

# M-L2O: TOWARDS GENERALIZABLE LEARNING-TO-OPTIMIZE BY TEST-TIME FAST SELF-ADAPTATION

Junjie Yang , Xuxi Chen , Tianlong Chen , Zhangyang Wang , Yingbin Liang

The Ohio State University, University of Texas at Austin

{yang.4972, liang.889}@osu.edu

{xxchen, tianlong.chen, atlaswang}@utexas.edu

## ABSTRACT

Learning to Optimize (L2O) has drawn increasing attention as it often remarkably accelerates the optimization procedure of complex tasks by “overfitting” specific task types, leading to enhanced performance compared to analytical optimizers. Generally, L2O develops a parameterized optimization method (*i.e.*, “optimizer”) by learning from solving sample problems. This data-driven procedure yields L2O that can efficiently solve problems similar to those seen in training, that is, drawn from the same “task distribution”. However, such learned optimizers often struggle when new test problems come with a substantial deviation from the training task distribution. This paper investigates a potential solution to this open challenge, by meta-training an L2O optimizer that can perform fast *test-time self-adaptation* to an out-of-distribution task, in only a few steps. We theoretically characterize the generalization of L2O, and further show that our proposed framework (termed as **M-L2O**) provably facilitates rapid task adaptation by locating well-adapted initial points for the optimizer weight. Empirical observations on several classic tasks like LASSO, Quadratic and Rosenbrock demonstrate that M-L2O converges significantly faster than vanilla L2O with only steps of adaptation, echoing our theoretical results. Codes are available in <https://github.com/VITA-Group/M-L2O>.

## 1 INTRODUCTION

Deep neural networks are showing overwhelming performance on various tasks, and their tremendous success partly lies in the development of analytical gradient-based optimizers. Such optimizers achieve satisfactory convergence on general tasks, with manually-crafted rules. For example, SGD (Ruder, 2016) keeps updating towards the direction of gradients and Momentum (Qian, 1999) follows the smoothed gradient directions. However, the reliance on such fixed rules can limit the ability of analytical optimizers to leverage task-specific information and hinder their effectiveness.

Learning to Optimize (L2O), an alternative paradigm emerges recently, aims at *learning* optimization algorithms (usually parameterized by deep neural networks) in a data-driven way, to achieve faster convergence on specific *optimization task* or **optimizee**. Various fields have witnessed the superior performance of these learned **optimizers** over analytical optimizers (Cao et al., 2019; Lv et al., 2017; Wichrowska et al., 2017; Chen et al., 2021a; Zheng et al., 2022). Classic L2Os follow a two-stage pipeline: at the *meta-training* stage, an L2O optimizer is trained to predict updates for the parameters of optimizees, by learning from their performance on sample tasks; and at the *meta-testing* stage, the L2O optimizer freezes its parameters and is used to solve new optimizees. In general, L2O optimizers can efficiently solve optimizees that are similar to those seen during the meta-training stage, or are drawn from the same “task distribution”.

However, new unseen optimizees may substantially deviate from the training task distribution. As L2O optimizers predict updates to variables based on the dynamics of the optimization tasks, such as gradients, different task distributions can lead to significant dissimilarity in task dynamics. Therefore, L2O optimizers often incur inferior performance when faced with these distinct unseen optimizees.

Such challenges have been widely observed and studied in related fields. For example, in the domain of meta-learning (Finn et al., 2017; Nichol & Schulman, 2018), we aim to enable neural networks to

be fast adapted to new tasks with limited samples. Among these techniques, Model-Agnostic Meta Learning (MAML) (Finn et al., 2017) is one of the most widely-adopted algorithms. Specifically, in the meta-training stage, MAML makes inner updates for individual tasks and subsequently conducts back-propagation to aggregate the gradients of individual task gradients, which are used to update the meta parameters. This design enables the learned initialization (meta parameters) to be sensitive to each task, and well-adapted after few fine-tuning steps.

Motivated by this, we propose a novel algorithm, named M-L2O, that incorporates the meta-adaption design in the meta-training stage of L2O. In detail, rather than updating the L2O optimizer directly based on optimizee gradients, M-L2O introduces a nested structure to calculate optimizer updates by aggregating the gradients of meta-updated optimizees. By adopting such an approach, M-L2O is able to identify a well-adapted region, where only a few adaptation steps are sufficient for the optimizer to generalize well on unseen tasks. In summary, the contributions of this paper are outlined below:

- To address the unsatisfactory generalization of L2O on out-of-distribution tasks, we propose to incorporate a meta adaptation design into L2O training. It enables the learned optimizer to locate in well-adapted initial points, which can be fast adapted in only a few steps to new unseen optimizees.
- We theoretically demonstrate that our meta adaption design grants M-L2O optimizer faster adaption ability in out-of-distribution tasks, shown by better generalization errors. Our analysis further suggests that training-like adaptation tasks can yield better generalization performance, in contrast to the common practice of using testing-like tasks. Such theoretical findings are further substantiated by the experimental results.
- Extensive experiments consistently demonstrate that the proposed M-L2O outperforms various baselines, including vanilla L2O and transfer learning, in terms of the testing performance within a small number of steps, showing the ability of M-L2O to promptly adapt in practical applications.

## 2 RELATED WORKS

### 2.1 LEARNING TO OPTIMIZE

Learning to Optimize (L2O) captures optimization rules in a data-driven way, and the learned optimizers have demonstrated success on various tasks, including but not limited to black-box (Chen et al., 2017), Bayesian (Cao et al., 2019), minimax optimization problems (Shen et al., 2021), domain adaptation (Chen et al., 2020b; Li et al., 2020), and adversarial training (Jiang et al., 2018; Xiong & Hsieh, 2020). The success of L2O is based on the parameterized optimization rules, which are usually modeled through a long short-term memory network (Andrychowicz et al., 2016), and occasionally as multi-layer perceptrons (Vicol et al., 2021). Although the parameterization is practically successful, it comes with the “curse” of generalization issues. Researchers have established two major directions for improving L2O generalization ability: the first focuses on the generalization to similar optimization tasks but *longer* training iterations. For example, Chen et al. (2020a) customized training procedures with curriculum learning and imitation learning, and Lv et al. (2017); Li et al. (2020) designed rich input features for better generalization. Another direction focuses on the generalization to different optimization tasks: Chen et al. (2021b) studied the generalization for LISTA network on unseen problems, and Chen et al. (2020c) provided theoretical understandings to hybrid deep networks with learned reasoning layers. In comparison, our work theoretically studies general L2O and our proposals generalization performance under task distribution shifts.

### 2.2 FAST ADAPTATION

Fast adaptation is one of the major goals in the meta-learning area Finn et al. (2017); Nichol & Schulman (2018); Lee & Choi (2018) which often focuses on generalizing to new tasks with limited samples. MAML Finn et al. (2017), a famous and effective meta-learning algorithm, utilizes the nested loop for meta-adaption. Following this trend, numerous meta-learning algorithms chose to compute meta updates more efficiently. For example, FOMAML (Finn et al., 2017) only updated networks by first-order information; Reptile (Nichol & Schulman, 2018) introduced an extra intermediate variable to avoid Hessian computation; HF-MAML (Fallah et al., 2020) approximated the one-step meta update by Hessian-vector production; and Ji et al. (2022) adopted a multi-step approximation in updates. Meanwhile, many researchers designed algorithms to compute meta updates more wisely. For example, ANIL (Raghu et al., 2020) only updated the head of networks in the inner loop; HSML (Yao et al., 2019) tailored the transferable knowledge to different tasks; and MT-net (Lee & Choi, 2018) enabled meta-learner to learn on each layer’s activation space. In terms of theories,

(Fallah et al., 2021) measured the generalization error of MAML; Fallah et al. (2020) captured the single inner step MAML convergence rate by Hessian vector approximation. Furthermore, Ji et al. (2022) characterized the multiple-step MAML convergence rate. Recently, LFT Zhao et al. (2022) combines the meta-learning design in Learning to Optimize and demonstrates its better performance for adversarial attack applications.

