

An Improved Finite-time Analysis of Temporal Difference Learning with Deep Neural Networks

Zhifa Ke¹ Zaiwen Wen² Junyu Zhang³

Abstract

Temporal difference (TD) learning algorithms with neural network function parameterization have well-established empirical success in many practical large-scale reinforcement learning tasks. However, theoretical understanding of these algorithms remains challenging due to the nonlinearity of the action-value approximation. In this paper, we develop an improved non-asymptotic analysis of the neural TD method with a general L -layer neural network. New proof techniques are developed and an improved new $\tilde{O}(\epsilon^{-1})$ sample complexity is derived. To our best knowledge, this is the first finite-time analysis of neural TD that achieves an $\tilde{O}(\epsilon^{-1})$ complexity under the Markovian sampling, as opposed to the best known $\tilde{O}(\epsilon^{-2})$ complexity in the existing literature.

1. Introduction

The temporal difference (TD) learning method, firstly designed for policy evaluation (Sutton, 1988), is a fundamental building block of many popular Reinforcement Learning (RL) algorithms. In standard TD learning algorithms for tabular MDP, based on the Bellman operator, the agent iteratively obtains a state-action-reward-transition tuple and then updates the Q values by a weighted average of the current value and the TD target. Once the algorithm converges, the Q function is considered to be the final return obtained by executing the target policy given some initial action-state pair.

¹Center for Data Science, Peking University, China. ²Beijing International Center for Mathematical Research, Center for Machine Learning Research and Changsha Institute for Computing and Digital Economy, Beijing, China ³Department of Industrial Systems Engineering and Management, National University of Singapore, Singapore. Correspondence to: Zhifa Ke <kezhifa@stu.pku.edu.cn>, Zaiwen Wen <wenzw@pku.edu.cn>, Junyu Zhang <junyuz@nus.edu.sg>.

For large-scale reinforcement learning (RL) problems, appropriate parameterization of the Q function is crucial for better scalability of the TD algorithms. Common examples include linear (Tesauro et al., 1995), general smooth non-linear (Maei et al., 2009), and neural network (Mnih et al., 2013) function approximations. However, it is well known that the naive extension of TD learning and Q-learning algorithms can diverge under the general function approximation Tsitsiklis & Van Roy (1996). To encourage convergence, numerous variants of TD and Q-learning have been proposed, including Least-squares TD (LSTD) (Bradtke & Barto, 1996; Boyan, 2002) and gradient TD (GTD) (Sutton et al., 2009a;b), to name a few.

The applications of neural network function approximation have witnessed huge empirical success in many real-world tasks, including Deep Q-network (DQN) algorithms (Mnih et al., 2013; Van Hasselt et al., 2016), policy improvement method (Sutton et al., 1999), trust region policy optimization (Schulman et al., 2015) and the actor-critic algorithms (Konda & Tsitsiklis, 1999; Lillicrap et al., 2015; Fujimoto et al., 2018), etc. However, due to the analysis difficulties brought by the function approximation, a significant gap exists between the empirical success and the theoretical understanding of these algorithms. Hence analyzing the convergence and sample complexity of TD learning and Q-learning under various Q function parameterizations has always been an active topic in the RL community during the past decades.

Early works focus on the asymptotic convergence of the algorithms with tabular or linear function approximation. For the tabular (stochastic) TD or Q learning method, Jaakkola et al. (1993) established the asymptotic convergence for the first time. Later on, the asymptotic convergence of algorithms with linear function approximation has been extensively discussed using ODE-based methods, see e.g. Tsitsiklis & Van Roy (1996); Perkins & Pendrith (2002); Borkar (2009). Meanwhile, in contrast to the convergent results for RL algorithms under the tabular or linear settings, TD with nonlinear function approximation is known to diverge in general Tsitsiklis & Van Roy (1996); Brandfonbrener & Bruna (2019). To overcome this issue, Maei et al. (2009) proposed to optimize the Mean Squared Projected Bellman

Error (MSPBE) via a gradient-based algorithm. Due to the problem nonconvexity, only asymptotic convergence to stationary points can be guaranteed.

More recently, benefiting from the improved techniques for analyzing stochastic optimization algorithms, there has been a growing number of research on providing finite-time analysis for TD and Q-learning algorithms with function approximations.

For linear function approximation, the non-asymptotic results of TD learning and its variants are relatively well-understood, including TD Bhandari et al. (2018); Dalal et al. (2018); Zou et al. (2019), gradient TD Dalal et al. (2018); Touati et al. (2018); Liu et al. (2020a), and Least-Squares TD Lazaric et al. (2010); Prashanth et al. (2014); Tagorti & Scherrer (2015), etc. In particular, Bhandari et al. (2018) established the first finite-time analysis of linear Q-learning under both i.i.d. sampling and Markovian sampling settings.

For neural network function approximation, which is directly related to this paper, we provide a more detailed discussion. Based on the recent advances in the understanding of optimizing ReLU network Jacot et al. (2018); Du et al. (2018); Allen-Zhu et al. (2019a;b); Cao & Gu (2019; 2020), a few recent works have successfully developed the finite-time analysis of the neural TD and neural Q-learning algorithms, as long as the Q network is sufficiently wide. Let Q^* be the true action-value function and let $Q(s, a; \theta)$ denote the action-value function parameterized by a neural network with weights θ , at any state action pair (s, a) . Then we aim to find some ϵ -optimal parameter $\bar{\theta}$ such that $\mathbb{E}[(Q(s, a; \bar{\theta}) - Q^*(s, a))^2] \leq \epsilon + \epsilon_{\mathcal{F}}$, where the expectation is taken over the possible randomness in the output $\bar{\theta}$ as well as the distribution over the state-action pairs (s, a) , and $\epsilon_{\mathcal{F}}$ is the optimal approximation error of the parameterization function class. In (Xu & Gu, 2020), a neural Q-learning algorithm with a general L -layer ReLU network is analyzed, and an $\tilde{O}(\epsilon^{-2})$ sample complexity is guaranteed given that the network is sufficiently wide. In (Cai et al., 2023), the authors studied both the neural TD learning and neural Q-learning algorithms for minimizing the MSPBE for policy evaluation and policy optimization, respectively. For policy evaluation, the Q^* in the definition of an ϵ -optimal solution is defaulted as Q^π with π being the policy to be evaluated. For both cases, an $\tilde{O}(\epsilon^{-2})$ sample complexity is guaranteed for wide two-layer ReLU networks. In (Sun et al., 2022), an $\tilde{O}(\epsilon^{-\frac{2}{2-\alpha}})$ complexity has been achieved by an adaptive neural TD algorithm with multi-layer ReLU networks, where $\alpha \in (0, 1]$ is a constant that characterizes the sparsity and decay rate of the stochastic semi-gradients. However, without additional assumption, only an $\tilde{O}(\epsilon^{-2})$ complexity with $\alpha = 1$ can be theoretically guaranteed. Finally, for policy evaluation problems, there are also several works that aim at reducing the width of the over-parameterized

Q networks in the existing works (Tian et al., 2022; Cayci et al., 2023). In terms of complexity, both of them requires $\tilde{O}(\epsilon^{-2})$ samples to obtain an ϵ -optimal solution.

Despite the fact that existing analysis of the neural TD or neural Q-learning algorithms merely provides the $\tilde{O}(\epsilon^{-2})$ sample complexity under various settings, an $\tilde{O}(\epsilon^{-1})$ sample complexity should be expected. In fact, a double-loop fitted Q-iteration (FQI) method (Fan et al., 2020) and its single-loop Gauss-Newton variant (Ke et al., 2023) can achieve an $\tilde{O}(\epsilon^{-1})$ sample complexity is obtained for two-layer Q networks. Let $\mathcal{T}$ be the Bellman (optimality) operator, then the FQI method repeatedly solves a non-linear least square subproblem to obtain the next iteration: $\theta_{k+1} \approx \operatorname{argmin}_{\theta \in \Theta} \mathbb{E}[(Q(s, a; \theta) - \mathcal{T}Q(s, a; \theta_k))^2]$. Compared to the single-loop neural TD or neural Q-learning method that takes only one sample (or a mini-batch) to update the weights of Q networks, the update scheme of FQI requires repeatedly solving a subproblem to sufficiently high accuracy to enable convergence, which makes it inefficient and less favorable in practice. Therefore, we would like to raise a question:

Can we improve the existing analysis of the neural temporal difference learning algorithm and obtain an $\tilde{O}(\epsilon^{-1})$ sample complexity under general multi-layer Q neural networks?

To answer this question, we revisit the convergence analysis of the neural TD learning or Q-learning algorithms under the non-i.i.d. Markovian observations where a general L -layer neural network is used for Q function parameterization. By proposing a new subspace analysis technique, under suitable conditions, we derive a brand new $\tilde{O}(\epsilon^{-1})$ sample complexity for neural TD learning or Q-learning, improving the state-of-the-art $\tilde{O}(\epsilon^{-2})$ sample complexity in the existing works. Our contributions are summarized as follows.

- Under the non-i.i.d. Markovian sampling setting, we derive an $\tilde{O}(\epsilon^{-1})$ sample complexity for both neural TD learning and Q-learning methods under the multi-layer network approximation for Q functions. Our result also improves the best known $\tilde{O}(\epsilon^{-2})$ sample complexity in the existing works.
- Based on our newly developed techniques, we further provide a finite-sample analysis for a minimax neural Q-learning algorithm that solves two-player zero-sum Markov games. An $\tilde{O}(\epsilon^{-1})$ sample complexity is obtained under the non-i.i.d. Markovian sampling setting.

Technically, the subspace analysis approach that we propose to establish the $\tilde{O}(\epsilon^{-1})$ sample complexity is by itself of independent interest. We believe this technique can potentially

	Neural Approximation	Network Depth	Network Width	Activation	Sample Complexity
(Bhandari et al., 2018)	No	NA	NA	NA	$\mathcal{O}(1/\epsilon)$
(Cai et al., 2023)	Yes	2	$\Omega(1/\epsilon^4)$	ReLU	$\mathcal{O}(1/\epsilon^2)$
(Xu & Gu, 2020)	Yes	L	$\Omega(1/\epsilon^6)$	ReLU	$\mathcal{O}(1/\epsilon^2)$
(Sun et al., 2022)	Yes	L	$\Omega(1/\epsilon^6)$	ReLU	$\mathcal{O}(1/\epsilon^{\frac{2}{2-\alpha}}), \alpha \in (0, 1]$
(Tian et al., 2022)	Yes	L	$\Omega(1/\epsilon^2)$	ELU, GeLU	$\mathcal{O}(1/\epsilon^2)$
Ours	Yes	L	$\Omega(1/\epsilon^2)$	ELU, GeLU	$\mathcal{O}(1/\epsilon)$

Table 1. Sample complexity for parameterized Q learning to find some $\bar{\theta}$ such that $\mathbb{E} [\|Q(s, a; \bar{\theta}) - Q^*(s, a)\|_\mu^2] \leq \epsilon$, where $\|f\|_\mu^2 := \int |f|^2 d\mu$ and $Q^*(s, a)$ satisfies the Bellman optimality equation $Q^*(s, a) = \mathcal{T}Q^*(s, a)$.

be applied to linear Q-learning algorithms and linear Actor-Critic algorithms without requiring the positive definiteness assumption of the feature covariance matrix (Bhandari et al., 2018; Zou et al., 2019; Barakat et al., 2022), while maintaining the $\tilde{O}(\epsilon^{-1})$ complexity.

In summary, we provide a comprehensive comparison between our work and the most related works in their respective settings and sample complexity in Table 1. Our work establishes an optimal sample complexity analysis within a broader contextual framework.

2. Preliminaries

We consider the infinite-horizon discounted Markov decision process (MDP), which is denoted as $\mathcal{M} = (\mathcal{S}, \mathcal{A}, \mathbb{P}, r, \gamma)$. We consider a general state space $\mathcal{S}$ and a finite action space $\mathcal{A}$. At any state $s \in \mathcal{S}$, if the agent takes an action $a \in \mathcal{A}$, it will receive a reward $r(s, a) \in [-R_{\max}, R_{\max}]$ and transition to the next state $s' \in \mathcal{S}$ with probability $\mathbb{P}(s'|s, a)$. We call r the reward function and $\mathbb{P}$ the transition kernel. Let $\gamma \in (0, 1)$ be a discount factor, then an MDP aims to find a sequence of actions $\{a_t\}_{t \geq 0}$ to maximize the expected and discounted cumulative reward $\mathbb{E}[\sum_{t=0}^{\infty} \gamma^t \cdot r(s_t, a_t) | s_0 \sim \mu]$, where μ is the distribution of the initial state s_0 .

Let $\Delta_{\mathcal{A}}$ denote the set of all probability distributions over the action space $\mathcal{A}$, and let a policy $\pi : \mathcal{S} \mapsto \Delta_{\mathcal{A}}$ be a mapping that returns a probability distribution $\pi(\cdot|s) \in \Delta_{\mathcal{A}}$ given any state $s \in \mathcal{S}$. If an agent follows a policy π , then at any state s_t , it will act by sampling an action $a_t \sim \pi(\cdot|s_t)$. Therefore, the action-value function (Q-function) under the policy π is

$$Q^\pi(s, a) := \mathbb{E}_\pi \left[\sum_{t=0}^{\infty} \gamma^t \cdot r(s_t, a_t) | s_0 = s, a_0 = a \right],$$

for $\forall (s, a) \in \mathcal{S} \times \mathcal{A}$, where all actions except a_0 are sampled according to π . For any mapping $Q : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$, let the Bellman operator $\mathcal{T}^\pi$ be

$$\mathcal{T}^\pi Q(s, a) := r(s, a) + \gamma \mathbb{E} [Q(s', a') | s' \sim \mathbb{P}(\cdot | s, a), a' \sim \pi(\cdot | s')], \forall s, a.$$

Then $\mathcal{T}^\pi$ is a γ -contraction under the infinity norm and Q^π is the unique solution to the fixed-point equation $Q = \mathcal{T}^\pi Q$ (Bertsekas, 2012). If the Q function is parameterized by some function $Q(s, a; \theta)$ to gain better scalability for large-scale RL problems, popular approaches for finding a good θ include minimizing the the Mean-Squared Bellman Error (MSBE):

$$\min_{\theta \in \Theta} \mathbb{E}_{(s,a) \sim \mu} [(Q(s, a; \theta) - \mathcal{T}^\pi Q(s, a; \theta))^2], \quad (1)$$

and minimizing the Mean-Squared Projected Bellman Error (MSPBE):

$$\min_{\theta \in \Theta} \mathbb{E}_{(s,a) \sim \mu} [(Q(s, a; \theta) - \Pi_{\mathcal{F}} \mathcal{T}^\pi Q(s, a; \theta))^2], \quad (2)$$

where Θ is a feasible domain of the parameter θ , μ is some distribution over state action pairs, and $\Pi_{\mathcal{F}}$ is the projection onto some function class $\mathcal{F}$. Typical choices of $\mathcal{F}$ includes the Q function parameterization class itself $\mathcal{F} := \{Q(\cdot; \theta) : \theta \in \Theta\}$ (Maei et al., 2009), and some local linearization of the parameterization function class (Cai et al., 2023).

