training steps given the fixed batch size. The Chinchilla law is exclusively applicable to the final training loss because the Chinchilla law is LR-independent and the mid-training parts of loss curves commonly bear insufficient learning rate decay compared to the final loss (Hoffmann et al., 2022). As shown on the left of Figure 7, to extrapolate the final loss, we first need to generate several (typically more than 10 (Hoffmann et al., 2022; Dubey et al., 2024)) loss curves given a specific schedule (typically the Cosine schedule). Then we could fit CDSL over the final losses. Noticeably, the validation loss decreasing by 0.001 matters in the LLM scenario, because slight progress in loss may require intense computation practically, especially on a large scale. Moreover, the validation loss probably correlates with the emergent ability in downstream tasks. A little difference in the loss scale may indicate a steep deviation in the downstream performance (Du et al., 2024).

D.2 Motivation: Continuous Approximations of the Training Dynamics.

The rationale behind this approach is that matching the learning rate sum between the two training processes should result in similar training losses, and thus a more accurate approximation can be

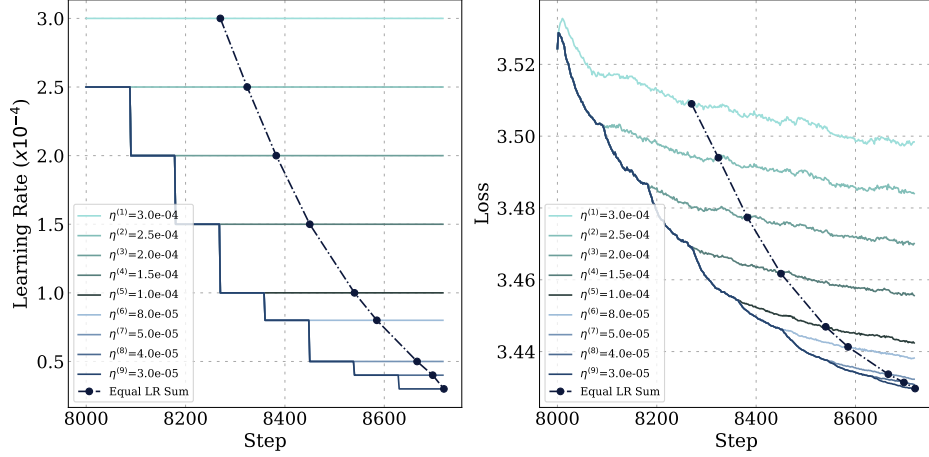


Figure 14: **Left:** Multi-stage schedule and interpolated LR schedules between the multi-stage LR schedule and the auxiliary schedule. There are 9 stages in our case, and the length of each stage, except the first one, is 90. The step points with the equal LR sum as the final step are marked in black and linked with the dash-point line. The learning rates before 8000 steps are constant at 3×10^{-4} . **Right:** Corresponding training curves for the actual multi-stage training curve, the auxiliary schedule as well as their interpolation.

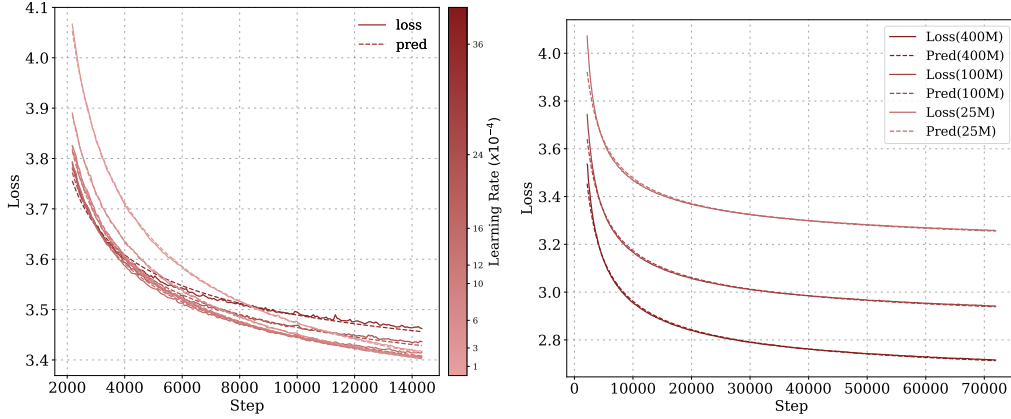


Figure 15: Loss curves trained with constant LR schedules. **Left:** The learning rates of schedules range from 3.0×10^{-4} to 3.6×10^{-3} . The total step number is 14400. The loss curves are all fit with Equation (5). The mean MSE is 1.55×10^{-5} and the mean coefficient of determination (R^2) is 0.9976. **Right:** Model sizes include 25M, 100M, and 400M. The total step number is 72000 and LR is 3.0×10^{-4} . The mean MSE is 8.04×10^{-5} and the mean R^2 is 0.9947.

