

Understanding Deep Architectures with Reasoning Layer

Xinshi Chen

Georgia Institute of Technology
xinshi.chen@gatech.edu

Yufei Zhang

University of Oxford
yufei.zhang@maths.ox.ac.uk

Christoph Reisinger

University of Oxford
christoph.reisinger@maths.ox.ac.uk

Le Song

Georgia Institute of Technology
lsong@cc.gatech.edu

Abstract

Recently, there has been a surge of interest in combining deep learning models with reasoning in order to handle more sophisticated learning tasks. In many cases, a reasoning task can be solved by an iterative algorithm. This algorithm is often unrolled, and used as a specialized layer in the deep architecture, which can be trained end-to-end with other neural components. Although such hybrid deep architectures have led to many empirical successes, the theoretical foundation of such architectures, especially the interplay between algorithm layers and other neural layers, remains largely unexplored. In this paper, we take an initial step towards an understanding of such hybrid deep architectures by showing that properties of the algorithm layers, such as convergence, stability and sensitivity, are intimately related to the approximation and generalization abilities of the end-to-end model. Furthermore, our analysis matches closely our experimental observations under various conditions, suggesting that our theory can provide useful guidelines for designing deep architectures with reasoning layers.

1 Introduction

Many real world applications require perception and reasoning to work together to solve a problem. Perception refers to the ability to understand and represent inputs, while reasoning refers to the ability to follow prescribed steps and derive answers satisfying certain constraints. To tackle such sophisticated learning tasks, recently, there has been a surge of interests in combining deep perception models with reasoning modules.

Typically, a **reasoning module** is stacked on top of a **neural module**, and treated as an additional layer of the overall deep architecture; then all the parameters in the architecture are optimized end-to-end with loss gradients (Fig 1). Very often these reasoning modules can be implemented as unrolled *iterative algorithms*, which can solve more sophisticated tasks with carefully designed and interpretable operations. For instance, SATNet [1] integrated a satisfiability solver into its deep model as a reasoning module; E2Efold [2] used a constrained optimization algorithm on top of a neural energy network to predict and reason about RNA structures, while [3] used optimal transport algorithm as a reasoning module for learning to sort. Other algorithms such as ADMM [4, 5], Langevin dynamics [6], inductive logic programming [7], DP [8], k-means clustering [9], message passing [10, 11], power iterations [12] are

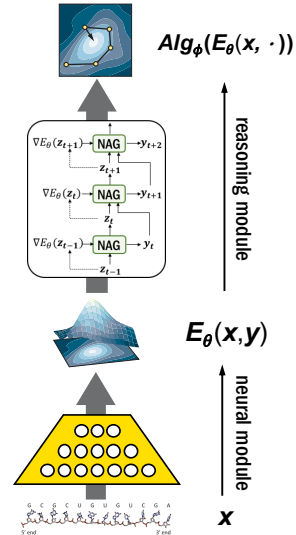


Figure 1: Hybrid architecture.

also used as differentiable reasoning modules in deep models for various learning tasks. Thus in the reminder of the paper, we will use reasoning layer and algorithm layer interchangeably.

While these previous works have demonstrated the effectiveness of combining deep learning with reasoning, the theoretical underpinning of such hybrid deep architectures remains largely unexplored. For instance, what is the benefit of using a reasoning module based on unrolled algorithms compared to generic architectures such as recurrent neural networks (RNN)? How exactly will the reasoning module affect the generalization ability of the deep architecture? For different algorithms which can solve the same task, what are their differences when used as reasoning modules in deep models? Despite the rich literature on rigorous analysis of algorithm properties, there is a paucity of work leveraging these analyses to formally study the learning behavior of deep architectures containing algorithm layers. This motivates us to ask the crucial and timely question of

How will the algorithm properties of a reasoning layer affect the learning behavior of deep architectures containing such layers?

In this paper, we provide a first step towards an answer to this question by analyzing the approximation and generalization abilities of such hybrid deep architectures. To the best of our knowledge, such an analysis has not been done before and faces several difficulties: 1) The analysis of certain algorithm properties such as convergence can be complex by itself; 2) Models based on highly structured iterative algorithms have rarely been analyzed before; 3) The bound needs to be sharp enough to match empirical observations. In this new setting, the complexities of the algorithm’s analysis and generalization analysis are intertwined together, making the analysis even more challenging.

Summary of results. We find that standard Rademacher complexity analysis, widely used for neural networks [13, 14, 15], is insufficient for explaining the behavior of these hybrid architectures. Thus we resort to a more refined local Rademacher complexity analysis [16, 17], and find the following:

- **Relation to algorithm properties.** Algorithm properties such as convergence, stability and sensitivity all play important roles in the generalization ability of the hybrid architecture. Generally speaking, an algorithm layer that is faster converging, more stable and less sensitive will be able to better approximate the joint perception and reasoning task, while at the same time generalize better.
- **Which algorithm?** There is a tradeoff that a faster converging algorithm has to be less stable [18]. Therefore, depending on the precise setting, the best choice of algorithm layer may be different. Our theorem reveals that when the neural module is over- or under-parameterized, stability of the algorithm layer can be more important than its convergence; but when the neural module is has an ‘about-right’ parameterization, a faster converging algorithm layer may give a better generalization.
- **What depth?** With deeper algorithm layers, the representation ability gets better, but the generalization becomes worse if the neural module is over/under-parameterized. Only when it has ‘about-right’ complexity, deeper algorithm layers can induce both better representation and generalization.
- **What if RNN?** It has been shown that RNN (or graph neural networks, GNN) can represent reasoning and iterative algorithms [19, 15]. On the example of RNN we demonstrate in Sec 6 that these generic reasoning modules can also be analyzed under our framework, revealing that RNN layers induce better representation but worse generalization compared to traditional algorithm layers.
- **Experiments.** We conduct empirical studies to validate our theory and show that it matches well with experimental observations under various conditions. These results suggest that our theory can provide useful practical guidelines for designing deep architectures with reasoning layers.

Contributions and limitations. To the best of our knowledge, this is the first result to quantitatively characterize the effects of algorithm properties on the learning behavior of hybrid deep architectures with reasoning layers, showing that algorithm biases can help reduce sample complexity of such architectures. Our result also reveals a subtle and previously unknown interplay between algorithm convergence, stability and sensitivity when affecting model generalization, and thus provides design principles for deep architectures with reasoning layers. To simplify the analysis, our initial study is limited to a setting where the reasoning module is an unconstrained optimization algorithm and the neural module outputs a quadratic energy function. However, our analysis framework can be extended to more complicated cases and the insights can be expected to apply beyond our current setting.

Related theoretical works. Our analysis borrows proof techniques for analyzing algorithm properties from the optimization literature [18, 20] and for bounding Rademacher complexity from the statistical learning literature [13, 16, 17, 21, 22], but our focus and results are new. More precisely, the ‘leave-

one-out' stability of optimization algorithms have been used to derive generalization bounds [23, 24, 25, 18, 26, 27]. However, all existing analyses are in the context where the optimization algorithms are used to train and select the model, while our analysis is based on a fundamentally different viewpoint where the algorithm itself is unrolled and integrated as a layer in the deep model. Also, existing works on the generalization of deep learning mainly focus on generic neural architectures such as feed-forward neural networks, RNN, GNN, etc [13, 14, 15]. The complexity of models based on highly structured iterative algorithms and the relation to algorithm properties have not been investigated. Furthermore, we are not aware of any previous use of local Rademacher complexity analysis for deep learning models.

2 Setting: Optimization Algorithms as Reasoning Modules

In many applications, reasoning can be accomplished by solving an optimization problem defined by a neural perceptual module. For instance, a visual SUDOKU puzzle can be solved using a neural module to perceive the digits followed by a quadratic optimization module to maximize a logic satisfiability objective [1]. The RNA folding problem can be tackled by a neural energy model to capture pairwise relations between RNA bases and a constrained optimization module to minimize the energy, with additional pairing constraints, to obtain a folding [2]. In a broader context, MAML [28, 29] also has a neural module for joint initialization and a reasoning module that performs optimization steps for task-specific adaptation. Other examples include [6, 30, 31, 32, 33, 34, 35, 36, 37, 38, 39]. More specifically, perception and reasoning can be jointly formulated in the form

$$\mathbf{y}(\mathbf{x}) = \arg \min_{\mathbf{y} \in \mathcal{Y}} E_{\theta}(\mathbf{x}, \mathbf{y}), \quad (1)$$

where $\mathbf{x}$ is an input, and $E_{\theta}(\mathbf{x}, \mathbf{y})$ is a neural energy function with parameters θ , which specifies the type of information needed for performing reasoning, and together with constraints $\mathcal{Y}$ on the output $\mathbf{y}$, specifies the style of reasoning. Very often, the optimizer can be approximated by iterative algorithms, so the mapping in Eq. 1 can be approximated by the following end-to-end hybrid model

$$f_{\phi, \theta}(\mathbf{x}) := \text{Alg}_{\phi}^k(E_{\theta}(\mathbf{x}, \cdot)) : \mathcal{X} \mapsto \mathcal{Y}. \quad (2)$$

Alg_{ϕ}^k is the reasoning module with parameters ϕ . Given a neural energy, it performs k -step iterative updates to produce the output (Fig 1). When k is finite, Alg_{ϕ}^k corresponds to approximate reasoning. As an initial attempt to analyze deep architectures with reasoning layers, we will restrict our analysis to a simple case where $E_{\theta}(\mathbf{x}, \mathbf{y})$ is quadratic in $\mathbf{y}$. A reason is that the analysis of advanced algorithms such as Nesterov accelerated gradients will become very complex for general cases. Similar problems occur in [18] which also restricts the proof to quadratic objectives. Specifically:

Problem setting: Consider a hybrid architecture where the neural module is an energy function of the form $E_{\theta}((\mathbf{x}, \mathbf{b}), \mathbf{y}) = \frac{1}{2} \mathbf{y}^{\top} Q_{\theta}(\mathbf{x}) \mathbf{y} + \mathbf{b}^{\top} \mathbf{y}$, with Q_{θ} a neural network that maps $\mathbf{x}$ to a matrix. Each energy can be uniquely represented by $(Q_{\theta}(\mathbf{x}), \mathbf{b})$, so we can write the overall architecture as

$$f_{\phi, \theta}(\mathbf{x}, \mathbf{b}) := \text{Alg}_{\phi}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}). \quad (3)$$

Assume we are given a set of n i.i.d. samples $S_n = \{((\mathbf{x}_1, \mathbf{b}_1), \mathbf{y}_1^*), \dots, ((\mathbf{x}_n, \mathbf{b}_n), \mathbf{y}_n^*)\}$, where the labels $\mathbf{y}^*$ are given by the *exact minimizer* Opt of the corresponding Q^* , i.e.,

$$\mathbf{y}^* = \text{Opt}(Q^*(\mathbf{x}), \mathbf{b}). \quad (4)$$

Then the learning problem is to find the best model $f_{\phi, \theta}$ from the space $\mathcal{F} := \{f_{\phi, \theta} : (\phi, \theta) \in \Phi \times \Theta\}$ by minimizing the empirical loss function

$$\min_{f_{\phi, \theta} \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \ell_{\phi, \theta}(\mathbf{x}_i, \mathbf{b}_i), \quad (5)$$

where $\ell_{\phi, \theta}(\mathbf{x}, \mathbf{b}) := \|\text{Alg}_{\phi}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2$. Furthermore, we assume:

- We have $\mathcal{Y} = \mathbb{R}^d$, and both Q_{θ} and Q^* map $\mathcal{X}$ to $\mathcal{S}_{\mu, L}^{d \times d}$, the space of symmetric positive definite (SPD) matrices with $\mu, L > 0$ as its smallest and largest singular values, respectively. Thus the induced energy function E_{θ} will be μ -strongly convex and L -smooth, and the output of Opt is unique.
- The input $(\mathbf{x}, \mathbf{b})$ is a pair of random variables where $\mathbf{x} \in \mathcal{X} \subseteq \mathbb{R}^m$ and $\mathbf{b} \in \mathcal{B} \subseteq \mathbb{R}^d$. Assume $\mathbf{b}$ satisfies $\mathbb{E}[\mathbf{b}\mathbf{b}^{\top}] = \sigma_b^2 I$. Assume $\mathbf{x}$ and $\mathbf{b}$ are independent, and their joint distribution follows a probability measure P . Assume samples in S_n are drawn i.i.d. from P .

- Assume $\mathcal{B}$ is bounded, and let $M = \sup_{(Q, \mathbf{b}) \in \mathcal{S}_{\mu, L}^{d \times d} \times \mathcal{B}} \|\text{Opt}(Q, \mathbf{b})\|_2$.

Though this setting does not encompass the full complexity of hybrid deep architectures, it already reveals interesting connections between algorithm properties of the reasoning module and the learning behaviors of hybrid architectures.

3 Properties of Algorithms

In this section, we formally define the algorithm properties of the reasoning module Alg_{ϕ}^k , under the problem setting presented in Sec 2. After that, we compare the corresponding properties of gradient descent, GD_{ϕ}^k , and Nesterov’s accelerated gradients, NAG_{ϕ}^k , as concrete examples.

(I) The convergence rate of an algorithm expresses how fast the optimization error decreases as k grows. Formally, we say Alg_{ϕ}^k has a convergence rate $\text{Cvg}(k, \phi)$ if for any $Q \in \mathcal{S}_{\mu, L}^{d \times d}$, $\mathbf{b} \in \mathcal{B}$,

$$\|\text{Alg}_{\phi}^k(Q, \mathbf{b}) - \text{Opt}(Q, \mathbf{b})\|_2 \leq \text{Cvg}(k, \phi) \|\text{Alg}_{\phi}^0(Q, \mathbf{b}) - \text{Opt}(Q, \mathbf{b})\|_2. \quad (6)$$

(II) Stability of an algorithm characterizes its robustness to small *perturbations in the optimization objective*, which corresponds to the perturbation of Q and $\mathbf{b}$ in the quadratic case. For the purpose of this paper, we say an algorithm Alg_{ϕ}^k is $\text{Stab}(k, \phi)$ -stable if for any $Q, Q' \in \mathcal{S}_{\mu, L}^{d \times d}$ and $\mathbf{b}, \mathbf{b}' \in \mathcal{B}$,

$$\|\text{Alg}_{\phi}^k(Q, \mathbf{b}) - \text{Alg}_{\phi}^k(Q', \mathbf{b}')\|_2 \leq \text{Stab}(k, \phi) \|Q - Q'\|_2 + \text{Stab}(k, \phi) \|\mathbf{b} - \mathbf{b}'\|_2, \quad (7)$$

where $\|Q - Q'\|_2$ is the spectral norm of the matrix $Q - Q'$.

(III) Sensitivity characterizes the robustness to small *perturbations in the algorithm parameters* ϕ . We say the sensitivity of Alg_{ϕ}^k is $\text{Sens}(k)$ if it holds for all $Q \in \mathcal{S}_{\mu, L}^{d \times d}$, $\mathbf{b} \in \mathcal{B}$, and $\phi, \phi' \in \Phi$ that

$$\|\text{Alg}_{\phi}^k(Q, \mathbf{b}) - \text{Alg}_{\phi'}^k(Q, \mathbf{b})\|_2 \leq \text{Sens}(k) \|\phi - \phi'\|_2. \quad (8)$$

This concept is referred in the deep learning community to “parameter perturbation error” or “sharpness” [40, 41, 42]. It has been used for deriving generalization bounds of neural networks, both in the Rademacher complexity framework [13] and PAC-Bayes framework [43].

(IV) The stable region is the range Φ of the parameters ϕ where the algorithm output will remain bounded as k grows to infinity, i.e., numerically stable. Only when the algorithms operate in the stable region, the corresponding $\text{Cvg}(k, \phi)$, $\text{Stab}(k, \phi)$ and $\text{Sens}(k)$ will remain finite for all k . It is usually very difficult to identity the exact stable region, but a sufficient range can be provided.

GD and NAG. Now we will compare the above four algorithm properties for gradient descent and Nesterov’s accelerated gradient method, both of which can be used to solve the quadratic optimization in our problem setting. First, the algorithm update steps are summarized below:

$$\text{GD}_{\phi} : \mathbf{y}_{k+1} \leftarrow \mathbf{y}_k - \phi(Q\mathbf{y}_k + \mathbf{b}) \quad \text{NAG}_{\phi} : \begin{cases} \mathbf{y}_{k+1} \leftarrow \mathbf{z}_k - \phi(Q\mathbf{z}_k + \mathbf{b}) \\ \mathbf{z}_{k+1} \leftarrow \mathbf{y}_{k+1} + \frac{1-\sqrt{\mu\phi}}{1+\sqrt{\mu\phi}}(\mathbf{y}_{k+1} - \mathbf{y}_k) \end{cases} \quad (9)$$

where the hyperparameter ϕ corresponds to the step size. The initializations $\mathbf{y}_0, \mathbf{z}_0$ are set to zero vectors throughout this paper. Denote the results of k -step update, $\mathbf{y}_k$, of GD and NAG by $\text{GD}_{\phi}^k(Q, \mathbf{b})$ and $\text{NAG}_{\phi}^k(Q, \mathbf{b})$, respectively. Then their algorithm properties are summarized in Table 1.

Table 1: Comparison of algorithm properties between GD and NAG. For simplicity, only the order in k is presented. Complete statements with detailed coefficients and proofs are given in Appendix A.

Alg	$\text{Cvg}(k, \phi)$	$\text{Stab}(k, \phi)$	$\text{Sens}(k)$	Stable region Φ
GD_{ϕ}^k	$\mathcal{O}((1 - \phi\mu)^k)$	$\mathcal{O}(1 - (1 - \phi\mu)^k)$	$\mathcal{O}(k(1 - c_0\mu)^{k-1})$	$[c_0, \frac{2}{\mu+L}]$
NAG_{ϕ}^k	$\mathcal{O}(k(1 - \sqrt{\phi\mu})^k)$	$\mathcal{O}(1 - (1 - \sqrt{\phi\mu})^k)$	$\mathcal{O}(k^3(1 - \sqrt{c_0\mu})^k)$	$[c_0, \frac{4}{\mu+3L}]$

Table 1 shows: (i) *Convergence*: NAG converges faster than GD, especially when μ is very small, which is a well-known result. (ii) *Stability*: However, as k grows, NAG is less stable than GD for a fixed k , in contrast to their convergence behaviors. This is pointed out in [18], which proves that a faster converging algorithm has to be less stable. (iii) *Sensitivity*: The sensitivity behaves similar to

the convergence, where NAG is less sensitive to step-size perturbation than GD. Also, the sensitivity of both algorithms gets smaller as k grows larger. (iv): *Stable region*: Since $\mu < L$, the stable region of GD is larger than that of NAG. It means a larger step size is allowable for GD that will not lead to exploding outputs even if k is large. Note that all the other algorithm properties are based on the assumption that ϕ is in the stable region Φ . Furthermore, as k goes to infinity, the space $\{\text{Alg}_\phi^k : \phi \in \Phi\}$ will finally shrink to a single function, which is the exact minimizer $\{\text{Opt}\}$.