### 3 PROBLEM FORMULATION AND ALGORITHM

In this section, we firstly introduce the formulation of L2O, and subsequently propose M-L2O for generalizable self-adaptation.

#### 3.1 L2O PROBLEM DEFINITION

Most machine learning algorithms adopt analytical optimizer, *e.g.* SGD, to compute parameter updates for general loss functions (we call it **optimizee** or **task**). Instead, L2O aims to estimate such updates by a model (usually a neural network), which we call **optimizer**. Specifically, the L2O optimizer takes the optimizee information (such as loss values and gradients) as input and generates updates to the optimizee. In this work, our objective is to learn the initialization of the L2O optimizer on training optimizees and subsequently finetune it on adaptation optimizees. Finally, we apply such an adapted optimizer to optimize the testing optimizees and evaluate their performance.

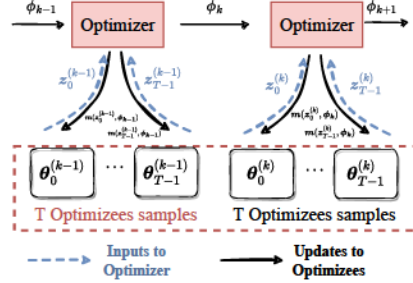


Figure 1: The pipeline of L2O problems.

We define  $\mathcal{L}(\theta; \phi)$  as the loss function where  $\theta$  is the **optimizee**'s parameter,  $\phi$  denotes the data sample, then the **optimizee** empirical and population risks are defined as below:

In L2O, the **optimizee**'s parameter  $\theta$  is updated by the **optimizer**, an update rule parameterized by  $\phi$  and we formulate it as  $\theta \leftarrow \mathcal{U}(\theta; \phi)$ . Specifically,  $\mathcal{U}$  denotes the optimizer model input. It captures the  $t$ -th iteration's optimizee information and parameterized by  $\theta_t$  with the data batch  $\phi_t$ . Then, the update rule of  $\theta$  in  $t$ -th iteration is shown as below:

(1)

The above pipeline is also summarized in Figure 1 where  $k$  denotes the update epoch of the optimizer. Note that  $\phi_t$  refers to the data batch of size  $T$  used at  $t$ -th iteration while  $\theta_t$  refers to the  $t$ -th single sample in data batch  $\phi_t$ . For theoretical analysis, we only consider taking the optimizee's gradients as input to the optimizer for update, i.e.,  $\nabla_{\theta} \mathcal{L}(\theta; \phi_t)$ . Therefore, the optimizee's gradient at  $t$ -th iteration over the optimizer  $\mathcal{U}$  takes a form of

(2)

where we assume that  $\nabla_{\theta} \mathcal{L}(\theta; \phi_t)$  is independent from the optimizee's initial parameter  $\theta_0$  and all samples are independent. The detailed derivation is shown in Lemma 1 in the Appendix.

Next, we consider a common initial parameter  $\theta_0$  for all optimizees and define the **optimizer** empirical and population risks w.r.t.  $\theta_0$  as below:

(3)

where  $\mathcal{U}$  updates in such a fashion:  $\theta_t \leftarrow \mathcal{U}(\theta_{t-1}; \phi_t)$ . Typically, the L2O optimizer is evaluated and updated after updating the optimizees for  $T$ -th iterations. Therefore, the optimal points of the optimizer risks in Equation (3) are defined as below:

(4)

### 3.2 M-L2O ALGORITHM

To improve the generalization ability of L2O, we aim at learning a well-adapted optimizer initialization by introducing a nested meta-update architecture. Instead of directly updating  $\theta$ , we adapt it for one step (namely an inner update) and define a new empirical risk for the optimizer as follows:

(5)

where  $\eta$  is the step size for inner updates. Consequently, the optimal point of the corresponding updated optimizer is:

(6)

Based on such an optimizer loss, we introduce M-L2O in Algorithm 1. Such a nested update design has been proved effective in the field of meta-learning, particularly MAML (Finn et al., 2017). Note that Algorithm 1 only returns the well-adapted optimizer initial point, which would require further adaptation in practice. We first denote the optimizers for training, adaptation, and testing by  $\theta_t$ ,  $\theta_a$ , and  $\theta_e$ , respectively, to distinguish tasks seen in different stages. Next, we obtain the results of meta training, denoted by  $\theta^*$ , via Algorithm 1, and we further adapt it based on  $\theta^*$ . The testing loss of M-L2O can be expressed as follows:

(7)

where  $\mathcal{L}_e$  denotes the meta testing loss. Note that  $\mathcal{L}_t$  refers to empirical risk and  $\mathcal{L}_p$  refers to population risk.

---

**Algorithm 1** Our Proposed M-L2O.

---

```

1: Input: Inner step size  $\eta$ , Outer learning stepsize  $\epsilon$ , Total epochs  $E$ , Epoch number per task  $E_t$ ,
   Optimizer initial point  $\theta_0$ , Training task  $\mathcal{T}$ , Adaptation task  $\mathcal{A}$ , Testing task  $\mathcal{E}$ 
2: for  $e = 1$  do
3:   if mod  $(e, E_t) = 0$  :  $\theta = \theta_0$  (random initial) else
4:     for  $t = 1$  do
5:        $\theta_t = \text{Update}(\theta, \mathcal{T}, \eta)$ 
6:     end for
7:     Compute  $\mathcal{L}_t(\theta_t)$ 
8:   update:  $\theta = \theta_t$ 
9: end for
10: Output:  $\theta^*$  Testing Loss:  $\mathcal{L}_e(\theta^*)$ 

```

---

## 4 GENERALIZATION THEOREM OF M-L2O

In this section, we introduce several assumptions and characterize M-L2O’s generalization ability.

### 4.1 TECHNICAL ASSUMPTIONS AND DEFINITIONS

To characterize the generalization of meta adaptation, it is necessary to make strong convexity assumptions which have been widely observed (Li & Yuan, 2017; Du et al., 2019) under over-parameterization condition and adopted in the geometry of functions (Fallah et al., 2021; Finn et al., 2019; Ji et al., 2021).

**Assumption 1.** We assume that the function  $\mathcal{L}_t$  is strongly convex. This assumption also holds for stochastic  $\mathcal{L}_t$ .

To capture the relationship between the L2O function  $\mathcal{L}_t$  and the M-L2O function  $\mathcal{L}_e$ , we make the following optimal point assumptions.

**Assumption 2.** We assume there exists a non-trivial optimizer optimal point  $\theta^*$  which is defined in Equation (6). Non-trivial means that  $\theta^* \neq \theta_0$  and  $\theta^* \neq \theta_e$ . Then, based on the definition of  $\theta^*$  and the existence of trivial solutions, for any  $\theta$ , we have the following equation:

where  $\eta$  is the step size for inner update of the optimizer. Note that we have defined the strongly-convex of landscape in Assumption 1 which validates the uniqueness of  $\theta^*$ .

The aforementioned assumption claims that there exist meta optimal points which are different from original task optimal points. Such an assumption is natural in experimental view. In MAML experiments where a single training task is considered, it is reasonable to expect that the solution would converge towards a well-adapted point instead of the task optimal point. Otherwise, MAML would be equivalent to simple transfer learning, where the learned point may not generalize well to new tasks.