obtained by further exploring the loss reduction term $LD(t)$. To see this, we take SGD as an example. In theory, if the learning rates used in training $\eta_1, \dots, \eta_T$ are small, then SGD is known to be a first-order approximation of its continuous counterpart (Li et al., 2017; Cheng et al., 2020; Elkabetz & Cohen, 2021), gradient flow, which evolves the parameters $\theta(\tau)$ according to the differential equation, $\frac{d\theta(\tau)}{d\tau} = -\nabla \mathcal{L}(\theta(\tau))$, where $\nabla \mathcal{L}(\theta)$ stands for the gradient of the loss function at θ , and τ is a continuous time variable. In this continuous approximation, the step t of SGD corresponds to the evolution of $\theta(\tau)$ for a small continuous time interval of length η_t . When the LRs are extremely small, the parameter after t steps of SGD should be close to $\theta(\tau)$ with $\tau = \sum_{k=1}^t \eta_k$. This motivates us to match the learning rate sum $\sum_{k=1}^t \eta_k$ between the two training processes to approximate the loss curve. This argument can be extended to other optimization algorithms, such as Adam (Ma et al., 2022). However, these continuous approximations can be loose for realistic LR schedules, thus necessitating a more detailed analysis of the loss reduction term $LD(t)$.

D.3 Incorporating the Warmup Stage.

In practice, many learning rate schedules include a warmup stage where the learning rate gradually increases from zero to a peak value. To incorporate the effect of this stage, we can change the definition of auxiliary process to include the same warmup stage before using the constant learning rate. Then similar to the argument in Section 2.1, letting W be the sum of learning rates in the warmup stage, we can change the auxiliary loss formula to $\hat{\mathcal{L}}_{\text{const}}(Z(t)) = L_0 + A \cdot (W + S_1(t))^{-\alpha}$. Our multi-power law then becomes

$$\mathcal{L}(T) \approx L_0 + A \cdot (W + S_1(t))^{-\alpha} - \sum_{k=2}^t B(\eta_{k-1} - \eta_k)(1 - (C\eta_k^{-\gamma} S_k(t) + 1)^{-\beta}). \quad (18)$$

This is the actual law we use in experiments since practical schedules often include a warmup stage.

E Two-Stage Experiments (Section 2.2)

In this section, we show the details of the investigation of the variation of coefficients of the power law of two-stage LR schedules.

Experiment Setting and Law Fitting. The default setting is $\eta_A = 3 \times 10^{-4}$, $\eta_B = 3 \times 10^{-5}$, $T_A = 8000$. In the ablation experiment, η_A ranges from 5×10^{-5} to 1×10^{-3} , η_B ranges from 4×10^{-5} to 2.9×10^{-4} , and T_A ranges from 4000 to 28000. The second stage lengths range from 1000 to over 6000. We follow Hoffmann et al. (2022) to utilize Huber loss as the objection function (Huber, 1992),

$$\min_{\Theta} \sum_x \text{Huber}_{\delta}(\log \widehat{\text{LD}}_{\Theta}(T_A + x) - \log \text{LD}(T_A + x)),$$

where $\Theta = \{\tilde{B}, \tilde{C}, \beta\}$, and we set $\delta = 1 \times 10^{-2}$. For each experiment, we use the Adam optimizer with a learning rate at 1×10^{-4} , and total steps of 20000. Here we do not conform to the L-BFGS algorithm due to the function form of $U(s)$. The parameters are initialized based on the estimation of asymptotic values of loss reduction and the slopes at the beginning of the second stage.