Our purpose of comparing the algorithm properties of GD and NAG is to show in a later section their difference as a reasoning layer in deep architectures. However, some results in Table 1 are new by themselves, which may be of independent interest. For instance, we are not aware of other analysis of the sensitivity of GD and NAG to their step-size perturbation. Besides, for the stability results, we provide a proof with a weaker assumption where ϕ can be larger than $1/L$, which is not allowed in [18]. This is necessary since in practice the learned step size ϕ is usually larger than $1/L$.

4 Approximation Ability

How will the algorithm properties affect the approximation ability of deep architecture with reasoning layers? Given a model space $\mathcal{F} := \{\text{Alg}_\phi^k(Q_\theta(\cdot), \cdot) : \phi \in \Phi, \theta \in \Theta\}$, we are interested in its approximation ability to functions of the form $\text{Opt}(Q^*(x), b)$. More specifically, we define the loss

$$\ell_{\phi, \theta}(x, b) := \|\text{Alg}_\phi^k(Q_\theta(x), b) - \text{Opt}(Q^*(x), b)\|_2, \quad (10)$$

and measure the approximation ability by $\inf_{\phi \in \Phi, \theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\ell_{\phi, \theta}$, where $\mathcal{Q}^* := \{\mathcal{X} \mapsto \mathcal{S}_{\mu, L}^{d \times d}\}$ and $P\ell_{\phi, \theta} = \mathbb{E}_{x, b}[\ell_{\phi, \theta}(x, b)]$. Intuitively, using a faster converging algorithm, the model Alg_ϕ^k could represent the reasoning-task structure, Opt , better and improve the overall approximation ability. Indeed we can prove the following lemma confirming this intuition.

Lemma 4.1. (Faster Convergence $\Rightarrow$ Better Approximation Ability). *Assume the problem setting in Sec 2. The approximation ability can be bounded by two terms:*

$$\inf_{\phi \in \Phi, \theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\ell_{\phi, \theta} \leq \sigma_b \mu^{-2} \underbrace{\inf_{\theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\|Q_\theta - Q^*\|_F}_{\text{approximation ability of the neural module}} + M \underbrace{\inf_{\phi \in \Phi} \text{Cvg}(k, \phi)}_{\text{best convergence}}. \quad (11)$$

With Lemma 4.1, we conclude that: A faster converging algorithm can define a model with better approximation ability. For example, for a fixed k and Q_θ , NAG converges faster than GD, so NAG_ϕ^k can approximate Opt more accurately than GD_ϕ^k , which is experimentally validated in Sec 7.

Similarly, we can also reverse the reasoning, and ask the question that, given two hybrid architectures with the same approximation error, which architecture has a smaller error in representing the energy function Q^* ? We show that this error is also intimately related to the convergence of the algorithm.

Lemma 4.2. (Faster Convergence $\Rightarrow$ Better Representation of Q^*). *Assume the problem setting in Sec 2. Then $\forall \phi \in \Phi, \theta \in \Theta, Q^* \in \mathcal{Q}^*$ it holds true that*

$$P\|Q_\theta - Q^*\|_F^2 \leq \sigma_b^{-2} L^4 (\sqrt{P\ell_{\phi, \theta}^2} + M \cdot \text{Cvg}(k, \phi))^2. \quad (12)$$

Lemma 4.2 highlights the benefit of using an algorithmic layer that aligns with the reasoning-task structure. Here the task structure is represented by Opt , the minimizer, and convergence measures how well Alg_ϕ^k is aligned with Opt . Lemma 4.2 essentially indicates that *if the structure of a reasoning module can better align with the task structure, then it can better constrain the search space of the underlying neural module Q_θ , making it easier to learn, and further lead to better sample complexity, which we will explain more in the next section.*

As a concrete example for Lemma 4.2, if $\text{GD}_\phi^k(Q_\theta, \cdot)$ and $\text{NAG}_\phi^k(Q_\theta, \cdot)$ achieve the **same** accuracy for approximating $\text{Opt}(Q^*, \cdot)$, then the neural module Q_θ in $\text{NAG}_\phi^k(Q_\theta, \cdot)$ will have a **better** accuracy for approximating Q^* than Q_θ in $\text{GD}_\phi^k(Q_\theta, \cdot)$. In other words, a faster converging algorithm imposes more constraints on the energy function Q_θ , making it approach Q^* faster.

5 Generalization Ability

How will algorithm properties affect the generalization ability of deep architectures with reasoning layers? We theoretically showed that the generalization bound is determined by both the algorithm properties and the complexity of the neural module. Moreover, it induces interesting implications - when the neural module is over- or under- parameterized, the generalization bound is dominated by algorithm stability; but when the neural module has an about-right parameterization, the bound is dominated by the product of algorithm stability and convergence.

More specifically, we will analyze *generalization gap* between the expected loss and empirical loss,

$$P\ell_{\phi,\theta} = \mathbb{E}_{\mathbf{x},\mathbf{b}}\ell_{\phi,\theta}(\mathbf{x},\mathbf{b}) \text{ and } P_n\ell_{\phi,\theta} = \frac{1}{n} \sum_{i=1}^n \ell_{\phi,\theta}(\mathbf{x}_i, \mathbf{b}_i), \text{ respectively,} \quad (13)$$

where P_n is the empirical probability measure induced by the samples S_n . Let $\ell_{\mathcal{F}} := \{\ell_{\phi,\theta} : \phi \in \Phi, \theta \in \Theta\}$ be the function space of losses of the models. The generalization gap, $P\ell_{\phi,\theta} - P_n\ell_{\phi,\theta}$, can be bounded by the Rademacher complexity, $\mathbb{E}R_n\ell_{\mathcal{F}}$, which is defined as the expectation of the empirical Rademacher complexity, $R_n\ell_{\mathcal{F}} := \mathbb{E}_{\sigma} \sup_{\phi \in \Phi, \theta \in \Theta} \frac{1}{n} \sum_{i=1}^n \sigma_i \ell_{\phi,\theta}(\mathbf{x}_i, \mathbf{b}_i)$, where $\{\sigma_i\}_{i=1}^n$ are n independent Rademacher random variables uniformly distributed over $\{\pm 1\}$. Generalization bounds derived from Rademacher complexity have been studied in many works [44, 45, 46, 47].

However, deriving the Rademacher complexity of $\ell_{\mathcal{F}}$ is highly nontrivial in our case, and we are not aware of prior bounds for deep learning models with reasoning layers. Aiming at bridging the relation between algorithm properties and generalization ability that can explain experimental observations, we find that standard Rademacher complexity analysis is insufficient. The shortcoming of the standard Rademacher complexity is that it provides *global* estimates of the complexity of the model space, which ignores the fact that the training process will likely pick models with small errors. Taking this factor into account, we resort to more refined analysis using *local* Rademacher complexity [13, 16, 17]. Remarkably, we found that the bounds derived via *global* and *local* Rademacher complexity will lead to *different* conclusions about the effects of algorithm layers. That is, an algorithm that converges faster could lead to a model space that has a larger global Rademacher complexity but a smaller local Rademacher complexity. Also, the *global* Rademacher complexity is dominated by algorithmic stability. However, in the *local* counterpart, there is a trade-off term between stability and convergence, which aligns better with the experimental observations.

Main Result. More specifically, the local Rademacher complexity of $\ell_{\mathcal{F}}$ at level r is defined as

$$\mathbb{E}R_n\ell_{\mathcal{F}}^{loc}(r) \text{ where } \ell_{\mathcal{F}}^{loc}(r) := \{\ell_{\phi,\theta} : \phi \in \Phi, \theta \in \Theta, P\ell_{\phi,\theta}^2 \leq r\}. \quad (14)$$

This notion is less general than the one defined in [16, 17] but is sufficient for our purpose. Here we also define a loss function space $\ell_{\mathcal{Q}} := \{\|Q_{\theta} - Q^*\|_F : \theta \in \Theta\}$ for the neural module Q_{θ} , and introduce its local Rademacher complexity $\mathbb{E}R_n\ell_{\mathcal{Q}}^{loc}(r_q)$, where $\ell_{\mathcal{Q}}^{loc}(r_q) = \{\|Q_{\theta} - Q^*\|_F \in \ell_{\mathcal{Q}} : P\|Q_{\theta} - Q^*\|_F^2 \leq r_q\}$. With these definitions, we can show that the local Rademacher complexity of the hybrid architecture is explicitly related to all considered algorithm properties, namely convergence, stability and sensitivity, and there is an intricate trade-off.

Theorem 5.1. Assume the problem setting in Sec 2. Then we have for any $t > 0$ that

$$\mathbb{E}R_n\ell_{\mathcal{F}}^{loc}(r) \leq \sqrt{2}dn^{-\frac{1}{2}} \text{Stab}(k) \left(\sqrt{(\text{Cvg}(k)M + \sqrt{r})^2 C_1(n) + C_2(n, t) + C_3(n, t) + 4} \right) \quad (15)$$

$$+ \text{Sens}(k)B_{\Phi}, \quad (16)$$

where $\text{Stab}(k) = \sup_{\phi} \text{Stab}(k, \phi)$ and $\text{Cvg}(k) = \sup_{\phi} \text{Cvg}(k, \phi)$ are worst-case stability and convergence, $B_{\Phi} = \frac{1}{2} \sup_{\phi, \phi' \in \Phi} \|\phi - \phi'\|_2$, $C_1(n) = \mathcal{O}(\log N(n))$, $C_3(n, t) = \mathcal{O}(\frac{\log N(n)}{\sqrt{n}} + \frac{\sqrt{\log N(n)}}{e^t})$, $C_2(n, t) = \mathcal{O}(\frac{t \log N(n)}{n} + (C_3(n, t) + 1) \frac{\log N(n)}{\sqrt{n}})$, and $N(n) = \mathcal{N}(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}, L_{\infty})$ is the covering number of $\ell_{\mathcal{Q}}$ with radius $\frac{1}{\sqrt{n}}$ and L_{∞} norm.

Proof Sketch. We will explain the key steps here, and the full proof details are deferred to Appendix C. The essence of the proof is to find the relation between $R_n\ell_{\mathcal{F}}^{loc}(r)$ and $R_n\ell_{\mathcal{Q}}^{loc}(r_q)$, and also the relation between the local level r and r_q . Then the analysis of the local Rademacher complexity of the end-to-end model $\text{Alg}_{\phi}^k(Q_{\theta}, \cdot)$ can be reduced to that of the neural module Q_{θ} .

More specifically, we first show that the loss $\ell_{\phi, \theta}$ is $Stab(k)$ -Lipschitz in Q_θ and $Sens(k)$ -Lipschitz in ϕ . By the triangle inequality and algorithm properties, we can bound the sensitivity of the loss by

$$|\ell_{\phi, \theta}(\mathbf{x}) - \ell_{\phi', \theta'}(\mathbf{x})| \leq Stab(k) \|Q_\theta(\mathbf{x}) - Q_{\theta'}(\mathbf{x})\|_2 + Sens(k) \|\phi - \phi'\|_2. \quad (17)$$

Second, by leveraging vector-contraction inequality for Rademacher complexity of vector-valued hypothesis [21, 22] and our previous observations in Lemma 4.2, we can turn the sensitivity bound on the loss function in Eq. 17 to a local Rademacher complexity bound

$$R_n \ell_{\mathcal{F}}^{loc}(r) \leq \sqrt{2d} \textcolor{red}{Stab}(k) R_n \ell_{\mathcal{Q}}^{loc}(r_q) + \textcolor{blue}{Sens}(k) B_\Phi \text{ with } r_q = \sigma_b^{-2} L^4 (\sqrt{r} + M \textcolor{blue}{Cvg}(k))^2. \quad (18)$$

Therefore, bounding the local Rademacher complexity of $\ell_{\mathcal{F}}^{loc}$ at level r resorts to bounding that of $\ell_{\mathcal{Q}}^{loc}$ at level r_q . This inequality has already revealed the role of stability, convergence, and sensitivity in bounding local Rademacher complexity, and is the key step in the proof.

Third, based on an extension of Talagrand’s inequality for empirical processes [48, 16], we can bound the empirical error $P_n \|Q_\theta - Q^*\|_F^2$ using r_q and some other terms with high probability. Then $R_n \ell_{\mathcal{Q}}^{loc}(r_q)$ can be bounded using the covering number of $\ell_{\mathcal{Q}}$ via the classical Dudley entropy integral [49], where the upper integration bound is given by the upper bound of $P_n \|Q_\theta - Q^*\|_F^2$. $\square$

Trade-offs between convergence, stability and sensitivity. Generally speaking, the convergence rate $\textcolor{blue}{Cvg}(k)$ and sensitivity $\textcolor{blue}{Sens}(k)$ have similar behavior, but $\textcolor{red}{Stab}(k)$ behaves opposite to them; see illustrations in Fig 2. Therefore, the way these three quantities interact in Theorem 5.1 suggests that in different regimes one may see different generalization behavior. More specially, depending on the parameterization of Q_θ , the coefficients C_1 , C_2 , and C_3 in Eq. 15 may have different scale, making the local Rademacher complexity bound dominated by different algorithm properties. Since the coefficients C_i are monotonely increasing in the covering number of $\ell_{\mathcal{Q}}$, we expect that:

- (i) When Q_θ is **over-parameterized**, the covering number of $\ell_{\mathcal{Q}}$ becomes large, as do the three coefficients. Large C_i will reduce the effect of $\textcolor{blue}{Cvg}(k)$ and make Eq. 15 dominated by $\textcolor{red}{Stab}(k)$;
- (ii) When Q_θ is **under-parameterized**, the three coefficients get small, but they still reduce the effect of $\textcolor{blue}{Cvg}(k)$ given the constant 4 in Eq. 15, again making it dominated by $\textcolor{red}{Stab}(k)$;
- (iii) When the parametrization of Q_θ is **about-right**, we can expect $\textcolor{blue}{Cvg}(k)$ to play a critical role in Eq. 15, which will then behave similar to the product $\textcolor{red}{Stab}(k) \textcolor{blue}{Cvg}(k)$, as illustrated schematically in Fig 2. We experimentally validate these implications in Sec 7.

Trade-off of the depth. Combining the above implications with the approximation ability analysis in Sec 4, we can see that in the above-mentioned cases (i) and (ii), deeper algorithm layers will lead to better approximation accuracy but worse generalization. Only in the ideal case (iii), a deeper reasoning module can induce both better representation and generalization abilities. This result provides practical guidelines for some recently proposed infinite-depth models [50, 51].

Standard Rademacher complexity analysis. If we consider the standard Rademacher complexity and directly bound it by the covering number of $\ell_{\mathcal{F}}$ via Dudley’s entropy integral in the way some existing generalization bounds of deep learning are derived [13, 14, 15], we will get the following upper bound for the covering number, where $\textcolor{blue}{Cvg}(k)$ does not play a role:

$$\mathcal{N}(\epsilon, \ell_{\mathcal{F}}, L_2(P_n)) \leq \mathcal{N}(\epsilon / (2\textcolor{red}{Stab}(k)), \mathcal{Q}, L_2(P_n)) \cdot \mathcal{N}(\epsilon / (2\textcolor{blue}{Sens}(k)), \Phi, \|\cdot\|_2). \quad (19)$$

Since Φ only contains the hyperparameters in the algorithm and $\mathcal{Q} := \{Q_\theta, \theta \in \Theta\}$ is often highly expressive, typically stability will dominate this bound. Or, consider the case when algorithm layers are fixed so Φ only contains one element. Then this covering number is determined by stability, which infers that $\text{NAG}_1^k(Q_\theta, \cdot)$ has a **larger** Rademacher complexity than $\text{GD}_1^k(Q_\theta, \cdot)$ since it is less stable. However, in the local Rademacher complexity bound in Theorem 5.1, even if $\textcolor{blue}{Sens}(k)$ in Eq. 16 is ignored, there is still a trade-off between convergence and stability which implies $\text{NAG}_1^k(Q_\theta, \cdot)$ can have a **smaller** local Rademacher complexity than $\text{GD}_1^k(Q_\theta, \cdot)$, leading to a different conclusion. Our experiments show the local Rademacher complexity bound is better for explaining the actual observations.

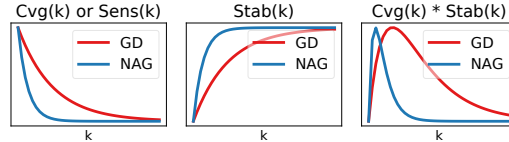


Figure 2: Overall trend of algorithm properties.

6 Pros and Cons of RNN as a Reasoning Layer

It has been shown that RNN (or GNN) can represent reasoning and iterative algorithms over structures [19, 15]. Can our analysis framework also be used to understand RNN (or GNN)? How will its behavior compare with more interpretable algorithm layers such as GD_ϕ^k and NAG_ϕ^k ? In the case of RNN, the algorithm update steps in each iteration are given by an RNN cell

$$\mathbf{y}_{k+1} \leftarrow \text{RNNcell}(Q, \mathbf{b}, \mathbf{y}_k) := V\sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_t + W_2^1\mathbf{g}_t))). \quad (20)$$

where the activation function $\sigma = \text{ReLU}$ takes $\mathbf{y}_k$ and the gradient $\mathbf{g}_t = Q\mathbf{y}_t + \mathbf{b}$ as inputs. Then a recurrent neural network RNN_ϕ^k having k unrolled RNN cells can be viewed as a neural algorithm. The algorithm properties of RNN_ϕ^k are summarized in Table 2. Assume $\phi = \{V, W_1^1, W_2^1, W^{2:L}\}$ is in a stable region with $c_\phi := \sup_Q \|V\|_2 \|W_1^1 + W_2^1 Q\|_2 \prod_{l=2}^L \|W^l\|_2 < 1$, so that the operations in RNNcell are strictly contractive, i.e., $\|\mathbf{y}_{k+1} - \mathbf{y}_k\|_2 < \|\mathbf{y}_k - \mathbf{y}_{k-1}\|_2$. In this case, the stability and sensitivity of RNN_ϕ^k are guaranteed to be bounded.

However, the fundamental disadvantage of RNN is its lack of worst-case guarantee for convergence. In general the outputs of RNN_ϕ^k may not converge to the minimizer Opt , meaning that its worst-case convergence rate can be much larger than 1. This will lead to worse generalization bound according to our theory compared to GD_ϕ^k and NAG_ϕ^k .

Table 2: Properties of RNN_ϕ^k . (Details are given in Appendix D.)

Stable region Φ	$c_\phi < 1$
$\text{Stab}(k, \phi)$	$\mathcal{O}(1 - c_\phi^k)$
$\text{Sens}(k)$	$\mathcal{O}(1 - (\inf_\phi c_\phi)^k)$
$\min_\phi \text{Cvg}(k, \phi)$	$\mathcal{O}(\rho^k)$ with $\rho < 1$

The advantage of RNN is its expressiveness, especially given the universal approximation ability of MLP in the RNNcell . One can show that RNN_ϕ^k can express GD_ϕ^k or NAG_ϕ^k with suitable choices of ϕ . Therefore, its best-case convergence can be as small as $\mathcal{O}(\rho^k)$ for some $\rho < 1$. When the needed types of reasoning is unknown or beyond what existing algorithms are capable of, RNN has the potential to learn new reasoning types given sufficient data.