**Assumption 3.** We assume that  $\mathcal{L}_t$  is  $L_t$ -Lipschitz,  $\mathcal{L}_a$  is  $L_a$ -Lipschitz and  $\mathcal{L}_m$  is  $L_m$ -Lipschitz for each loss function  $\mathcal{L}$ . This assumption also holds for stochastic  $\mathcal{L}$ , and  $\mathcal{L}_t$  is  $L_t$ -Lipschitz w.r.t.  $\theta$ . We further assume that  $\mathcal{L}_a$  is  $L_a$ -Lipschitz w.r.t.  $\theta$ ,  $\mathcal{L}_m$  is  $L_m$ -Lipschitz w.r.t.  $\theta$  and  $\mathcal{L}_m$  is  $L_m$ -Lipschitz.

The above Lipschitz assumptions are widely adopted in previous optimization works (Fallah et al., 2021; 2020; Ji et al., 2022; Zhao et al., 2022). To characterize the difference between tasks for meta training and meta adaptation, we define  $\Delta \mathcal{L}_t$  and  $\Delta \mathcal{L}_a$  as follows:

**Assumption 4.** We assume there exist union bounds  $\Delta \mathcal{L}_t$  and  $\Delta \mathcal{L}_a$  to capture the gradient and Hessian differences respectively between meta training task and adaptation task:

Such an assumption has been made similarly in MAML generalization works (Fallah et al., 2021).

#### 4.2 MAIN THEOREM

In this section, we theoretically analyze the generalization error of M-L2O and compare it with the vanilla L2O approach (Chen et al., 2017). Firstly, we characterize the difference of optimizer gradients between any two tasks ( $\mathcal{L}_t$  and  $\mathcal{L}_a$ ) in the following form:

**Proposition 1.** Based on Equation (2), Assumptions 3 and 4, we obtain

where  $\mathcal{O}_1$ ,  $\mathcal{O}_2$ ,  $\mathcal{O}_3$  and  $\mathcal{O}_4$  are defined in Assumption 4. Furthermore, we characterize the task difference of optimizer gradient as follows:

$$\mathcal{O}_1$$

where  $\mathcal{O}_1$ ,  $\mathcal{O}_2$  and  $\mathcal{O}_3$  are defined in Assumption 4.

Proposition 1 shows that the difference in optimizer gradient landscape scales exponentially with the optimizer iteration number  $k$ . Specifically, it involves the  $\mathcal{O}_1$  term with gradient difference and the  $\mathcal{O}_2$  term with Hessian difference  $\Delta \mathcal{H}_t$ . Clearly,  $\mathcal{O}_1$  dominates when  $k$  increases, which implies that the gradient gap between optimizers is the key component of the difference in optimizer gradient.

**Theorem 1.** Suppose that Assumptions 1, 2, 3 and 4 hold. Considering Algorithm 1 and Equation (7), if we define  $\mathcal{O}_1$ , set  $\mathcal{O}_2 = \mathcal{O}_3 = \mathcal{O}_4 = 0$  for  $k \geq 1$ . Then, with a probability at least  $1 - \epsilon$ , we obtain

$$\mathcal{O}_1 = \frac{\Delta \mathcal{L}_t}{L_t} \left( \frac{1 - \eta L_t}{1 - \eta L_m} \right)^k$$

$$\mathcal{O}_2$$

where  $\mathcal{O}_1$ ,  $\mathcal{O}_2$ ,  $\mathcal{O}_3$ ,  $\mathcal{O}_4$  and  $\mathcal{O}_5$  are defined in Assumption 4,  $\mathcal{O}_1$  is total epoch number for meta training,  $\mathcal{O}_2$  is the batch size for optimizer training.

To provide further understanding of generalization errors, we first make the following remark:

**Remark 1** (The choice of  $\eta$ ). *In Theorem 1, we set  $\eta = \frac{1}{\sqrt{m}}$ , thus the error term  $\mathcal{O}(\frac{1}{\sqrt{m}})$  vanishes with larger  $m$ . If we fix the iteration number  $k$ , then such an error term is determined by the gradient bound  $G$  and the Hessian bound  $H$ .*

The key components that lead to  $\frac{1}{\sqrt{m}}$  dependency are the Lipschitz properties to characterize the L2O loss landscape, e.g. the Lipschitz term  $\mathcal{O}(\frac{1}{\sqrt{m}})$  defined for  $\eta = \frac{1}{\sqrt{m}}$ . The reason is due to our nested update procedure  $\eta = \frac{1}{\sqrt{m}}$ . If we take the gradient of the last update term  $\eta = \frac{1}{\sqrt{m}}$  over  $\eta$ , then it requires us to compute gradients iteratively for all  $\eta = \frac{1}{\sqrt{m}}$ , which leads to the exponential term.

Consequently, it can be observed that the generalization error of M-L2O can be decomposed into three components: (i) The first term determined by  $\frac{1}{\sqrt{m}}$  is dominated by the training epoch  $k$ . Such an error term characterizes how meta training influences the generalization; (ii) The second term  $\frac{1}{\sqrt{m}}$  reflects the empirical error introduced by limited samples; hence it is controlled by the sample size  $m$ . (iii) The last two error terms capture task differences. Specifically,  $\frac{1}{\sqrt{m}}$  measures the gap between training and testing optimal points, while  $\frac{1}{\sqrt{m}}$  and  $\frac{1}{\sqrt{m}}$ , which dominate the last error term and represent the gradient and Hessian union bounds, respectively, reflect the geometry difference between training and adaptation tasks.

For better comparison with L2O, we make the following remark about generalization of M-L2O and Transfer Learning.

**Remark 2** (Comparison with Transfer Learning). *We can rewrite the generalization error of M-L2O in Theorem 1 in the following form:*

*For L2O Transfer Learning, the generalization error is shown as below:*

where  $\theta^*$  and  $\theta^*$  represents transfer learning L2O learned point after  $k$  epochs.

The generalization error gap between M-L2O and Transfer Learning can be categorized into two parts:

(i) Difference between  $\theta^*$  and  $\theta^*$ ; (ii) Difference between  $\theta^*$  and  $\theta^*$ . If we assume  $\theta^*$ , then both differences can be characterized by the gap between  $\theta^*$  and  $\theta^*$ . Since  $\theta^*$  is trained to converge to  $\theta^*$  as  $k$  increases, it is natural to see that M-L2O ( $\theta^*$ ) enjoys smaller generalization error compared to Transfer Learning ( $\theta^*$ ).

We further distinguish our theory from previous theoretical works. In the L2O area, Heaton et al. (2020) analyzed the convergence of proposed safe-L2O but not generalization while Chen et al. (2020c) analyzed the generalization of quadratic-based L2O. Instead, we develop the generalization on a general class of L2O problems. In the meta-learning area, the previous works have demonstrated the convergence and generalization of MAML (Ji et al., 2022; Fallah et al., 2021). Instead, we leverage the MAML results in L2O domain to measure the learned point and training optimal point distance. Then, our L2O theory further characterizes transferability of learned point on meta testing tasks. Overall, our developed theorem is based on both L2O and meta learning results. In conclusion, our theoretical novelty lies in three aspects: ① Rigorously characterizing a generic class of L2O generalization. ② Incorporating the MAML results in our meta-learning analysis. ③ Theoretically proving that both training-like and testing-like adaptation contribute to better generalization in L2O.

## 5 EXPERIMENTS

In this section, we provide a comprehensive description of the experimental settings and present the results we obtained. Our findings demonstrate a high degree of consistency between the empirical observations and the theoretical outcomes.

### 5.1 EXPERIMENTAL CONFIGURATIONS

**Backbones and observations.** For all our experiments, we use a single-layer LSTM network with hidden units as the backbone. We adopt the methodology proposed by [Lv et al. \(2017\)](#) and [Chen et al. \(2020a\)](#) to utilize the parameters’ gradients and their corresponding normalized momentum to construct the observation vectors.