In this paper, we study the neural temporal difference learning method where the action-value function is parameterized by some multi-layer neural network. Let us define a feed-forward neural network by the following recursion:

$$\mathbf{x}^{(l)} = \frac{1}{\sqrt{m}} \sigma(\mathbf{W}_l \mathbf{x}^{(l-1)}), \quad l \in \{1, 2, \dots, L\}, \quad (3)$$

where $\mathbf{W}_1 \in \mathbb{R}^{m \times d}$, $\mathbf{W}_l \in \mathbb{R}^{m \times m}$ for $2 \leq l \leq L$ are the weight matrices of the network, $\sigma(\cdot)$ is an activation function, and the input is a feature map $\mathbf{x}^{(0)} = \phi(s, a) \in \mathbb{R}^d$ for any state action pair (s, a) . For simplicity of notation, we write $\mathbf{x} = \mathbf{x}^{(0)}$, then $Q(s, a)$ is parameterized by

$$Q(\mathbf{x}; \theta) = \frac{1}{\sqrt{m}} \mathbf{b}^\top \mathbf{x}^{(L)}, \quad (4)$$

where the parameter $\theta = (\text{Vec}(\mathbf{W}_1); \dots; \text{Vec}(\mathbf{W}_L))$ denotes the collection of all weight matrices, and $\mathbf{b}$ is given by a random initialization. $\text{Vec}(\cdot)$ stands for the vectorization operator that reshapes a matrix to a column vector by stacking its columns one by one and the “;” separator in θ stands

for the vertical stacking of the elements. That is, we reshape θ to a long column vector for the notational convenience in later discussion.

Assumption 2.1. The activation function $\sigma(\cdot)$ is L_1 -Lipschitz and L_2 -smooth, i.e., for $\forall y_1, y_2 \in \mathbb{R}$:

$$|\sigma(y_1) - \sigma(y_2)| \leq L_1 |y_1 - y_2|$$

and

$$|\sigma'(y_1) - \sigma'(y_2)| \leq L_2 |y_1 - y_2|.$$

Assumption 2.1 indicates that our results below are not based on the popular ReLU activation function. However, we primarily focus on some twice-differentiable activation functions (such as Sigmoid, ELU, GeLU, etc.), which are smooth approximations of the ReLU function and are frequently utilized in practical problems (Devlin et al., 2018; Godfrey, 2019). Such a setup aligns with (Liu et al., 2020b), and provides a $\mathcal{O}(m^{-\frac{1}{2}})$ -smooth property for the neural Q-function.

Let $\theta^0 = (\text{Vec}(\mathbf{W}_1^0); \dots; \text{Vec}(\mathbf{W}_L^0))$ be the initial solution. For each l , we initialize the weights of $\mathbf{W}_l^0$ element-wise from a normal distribution $\mathcal{N}(0, 1)$ and each element of $\mathbf{b}$ is drawn uniformly from $\{-1, +1\}$. The parameter $\mathbf{b}$ will not be optimized during training. For regularity purpose, we would like to restrict the iterations to a bounded set around θ^0 , which is defined as

$$S_\omega := \{\theta = (\text{Vec}(\mathbf{W}_1); \dots; \text{Vec}(\mathbf{W}_L)) : \|\theta - \theta^0\|_2 \leq \omega, 1 \leq l \leq L\}.$$

In each iteration t , the neural Q-learning algorithm obtains a sample of state-action-reward-transition tuple $(s_t, a_t, r_t, s_{t+1}, a_{t+1})$ and computes the TD error by

$$\Delta_t = Q(\mathbf{x}_t; \theta^t) - (r_t + \gamma Q(\mathbf{x}_{t+1}; \theta^t)) \quad (5)$$

with $\mathbf{x}_t = \phi(s_t, a_t)$, $\mathbf{x}_{t+1} = \phi(s_{t+1}, a_{t+1})$. Then a projected stochastic semi-gradient step is performed to update the weight matrices:

$$\theta^{t+1} = \Pi_{S_\omega}(\theta^t - \eta_t g(\theta^t)) \quad (6)$$

with

$$g(\theta^t) = \Delta_t \cdot \nabla_\theta Q(\mathbf{x}_t; \theta^t).$$

We formally describe the neural TD learning method in Algorithm 1.

Algorithm 1 Neural Temporal Difference Learning with Markovian Sampling

Input: A learning policy π , a discount factor $\gamma \in (0, 1)$, a sequence of learning rates $\{\eta_t\}_{t \geq 0}$, a maximum iteration number T , a projection radius $\omega > 0$, a Q network with architecture (4).

Initialization: Generate each entry of $\mathbf{W}_l^0$ independently from $\mathcal{N}(0, 1)$, for $l = 1, 2, \dots, L$, and each entry of $\mathbf{b}$ independently from $\text{Unif}\{-1, +1\}$. Generate $s_0 \sim \mu$, $a_0 \sim \pi(\cdot | s_0)$.

for $t = 0, 1, \dots, T - 1$ **do**

 Sample $(s_t, a_t, r_t, s_{t+1}, a_{t+1})$ from the learning policy π with $a_{t+1} \sim \pi(\cdot | s_{t+1})$.

 Compute the TD error Δ_t by (5).

 Update θ^{t+1} by the projected stochastic semi-gradient step (6).

end for

Output: θ^T .

One remark is that, under the non-i.i.d. Markovian sampling setting, the agent is only able to generate a trajectory of samples following some given learning policy π , which is very common in the offline RL (Wu et al., 2019; Levine et al., 2020; Kostrikov et al., 2021) where the data trajectories are generated by some learning policy.

In later sections, we will revisit the Algorithm 1 and design a novel subspace analysis technique for this method and achieve an improved sample complexity of $\tilde{\mathcal{O}}(\epsilon^{-1})$. Moreover, by replacing the TD error induced by the Bellman operator (5) with TD error induced by the Bellman optimality operator: $\Delta_t = Q(\mathbf{x}_t; \theta^t) - (r_t + \gamma \max_{b \in \mathcal{A}} Q(s', b; \theta^t))$, Algorithm 1 can be reduced to the neural Q-learning method for finding optimal state-action value Q^* . Our analysis for neural TD learning can be extended to the neural Q-learning analogously and obtain the same $\tilde{\mathcal{O}}(\epsilon^{-1})$ sample complexity.

3. Convergence of Neural Temporal Difference Learning

3.1. Basic Settings and Assumptions

To analyze Algorithm 1, let us first define the local linearization function class of the multi-layer Q network (4) at the random initialization θ^0 :

$$\mathcal{F}_{\omega, m} := \left\{ \hat{Q}(\cdot; \theta) = Q(\cdot; \theta^0) + \langle \nabla_\theta Q(\cdot; \theta^0), \theta - \theta^0 \rangle \right\} \quad (7)$$

for any $\theta \in S_\omega$. Consider the MSPBE minimization problem:

$$\min_{\theta \in S_\omega} \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[(Q(\mathbf{x}; \theta) - \Pi_{\mathcal{F}_{\omega, m}} \mathcal{T}^\pi Q(\mathbf{x}; \theta))^2 \right], \quad (8)$$

where μ is the initial state distribution, π is the learning policy, and $\mathbb{P}$ is the transition kernel, the expectation $\mathbb{E}_{\mu, \pi, \mathbb{P}}[\cdot]$ is taken over $s \sim \mu$, $a \sim \pi(\cdot|s)$, and $s' \sim \mathbb{P}(\cdot|s, a)$, $a' \sim \pi(\cdot|s')$ in $\mathcal{T}^\pi$. Define the set Ξ_β as

$$\Xi_\beta := \{\theta \in S_\beta : \widehat{Q}(x; \theta) = \Pi_{\mathcal{F}_{\omega, m}} \mathcal{T}^\pi \widehat{Q}(x; \theta), \forall x = \phi(s, a)\}. \quad (9)$$

Then the set Ξ_ω consists of the points θ with which $\widehat{Q}(\cdot; \theta)$ forms a fixed point of the projected Bellman operator $\Pi_{\mathcal{F}_{\omega, m}} \mathcal{T}^\pi$ for the problem (8). By Section 4.1 in (Cai et al., 2023), the fixed point of $\Pi_{\mathcal{F}_{\omega, m}} \mathcal{T}^\pi$ is unique for $\theta \in S_\omega$. Therefore, the following relationship holds

$$\widehat{Q}(x; \theta) = \widehat{Q}(x; \theta'). \quad (10)$$

for $\forall x = \phi(s, a)$, $\forall (s, a) \in \mathcal{S} \times \mathcal{A}$, $\forall \theta, \theta' \in \Xi_\beta$, $\forall \beta \geq \omega$. Moreover, it is also shown that a point $\theta^* \in \Xi_\omega$ if and only if it satisfies the stationarity condition:

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(x, x'; \theta^*) \langle \nabla_{\theta} \widehat{Q}(x; \theta^*), \theta - \theta^* \rangle \right] \geq 0, \quad (11)$$

where $\widehat{Q}(\cdot; \theta^*) \in \mathcal{F}_{\omega, m}$ is a local linearization provided by (7) and $\widehat{\Delta}$ is defined as

$$\widehat{\Delta}(x, x'; \theta^*) = \widehat{Q}(x; \theta^*) - (r(s, a) + \gamma \widehat{Q}(x'; \theta^*)).$$

Hence people may analyze the gap between $Q^\pi(\cdot)$ and $Q(\cdot; \theta^T)$ by first connecting it to $\widehat{Q}(\cdot; \theta^*)$. Based on this, Cai et al. (2023) derived an $\tilde{O}(\epsilon^{-2})$ sample complexity for the neural TD method. Now we define

$$\Sigma_\pi = \mathbb{E}_{\mu, \pi} [\nabla_{\theta} Q(x; \theta^0) \nabla_{\theta} Q(x; \theta^0)^\top]. \quad (12)$$

It is worth noting that the matrix Σ_π only depends on π and θ^0 . In the original assumption about (12), (Zou et al., 2019; Xu & Gu, 2020) in fact assumed positive definiteness ($\succ 0$) of Σ_π , which can be viewed as a generalized version of the positive definite feature covariance matrix assumption in the analysis of linear TD and linear Q-learning, see e.g. (Zou et al., 2019). However, in this paper we adopt the following weaker regularity assumption.

Assumption 3.1. Let $\overline{\sigma}_{\min}(\Sigma_\pi)$ denote the minimum non-zero singular value of the matrix Σ_π , then there exist constants $\lambda_0, m^* > 0$ such that $\overline{\sigma}_{\min}(\Sigma_\pi) \geq \lambda_0$ as long as the Q network width $m \geq m^*$.

For neural Q function approximation, a sufficient but not necessary condition for Assumption 3.1 can be obtained by exploiting the theory of over-parameterized neural networks. Roughly speaking, for a finite MDP with an L -layer ReLU Q network, if the feature map satisfies $\phi(s, a) \not\parallel \phi(s', a')$ for $\forall (s, a) \neq (s', a')$, the results of (Jacot et al., 2018; Allen-Zhu et al., 2019a;b; Cao & Gu, 2019; 2020) suggest

that there exist $\lambda', m^* > 0$ such that with high probability $\text{Gram}(\theta_0) \succ \lambda' \cdot \mathbf{I}$ for networks with width $m \geq m^*$. Here $\text{Gram}(\theta_0)$ stands for the Gram matrix of the network at the initialization θ_0 . A lower bound on $\overline{\sigma}_{\min}(\Sigma_\pi)$ can then be constructed with λ' , refer to Remark D.6 in Appendix D.

Finally, to facilitate the sample complexity analysis under the non-i.i.d. Markovian sampling setting, let us make the following assumption on the fast mixing rate of the MDP sample trajectories, which is widely adopted in the related analysis (Zou et al., 2019; Xu & Gu, 2020; Cai et al., 2023).

Assumption 3.2. We assume that the Markov chain $\{s_t\}_{t=0,1,\dots}$ induced by the learning policy π and the transition kernel $\mathbb{P}$ is uniformly ergodic with its invariant measure $\mathbb{P}^\pi$. Furthermore, we assume that there are constants $\kappa > 0, \rho \in (0, 1)$ such that

$$\sup_{s \in \mathcal{S}} d_{TV}(\mathbb{P}(s_t \in \cdot | s_0 = s), \mathbb{P}^\pi) \leq \kappa \rho^t$$

for all $t \geq 0$.

Without loss of generality, we also make the following technical assumption, which is not fundamental as opposed to Assumption 3.1, and 3.2.

Assumption 3.3. We assume the initial state distribution μ to be the stationary state distribution under policy π .

This assumption is in fact very natural. Concerning the stationarity of μ , it can always be guaranteed by abandoning the first $\tilde{O}(t_{\text{mix}})$ samples while Assumption 3.2 indicates that the mixing time $t_{\text{mix}} = \tilde{O}(1)$. This assumption guarantees that the operator $\mathcal{T}^\pi$ is γ -contractive w.r.t. $\|\cdot\|_\mu$ in policy evaluation. Similar assumptions are included in (Bhandari et al., 2018; Cai et al., 2023).

3.2. An Improved Complexity of Neural TD Learning

To derive the $\tilde{O}(\epsilon^{-1})$ sample complexity, we rely on the following key observation on subspace decomposition, which is beyond the existing analysis framework.

Proposition 3.4. Let $\mathcal{R}(\Sigma_\pi)$ and $\mathcal{K}(\Sigma_\pi)$ denote the range space and kernel space of the matrix Σ_π , respectively. Then for any parameter $\theta \in S_\omega$, there exists θ_* such that

$$\theta_* \in \Xi_{2\omega} \quad \text{and} \quad \theta - \theta_* \in \mathcal{R}(\Sigma_\pi),$$

which also implies that the projections of θ and θ_* onto the subspace $\mathcal{K}(\Sigma_\pi)$ are identical.

Based on this argument, for the iteration sequence $\{\theta^t\}_{t \geq 0}$ generated by Algorithm 1, there exists a sequence $\{\theta_*^t\}_{t \geq 0} \subseteq \Xi_{2\omega}$ such that $\{\theta^t - \theta_*^t\}_{t \geq 0} \subseteq \mathcal{R}(\Sigma_\pi)$. Therefore, unlike the existing works that analyze $\|\theta^t - \theta^*\|^2$ for some $\theta^* \in \Xi_\omega$, c.f. (Cai et al., 2023; Xu & Gu, 2020), we will prove a much faster convergence in $\|\theta^t - \theta_*^t\|^2$. Combined with (10), this further indicates the improved sample

complexity in this paper. The proof of this proposition is presented as follows.

Proof. For the ease of discussion, let us denote the dimension of weight parameter θ as n . Then we may denote $\Sigma_\pi \in \mathbb{R}^{n \times n}$ and $\theta \in \mathbb{R}^n$. First of all let us fix an arbitrary $\bar{\theta} \in \Xi_\omega$, then we may decompose it into two orthogonal components:

$$\bar{\theta} = \bar{\theta}_\parallel + \bar{\theta}_\perp \quad \text{s.t.} \quad \bar{\theta}_\parallel \in \mathcal{R}(\Sigma_\pi) \text{ and } \bar{\theta}_\perp \in \mathcal{K}(\Sigma_\pi).$$

Similarly, we can decompose the currently considered vector θ as

$$\theta = \theta_\parallel + \theta_\perp \quad \text{s.t.} \quad \theta_\parallel \in \mathcal{R}(\Sigma_\pi) \text{ and } \theta_\perp \in \mathcal{K}(\Sigma_\pi).$$

Note that having an arbitrary vector $v \in \mathbb{R}^n$ in the kernel space of Σ_π means that $\Sigma_\pi v = 0$, which further indicates that

$$\begin{aligned} 0 &= v^\top \Sigma_\pi v \\ &= v^\top \mathbb{E}_{\mu, \pi} [\nabla_\theta Q(x; \theta^0) \nabla_\theta Q(x; \theta^0)^\top] v \\ &= \mathbb{E}_{\mu, \pi} [\langle \nabla_\theta Q(x; \theta^0), v \rangle^2]. \end{aligned}$$

Therefore, under the measure $(s, a) \sim \mu \times \pi$, we have

$$v \in \mathcal{K}(\Sigma_\pi) \implies \langle \nabla_\theta Q(x; \theta^0), v \rangle = 0 \quad a.s. \quad (13)$$

where *a.s.* stands for almost surely. Therefore, define $\theta_* = \bar{\theta}_\parallel + \theta_\perp$, we can check the stationarity condition (11) for θ_* by establishing:

$$\begin{aligned} &\mathbb{E}_{\mu, \pi, \mathbb{P}} [\hat{\Delta}(x, x'; \theta_*) \cdot \langle \nabla_\theta \hat{Q}(x; \theta_*), \theta' - \theta_* \rangle] \\ &= \mathbb{E}_{\mu, \pi, \mathbb{P}} [\hat{\Delta}(x, x'; \bar{\theta}) \cdot \langle \nabla_\theta \hat{Q}(x; \theta_0), \theta' - \bar{\theta} \rangle] \geq 0 \end{aligned} \quad (14)$$

The proof of (14) is lengthy and is thus moved to Appendix A.1 for succinctness. As a result we have $\theta_* \in \Xi_{2\omega}$ and $\theta - \theta_* = \bar{\theta}_\perp \in \mathcal{K}(\Sigma_\pi)$. Note that for any $\theta \in S_\omega$,

$$\begin{aligned} \|\theta_* - \theta^0\| &\leq \|\bar{\theta}_\parallel - \theta_\parallel^0\| + \|\theta_\perp - \theta_\perp^0\| \\ &\leq \|\bar{\theta} - \theta^0\| + \|\bar{\theta} - \theta^0\| \leq 2\omega, \end{aligned}$$

which completes the proof. $\square$

Following basic linear algebra analysis, we also have the following proposition.