7 Experimental Validation

Our experiments aim to validate our theoretical prediction with computational simulations, rather than obtaining state-of-the-art results. We hope the theory together with these experiments can lead to practical guidelines for designing deep architectures with reasoning layers. We conduct two sets of experiments, where the first set of experiments strictly follows the problem setting described in Sec 2 and the second is conducted on BSD500 dataset [52] to demonstrate the possibility of generalizing the theorem to more realistic applications. Implementations in Python are released¹.

7.1 Synthetic Experiments

The experiments follow the problem setting in Sec 2. We sample 10000 pairs of $(\mathbf{x}, \mathbf{b})$ uniformly as overall dataset. During training, n samples are randomly drawn from these 10000 data points as the training set. Each $Q^*(\mathbf{x})$ is produced by a rotation matrix and a vector of eigenvalues parameterized by a randomly fixed 2-layer dense neural network with hidden dimension 3. Then the labels are generated according to $\mathbf{y} = \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})$. We train the model $\text{Alg}_\phi^k(Q_\theta, \cdot)$ on S_n using the loss in Eq. 10. Here, Q_θ has the same overall architecture as Q^* but the hidden dimension could vary. Note that in all figures, each k corresponds to an **independently trained model** with k iterations in the algorithm layer, instead of the sequential outputs of a single model. Each model is trained by ADAM and SGD with learning rate grid-searched from $[1\text{e-}2, 5\text{e-}3, 1\text{e-}3, 5\text{e-}4, 1\text{e-}4]$, and only the best result is reported. Furthermore, error bars are produced by 20 independent instantiations of the experiments. See Appendix E for more details.

Approximation ability. To validate Lemma 4.1, we compare $\text{GD}_\phi^k(Q_\theta, \cdot)$ and $\text{NAG}_\phi^k(Q_\theta, \cdot)$ in terms of approximation accuracy. For various hidden sizes of Q_θ , the results are similar, so we report one representative case in Fig 3. The approximation accuracy aligns with the convergence of the algorithms, showing that faster converging algorithm can induce better approximation ability.

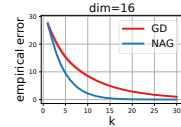


Figure 3

Faster convergence $\Rightarrow$ better Q_θ . We report the error of the neural module Q_θ in Fig 4. Note that $\text{Alg}_\phi^k(Q_\theta, \cdot)$ is trained end-to-end, without supervision on Q_θ . In Fig 4, the error of Q_θ decreases as

¹<https://github.com/xinshi-chen/Deep-Architecture-With-Reasoning-Layer>

k grows, in a rate similar to algorithm convergence. This validates the implication of Lemma 4.2 that, when Alg_ϕ^k is closer to Opt , it can help the underlying neural module Q_θ to get closer to Q^* .

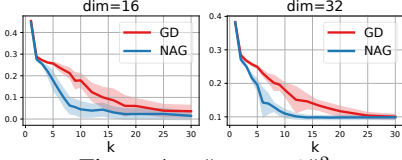


Figure 4: $P\|Q_\theta - Q^*\|_F^2$

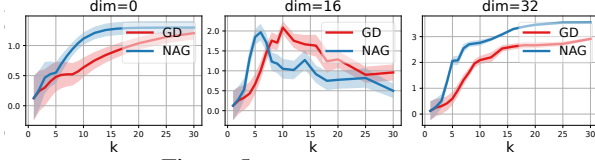


Figure 5: Generalization gap

Generalization gap. In Fig 5, we report the generalization gaps, with hidden sizes of Q_θ being 0, 16, and 32, which corresponds to the three cases (ii), (iii), and (i) discussed under Theorem 5.1, respectively. Comparing Fig 5 to Fig 2, we can see that the experimental results match very well with the theoretical implications.

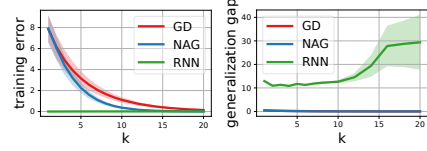


Figure 6: Algorithm layers vs RNN.

RNN. As discussed in Sec 6, RNN can be viewed as neural algorithms. To have a cleaner comparison, we report their behaviors under the ‘learning to optimize’ scenario where the objectives $(Q, \mathbf{b})$ are given. Fig 6 shows that RNN has a better representation power but worse generalization ability.

7.2 Experiments on Real Dataset

To show the real world applicability of our theoretical framework, we consider the **local adaptive image denoising** task. Details are given below.

Dataset. We split BSD500 (400 images) into a training set (100 images) and a test set (300 images). Gaussian noises are added to each pixel with noise levels depending on image local smoothness, making the noise levels on edges lower than non-edge regions. The task is to restore the original image from the noisy version $X \in [0, 1]^{180 \times 180}$.

Architecture. In $\text{Alg}_\phi^k(E_\theta(X, \cdot))$, Alg_ϕ^k is a k -step unrolled minimization algorithm to the ℓ_2 -regularized reconstruction objective $E_\theta(X, Y) := \frac{1}{2}\|Y + g_\theta(X) - X\|_F^2 + \frac{1}{2}\sum_{i,j} |[f_\theta(X)]_{i,j} Y_{i,j}|^2$, and the residual $g_\theta(X)$ and position-wise regularization coefficient $f_\theta(X)$ are both DnCNN networks as in [53]. The optimization objective, $E_\theta(X, Y)$, is quadratic in Y .

Generalization gap. We instantiate the hybrid architecture into different models using GD and NAG algorithms with different unrolled steps k . Each model is trained with 3000 epochs, and the *generalization gaps* are reported in Fig. 7. The results also show good consistency with our theory, where stabler algorithm (GD) can generalize better given *over/under*-parameterized neural module, and for the *about-right* parameterization case, the generalization gap behaviors are similar to $\text{Stab}(k) * \text{Cvg}(k)$.

Visualization. To show that the learned hybrid model has a good performance in this real application, we include a visualization of the original, noisy, and denoised images.

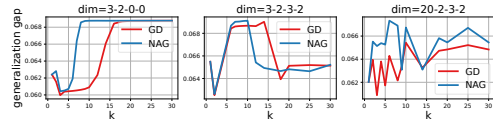
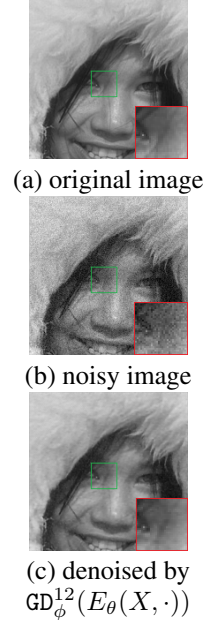


Figure 7: Generalization gap. Each k corresponds to a separately trained model. Left (under-parameterized): f_θ is a DnCNN with 3 channels and 2 hidden layers and $g_\theta = 0$. Middle (about-right): both f_θ and g_θ have 3 channels and 2 hidden layers. Right (over-parameterized): f_θ has 20 channels.

8 Conclusion and Discussion

In this paper, we take an initial step towards the theoretical understanding of deep architectures with reasoning layers. Our theorem indicates intriguing relation between algorithm properties of the reasoning module and the approximation and generalization of the end-to-end model, which in turn provides practical guideline for designing reasoning layers. The current analysis is limited due to the simplified problem setting. However, assumptions we made are only for avoiding the non-uniqueness of the reasoning solution and the instability of the mapping from the reasoning solution to the neural module. The assumptions could be relaxed if we can involve other techniques to resolve these issues. These additional efforts could potentially generalize the results to more complex cases.

Broader Impact

A common ethical concern of deep learning models is that they may not perform well on unseen examples, which could lead to the risk of producing biased content reflective of the training data. Our work, which learns an energy optimization model from the data, is not an exception. The approach we adopt to address this issue is to design hybrid deep architectures containing specialized reasoning modules. In the setting of quadratic energy functions, our theoretical analysis and numerical experiments show that hybrid deep models produce more reliable results than generic deep models on unseen data sets. More work is needed to determine the extent to which such hybrid model prevents biased outputs in more sophisticated tasks

Acknowledgement

We would like to thank Professor Vladimir Koltchinskii for providing valuable suggestions and thank anonymous reviewers for providing constructive feedbacks. This work is supported in part by NSF grants CDS&E-1900017 D3SC, CCF-1836936 FMitF, IIS-1841351, CAREER IIS-1350983 to L.S.

References

- [1] Po-Wei Wang, Priya Donti, Bryan Wilder, and Zico Kolter. Satnet: Bridging deep learning and logical reasoning using a differentiable satisfiability solver. In *International Conference on Machine Learning*, pages 6545–6554, 2019.
- [2] Xinshi Chen, Yu Li, Ramzan Umarov, Xin Gao, and Le Song. Rna secondary structure prediction by learning unrolled algorithms. *arXiv preprint arXiv:2002.05810*, 2020.
- [3] Marco Cuturi, Olivier Teboul, and Jean-Philippe Vert. Differentiable sorting using optimal transport: The sinkhorn cdf and quantile operator. *arXiv preprint arXiv:1905.11885*, 2019.
- [4] Harsh Shrivastava, Xinshi Chen, Binghong Chen, Guanghui Lan, Srinivas Aluru, Han Liu, and Le Song. GLAD: Learning sparse graph recovery. In *International Conference on Learning Representations*, 2020.
- [5] Y Yang, J Sun, H Li, and Z Xu. Admm-net: A deep learning approach for compressive sensing mri. corr. *arXiv preprint arXiv:1705.06869*, 2017.
- [6] John Ingraham, Adam Riesselman, Chris Sander, and Debora Marks. Learning protein structure with a differentiable simulator. In *International Conference on Learning Representations*, 2019.
- [7] Robin Manhaeve, Sebastijan Dumancic, Angelika Kimmig, Thomas Demeester, and Luc De Raedt. Deepproblog: Neural probabilistic logic programming. In *Advances in Neural Information Processing Systems*, pages 3749–3759, 2018.
- [8] Arthur Mensch and Mathieu Blondel. Differentiable dynamic programming for structured prediction and attention. In *35th International Conference on Machine Learning*, volume 80, 2018.
- [9] Bryan Wilder, Eric Ewing, Bistra Dilkina, and Milind Tambe. End to end learning and optimization on graphs. In *Advances in Neural Information Processing Systems*, pages 4674–4685, 2019.
- [10] Justin Domke. Parameter learning with truncated message-passing. In *CVPR 2011*, pages 2937–2943. IEEE, 2011.
- [11] Despoina Paschalidou, Osman Ulusoy, Carolin Schmitt, Luc Van Gool, and Andreas Geiger. Raynet: Learning volumetric 3d reconstruction with ray potentials. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3897–3906, 2018.
- [12] Wei Wang, Zheng Dang, Yinlin Hu, Pascal Fua, and Mathieu Salzmann. Backpropagation-friendly eigendecomposition. In *Advances in Neural Information Processing Systems*, pages 3156–3164, 2019.
- [13] Peter L Bartlett, Dylan J Foster, and Matus J Telgarsky. Spectrally-normalized margin bounds for neural networks. In *Advances in Neural Information Processing Systems*, pages 6240–6249, 2017.

- [14] Minshuo Chen, Xingguo Li, and Tuo Zhao. On generalization bounds of a family of recurrent neural networks. *arXiv preprint arXiv:1910.12947*, 2019.
- [15] Vikas K Garg, Stefanie Jegelka, and Tommi Jaakkola. Generalization and representational limits of graph neural networks. *arXiv preprint arXiv:2002.06157*, 2020.
- [16] Peter L Bartlett, Olivier Bousquet, Shahar Mendelson, et al. Local rademacher complexities. *The Annals of Statistics*, 33(4):1497–1537, 2005.
- [17] Vladimir Koltchinskii et al. Local rademacher complexities and oracle inequalities in risk minimization. *The Annals of Statistics*, 34(6):2593–2656, 2006.
- [18] Yuansi Chen, Chi Jin, and Bin Yu. Stability and convergence trade-off of iterative optimization algorithms. *arXiv preprint arXiv:1804.01619*, 2018.
- [19] 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*, pages 3981–3989, 2016.
- [20] Yurii Nesterov. *Introductory lectures on convex optimization: A basic course*, volume 87. Springer Science & Business Media, 2013.
- [21] Andreas Maurer. A vector-contraction inequality for rademacher complexities. In *International Conference on Algorithmic Learning Theory*, pages 3–17. Springer, 2016.
- [22] Corinna Cortes, Vitaly Kuznetsov, Mehryar Mohri, and Scott Yang. Structured prediction theory based on factor graph complexity. In *Advances in Neural Information Processing Systems*, pages 2514–2522, 2016.
- [23] Olivier Bousquet and André Elisseeff. Stability and generalization. *Journal of machine learning research*, 2(Mar):499–526, 2002.
- [24] Shivani Agarwal and Partha Niyogi. Generalization bounds for ranking algorithms via algorithmic stability. *Journal of Machine Learning Research*, 10(Feb):441–474, 2009.
- [25] Moritz Hardt, Benjamin Recht, and Yoram Singer. Train faster, generalize better: Stability of stochastic gradient descent. *arXiv preprint arXiv:1509.01240*, 2015.
- [26] Omar Rivasplata, Emilio Parrado-Hernández, John S Shawe-Taylor, Shiliang Sun, and Csaba Szepesvári. Pac-bayes bounds for stable algorithms with instance-dependent priors. In *Advances in Neural Information Processing Systems*, pages 9214–9224, 2018.
- [27] Saurabh Verma and Zhi-Li Zhang. Stability and generalization of graph convolutional neural networks. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1539–1548, 2019.
- [28] Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pages 1126–1135. JMLR. org, 2017.
- [29] Aravind Rajeswaran, Chelsea Finn, Sham M Kakade, and Sergey Levine. Meta-learning with implicit gradients. In *Advances in Neural Information Processing Systems*, pages 113–124, 2019.
- [30] David Belanger, Bishan Yang, and Andrew McCallum. End-to-end learning for structured prediction energy networks. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pages 429–439. JMLR. org, 2017.
- [31] Priya Donti, Brandon Amos, and J Zico Kolter. Task-based end-to-end model learning in stochastic optimization. In *Advances in Neural Information Processing Systems*, pages 5484–5494, 2017.
- [32] Brandon Amos and J Zico Kolter. Optnet: Differentiable optimization as a layer in neural networks. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pages 136–145. JMLR. org, 2017.
- [33] Marin Vlastelica Pogančič, Anselm Paulus, Vit Musil, Georg Martius, and Michal Rolinek. Differentiation of blackbox combinatorial solvers. In *International Conference on Learning Representations*, 2019.

- [34] Michal Rolínek, Paul Swoboda, Dominik Zietlow, Anselm Paulus, Vít Musil, and Georg Martius. Deep graph matching via blackbox differentiation of combinatorial solvers. *arXiv preprint arXiv:2003.11657*, 2020.
- [35] Quentin Berthet, Mathieu Blondel, Olivier Teboul, Marco Cuturi, Jean-Philippe Vert, and Francis Bach. Learning with differentiable perturbed optimizers. *arXiv preprint arXiv:2002.08676*, 2020.
- [36] Xiaohan Chen, Jialin Liu, Zhangyang Wang, and Wotao Yin. Theoretical linear convergence of unfolded ista and its practical weights and thresholds. In *Advances in Neural Information Processing Systems*, pages 9061–9071, 2018.
- [37] Aaron Ferber, Bryan Wilder, Bistra Dilkina, and Milind Tambe. Mipaal: Mixed integer program as a layer. In *AAAI*, pages 1504–1511, 2020.
- [38] Patrick Knobelreiter, Christian Reinbacher, Alexander Shekhovtsov, and Thomas Pock. End-to-end training of hybrid cnn-crf models for stereo. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 2339–2348, 2017.
- [39] Vlad Niculae, Andre Martins, Mathieu Blondel, and Claire Cardie. Sparsemap: Differentiable sparse structured inference. In *International Conference on Machine Learning*, pages 3799–3808, 2018.
- [40] Nitish Shirish Keskar, Dheevatsa Mudigere, Jorge Nocedal, Mikhail Smelyanskiy, and Ping Tak Peter Tang. On large-batch training for deep learning: Generalization gap and sharp minima. *arXiv preprint arXiv:1609.04836*, 2016.
- [41] Kenji Kawaguchi, Leslie Pack Kaelbling, and Yoshua Bengio. Generalization in deep learning. *arXiv preprint arXiv:1710.05468*, 2017.
- [42] Behnam Neyshabur, Srinadh Bhojanapalli, David McAllester, and Nati Srebro. Exploring generalization in deep learning. In *Advances in Neural Information Processing Systems*, pages 5947–5956, 2017.
- [43] Behnam Neyshabur, Srinadh Bhojanapalli, and Nathan Srebro. A PAC-bayesian approach to spectrally-normalized margin bounds for neural networks. In *International Conference on Learning Representations*, 2018.
- [44] Vladimir Koltchinskii and Dmitriy Panchenko. Rademacher processes and bounding the risk of function learning. In *High dimensional probability II*, pages 443–457. Springer, 2000.
- [45] Vladimir Koltchinskii. Rademacher penalties and structural risk minimization. *IEEE Transactions on Information Theory*, 47(5):1902–1914, 2001.
- [46] Vladimir Koltchinskii, Dmitry Panchenko, et al. Empirical margin distributions and bounding the generalization error of combined classifiers. *The Annals of Statistics*, 30(1):1–50, 2002.
- [47] Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3(Nov):463–482, 2002.
- [48] Michel Talagrand. Sharper bounds for gaussian and empirical processes. *The Annals of Probability*, pages 28–76, 1994.
- [49] Richard M Dudley. *Uniform central limit theorems*, volume 142. Cambridge university press, 2014.
- [50] Shaojie Bai, J Zico Kolter, and Vladlen Koltun. Deep equilibrium models. In *Advances in Neural Information Processing Systems*, pages 688–699, 2019.
- [51] Laurent El Ghaoui, Fangda Gu, Bertrand Travacca, and Armin Askari. Implicit deep learning. *arXiv preprint arXiv:1908.06315*, 2019.
- [52] Pablo Arbelaez, Michael Maire, Charless Fowlkes, and Jitendra Malik. Contour detection and hierarchical image segmentation. *IEEE transactions on pattern analysis and machine intelligence*, 33(5):898–916, 2010.
- [53] Kai Zhang, Wangmeng Zuo, Yunjin Chen, Deyu Meng, and Lei Zhang. Beyond a gaussian denoiser: Residual learning of deep cnn for image denoising. *IEEE Transactions on Image Processing*, 26(7):3142–3155, 2017.

A Proof of Algorithm Properties

In this section, we study several important properties of gradient descent algorithm (GD) and Nesterov's accelerated gradient algorithm (NAG), which have already been summarized in Table 1 of Section 3. To simplify the presentation, we shall focus on quadratic minimization problems as in Section 2 and estimate the sharp dependence on the iteration number k .