**Optimizees.** We conduct experiments on three distinct optimizees, namely LASSO, Quadratic, and Rosenbrock ([Rosenbrock, 1960](#)). The formulation of the Quadratic problem is  $\mathbf{x} - \mathbf{A} \mathbf{b}$  and the formulation of the LASSO problem is  $\mathbf{x} - \mathbf{A} \mathbf{b}$ , where  $\mathbf{A}$ ,  $\mathbf{b}$ . We set . The precise formulation of the Rosenbrock problem is available in Section A.6. During the meta-training and testing stage, the optimizees  $\text{train}$  and  $\text{test}$  are drawn from the pre-specified distributions  $\text{train}$  and  $\text{test}$ , respectively. Similarly, the optimizees  $\text{adapt}$  used during adaptation are sampled from the distribution  $\text{adapt}$ .

**Baselines and Training Settings.** We compare M-L2O against three baselines: ( ) *Vanilla L2O*, where we train a randomly initialized L2O optimizer on  $\text{adapt}$  for only steps; ( ) *Transfer Learning* (TL), where we first meta-train a randomly initialized L2O optimizer on  $\text{train}$ , and then fine-tune on  $\text{adapt}$  for steps; ( ) *Direct Transfer* (DT), where we meta-train a randomly initialized L2O optimizer on  $\text{train}$  only. M-L2O adopts a fair experimental setting, whereby we meta-train on  $\text{train}$  and adapt on the same  $\text{adapt}$ . We evaluate these methods using the same set of optimizees for testing (*i.e.*,  $\text{test}$ ), and report the minimum *logarithmic* value of the objective functions achieved for these optimizees.

For all experiments, we set the number of optimizee iterations, denoted by , to when meta-training the L2O optimizers and adapting to optimizees. Notably, in large scale experiments involving neural networks as tasks, the common choice for is ([Zheng et al., 2022](#); [Shen et al., 2021](#)). However, in our experiments, we set to to achieve better experimental performance. The value of the total epochs, denoted by , is set to , and we adopt the curriculum learning technique ([Chen et al., 2020a](#)) to dynamically adjust the number of epochs per task, denoted by . To update the weights of the optimizers ( ), we use Adam ([Kingma & Ba, 2014](#)) with a fixed learning rate of .

### 5.2 FAST ADAPTATION RESULTS OF M-L2O

**Experiments on LASSO optimizees.** We begin with experiments on LASSO optimizees. Specifically, for  $\text{train}$ , the coefficient matrix  $\mathbf{A}$  is generated by sampling from a mixture of uniform distributions comprising . In contrast, for  $\text{test}$  and  $\text{adapt}$ , the coefficient matrices  $\mathbf{A}$  are obtained by sampling from a normal distribution with a standard deviation of . We conduct experiments with and , and report the results in Figures 2a and 2b. Our findings demonstrate that:

- ① The Vanilla L2O approach, which relies on only five steps of adaptation from initialization on  $\text{adapt}$ , exhibits the weakest performance, as evidenced by the largest values of the objective function.
- ② Although Direct Transfer (DT) is capable of learning optimization rules from the training optimizees, the larger variance in coefficients among the testing optimizees renders the learned rules inadequate for generalization.
- ③ The superiority of Transfer Learning (TL) over DT highlights the values of adaptation when the testing optimizees deviates significantly from those seen in training, as the optimizer is presumably able to acquire new knowledge during the adaptation process.
- ④ Finally, M-L2O exhibits consistent and notably faster convergence speed compared to other baseline methods. Moreover, it demonstrates the best performance overall, reducing the logarithm of the objective values by approximately and when and , respectively. M-L2O’s superior performance can be attributed to its ability to learn well-adapted initial weights for optimizers, which enables rapid self-adaptation, thus leading to better performance in comparison to the baseline methods.

Table 1: Minimum logarithm loss of different methods on LASSO at different levels of  $\sigma$ . We report the 95% confidence interval from 10 repeated runs.

|     | Methods     |               |               |        |
|-----|-------------|---------------|---------------|--------|
|     | Vanilla L2O | M-L2O         | DT            | TL     |
| 10  | 0.033       | -3.712        | <b>-4.233</b> | -4.077 |
| 25  | 1.559       | -3.433        | <b>-4.125</b> | -4.019 |
| 50  | 0.550       | <b>-3.285</b> | -3.098        | -3.181 |
| 100 | 2.435       | <b>-2.408</b> | -1.775        | -1.961 |
| 200 | 4.104       | <b>-1.396</b> | -0.453        | -0.982 |

**Experiments on Quadratic optimizers.** We continue to assess the performance of our approach on a different optimizer, *i.e.*, the Quadratic problem. The coefficient matrices  $\mathbf{A}$  of the optimizers are also randomly sampled from a normal distribution. We conduct two evaluations, with values of  $\sigma = 50$  and  $\sigma = 100$ , respectively, and present the outcomes in Figure 3. Notably, the results show a similar trend to those we obtained in the previous experiments.

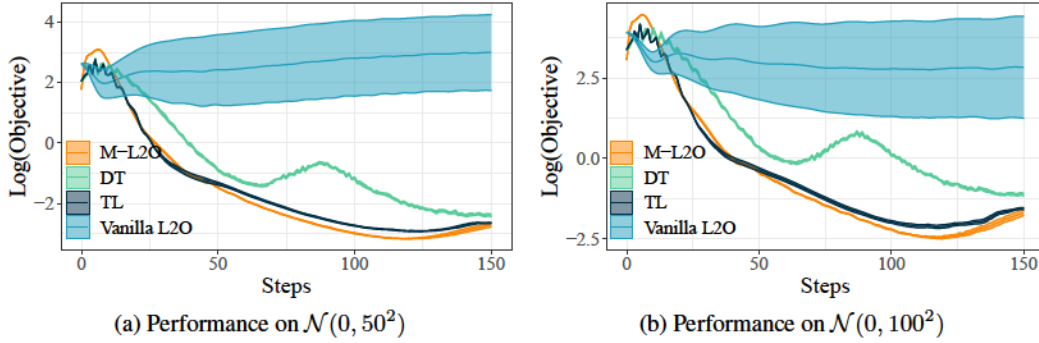


Figure 2: Comparison of convergence speeds on target distribution of LASSO optimizers. We repeat the experiments for 10 times and show the 95% confidence intervals in the figures.

**More LASSO experiments.** We proceed to investigate the impact of varying the standard deviation of the distributions we used to sample the coefficient matrices  $\mathbf{A}$  for  $\text{adapt}$  and  $\text{test}$ . The minimum logarithm of the objective value for each method is reported in Table 1. Our findings reveal that:

① At lower levels of  $\sigma$ , it is not always necessary, and may even be unintentionally harmful, to use adaptation for normally trained L2O. Although M-L2O produces satisfactory results, it exhibits slightly lower performance than TL, which could be due to the high similarity between the training and testing tasks. Since M-L2O’s objective is to identify optimal general initial points, L2O optimizers trained directly on related and similar tasks may effectively generalize. However, after undergoing adaptation on a single task, L2O optimizers may discard certain knowledge acquired during meta-training that could be useful for novel but similar tasks.

② Nevertheless, as the degree of similarity between training and testing tasks is declines, as characterized by an increasing value of  $\sigma$ , M-L2O begins to demonstrate considerable advantage. For values of  $\sigma$  greater than 100, M-L2O exhibits consistent performance advantages that exceed 1 in terms of logarithmic loss. This observation empirically supports that the learned initial weights facilitate rapid adaptation to new tasks that are “out-of-distribution”, and that manifest large deviations.