Proposition 3.5. *Under Assumption 3.1, suppose the adopted Q network is sufficiently wide so that $m \geq m^*$, then for any $\theta \in \mathcal{R}(\Sigma_\pi)$, we have $\theta^\top \Sigma_\pi \theta \geq \lambda_0 \|\theta\|_2^2$.*

Proposition 3.4 indicates that the variations in the local linearization of Q -function values solely depend on the variations in parameters within the subspace

$\mathcal{R}(\Sigma_\pi)$. In the mean while, Proposition 3.5 indicates that such local linearization is non-singular within $\mathcal{R}(\Sigma_\pi)$. Based on these observations, we can first provide a fast convergence of $\|\theta^t - \theta_*^t\|^2 = \mathcal{O}(1/T)$ and then show that $\mathbb{E}[(\hat{Q}(x; \theta^T) - \hat{Q}(x; \theta_*^T))^2] = \mathbb{E}[(\hat{Q}(x; \theta^T) - \hat{Q}(x; \theta_*^T))^2] \leq \mathcal{O}(1/T)$ for any $\theta_* \in \Xi_\omega$. We summarize this result in Theorem 3.6 while presenting its proof in Appendix A.2.

Theorem 3.6. *Suppose Assumptions 3.1, 3.2 and 3.3 hold. We set $\omega = \tilde{C}_1$ and the learning rate $\eta_t = \frac{1}{2(1-\gamma)\lambda_0(t+1)}$. If the feature map $\|\phi(s, a)\| = 1$ for each state-action pair (s, a) and the network width $m \geq m^*$, then the output θ^T of Algorithm 1 satisfies*

$$\begin{aligned} &\mathbb{E}[(\hat{Q}(x; \theta^T) - \hat{Q}(x; \theta_*^T))^2 | \theta^0] \\ &\leq \frac{\tilde{C}_3(\log T + 1)}{(1-\gamma)^2 \lambda_0^2 T} + \frac{\tilde{C}_4 m^{-1/2}}{(1-\gamma)\lambda_0} \cdot \sqrt{\log(T/\delta)} \\ &+ \frac{\tilde{C}_5 \tau^* (\log(T/\delta) + 1) \log T}{(1-\gamma)^2 \lambda_0^2 T}, \end{aligned}$$

with probability at least $1 - 2\delta - 2L \exp(-\tilde{C}_2 m)$, where τ^* is the mixing time of Markov chain in Assumption 3.2, and $\tilde{C}_1, \dots, \tilde{C}_5 > 0$ are universal constants.

Let Q^π be the true state-action value function that satisfies the Bellman equation $Q^\pi = \mathcal{T}^\pi Q^\pi$. Then based on the convergence of the local linearization in Theorem 3.6, we establish the global convergence of neural temporal difference learning as Theorem 3.7.

Theorem 3.7. *Suppose the conditions in Theorem 3.6 hold. Then the output of Algorithm 1 satisfies*

$$\begin{aligned} &\mathbb{E}[(Q(\phi(s, a); \theta^T) - Q^\pi(s, a))^2 | \theta^0] \\ &\leq \frac{3\mathbb{E}[(Q^\pi(s, a) - \Pi_{\mathcal{F}_{\omega, m}} Q^\pi(s, a))^2]}{(1-\gamma)^2} + \tilde{C}_6 m^{-1} \\ &+ \frac{\tilde{C}_7(\log T + 1)}{(1-\gamma)^2 \lambda_0^2 T} + \frac{\tilde{C}_8 m^{-1/2}}{(1-\gamma)\lambda_0} \sqrt{\log(T/\delta)} \\ &+ \frac{\tilde{C}_9 \tau^* (\log(T/\delta) + 1) \log T}{(1-\gamma)^2 \lambda_0^2 T} \end{aligned} \quad (15)$$

w.p. $1 - 2\delta - 2L \exp(-\tilde{C}_2 m)$, where $\tilde{C}_6, \dots, \tilde{C}_9 > 0$ are universal constants.

Let $\epsilon_{\mathcal{F}} := \frac{3}{(1-\gamma)^2} \mathbb{E}[(Q^\pi(s, a) - \Pi_{\mathcal{F}_{\omega, m}} Q^\pi(s, a))^2]$ be the optimal approximation error of the function class $\mathcal{F}_{\omega, m}$. Then Theorem 3.7 demonstrates that under suitable parameter choices, neural TD learning method identify an approximation error bound of $\mathcal{O}(\epsilon_{\mathcal{F}} + \epsilon + m^{-\frac{1}{2}})$ within $\mathcal{O}(\epsilon^{-1})$ samples. Existing works include Cai et al. (2023); Xu &

Gu (2020); Tian et al. (2022); Cayci et al. (2023) achieve $\tilde{O}(\epsilon^{-2})$ sample complexity, and (Sun et al., 2022) achieves $\mathcal{O}(\epsilon^{-\frac{2}{2-a}})$, $a \in (0, 1]$ with additional assumptions.

Following a similar analysis while adopting an additional regularity assumption on the matrix Σ_π , one can further extend the above analysis to the neural Q-learning by substituting the Bellman operator with the Bellman optimality operator. A similar $\mathcal{O}(\epsilon^{-1})$ sample complexity can still be achieved, which is relegated to Appendix for succinctness.

4. Convergence of Minimax Neural Q-Learning

A two-player zero-sum Markov game (Littman, 1994; Bowling & Veloso, 2001; Perolat et al., 2018), as a simple variant of MDP, is defined as a six-tuple $\mathcal{M} = (\mathcal{S}, \mathcal{A}_1, \mathcal{A}_2, \mathbb{P}, r, \gamma)$. Here $\mathcal{S}$ is state space, $\mathcal{A}_1$ and $\mathcal{A}_2$ are the action space of the first and second player, respectively, $\mathbb{P} : \mathcal{A}_1 \times \mathcal{A}_2 \rightarrow \mathcal{P}(\mathcal{S})$ is the transition probability, $r : \mathcal{S} \times \mathcal{A}_1 \times \mathcal{A}_2 \rightarrow \mathbb{R}$ is the reward function and γ is the discounted factor. At time t , player 1 and player 2 take actions ($a_t^1 \in \mathcal{A}_1$ and $a_t^2 \in \mathcal{A}_2$) simultaneously. Player 1 obtains the reward $r(s_t, a_t^1, a_t^2)$, while player 2 obtains $-r(s_t, a_t^1, a_t^2)$. The goal of the two players is to maximize their cumulative rewards respectively. For a policy pair (π_1, π_2) , we can define the state-action value function as follows:

$$Q^{\pi_1, \pi_2}(s, a^1, a^2) = \mathbb{E}_{\pi_1, \pi_2} \left[\sum_{t=0}^{\infty} \gamma^t \cdot r(s_t, a_t^1, a_t^2) \mid s_0 = s, a_0^1 = a^1, a_0^2 = a^2 \right], \forall s, a^1, a^2.$$

The optimal state-action value function Q^* is defined as

$$\begin{aligned} Q^*(s, a^1, a^2) &= \max_{\pi_1} \min_{\pi_2} Q^{\pi_1, \pi_2}(s, a^1, a^2) \\ &= \min_{\pi_2} \max_{\pi_1} Q^{\pi_1, \pi_2}(s, a^1, a^2). \end{aligned}$$

We denote the optimal policy pair $\pi^* = \{\pi_1^*, \pi_2^*\}$ if $Q^*(s, a^1, a^2) = Q^{\pi_1^*, \pi_2^*}(s, a^1, a^2)$. Moreover, the Minimax Bellman operator $\mathcal{H}$ for the Markov game is defined as

$$\mathcal{H}Q(s, a^1, a^2) = r(s, a^1, a^2) + \gamma \mathbb{E} \left[\min_{b^1} \max_{b^2} Q(s', b^1, b^2) \mid s' \sim \mathbb{P}(\cdot \mid s, a^1, a^2) \right], \forall s, a^1, a^2.$$

Thus $\mathcal{H}Q^* = Q^*$. Let the feature map $\mathbf{x} = \phi(s, a^1, a^2)$ and $\pi = \{\pi^1, \pi^2\}$ be a given learning policy for players 1 and 2. Assume that $\{s_t, a_t^1, a_t^2, r_t\}_{t=0}^T$ is a sampled trajectory of states, actions and rewards obtained from the environment using policy π . Let us recall the definition of the local linearization function class $\mathcal{F}_{\omega, m}$ introduced in (7). Consider the MSPBE minimization problem with multi-layer neural network approximation:

$$\min_{\theta \in S_\omega} \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[(Q(\mathbf{x}; \theta) - \Pi_{\mathcal{F}_{\omega, m}} \mathcal{H}Q(\mathbf{x}; \theta))^2 \right].$$

To solve this problem, we still adopt the projected stochastic semi-gradient iteration method is provided described by (6), that is,

$$\theta^{t+1} = \Pi_{S_\omega} \left(\theta^t - \eta_t g(\theta^t) \right), \quad (16)$$

while redefining the stochastic semi-gradient estimator $g(\theta^t)$ as

$$g(\theta^t) = \Delta(s_t, a_t^1, a_t^2, s_{t+1}; \theta^t) \cdot \nabla_\theta Q(\mathbf{x}_t; \theta^t),$$

where $\mathbf{x}_t := \phi(s_t, a_t^1, a_t^2)$ and

$$\begin{aligned} \Delta(s_t, a_t^1, a_t^2, s_{t+1}; \theta^t) &= Q(\mathbf{x}_t; \theta^t) - \left(r(s_t, a_t^1, a_t^2) + \right. \\ &\quad \left. \gamma \max_{b^1 \in \mathcal{A}_1} \min_{b^2 \in \mathcal{A}_2} Q(\phi(s_{t+1}, b^1, b^2); \theta^t) \right). \end{aligned} \quad (17)$$

Now we redefine the function class $\mathcal{F}_{\omega, m}$ as a collection of all local linearization of $Q(\mathbf{x}; \theta)$ at the initial point θ^0 :

$$\mathcal{F}_{\omega, m} = \left\{ \widehat{Q}(\mathbf{x}; \theta) = Q(\mathbf{x}; \theta^0) + \langle \nabla_\theta Q(\mathbf{x}; \theta^0), \theta - \theta^0 \rangle, \theta \in S_\omega \right\}.$$

To analyze this method, for any $\beta > 0$, we redefine the set Ξ_β introduced in (9) by replacing the Bellman operator $\mathcal{T}^\pi$ with the Minimax Bellman operator $\mathcal{H}$. Similar to (10), we still have a point $\theta^* \in \Xi_\omega$ if and only if

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(s, a^1, a^2, s'; \theta^*) \langle \nabla_\theta \widehat{Q}(\phi(s, a^1, a^2); \theta^*), \theta - \theta^* \rangle \right] \geq 0,$$

where $\widehat{Q}(\cdot; \theta^*) \in \mathcal{F}_{\omega, m}$, and $\widehat{\Delta}(s, a^1, a^2, s'; \theta)$ has the same structure as $\Delta(s, a^1, a^2, s'; \theta)$ expect that the function $Q(\cdot; \theta)$ is replaced by $\widehat{Q}(\cdot; \theta)$.

Unlike the neural temporal difference learning method that aims at evaluating the state-action values of a fixed learning policy. The Minimax Bellman operator significantly sophisticates the analysis. Let us redefine the feature covariance matrix Σ_π with respect to the learning policy $\pi = \{\pi^1, \pi^2\}$, that is

$$\Sigma_\pi = \mathbb{E}_\pi \left[\nabla_\theta Q(s, a^1, a^2; \theta^0) \nabla_\theta Q(s, a^1, a^2; \theta^0)^\top \right].$$

Let the actions (a_θ^1, a_θ^2) satisfies $\langle \nabla_\theta Q(s, a^1, a^2; \theta^0), \theta \rangle = \max_{a^1 \in \mathcal{A}_1} \min_{a^2 \in \mathcal{A}_2} \langle \nabla_\theta Q(s, a^1, a^2; \theta^0), \theta \rangle$. For each parameter pair $\theta_s = (\theta_1, \theta_2)$, we define the action pair $(a_{\theta_s}^1, a_{\theta_s}^2)$ that satisfies

$$\begin{aligned} (a_{\theta_s}^1, a_{\theta_s}^2) &= \arg \max_{(a^1, a^2) \in \{(a_{\theta_1}^1, a_{\theta_2}^2), (a_{\theta_2}^1, a_{\theta_1}^2)\}} \\ &\quad \left\{ |\langle \nabla_\theta Q(s, a^1, a^2; \theta^0), \theta_1 - \theta_2 \rangle| \right\}. \end{aligned}$$

Then for any θ_1, θ_2 , the minimax feature covariance matrix is defined as follows:

$$\begin{aligned} \Sigma_\pi^*(\theta_1, \theta_2) &= \mathbb{E}_\pi \left[\nabla_\theta Q(s, a_{\theta_s}^1, a_{\theta_s}^2; \theta^0) \right. \\ &\quad \left. \nabla_\theta Q(s, a_{\theta_s}^1, a_{\theta_s}^2; \theta^0)^\top \right]. \end{aligned}$$

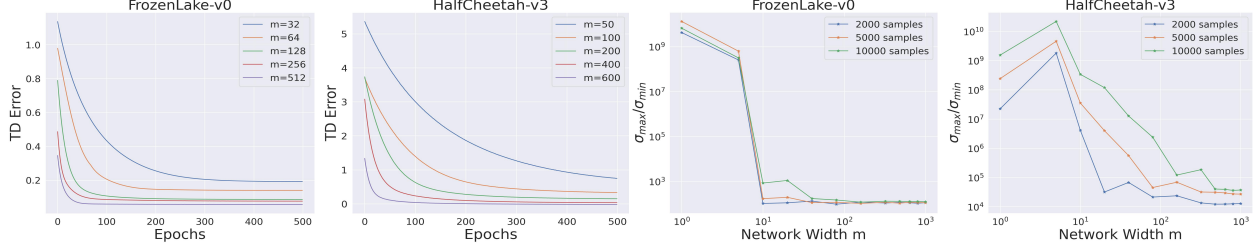


Figure 1. Training curves and the ratio of the largest and smallest non-zero singular values of Σ_π over different network widths m .

Assumption 4.1. For any θ^1, θ^2 , there exists a constant $\nu \in (0, 1)$ such that $(1 - \nu)^2 \Sigma_\pi - \gamma^2 \Sigma_\pi^*(\theta^1, \theta^2) \succeq 0$.