More precisely, in the subsequent analysis, we shall fix the constants $L \geq \mu > 0$ and assume the objective function is in the function class $\mathcal{Q}_{\mu,L}$, which contains all μ -strongly convex and L -smooth quadratic functions on $\mathbb{R}^d$. Then, for any given $f \in \mathcal{Q}_{\mu,L}$, the eigenvalue decomposition enables us to represent the Hessian matrix of f , denoted by Q , as $Q = U\Lambda U^\top$, where Λ is a diagonal matrix comprising of the eigenvalues $(\lambda_i)_{i=1}^d$ of Q sorted in ascending order, i.e., $\mu \leq \lambda_1 \leq \dots \leq \lambda_d \leq L$, and $U \in \mathbb{R}^{d \times d}$ is an orthogonal matrix whose columns constitute an orthonormal basis of corresponding eigenvectors of Q . Moreover, we shall denote by $\mathbb{I}_d$ the $d \times d$ identity matrix, and by $\|A\|_2$ the spectral norm of a given matrix $A \in \mathbb{R}^{d \times d}$.

We start with the GD algorithm. Let $f \in \mathcal{Q}_{\mu,L}$, $s \geq 0$ be the stepsize, and $x_0 \in \mathbb{R}^d$ be the initial guess. For each $k \in \mathbb{N} \cup \{0\}$, we denote by x_{k+1} the $k+1$ -th iterate generated by the following recursive formula (cf. the output $\mathbf{y}_{k+1}$ of GD_ϕ in Section 3):

$$x_{k+1} = x_k - s\nabla f(x_k). \quad (21)$$

The following theorem establishes the convergence of Eq. 21 as k tends to infinity, and the Lipschitz dependence of the iterates $(x_k^s)_{k \in \mathbb{N}}$ in terms of the stepsize s (i.e., the sensitivity of GD). Similar results can be established for general μ -strongly convex and L -smooth objective functions.

Theorem A.1. *Let $f \in \mathcal{Q}_{\mu,L}$ admit the minimiser $x^* \in \mathbb{R}^d$, $x_0 \in \mathbb{R}^d$ and for each $s \geq 0$ let $(x_k^s)_{k \in \mathbb{N} \cup \{0\}}$ be the iterates generated by Eq. 21 with stepsize s . Then we have for all $k \in \mathbb{N}$, $c_0 > 0$, $s, t \in [c_0, \frac{2}{\mu+L}]$ that*

$$\|x_k^s - x^*\|_2 \leq (1 - s\mu)^k \|x_0 - x^*\|_2, \quad \|x_k^t - x_k^s\|_2 \leq Lk(1 - c_0\mu)^{k-1} |t - s| \|x_0 - x^*\|_2. \quad (22)$$

Proof. Let Q be the Hessian matrix of f and $(\lambda_i)_{i=1}^d$ be the eigenvalues of Q . By using the fact that $\nabla f(x^*) = 0$ and Eq. 21, we can obtain for all $k \in \mathbb{N} \cup \{0\}$ and $s \geq 0$ that $x_k^s - x^* = (\mathbb{I}_d - sQ)(x_{k-1}^s - x^*) = (\mathbb{I}_d - sQ)^k(x_0 - x^*)$.

Since the spectral norm of a matrix is invariant under orthogonal transformations, we have for all $s \in [c_0, \frac{2}{\mu+L}]$ that

$$\begin{aligned} \|\mathbb{I}_d - sQ\|_2 &= \|\mathbb{I}_d - s\Lambda\|_2 = \max_{i=1,\dots,d} |1 - s\lambda_i| = \max(|1 - s\mu|, |1 - sL|) \\ &\leq 1 - s\mu. \end{aligned} \quad (23)$$

Hence, for any given $k \in \mathbb{N} \cup \{0\}$, the inequality that $\|x_k^s - x^*\|_2 \leq (\|\mathbb{I}_d - sQ\|_2)^k \|x_0 - x^*\|_2$ leads us to the desired estimate for $(\|x_k^s - x^*\|_2)_{k \in \mathbb{N} \cup \{0\}}$.

Now let $t, s \in [c_0, \frac{2}{\mu+L}]$ be given, by using the fact that $\frac{d}{ds} x_k^s = k(\mathbb{I}_d - sQ)^{k-1} Q(x_0 - x^*)$ for all $s > 0$, we can deduce from the mean value theorem that

$$\begin{aligned} \|x_k^s - x_k^t\|_2 &\leq \left(\sup_{r \in (c_0, \frac{2}{\mu+L})} \left\| \frac{d}{dr} x_k^r \right\|_2 \right) |t - s| \\ &\leq \left(\sup_{r \in (c_0, \frac{2}{\mu+L})} k(\|\mathbb{I}_d - rQ\|_2)^{k-1} \|Q\|_2 \|x_0 - x^*\|_2 \right) |t - s| \\ &\leq k \left(\sup_{r \in [c_0, \frac{2}{\mu+L}]} \|\mathbb{I}_d - rQ\|_2 \right)^{k-1} L |t - s| \|x_0 - x^*\|_2, \end{aligned}$$

which along with Eq. 23 finishes the proof of the desired sensitivity estimate. $\square$

The next theorem shows that Eq. 21 with stepsize $s \in (0, \frac{2}{\mu+L}]$ is Lipschitz stable in terms of the perturbations of f . In particular, for a quadratic function $f \in \mathcal{Q}_{\mu,L}$, we shall establish the Lipschitz

stability with respect to the perturbations in the parameters of f . For notational simplicity, we assume $x_0 = 0$ as in Section 3, but it is straightforward to extend the results to an arbitrary initial guess $x_0 \in \mathbb{R}^d$.

Theorem A.2. *Let $x_0 = 0$, for each $i \in \{1, 2\}$ let $f_i \in \mathcal{Q}_{\mu, L}$ admit the minimizer $x^{*,i} \in \mathbb{R}^d$ and satisfy $\nabla f_i(x) = Q_i x + b_i$ for a symmetric matrix $Q_i \in \mathbb{R}^{d \times d}$ and $b_i \in \mathbb{R}^d$, for each $i \in \{1, 2\}$, $s > 0$ let $(x_{k,i}^s)_{k \in \mathbb{N} \cup \{0\}}$ be the iterates generated by Eq. 21 with $f = f_i$ and stepsize s , and let $M = \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2)$. Then we have for all $k \in \mathbb{N}$, $c_0 > 0$, $s \in [c_0, \frac{2}{\mu+L}]$ that:*

$$\begin{aligned} \|x_{k,1}^s - x_{k,2}^s\|_2 &\leq \left[\frac{1}{\mu} (1 - (1 - s\mu)^k) + sk(1 - s\mu)^{k-1} \right] M \|Q_1 - Q_2\|_2 \\ &\quad + \frac{1}{\mu} (1 - (1 - s\mu)^k) \|b_1 - b_2\|_2. \end{aligned}$$

Proof. Let us assume without loss of generality that $\|x^{*,2}\|_2 \leq \|x^{*,1}\|_2$ and $c_0 \leq \frac{2}{\mu+L}$. We write $\delta x_k = x_{k,1}^s - x_{k,2}^s$ for each $k \in \mathbb{N} \cup \{0\}$. Then, by using Eq. 21 and the fact that $\nabla f_1(x) - \nabla f_1(y) = Q_1(x - y)$ for all $x, y \in \mathbb{R}^d$, we can deduce that $\delta x_0 = 0$ and for all $k \in \mathbb{N} \cup \{0\}$ that

$$\delta x_{k+1} = (\mathbb{I}_d - sQ_1)\delta x_k + e_k = \sum_{i=0}^k (\mathbb{I}_d - sQ_1)^i e_{k-i},$$

where $e_k = -s(\nabla f_1 - \nabla f_2)(x_{k,2}^s)$ for each $k \in \mathbb{N} \cup \{0\}$. Note that it holds for all $k \in \mathbb{N} \cup \{0\}$ that

$$\begin{aligned} \|e_k\|_2 &\leq s\|(\nabla f_1 - \nabla f_2)(x_{k,2}^s)\|_2 \leq s(\|Q_2 - Q_1\|_2 \|x_{k,2}^s\|_2 + \|b_1 - b_2\|_2) \\ &\leq s(\|Q_2 - Q_1\|_2 (\|x^{*,2}\|_2 + \|x_{k,2}^s - x^{*,2}\|_2) + \|b_1 - b_2\|_2) \\ &\leq s(\|Q_2 - Q_1\|_2 (\|x^{*,2}\|_2 + (1 - s\mu)^k \|x_0 - x^{*,2}\|_2) + \|b_1 - b_2\|_2), \end{aligned}$$

where we have applied Theorem A.1 for the last inequality. Thus for each $k \in \mathbb{N}$, we can obtain from Eq. 23 and $x_0 = 0$ that

$$\begin{aligned} \|\delta x_k\|_2 &\leq \sum_{i=0}^{k-1} (\|\mathbb{I}_d - sQ_1\|_2)^i \|e_{k-1-i}\|_2 \\ &\leq \sum_{i=0}^{k-1} (1 - s\mu)^i s [(1 + (1 - s\mu)^{k-1-i}) \|x^{*,2}\|_2 \|Q_2 - Q_1\|_2 + \|b_1 - b_2\|_2] \\ &= \left[\frac{1}{\mu} (1 - (1 - s\mu)^k) + sk(1 - s\mu)^{k-1} \right] \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2) \|Q_2 - Q_1\|_2 \\ &\quad + \frac{1}{\mu} (1 - (1 - s\mu)^k) \|b_1 - b_2\|_2. \end{aligned}$$

which leads to the desired conclusion due to the fact that $M = \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2)$. $\square$

We now proceed to investigate similar properties of the NAG algorithm, whose proofs are more involved due to the fact that NAG is a multi-step method.

Recall that for any given $f \in \mathcal{Q}_{\mu, L}$, initial guess $x_0 \in \mathbb{R}^d$ and stepsize $s \geq 0$, the NAG algorithm generates iterates $(x_k, y_k)_{k \in \mathbb{N} \cup \{0\}}$ as follows: $y_0 = x_0$ and for each $k \in \mathbb{N} \cup \{0\}$,

$$x_{k+1} = y_k - s\nabla f(y_k), \quad y_{k+1} = x_{k+1} + \frac{1 - \sqrt{\mu s}}{1 + \sqrt{\mu s}}(x_{k+1} - x_k). \quad (24)$$

Note that x_{k+1}, y_{k+1} are denoted by $\mathbf{y}_{k+1}, \mathbf{z}_{k+1}$, respectively, in Section 3.

We first introduce the following matrix $R_{\text{NAG}, s}$ for Eq. 24 for any given function $f \in \mathcal{Q}_{\mu, L}$ and stepsize $s \in [0, \frac{4}{3L+\mu}]$:

$$R_{\text{NAG}, s} := \begin{pmatrix} (1 + \beta_s)(\mathbb{I}_d - sQ) & -\beta_s(\mathbb{I}_d - sQ) \\ \mathbb{I}_d & 0 \end{pmatrix} \quad (25)$$

where $\beta_s = \frac{1-\sqrt{\mu s}}{1+\sqrt{\mu s}}$ and Q is the Hessian matrix of f . The following lemma establishes an upper bound of the spectral norm of the k -th power of $R_{\text{NAG},s}$, which extends [18, Lemma 22] to block matrices, a wider range of stepsize (s is allowed to be larger than $1/L$) and a momentum parameter β_s depending on the stepsize s .

Lemma A.1. *Let $f \in \mathcal{Q}_{\mu,L}$, $s \in (0, \frac{4}{3L+\mu}]$, $\beta_s = \frac{1-\sqrt{\mu s}}{1+\sqrt{\mu s}}$ and $R_{\text{NAG},s}$ be defined as in Eq. 25. Then we have for all $k \in \mathbb{N}$ that $\|R_{\text{NAG},s}^k\|_2 \leq 2(k+1)(1-\sqrt{\mu s})^k$.*

Proof. Let $Q = U\Lambda U^T$ be the eigenvalue decomposition of the Hessian matrix Q of f , where Λ is a diagonal matrix comprising of the corresponding eigenvalues of Q sorted in ascending order, i.e., $0 < \mu \leq \lambda_1 \leq \dots \leq \lambda_d \leq L$. Then we have that

$$R_{\text{NAG},s} = \begin{pmatrix} U & 0 \\ 0 & U \end{pmatrix} \begin{pmatrix} (1+\beta_s)(\mathbb{I}_d - s\Lambda) & -\beta_s(\mathbb{I}_d - s\Lambda) \\ \mathbb{I}_d & 0 \end{pmatrix} \begin{pmatrix} U^T & 0 \\ 0 & U^T \end{pmatrix},$$

which together with the facts that any permutation matrix is orthogonal, and the spectral norm of a matrix is invariant under orthogonal transformations, gives us the identity that: for all $k \in \mathbb{N}$,

$$\|R_{\text{NAG},s}^k\|_2 = \left\| \begin{pmatrix} (1+\beta_s)(\mathbb{I}_d - s\Lambda) & -\beta_s(\mathbb{I}_d - s\Lambda) \\ \mathbb{I}_d & 0 \end{pmatrix}^k \right\|_2 = \max_{i=1,\dots,d} \|T_{s,i}^k\|_2, \quad (26)$$

where $T_{s,i} = \begin{pmatrix} (1+\beta_s)(1-s\lambda_i) & -\beta_s(1-s\lambda_i) \\ 1 & 0 \end{pmatrix}$ for all $i = 1, \dots, d$.

Now let $s \in (0, \frac{4}{3L+\mu}]$ and $i = 1, \dots, d$ be fixed. If $1 - s\lambda_i \geq 0$, by using [18, Lemma 22] (with $\alpha = \mu$, $\beta = 1/s$, $h = 1 - s\lambda_i$ and $\kappa = \beta/\alpha = 1/(\mu s)$), we can obtain that

$$\|T_{s,i}^k\|_2 \leq 2(k+1) \left(\frac{1-\sqrt{\mu s}}{1+\sqrt{\mu s}} (1-\mu s) \right)^{k/2} \leq 2(k+1)(1-\sqrt{\mu s})^k.$$

We then discuss the case where $1 - s\lambda_i < 0$. Let us write $T_{s,i}^k = \begin{pmatrix} a_k & b_k \\ c_k & d_k \end{pmatrix}$ for each $k \in \mathbb{N} \cup \{0\}$, then we have for all $k \in \mathbb{N}$ that

$$\begin{aligned} a_k &= (1+\beta_s)(1-s\lambda_i)a_{k-1} - \beta_s(1-s\lambda_i)c_{k-1}, & c_k &= a_{k-1}, \\ b_k &= (1+\beta_s)(1-s\lambda_i)b_{k-1} - \beta_s(1-s\lambda_i)d_{k-1}, & d_k &= b_{k-1}, \end{aligned}$$

with $a_1 = (1+\beta_s)(1-s\lambda_i)$, $b_1 = -\beta_s(1-s\lambda_i)$, $c_1 = 1$ and $d_1 = 0$. Since the conditions $1 - s\lambda_i < 0$ and $s \leq \frac{4}{3L+\mu}$ imply that $\lambda_i > \frac{1}{s} \geq \frac{3L+\mu}{4} \geq \mu$, we see the discriminant of the characteristic polynomial satisfies that

$$\Delta = (1+\beta_s)^2(1-s\lambda_i)^2 - 4\beta_s(1-s\lambda_i) = \frac{4(1-s\lambda_i)}{(1+\sqrt{\mu s})^2} s(\mu - \lambda_i) > 0,$$

which implies that there exist $l_1, l_2, l_3, l_4 \in \mathbb{R}$ such that it holds for all $k \in \mathbb{N} \cup \{0\}$ that $a_k = l_1\tau_+^{k+1} + l_2\tau_-^{k+1}$ and $b_k = l_3\tau_+^{k+1} + l_4\tau_-^{k+1}$, with $\tau_{\pm} = \frac{(1+\beta_s)(1-s\lambda_i) \pm \sqrt{\Delta}}{2}$, $l_1 = \frac{1}{\tau_+ - \tau_-}$, $l_2 = -\frac{1}{\tau_+ - \tau_-}$, $l_3 = \frac{-\tau_-}{\tau_+ - \tau_-}$ and $l_4 = \frac{\tau_+}{\tau_+ - \tau_-}$. Thus, by letting $\rho_i := \max(|\tau_+|, |\tau_-|)$, we have that $|a_k| = |\sum_{j=0}^k \tau_+^{k-j}\tau_-^j| \leq (k+1)\rho_i^k$ and $|b_k| = |(-\tau_+\tau_-)\sum_{j=0}^{k-1} \tau_+^{k-1-j}\tau_-^j| \leq k\rho_i^{k+1}$ for all $k \in \mathbb{N} \cup \{0\}$.

Now we claim that the conditions $1 - s\lambda_i < 0$ and $0 < s \leq \frac{4}{3L+\mu}$ imply the estimate that $\rho_i \leq 1 - \sqrt{\mu s} < 1$. In fact, the inequality $s \leq \frac{4}{3L+\mu}$ gives us that $\mu s \leq \frac{4\mu}{3L+\mu} \leq 1$, which implies that $\beta_s = \frac{1-\sqrt{\mu s}}{1+\sqrt{\mu s}} \geq 0$. Hence we can deduce from $1 - s\lambda_i < 0$ that $\sqrt{\Delta} \geq (1+\beta_s)(s\lambda_i - 1)$ and

$$|\tau_+| \leq |\tau_-| \leq \frac{s\lambda_i - 1 + \sqrt{(s\lambda_i - 1)s(\lambda_i - \mu)}}{1 + \sqrt{\mu s}} \leq \frac{sL - 1 + \sqrt{(sL - 1)s(L - \mu)}}{1 + \sqrt{\mu s}}.$$

Note that $2 - (\mu + L)s \geq 2 - \frac{4(\mu+L)}{3L+\mu} \geq 0$, we see that

$$\begin{aligned} \rho_i \leq 1 - \sqrt{\mu s} &\iff |\tau_-| \leq 1 - \sqrt{\mu s} \iff sL - 1 + \sqrt{(sL - 1)s(L - \mu)} \leq 1 - \mu s \\ &\iff (sL - 1)s(L - \mu) \leq (2 - (\mu + L)s)^2 \\ &\iff (us - 1)((3L + \mu)s - 4) \geq 0. \end{aligned}$$

Therefore, we have that $\max(|a_k|, |b_k|, |c_k|, |d_k|) \leq (k+1)(1 - \sqrt{\mu s})^k$, which, along with the relationship between the spectral norm and Frobenius norm, gives us that $\|T_{s,i}^k\|_2 \leq \|T_{s,i}^k\|_F \leq 2(k+1)(1 - \sqrt{\mu s})^k$, and finishes the proof of the desired estimate for the case with $1 - s\lambda_i < 0$. $\square$

As an important consequence of Lemma A.1, we now obtain the following upper bound of the error $(\|x_k - x^*\|_2)_{k \in \mathbb{N}}$ for any given objective function $f \in \mathcal{Q}_{\mu,L}$ and stepsize $s \in (0, \frac{4}{3L+\mu}]$.