### 5.3 ADAPTATION WITH SAMPLES FROM DIFFERENT TASK DISTRIBUTION

In Section 5.2, we impose a constraint on the standard deviation of the distribution used to sample  $\mathbf{A}$ , ensuring that is identical for both the optimizers for adaptation and testing. However, it is noteworthy that this constraint is not mandatory, given that our theory can accommodate adaptation and testing optimizers with different distributions. Consequently, we conduct an experiment on LASSO optimizers with varying standard deviations of the distribution, from which the matrices  $\mathbf{A}$  for optimizer  $\text{adapt}$  is drawn. Specifically, we sample  $\sigma$  with smaller values that more resemble the training tasks, as well as larger values that are more similar to the testing task ( $\sigma = 100$ ).

In Theorem 1, we have characterized the generalization of M-L2O with flexible distribution adaptation tasks. The theoretical analysis suggests that a similar geometry landscape (smaller

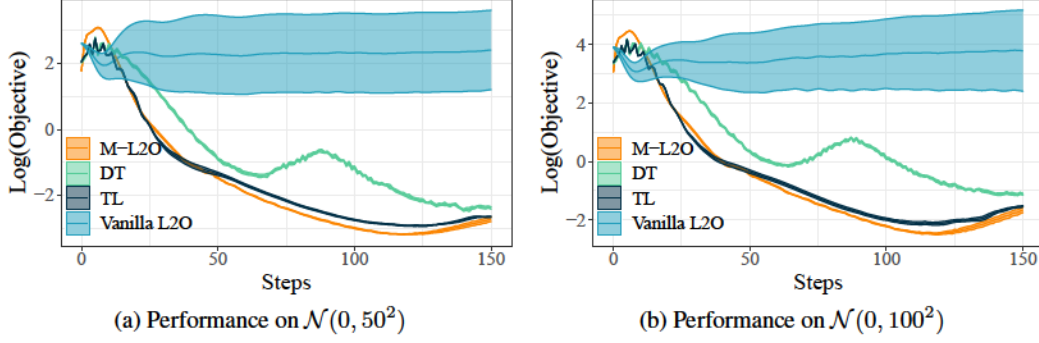


Figure 3: Comparison of convergence speeds on target distribution of Quadratic optimizers. We repeat the experiments for 10 times and show the 95% confidence intervals are shown in the figures.

$\Delta_{12}, \tilde{\Delta}_{12}$ ) between the training and adaptation tasks, can lead to a reduction in generalization loss as defined in Equation (7). This claim has been corroborated by the results of our experiments, as presented in Figure 4. When the  $\sigma$  is similar to the training tasks (e.g., 10), implying a smaller  $\Delta_{12}$  and  $\tilde{\Delta}_{12}$ , M-L2O demonstrates superior testing performance. In conclusion, incorporating training-like adaptation tasks can lead to better generalization performance.

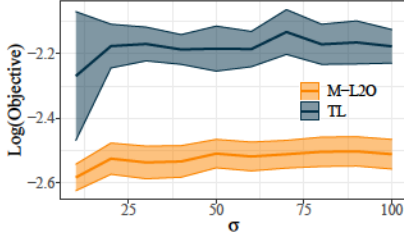


Figure 4: Performance on LASSO optimizers. We vary the standard deviation of the distribution used for sampling the weight matrix  $A$  for adaptation optimizers. We visualize both the mean and the confidence interval in the figure.

Meanwhile, it is reasonable to suggest that the task differences between adaptation and testing, denoted by  $(\Delta_{23}, \tilde{\Delta}_{23})$ , may also have an impact on M-L2O’s generalization ability. Intuitively, if the optimizer is required to adapt to testing optimizers, the adapted optimizer should demonstrate strong generalization ability on other optimizers that are similar. In order to have a deeper understanding of the relationship between the generalization ability and the difference between adaptation and testing tasks, we rewrite M-L2O generalization error in Theorem 1 in the

following form with  $\tilde{\phi}_*^3 = \tilde{\phi}_{M*}^3 - \alpha \nabla_{\phi} \hat{g}_T^3(\tilde{\phi}_{M*}^3)$  and  $\tilde{\delta}_{13}^M = \|\tilde{\phi}_{M*}^1 - \tilde{\phi}_{M*}^3\|$ :

$$\begin{aligned} & g_T^3(\tilde{\phi}_{MK}^1 - \alpha \nabla_{\phi} \hat{g}_T^2(\tilde{\phi}_{MK}^1)) - g_T^3(\phi_*^3) \\ & \leq M_{g_T} (\|\tilde{\phi}_{MK}^1 - \tilde{\phi}_{M*}^1\| + \|\tilde{\phi}_*^3 - \phi_*^3\| + \alpha \|\nabla_{\phi} \hat{g}_T^2(\tilde{\phi}_{MK}^1) - \nabla_{\phi} \hat{g}_T^3(\tilde{\phi}_{M*}^3)\| + \tilde{\delta}_{13}^M). \end{aligned} \quad (8)$$

In Equation (8), M-L2O generalization error is partly captured by  $\|\nabla_{\phi} \hat{g}_T^2(\tilde{\phi}_{MK}^1) - \nabla_{\phi} \hat{g}_T^3(\tilde{\phi}_{M*}^3)\|$  which is controlled by difference in optimizers (i.e.,  $\|\nabla_{\phi} \hat{g}_T^2(\phi) - \nabla_{\phi} \hat{g}_T^3(\phi)\|$ ). From Proposition 1, we know that this term is determined by difference in optimizers, denoted by  $\Delta_{23}$  and  $\tilde{\Delta}_{23}$ . Similar to the results established in Theorem 1, we can deduce that superior testing performance is connected with a smaller difference between testing and adaptation optimizers. This result has been demonstrated in Figure 4 where TL generalizes well with larger  $\sigma$  (more testing-like). Moreover, M-L2O also benefits from larger  $\sigma$  values (e.g.,  $\sigma = 100$ ) in certain scenarios.

To summarize, both training-like and testing-like adaptation task can lead to improved testing performance. As shown in Figure 4, training-like adaptation results in better generalization in L2O. One possible explanation is that when the testing task significantly deviates from the training tasks, it becomes highly challenging for the optimizer to generalize well within limited adaptation steps. In such scenarios, the training-like adaptation provides a more practical solution.

## 6 CONCLUSION AND DISCUSSION

In this paper, we propose a self-adapted L2O algorithm (M-L2O), which is incorporated with meta adaptation. Such a design enables the optimizer to reach a well-adapted initial point, facilitating its adaptation ability with only a few updates. Our superior generalization performances in out-of-distribution tasks have been theoretically characterized and empirically validated across various scenarios. Furthermore, the comprehensive empirical results demonstrate that training-like adaptation tasks can contribute to better testing generalization, which is consistent with our theoretical analysis. One potential future direction is to develop a convergence analysis for L2O. It will be more interesting to consider meta adaptation in analyzing L2O convergence from a theoretical view.

## ACKNOWLEDGEMENTS

The work of Y. Liang was supported in part by the U.S. National Science Foundation under the grants ECCS-2113860 and DMS-2134145. The work of Z. Wang was supported in part by the U.S. National Science Foundation under the grant ECCS-2145346 (CAREER Award).