Note the original version of this assumption in (Zou et al., 2019) in fact requires a strict positive definite condition: $((1 - \nu)^2 \Sigma_\pi - \gamma^2 \Sigma_\pi^*(\theta^1, \theta^2)) \succ 0$. Under this additional assumption, (Zou et al., 2019) obtained an $\tilde{O}(\epsilon^{-1})$ sample complexity for minimax Q-learning with linear function approximation. With the help of our subspace analysis technique, in this paper, we relax it to the positive semi-definiteness ($\succeq 0$). Now we are ready to state our result for minimax neural Q-learning.

Theorem 4.2. Suppose Assumptions 3.1, 3.2 and 4.1 hold. We set $\omega = \tilde{C}_1$ and the learning rate $\eta_t = \frac{1}{2\nu\lambda_0(t+1)}$. If the feature map $\|\phi(s, a^1, a^2)\| = 1$ for each state-action pair (s, a^1, a^2) and the network width $m \geq m^*$, then the output θ^T of neural minimax Q-learning Algorithm 3 satisfies

$$\begin{aligned} & \mathbb{E} \left[\left(\hat{Q}(x; \theta^T) - \hat{Q}(x; \theta^*) \right)^2 \mid \theta^0 \right] \\ & \leq \frac{\tilde{C}_3(\log T + 1)}{\nu^2 \lambda_0^2 T} + \frac{\tilde{C}_4 m^{-1/2}}{\nu \lambda_0} \cdot \sqrt{\log(T/\delta)} \\ & \quad + \frac{\tilde{C}_5 \tau^* (\log(T/\delta) + 1) \log T}{\nu^2 \lambda_0^2 T}, \end{aligned}$$

with probability at least $1 - 2\delta - 2L \exp(-\tilde{C}_2 m)$, where τ^* is the mixing time of Markov chain in Assumption 3.2, and $\{\tilde{C}_i > 0\}_{i=1, \dots, 5}$ are universal constants.

Theorem 4.2 establishes a finite-time analysis of $\tilde{O}(\epsilon^{-1})$ -sample complexity for minimax neural Q-learning in terms of the function class $\mathcal{F}_{\omega, m}$. For a more specific description and theorem proof, see Appendix C. To the best of our knowledge, this is the first analysis of minimax Q-learning with neural network function approximation, characterized by a complexity bound of $\tilde{O}(\epsilon^{-1})$.

5. Experiments

Finally, we construct several experiments over the OpenAI Gym (Brockman et al., 2016) tasks and validate our theoret-

ical findings. We consider a two-layer neural network, as follows:

$$Q(s, a; \theta) := \frac{1}{\sqrt{m}} \sum_{r=1}^m b_r \sigma(\theta_r^\top \phi(s, a)),$$

where $\sigma(\cdot)$ is ELU activation in this section. Furthermore, details regarding the initialization and iteration methods for the parameters can be found in Section 2. For all experiments, we generate samples based on a prescribed ϵ -greedy policy with $\epsilon = 0.1$. To prevent redundancy in the features $\phi(s, a)$, we employ one-hot encoding for discrete action-state spaces and implement a fixed grid discretization for continuous spaces. When both $\phi(s, a)$ and $\phi(s', a')$ belong to the same one-hot encoding or grid, we treat them as the same sample point. Our investigation into the impact of network width on the TD learning algorithm will be conducted from two perspectives: (i) examining whether the network width m is correlated with the TD error, and (ii) exploring the existence of constants m^* and λ_0 that satisfy Assumption 3.1.

The four subfigures in Figure 1 represent two types of environments: one with a discrete state space and the other with a continuous state space. The first two subfigures depict the convergence performance of the TD algorithm at different network widths. We generate 2,000 sample points and run for 500 epochs. Notably, as the parameter m increases, the TD algorithm demonstrates faster convergence, resulting in smaller final TD errors. The latter two subfigures illustrate the existence of m^* and λ_0 . Specifically, we compute the largest non-zero singular value $\sigma_{\max}$ and smallest non-zero singular value $\sigma_{\min}$ of the matrix Σ_π . To mitigate the absolute magnitude of $\sigma_{\min}$, we introduce the ratio $r = \sigma_{\max}/\sigma_{\min}$ as a metric to validate Assumption 3.1. It can be observed that the value of r approaches a constant as m increases for all cases, providing empirical support for the validity of the assumption.

6. Conclusion

We study the finite-time analysis of the TD and Q learning methods with neural network approximation, where the

state-action pairs are generated by a given policy under the Markovian sampling. Besides the convergence to the true action-value function except for an inevitable function approximation error, an improved analysis technique is introduced to establish an $\tilde{O}(\epsilon^{-1})$ complexity for the neural TD and Q learning methods, which improves the existing $\tilde{O}(\epsilon^{-2})$ complexity. For future work, it is also interesting to investigate if the proposed technique can improve the current complexity estimate of the actor-critic methods, which are partially built upon the neural TD methods.

7. Acknowledgements

Dr. Zaiwen Wen is supported in part by the NSFC grant 12331010. Dr. Junyu Zhang is supported in part by the MOE AcRF grant A-0009530-05-00.

References

- Allen-Zhu, Z., Li, Y., and Liang, Y. Learning and generalization in overparameterized neural networks, going beyond two layers. *Advances in neural information processing systems*, 32, 2019a.
- Allen-Zhu, Z., Li, Y., and Song, Z. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pp. 242–252. PMLR, 2019b.
- Barakat, A., Bianchi, P., and Lehmann, J. Analysis of a target-based actor-critic algorithm with linear function approximation. In *International Conference on Artificial Intelligence and Statistics*, pp. 991–1040. PMLR, 2022.
- Bertsekas, D. *Dynamic programming and optimal control: Volume I*, volume 1. Athena scientific, 2012.
- Bhandari, J., Russo, D., and Singal, R. A finite time analysis of temporal difference learning with linear function approximation. In *Conference on learning theory*, pp. 1691–1692. PMLR, 2018.
- Borkar, V. S. *Stochastic approximation: a dynamical systems viewpoint*, volume 48. Springer, 2009.
- Bowling, M. and Veloso, M. Rational and convergent learning in stochastic games. In *International joint conference on artificial intelligence*, volume 17, pp. 1021–1026. Cite-seer, 2001.
- Boyan, J. A. Technical update: Least-squares temporal difference learning. *Machine learning*, 49(2):233–246, 2002.
- Bradtke, S. J. and Barto, A. G. Linear least-squares algorithms for temporal difference learning. *Machine learning*, 22(1):33–57, 1996.
- Brandfonbrener, D. and Bruna, J. Geometric insights into the convergence of nonlinear td learning. *arXiv preprint arXiv:1905.12185*, 2019.
- Brockman, G., Cheung, V., Pettersson, L., Schneider, J., Schulman, J., Tang, J., and Zaremba, W. Openai gym. *arXiv preprint arXiv:1606.01540*, 2016.
- Cai, Q., Yang, Z., Lee, J. D., and Wang, Z. Neural temporal difference and q learning provably converge to global optima. *Mathematics of Operations Research*, 2023.
- Cao, Y. and Gu, Q. Generalization bounds of stochastic gradient descent for wide and deep neural networks. *Advances in neural information processing systems*, 32, 2019.
- Cao, Y. and Gu, Q. Generalization error bounds of gradient descent for learning over-parameterized deep relu networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pp. 3349–3356, 2020.
- Cayci, S., Satpathi, S., He, N., and Srikant, R. Sample complexity and overparameterization bounds for temporal difference learning with neural network approximation. *IEEE Transactions on Automatic Control*, 2023.
- Dalal, G., Szörényi, B., Thoppe, G., and Mannor, S. Finite sample analyses for td (0) with function approximation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 32, 2018.
- Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- Du, S., Lee, J., Li, H., Wang, L., and Zhai, X. Gradient descent finds global minima of deep neural networks. In *International conference on machine learning*, pp. 1675–1685. PMLR, 2019.
- Du, S. S., Zhai, X., Póczos, B., and Singh, A. Gradient descent provably optimizes over-parameterized neural networks. *arXiv preprint arXiv:1810.02054*, 2018.
- Fan, J., Wang, Z., Xie, Y., and Yang, Z. A theoretical analysis of deep q-learning. In *Learning for dynamics and control*, pp. 486–489. PMLR, 2020.
- Fujimoto, S., Hoof, H., and Meger, D. Addressing function approximation error in actor-critic methods. In *International conference on machine learning*, pp. 1587–1596. PMLR, 2018.
- Godfrey, L. B. An evaluation of parametric activation functions for deep learning. In *2019 IEEE International Conference on Systems, Man and Cybernetics (SMC)*, pp. 3006–3011. IEEE, 2019.

- Jaakkola, T., Jordan, M., and Singh, S. Convergence of stochastic iterative dynamic programming algorithms. *Advances in neural information processing systems*, 6, 1993.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. *Advances in neural information processing systems*, 31, 2018.
- Ke, Z., Wen, Z., and Zhang, J. Provably efficient gauss-newton temporal difference learning method with function approximation. *arXiv preprint arXiv:2302.13087*, 2023.
- Konda, V. and Tsitsiklis, J. Actor-critic algorithms. *Advances in neural information processing systems*, 12, 1999.
- Kostrikov, I., Fergus, R., Tompson, J., and Nachum, O. Offline reinforcement learning with fisher divergence critic regularization. In *International Conference on Machine Learning*, pp. 5774–5783. PMLR, 2021.
- Lazaric, A., Ghavamzadeh, M., and Munos, R. Finite-sample analysis of lstd. In *ICML-27th International Conference on Machine Learning*, pp. 615–622, 2010.
- Levine, S., Kumar, A., Tucker, G., and Fu, J. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- Lillicrap, T. P., Hunt, J. J., Pritzel, A., Heess, N., Erez, T., Tassa, Y., Silver, D., and Wierstra, D. Continuous control with deep reinforcement learning. *arXiv preprint arXiv:1509.02971*, 2015.
- Littman, M. L. Markov games as a framework for multi-agent reinforcement learning. In *Machine learning proceedings 1994*, pp. 157–163. Elsevier, 1994.
- Liu, B., Liu, J., Ghavamzadeh, M., Mahadevan, S., and Petrik, M. Finite-sample analysis of proximal gradient td algorithms. *arXiv preprint arXiv:2006.14364*, 2020a.
- Liu, C., Zhu, L., and Belkin, M. On the linearity of large non-linear models: when and why the tangent kernel is constant. *Advances in Neural Information Processing Systems*, 33:15954–15964, 2020b.
- Maei, H., Szepesvári, C., Bhatnagar, S., Precup, D., Silver, D., and Sutton, R. S. Convergent temporal-difference learning with arbitrary smooth function approximation. *Advances in neural information processing systems*, 22, 2009.
- Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., and Riedmiller, M. Playing atari with deep reinforcement learning. *arXiv preprint arXiv:1312.5602*, 2013.
- Perkins, T. J. and Pendrith, M. D. On the existence of fixed points for q-learning and sarsa in partially observable domains. In *ICML*, pp. 490–497, 2002.
- Perolat, J., Piot, B., and Pietquin, O. Actor-critic fictitious play in simultaneous move multistage games. In *International Conference on Artificial Intelligence and Statistics*, pp. 919–928. PMLR, 2018.
- Prashanth, L., Korda, N., and Munos, R. Fast lstd using stochastic approximation: Finite time analysis and application to traffic control. In *Joint European conference on machine learning and knowledge discovery in databases*, pp. 66–81. Springer, 2014.
- Schulman, J., Levine, S., Abbeel, P., Jordan, M., and Moritz, P. Trust region policy optimization. In *International conference on machine learning*, pp. 1889–1897. PMLR, 2015.
- Sun, T., Li, D., and Wang, B. Finite-time analysis of adaptive temporal difference learning with deep neural networks. *Advances in Neural Information Processing Systems*, 35:19592–19604, 2022.
- Sutton, R. S. Learning to predict by the methods of temporal differences. *Machine learning*, 3(1):9–44, 1988.
- Sutton, R. S., McAllester, D., Singh, S., and Mansour, Y. Policy gradient methods for reinforcement learning with function approximation. *Advances in neural information processing systems*, 12, 1999.
- Sutton, R. S., Maei, H. R., Precup, D., Bhatnagar, S., Silver, D., Szepesvári, C., and Wiewiora, E. Fast gradient-descent methods for temporal-difference learning with linear function approximation. In *Proceedings of the 26th annual international conference on machine learning*, pp. 993–1000, 2009a.
- Sutton, R. S., Szepesvári, C., and Maei, H. R. A convergent $O(n)$ algorithm for off-policy temporal-difference learning with linear function approximation. *Advances in neural information processing systems*, 21(21):1609–1616, 2009b.
- Tagorti, M. and Scherrer, B. On the rate of convergence and error bounds for lstd (λ). In *International Conference on Machine Learning*, pp. 1521–1529. PMLR, 2015.
- Tesauro, G. et al. Temporal difference learning and td-gammon. *Communications of the ACM*, 38(3):58–68, 1995.

- Tian, H., Paschalidis, I., and Olshevsky, A. On the performance of temporal difference learning with neural networks. In *The Eleventh International Conference on Learning Representations*, 2022.
- Touati, A., Bacon, P.-L., Precup, D., and Vincent, P. Convergent tree backup and retrace with function approximation. In *International Conference on Machine Learning*, pp. 4955–4964. PMLR, 2018.
- Tsitsiklis, J. and Van Roy, B. Analysis of temporal-difference learning with function approximation. *Advances in neural information processing systems*, 9, 1996.
- Van Hasselt, H., Guez, A., and Silver, D. Deep reinforcement learning with double q-learning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 30, 2016.
- Wu, Y., Tucker, G., and Nachum, O. Behavior regularized offline reinforcement learning. *arXiv preprint arXiv:1911.11361*, 2019.
- Xu, P. and Gu, Q. A finite-time analysis of q-learning with neural network function approximation. In *International Conference on Machine Learning*, pp. 10555–10565. PMLR, 2020.
- Zou, S., Xu, T., and Liang, Y. Finite-sample analysis for sarsa with linear function approximation. *Advances in neural information processing systems*, 32, 2019.

A. Details of Section 3

A.1. Proof of (14)

$$\begin{aligned}
 & \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \boldsymbol{\theta}_*) \cdot \langle \nabla_{\boldsymbol{\theta}} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_*), \boldsymbol{\theta}' - \boldsymbol{\theta}_* \rangle \right] \\
 \stackrel{(i)}{=} & \mathbb{E}_{\mu, \pi} \left[\mathbb{E}_{\mathbb{P}} \left[\widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \boldsymbol{\theta}_*) \right] \cdot \langle \nabla_{\boldsymbol{\theta}} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_0), \boldsymbol{\theta}' - \bar{\boldsymbol{\theta}} \rangle \right] \\
 \stackrel{(ii)}{=} & \mathbb{E}_{\mu, \pi} \left[\mathbb{E}_{\mathbb{P}} \left[\widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \bar{\boldsymbol{\theta}}) \right] \cdot \langle \nabla_{\boldsymbol{\theta}} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_0), \boldsymbol{\theta}' - \bar{\boldsymbol{\theta}} \rangle \right] \\
 = & \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \bar{\boldsymbol{\theta}}) \cdot \langle \nabla_{\boldsymbol{\theta}} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_0), \boldsymbol{\theta}' - \bar{\boldsymbol{\theta}} \rangle \right]
 \end{aligned}$$

where (i) is because $\nabla_{\boldsymbol{\theta}} \widehat{Q}(\cdot; \boldsymbol{\theta}_*) = \widehat{Q}(\cdot; \boldsymbol{\theta}_0)$, the decomposition

$$\boldsymbol{\theta}' - \boldsymbol{\theta}_* = \boldsymbol{\theta}' - \bar{\boldsymbol{\theta}} + \bar{\boldsymbol{\theta}} - \boldsymbol{\theta}_* = \boldsymbol{\theta}' - \bar{\boldsymbol{\theta}} + (\bar{\boldsymbol{\theta}}_{\perp} - \boldsymbol{\theta}_{\perp})$$

the fact that $(\bar{\boldsymbol{\theta}}_{\perp} - \boldsymbol{\theta}_{\perp}) \in \mathcal{K}(\Sigma_{\pi})$, and (13), (ii) is because $(\bar{\boldsymbol{\theta}} - \boldsymbol{\theta}^*) \in \mathcal{K}(\Sigma_{\pi})$.