Theorem A.3. *Let $f \in \mathcal{Q}_{\mu,L}$ admit the minimizer $x^* \in \mathbb{R}^d$, $x_0 \in \mathbb{R}^d$, $s \in (0, \frac{4}{3L+\mu}]$ and $(x_k^s, y_k^s)_{k \in \mathbb{N} \cup \{0\}}$ be the iterates generated by Eq. 24 with stepsize s . Then we have for all $k \in \mathbb{N} \cup \{0\}$ that*

$$\|x_{k+1}^s - x^*\|_2^2 + \|x_k^s - x^*\|_2^2 \leq 8(1+k)^2(1 - \sqrt{\mu s})^{2k} \|x_0 - x^*\|_2^2.$$

Proof. For any $f \in \mathcal{Q}_{\mu,L}$, and $s \in (0, \frac{4}{3L+\mu}]$, by letting $\beta_s = \frac{1-\sqrt{\mu s}}{1+\sqrt{\mu s}}$, we can rewrite Eq. 24 as follows: $x_0^s = x_0$, $x_1^s = x_0 - s\nabla f(x_0)$ and for all $k \in \mathbb{N}$,

$$x_{k+1}^s = (1 + \beta_s)x_k^s - \beta_s x_{k-1} - s\nabla f((1 + \beta_s)x_k^s - \beta_s x_{k-1}), \quad (27)$$

which together with the fact that $\nabla f(x^*) = 0$ shows that

$$\begin{pmatrix} x_{k+1}^s - x^* \\ x_k^s - x^* \end{pmatrix} = R_{\text{NAG},s} \begin{pmatrix} x_k^s - x^* \\ x_{k-1}^s - x^* \end{pmatrix} = R_{\text{NAG},s}^k \begin{pmatrix} x_1^s - x^* \\ x_0^s - x^* \end{pmatrix}$$

where $R_{\text{NAG},s}$ is defined as in Eq. 25. Hence by using $x_1^s = x_0 - s\nabla f(x_0)$ and Theorem A.1, we can obtain that

$$\begin{aligned} \|x_{k+1}^s - x^*\|_2^2 + \|x_k^s - x^*\|_2^2 &\leq \|R_{\text{NAG},s}^k\|_2^2 (\|x_1^s - x^*\|_2^2 + \|x_0^s - x^*\|_2^2) \\ &\leq \|R_{\text{NAG},s}^k\|_2^2 2\|x_0 - x^*\|_2^2, \end{aligned}$$

which together with Lemma A.1 leads to the desired convergence result. $\square$

Remark A.1. It is well-known that for a general μ -strongly convex and L -smooth objective function f , one can employ a Lyapunov argument and establish that the iterates obtained by Eq. 24 with stepsize $s \in [0, \frac{1}{L}]$ satisfy the estimate that $\|x_k - x^*\|_2^2 \leq \frac{2L}{\mu}(1 - \sqrt{\mu s})^k \|x_0 - x^*\|_2^2$. Here by taking advantage of the affine structure of ∇f , we have obtained a sharper estimate of the convergence rate for a wider range of stepsize $s \in (0, \frac{4}{3L+\mu}]$.

We also would like to emphasize that the upper bound in Theorem A.3 is tight, in the sense that the additional quadratic dependence on k in the error estimate is inevitable. In fact, one can derive a *closed-form expression* of $R_{\text{NAG},s}^k$ and show that, for an index i such that the eigenvalue λ_i is sufficiently close to μ , the squared error for that component is of the magnitude $\mathcal{O}((k\sqrt{\mu s} + 1)^2(1 - \sqrt{\mu s})^{2k})$.

We then proceed to analyze the sensitivity of Eq. 24 with respect to the stepsize. The following theorem shows that the iterates $(x_k, y_k)_{k \in \mathbb{N} \cup \{0\}}$ generated by Eq. 24 depend Lipschitz continuously on the stepsize s .

Theorem A.4. *Let $f \in \mathcal{Q}_{\mu,L}$ admit the minimiser $x^* \in \mathbb{R}^d$, $x_0 \in \mathbb{R}^d$, and for each $s \in (0, \frac{4}{3L+\mu}]$ let $(x_k^s, y_k^s)_{k \in \mathbb{N} \cup \{0\}}$ be the iterates generated by Eq. 24 with stepsize s . Then we have for all $k \in \mathbb{N}$, $c_0 > 0$ and $t, s \in [c_0, \frac{4}{3L+\mu}]$ that:*

$$\|x_k^t - x_k^s\|_2 \leq \left(2L(1+k) + \frac{4}{3}k(k+1)(k+5) \left(\sqrt{\frac{\mu}{c_0}} + 2L\right)\right) (1 - \sqrt{\mu c_0})^k |t - s| \|x_0 - x^*\|_2.$$

Proof. Throughout this proof we assume without loss of generality that $c_0 \leq s < t \leq \frac{4}{3L+\mu}$. Let Q be the Hessian matrix of f , for each $r \in [c_0, \frac{4}{3L+\mu}]$ let $\beta_r = \frac{1-\sqrt{\mu r}}{1+\sqrt{\mu r}}$, and for each $k \in \mathbb{N} \cup \{0\}$ let $\delta x_k = x_k^t - x_k^s$. Then we can deduce from Eq. 27 that $\delta x_0 = 0$, $\delta x_1 = -(t-s)\nabla f(x_0)$ and for all $k \in \mathbb{N}$ that

$$\begin{aligned} x_{k+1}^t - x_{k+1}^s &= [(1 + \beta_t)x_k^t - \beta_t x_{k-1}^t - t\nabla f((1 + \beta_t)x_k^t - \beta_t x_{k-1}^t)] \\ &\quad - [(1 + \beta_s)x_k^s - \beta_s x_{k-1}^s - s\nabla f((1 + \beta_s)x_k^s - \beta_s x_{k-1}^s)], \end{aligned}$$

which together with the fact that $\nabla f(x) - \nabla f(y) = Q(x - y)$ for all $x, y \in \mathbb{R}^d$ shows that

$$\begin{pmatrix} \delta x_{k+1} \\ \delta x_k \end{pmatrix} = R_{\text{NAG},t} \begin{pmatrix} \delta x_k \\ \delta x_{k-1} \end{pmatrix} + \begin{pmatrix} e_k \\ 0 \end{pmatrix}$$

with $R_{\text{NAG},t}$ defined as in Eq. 25 and the following residual term

$$\begin{aligned} e_k &:= [(1 + \beta_t)x_k^s - \beta_t x_{k-1}^s - t \nabla f((1 + \beta_t)x_k^s - \beta_t x_{k-1}^s)] \\ &\quad - [(1 + \beta_s)x_k^s - \beta_s x_{k-1}^s - s \nabla f((1 + \beta_s)x_k^s - \beta_s x_{k-1}^s)]. \end{aligned}$$

Hence we can obtain by induction that: for all $k \in \mathbb{N}$,

$$\begin{pmatrix} \delta x_{k+1} \\ \delta x_k \end{pmatrix} = R_{\text{NAG},t}^k \begin{pmatrix} \delta x_1 \\ \delta x_0 \end{pmatrix} + \sum_{i=0}^{k-1} R_{\text{NAG},t}^i \begin{pmatrix} e_{k-i} \\ 0 \end{pmatrix}. \quad (28)$$

Now the facts that $\nabla f(x^*) = 0$ and $\nabla^2 f \equiv Q$ gives us that

$$\begin{aligned} e_k &= (\beta_t - \beta_s)(x_k^s - x_{k-1}^s) - t \nabla f((1 + \beta_t)x_k^s - \beta_t x_{k-1}^s) + s \nabla f((1 + \beta_s)x_k^s - \beta_s x_{k-1}^s) \\ &= (\beta_t - \beta_s)((x_k^s - x^*) - (x_{k-1}^s - x^*)) - tQ((1 + \beta_t)(x_k^s - x^*) - \beta_t(x_{k-1}^s - x^*)) \\ &\quad + sQ((1 + \beta_s)(x_k^s - x^*) - \beta_s(x_{k-1}^s - x^*)) \\ &= [(\beta_t - \beta_s) - (t + t\beta_t - s - s\beta_s)Q](x_k^s - x^*) - [(\beta_t - \beta_s) - (t\beta_t - s\beta_s)Q](x_{k-1}^s - x^*). \end{aligned}$$

Note that one can easily verify that the function $g_1(r) = \beta_r$ is $\sqrt{\mu/c_0}$ -Lipschitz on $[c_0, \frac{4}{3L+\mu}]$, and the function $g_2(r) = r\beta_r$ is 1-Lipschitz on $[0, \frac{4}{3L+\mu}]$. Moreover, the fact that $f \in \mathcal{Q}_{\mu,L}$ implies that $\|Q\|_2 \leq L$. Thus we can obtain from Theorem A.3 that

$$\begin{aligned} \|e_k\|_2 &\leq \left(\sqrt{\frac{\mu}{c_0}} + 2L \right) |t - s| \|x_k^s - x^*\|_2 + \left(\sqrt{\frac{\mu}{c_0}} + L \right) |t - s| \|x_{k-1}^s - x^*\|_2 \\ &\leq \left(\sqrt{\frac{\mu}{c_0}} + 2L \right) |t - s| \sqrt{2(\|x_k^s - x^*\|_2^2 + \|x_{k-1}^s - x^*\|_2^2)} \\ &\leq \left(\sqrt{\frac{\mu}{c_0}} + 2L \right) |t - s| 4(1 + k)(1 - \sqrt{\mu s})^k \|x_0 - x^*\|_2. \end{aligned}$$

This, along with Eq. 28, Lemma A.1 and $s < t$, gives us that

$$\begin{aligned} \sqrt{\|\delta x_{k+1}\|_2^2 + \|\delta x_k\|_2^2} &\leq \|R_{\text{NAG},t}^k\|_2 \|\delta x_1\|_2 + \sum_{i=0}^{k-1} \|R_{\text{NAG},t}^i\|_2 \|e_{k-i}\|_2 \\ &\leq 2(1 + k)(1 - \sqrt{\mu t})^k |t - s| L \|x_0 - x^*\|_2 \\ &\quad + \sum_{i=0}^{k-1} 2(1 + i)(1 - \sqrt{\mu t})^i \left(\sqrt{\frac{\mu}{c_0}} + 2L \right) |t - s| 4(1 + k - i)(1 - \sqrt{\mu s})^{k-i} \|x_0 - x^*\|_2 \\ &= \left(2L(1 + k) + \frac{4}{3}k(k + 1)(k + 5) \left(\sqrt{\frac{\mu}{c_0}} + 2L \right) \right) |t - s| (1 - \sqrt{\mu s})^k \|x_0 - x^*\|_2, \end{aligned}$$

which finishes the proof of the desired estimate due to the fact that $s \geq c_0$. $\square$

The next theorem is an analog of Theorem A.2 for the NAC scheme Eq. 24, which shows that the outputs of Eq. 24 with stepsize $s \in (0, \frac{4}{3L+\mu}]$ is Lipschitz stable with respect to the perturbations of the parameters in f .

Theorem A.5. *Let $x_0 = 0$, for each $i \in \{1, 2\}$ let $f_i \in \mathcal{Q}_{\mu,L}$ admit the minimizer $x^{*,i} \in \mathbb{R}^d$ and satisfy $\nabla f_i(x) = Q_i x + b_i$ for a symmetric matrix $Q_i \in \mathbb{R}^{d \times d}$ and $b_i \in \mathbb{R}^d$, for each $i \in \{1, 2\}$, $s > 0$ let $(x_{k,i}^s)_{k \in \mathbb{N} \cup \{0\}}$ be the iterates generated by Eq. 24 with $f = f_i$ and stepsize s , and let $M = \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2)$. Then we have for all $k \in \mathbb{N}$, $s \in [c_0, \frac{4}{3L+\mu}]$ that:*

$$\begin{aligned} \|x_{k,1}^s - x_{k,2}^s\|_2 &\leq \left[\frac{2}{\mu} (1 - (1 - \sqrt{\mu s})^{k-1}) + s \frac{8(k-1)k(k+4)}{3} (1 - \sqrt{\mu s})^{k-1} \right] M \|Q_1 - Q_2\|_2 \\ &\quad + \frac{2}{\mu} (1 - (1 - \sqrt{\mu s})^k) \|b_1 - b_2\|_2. \end{aligned}$$

Proof. Let us assume without loss of generality that $\|x^{*,2}\|_2 \leq \|x^{*,1}\|_2$. We first fix an arbitrary $s \in [c_0, \frac{4}{3L+\mu}]$ and write $\delta x_k = x_{k,1}^s - x_{k,2}^s$ for each $k \in \mathbb{N} \cup \{0\}$. Then, by using Eq. 27 and the fact that $\nabla f_1(x) - \nabla f_1(y) = Q_1(x - y)$ for all $x, y \in \mathbb{R}^d$, we can deduce that $\delta x_0 = 0$, $\delta x_1 = -s(\nabla f_1 - \nabla f_2)(x_0)$ and for all $k \in \mathbb{N}$,

$$\begin{pmatrix} \delta x_{k+1} \\ \delta x_k \end{pmatrix} = R_{\text{NAG},s} \begin{pmatrix} \delta x_k \\ \delta x_{k-1} \end{pmatrix} + \begin{pmatrix} e_k \\ 0 \end{pmatrix} = R_{\text{NAG},s}^k \begin{pmatrix} \delta x_1 \\ \delta x_0 \end{pmatrix} + \sum_{j=0}^{k-1} R_{\text{NAG},s}^j \begin{pmatrix} e_{k-j} \\ 0 \end{pmatrix}, \quad (29)$$

where $R_{\text{NAG},s}$ is defined as in Eq. 25 (with $Q = Q_1$) and the residual term e_k is given by

$$e_k := -s(\nabla f_1 - \nabla f_2)((1 + \beta_s)x_{k,2}^s - \beta_s x_{k-1,2}^s) \quad \forall k \in \mathbb{N}.$$

Note that, by using Theorem A.3 and the inequality that $x + y \leq \sqrt{2(x^2 + y^2)}$ for all $x, y \in \mathbb{R}$, we have for each $k \in \mathbb{N}$ that

$$\begin{aligned} \|e_k\|_2 &= s\|(Q_1 - Q_2)((1 + \beta_s)x_{k,2}^s - \beta_s x_{k-1,2}^s) + (b_1 - b_2)\|_2 \\ &\leq s\|Q_1 - Q_2\|_2(\|x^{*,2}\|_2 + 2\|x_{k,2}^s - x^{*,2}\|_2 + \|x_{k-1,2}^s - x^{*,2}\|_2) + s\|b_1 - b_2\|_2 \\ &\leq s\|Q_1 - Q_2\|_2(\|x^{*,2}\|_2 + 2\|x_{k,2}^s - x^{*,2}\|_2 + \|x_{k-1,2}^s - x^{*,2}\|_2) + s\|b_1 - b_2\|_2 \\ &\leq s\|Q_1 - Q_2\|_2(\|x^{*,2}\|_2 + 8(1 + k)(1 - \sqrt{\mu s})^k\|x_0 - x^{*,2}\|_2) + s\|b_1 - b_2\|_2. \end{aligned}$$

Hence we can obtain from Eq. 29, Lemma A.1 and $x_0 = 0$ that

$$\begin{aligned} \sqrt{\|\delta x_{k+1}\|_2^2 + \|\delta x_k\|_2^2} &\leq 2(k+1)(1 - \sqrt{\mu s})^k \|\delta x_1\|_2 + \sum_{j=0}^{k-1} 2(j+1)(1 - \sqrt{\mu s})^j \|e_{k-j}\|_2 \\ &\leq 2(k+1)(1 - \sqrt{\mu s})^k s\|b_1 - b_2\|_2 + \sum_{j=0}^{k-1} 2(j+1)(1 - \sqrt{\mu s})^j [s\|b_1 - b_2\|_2 \\ &\quad + s\|Q_1 - Q_2\|_2(1 + 8(1 + k - j)(1 - \sqrt{\mu s})^{k-j})\|x^{*,2}\|_2] \\ &\leq 2s \sum_{j=0}^k (j+1)(1 - \sqrt{\mu s})^j \|b_1 - b_2\|_2 + 2s \sum_{j=0}^{k-1} [(j+1)(1 - \sqrt{\mu s})^j \\ &\quad + 8(j+1)(1 + k - j)(1 - \sqrt{\mu s})^k] \|Q_1 - Q_2\|_2 \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2). \end{aligned}$$

Let $p = 1 - \sqrt{\mu s} \in [0, 1]$, then we can easily show for each $k \in \mathbb{N} \cup \{0\}$ that $(1-p) \sum_{j=0}^k (j+1)p^j = \sum_{j=0}^k p^j - p^{k+1}$, which implies that $\sum_{j=0}^k (j+1)(1 - \sqrt{\mu s})^j \leq \frac{1 - (1 - \sqrt{\mu s})^{k+1}}{\mu s}$. Moreover, we have that $\sum_{j=0}^{k-1} (j+1)(1 + k - j) = \frac{k(k+1)(k+5)}{6}$ for all $k \in \mathbb{N}$. Thus we can simplify the above estimate and deduce for each $k \in \mathbb{N}$ that

$$\begin{aligned} \|\delta x_{k+1}\|_2 &\leq \frac{2}{\mu} (1 - (1 - \sqrt{\mu s})^{k+1}) \|b_1 - b_2\|_2 + \left[\frac{2}{\mu} (1 - (1 - \sqrt{\mu s})^k) \right. \\ &\quad \left. + s \frac{8k(k+1)(k+5)}{3} (1 - \sqrt{\mu s})^k \right] \|Q_1 - Q_2\|_2 \min(\|x^{*,1}\|_2, \|x^{*,2}\|_2). \end{aligned}$$

Moreover, the condition that $s \leq \frac{4}{3L+\mu} \leq \frac{1}{\mu}$ implies that $\|\delta x_1\|_2 = s\|b_1 - b_2\|_2 \leq \frac{2}{\mu} (1 - (1 - \sqrt{\mu s})) \|b_1 - b_2\|_2$, which shows that the same upper bound also holds for $\|\delta x_1\|_2$ and finishes the proof of the desired estimate. $\square$

B Approximation Ability

Lemma 4.1. (Faster Convergence $\Rightarrow$ Better Approximation Ability). Assume the problem setting in Sec 2. The approximation ability can be bounded by two terms:

$$\inf_{\phi \in \Phi, \theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\ell_{\phi, \theta} \leq \underbrace{\sigma_b \mu^{-2} \inf_{\theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\|Q_\theta - Q^*\|_F}_{\text{approximation ability of the neural module}} + \underbrace{M \inf_{\phi \in \Phi} \text{Cvg}(k, \phi)}_{\text{best convergence}}. \quad (11)$$

Proof. For each $\phi \in \Phi, \theta \in \Theta, Q^* \in \mathcal{Q}^*$,

$$\ell_{\phi, \theta}(\mathbf{x}, \mathbf{b}) = \|\text{Alg}_\phi^k(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 \quad (30)$$

$$\leq \|\text{Alg}_\phi^k(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q_\theta(\mathbf{x}), \mathbf{b})\|_2 + \|\text{Opt}(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 \quad (31)$$

$$\leq \text{Cvg}(k, \phi) \|\text{Alg}_\phi^0(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 + \|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2 \quad (32)$$

$$\leq \text{Cvg}(k, \phi) \cdot M + \|(Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1}) \mathbf{b}\|_2, \quad (33)$$

where in the last inequality we have used the facts that the initialization is assumed to be zero vector, i.e., $\text{Alg}_\phi^0(Q_\theta(\mathbf{x}), \mathbf{b}) = \mathbf{0}$, and that $M \geq \sup_{\mathbf{x} \in \mathcal{X}, \mathbf{b} \in \mathcal{B}} \|\text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2$. Note that the independence of $(\mathbf{x}, \mathbf{b})$ and the fact that $\mathbb{E} \mathbf{b} \mathbf{b}^\top = \sigma_b^2 I$ imply that

$$\mathbb{E}_{\mathbf{b}} \|(Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1}) \mathbf{b}\|_2^2 \quad (34)$$

$$= \text{Tr}((Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1})^\top (Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1}) \sigma_b^2 I) \quad (35)$$

$$= \sigma_b^2 \|Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1}\|_F^2 \quad (36)$$

$$= \sigma_b^2 \|Q_\theta(\mathbf{x})^{-1} (Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})) Q^*(\mathbf{x})^{-1}\|_F^2 \quad (37)$$

$$\leq \mu^{-4} \sigma_b^2 \|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F^2 \quad (38)$$