## REFERENCES

- Marcin Andrychowicz, Misha Denil, Sergio Gomez, Matthew W Hoffman, David Pfau, Tom Schaul, Brendan Shillingford, and Nando De Freitas. Learning to learn by gradient descent by gradient descent. In *Advances in neural information processing systems (NeurIPS)*, 2016.
- Yue Cao, Tianlong Chen, Zhangyang Wang, and Yang Shen. Learning to optimize in swarms. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 15018–15028, 2019.
- Tianlong Chen, Weiyi Zhang, Jingyang Zhou, Shiyu Chang, Sijia Liu, Lisa Amini, and Zhangyang Wang. Training stronger baselines for learning to optimize. *arXiv preprint arXiv:2010.09089*, 2020a.
- Tianlong Chen, Xiaohan Chen, Wuyang Chen, Howard Heaton, Jialin Liu, Zhangyang Wang, and Wotao Yin. Learning to optimize: A primer and a benchmark. *arXiv preprint arXiv:2103.12828*, 2021a.
- Wuyang Chen, Zhiding Yu, Zhangyang Wang, and Animashree Anandkumar. Automated synthetic-to-real generalization. In *International Conference on Machine Learning (ICML)*, pp. 1746–1756, 2020b.
- Xiaohan Chen, Jialin Liu, Zhangyang Wang, and Wotao Yin. Hyperparameter tuning is all you need for lista. *Advances in Neural Information Processing Systems*, 34:11678–11689, 2021b.
- Xinshi Chen, Yufei Zhang, Christoph Reisinger, and Le Song. Understanding deep architecture with reasoning layer. *Advances in Neural Information Processing Systems*, 33:1240–1252, 2020c.
- Yutian Chen, Matthew W Hoffman, Sergio Gómez Colmenarejo, Misha Denil, Timothy P Lillicrap, Matt Botvinick, and Nando De Freitas. Learning to learn without gradient descent by gradient descent. In *International Conference on Machine Learning (ICML)*, pp. 748–756, 2017.
- Simon S Du, Xiyu Zhai, Barnabas Poczos, and Aarti Singh. Gradient descent provably optimizes over-parameterized neural networks. In *International Conference on Learning Representations (ICLR)*, 2019.
- Alireza Fallah, Aryan Mokhtari, and Asuman Ozdaglar. On the convergence theory of gradient-based model-agnostic meta-learning algorithms. In *International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2020.
- Alireza Fallah, Aryan Mokhtari, and Asuman Ozdaglar. Generalization of model-agnostic meta-learning algorithms: Recurring and unseen tasks. *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *International Conference on Machine Learning (ICML)*, 2017.
- Chelsea Finn, Aravind Rajeswaran, Sham Kakade, and Sergey Levine. Online meta-learning. In *International Conference on Machine Learning (ICML)*, pp. 1920–1930, 2019.
- Howard Heaton, Xiaohan Chen, Zhangyang Wang, and Wotao Yin. Safeguarded learned convex optimization. *arXiv preprint arXiv:2003.01880*, 2020.
- Kaiyi Ji, Junjie Yang, and Yingbin Liang. Bilevel optimization: Convergence analysis and enhanced design. In *International Conference on Machine Learning (ICML)*, pp. 4882–4892, 2021.
- Kaiyi Ji, Junjie Yang, and Yingbin Liang. Theoretical convergence of multi-step model-agnostic meta-learning. *Journal of Machine Learning Research (JMLR)*, 2022.

- Haoming Jiang, Zhehui Chen, Yuyang Shi, Bo Dai, and Tuo Zhao. Learning to defense by learning to attack. *arXiv preprint arXiv:1811.01213*, 2018.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Yoonho Lee and Seungjin Choi. Gradient-based meta-learning with learned layerwise metric and subspace. In *International Conference on Machine Learning (ICML)*, 2018.
- Chaojian Li, Tianlong Chen, Haoran You, Zhangyang Wang, and Yingyan Lin. Halo: Hardware-aware learning to optimize. In *European Conference on Computer Vision (ECCV)*, pp. 500–518. Springer, 2020.
- Yuanzhi Li and Yang Yuan. Convergence analysis of two-layer neural networks with relu activation. *Advances in neural information processing systems (NeurIPS)*, 2017.
- Kaifeng Lv, Shunhua Jiang, and Jian Li. Learning gradient descent: Better generalization and longer horizons. In *International Conference on Machine Learning (ICML)*, pp. 2247–2255, 2017.
- Alex Nichol and John Schulman. Reptile: a scalable metalearning algorithm. *arXiv preprint arXiv:1803.02999*, 2(3):4, 2018.
- Ning Qian. On the momentum term in gradient descent learning algorithms. *Neural networks*, 12(1): 145–151, 1999.
- Aniruddh Raghu, Maithra Raghu, Samy Bengio, and Oriol Vinyals. Rapid learning or feature reuse? towards understanding the effectiveness of maml. In *International Conference on Learning Representations*, 2020.
- HoHo Rosenbrock. An automatic method for finding the greatest or least value of a function. *The computer journal*, 3(3):175–184, 1960.
- Sebastian Ruder. An overview of gradient descent optimization algorithms. *arXiv preprint arXiv:1609.04747*, 2016.
- Shai Shalev-Shwartz, Ohad Shamir, Nathan Srebro, and Karthik Sridharan. Learnability, stability and uniform convergence. *The Journal of Machine Learning Research (JMLR)*, 2010.
- Jiayi Shen, Xiaohan Chen, Howard Heaton, Tianlong Chen, Jialin Liu, Wotao Yin, and Zhangyang Wang. Learning a minimax optimizer: A pilot study. In *International Conference on Learning Representations (ICLR)*, 2021.
- Laurits Tani, Diana Rand, Christian Veelken, and Mario Kadastik. Evolutionary algorithms for hyperparameter optimization in machine learning for application in high energy physics. *The European Physical Journal C*, 81(2):1–9, 2021.
- Paul Vicol, Luke Metz, and Jascha Sohl-Dickstein. Unbiased gradient estimation in unrolled computation graphs with persistent evolution strategies. In Marina Meila and Tong Zhang (eds.), *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pp. 10553–10563. PMLR, 18–24 Jul 2021. URL <https://proceedings.mlr.press/v139/vicol21a.html>.
- Olga Wichrowska, Niru Maheswaranathan, Matthew W Hoffman, Sergio Gomez Colmenarejo, Misha Denil, Nando de Freitas, and Jascha Sohl-Dickstein. Learned optimizers that scale and generalize. In *International Conference on Machine Learning (ICML)*, 2017.
- Yuanhao Xiong and Cho-Jui Hsieh. Improved adversarial training via learned optimizer, 2020.
- Huaxiu Yao, Ying Wei, Junzhou Huang, and Zhenhui Li. Hierarchically structured meta-learning. In *International Conference on Machine Learning (ICML)*, 2019.
- Pu Zhao, Parikshit Ram, Songtao Lu, Yuguang Yao, Djallel Bouneffouf, Xue Lin, and Sijia Liu. Learning to generate image source-agnostic universal adversarial perturbations. In *International Joint Conference on Artificial Intelligence (IJCAI)*, 2022.
- Wenqing Zheng, Tianlong Chen, Ting-Kuei Hu, and Zhangyang Wang. Symbolic learning to optimize: Towards interpretability and scalability. *arXiv preprint arXiv:2203.06578*, 2022.