Fixed β Results. We propose a function form in Equation (6) to fit the loss reduction curve. The form guarantees the loss reduction zero when there is no LR reduction and fits with a power law. To explore the coefficient relation with the learning rate, we first investigate the coefficients fit in the ablation experiments. For the sake of further derivation and based on the fitted coefficients in the experiments, we fix the exponent β as LR-independent parameter 0.4. Then we re-fit the loss curves given $\beta = 0.4$ to validate the power form holds and further investigate the dependency of different parameters on the η_A , η_B , and T_A . The relation pattern is presented in Figure 4. Part of the two-stage schedules experiments are shown in Figure 16, including the ablations over the first stage as well as the second stage learning rates. Although β is fixed, the error margin is feasible for further derivation.

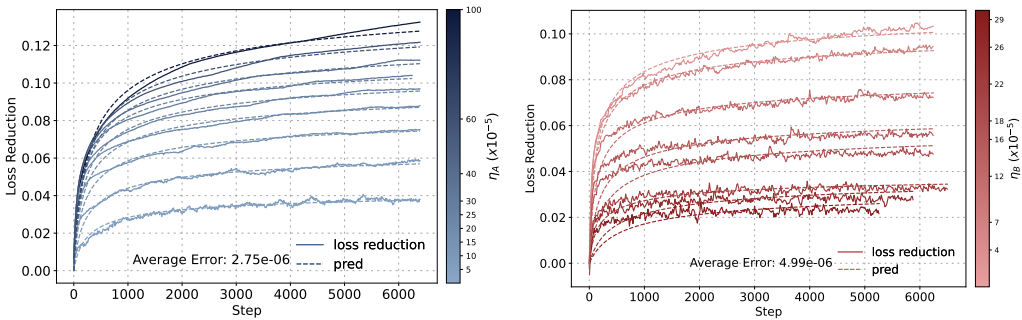


Figure 16: **Left:** Varying η_A , the loss reductions over steps x of two-stage LR schedules; **Right:** Varying η_B , the loss reductions over steps x of two-stage LR schedules.

F Details of Validation Experiments (Section 3)

Training Set, Test Set, and Model Training Settings. The validation experiments are framed as a machine learning task, where the loss curves in the training set are used to fit the multi-power law, which is then evaluated on the test set to report prediction accuracy.

The default training set consists of three loss curves: one trained with a cosine learning rate schedule, one with a constant learning rate schedule, and one with a two-stage learning rate schedule. The default test set includes unseen learning rate schedule types, loss curves over longer horizons, and more extreme two-stage learning rate schedules. Detailed descriptions of the training and test sets are provided in Table 6. Unless otherwise specified, the ending learning rate is set to $1/10$ of the peak learning rate (3×10^{-5} by default). The warmup phase spans 2,160 steps, but as the focus is on the post-warmup phase, only the post-warmup sections are used for fitting (Hu et al., 2024; Tissue et al., 2024).

The loss curves used in these experiments are generated from training the Llama2 model. The batch size is fixed at 128, and the sequence length is set to 4,096 across all configurations, resulting in 0.5M tokens per step. To simplify, data volume is described in terms of steps, where 10,000 steps consume 5B tokens. Validation loss is used as the default performance measure. Detailed model training hyperparameters are listed in Table 7, and a summary of the model series parameters used in the experiments is presented in Table 8.

Set	Schedule Type	Total Lengths	η_B/η_A
Training	Constant	24000	0.3
	Cosine	24000	
	Two-stage	16000	
Test	WSD	24000	0.1
	WSDLD	24000	
	Two-stage	16000	
	Two-stage	16000	0.6
	Constant	72000	
	Cosine	72000	

Table 6: Summary of training and test sets.

Default Hyperparameter	Value
Sequence Batch Size	128
Sequence Length	4096
Optimizer Type	AdamW
Beta1	0.9
Beta2	0.95
Epsilon	1×10^{-8}
Weight Decay	0.1
Gradient Clipping	1.0
Peak Learning Rate	3×10^{-4}
Final Learning Rate	3×10^{-5}
Warmup Steps	2160

Table 7: Hyperparameters related to model training.