Therefore, we see from Hölder's inequality that

$$\mathbb{E}_{\mathbf{b}} \|(Q_\theta(\mathbf{x})^{-1} - Q^*(\mathbf{x})^{-1}) \mathbf{b}\|_2 \leq \mu^{-2} \sigma_b \|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F. \quad (39)$$

Collecting all the above inequalities, we have

$$P\ell_{\phi, \theta} \leq \text{Cvg}(k, \phi) \cdot M + \sigma_b \mu^{-2} P\|Q_\theta - Q^*\|_F. \quad (40)$$

Taking supremum over Q^* , we have

$$\sup_{Q^* \in \mathcal{Q}^*} P\ell_{\phi, \theta} \leq \text{Cvg}(k, \phi) \cdot M + \sigma_b \mu^{-2} \sup_{Q^* \in \mathcal{Q}^*} P\|Q_\theta - Q^*\|_F. \quad (41)$$

Taking infimum over ϕ and θ , we have

$$\inf_{\phi \in \Phi, \theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\ell_{\phi, \theta} \leq \inf_{\phi \in \Phi} \text{Cvg}(k, \phi) \cdot M + \sigma_b \mu^{-2} \inf_{\theta \in \Theta} \sup_{Q^* \in \mathcal{Q}^*} P\|Q_\theta - Q^*\|_F. \quad (42)$$

□

Lemma 4.2. (Faster Convergence $\Rightarrow$ Better Representation of Q^*). Assume the problem setting in Sec 2. Then $\forall \phi \in \Phi, \theta \in \Theta, Q^* \in \mathcal{Q}^*$ it holds true that

$$P\|Q_\theta - Q^*\|_F^2 \leq \sigma_b^{-2} L^4 (\sqrt{P\ell_{\phi, \theta}^2} + M \cdot \text{Cvg}(k, \phi))^2. \quad (12)$$

Proof. Let us assume without loss of generality that $P\ell_{\phi, \theta}^2 = \epsilon$ for some $\epsilon \geq 0$. For any $\mathbf{x} \in \mathcal{X}, \mathbf{b} \in \mathcal{B}$, we have

$$\begin{aligned} \ell_{\phi, \theta}(\mathbf{x}) &\geq \|\text{Opt}(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 - \|\text{Alg}_\phi^k(Q_\theta(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q_\theta(\mathbf{x}), \mathbf{b})\|_2 \\ &\geq \|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2 - \text{Cvg}(k, \phi) \|\text{Opt}(Q_\theta(\mathbf{x}), \mathbf{b})\|_2 \end{aligned} \quad (43)$$

$$\geq \|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2 - M \cdot \text{Cvg}(k, \phi). \quad (44)$$

Rearranging the terms in the above inequality, we have

$$\|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2 \leq \ell_{\phi, \theta}(\mathbf{x}) + M \cdot \text{Cvg}(k, \phi). \quad (45)$$

By Eq. 37 and the inequality that $\|AB\|_F \leq \|A\|_2 \|B\|_F$ for any given $A \in \mathbb{R}^{m \times r}$ and $B \in \mathbb{R}^{r \times n}$, we have that

$$\mathbb{E}_{\mathbf{b}} \|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2^2 \quad (46)$$

$$= \sigma_b^2 \|Q_\theta(\mathbf{x})^{-1} (Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})) Q^*(\mathbf{x})^{-1}\|_F^2 \quad (47)$$

$$\geq \sigma_b^2 \frac{\|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F^2}{\|Q^*(\mathbf{x})\|_2^2 \|Q_\theta(\mathbf{x})\|_2^2} \quad (48)$$

$$\geq \sigma_b^2 \|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F^2 / L^4, \quad (49)$$

which implies that,

$$\|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F^2 \leq \sigma_b^{-2} L^4 \mathbb{E}_{\mathbf{b}} \|Q_\theta(\mathbf{x})^{-1} \mathbf{b} - Q^*(\mathbf{x})^{-1} \mathbf{b}\|_2^2. \quad (50)$$

Combining it with Eq. 45 and the fact that $(P\ell_{\phi,\theta})^2 \leq P\ell_{\phi,\theta}^2$, we have

$$P\|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F^2 \leq \sigma_b^{-2} L^4 P(\ell_{\phi,\theta} + M \cdot \text{Cvg}(k, \phi))^2 \quad (51)$$

$$= \sigma_b^{-2} L^4 (P\ell_{\phi,\theta}^2 + (M \cdot \text{Cvg}(k, \phi))^2 + 2(M \cdot \text{Cvg}(k, \phi))P\ell_{\phi,\theta}) \quad (52)$$

$$\leq \sigma_b^{-2} L^4 (\varepsilon + (M \cdot \text{Cvg}(k, \phi))^2 + 2(M \cdot \text{Cvg}(k, \phi))\sqrt{\varepsilon}) \quad (53)$$

$$= \sigma_b^{-2} L^4 (\sqrt{\varepsilon} + M \cdot \text{Cvg}(k, \phi))^2, \quad (54)$$

which completes the proof. $\square$

C Generalization Ability

In this section, we shall prove the following result, which is a refined version of Theorem 5.1.

Theorem C.1. *Assume the problem setting in Sec 2 and let $r > 0$. Then for any $t > 0$, with probability at least $1 - e^{-t}$, the empirical Rademacher complexity of $\ell_{\mathcal{F}}^{loc}(r)$ can be bounded by*

$$R_n \ell_{\mathcal{F}}^{loc}(r) \leq \sqrt{2} d n^{-\frac{1}{2}} \text{Stab}(k) \left(\sqrt{(\sqrt{r} + M \text{Cvg}(k))^2 C_1(n) + C_2(n, t, k, r) + 4} \right) + \text{Sens}(k) B_{\Phi},$$

where

$$C_1(n) = 216 \sigma_b^{-2} L^4 \log \mathcal{N}(n^{-\frac{1}{2}}, \ell_Q, L_2(P_n))$$

$$C_2(n, t, k, r) = \left(\frac{768 B_Q^2 t}{n} + 720 B_Q \mathbb{E} R_n \ell_Q^{loc}(r_q) \right) \log \mathcal{N}(n^{-\frac{1}{2}}, \ell_Q, L_2(P_n)),$$

$$r_q = \sigma_b^{-2} L^4 (\sqrt{r} + M \text{Cvg}(k))^2, \ell_Q^{loc}(r_q) = \{\|Q_{\theta} - Q^*\|_F : \theta \in \Theta, P\|Q_{\theta} - Q^*\|_F^2 \leq r_q\},$$

$$B_Q = 2L\sqrt{d}, \text{ and } B_{\Phi} = \frac{1}{2} \sup_{\phi_1, \phi_2 \in \Phi} \|\phi_1 - \phi_2\|_2.$$

Furthermore, for any $t > 0$, the expected Rademacher complexity of $\ell_{\mathcal{F}}^{loc}(r)$ can be bounded by

$$\mathbb{E} R_n \ell_{\mathcal{F}}^{loc}(r) \leq \sqrt{2} d n^{-\frac{1}{2}} \text{Stab}(k) \left(\sqrt{(\sqrt{r} + M \text{Cvg}(k))^2 \bar{C}_1(n) + \bar{C}_2(n, t) + \bar{C}_3(n, t) + 4} \right) + \text{Sens}(k) B_{\Phi},$$

where

$$\bar{C}_1(n) = 216 \sigma_b^{-2} L^4 \log \mathcal{N}_Q,$$

$$\bar{C}_2(n, t) = \left(1 + 3 B_Q e^{-t} \sqrt{\log \mathcal{N}_Q} + \frac{45}{\sqrt{n}} B_Q \log \mathcal{N}_Q \right) \frac{2880}{\sqrt{n}} B_Q \log \mathcal{N}_Q + t \frac{768 B_Q^2}{n} \log \mathcal{N}_Q,$$

$$\bar{C}_3(n, t) = 12 B_Q e^{-t} \sqrt{\log \mathcal{N}_Q} + \frac{360}{\sqrt{n}} B_Q \log \mathcal{N}_Q,$$

and $\mathcal{N}_Q = \mathcal{N}(n^{-\frac{1}{2}}, \ell_Q, L_{\infty})$.

In order to prove Theorem C.1, we first prove the following theorem, which reduces bounding the empirical Rademacher complexity of $\ell_{\mathcal{F}}^{loc}(r)$ to that of $\ell_Q^{loc}(r_q)$, and plays an important role in our complexity analysis.

Theorem C.2. *Assume the problem setting in Sec 2. Then it holds for any $r > 0$ that*

$$R_n \ell_{\mathcal{F}}^{loc}(r) \leq \sqrt{2} d \text{Stab}(k) R_n \ell_Q^{loc}(r_q) + \text{Sens}(k) B_{\Phi}, \quad (55)$$

with $r_q = \sigma_b^{-2} L^4 (\sqrt{r} + M \text{Cvg}(k))^2$, $\ell_Q^{loc}(r_q) = \{\|Q_{\theta} - Q^*\|_F : \theta \in \Theta, P\|Q_{\theta} - Q^*\|_F^2 \leq r_q\}$ and $B_{\Phi} = \frac{1}{2} \sup_{\phi_1, \phi_2 \in \Phi} \|\phi_1 - \phi_2\|_2$.

Proof. Let $k \in \mathbb{N}$ be fixed throughout this proof. We first show that the loss $\ell_{\phi, \theta}$ is $\text{Stab}(k)$ -Lipschitz in Q_{θ} and $\text{Sens}(k)$ -Lipschitz in ϕ . For any $(\mathbf{x}, \mathbf{b}) \in \mathcal{X} \times \mathcal{B}$, by using the triangle inequality and the definitions of $\text{Stab}(k, \phi')$ and $\text{Sens}(k)$, we can obtain the following estimate of the loss:

$$\begin{aligned} & |\ell_{\phi, \theta}(\mathbf{x}) - \ell_{\phi', \theta'}(\mathbf{x})| \\ &= \|\text{Alg}_{\phi}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 - \|\text{Alg}_{\phi'}^k(Q_{\theta'}(\mathbf{x}), \mathbf{b}) - \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})\|_2 \\ &\leq \|\text{Alg}_{\phi}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}) - \text{Alg}_{\phi'}^k(Q_{\theta'}(\mathbf{x}), \mathbf{b})\|_2 \\ &\leq \|\text{Alg}_{\phi'}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}) - \text{Alg}_{\phi'}^k(Q_{\theta'}(\mathbf{x}), \mathbf{b})\|_2 + \|\text{Alg}_{\phi}^k(Q_{\theta}(\mathbf{x}), \mathbf{b}) - \text{Alg}_{\phi'}^k(Q_{\theta}(\mathbf{x}), \mathbf{b})\|_2 \\ &\leq \text{Stab}(k, \phi') \|Q_{\theta}(\mathbf{x}) - Q_{\theta'}(\mathbf{x})\|_2 + \text{Sens}(k) \|\phi - \phi'\|_2 \\ &\leq \text{Stab}(k) \|Q_{\theta}(\mathbf{x}) - Q_{\theta'}(\mathbf{x})\|_2 + \text{Sens}(k) \|\phi - \phi'\|_2. \end{aligned} \quad (56)$$

where we write $Stab(k) = \sup_{\phi \in \Phi} Stab(k, \phi)$ for each $k \in \mathbb{N}$.

We then establish a vector contraction inequality, which is a modified version of Corollary 4 in [21] and Lemma 5 in [22]. Note that the empirical Rademacher complexity of $\ell_{\mathcal{F}}^{loc}$ can be written as:

$$R_n \ell_{\mathcal{F}}^{loc}(r) = \frac{1}{n} \mathbb{E}_{\sigma} \sup_{\phi, \theta} \sum_{i=1}^n \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) = \frac{1}{n} \mathbb{E}_{\sigma_{1:n-1}} \mathbb{E}_{\sigma_n} \sup_{\phi, \theta} \sum_{i=1}^{n-1} \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) + \sigma_n \ell_{\phi, \theta}(\mathbf{x}_n), \quad (57)$$

where the supremum is taken over the parameter space $\{(\phi, \theta) : \phi \in \Phi, \theta \in \Theta, P\ell_{\phi, \theta}^2 \leq r\}$.

Let $U_{n-1}(\phi, \theta) = \sum_{i=1}^{n-1} \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i)$ for each (ϕ, θ) . We now assume without loss of generality that the supremum can be attained and let

$$\phi_1, \theta_1 = \arg \sup_{\phi, \theta} (U_{n-1}(\phi, \theta) + \ell_{\phi, \theta}(\mathbf{x}_n)),$$

$$\phi_2, \theta_2 = \arg \sup_{\phi, \theta} (U_{n-1}(\phi, \theta) - \ell_{\phi, \theta}(\mathbf{x}_n)),$$

since otherwise we can consider (ϕ_1, θ_1) and (ϕ_2, θ_2) that are ϵ -close to the suprema for any $\epsilon > 0$ and conclude the same result. Then we can deduce from Eq. 56 that

$$\begin{aligned} & \mathbb{E}_{\sigma_n} \sup_{\phi, \theta} \sum_{i=1}^{n-1} \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) + \sigma_n \ell_{\phi, \theta}(\mathbf{x}_n) \\ &= \frac{1}{2} (U_{n-1}(\phi_1, \theta_1) + \ell_{\phi_1, \theta_1}(\mathbf{x}_n) + U_{n-1}(\phi_2, \theta_2) - \ell_{\phi_2, \theta_2}(\mathbf{x}_n)) \\ &= \frac{1}{2} (U_{n-1}(\phi_1, \theta_1) + U_{n-1}(\phi_2, \theta_2) + (\ell_{\phi_1, \theta_1}(\mathbf{x}_n) - \ell_{\phi_2, \theta_2}(\mathbf{x}_n))) \\ &\leq \frac{1}{2} (U_{n-1}(\phi_1, \theta_1) + U_{n-1}(\phi_2, \theta_2)) + \frac{1}{2} (Stab(k) \|Q_{\theta_1}(\mathbf{x}_n) - Q_{\theta_2}(\mathbf{x}_n)\|_2 + Sens(k) \|\phi_1 - \phi_2\|_2) \\ &\leq \frac{1}{2} (U_{n-1}(\phi_1, \theta_1) + U_{n-1}(\phi_2, \theta_2)) + \frac{1}{2} Stab(k) \|Q_{\theta_1}(\mathbf{x}_n) - Q_{\theta_2}(\mathbf{x}_n)\|_F + Sens(k) B_{\Phi}, \end{aligned}$$

where $B_{\Phi} = \frac{1}{2} \sup_{\phi_1, \phi_2 \in \Phi} \|\phi_1 - \phi_2\|_2$.

For each $\mathbf{x} \in \mathcal{X}$, $\theta \in \Theta$ and $1 \leq j, k \leq d$, let $Q_{\theta}^{j,k}(\mathbf{x})$ be the j, k -th entry of the matrix $Q_{\theta}(\mathbf{x})$. The the Khintchine-Kahane inequality (see e.g. [21]) gives us that

$$\mathbb{E}_{\sigma_n} \sup_{\phi, \theta} \sum_{i=1}^n \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) \leq \frac{1}{2} (U_{n-1}(\phi_1, \theta_1) + U_{n-1}(\phi_2, \theta_2)) + Sens(k) B_{\Phi} \quad (58)$$

$$+ \frac{1}{2} Stab(k) \sqrt{2} \mathbb{E}_{\epsilon_n} \left| \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta_1}^{j,k}(\mathbf{x}_n) - Q_{\theta_2}^{j,k}(\mathbf{x}_n) \right) \right|, \quad (59)$$

where $\epsilon_n = (\epsilon_n^{j,k})_{j,k=1}^n$ are independent Rademacher variables. Hence, if we denote by $s(\epsilon_n)$ the sign of $\sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta_1}^{j,k}(\mathbf{x}_n) - Q_{\theta_2}^{j,k}(\mathbf{x}_n) \right)$ and by $Q^{*j,k}(\mathbf{x})$ be the j, k -th entry of the matrix $Q^*(\mathbf{x})$, then we can obtain that

$$\begin{aligned} & \mathbb{E}_{\sigma_n} \sup_{\phi, \theta} \sum_{i=1}^n \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) \\ &\leq \mathbb{E}_{\epsilon_n} \frac{1}{2} \left[\left(U_{n-1}(\phi_1, \theta_1) + Stab(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} Q_{\theta_1}^{j,k}(\mathbf{x}_n) \right) \right. \\ &\quad \left. + \left(U_{n-1}(\phi_2, \theta_2) - Stab(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} Q_{\theta_2}^{j,k}(\mathbf{x}_n) \right) \right] + Sens(k) B_{\Phi} \\ &= \mathbb{E}_{\epsilon_n} \frac{1}{2} \left[\left(U_{n-1}(\phi_1, \theta_1) + Stab(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta_1}^{j,k}(\mathbf{x}_n) - Q^{*j,k}(\mathbf{x}_n) \right) \right) \right. \\ &\quad \left. + \left(U_{n-1}(\phi_2, \theta_2) - Stab(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta_2}^{j,k}(\mathbf{x}_n) - Q^{*j,k}(\mathbf{x}_n) \right) \right) \right] + Sens(k) B_{\Phi}. \end{aligned}$$

Then by taking the supremum over (ϕ, θ) and using the fact that σ_n is an independent Rademacher variable, we can deduce that

$$\begin{aligned}
& \mathbb{E}_{\sigma_n} \sup_{\phi, \theta} \sum_{i=1}^n \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) \\
& \leq \mathbb{E}_{\epsilon_n} \frac{1}{2} \left[\sup_{\phi, \theta} \left(U_{n-1}(\phi, \theta) + \text{Stab}(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right) \right. \\
& \quad \left. + \sup_{\phi, \theta} \left(U_{n-1}(\phi, \theta) - \text{Stab}(k) \sqrt{2} s(\epsilon_n) \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right) \right] + \text{Sens}(k) B_{\Phi} \\
& = \mathbb{E}_{\epsilon_n} \mathbb{E}_{\sigma_n} \left[\sup_{\phi, \theta} \left(U_{n-1}(\phi, \theta) + \text{Stab}(k) \sqrt{2} \sigma_n \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right) \right] + \text{Sens}(k) B_{\Phi} \\
& = \mathbb{E}_{\epsilon_n} \left[\sup_{\phi, \theta} \left(U_{n-1}(\phi, \theta) + \text{Stab}(k) \sqrt{2} \sum_{j,k} \epsilon_n^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right) \right] + \text{Sens}(k) B_{\Phi},
\end{aligned}$$

where we have used the fact that $\sum_{j,k} \epsilon_n^{j,k} (Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n))$ is a symmetric random variable in the last line.