A.2. Proof of Theorem 3.6

Proof. Recall the definition of the semi-gradient in Section (6). We denote $\bar{\mathbf{g}}(\boldsymbol{\theta})$ as its expectation. Let $\bar{\mathbf{m}}(\boldsymbol{\theta})$ and $\mathbf{m}(\boldsymbol{\theta})$ also be the corresponding semi-gradients based on the linearized function $\widehat{Q}(\cdot; \boldsymbol{\theta})$, that is,

$$\begin{aligned}
 \mathbf{g}(\boldsymbol{\theta}^t) &= \Delta(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^t), \quad \bar{\mathbf{g}}(\boldsymbol{\theta}^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{g}(\boldsymbol{\theta}^t)] \\
 \mathbf{m}(\boldsymbol{\theta}^t) &= \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0), \quad \bar{\mathbf{m}}(\boldsymbol{\theta}^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{m}(\boldsymbol{\theta}^t)],
 \end{aligned}$$

where

$$\begin{aligned}
 \Delta(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t) &= Q(\mathbf{x}_t; \boldsymbol{\theta}^t) - (r(s_t, a_t) + \gamma \cdot Q(\mathbf{x}_{t+1}; \boldsymbol{\theta}^t)), \\
 \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t) &= \widehat{Q}(\mathbf{x}_t; \boldsymbol{\theta}^t) - (r(s_t, a_t) + \gamma \cdot \widehat{Q}(\mathbf{x}_{t+1}; \boldsymbol{\theta}^t)).
 \end{aligned}$$

To simplify the notation, let $\Delta_t := \Delta(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t)$ and $\widehat{\Delta}_t := \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}; \boldsymbol{\theta}^t)$. Recall the definition of the range space $\mathcal{R}(\Sigma_{\pi})$ and the kernel space $\mathcal{K}(\Sigma_{\pi})$. By Proposition 3.4, we know that $\mathbf{v}_1^{\top} \mathbf{v}_2 = 0$ for any vector $\mathbf{v}_1 \in \mathcal{R}(\Sigma_{\pi})$, $\mathbf{v}_2 \in \mathcal{K}(\Sigma_{\pi})$ thus $\langle \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}; \boldsymbol{\theta}^0), \boldsymbol{\theta}_{\perp} \rangle = 0$ for any feature map $\mathbf{x}$ and parameter $\boldsymbol{\theta}_{\perp} \in \mathcal{K}(\Sigma_{\pi})$. Then we can decompose $\|\boldsymbol{\theta}^{t+1} - \boldsymbol{\theta}_*^{t+1}\|^2$ as

$$\begin{aligned}
 \|\boldsymbol{\theta}^{t+1} - \boldsymbol{\theta}_*^{t+1}\|^2 &= \|\Pi_{S_{2\omega}}(\boldsymbol{\theta}^t - \eta_t \mathbf{g}(\boldsymbol{\theta}^t)) - \Pi_{S_{2\omega}}(\boldsymbol{\theta}_*^{t+1})\|^2 \\
 &\leq \|\boldsymbol{\theta}^t - \eta_t \mathbf{g}(\boldsymbol{\theta}^t) - \boldsymbol{\theta}_*^{t+1}\|^2 \\
 &= \|\boldsymbol{\theta}^t - \eta_t \mathbf{g}(\boldsymbol{\theta}^t) - \boldsymbol{\theta}_*^t + \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1}\|^2 \\
 &= \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2 + \eta_t^2 \|\mathbf{g}(\boldsymbol{\theta}^t)\|^2 + \|\boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1}\|^2 - 2\eta_t \langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \mathbf{g}(\boldsymbol{\theta}^t) \rangle \\
 &\quad - 2\eta_t \langle \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1}, \mathbf{g}(\boldsymbol{\theta}^t) \rangle + 2\eta_t \langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1} \rangle \\
 \stackrel{(i)}{=} &\|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2 - 2\eta_t \langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \mathbf{g}(\boldsymbol{\theta}^t) \rangle \\
 &\quad - 2\eta_t \langle \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1}, \mathbf{g}(\boldsymbol{\theta}^t) - \mathbf{m}(\boldsymbol{\theta}^t) \rangle + \eta_t^2 \|\mathbf{g}(\boldsymbol{\theta}^t)\|^2,
 \end{aligned}$$

where (i) follows

$$\langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1} \rangle = \langle \boldsymbol{\theta}_{\parallel}^t - \boldsymbol{\theta}_{\parallel}^*, \boldsymbol{\theta}_{\perp}^t - \boldsymbol{\theta}_{\perp}^{t+1} \rangle = 0,$$

and

$$\begin{aligned}
 \langle \boldsymbol{\theta}_*^t - \boldsymbol{\theta}_*^{t+1}, \mathbf{m}(\boldsymbol{\theta}^t) \rangle &= \langle \boldsymbol{\theta}_{\perp}^t - \boldsymbol{\theta}_{\perp}^{t+1}, \mathbf{m}(\boldsymbol{\theta}^t) \rangle \\
 &= \widehat{\Delta}_t \cdot \langle \boldsymbol{\theta}_{\perp}^t - \boldsymbol{\theta}_{\perp}^{t+1}, \nabla Q(\mathbf{x}_t; \boldsymbol{\theta}^0) \rangle = 0.
 \end{aligned}$$

Recall the stationarity condition (11), for any $t \in \{1, 2, \dots, T\}$,

$$\begin{aligned}
 0 &\leq \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(x_t, x_{t+1}; \theta^*) \langle \nabla_{\theta} \widehat{Q}(x_t; \theta^*), \theta^t - \theta^* \rangle \right] \\
 &= \langle \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(x_t, x_{t+1}; \theta_*^t) \nabla_{\theta} \widehat{Q}(x_t; \theta_*^t) \right], \theta^t - \theta^* \rangle \\
 &\stackrel{(i)}{=} \langle \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\widehat{\Delta}(x_t, x_{t+1}; \theta_*^t) \nabla_{\theta} \widehat{Q}(x_t; \theta_*^t) \right], \theta^t - \theta_*^t \rangle = \langle \bar{\mathbf{m}}(\theta_*^t), \theta^t - \theta_*^t \rangle,
 \end{aligned}$$

where (i) is the same as the proof in Section A.1. Therefore,

$$\begin{aligned}
 &\|\theta^{t+1} - \theta_*^{t+1}\|^2 \\
 &\leq \|\theta^t - \theta_*^t\|^2 - 2\eta_t \langle \theta^t - \theta_*^t, \mathbf{g}(\theta^t) \rangle - 2\eta_t \langle \theta_*^t - \theta_*^{t+1}, \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t) \rangle + \eta_t^2 \|\mathbf{g}(\theta^t)\|^2 \\
 &= \|\theta^t - \theta_*^t\|^2 - 2\eta_t \langle \theta^t - \theta_*^t, \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t) \rangle - 2\eta_t \langle \theta^t - \theta_*^t, \mathbf{m}(\theta^t) - \bar{\mathbf{m}}(\theta^t) \rangle \\
 &\quad - 2\eta_t \langle \theta^t - \theta_*^t, \bar{\mathbf{m}}(\theta^t) \rangle - 2\eta_t \langle \theta_*^t - \theta_*^{t+1}, \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t) \rangle + \eta_t^2 \|\mathbf{g}(\theta^t)\|^2 \\
 &\stackrel{(i)}{\leq} \|\theta^t - \theta_*^t\|^2 + \eta_t^2 \underbrace{\|\mathbf{g}(\theta^t)\|^2}_{\text{I}_1: \text{Gradient Bound}} - 2\eta_t \underbrace{\langle \theta^t - \theta_*^{t+1}, \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t) \rangle}_{\text{I}_2: \text{Gradient Gap}} \\
 &\quad - 2\eta_t \underbrace{\langle \theta^t - \theta_*^t, \mathbf{m}(\theta^t) - \bar{\mathbf{m}}(\theta^t) \rangle}_{\text{I}_3: \text{Markov Sampling Error}} - 2\eta_t \underbrace{\langle \theta^t - \theta_*^t, \bar{\mathbf{m}}(\theta^t) - \bar{\mathbf{m}}(\theta_*^t) \rangle}_{\text{I}_4: \text{Gradient Decent}}, \tag{18}
 \end{aligned}$$

where (i) follows $\langle \theta^t - \theta_*^t, \bar{\mathbf{m}}(\theta_*^t) \rangle \geq 0$ for any $0 \leq t \leq T-1$.

Next, we analyze the upper bounds of I_1 , I_2 , I_3 and I_4 item by item. To simplify the notation, let $\{C_i > 0\}_{i=1, \dots, 7}$ be universal constants in this section. We set $\omega = C_1$ and $\delta \in (0, 1)$. By Lemma D.5, we have

$$\|\mathbf{g}(\theta^t)\|^2 \leq C_2 \sqrt{\log(T/\delta)} \tag{19}$$

and

$$\begin{aligned}
 \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[|\langle \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t), \theta^t - \theta_*^{t+1} \rangle| \mid \theta^0 \right] &\leq \left(C_3 \omega m^{-\frac{1}{2}} \sqrt{\log(T/\delta)} + C_4 m^{-\frac{1}{2}} \right) \|\theta^t - \theta_*^{t+1}\| \\
 &\stackrel{(i)}{\leq} C_5 m^{-1/2} \sqrt{\log(T/\delta)}, \tag{20}
 \end{aligned}$$

with probability at least $1 - 2\delta - 2\exp(-C_6 m)$, where (i) follows $\omega = C_1$ and

$$\begin{aligned}
 \|\theta^t - \theta_*^{t+1}\| &\leq \|\theta^t - \theta^0\| + \|\theta^0 - \theta_*^0\| + \|\theta_*^0 - \theta_*^{t+1}\| \\
 &= \|\theta^t - \theta^0\| + \|\theta^0 - \theta_*^0\| + \|\theta_*^0 - \theta_*^{t+1}\| \\
 &\leq \|\theta^t - \theta^0\| + \|\theta^0 - \theta_*^0\| + \|\theta^0 - \theta^{t+1}\| \leq 3\omega.
 \end{aligned}$$

Thus (19) and (20) provide upper bounds on I_1 and I_2 , respectively. The next lemma provides an estimate of the Markov sampling error.

Lemma A.1. Suppose the learning rate sequence $\{\eta_0, \eta_1, \dots, \eta_T\}$ is non-increasing. Under Assumption 3.2, it holds that

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\langle \mathbf{m}(\theta^t) - \bar{\mathbf{m}}(\theta^t), \theta^t - \theta^* \rangle \mid \theta^0 \right] \leq C_7 (\log(T/\delta) + C_1^2) \tau^* \eta_{\max\{0, t-\tau^*\}}, \tag{21}$$

for any fixed $t \leq T$, where

$$\tau^* = \min \{t = 0, 1, 2, \dots \mid \kappa \rho^t \leq \eta_T\}$$

is the mixing time of the Markov chain $\{s_t, a_t\}_{t=0, 1, \dots}$.

Proof. We adopt the proof framework outlined in Lemma 6.2 of Xu & Gu (2020). However, variations in the neural network settings lead to differences in the norms of gradients and parameters, thereby resulting in slight variations in the results. Thereby we have

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} [\langle \mathbf{m}(\boldsymbol{\theta}^t) - \bar{\mathbf{m}}(\boldsymbol{\theta}^t), \boldsymbol{\theta}^t - \boldsymbol{\theta}^* \rangle \mid \boldsymbol{\theta}^0] \leq C_7 (\log(T/\delta) + \omega^2) \tau^* \eta_{\max\{0, t-\tau^*\}}.$$

□

Looking back at the definitions of λ_0 and Σ_π , and the discussion in Section 3.2, we derive Lemmas A.2 and A.3 to estimate $\mathbf{I}_4$.

Lemma A.2. *Let λ_0 as the minimum nonzero singular value of Σ_π . For any $\boldsymbol{\theta} \in \mathcal{R}(\Sigma_\pi)$, we have*

$$\boldsymbol{\theta}^\top \Sigma_\pi \boldsymbol{\theta} \geq \lambda_0 \|\boldsymbol{\theta}\|_2^2.$$

Lemma A.3. *Under Assumption 3.3, we have that*

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} [\langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \bar{\mathbf{m}}(\boldsymbol{\theta}^t) - \bar{\mathbf{m}}(\boldsymbol{\theta}_*^t) \rangle \mid \boldsymbol{\theta}^0] \geq (1 - \gamma) \lambda_0 \cdot \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2. \quad (22)$$

Proof. Define $d \sim \mu \times \pi$. To begin with, the Bellman operator $\mathcal{T}^\pi$ is a γ -contraction with ℓ_2 -norm since d is the stationary distribution of (s, a) corresponding to the policy π . In details, consider

$$\begin{aligned} \mathbb{E}_{(s,a) \sim d} [\mathcal{T}^\pi Q_1(\mathbf{x}) - \mathcal{T}^\pi Q_2(\mathbf{x})]^2 &= \gamma^2 \mathbb{E}_{(s,a) \sim d} \left[\mathbb{E} \left[(Q_1(\mathbf{x}') - Q_2(\mathbf{x}'))^2 \mid s' \sim \mathbb{P}(\cdot \mid s, a), a' \sim \pi(\cdot \mid s') \right] \right] \\ &\stackrel{(i)}{\leq} \gamma^2 \mathbb{E}_{(s,a) \sim d} [(Q_1(\mathbf{x}) - Q_2(\mathbf{x}))^2], \end{aligned} \quad (23)$$

where (i) follows that $\mathbf{x}$ and $\mathbf{x}'$ have the same stationary distribution. To simplify the notation, we denote $\mathbb{E}[\cdot]$ as $\mathbb{E}_{\mu, \pi, \mathbb{P}}[\cdot]$ in the proof of this lemma. Then we compute

$$\begin{aligned} &\mathbb{E} [\langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \bar{\mathbf{m}}(\boldsymbol{\theta}^t) - \bar{\mathbf{m}}(\boldsymbol{\theta}_*^t) \rangle \mid \boldsymbol{\theta}^0] \\ &= \mathbb{E} \left[\left(\widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \boldsymbol{\theta}^t) - \widehat{\Delta}(\mathbf{x}, \mathbf{x}'; \boldsymbol{\theta}_*^t) \right) \langle \nabla Q(\mathbf{x}; \boldsymbol{\theta}^0), \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t \rangle \mid \boldsymbol{\theta}^0 \right] \\ &= \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_*^t) \right) \langle \nabla Q(\mathbf{x}; \boldsymbol{\theta}^0), \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t \rangle \mid \boldsymbol{\theta}^0 \right] \\ &\quad - \gamma \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}_*^t) \right) \langle \nabla Q(\mathbf{x}; \boldsymbol{\theta}^0), \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t \rangle \mid \boldsymbol{\theta}^0 \right] \\ &= \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right] - \gamma \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right) \cdot \left(\widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}_*^t) \right) \right] \\ &\stackrel{(i)}{\geq} \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right] - \gamma \sqrt{\mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right]} \cdot \sqrt{\mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}'; \boldsymbol{\theta}_*^t) \right)^2 \right]} \\ &\stackrel{(i)}{\geq} (1 - \gamma) \mathbb{E} \left[\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right]^2 \\ &= (1 - \gamma) (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\ &\stackrel{(ii)}{\geq} (1 - \gamma) \lambda_0 \cdot \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2, \end{aligned} \quad (24)$$

where (i) follows the Cauchy-Schwarz inequality, (ii) follows (23), and (iii) follows $\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t = \boldsymbol{\theta}_\parallel^t - \bar{\boldsymbol{\theta}}_\parallel \in \mathcal{R}(\Sigma_\pi)$ and Lemma A.2, which provides λ_0 -strong convexity. Thus we complete the proof of Lemma A.3. □