Fit the Multi-power Law. Similar to the two-stage fitting, we utilize the Huber loss as the objective function (Huber, 1992),

$$\min_{\Theta} \sum_t \text{Huber}_{\delta}(\log \mathcal{L}_{\Theta}(t) - \log \mathcal{L}_{\text{gt}}(t)), \quad (19)$$

where $\Theta = \{A, B, C, \alpha, \beta, \gamma, L_0\}$, $\delta = 1 \times 10^{-3}$ and $\mathcal{L}_{\text{gt}}(t)$ denotes the ground truth of validation losses. We adopt the Adam optimizer, with a learning rate at 5×10^{-3} for the index parameters that

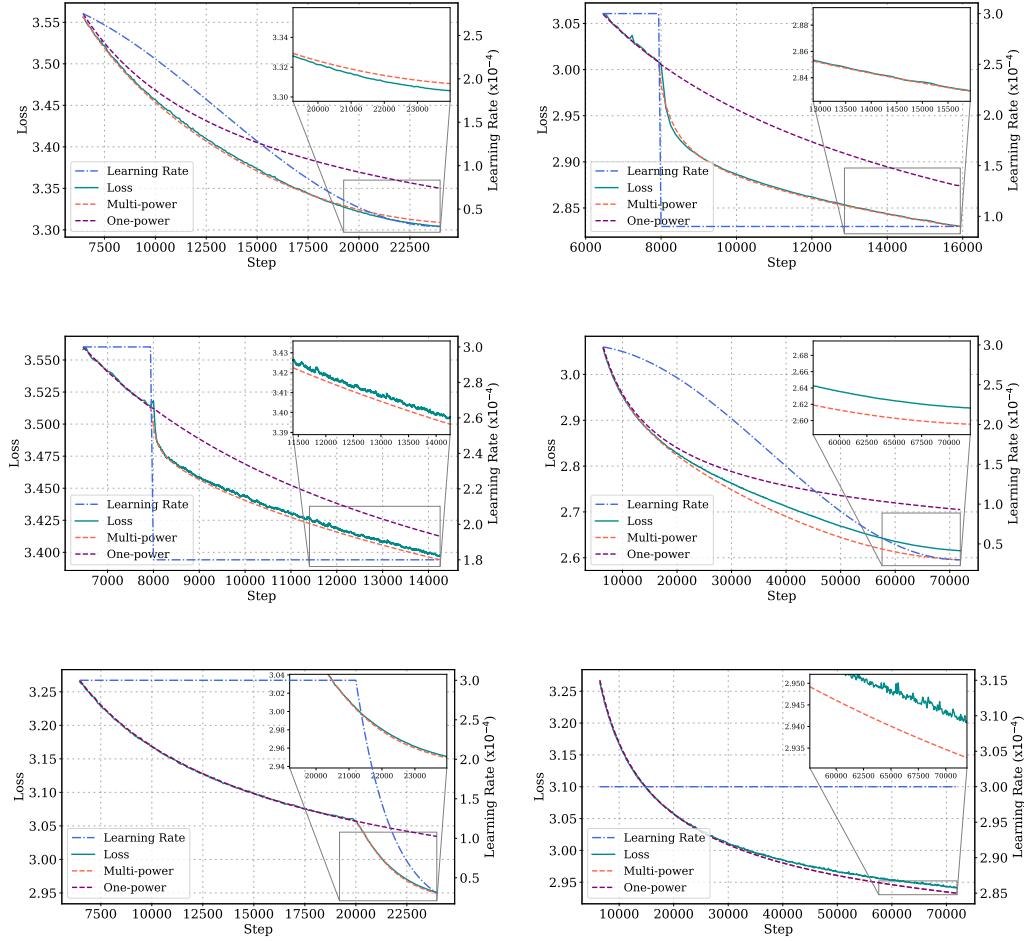


Figure 17: **Details of fitting and prediction.** The subfigures illustrate loss curve fitting (training set) and prediction (test set) for various configurations. (X, Y) indicates the subfigure at row X , column Y . The columns in the accompanying table describe: **F/P** for Fitting (F) or Prediction (P), **Model Size (M)**, **Step Length (S)**, and **Learning Rate Schedule (LRS)**. Details of each subfigure are provided below:

(X, Y)	F/P	Model Size (M)	Step Length (S)	LR Schedule (LRS)
(1, 1)	F	25M	24,000	Cosine
(1, 2)	F	400M	16,000	2-stage ($3 \times 10^{-4} \rightarrow 9 \times 10^{-5}$)
(2, 1)	P	25M	16,000	2-stage ($3 \times 10^{-4} \rightarrow 1.8 \times 10^{-4}$)
(2, 2)	P	400M	72,000	Cosine
(3, 1)	P	100M	24,000	WSD
(3, 2)	P	100M	72,000	Constant

Codename	Embedding Dimension	#Heads	#Layers	#Non-embeddings	#Params
25M	640	5	5	25	89
100M	1024	8	8	101	205
400M	1536	12	12	340	493
1B	2048	36	16	822	1026

Table 8: The model series used in all the experiments. Hoffmann et al. (2022) utilizes the number of non-embedding parameters (#Non-embeddings) to count model sizes, while Kaplan et al. (2020) counts the total number of parameters (#Params). The unit of the Parameter is M in this table.

are α , β , and γ in our law, and 5×10^{-2} for the coefficient or constant parameters, that are A , B , C and L_0 in our law. We also take a learning rate at 1×10^{-5} and 1×10^{-6} with initialization as the previous fitting parameters. We select the result with lower training loss from the first optimization result and the second one. Each optimization takes over 5×10^4 steps. For 400M results, we find that $A = 0.658$, $B = 614.3$, $C = 0.164$, $\alpha = 0.421$, $\beta = 0.883$, $\gamma = 0.564$, $L_0 = 2.524$. For 100M cases, we have $A = 0.592$, $B = 521.4$, $C = 0.242$, $\alpha = 0.460$, $\beta = 0.604$, $\gamma = 0.647$, $L_0 = 2.792$. Part of fitting and prediction examples are shown in Figure 17.

Fit the Momentum Law. We mainly follow the approach proposed by Tissue et al. (2024). The objective function follows Equation (19) and we adopt the L-BFGS algorithm to minimize it. For a fair comparison, we grid search over its hyperparameter λ in $\{0.95, 0.99, 0.995, 0.999, 0.9995\}$ and select the best hyperparameter based on the fitting accuracy over the training set. We also evaluate the law over the full test set listed in Table 6. The prediction accuracy comparison between our multi-power law and the momentum law is shown in Table 5.

G Details of Optimized LR schedule (Section 4)

Details of Optimizing the Surrogate Objective. To make the optimization more stable, we define the following quantities $d\eta := \{d\eta_1, d\eta_2, \dots, d\eta_T\}$, where $d\eta_i := \eta_{i-1} - \eta_i$. Thus we can conduct optimization with an easier constraint over $\min_{d\eta} \tilde{\mathcal{L}}_{\hat{\Theta}}(d\eta)$. Notice that $\eta_i = \eta_0 - \sum_{k=1}^i d\eta_k$ and η is one-to-one with $d\eta$. We can denote $\tilde{\mathcal{L}}_{\hat{\Theta}}(d\eta) = \mathcal{L}_{\hat{\Theta}}(\eta)$. So now, instead of directly optimizing $\min_{\eta} \mathcal{L}_{\hat{\Theta}}(\eta)$ in Equation (12), we can conduct optimization with an easier constraint over $\min_{d\eta} \tilde{\mathcal{L}}_{\hat{\Theta}}(d\eta)$, which is,

$$\begin{aligned} & \min_{d\eta} \tilde{\mathcal{L}}_{\hat{\Theta}}(d\eta) \\ & s.t., \sum_{i=1}^T d\eta_i \leq \eta_0, \forall 1 \leq i \leq T, \\ & 0 \leq d\eta_i. \end{aligned}$$

In practice, we find we can solve the optimization problem with a relaxed constraint,

$$\begin{aligned} & \min_{d\eta} \tilde{\mathcal{L}}_{\hat{\Theta}}(d\eta) \\ & s.t., 0 \leq d\eta_i \leq \eta_0. \end{aligned}$$

The optimized results $d\eta$ also satisfy the constraint $\sum_{i=1}^T d\eta_i \leq \eta_0, \forall 1 \leq i \leq T$. The constraints are forced by clipping. Our optimization is applied to the law fitted over the training set mentioned in Appendix F. Regarding the optimization details, we use the Adam optimizer with a constant learning rate. The learning rate scale is searched ranging from 2×10^{-8} to 1×10^{-9} and the optimization step number ranges from 50000 to 200000 for better convergence.

