
Federated Reinforcement Learning: Linear Speedup Under Markovian Sampling

Sajad Khodadadian¹ Pranay Sharma² Gauri Joshi² Siva Theja Maguluri¹

Abstract

Since reinforcement learning algorithms are notoriously data-intensive, the task of sampling observations from the environment is usually split across multiple agents. However, transferring these observations from the agents to a central location can be prohibitively expensive in terms of the communication cost, and it can also compromise the privacy of each agent’s local behavior policy. In this paper, we consider a federated reinforcement learning framework where multiple agents collaboratively learn a global model, without sharing their individual data and policies. Each agent maintains a local copy of the model and updates it using locally sampled data. Although having N agents enables the sampling of N times more data, it is not clear if it leads to proportional convergence speedup. We propose federated versions of on-policy TD, off-policy TD and Q-learning, and analyze their convergence. For all these algorithms, to the best of our knowledge, we are the first to consider Markovian noise and multiple local updates, and prove a linear convergence speedup with respect to the number of agents. To obtain these results, we show that federated TD and Q-learning are special cases of a general framework for federated stochastic approximation with Markovian noise, and we leverage this framework to provide a unified convergence analysis that applies to all the algorithms.

1. Introduction

Reinforcement Learning (RL) is an online sequential decision-making paradigm that is typically modeled as a Markov Decision Process (MDP) (Sutton & Barto, 2018). In an RL task, the agent aims to learn the optimal policy of the MDP that maximizes long-term reward, without knowledge of its parameters. The agent performs this task by repeatedly interacting with the environment according to a behavior policy, which in turn provides data samples that can be used to improve the policy. This MDP-based RL framework has a vast array of applications including self-driving cars (Yurtsever et al., 2020), robotic systems (Kober et al., 2013), games (Silver et al., 2016), UAV-based surveillance (Yun et al., 2022), and Internet of Things (IoT) (Lim et al., 2020).

Due to the high-dimensional state and action spaces that are typical in these applications, RL algorithms are extremely data hungry (Duan et al., 2016; Kalashnikov et al., 2018; Akkaya et al., 2019), and training RL models with limited data can result in low accuracy and high output variance (Islam et al., 2017; Xu et al., 2021). However, generating massive amounts of training data sequentially can be extremely time consuming (Nair et al., 2015). Hence, many practical implementations of RL algorithms from Atari domain to Cyber-Physical Systems rely on parallel sampling of the data from the environment using multiple agents (Mnih et al., 2016; Espeholt et al., 2018; Chen et al., 2021a; Xu et al., 2021). It was empirically shown in (Mnih et al., 2016) that the federated version of these algorithms yields faster training time and improved accuracy. A naive approach would be to transfer all the agents’ locally collected data to a central server that uses it for training. However, in applications such as IoT (Chen & Giannakis, 2018), autonomous driving (Shalev-Shwartz et al., 2016) and robotics (Kalashnikov et al., 2018), communicating high-dimensional data over low bandwidth network link can be prohibitively slow. Moreover, sharing individual data of the agents with the server might also be undesirable due to privacy concerns (Yang et al., 2019; Mothukuri et al., 2021).

Federated Learning (FL) (Kairouz et al., 2019) is an emerging distributed learning framework, where multiple agents seek to collaboratively train a shared model, while complying with the privacy and data confidentiality requirements

¹H. Milton Stewart School of Industrial & Systems Engineering, Georgia Institute of Technology, Atlanta, GA, 30332, USA

²Electrical and Computer Engineering, Carnegie Mellon University, Pittsburgh, PA, 15213, USA,. Correspondence to: Sajad Khodadadian <skhodadadian3@gatech.edu>.

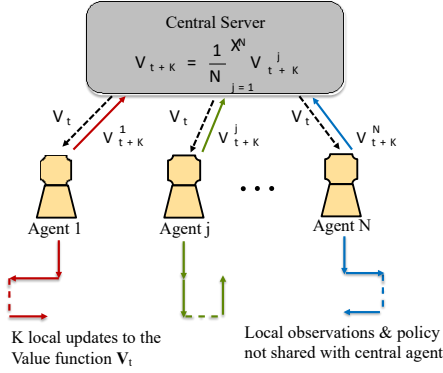


Figure 1. Schematic representation of FRL where N agents follow a Markovian trajectory and synchronize their parameter every K time steps.