Given $\boldsymbol{\theta}^0$, taking the expectation on both sides of (18) and plugging (19)~(22) into (18) yields that

$$\begin{aligned} \mathbb{E}_{\mu, \pi, \mathbb{P}} [\|\boldsymbol{\theta}^{t+1} - \boldsymbol{\theta}_*^{t+1}\|^2 \mid \boldsymbol{\theta}^0] &\leq (1 - 2\eta_t(1 - \gamma)\lambda_0) \mathbb{E} [\|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2 \mid \boldsymbol{\theta}^0] + C_2 \eta_t^2 \\ &\quad + 2\eta_t C_5 m^{-1/2} \sqrt{\log(T/\delta)} + 2\eta_t C_7 (\log(T/\delta) + C_1^2) \tau^* \eta_{\max\{0, t-\tau^*\}}. \end{aligned}$$

We choose $\eta_t = \frac{1}{2(1-\gamma)\lambda_0(t+1)}$ and have that

$$(1-\gamma)\lambda_0(t+1)\mathbb{E}_{\mu,\pi,\mathbb{P}}[\|\boldsymbol{\theta}^{t+1} - \boldsymbol{\theta}_*^{t+1}\|^2 \mid \boldsymbol{\theta}^0] \leq (1-\gamma)\lambda_0 t \mathbb{E}[\|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2 \mid \boldsymbol{\theta}^0] + C_2\eta_t \\ + C_5m^{-1/2}\sqrt{\log(T/\delta)} + C_7(\log(T/\delta) + C_1^2)\tau^*\eta_{\max\{0,t-\tau^*\}}.$$

Summing (25) from $t = 0, 1, \dots, T-1$ yields that

$$\mathbb{E}_{\mu,\pi,\mathbb{P}}[\|\boldsymbol{\theta}^T - \boldsymbol{\theta}_*^T\|^2 \mid \boldsymbol{\theta}^0] \leq \frac{1}{(1-\gamma)\lambda_0 T} \sum_{t=0}^{T-1} (C_2\eta_t + C_5m^{-1/2}\sqrt{\log(T/\delta)} \\ + C_7(\log(T/\delta) + C_1^2)\tau^*\eta_{\max\{0,t-\tau^*\}}) \\ \leq \frac{C_2(\log T + 1)}{2(1-\gamma)^2\lambda_0^2 T} + \frac{C_5m^{-1/2}\sqrt{\log(T/\delta)}}{(1-\gamma)\lambda_0} + \frac{C_8\tau^*(\log(T/\delta) + 1)\log T}{2(1-\gamma)^2\lambda_0^2 T}.$$

Therefore, according to the gradient bound (19), we have

$$\mathbb{E}_{\mu,\pi} \left[\left(\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^T) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) \right)^2 \mid \boldsymbol{\theta}^0 \right] = \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^T) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}_*^T) \right)^2 \mid \boldsymbol{\theta}^0 \right] \\ \leq C_3^2 m \mathbb{E}[\|\boldsymbol{\theta}^T - \boldsymbol{\theta}_*^T\|^2 \mid \boldsymbol{\theta}^0] \\ \leq \frac{C_2^3(\log T + 1)}{2(1-\gamma)^2\lambda_0^2 T} + \frac{C_2^2 C_5 m^{-1/2} \sqrt{\log(T/\delta)}}{(1-\gamma)\lambda_0} \\ + \frac{C_2^2 C_8 \tau^* (\log(T/\delta) + 1) \log T}{2(1-\gamma)^2\lambda_0^2 T}$$

with probability at least $1 - 2\delta - 2L \exp(-C_6 m)$. Let $\tilde{C}_1 = \max\{1, C_1\}$, $\tilde{C}_2 = C_6$, $\tilde{C}_3 = \frac{C_2^3}{2}$, $\tilde{C}_4 = \frac{C_2^2 C_5}{2}$, and $\tilde{C}_5 = \frac{C_2^2 C_8}{2}$, and we complete the proof. $\square$

A.3. Proof of Theorem 3.7

Proof. Let $(s, a) \sim \mu \times \pi =: d$. To simplify the notation, we denote $\mathbb{E}[\cdot]$ as $\mathbb{E}_{(s,a) \sim d}[\cdot]$ in this subsection. Note that

$$\mathbb{E} \left[\left(Q(\mathbf{x}; \boldsymbol{\theta}^T) - Q^*(s, a) \right)^2 \mid \boldsymbol{\theta}^0 \right] \leq 3\mathbb{E} \left[\left(Q(\mathbf{x}; \boldsymbol{\theta}^T) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^T) \right)^2 \mid \boldsymbol{\theta}^0 \right] \\ + 3\mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^T) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) \right)^2 \mid \boldsymbol{\theta}^0 \right] + 3\mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - Q^*(s, a) \right)^2 \mid \boldsymbol{\theta}^0 \right].$$

By Lemma D.4, we have

$$\mathbb{E} \left[\left(Q(\mathbf{x}; \boldsymbol{\theta}^T) - \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^T) \right)^2 \mid \boldsymbol{\theta}^0 \right] \leq C_8 m^{-1} \quad (25)$$

with probability at least $1 - \delta$. Recall that $\widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*)$ is the fixed point of $\Pi_{\mathcal{F}_{\omega,m}} \mathcal{T}$ and $Q^*(s, a)$ is the fixed point of $\mathcal{T}$. We define the ℓ_2 -norm $\|f(s, a)\|_d^2 = \mathbb{E}_{(s,a) \sim d}[f(s, a)^2]$. Thus

$$\left\| \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - Q^*(s, a) \right\|_d = \left\| \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) + \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) - Q^*(s, a) \right\|_d \\ \stackrel{(i)}{=} \left\| \Pi_{\mathcal{F}_{\omega,m}} \mathcal{T} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - \Pi_{\mathcal{F}_{\omega,m}} \mathcal{T} Q^*(s, a) + \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) - Q^*(s, a) \right\|_d \\ \leq \left\| \Pi_{\mathcal{F}_{\omega,m}} \mathcal{T} \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - \Pi_{\mathcal{F}_{\omega,m}} \mathcal{T} Q^*(s, a) \right\|_d + \left\| \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) - Q^*(s, a) \right\|_d \\ \stackrel{(ii)}{\leq} \gamma \left\| \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - Q^*(s, a) \right\|_d + \left\| \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) - Q^*(s, a) \right\|_d,$$

where (i) is due to the properties of the fixed point, and (ii) is due to $\Pi_{\mathcal{F}_{\omega,m}} \mathcal{T}$ is γ -contractive on the ∞ -norm. This further means that

$$\left\| \widehat{Q}(\mathbf{x}; \boldsymbol{\theta}^*) - Q^*(s, a) \right\|_d^2 \leq \frac{1}{(1-\gamma)^2} \left\| \Pi_{\mathcal{F}_{\omega,m}} Q^*(s, a) - Q^*(s, a) \right\|_d^2. \quad (26)$$

Plugging (25) and (26) into (25) and using Theorem 3.6, we complete the proof. $\square$

B. Convergence Results of Neural Q-learning

B.1. Neural Q-Learning Algorithm

For neural Q-learning, let us redefine some of the above notations. Let the optimal Q-function be $Q^*(s, a) = \sup_{\pi} Q^{\pi}(s, a)$ for all state action pairs (s, a) , then the optimal sequence of actions that maximizes the expected cumulative reward will follow $a_t = \operatorname{argmax}_{a' \in \mathcal{A}} Q^*(s_t, a')$, $t \geq 0$. Therefore, to obtain a near-optimal policy, it is sufficient to find some $\hat{Q}$ that approximates Q^* well. Define the Bellman optimality operator $\mathcal{T}$ as

$$\mathcal{T}Q(s, a) := r(s, a) + \gamma \mathbb{E} \left[\max_{a'} Q(s', a') \mid s' \sim \mathbb{P}(\cdot \mid s, a) \right],$$

for any (s, a) . Let us remain the definition of the local linearization function class $\mathcal{F}_{\omega, m}$ introduced in (7). Consider the MSPBE minimization problem with multi-layer neural network approximation:

$$\min_{\theta \in S_{\omega}} \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[(Q(\mathbf{x}; \theta) - \Pi_{\mathcal{F}_{\omega, m}} \mathcal{T}Q(\mathbf{x}; \theta))^2 \right].$$

Then the projected neural Q-learning algorithm can be written as follows:

$$\theta^{t+1} = \Pi_{S_{\omega}} \left(\theta^t - \eta_t g(\theta^t) \right), \quad \text{with} \quad g(\theta^t) = \Delta(s_t, a_t, s_{t+1}; \theta^t) \cdot \nabla_{\theta} Q(\phi(s_t, a_t); \theta^t) \quad (27)$$

where

$$\Delta(s, a, s'; \theta^t) = Q(\phi(s, a); \theta^t) - \left(r(s, a) + \gamma \max_{b \in \mathcal{A}} Q(\phi(s', b); \theta^t) \right). \quad (28)$$

The algorithm details can be described by Algorithm 2 as follows.

Algorithm 2 Neural Q-Learning with Markovian Sampling

Input: A learning policy π , a discount factor $\gamma \in (0, 1)$, a sequence of learning rates $\{\eta_t\}_{t \geq 0}$, a maximum iteration number T , a projection radius $\omega > 0$, a Q network with architecture (4).

Initialization: Generate each entry of $\mathbf{W}_l^0$ independently from $\mathcal{N}(0, 1)$, for $l = 1, 2, \dots, L$, and each entry of $\mathbf{b}$ independently from $\text{Unif}\{-1, +1\}$.

for $t = 0, 1, \dots, T - 1$ **do**

Sample (s_t, a_t, r_t, s_{t+1}) from the learning policy π with $a_t \sim \pi(\cdot \mid s_t)$.

Compute the TD error Δ_t by (28).

Update θ^{t+1} by the projected stochastic semi-gradient step (27).

end for

Output: θ^T .

B.2. Global Convergence

Similar to Section 3, we define the function class $\mathcal{F}_{\omega, m}$ as a collection of all local linearization of $Q(\mathbf{x}; \theta)$ at the initial point θ^0 :

$$\mathcal{F}_{\omega, m} := \left\{ \hat{Q}(\mathbf{x}; \theta) = Q(\mathbf{x}; \theta^0) + \langle \nabla_{\theta} Q(\mathbf{x}; \theta^0), \theta - \theta^0 \rangle, \theta \in S_{\omega} \right\}.$$

Let $\hat{Q}(\cdot; \theta^*) \in \mathcal{F}_{\omega, m}$, and $\hat{\Delta}(s, a, s'; \theta)$ has the same structure as $\Delta(s, a, s'; \theta)$ expect that the function $Q(\cdot; \theta)$ is replaced by $\hat{Q}(\cdot; \theta)$. The stationary point θ^* satisfies $\hat{Q}(\mathbf{x}; \theta^*) = \Pi_{\mathcal{F}_{\omega, m}} \mathcal{T} \hat{Q}(\mathbf{x}; \theta^*)$ for neural Q-learning. We redefine Ξ_{β} by replacing the Bellman operator $\mathcal{T}^{\pi}$ in Section 3 with the Bellman optimality operator $\mathcal{T}$. A point $\theta^* \in \Xi_{\omega}$ if and only if

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\hat{\Delta}(s, a, s'; \theta^*) \langle \nabla_{\theta} \hat{Q}(\phi(s, a); \theta^*), \theta - \theta^* \rangle \right] \geq 0.$$

The maximum operator introduced by the Bellman optimality operator significantly sophisticates the analysis. Let us remain the definition of Σ_{π} in (12), and we define $\Sigma_{\pi}^*(\theta)$ as follows:

$$\mathbb{E}_{\mu, \pi} \left[\nabla_{\theta} Q(\phi(s, a_{\max}^{\theta}); \theta^0) \nabla_{\theta} Q(\phi(s, a_{\max}^{\theta}); \theta^0)^{\top} \right], \quad (29)$$

where $a_{\max}^{\theta} = \operatorname{argmax}_{a \in \mathcal{A}} |\langle \nabla_{\theta} Q(s, a; \theta^0), \theta \rangle|$. To facilitate the analysis of neural Q-learning, we further assume the following regularity condition introduced by (Xu & Gu, 2020).

Assumption B.1. $\exists \nu \in (0, 1)$ such that $(1 - \nu)^2 \Sigma_\pi - \gamma^2 \Sigma_\pi^*(\theta) \succeq 0$ for any θ^0 and $\theta \in S_\omega$.

The original version of this assumption comes from (Xu & Gu, 2020), which requires a strict positive definite condition: $(1 - \nu)^2 \Sigma_\pi - \gamma^2 \Sigma_\pi^*(\theta) \succ 0$. Under this additional assumption, (Xu & Gu, 2020) obtained an $\tilde{O}(\epsilon^{-2})$ sample complexity for neural Q-learning. A similar complexity result was also derived in (Cai et al., 2023) under a similar regularity condition on the learning policy π . At this time, we relax it to the positive semi-definiteness ($\succeq 0$) and provide a convergence result of neural Q-learning. See Theorem B.2.

Theorem B.2. Suppose Assumptions 3.1, 3.2 and B.1 hold. We set $\omega = \tilde{C}_1$ and the learning rate $\eta_t = \frac{1}{2\nu\lambda_0(t+1)}$. If the feature map $\|\phi(s, a)\| = 1$ for each state-action pair (s, a) and the network width $m \geq m^*$, then the output θ^T of neural Q-learning algorithm (i.e. (27)) satisfies

$$\mathbb{E} \left[(\hat{Q}(x; \theta^T) - \hat{Q}(x; \theta^*))^2 \mid \theta^0 \right] \leq \frac{\tilde{C}_3(\log T + 1)}{\nu^2 \lambda_0^2 T} + \frac{\tilde{C}_4 m^{-1/2}}{\nu \lambda_0} \cdot \sqrt{\log(T/\delta)} + \frac{\tilde{C}_5 \tau^* (\log(T/\delta) + 1) \log T}{\nu^2 \lambda_0^2 T},$$

with probability at least $1 - 2\delta - 2L \exp(-\tilde{C}_2 m)$, where τ^* is the mixing time of Markov chain in Assumption 3.2, and $\tilde{C}_1, \dots, \tilde{C}_5 > 0$ are universal constants.

Proof. For a little notation abuse, we redefine

$$\begin{aligned} \mathbf{g}(\theta^t) &= \Delta(s_t, a_t, s'_t, \theta^t) \cdot \nabla_{\theta} Q(x_t; \theta^t), \quad \bar{\mathbf{g}}(\theta^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{g}(\theta^t)] \\ \mathbf{m}(\theta^t) &= \hat{\Delta}(s_t, a_t, s'_t, \theta^t) \cdot \nabla_{\theta} Q(x_t; \theta^0), \quad \bar{\mathbf{m}}(\theta^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{m}(\theta^t)], \end{aligned}$$

where

$$\begin{aligned} \Delta(s_t, a_t, s'_t, \theta^t) &= Q(x_t; \theta^t) - \left(r(s_t, a_t) + \gamma \max_{b \in \mathcal{A}} Q(\phi(s_{t+1}, b); \theta^t) \right), \\ \hat{\Delta}(s_t, a_t, s'_t, \theta^t) &= \hat{Q}(x_t; \theta^t) - \left(r(s_t, a_t) + \gamma \max_{b \in \mathcal{A}} \hat{Q}(\phi(s_{t+1}, b); \theta^t) \right). \end{aligned}$$

Let $\Delta_t = \Delta(s_t, a_t, s'_t, \theta^t)$ and $\hat{\Delta}_t = \hat{\Delta}(s_t, a_t, s'_t, \theta^t)$. Similarly, (18) can be derived in neural Q-learning. To estimate the terms $\mathbf{I}_1 \sim \mathbf{I}_4$, we can apply Lemmas D.5 and A.1. However, due to the utilization of the Bellman optimality operator in neural Q-learning, some modifications based on Lemma A.3 are required.