Decay Ratio Details. Following Hu et al. (2024), we take the decay function of both exponential decay and linear decay. We grid search over 3000, 4000, 5000, 6000 and 7000 to find that the best decay step number is 6000, with a total steps of 24000. The ending learning rate is set to 1/10 of the peak learning rate following Hu et al. (2024). According to Figure 1, we find that the decay ratio of our optimized learning rate schedule aligns with the grid-searched WSD. The ending learning rate is lower than the empirical one and the decay shape is between linear and exponential functions.

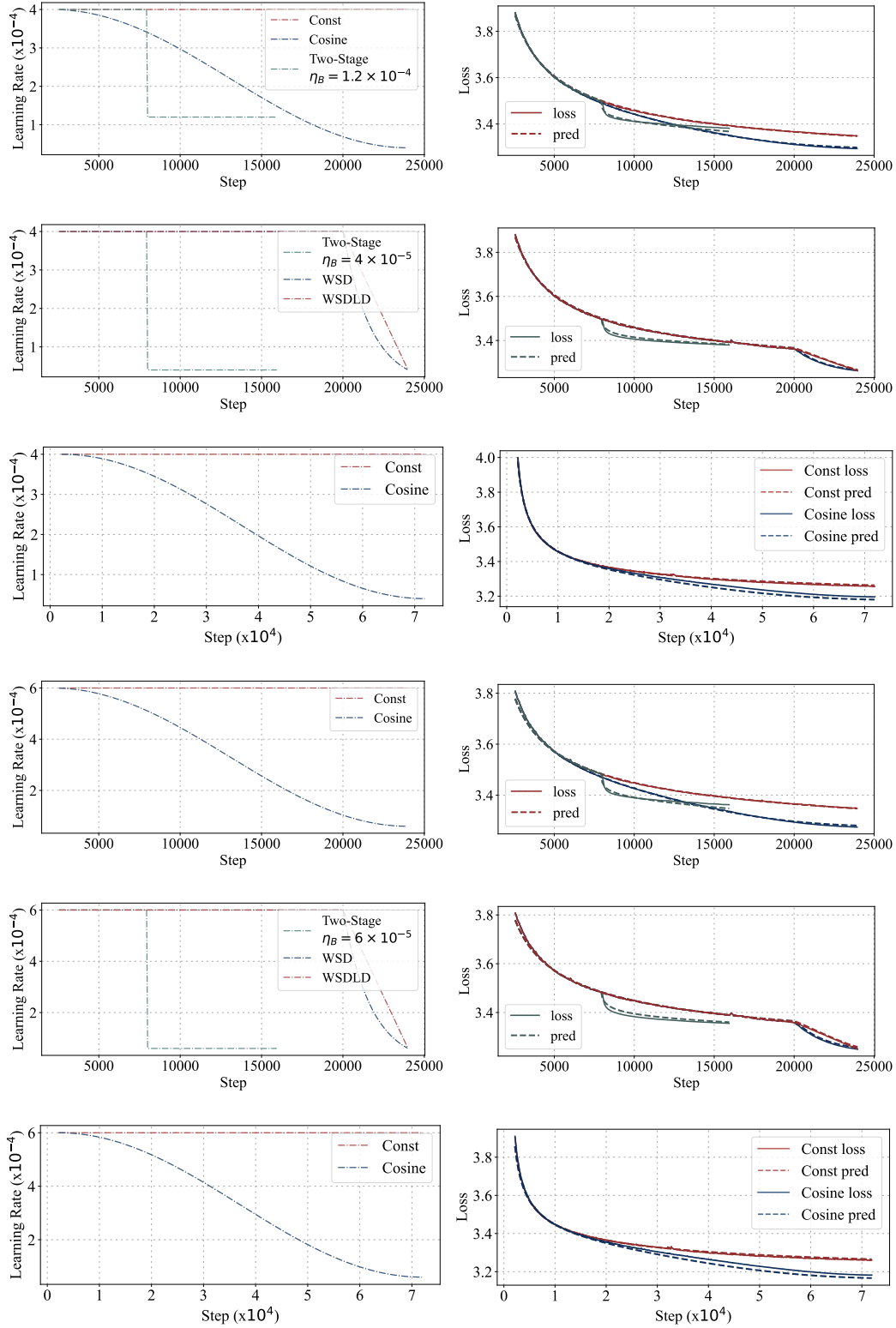


Figure 18: Ablation over peak learning rates. **Left:** the learning rates of the schedules; **Right:** the loss curves of schedules. **Updown:** the first three rows are the results for the peak learning rate at 4×10^{-4} and the last three rows are for the peak learning rate at 6×10^{-4} . For each set of the three rows, the first row shows the fitting on the training set, the second row shows the prediction over unseen schedules and the third row shows the extrapolation on a long horizon loss curve.