## A APPENDIX

### A.1 RESTATEMENT OF ASSUMPTION 3

**Assumption 5** (Restatement of Assumption 3). *We assume Lipschitz properties for all functions as follows:*

- a)  $\mathcal{F}$  is  $L$ -Lipschitz, i.e., for any  $\mathbf{x}, \mathbf{y}$  and  $\mathbf{z}$ ,  $\|\mathcal{F}(\mathbf{x}) - \mathcal{F}(\mathbf{y})\| \leq L \|\mathbf{x} - \mathbf{y}\|$ .
- b)  $\mathcal{G}$  is  $L$ -Lipschitz, i.e., for any  $\mathbf{x}, \mathbf{y}$  and  $\mathbf{z}$ ,  $\|\mathcal{G}(\mathbf{x}) - \mathcal{G}(\mathbf{y})\| \leq L \|\mathbf{x} - \mathbf{y}\|$ .
- c)  $\mathcal{H}$  is  $L$ -Lipschitz, i.e., for any  $\mathbf{x}, \mathbf{y}$  and  $\mathbf{z}$ ,  $\|\mathcal{H}(\mathbf{x}) - \mathcal{H}(\mathbf{y})\| \leq L \|\mathbf{x} - \mathbf{y}\|$ .
- d)  $\mathcal{I}$  is  $L$ -Lipschitz w.r.t.  $\mathbf{x}$  and  $L$ -Lipschitz w.r.t.  $\mathbf{z}$ , i.e.,  
 $\|\mathcal{I}(\mathbf{x}, \mathbf{z}) - \mathcal{I}(\mathbf{y}, \mathbf{z})\| \leq L \|\mathbf{x} - \mathbf{y}\|$  and  $\|\mathcal{I}(\mathbf{x}, \mathbf{z}) - \mathcal{I}(\mathbf{x}, \mathbf{y})\| \leq L \|\mathbf{z} - \mathbf{y}\|$ .
- e)  $\mathcal{J}$  is  $L$ -Lipschitz, i.e., for any  $\mathbf{x}, \mathbf{y}$  and  $\mathbf{z}$ ,  $\|\mathcal{J}(\mathbf{x}, \mathbf{y}, \mathbf{z}) - \mathcal{J}(\mathbf{x}', \mathbf{y}, \mathbf{z})\| \leq L \|\mathbf{x} - \mathbf{x}'\|$ .

The above Assumptions (a)(b)(c) also hold for stochastic  $\mathcal{F}, \mathcal{G}, \mathcal{H}$  and  $\mathcal{I}$ .

### A.2 PROOF OF SUPPORTING LEMMAS (LEMMA 12 CORRESPONDS TO PROPOSITION 1)

**Lemma 1.** *Based on update procedure of  $\mathcal{I}$ , we obtain*

*Proof.* The update process is shown below:

If we only consider  $\mathcal{I}$ , then we obtain

If we iterate the above equation from  $t$  to  $t+1$ , then we obtain

We assume  $\mathbf{z}$  is randomly sampled and independent from  $\mathbf{x}, \mathbf{y}$ , then we obtain

□

**Lemma 2.** *If we assume that  $\mathcal{F}, \mathcal{G}, \mathcal{H}$  are Lipschitz continuous, based on Assumption 3, then we obtain*

$$\| \mathbf{x} - \mathbf{y} \| \leq \dots \quad (9)$$

*Proof.* Based on the iterate procedure of  $\mathcal{A}$ , we obtain

where  $\mathcal{A}^k$  follows from Equation (1),  $\mathcal{A}^{k+1}$  and  $\mathcal{A}^{k+2}$  from Assumption 3. If we further iterate it from  $\mathcal{A}^{k+2}$  to  $\mathcal{A}^{k+3}$ , we obtain

—

□

**Lemma 3.** *If we define  $\mathcal{A}^k$  as above, based on Assumption 3 and Lemma 2, we obtain*

where  $\mathcal{A}^k$  follows from Equation (1).

*Proof.* Based on the definition of  $\mathcal{A}^k$ , we have

where  $\mathcal{A}^k$  follows from Assumption 3,  $\mathcal{A}^{k+1}$  follows from Lemma 2. □

**Lemma 4.** *We first define  $\mathcal{A}^k$  as above. Based on the Lemma 2 and Assumption 3, we obtain*

where  $\mathcal{A}^k$  follows from Equation (1).

*Proof.* Based on the definition of  $\mathcal{L}_\text{KL}$ , we have

where  $\mathcal{L}_\text{KL}$  and  $\mathcal{L}_\text{KL}$  follows from Assumption 3,  $\mathcal{L}_\text{KL}$  follows from Lemma 2. □

**Lemma 5.** *Based on Assumption 3 and Lemmas 1, 3 and 4, then we obtain*

(10)

where

*Proof.* Based on the definition of  $\mathcal{L}_\text{KL}$  in Lemma 1, we obtain

where  $\gamma$  is based on Lemma 1,  $\beta$  is based on the fact that  $\gamma \leq 1$ ,  $\alpha$  follows from Assumption 3 and  $\beta$  follows from Lemma 3 and 4.  $\square$

**Lemma 6.** *Based on Lemmas 2, 5 and Assumption 3, we obtain*

where  $\gamma$  is defined in Lemma 5,  $\beta$  is defined in Lemma 2.

*Proof.* We assume all functions share the same starting point  $\mathbf{x}_0$ , then we have

where  $\gamma$  from Assumption 3,  $\beta$  from Lemma 2 and 5.  $\square$

**Lemma 7.** *Based on the Lemma 2, 5 and Assumption 3, we obtain*

(11)

where  $\gamma$  is defined in Lemma 5.

*Proof.* We first compute the Lipschitz condition of  $\mathbf{f}$  as follows

where  $\gamma$  follows from Lemma 2, 5 and Assumption 3. Then, based on the definition of  $\mathbf{f}$ , we have

where  $\epsilon_k$  follows from Lemma 2 and 5.

□

**Lemma 8.** *If we assume  $\epsilon_k \leq \epsilon$ , based on Assumption 3 and 4, we obtain*

where  $k$  is the iteration number and  $\epsilon_k = \epsilon - \frac{\epsilon}{k}$ .

*Proof.* Based on the iterative process of  $\epsilon_k$ , we obtain

where  $\epsilon_k$  follows from Equation (1),  $\epsilon_k$  follows from Assumption 3,  $\epsilon_k$  follows from Assumption 4. If we iterate above inequalities from  $k=1$  to  $k=k$ , then we obtain:

$$\epsilon_k = \epsilon - \frac{\epsilon}{k}$$

□

**Lemma 9.** *Based on Assumptions 3 and 4, Lemma 8, we have following inequality:*

where  $\epsilon_k = \epsilon - \frac{\epsilon}{k}$  and  $\epsilon_k = \epsilon - \frac{\epsilon}{k}$ .

*Proof.* Based on the definition of  $\epsilon_k$ , we obtain

where  $\epsilon_k$  follows from Assumption 3 and 4,  $\epsilon_k$  follows from Lemma 8.

□

**Lemma 10.** *Then based on Assumptions 3 and 4, Lemma 8, we have following inequality:*

where  $\alpha$  and  $\beta$  is defined in Lemma 8. ,

*Proof.* Based on the definition of  $\alpha$  , we obtain

where  $\alpha$  follow from Lemma 8. □

**Lemma 11.** *Based on Assumptions 3, 4 and Lemma 1, we obtain*

where  $\alpha$  and  $\beta$  have been defined in Lemmas 9 and 10.

*Proof.* Based on the Lemma 1, we obtain

where  $\mathcal{O}(\dots)$  follows from the definitions that  $\mathcal{O}(\dots)$ ,  $\mathcal{O}(\dots)$  follows from Assumption 3 and  $\mathcal{O}(\dots)$  follows from Lemma 9 and 10.  $\square$

**Lemma 12.** (Correspond to Proposition 1) Based on Assumptions 3 and 4, Lemmas 8 and 11, we obtain

$$\mathcal{O}(\dots)$$

where  $\mathcal{O}(\dots)$ .

*Proof.* We first consider  $\mathcal{O}(\dots)$  and  $\mathcal{O}(\dots)$ , we obtain

$$\mathcal{O}(\dots) \leq \mathcal{O}(\dots) \quad (12)$$

$$\mathcal{O}(\dots) \leq \mathcal{O}(\dots) \quad (13)$$

where  $\mathcal{O}(\dots)$  follows because  $\mathcal{O}(\dots) \leq \mathcal{O}(\dots)$ .