Lemma B.3. Under Assumption B.1, we have that

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} [\langle \theta^t - \theta_*^t, \bar{\mathbf{m}}(\theta^t) - \bar{\mathbf{m}}(\theta_*^t) \rangle \mid \theta^0] \geq \nu \lambda_0 \cdot \|\theta^t - \theta_*^t\|^2. \quad (30)$$

Proof. To simplify the notation, we denote $\mathbb{E}[\cdot]$ as $\mathbb{E}_{\mu, \pi, \mathbb{P}}[\cdot]$ in the proof of this lemma. Define $\hat{Q}^\#(s; \theta) := \max_{a \in \mathcal{A}} \hat{Q}(\phi(s, a); \theta)$. Then we have

$$\begin{aligned} & \mathbb{E} [\langle \theta^t - \theta_*^t, \bar{\mathbf{m}}(\theta^t) - \bar{\mathbf{m}}(\theta_*^t) \rangle \mid \theta^0] \\ &= \mathbb{E} \left[\left(\hat{\Delta}(s_t, a_t, s'_t, \theta^t) - \hat{\Delta}(s_t, a_t, s'_t, \theta_*^t) \right) \langle \nabla Q(x; \theta^0), \theta^t - \theta_*^t \rangle \mid \theta^0 \right] \\ &= \mathbb{E} \left[\left(\hat{Q}(x; \theta^t) - \hat{Q}(x; \theta_*^t) \right) \langle \nabla Q(x; \theta^0), \theta^t - \theta_*^t \rangle \mid \theta^0 \right] \\ & \quad - \gamma \mathbb{E} \left[\left(\hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*^t) \right) \langle \nabla Q(x; \theta^0), \theta^t - \theta_*^t \rangle \mid \theta^0 \right] \\ &= \mathbb{E} \left[\left(\hat{Q}(x, \theta^t) - \hat{Q}(x, \theta_*^t) \right)^2 \right] - \gamma \mathbb{E} \left[\left(\hat{Q}(x, \theta^t) - \hat{Q}(x, \theta_*^t) \right) \cdot \left(\hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*^t) \right) \right]. \end{aligned} \quad (31)$$

For the second term of (31), we consider

$$\begin{aligned}
 \mathbb{E} \left[\left(\widehat{Q}^\#(s; \boldsymbol{\theta}^t) - \widehat{Q}^\#(s; \boldsymbol{\theta}_*^t) \right)^2 \right] &\leq \mathbb{E} \left[\max_{a \in \mathcal{A}} \left| \widehat{Q}(s, a; \boldsymbol{\theta}^t) - \widehat{Q}(s, a; \boldsymbol{\theta}_*^t) \right|^2 \right] \\
 &\stackrel{(i)}{=} \mathbb{E} \left[\max_{a \in \mathcal{A}} \left| \widehat{Q}(s, a; \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \right|^2 \right] \\
 &= (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi^* (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\
 &\stackrel{(ii)}{\leq} \frac{(1-\nu)^2}{\gamma^2} (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi \cdot (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\
 &\stackrel{(iii)}{=} \frac{(1-\nu)^2}{\gamma^2} \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right], \tag{32}
 \end{aligned}$$

where (i) and (iii) follow that $\widehat{Q}(\mathbf{x}; \cdot)$ is linear, and (ii) follows Assumption B.1. Therefore,

$$\begin{aligned}
 &\mathbb{E} \left[\langle \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t, \bar{\mathbf{m}}(\boldsymbol{\theta}^t) - \bar{\mathbf{m}}(\boldsymbol{\theta}_*^t) \rangle \mid \boldsymbol{\theta}^0 \right] \\
 &= \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right] - \gamma \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right) \cdot \left(\widehat{Q}^\#(s; \boldsymbol{\theta}^t) - \widehat{Q}^\#(s; \boldsymbol{\theta}_*^t) \right) \right] \\
 &\geq \mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right] - \gamma \sqrt{\mathbb{E} \left[\left(\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right)^2 \right]} \cdot \sqrt{\mathbb{E} \left[\left(\widehat{Q}^\#(s; \boldsymbol{\theta}^t) - \widehat{Q}^\#(s; \boldsymbol{\theta}_*^t) \right)^2 \right]} \\
 &\stackrel{(i)}{\geq} \left(1 - \gamma \cdot \frac{1-\nu}{\gamma} \right) \mathbb{E} \left[\widehat{Q}(\mathbf{x}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}, \boldsymbol{\theta}_*^t) \right] \\
 &= \nu \cdot (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\
 &\stackrel{(ii)}{\geq} \nu \lambda_0 \cdot \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2,
 \end{aligned}$$

where (i) follows (32), and (ii) follows $\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t = \boldsymbol{\theta}_\parallel^t - \boldsymbol{\theta}^* \in \mathcal{R}(\Sigma_\pi)$ and Lemma A.2, which provides λ_0 -strong convexity. $\square$

Now given $\boldsymbol{\theta}^0$, we can deduce that

$$\begin{aligned}
 \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\|\boldsymbol{\theta}^{t+1} - \boldsymbol{\theta}_*^{t+1}\|^2 \mid \boldsymbol{\theta}^0 \right] &\leq (1 - 2\eta_t \nu \lambda_0) \mathbb{E} \left[\|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t\|^2 \mid \boldsymbol{\theta}^0 \right] + C_1 \eta_t^2 \\
 &\quad + 2\eta_t C_2 m^{-1/2} \sqrt{\log(T/\delta)} + 2\eta_t C_3 (\log(T/\delta) + C_1^2) \tau^* \eta_{\max\{0, t-\tau^*\}},
 \end{aligned}$$

with probability at least $1 - 2\delta - 2L \exp(-C_4 m)$, where $\{C_i > 0\}_{i=1, \dots, 4}$ are universal constants in this subsection. Choosing $\eta_t = \frac{1}{2\nu \lambda_0 (t+1)}$ can derive the similar results as (25). This suggests that we can utilize the techniques outlined in Section A.2 to finalize the remaining proof of Theorem B.2. As a result, we conclude the proof. $\square$

C. Details of Section 4

We formally describe the minimax neural Q-learning method in Algorithm 3.

Algorithm 3 Minimax Neural Q-Learning with Gaussian Initialization

Input: A learning policy pair $\pi = (\pi^1, \pi^2)$, a discount factor $\gamma \in (0, 1)$, a sequence of learning rates $\{\eta_t\}_{t \geq 0}$, a maximum iteration number T , a projection radius $\omega > 0$, a Q network with architecture (4).

Initialization: Generate each entry of $\mathbf{W}_l^0$ independently from $\mathcal{N}(0, 1)$, for $l = 1, 2, \dots, L$, and each entry of $\mathbf{b}$ independently from $\text{Unif}\{-1, +1\}$.

for $t = 0, 1, \dots, T - 1$ **do**

Sample $(s_t, a_t^1, a_t^2, r_t, s_{t+1})$ from the learning policy pair π with $a_t^1 \sim \pi^1(\cdot | s_t)$, $a_t^2 \sim \pi^2(\cdot | s_t)$.

Compute the TD error Δ_t by (17).

Update θ^{t+1} by the projected stochastic semi-gradient step (16).

end for

Output: θ^T .

C.1. Proof of Theorem 4.2

The proof of Theorem 4.2 is similar to Sections A.2 and B.2. However, due to the difference in Bellman operators, we still need to make some modifications to Lemma A.3 or Lemma B.3. See Lemma C.1.

Lemma C.1. *Under Assumption 4.1, we have that*

$$\mathbb{E}_{\mu, \pi, \mathbb{P}} [\langle \theta^t - \theta_*, \bar{\mathbf{m}}(\theta^t) - \bar{\mathbf{m}}(\theta_*) \rangle | \theta^0] \geq \nu \lambda_0 \cdot \|\theta^t - \theta_*\|^2. \quad (33)$$

Proof. To simplify the notation, we denote $\mathbb{E}[\cdot]$ as $\mathbb{E}_{\mu, \pi, \mathbb{P}}[\cdot]$ in the proof of this lemma. Define $\hat{Q}^\#(s; \theta) := \max_{a^1 \in \mathcal{A}} \min_{a^2 \in \mathcal{A}} \hat{Q}(\phi(s, a^1, a^2); \theta)$. Define the sets $\mathcal{S}_+ = \{s : \hat{Q}^\#(s; \theta^t) > \hat{Q}^\#(s; \theta_*)\}$ and $\mathcal{S}_- = \mathcal{S} / \mathcal{S}_+$. For each $s \in \mathcal{S}_+$,

$$\begin{aligned} \hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*) &= \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t \rangle - \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta_* \rangle \\ &= \left(\langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t \rangle - \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta_* \rangle \right) - \\ &\quad \left(\langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta_* \rangle - \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta^t - \theta_* \rangle \right) \\ &\leq \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t - \theta_* \rangle \end{aligned}$$

and

$$\begin{aligned} \hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*) &= \left(\langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t \rangle - \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta_* \rangle \right) - \\ &\quad \left(\langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t - \theta_* \rangle - \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta^t - \theta_* \rangle \right) \\ &\geq \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t - \theta_* \rangle. \end{aligned}$$

In the same way, for each $s \in \mathcal{S}_-$,

$$\begin{aligned} \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t - \theta_* \rangle &\leq \hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*) \\ &\leq \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta^t - \theta_* \rangle. \end{aligned}$$

Therefore,

$$\begin{aligned} |\hat{Q}^\#(s; \theta^t) - \hat{Q}^\#(s; \theta_*)| &\leq \max \left\{ \left| \langle \nabla_{\theta} Q(\phi(s, a_{\theta^t}^1, a_{\theta^t}^2); \theta^0), \theta^t - \theta_* \rangle \right|, \right. \\ &\quad \left. \left| \langle \nabla_{\theta} Q(\phi(s, a_{\theta_*}^1, a_{\theta_*}^2); \theta^0), \theta^t - \theta_* \rangle \right| \right\}. \end{aligned} \quad (34)$$

By Assumption 4.1, we compute

$$\begin{aligned}
 \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\left| \widehat{Q}^\#(s; \boldsymbol{\theta}^t) - \widehat{Q}^\#(s; \boldsymbol{\theta}_*^t) \right|^2 \right] &\stackrel{(i)}{\leq} (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi^*(\boldsymbol{\theta}^t, \boldsymbol{\theta}_*^t) (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\
 &\stackrel{(ii)}{\leq} \frac{(1-\nu)^2}{\gamma^2} \cdot (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t)^\top \Sigma_\pi (\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^t) \\
 &= \frac{(1-\nu)^2}{\gamma^2} \cdot \mathbb{E}_{\mu, \pi, \mathbb{P}} \left[\left| \widehat{Q}(s, a^1, a^2; \boldsymbol{\theta}^t) - \widehat{Q}(s, a^1, a^2; \boldsymbol{\theta}_*^t) \right|^2 \right],
 \end{aligned} \tag{35}$$

where (i) is due to (34), and (ii) is due to Assumption 4.1. Similar to Lemma B.3, we can also obtain (31) in this lemma. By substituting (35) into (31), the proof can be completed. $\square$

Now we are ready to prove Theorem 4.2. For a little notation abuse, we redefine

$$\begin{aligned}
 \mathbf{g}(\boldsymbol{\theta}^t) &= \Delta(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^t), \quad \bar{\mathbf{g}}(\boldsymbol{\theta}^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{g}(\boldsymbol{\theta}^t)] \\
 \mathbf{m}(\boldsymbol{\theta}^t) &= \widehat{\Delta}(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0), \quad \bar{\mathbf{m}}(\boldsymbol{\theta}^t) = \mathbb{E}_{\mu, \pi, \mathbb{P}} [\mathbf{m}(\boldsymbol{\theta}^t)],
 \end{aligned}$$

where

$$\begin{aligned}
 \Delta(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t) &= Q(\mathbf{x}_t; \boldsymbol{\theta}^t) - \left(r(s_t, a_t) + \gamma \max_{b^1 \in \mathcal{A}} \min_{b^2 \in \mathcal{A}} Q(\phi(s_{t+1}, b^1, b^2); \boldsymbol{\theta}^t) \right), \\
 \widehat{\Delta}(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t) &= \widehat{Q}(\mathbf{x}_t; \boldsymbol{\theta}^t) - \left(r(s_t, a_t) + \gamma \max_{b^1 \in \mathcal{A}} \min_{b^2 \in \mathcal{A}} \widehat{Q}(\phi(s_{t+1}, b^1, b^2); \boldsymbol{\theta}^t) \right).
 \end{aligned}$$

Let $\Delta_t = \Delta(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t)$ and $\widehat{\Delta}_t = \widehat{\Delta}(s_t, a_t^1, a_t^2, s'_t, \boldsymbol{\theta}^t)$. After redefining the corresponding notation, we can similarly derive (18) and adopt the associated lemmas. Due to the introduction of the additional Assumption 4.1, we provide Lemma C.1, ensuring that terms $\mathbf{I}_1 \sim \mathbf{I}_4$ can be estimated. The remainder of the proof is entirely analogous to Sections A.2 and B.2. Thus, we conclude the proof.

D. Supporting Lemmas for Multi-layer Neural Network

Recalling the definition of the parameterized Q-function, we present the following lemmas related to neural network functions, which play a crucial role in illustrating the main results of our paper. $\{\tau_i > 0\}_{i=1, \dots, 10}$ mentioned below are universal constants.

Lemma D.1. *For any $t \in \{1, 2, \dots, T\}$, we have*

$$\|\boldsymbol{\theta}^t\| \leq \tau_1 \sqrt{m}, \quad w.p. 1 - L \exp(-\tau_2 m).$$

Proof. By Lemma G.2 in (Du et al., 2019), $\|\boldsymbol{\theta}^0\| \leq \mathcal{O}(\sqrt{m})$ with probability at least $1 - L \exp(-\tau_2 m)$. Therefore,

$$\begin{aligned}
 \|\boldsymbol{\theta}^t\|_2 &\leq \|\boldsymbol{\theta}^t - \boldsymbol{\theta}^0\| + \|\boldsymbol{\theta}^0\| \\
 &\leq \omega + \|\boldsymbol{\theta}^0\| \\
 &\leq \mathcal{O}(\sqrt{m}).
 \end{aligned}$$

$\square$

Lemma D.2. *For any $l \in \{1, 2, \dots, L\}$, we have*

$$\|\mathbf{x}^{(l)}\| \leq \tau_3 \sqrt{m}, \quad \text{and} \quad \|\nabla_{\boldsymbol{\theta}} Q(\mathbf{x}; \boldsymbol{\theta})\| \leq \tau_4, \quad w.p. 1 - L \exp(-\tau_2 m).$$

Lemma D.2 has been proved by (Tian et al., 2022) (Lemmas A.6~A.10).