(Qi et al., 2021; Yang et al., 2019). The key idea is that the agents collect data, use on-device computation capabilities to locally train the model, and only share the model updates with the central server. Not sharing data reduces communication cost and also alleviates privacy concerns.

Recently, there is a growing interest in employing FL for RL algorithms (also known as FRL) (Nadiger et al., 2019; Liu et al., 2019; Ren et al., 2019; Xu et al., 2021; Zhang et al., 2022). Unlike standard supervised learning where data is collected before training begins, in FRL, each agent collects data by following its own Markovian trajectory, while simultaneously updating the model parameters.

To ensure convergence, after every K time steps, the agents communicate with the central server to synchronize their parameters (see Figure 1). Intuitively, using more agents and a higher synchronization frequency should improve the convergence of training algorithm. However, the following questions remain to be concretely answered:

1. With N agents, do we get an N -fold (linear) speedup in the convergence of FRL algorithms?
2. How does the convergence speed and the final error scale with synchronization frequency K ?

While these questions are well-studied (Wang & Joshi, 2021; Stich, 2018; Qu et al., 2020; Li et al., 2019) in federated supervised learning, only a few works (Wai, 2020; Shen et al., 2020) have attempted to answer them in the context of FRL. However, none of them have established the convergence analysis of FRL algorithms by considering Markovian local trajectories and multiple local updates (see Table 1).

In this paper, we tackle this challenging open problem and answer both the questions listed above. We propose communication-efficient federated versions of on-policy TD,

off-policy TD, and Q-learning algorithms. In addition, we are the first to establish the convergence bounds for these algorithms in the realistic Markovian setting, showing a linear speedup in the number of agents. Previous works (Liu & Olshesky, 2021; Shen et al., 2020) on distributed RL have only shown such a speedup by assuming i.i.d. noise. Moreover, based on experiments, (Shen et al., 2020) conjectures that linear speedup may be possible under the realistic Markov noise setting, which we establish analytically. The main contributions and organization of the paper are summarized below.

- In the on-policy setting, in Section 4 we propose and analyze federated TD-learning with linear function approximation, where the agents' goal is to evaluate a common policy using on-policy samples collected in parallel from their environments. The agents only share the updated value function (not data) with the central server, thus saving communication cost. We prove a linear convergence speedup with the number of agents and also characterize the impact of communication frequency on the convergence.
- In the off-policy setting, in Section 5 we propose and analyze the federated off-policy TD-learning and federated Q-learning algorithms. Again, we establish a linear speedup in their convergence with respect to the number of agents and characterize the impact of synchronization frequency on the convergence. Since every agent samples data using a private policy and only communicates the updated value or Q-function, off-policy FRL helps keep both the data as well as the policy private.
- In Section 6, we propose a general Federated Stochastic Approximation framework with Markovian noise (FedSAM) which subsumes both federated TD-learning and federated Q-learning algorithms proposed above. Considering Markovian sampling noise poses a significant challenge in the analysis of this algorithm. The convergence result for FedSAM serves as a workhorse that enables us to analyze both federated TD-learning and federated Q-learning. We characterize the convergence of FedSAM with a refined analysis of general stochastic approximation algorithms, fundamentally improving upon prior work.

2. Related Work

Single node TD-learning and Q-learning. Most existing RL literature is focused on designing and analyzing algorithms that run at a single computing node. In the on-policy setting, the asymptotic convergence of TD-learning was established in (Tsitsiklis & Van Roy, 1997; Tadić, 2001; Borkar, 2009), and the finite-sample bounds were studied

Table 1. Comparison of sample complexity results for federated supervised learning (local SGD) and reinforcement learning algorithms. The possible distributed architectures are: 1) Worker-server, with a central server that coordinates with N agents; 2) Decentralized, where each agent directly communicates with its neighboring agents, without a central server; and 3) Shared memory, where each agent modifies a subset of the parameters of a global model held in a shared memory, that is accessible to all agents. (Recht et al., 2011).