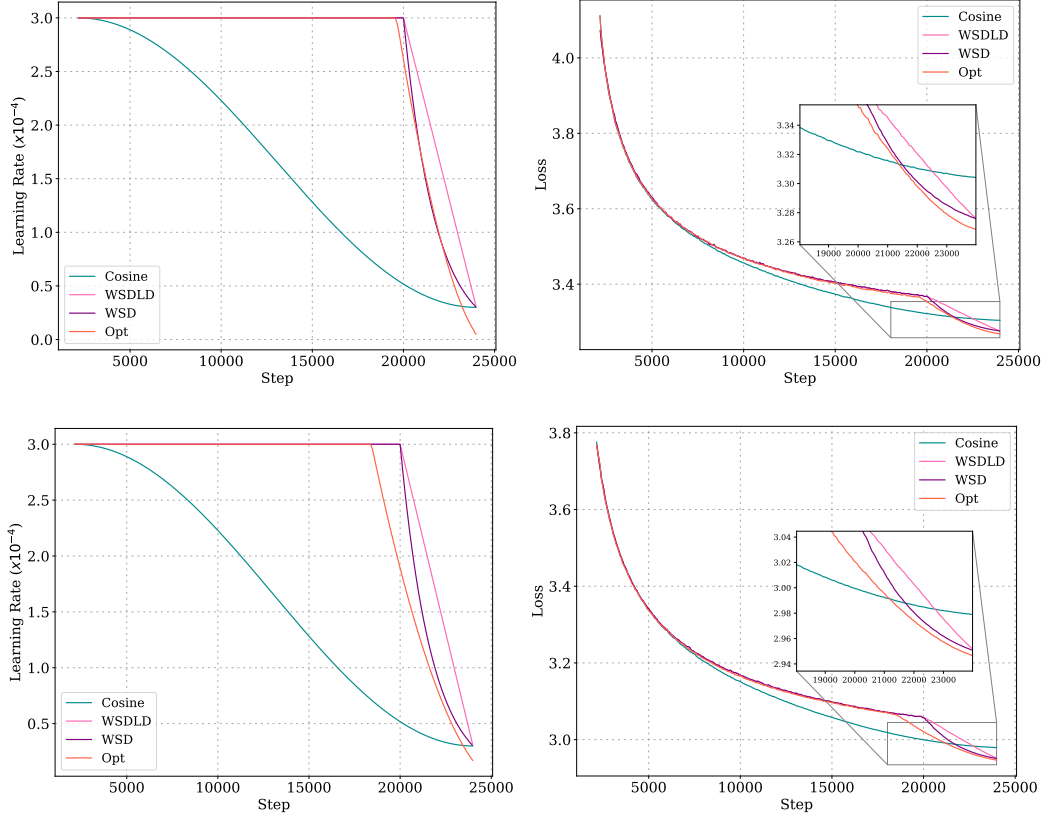


Figure 19: Our optimized LR schedules and their loss curve compared with Cosine, WSD, and WSDLD schedules. The total step number is 24000. The decay step number of WSD and its variant is 4000. **Upper:** 25M; **Lower:** 100M; **Left:** Learning rates over step; **Right:** Losses over step.

Optimized Schedule of Longer Horizons and Different Model Sizes Apart from the optimized LR schedules shown in Figure 1 and Figure 19, we further validate the optimized schedules of longer horizons and different model sizes. We optimize the LR schedules of 72000 steps based on the multi-power law fit over the training set. The training set conforms to the default setting only containing curves with lengths no longer than 24000, and we conduct experiments from 25M to 400M. As shown in Figure 20, the resulting schedules are also in the shape of WSD schedules, consisting of a stable phase and a decay phase. We compare the loss curves of the optimized LR schedules with those of commonly used Cosine LR schedules, we find that the optimized LR schedules outperform the Cosine LR schedules across different model sizes.