By proceeding in the same way for all other $\sigma_{n-1}, \dots, \sigma_1$, we can obtain the following vector-contraction inequality:

$$\begin{aligned}
\mathbb{E}_{\sigma} \sup_{\phi, \theta} \sum_{i=1}^n \sigma_i \ell_{\phi, \theta}(\mathbf{x}_i) & \leq \sqrt{2} \text{Stab}(k) \mathbb{E}_{\epsilon_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sum_{j,k} \epsilon_i^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right] \\
& \quad + n \text{Sens}(k) B_{\Phi}.
\end{aligned} \tag{60}$$

The first term on the right-hand side can be bounded by using the Cauchy-Schwarz inequality as follows:

$$\begin{aligned}
& \mathbb{E}_{\epsilon_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sum_{j,k} \epsilon_i^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right] \\
& = \mathbb{E}_{\sigma_{1:n}} \mathbb{E}_{\epsilon_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sigma_i \sum_{j,k} \epsilon_i^{j,k} \left(Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n) \right) \right] \\
& \leq \mathbb{E}_{\sigma_{1:n}} \mathbb{E}_{\epsilon_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sigma_i \sqrt{\sum_{j,k} (\epsilon_i^{j,k})^2} \sqrt{\sum_{j,k} |Q_{\theta}^{j,k}(\mathbf{x}_n) - Q^{*,j,k}(\mathbf{x}_n)|^2} \right] \\
& = \mathbb{E}_{\sigma_{1:n}} \mathbb{E}_{\epsilon_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sigma_i d \|Q_{\theta}(\mathbf{x}_n) - Q^*(\mathbf{x}_n)\|_F \right] \\
& = d \mathbb{E}_{\sigma_{1:n}} \left[\sup_{\theta} \sum_{i=1}^n \sigma_i \|Q_{\theta}(\mathbf{x}_n) - Q^*(\mathbf{x}_n)\|_F \right].
\end{aligned} \tag{61}$$

Therefore, bounding the Rademacher complexity of $\ell_{\mathcal{F}}^{\text{loc}}(r)$ reduces to bounding the Rademacher complexity of the space of functions $\|Q_{\theta} - Q^*\|_F$. Recall that the supremum is taken over the parameter space where $(\phi, \theta) \in \Phi \times \Theta$ satisfies $P \ell_{\phi, \theta}^2 \leq r$. Note that Lemma 4.2 implies that,

$$P \|Q_{\theta} - Q^*\|_F^2 \leq r_q := \sigma_b^{-2} L^4 (\sqrt{\varepsilon} + M \cdot \text{Cvg}(k, \phi))^2. \tag{62}$$

Hence, by defining the following function space:

$$\ell_{\mathcal{Q}}^{\text{loc}}(r_q) := \{ \|Q_{\theta} - Q^*\|_F : \theta \in \Theta, P \|Q_{\theta} - Q^*\|_F^2 \leq r_q \}, \tag{63}$$

we can conclude the desired relationship between $R_n \ell_{\mathcal{F}}^{\text{loc}}(r)$ and $R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$ from the inequalities Eq. 60 and Eq. 61.

□

With Theorem C.2 in hand, we see that, for each $r > 0$, in order to obtain the upper bounds of $R_n \ell_{\mathcal{F}}^{\text{loc}}(r)$ in Theorem C.1, it suffices to estimate $R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$, i.e., the Rademacher complexity of the function space $\ell_{\mathcal{Q}}^{\text{loc}}(r_q)$.

The following theorem summarizes the estimates for the empirical and expected Rademacher complexity of the local class $\ell_{\mathcal{Q}}^{\text{loc}}$, which will be established in Propositions C.1 and C.2, respectively.

Recall that, for any given $\epsilon > 0$, a class of functions $\mathcal{F}$ and pseudometric $\|\cdot\|$, the covering number $\mathcal{N}(\epsilon, \mathcal{F}, \|\cdot\|)$ is defined as the cardinality of the smallest subset $\hat{\mathcal{F}}$ of $\mathcal{F}$ for which every element of $\mathcal{F}$ is within the ϵ -neighbourhood of some element of $\hat{\mathcal{F}}$ with respect to the pseudometric $\|\cdot\|$.

Theorem C.3. *Assume the problem setting in Sec 2. Let $r > 0$, $r_q = \sigma_b^{-2} L^4 (\sqrt{r} + MCvg(k))^2$ and $\ell_{\mathcal{Q}}^{\text{loc}}(r_q) = \{\|Q_\theta - Q^*\|_F : \theta \in \Theta, P\|Q_\theta - Q^*\|_F^2 \leq r_q\}$. Then for all $t > 0$, we have with probability at least $1 - e^{-t}$ that*

$$R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq n^{-\frac{1}{2}} \left[\left(C_1(n)(\sqrt{r} + MCvg(k))^2 + C_2(n, t, k, r) \right)^{\frac{1}{2}} + 4 \right], \quad (64)$$

where

$$\begin{aligned} C_1(n) &= 216 \sigma_b^{-2} L^4 \log \mathcal{N}(n^{-\frac{1}{2}}, \ell_{\mathcal{Q}}, L_2(P_n)), \\ C_2(n, t, k, r) &= \left(\frac{768 B_Q^2 t}{n} + 720 B_Q \mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \right) \log \mathcal{N}(n^{-\frac{1}{2}}, \ell_{\mathcal{Q}}, L_2(P_n)), \end{aligned}$$

and $B_Q = 2L\sqrt{d}$.

Moreover, for all $t > 0$, we have that

$$\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq n^{-\frac{1}{2}} \left[\left(\bar{C}_1(n)(\sqrt{r} + MCvg(k))^2 + \bar{C}_2(n, t) \right)^{\frac{1}{2}} + \bar{C}_3(n, t) + 4 \right], \quad (65)$$

where

$$\begin{aligned} \bar{C}_1(n) &= 216 \sigma_b^{-2} L^4 \log \mathcal{N}_Q, \\ \bar{C}_2(n, t) &= \left(1 + 3B_Q e^{-t} \sqrt{\log \mathcal{N}_Q} + \frac{45}{\sqrt{n}} B_Q \log \mathcal{N}_Q \right) \frac{2880}{\sqrt{n}} B_Q \log \mathcal{N}_Q + t \frac{768 B_Q^2}{n} \log \mathcal{N}_Q, \\ \bar{C}_3(n, t) &= 12 B_Q e^{-t} \sqrt{\log \mathcal{N}_Q} + \frac{360}{\sqrt{n}} B_Q \log \mathcal{N}_Q \end{aligned}$$

and $\mathcal{N}_Q = \mathcal{N}(n^{-\frac{1}{2}}, \ell_{\mathcal{Q}}, L_\infty)$.

We first establish the estimate for the empirical Rademacher complexity $R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$, i.e., Eq. 64 in Theorem C.3.

Proposition C.1. *Assume the problem setting in Sec 2. Let $B_Q = \sup_{(\theta, \mathbf{x}) \in \Theta \times \mathcal{X}} \|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F$, and for each $r > 0$ let r_q and $\ell_{\mathcal{Q}}^{\text{loc}}(r_q)$ be defined as in Theorem C.2. Then we have that*

$$R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq \frac{4}{\sqrt{n}} \left(1 + 3B_Q \sqrt{\log \mathcal{N}\left(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)\right)} \right). \quad (66)$$

Moreover, for all $t > 0$, it holds with probability at least $1 - e^{-t}$ that

$$R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq \frac{4}{\sqrt{n}} \left(1 + 3C(r_q, t) \sqrt{\log \mathcal{N}\left(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)\right)} \right), \quad (67)$$

with the constant $C(r_q, t) = \left(\frac{3r_q}{2} + \frac{16B_Q^2 t}{3n} + 5B_Q \mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \right)^{1/2}$.

Proof. The classical Dudley's entropy integral bound for the empirical Rademacher complexity gives us that

$$R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq \inf_{\alpha > 0} \left(4\alpha + \frac{12}{\sqrt{n}} \int_{\alpha}^{\infty} \sqrt{\log \mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n))} d\epsilon \right). \quad (68)$$

Observe that all functions in $\ell_{\mathcal{Q}}^{\text{loc}}(r_q)$ take value in $[0, B_Q]$, which implies for all $\epsilon \geq B_Q$ that, $\mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)) \leq \mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_{\infty}(P_n)) = 1$ and consequently the integrand in Eq. 68 vanishes on $[B_Q, \infty)$. Hence we have that

$$\begin{aligned} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) &\leq \inf_{\alpha > 0} \left(4\alpha + \frac{12}{\sqrt{n}} \int_{\alpha}^{B_Q} \sqrt{\log \mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n))} d\epsilon \right) \\ &\leq \frac{4}{\sqrt{n}} + \frac{12}{\sqrt{n}} \int_{\frac{1}{\sqrt{n}}}^{B_Q} \sqrt{\log \mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n))} d\epsilon \\ &\leq \frac{4}{\sqrt{n}} + \frac{12}{\sqrt{n}} B_Q \sqrt{\log \mathcal{N}\left(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)\right)}, \end{aligned}$$

where we used the fact that $\mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n))$ is decreasing in terms of ϵ for the last inequality. This proves the estimate Eq. 66.

In order to establish the estimate Eq. 67, we shall bound the empirical error $P_n \|Q_{\theta} - Q^*\|_F^2$ with high probability. Let us consider the class of functions $\ell_{\mathcal{Q}^2}^{\text{loc}}(r_q) = \{\|Q_{\theta} - Q^*\|_F^2 : \theta \in \Theta, P\|Q_{\theta} - Q^*\|_F^2 \leq r_q\}$, whose element takes values in $[0, B_Q^2]$. Moreover, we see it holds for all $\|Q_{\theta} - Q^*\|_F^2 \in \ell_{\mathcal{Q}^2}^{\text{loc}}(r_q)$ that $P\|Q_{\theta} - Q^*\|_F^4 \leq B_Q^2 P\|Q_{\theta} - Q^*\|_F^2 \leq B_Q^2 r_q$. Hence, by applying Theorem 2.1 in [16] (with $\mathcal{F} = \ell_{\mathcal{Q}^2}^{\text{loc}}(r_q)$, $a = 0$, $b = B_Q^2$, $\alpha = 1/4$ and $r = B_Q^2 r_q$) and the Cauchy-Schwarz inequality, we can deduce that, for each $t > 0$, it holds with probability at least $1 - e^{-t}$ that

$$\begin{aligned} P_n \|Q_{\theta} - Q^*\|_F^2 &\leq P\|Q_{\theta} - Q^*\|_F^2 + \frac{5}{2} \mathbb{E} R_n \ell_{\mathcal{Q}^2}^{\text{loc}}(r_q) + \sqrt{\frac{2B_Q^2 r_q t}{n}} + B_Q^2 \frac{13t}{3n} \\ &\leq r_q + \frac{5}{2} \mathbb{E} R_n \ell_{\mathcal{Q}^2}^{\text{loc}}(r_q) + \frac{r_q}{2} + \frac{B_Q^2 t}{n} + B_Q^2 \frac{13t}{3n} \\ &\leq \frac{3r_q}{2} + 5B_Q \mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) + \frac{16B_Q^2 t}{3n}. \end{aligned}$$

Consequently, we see it holds with probability at least $1 - e^{-t}$ that, $\mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)) = 1$ for all $\epsilon \geq C(r_q, t)$, with the constant $C(r_q, t)$ defined as in the statement of Proposition C.1. Substituting this fact into the integral bound Eq. 68 and following the same argument as above, we can conclude Eq. 67 with probability at least $1 - e^{-t}$. $\square$

Now we proceed to prove the estimate of the expected Rademacher complexity $\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$, i.e., Eq. 65 in Theorem C.3.

Proposition C.2. *Assume the same setting as in Proposition C.1. Then it holds for any $r, t > 0$ that*

$$\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq n^{-\frac{1}{2}} \left[\left(C_1(n, t)(\sqrt{r} + MCv g(k))^2 + C_2(n, t) \right)^{\frac{1}{2}} + C_3(n, t) + 4 \right], \quad (69)$$

where $C_1(n, t)$, $C_2(n, t)$ and $C_3(n, t)$ the constants defined as in Eq. 72, Eq. 73 and Eq. 74, respectively.

Proof. Let $r, t > 0$ be fixed throughout this proof. Since it holds for all $\epsilon > 0$ and $n \in \mathbb{N}$ that $\mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_2(P_n)) \leq \mathcal{N}(\epsilon, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_{\infty})$, we can deduce from Proposition C.1 that

$$\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q) \leq \frac{4}{\sqrt{n}} \left(1 + 3[C(r_q, t)(1 - e^{-t}) + B_Q e^{-t}] \sqrt{\log \mathcal{N}\left(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_{\infty}\right)} \right), \quad (70)$$

with the constants B_Q and $C(r_q, t)$ defined as in the statement of Proposition C.1.

The above estimate gives an implicit upper bound of $\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$ since $C(r_q, t)$ also involves $\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$. Now we shall introduce the notation $\mathcal{N}_{\mathcal{Q}}^n = \mathcal{N}(\frac{1}{\sqrt{n}}, \ell_{\mathcal{Q}}^{\text{loc}}(r_q), L_{\infty})$ and derive an explicit upper bound of $\mathbb{E} R_n \ell_{\mathcal{Q}}^{\text{loc}}(r_q)$. By rearranging the terms in Eq. 70 and using the definition of $C(r_q, t)$,

we can obtain that

$$\begin{aligned} & \frac{\sqrt{n}}{4} \mathbb{E} R_n \ell_Q^{loc}(r_q) - 1 - 3B_Q e^{-t} \sqrt{\log \mathcal{N}_Q^n} \\ & \leq 3(1 - e^{-t}) \sqrt{\left(\frac{3r_q}{2} + \frac{16B_Q^2 t}{3n} + 5B_Q \mathbb{E} R_n \ell_Q^{loc}(r_q) \right) \log \mathcal{N}_Q^n}. \end{aligned} \quad (71)$$

We shall assume without loss of generality that $\mathbb{E} R_n \ell_Q^{loc}(r_q) \geq \frac{4}{\sqrt{n}} \left(1 + 3B_Q e^{-t} \sqrt{\log \mathcal{N}_Q^n} \right)$, since otherwise we have a trivial estimate that $\mathbb{E} R_n \ell_Q^{loc}(r_q) \leq 4n^{-\frac{1}{2}} A_1$, with $A_1 = 1 + 3B_Q e^{-t} \sqrt{\log \mathcal{N}_Q^n}$. Then by squaring both sides of Eq. 71 and rearranging the terms, we get that

$$\begin{aligned} & \frac{n}{16} (\mathbb{E} R_n \ell_Q^{loc}(r_q))^2 - \left(\frac{\sqrt{n}}{2} A_1 + 45(1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n \right) \mathbb{E} R_n \ell_Q^{loc}(r_q) \\ & + A_1^2 - 9(1 - e^{-t})^2 A_2 \log \mathcal{N}_Q^n \leq 0, \end{aligned}$$

with the constant $A_2 = \frac{3r_q}{2} + \frac{16B_Q^2 t}{3n}$. This implies that

$$\begin{aligned} \mathbb{E} R_n \ell_Q^{loc}(r_q) & \leq \frac{8}{n} \left[\frac{\sqrt{n} A_1}{2} + 45(1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n + \left(\left[\frac{\sqrt{n} A_1}{2} + 45(1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n \right]^2 \right. \right. \\ & \quad \left. \left. - \frac{n}{4} [A_1^2 - 9(1 - e^{-t})^2 A_2 \log \mathcal{N}_Q^n] \right)^{\frac{1}{2}} \right] \\ & = n^{-\frac{1}{2}} \left[4A_1 + \frac{360}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n + \left([4A_1 + \frac{360}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n]^2 \right. \right. \\ & \quad \left. \left. - 16[A_1^2 - 9(1 - e^{-t})^2 A_2 \log \mathcal{N}_Q^n] \right)^{\frac{1}{2}} \right]. \end{aligned}$$

Hence, for each $t > 0$, by introducing the following constants

$$C_1(n, t) = 216(1 - e^{-t})^2 \sigma_b^{-2} L^4 \log \mathcal{N}_Q^n, \quad (72)$$

$$\begin{aligned} C_2(n, t) & = [4A_1 + \frac{360}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n]^2 - 16A_1^2 + t(1 - e^{-t})^2 \frac{768B_Q^2}{n} \log \mathcal{N}_Q^n \\ & = \left(1 + 3B_Q e^{-t} \sqrt{\log \mathcal{N}_Q^n} + \frac{45}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n \right) \frac{2880}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n \\ & \quad + t(1 - e^{-t})^2 \frac{768B_Q^2}{n} \log \mathcal{N}_Q^n, \end{aligned} \quad (73)$$

$$C_3(n, t) = 12B_Q e^{-t} \sqrt{\log \mathcal{N}_Q^n} + \frac{360}{\sqrt{n}} (1 - e^{-t})^2 B_Q \log \mathcal{N}_Q^n, \quad (74)$$

with $B_Q = \sup_{(\theta, \mathbf{x}) \in \Theta \times \mathcal{X}} \|Q_\theta(\mathbf{x}) - Q^*(\mathbf{x})\|_F \leq 2\sqrt{d}L$ and $\mathcal{N}_Q^n = \mathcal{N}(\frac{1}{\sqrt{n}}, \ell_Q, L_\infty)$, we can deduce that

$$\mathbb{E} R_n \ell_Q^{loc}(r_q) \leq n^{-\frac{1}{2}} \left[\left(C_1(n, t)(\sqrt{r} + MCvg(k))^2 + C_2(n, t) \right)^{\frac{1}{2}} + C_3(n, t) + 4 \right].$$

□

D RNN as a Neural Algorithm

We denote by RNN_ϕ^k a recurrent neural network that has k unrolled RNN cells and view it as a neural algorithm. It has been proposed in [19] to use RNN to learn an optimization algorithm where the update steps in each iteration are given by the operations in an RNN cell

$$\mathbf{y}_{k+1} \leftarrow \text{RNNcell}(Q, \mathbf{b}, \mathbf{y}_k) := V\sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_k + W_2^1\mathbf{g}_k))). \quad (75)$$

In the above equation, we take a specific example where the RNNcell is a multi-layer perception (MLP) with activations $\sigma = \text{RELU}$ that takes the current iterate $\mathbf{y}_k$ and the gradient $\mathbf{g}_k = Q\mathbf{y}_k + \mathbf{b}$ as inputs.

(I) Stable Region. First, we show that when the parameters satisfy $c_\phi := \sup_Q \|V\|_2 \|W_1^1\| + W_2^1 Q\|_2 \prod_{l=2}^L \|W^l\|_2 < 1$, the operations in RNNcell are strictly contractive, i.e., $\|\mathbf{y}_{k+1} - \mathbf{y}_k\|_2 \leq c_\phi \|\mathbf{y}_k - \mathbf{y}_{k-1}\|_2$.