Algorithm	Architecture	References	Local Updates	Markov Noise	Linear Speedup
Local SGD	Worker-server	(Khaled et al., 2020)	✓	7	✓
Local SGD	Worker-server	(Spiridonoff et al., 2021)	✓	7	✓
TD(0)	Worker-server	(Liu & Olshevsky, 2021)	✓	7	✓
Stochastic Approximation	Decentralized	(Wai, 2020)	7	✓	7
A3C-TD(0)	Shared memory	(Shen et al., 2020)	7	7	✓
A3C-TD(0)	Shared memory	(Shen et al., 2020)	7	✓	7
TD & Q-learning	Worker-server	This paper	✓	✓	✓

in (Dalal et al., 2018; Lakshminarayanan & Szepesvari, 2018; Bhandari et al., 2018; Srikant & Ying, 2019; Hu & Syed, 2019; Chen et al., 2021c). In the off-policy setting, (Maei, 2018; Zhang et al., 2020) study the asymptotic and (Chen et al., 2020a; 2021c) characterize the finite time bound of TD-learning. The Q-learning algorithm was first proposed in (Watkins & Dayan, 1992). There has been a long line of work to establish the convergence properties of Q-learning. In particular, (Tsitsiklis, 1994; Jaakkola et al., 1994; Bertsekas & Tsitsiklis, 1996b; Borkar & Meyn, 2000; Borkar, 2009) characterize the asymptotic convergence of Q-learning, (Beck & Srikant, 2012b; 2013; Wainwright, 2019; Chen et al., 2020a; 2021c) study the finite-sample convergence bound in the mean-square sense, and (Even-Dar & Mansour, 2004; Li et al., 2020; Qu & Wierman, 2020) study the high-probability convergence bounds of Q-learning.

Federated Learning with i.i.d. Noise. When multiple agents are used to expedite sample collection, transferring the samples to a central server for the purpose of training can be costly in applications with high-dimensional data (Shao et al., 2019) and it may also compromise the agents’ privacy. Federated Learning (FL) is an emerging distributed optimization paradigm (Konečný et al., 2016; Kairouz et al., 2019) that utilizes local computation at the agents to train models, such that only model updates, not data, is shared with the central server. In local Stochastic Gradient Descent (Local SGD or FedAvg) (McMahan et al., 2017; Stich, 2018), the core algorithm in FL, locally trained models are periodically averaged by the central server in order to achieve consensus among the agents at a reduced communication cost. While the convergence of local SGD has been extensively studied in prior work (Khaled et al., 2020; Spiridonoff et al., 2021; Qu et al., 2020; Koloskova et al., 2020), these works assume i.i.d. noise in the gradients, which is acceptable for SGD but too restrictive for RL algorithms.

Distributed and Multi-agent RL. Some recent works have analyzed distributed and multi-agent RL algorithms in the presence of Markovian noise in various settings such as decentralized stochastic approximation (Doan et al., 2019; Sun et al., 2020; Wai, 2020; Zeng et al., 2020), TD learning with linear function approximation (Wang et al., 2020a), and off-policy TD in actor-critic algorithms (Chen et al., 2021e;f). However, all these works consider decentralized settings, where the agents communicate with their neighbors after every local update. On the other hand, we consider a federated setting, with each agent performing multiple local updates between successive communication rounds, thereby resulting in communication savings. In (Shen et al., 2020), a parallel implementation of asynchronous advantage actor-critic (A3C) algorithm (which does not have local updates) has been proposed under both i.i.d. and Markov sampling. However, the authors prove a linear speedup only for the i.i.d. case, and an almost linear speedup is observed experimentally for the Markovian case.