Zero-Ending Learning Rate Experiments. The optimized schedules consistently outperform WSD variants with “zero-ending” learning rates. As shown in Figure 21, we compare WSD(LD) variants with near-zero ending learning rates, the optimized schedules, and the original WSD(LD) schedules. In this experiment, the ending learning rate is set to 3×10^{-7} , which is 1/100 of the previous setting. Notably, a lower ending learning rate does not consistently lead to improved final loss. For example, the final loss of the WSD schedule increases with a near-zero ending learning rate. This suggests a complex interaction between the ending learning rate and the decay function, highlighting the challenges of jointly optimizing these hyperparameters in WSD schedules. In this context, the optimized schedule demonstrates its advantage by reducing the need for extensive hyperparameter tuning in WSD variants.

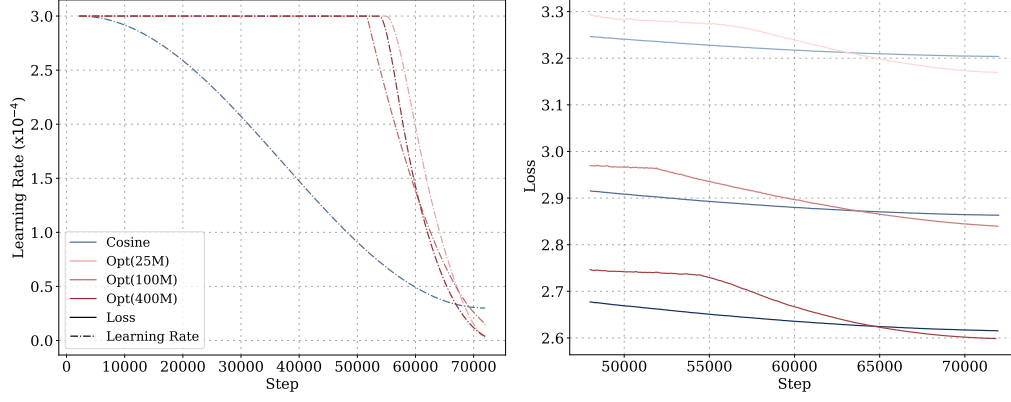


Figure 20: **Left:** optimized LR schedule vs Cosine LR schedule. The total step number is 72000, and model sizes range from 25M to 400M. **Right:** The loss curves of optimized schedules and Cosine schedules.

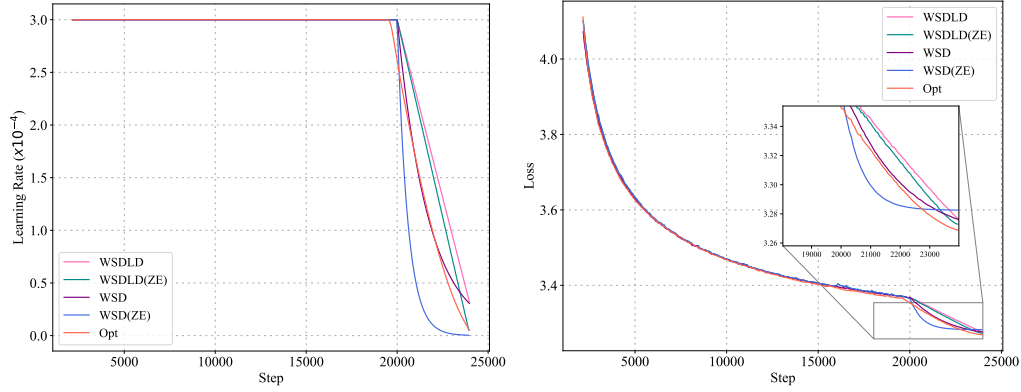


Figure 21: The comparison between the optimized schedules with the WSD variants with end LR as 0. WSD (ZE) and WSDL (ZE) represent the WSD and WSDL variants with ending learning rate as 3×10^{-7} , approximately “zero” ending LR compared to default 3×10^{-5} . **Left:** the learning rate comparison; **Right:** the loss comparison.