Furthermore, we consider the uniform bound for  $\mathcal{O}(\dots)$ , then we obtain

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

where  $\mathcal{O}(\dots)$  follows from Lemma 11,  $\mathcal{O}(\dots)$  follows from Equation (12) and Equation (13),  $\mathcal{O}(\dots)$  follows because  $\mathcal{O}(\dots)$ . Based on the formulation of  $\mathcal{O}(\dots)$  in Lemma 6, we have

$$\mathcal{O}(\dots)$$

where  $\mathcal{O}(\dots)$  follows because  $\mathcal{O}(\dots)$  defined in Lemma 2 satisfies that  $\mathcal{O}(\dots)$ .  $\square$

**Lemma 13.** Based on the Assumption 3 and Lemma 2, we obtain

where  $\mathcal{O}(\dots)$  and  $\mathcal{O}(\dots)$  is defined in Lemma 2.

*Proof.* Based on the definition of  $\mathcal{O}(\dots)$ , we have

where  $\mathcal{O}(\dots)$  is based on Assumption 3,  $\mathcal{O}(\dots)$  is based on Lemma 2.  $\square$

**Lemma 14.** Based on the proposition 1 in [Fallah et al. \(2021\)](#), Assumptions 1 and 3, if we set  $\dots$ ,  $\dots$  for  $\dots$  in Algorithm 1, then we have

$$\mathcal{O}(\dots)$$

where  $\dots$  is defined in Lemma 13,  $\dots$  is defined in Lemma 6,  $\dots$  is defined in Lemma 7.

*Proof.* Based on the Proposition 1 in [Fallah et al. \(2021\)](#), we obtain

$$\mathcal{O}(\dots)$$

where  $\dots$  is defined in Equation (5). Based on the Assumption 1 and the fact that  $\dots$ , we have

$$\mathcal{O}(\dots)$$

$\square$

**Lemma 15.** Based on Assumption 1 and Lemma 13, we have

$$\dots$$

where  $n$  is the sample size.

*Proof.* Based on Assumption 1 and Lemma 13, from Theorem 2 in [Shalev-Shwartz et al. \(2010\)](#), with probability at least  $\dots$ , we have

$$\dots$$

Furthermore, based on Assumption 1 and the fact that  $\dots$ , we obtain

$$\dots$$

We take the square root from both side and obtain:

$$\dots$$

with probability at least  $\dots$ .  $\square$

### A.3 PROOF OF THEOREM 1

Based on our definition of generalization error for the algorithm,

(14)

where  $\mathcal{O}(\dots)$  follows from Assumption 2,  $\mathcal{O}(\dots)$  follows from Lemma 13.

Furthermore, considering Algorithm 1, if we set  $\eta = \frac{1}{\sqrt{K}}$ ,  $\beta = \frac{1}{\sqrt{K}}$  for  $\eta, \beta$ , based on Lemma 12, 14, 15, with probability at least  $1 - \delta$ , we obtain

$$\mathcal{O}\left(\frac{1}{\sqrt{K}}\right) + \mathcal{O}\left(\frac{1}{\sqrt{K}}\right) + \mathcal{O}\left(\frac{1}{\sqrt{K}}\right) + \mathcal{O}\left(\frac{1}{\sqrt{K}}\right)$$

where  $\mathcal{O}(\dots)$ ,  $\mathcal{O}(\dots)$ ,  $\mathcal{O}(\dots)$  is the step number for update,  $\mathcal{O}(\dots)$  is the sample size for training.

Then for Lipschitz term  $\mathcal{O}(\dots)$  defined in Lemma 13,

$$\mathcal{O}(\dots)$$

where  $\mathcal{O}(\dots)$  defined in Lemma 2 satisfies  $\mathcal{O}(\dots)$ .

For Lipschitz term  $\mathcal{O}(\dots)$  defined in Lemma 6, we first compute the order for  $\mathcal{O}(\dots)$  which is defined in Lemma 5, then we obtain

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

where  $\mathcal{O}(\dots)$  follows from Lemmas 3 and 4. Then, we obtain

$$\mathcal{O}(\dots)$$

$$\mathcal{O}(\dots)$$

For Lipschitz term  $\mathcal{O}(\dots)$  defined in Lemma 7, we have

$$\mathcal{O}(\dots)$$

Then, the proof is complete.

#### A.4 PROOF OF REMARK 2

In terms of M-L2O generalization error, based on the Equation (14) in Appendix A.3, we have

where

In terms of Transfer Learning L2O generalization error with learned initial point , we have

where follows from , . Then, the proof is complete.

#### A.5 PROOF OF EQ. 8 IN SUBSECTION 5.3

We assume that , then we have

Then, the proof is complete.

#### A.6 ADDITIONAL EXPERIMENTS

**New Optimizees: Rosenbrock** We conduct additional experiments with substantially different optimizees, *i.e.* Rosenbrock (Rosenbrock, 1960). In this case, the optimizers are required to minimize a two-dimensional non-convex function taking the following formulation:

(15)

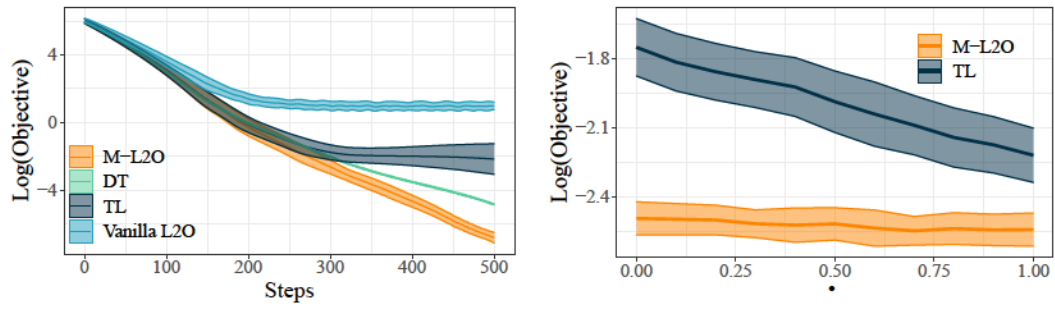
which is challenging for algorithms to converge to the global minimum (Tani et al., 2021).

We specify  $\mathcal{O}_{\text{adapt}}$  and  $\mathcal{O}_{\text{test}}$  to be the family of Rosenbrock optimizees with randomly sampled initial points from standard normal distribution. In contrast, the training optimizees are still LASSO with a mixture of uniform distribution from which the coefficient matrices are sampled. The experiments are repeated for times, with all the algorithms receiving identical adaptation and testing samples in each run. Figure A5a shows the curves of the logarithm of the objective values generated by different methods, where our proposed M-L2O outperforms other baselines significantly. At 500-th step, the (mean, standard deviation) of the logarithmic objective values for {Vanilla L2O, TL, DT, M-L2O} are , which provides numerical supports of the advantage of our methods.

#### New Evaluation: Interpolation

To obtain new optimize weights, we employ a linear interpolation strategy between two adapted optimizers. The first one is optimized on the optimizees that are similar to those used in training, and the second is optimized on the optimizees that are similar to those used in testing. We introduce a factor to control the interpolation between the two weights, denoted by and , respectively, and calculate the new weights as follows:

In Figure A5b, we present the mean values of the logarithmic loss, as well as the 95% confidence interval. The results of TL and M-L2O validate our claim that adapting to training-like optimizees tend to yield better performance than adapting to optimizees that more resemble the testing optimizees.



(a) Convergence speeds on Rosenbrock optimizees. We repeat the experiments for 10 times, and present the 95% confidence intervals are shown in the figure.

(b) Convergence speeds on LASSO optimizees, with different interpolation weights  $\alpha$ . Both the mean and the 95% confidence intervals are shown in the figure.

Figure A5: Visualization of additional experiment results.