Lemma D.3. For any $t \in \{1, 2, \dots, T\}$ and $\theta \in S_\omega$, we have

$$|Q(\mathbf{x}_t; \theta)| \leq \tau_5 \sqrt{\log(T/\delta)}, \quad w.p. 1 - \delta - L \exp(-\tau_2 m).$$

Proof. By Lemma D.2, we have $\frac{1}{\sqrt{m}} \|\mathbf{x}^{(L)}\| \leq \tau_3$, w.p. $1 - L \exp(-\tau_2 m)$. Recall the definition of the parameterized Q-function:

$$Q(\mathbf{x}; \theta) = \frac{1}{\sqrt{m}} \mathbf{b}^\top \mathbf{x}^{(L)},$$

where each element of $\mathbf{b}$ is generated from a uniform distribution over $\{-1, +1\}$. For each $\mathbf{x}$, by Hoeffding inequality, we have

$$\mathbb{P} \left(\left| \left\langle \frac{1}{m} \mathbf{x}^{(L)}, \mathbf{b} \right\rangle \right| \geq t \right) \leq 2 \exp \left(-\frac{2t^2}{4 \|\frac{1}{m} \mathbf{x}^{(L)}\|^2} \right) \stackrel{(i)}{\leq} 2 \exp \left(-\frac{t^2}{2\tau_3^2} \right),$$

where (i) follows Lemma D.2. Substituting $\frac{\delta}{T} = 2 \exp \left(-\frac{t^2}{2\tau_3^2} \right)$, we get

$$\mathbb{P} \left(\left| \left\langle \frac{1}{m} \mathbf{x}^{(L)}, \mathbf{b} \right\rangle \right| \geq \tau_3 \sqrt{2 \log(T/\delta)} \right) \leq \frac{\delta}{T}.$$

Now, by the union bound, if we set $\tau_5 = \tau_3 \sqrt{2}$, then

$$\mathbb{P} \left(\max_{t \in [T]} |Q(\mathbf{x}_t; \theta)| \geq \tau_5 \sqrt{\log(T/\delta)} \right) \leq \sum_{t=1}^T \mathbb{P} \left(|Q(\mathbf{x}_t; \theta)| \geq \tau_5 \sqrt{\log(T/\delta)} \right) \leq \delta,$$

which completes the proof. $\square$

Lemma D.4. Denote $\nabla_\theta^2 Q(\mathbf{x}; \theta)$ as the Hessian matrix of $Q(\mathbf{x}; \theta)$. Then for all $\mathbf{x}, \theta \in S_\omega$, we have w.p. $1 - \delta$ that

$$\|\nabla_\theta^2 Q(\mathbf{x}; \theta)\|_2 \leq \tau_6 m^{-\frac{1}{2}}$$

and

$$|Q(\mathbf{x}; \theta) - \widehat{Q}(\mathbf{x}; \theta)| \leq \tau_7 m^{-\frac{1}{2}}.$$

Proof. The first inequality in Lemma D.4 has been proved by (Liu et al., 2020b) (Theorem 3.2), which implies that $Q(\mathbf{x}; \theta)$ is $\mathcal{O}(m^{-\frac{1}{2}})$ -smoothness w.r.t. θ . Therefore,

$$|Q(\mathbf{x}; \theta) - \widehat{Q}(\mathbf{x}; \theta)| = |Q(\mathbf{x}; \theta) - Q(\mathbf{x}; \theta^0) - \langle \nabla_\theta Q(\mathbf{x}; \theta^0), \theta^0 - \theta \rangle| = \mathcal{O}(m^{-\frac{1}{2}}).$$

$\square$

Lemma D.5. Let $\theta \in S_\omega$ with the radius satisfying $\omega = \mathcal{O}(1)$. Then for all $\|\mathbf{x}\|_2 = 1$ and $\theta_*^t \in S_{2\omega}$ in the neural temporal difference learning algorithm 1, it holds that

$$|\langle \mathbf{g}(\theta^t) - \mathbf{m}(\theta^t), \theta^t - \theta_*^{t+1} \rangle| \leq \left(\tau_8 \tau_9 \omega m^{-\frac{1}{2}} \sqrt{\log(T/\delta)} + 2\tau_4 \tau_8 m^{-\frac{1}{2}} \right) \|\theta^t - \theta_*^{t+1}\|$$

with probability at least $1 - 2\delta - 2L \exp(-\tau_2 m)$ over the randomness of the initial point, and $\|\mathbf{g}(\theta^t)\|_2 \leq \tau_{10} \sqrt{\log(T/\delta)}$ holds with probability at least $1 - \delta - 2L \exp(-\tau_2 m)$.

Proof. Note that

$$\begin{aligned} \|\mathbf{g}(\theta^t) - \mathbf{m}(\theta^t)\| &= \left\| \Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \theta^t) \cdot \nabla_\theta Q(\mathbf{x}_t; \theta^t) - \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}, \theta^t) \cdot \nabla_\theta Q(\mathbf{x}_t; \theta^0) \right\| \\ &\leq \left\| \Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \theta^t) \cdot (\nabla_\theta Q(\mathbf{x}_t; \theta^t) - \nabla_\theta Q(\mathbf{x}_t; \theta^0)) \right\| \\ &\quad + \left\| \left(\Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \theta^t) - \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}, \theta^t) \right) \cdot \nabla_\theta Q(\mathbf{x}_t; \theta^0) \right\|, \end{aligned} \tag{36}$$

where

$$\begin{aligned}\Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t) &= Q(\mathbf{x}_t; \boldsymbol{\theta}^t) - (r(s_t, a_t) + \gamma Q(\mathbf{x}_{t+1}; \boldsymbol{\theta}^t)), \\ \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t) &= \widehat{Q}(\mathbf{x}_t; \boldsymbol{\theta}^t) - (r(s_t, a_t) + \gamma \widehat{Q}(\mathbf{x}_{t+1}; \boldsymbol{\theta}^t)).\end{aligned}$$

For the first term in (36), by Lemma D.3, we have $w.p. 1 - \delta - L \exp(-\tau_2 m)$ that

$$|\Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t)| \leq \tau_8 \sqrt{\log(T/\delta)}.$$

By Lemma D.4, we get that

$$\|\nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^t) - \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0)\| \leq \|\nabla_{\boldsymbol{\theta}}^2 Q(\mathbf{x}_t; \boldsymbol{\theta}^t)\|_2 \cdot \|\boldsymbol{\theta}^t - \boldsymbol{\theta}^0\| \leq \tau_9 \omega m^{-\frac{1}{2}}.$$

Therefore,

$$\|\Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t) \cdot (\nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^t) - \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0))\| \leq \tau_8 \tau_9 \omega m^{-\frac{1}{2}} \sqrt{\log(T/\delta)}. \quad (37)$$

For the second term in (36), we decompose it into

$$\begin{aligned}& \left\| \left(\Delta(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t) - \widehat{\Delta}(\mathbf{x}_t, \mathbf{x}_{t+1}, \boldsymbol{\theta}^t) \right) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0) \right\| \\ & \leq \left\| \left(Q(\mathbf{x}_t, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}_t, \boldsymbol{\theta}^t) \right) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0) \right\| + \left\| \left(Q(\mathbf{x}_{t+1}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}_{t+1}, \boldsymbol{\theta}^t) \right) \cdot \nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0) \right\| \\ & \leq \left| Q(\mathbf{x}_t, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}_t, \boldsymbol{\theta}^t) \right| \cdot \|\nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0)\| + \left| Q(\mathbf{x}_{t+1}, \boldsymbol{\theta}^t) - \widehat{Q}(\mathbf{x}_{t+1}, \boldsymbol{\theta}^t) \right| \cdot \|\nabla_{\boldsymbol{\theta}} Q(\mathbf{x}_t; \boldsymbol{\theta}^0)\| \\ & \stackrel{(i)}{\leq} 2\tau_4 \tau_8 m^{-\frac{1}{2}},\end{aligned} \quad (38)$$

with probability at least $w.p. 1 - \delta - L \exp(-\tau_2 m)$, where (i) is due to Lemmas D.3 and D.4. Plugging (37) and (38) into (36) yields that

$$\begin{aligned}|\langle \mathbf{g}(\boldsymbol{\theta}^t) - \mathbf{m}(\boldsymbol{\theta}^t), \boldsymbol{\theta}^t - \boldsymbol{\theta}_*^{t+1} \rangle| &\leq \|\mathbf{g}(\boldsymbol{\theta}^t) - \mathbf{m}(\boldsymbol{\theta}^t)\| \cdot \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^{t+1}\| \\ &\leq \left(\tau_8 \tau_9 \omega m^{-\frac{1}{2}} \sqrt{\log(T/\delta)} + 2\tau_4 \tau_8 m^{-\frac{1}{2}} \right) \|\boldsymbol{\theta}^t - \boldsymbol{\theta}_*^{t+1}\|.\end{aligned}$$

Thus we complete the proof. $\square$

Lemma D.5 provides the upper bounds on $\mathbf{I}_1$ and $\mathbf{I}_2$ in Section A.2. As discussed in Section 3.1, for a finite MDP, the Gram matrix of the L -layer neural network function is positive definite and has a minimum eigenvalue of $\mathcal{O}(1)$ when the network width is sufficiently large. This, in fact, serves as an upper bound for m^* in Assumption 3.1. Further details are provided in Remark D.6.

Remark D.6. In this special case, we assume that both state space and action space are finite. Let $|\mathcal{S}|$ and $|\mathcal{A}|$ represent the dimensions of the state space and action space, respectively. For simplicity of notation, we view $\mathbf{Q}(\boldsymbol{\theta})$ as an $|\mathcal{S}||\mathcal{A}| \times 1$ column vector, with (s, a) being a multi-index arranged in the lexicographical order. Let $d \sim \mu \times \pi$ and $\mathbf{D} = \text{diag}(d)$ be an $|\mathcal{S}||\mathcal{A}|$ -dimensional diagonal matrix, whose (s, a) -th diagonal entry is $d(s, a)$, and the order of (s, a) in $\mathbf{D}$ is the same as $\mathbf{Q}(\boldsymbol{\theta})$. Denote $\mathbf{J}$ as the Jacobian matrix of $\mathbf{Q}(\boldsymbol{\theta}^0)$ and $\mathbf{J}_{\mathbf{D}} = \mathbf{D}^{\frac{1}{2}} \mathbf{J}$. Thus we can rewrite $\Sigma_{\pi} = \mathbf{J}_{\mathbf{D}}^{\top} \mathbf{J}_{\mathbf{D}}$. Notice that Σ_{π} is different from the Gram matrix Jacot et al. (2018); Du et al. (2018; 2019); Cao & Gu (2019); Allen-Zhu et al. (2019b) in deep neural network. To derive the μ -weighted Gram matrix $\text{Gram}(\boldsymbol{\theta}^0)$, we provide the following definition.

Definition D.7. (Du et al. (2019); Cao & Gu (2019); Allen-Zhu et al. (2019b;a), Neural Tangent Kernel Matrix). For any $i, j \in [|\mathcal{S}||\mathcal{A}|]$, define

$$\begin{aligned}\widetilde{\Theta}_{i,j}^{(1)} &= \Sigma_{i,j}^{(1)} = \langle \widehat{\mathbf{x}}_i, \widehat{\mathbf{x}}_j \rangle, \quad \mathbf{A}_{i,j}^{(l)} = \begin{pmatrix} \Sigma_{i,i}^{(l)} & \Sigma_{i,j}^{(l)} \\ \Sigma_{j,i}^{(l)} & \Sigma_{j,j}^{(l)} \end{pmatrix}, \\ \Sigma_{i,j}^{(l+1)} &= 2 \cdot \mathbb{E}_{(u,v) \sim N(\mathbf{0}, \mathbf{A}_{i,j}^{(l)})} [\sigma(u)\sigma(v)], \\ \widetilde{\Theta}_{i,j}^{(l+1)} &= \widetilde{\Theta}_{i,j}^{(l)} \cdot 2 \cdot \mathbb{E}_{(u,v) \sim N(\mathbf{0}, \mathbf{A}_{i,j}^{(l)})} [\sigma'(u)\sigma'(v)] + \Sigma_{i,j}^{(l+1)},\end{aligned}$$

where $\widehat{\mathbf{x}} = \sqrt{d(s, a)}\phi(s, a)$. Then we call $\Theta^{(L)} = \left[\left(\widetilde{\Theta}_{i,j}^{(L)} + \Sigma_{i,j}^{(L)} \right) / 2 \right]_{|\mathcal{S}||\mathcal{A}| \times |\mathcal{S}||\mathcal{A}|}$ the μ -weighted neural tangent kernel matrix of an L -layer ReLU network on μ -weighted state-action pairs $\widehat{\mathbf{x}}_1, \dots, \widehat{\mathbf{x}}_{|\mathcal{S}||\mathcal{A}|}$.

There is a large body of work [Jacot et al. \(2018\)](#); [Du et al. \(2018; 2019\)](#); [Cao & Gu \(2019\)](#); [Allen-Zhu et al. \(2019b\)](#) exploring the positive definiteness of $\text{Gram}(\theta^0)$ in the literature. Suppose that for all pairs $(s, a), (s', a')$, $\|\phi(s, a)\| = 1, d(s, a) > 0$ and $\phi(s, a) \nparallel \phi(s', a')$. The results of Theorem 1 and Proposition 2 in [Jacot et al. \(2018\)](#) shows that for an L -layer neural network with Gaussian initialization parameters,

$$\langle \nabla_{\theta} Q(\hat{x}_i; \theta^0), \nabla_{\theta} Q(\hat{x}_j; \theta^0) \rangle \rightarrow \Theta_{i,j}^{(L)}, \quad \text{as } m \rightarrow \infty.$$

That is, under the NTK regime, the μ -weighted Gram matrix $\text{Gram}(\theta^0) = \mathbf{J}_D \mathbf{J}_D^\top$ converges to $\Theta^{(L)}$ when m is sufficiently large. Let $\lambda_{\min}(\Theta^{(L)}) = 2\lambda' = \mathcal{O}(1)$. For any $\delta \in (0, 1)$, there exists $m^* = \text{Poly}(|S||\mathcal{A}|, L, \delta, \lambda')$ such that if $m \geq m^*$, we have

$$\lambda_{\min}(\text{Gram}(\theta^0)) \geq \lambda' \quad \text{w.p. } 1 - \delta.$$

This signifies that if the network width $m \geq m^*$, then we have $\lambda_0 \geq \lambda_{\min}(\text{Gram}(\theta^0)) \geq \lambda'$, thereby substantiating our claim.

E. Additional Notes on the Experiments in Section 5

In this section, we further discuss the experimental setup introduced in Section 5. As mentioned in Section 5, our experiments mainly test the following two aspects: (1) how does the network width m affects the final error of the algorithm (first two subfigures in Figure 1); (2) the minimum nonzero singular value in Assumption 3.3 (latter two subfigures in Figure 1).

For point (1), we first generate 2000 samples according to a given policy π to imitate the Markov process of Algorithm 1. A two-layer neural network with ELU activation is introduced, and the parameters are initialized using Algorithm 1. We set the initial learning rate at 0.001 with linear decay (per epoch) and a batch size of 100. Notably, as the parameter m increases, the TD algorithm demonstrates smaller final TD errors.

For point (2), our experiments are based on three main points. First, note that the norm of feature map and parameter random initialization will affect the scaling of the gradient norm w.r.t. $Q(s, a; \theta)$. Thus we employ the ratio $r = \sigma_{\max}/\sigma_{\min}$ to characterize the minimum non-zero singular value in order to eliminate the impact of numerical scaling. Second, we set varying network widths to verify the existence of λ_0 . Finally, it's tough to directly obtain the joint distribution of (s, a) with a fixed learning policy. However, we have that $\left\| \frac{1}{N} \sum_{(s,a) \sim \mu \times \pi} \nabla_{\theta} Q(s, a; \theta) \nabla_{\theta} Q(s, a; \theta)^\top - \Sigma_{\pi} \right\|_2 = \mathcal{O}(1/N)$. To avoid the effects of sampling randomness, we estimate Σ_{π} from different samples. The experiments demonstrate that the minimum non-zero singular value λ_0 converges to a constant as m increases.