Proof. By definition,

$$\begin{aligned} \|\mathbf{y}_{k+1} - \mathbf{y}_k\|_2 &= \|V\sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_k + W_2^1\mathbf{g}_k))) \\ &\quad - V\sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_{k-1} + W_2^1\mathbf{g}_{k-1})))\|_2 \\ &\leq \|V\|_2 \|\sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_k + W_2^1\mathbf{g}_k))) \\ &\quad - \sigma(W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_{k-1} + W_2^1\mathbf{g}_{k-1})))\|_2 \end{aligned}$$

Since the activation function $\sigma = \text{RELU}$ satisfies the inequality that $\|\sigma(\mathbf{x}) - \sigma(\mathbf{x}')\|_2 \leq \|\mathbf{x} - \mathbf{x}'\|_2$ for any $\mathbf{x}, \mathbf{x}'$, we have

$$\begin{aligned} \|\mathbf{y}_{k+1} - \mathbf{y}_k\|_2 &\leq \|V\|_2 \|W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_k + W_2^1\mathbf{g}_k)) \\ &\quad - W^L\sigma(W^{L-1}\dots W^2\sigma(W_1^1\mathbf{y}_{k-1} + W_2^1\mathbf{g}_{k-1}))\|_2. \end{aligned}$$

Similarly, we can obtain

$$\begin{aligned} \|\mathbf{y}_{k+1} - \mathbf{y}_k\|_2 &\leq \|V\|_2 \|W^L\|_2 \dots \|W^2\|_2 \|W_1^1\mathbf{y}_k + W_2^1\mathbf{g}_k - (W_1^1\mathbf{y}_{k-1} + W_2^1\mathbf{g}_{k-1})\|_2 \\ &= \|V\|_2 \|W^L\|_2 \dots \|W^2\|_2 \|(W_1^1 + QW_2^1)(\mathbf{y}_k - \mathbf{y}_{k-1})\|_2 \\ &\leq \|V\|_2 \|W^L\|_2 \dots \|W^2\|_2 \|W_1^1 + QW_2^1\|_2 \|\mathbf{y}_k - \mathbf{y}_{k-1}\|_2 \\ &\leq c_\phi \|\mathbf{y}_k - \mathbf{y}_{k-1}\|_2. \end{aligned}$$

Therefore, if $c_\phi < 1$, then the operation is strictly contractive. $\square$

(II) Stability. We shall show the neural algorithm RNN_ϕ^k has a stability constant $\text{Stab}(k, \phi) = \mathcal{O}(1 - c_\phi^k)$ (see the definition of stability in Sec 3).

Proof. Let us consider two quadratic problems induced by $(Q, \mathbf{b})$ and $(Q', \mathbf{b}')$, and denote the corresponding outputs of RNN_ϕ^k as $\mathbf{y}_k = \text{RNN}_\phi^k(Q, \mathbf{b})$ and $\mathbf{y}'_k = \text{RNN}_\phi^k(Q', \mathbf{b}')$.

Denote $c_\phi^Q = \|V\|_2 \|W_1^1\| + W_2^1 Q\|_2 \prod_{l=2}^L \|W^l\|_2$, $c_\phi^{Q'} = \|V\|_2 \|W_1^1\| + W_2^1 Q'\|_2 \prod_{l=2}^L \|W^l\|_2$, and $\hat{c}_\phi := \|V\|_2 \|W_2^1\| \prod_{l=2}^L \|W^l\|_2$. First, we see that

$$\begin{aligned} \|\mathbf{y}_k\|_2 &\leq c_\phi^Q \|\mathbf{y}_{k-1}\|_2 + \hat{c}_\phi \|\mathbf{b}\|_2 \leq (c_\phi^Q)^k \|\mathbf{y}_0\|_2 + \hat{c}_\phi \|\mathbf{b}\|_2 \sum_{i=1}^k (c_\phi^Q)^{i-1} \\ &= \frac{\hat{c}_\phi \|\mathbf{b}\|_2 (1 - (c_\phi^Q)^k)}{1 - c_\phi^Q} \leq \frac{\hat{c}_\phi \|\mathbf{b}\|_2}{1 - c_\phi^Q}. \end{aligned} \quad (76)$$

Similar conclusion holds for $\mathbf{y}'_k$. Then, by following a similar argument as that for the proof of the stable region, we can deduce from $\mathbf{y}_0 = \mathbf{y}'_0$ that

$$\begin{aligned}
& \|\mathbf{y}_k - \mathbf{y}'_k\|_2 \\
& \leq \|V\|_2 \|W^L\|_2 \cdots \|W^2\|_2 \|(W_1^1 + W_1^2 Q) \mathbf{y}_{k-1} - (W_1^1 + W_1^2 Q') \mathbf{y}'_{k-1} + W_1^2 (\mathbf{b} - \mathbf{b}')\|_2 \\
& \leq \|V\|_2 \|W^L\|_2 \cdots \|W^2\|_2 (\|W_1^1 + W_1^2 Q\|_2 \|\mathbf{y}_{k-1} - \mathbf{y}'_{k-1}\|_2 + \|Q - Q'\|_2 \|W_1^2\|_2 \|\mathbf{y}'_{k-1}\|_2 \\
& \quad + \|W_1^2\|_2 \|\mathbf{b} - \mathbf{b}'\|_2) \\
& \leq c_\phi^Q \|\mathbf{y}_{k-1} - \mathbf{y}'_{k-1}\|_2 + \hat{c}_\phi \|Q - Q'\|_2 \frac{\hat{c}_\phi \|\mathbf{b}'\|_2}{1 - c_\phi^{Q'}} + \hat{c}_\phi \|\mathbf{b} - \mathbf{b}'\|_2 \\
& \leq (c_\phi^Q)^k \|\mathbf{y}_0 - \mathbf{y}'_0\|_2 + \left(\frac{\hat{c}_\phi^2 \|\mathbf{b}'\|_2}{1 - c_\phi^{Q'}} \|Q - Q'\|_2 + \hat{c}_\phi \|\mathbf{b} - \mathbf{b}'\|_2 \right) \sum_{i=1}^k (c_\phi^Q)^{i-1} \\
& = \frac{\hat{c}_\phi^2 \|\mathbf{b}'\|_2}{1 - c_\phi^{Q'}} \frac{1 - (c_\phi^Q)^k}{1 - c_\phi^Q} \|Q - Q'\|_2 + \hat{c}_\phi \frac{1 - (c_\phi^Q)^k}{1 - c_\phi^Q} \|\mathbf{b} - \mathbf{b}'\|_2.
\end{aligned}$$

Therefore, the stability constant is of the magnitude $\mathcal{O}(1 - c_\phi^k)$. $\square$

(III) Sensitivity. We now proceed to analyze the sensitivity of the neural algorithm RNN_ϕ^k as defined in Sec 3. Note that the strong non-linearity in the RNN cell and the high-dimensionality of the parameter space significantly complicate the analysis of the Lipschitz dependence of RNN_ϕ^k with respect to its parameter $\phi = \{W_1^1, W_1^2, \dots, W^L, V\}$. To simplify our presentation, we shall assume the parameter ϕ are constrained in a compact subset Φ of the stable region, and show the neural algorithm RNN_ϕ^k has a sensitivity $\text{Sens}(k) = \mathcal{O}(1 - (\inf_{\phi \in \Phi} c_\phi)^k)$. A rigorous sensitivity analysis of RNN with general weights is out of the scope of this paper.

Proof. Let the range of parameters Φ is a compact subset of the stable region, such that for all $\phi \in \Phi$, $c_\phi := \sup_Q \|V\|_2 \|W_1^1 + W_1^2 Q\|_2 \prod_{l=2}^L \|W^l\|_2 \leq c_0 < 1$ for some constant c_0 . Let $\phi, \phi' \in \Phi$ be two given sets of parameters. For each $k \in \mathbb{N}$, we denote $\mathbf{y}_k = \text{RNN}_\phi^k(Q, \mathbf{b})$ and $\mathbf{y}'_k = \text{RNN}_{\phi'}^k(Q, \mathbf{b})$ the outputs corresponding to the parameters ϕ and ϕ' , respectively. Then we have that

$$\begin{aligned}
\|\mathbf{y}_k - \mathbf{y}'_k\|_2 &= \|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}_{k-1}) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}'_{k-1})\|_2 \\
&\leq \|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}'_{k-1}) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}'_{k-1})\|_2 \\
&\quad + \|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}_{k-1}) - \text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}'_{k-1})\|_2 \\
&\leq \|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}'_{k-1}) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}'_{k-1})\|_2 + c_\phi \|\mathbf{y}_{k-1} - \mathbf{y}'_{k-1}\|_2
\end{aligned}$$

If there exists a constant K , independent of k, ϕ, ϕ' , such that

$$\|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}'_{k-1}) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}'_{k-1})\|_2 \leq K \|\phi - \phi'\|_2, \quad (77)$$

then we can obtain from $\mathbf{y}_0 = \mathbf{y}'_0$ that

$$\begin{aligned}
\|\mathbf{y}_k - \mathbf{y}'_k\|_2 &\leq v \|\phi - \phi'\|_2 + c_\phi \|\mathbf{y}_{k-1} - \mathbf{y}'_{k-1}\|_2 \\
&\leq K \|\phi - \phi'\|_2 \sum_{i=1}^k c_\phi^{i-1} = \frac{1 - c_\phi^k}{1 - c_\phi} K \|\phi - \phi'\|_2.
\end{aligned}$$

The fact that $c_\phi \leq c_0 < 1$ for some constant c_0 implies that the magnitude of sensitivity is $\mathcal{O}(1 - (\inf_{\phi \in \Phi} c_\phi)^k)$.

Now it remains to establish the estimate Eq. 77. For each $k \in \mathbb{N}$, $\phi = \{W_1^1, W_1^2, \dots, W^L, V\}$ and $l = 2, \dots, L$, we introduce the notation

$$f_\phi^l := W^l \sigma(W^{l-1} \dots W^2 \sigma(W_1^1 \mathbf{y}_k + W_2^1 \mathbf{g}_k)), \quad (78)$$

with $f_\phi^1 = W_1^1 \mathbf{y}_k + W_2^1 \mathbf{g}_k$. Then we have for each $l = 1, \dots, L$ that

$$\|f_\phi^l\|_2 \leq \prod_{j=2}^l \|W^j\|_2 (\|W_1^1 + W_2^1 Q\|_2 \|\mathbf{y}_k\|_2 + \|W_2^1\|_2 \|\mathbf{b}\|_2) = c_l \|\mathbf{y}_k\|_2 + \hat{c}_l \|\mathbf{b}\|_2, \quad (79)$$

with the constants $c_l := (\prod_{j=2}^l \|W^j\|_2) \|W_1^1 + W_2^1 Q\|_2$, $\hat{c}_l := (\prod_{j=2}^l \|W^j\|_2) \|W_2^1\|_2$ for all $l = 1, \dots, L$. Then by induction, we can see that

$$\begin{aligned}
\|f_\phi^L - f_{\phi'}^L\|_2 &= \|W^L \sigma(f_\phi^{L-1}) - W'^L \sigma(f_{\phi'}^{L-1})\|_2 \\
&\leq \|W^L - W'^L\|_2 \|f_{\phi'}^{L-1}\|_2 + \|W^L\|_2 \|f_\phi^{L-1} - f_{\phi'}^{L-1}\|_2 \\
&\leq \|W^L - W'^L\|_2 \|f_{\phi'}^{L-1}\|_2 + \|W^L\|_2 \left(\|W^{L-1} - W'^{L-1}\|_2 \|f_{\phi'}^{L-2}\|_2 \right. \\
&\quad \left. + \|W^{L-1}\|_2 \|f_\phi^{L-2} - f_{\phi'}^{L-2}\|_2 \right) \\
&\leq \sum_{l=2}^L \left(\prod_{j=l+1}^L \|W^j\|_2 \right) \|W^l - W'^l\|_2 \|f_{\phi'}^{l-1}\|_2 + \left(\prod_{l=2}^L \|W^l\|_2 \right) \|f_\phi^1 - f_{\phi'}^1\|_2.
\end{aligned}$$

Thus we have that

$$\begin{aligned}
\|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}_k) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}_k)\|_2 &= \|V \sigma(f_\phi^L) - V' \sigma(f_{\phi'}^L)\|_2 \\
&\leq \|V - V'\|_2 \|f_{\phi'}^L\|_2 + \|V\|_2 \|f_\phi^L - f_{\phi'}^L\|_2 \\
&\leq \|V - V'\|_2 \|f_{\phi'}^L\|_2 + \|V\|_2 \left[\sum_{l=2}^L \left(\prod_{j=l+1}^L \|W^j\|_2 \right) \|W^l - W'^l\|_2 \|f_{\phi'}^{l-1}\|_2 \right. \\
&\quad \left. + \left(\prod_{l=2}^L \|W^l\|_2 \right) \|f_\phi^1 - f_{\phi'}^1\|_2 \right].
\end{aligned}$$

Furthermore, we see that

$$\begin{aligned}
\|f_\phi^1 - f_{\phi'}^1\|_2 &= \|(W_1^1 + W_2^1 Q) \mathbf{y}_k + W_2^1 \mathbf{b} - (W_1'^1 + W_2'^1 Q) \mathbf{y}_k + W_2'^1 \mathbf{b}\|_2 \\
&\leq \|W_1^1 - W_1'^1 + (W_2^1 - W_2'^1) Q\|_2 \|\mathbf{y}_k\|_2 + \|W_2^1 - W_2'^1\|_2 \|\mathbf{b}\|_2 \\
&\leq \|W_1^1 - W_1'^1\| \|\mathbf{y}_k\|_2 + \|W_2^1 - W_2'^1\| (\|Q\|_2 \|\mathbf{y}_k\|_2 + \|\mathbf{b}\|_2),
\end{aligned}$$

from which we can conclude that

$$\begin{aligned}
&\|\text{RNNcell}_\phi(Q, \mathbf{b}, \mathbf{y}_k) - \text{RNNcell}_{\phi'}(Q, \mathbf{b}, \mathbf{y}_k)\|_2 \\
&\leq \|f_{\phi'}^L\|_2 \|V - V'\|_2 + \sum_{l=2}^L \left[\|V\|_2 \left(\prod_{j=l+1}^L \|W^j\|_2 \right) \|f_{\phi'}^{l-1}\|_2 \right] \|W^l - W'^l\|_2 \\
&\quad + \|V\|_2 \left(\prod_{l=2}^L \|W^l\|_2 \right) \left[\|W_1^1 - W_1'^1\| \|\mathbf{y}_k\|_2 + \|W_2^1 - W_2'^1\| (\|Q\|_2 \|\mathbf{y}_k\|_2 + \|\mathbf{b}\|_2) \right].
\end{aligned}$$

Note that we have assumed that the set of parameters Φ is a compact subset of the stable region and $(Q, \mathbf{b}) \in \mathcal{S}_{\mu, L}^{d \times d} \times \mathcal{B}$ are bounded, which imply that for all $\phi, \phi' \in \Phi$, the corresponding outputs $(\mathbf{y}_k)_{k \in \mathbb{N}}$ and $(\mathbf{y}'_k)_{k \in \mathbb{N}}$ are uniformly bound, and hence $\|f_{\phi'}^l\|_2$ is bounded for all k and $l = 1, \dots, L$ (see Eq. 79). Consequently, we see there exists a constant K such that Eq. 77 is satisfied. This finishes the proof of the desired sensitivity result. $\square$

(IV) Convergence. For the convergence of RNN_ϕ^k , we can only give the best case guarantee. It is easy to see that with the following choice of ϕ , RNN_ϕ^k can represent GD_s^k :

$$V = [I, -I], \quad W_1^1 = [I; -I]^\top, \quad W_2^1 = [-sI; sI]^\top, \quad W^l = I \text{ for } l = 2, \dots, L. \quad (80)$$

Therefore, for the best case, RNN_ϕ^k can converge at least as fast as GD_s^k .

E Experiment Details

Here we state the configuration details of the experiments.

- Convexity and smoothness. They are set to be $\mu = 0.1$ and $L = 1$, respectively.
- Dataset. 10000 pairs of $(\mathbf{x}, \mathbf{b})$ are generated in the following way: 10000 many $\mathbf{x}$ are uniformly sampled from $[-5, 5]^{10} \times \mathcal{U}^{5 \times 5}$, where $\mathcal{U}^{5 \times 5}$ denotes the space of all 5×5 unitary matrices. Each input $\mathbf{x}$ actually is a tuple $\mathbf{x} = (\mathbf{z}_x, U_x)$ where $\mathbf{z}_x \in [-5, 5]^{10}$ and U_x is unitary. 10000 many $\mathbf{b}$ are uniformly sampled from $[-5, 5]^5$. These 10000 pairs are viewed as the whole dataset.
- Training set S_n . During training, n samples are randomly drawn from these 10000 data points as the training set. The labels of these training samples are given by $\mathbf{y} = \text{Opt}(Q^*(\mathbf{x}), \mathbf{b})$.
- More details on $Q^*(\mathbf{x})$. As mentioned before, each $\mathbf{x}$ is a tuple $\mathbf{x} = (\mathbf{z}_x, U_x)$. Then we implement $Q^*(\mathbf{x}) = U_x \text{diag}([g^*(\mathbf{z}_x), \mu, L]) U_x^\top$, where g^* is a 2-layer dense neural network with hidden dimension 3, output dimension 3, and with randomly fixed parameters. Note that in the final layer of g^* , there is a sigmoid-activation that scales the output to the range $[0, 1]$ and then the range is further re-scaled to $[\mu, L]$. Finally, $g^*(\mathbf{z}_x)$ is concatenated with $[\mu, L]$ to form a 5-dimensional vector with smallest and largest value to be μ and L respectively. This vector represents the eigenvalues of $Q^*(\mathbf{x})$.
- Architecture of Q_θ . Q_θ has the same form as $Q^*(\mathbf{x})$, except that the network g^* in Q^* becomes g_θ in Q_θ . That is, $Q_\theta(\mathbf{x}) = U_x \text{diag}([g_\theta(\mathbf{z}_x), \mu, L]) U_x^\top$. Here g_θ is also a 2-layer dense neural network with output dimension 3, but the hidden dimension can vary. In the reported results, when we say hidden dimension=0, it means g_θ is a one-layer network.

For the experiments that compare RNN_ϕ^k with GD_ϕ^k and NAG_ϕ^k , they are conducted under the ‘learning to learn’ scenario, with the following modifications compared to the above setting.

- Dataset. Instead of sampling $(\mathbf{x}, \mathbf{b})$, here we directly sample the problem pairs $(Q, \mathbf{b})$. Similarly, 10000 pairs of $(Q, \mathbf{b})$ are sampled uniformly from $\mathcal{S}_{\mu, L}^{10 \times 10} \times [-5, 5]^{10}$.
- Architecture of RNN_ϕ^k . For each cell in RNN_ϕ^k , it is a 4-layer dense neural network with hidden dimension 20-20-20.

For all experiments, each model has been trained by both ADAM and SGD with learning rate searched over $[1\text{e-}2, 5\text{e-}3, 1\text{e-}3, 5\text{e-}4, 1\text{e-}4]$, and only the best result is reported. Furthermore, error bars are produced by 20 independent instantiations of the experiments. The experiments are mainly run parallelly (since we need to search the best learning rate) on clusters which have 416 nodes where on each node there are 24 Xeon 6226 CPU @ 2.70GHz with 192 GB RAM and 1x512 GB SSD.