3. Preliminaries: Single Node Setting

We model our RL setting with a Markov Decision Process (MDP) with 5 tuples $(S; A; P; R; \gamma)$, where S and A are finite sets of states and actions, P is the set of transition probabilities, R is the reward function, and $\gamma \in (0, 1)$ denotes the discount factor. At each time step t , the system is in some state S_t , and the agent takes some action A_t according to a policy (jS_t) in hand, which results in reward $R(S_t; A_t)$ for the agent. In the next time step, the system transitions to a new state S_{t+1} according to the state transition probability $P(jS_t; A_t)$. This series of states and actions $(S_t; A_t)_{t=0}^{\infty}$ constructs a Markov chain, which is the source of the Markovian noise in RL. Throughout this paper we assume that this Markov chain is irreducible and aperiodic (also known as ergodic). It is known that this Markov

chain asymptotically converges to a steady state, and we denote its stationary distribution with μ .

To measure the long-term reward achieved by following policy π , we define the value function

$$V(s) = E \left[\sum_{t=0}^{\infty} \gamma^t R(S_t; A_t) \mid S_0 = s; A_t = \pi(S_t) \right] \quad (1)$$

Equation (1) is the tabular representation of the value function. Sometimes, however, the size of the state space $|S|$ is large, and storing $V(s)$ for all $s \in S$ is computationally infeasible. Hence, a low dimensional vector $v \in \mathbb{R}^d$, where $d \ll |S|$, can be used to approximate the value function as $V(s) \approx \phi(s)^T v$ (Tsitsiklis & Van Roy, 1997). Here $\phi(s) \in \mathbb{R}^d$ is a given feature vector corresponding to the state s . Using a low-dimensional vector v to approximate a high-dimensional vector $(V(s))_{s \in S}$ is referred to as the function approximation paradigm in RL. For each $(s; a)$ pair, we also define the Q-function, $Q(s; a) = E \left[\sum_{t=0}^{\infty} \gamma^t R(S_t; A_t) \mid S_0 = s; A_0 = a \right]$, which will be employed in Q-learning.

3.1. Temporal Difference Learning

An intermediate goal in RL is to estimate the value function (either $(V(s))_{s \in S}$ or v) corresponding to a particular policy using data collected from the environment. This task is denoted as policy evaluation and one of the commonly-used approaches to accomplish this is Temporal Difference (TD)-learning (Sutton, 1988). TD-learning is an iterative algorithm where the elements of a d (or $|S|$, in the tabular setting) dimensional vector is updated until it converges to v (or V). This evaluated value function can be employed in different RL algorithms such as actor-critic (Konda & Tsitsiklis, 2000). In the on-policy function approximation setting, the update of the n -step TD-learning is as follows

$$\begin{aligned} \text{Sample } & A_{t+n} \sim \pi(\cdot | S_{t+n}); S_{t+n+1} \sim P(\cdot | S_{t+n}, A_{t+n}) \\ \text{update } & v_{t+1} = v_t + \alpha \left(\sum_{l=t}^{t+n} \gamma^{l-t} (R(S_l; A_l) + \gamma V(S_{l+1}) - V(S_l)) \right) \end{aligned} \quad (2)$$

where α is the step size. Note that in this setting, the evaluating policy and the sampling policy coincide. In contrast, in the off-policy setting these two policies can in general differ, and we need to account for this difference while running the algorithm. We will further expand on TD-learning and its variants in Sections 4.1 and 5.1.1.

3.2. Control Problem and Q-learning

Assuming some initial distribution μ_0 on the state space, the average value function corresponding to policy π is

defined as $V(\pi) = E_s[V(s)]$. This scalar quantity is a metric of average long-term rewards achieved by the agent, when it starts from distribution μ_0 and follows policy π . The ultimate goal of the agent is to obtain an optimal policy π^* which results in the maximum long-term rewards, i.e. $V(\pi^*) = \max_{\pi} V(\pi)$. Throughout the paper, we denote the parameters corresponding to the optimal policy with π^* , e.g., $V^* = V(\pi^*)$. The task of obtaining the optimal policy in RL is denoted as the control problem.

Q-learning (Watkins & Dayan, 1992) is one of the most widely used algorithms in RL to solve the control problem. At each time step t , Q-learning preserves a $|S| \times |A|$ dimensional table Q_t , and updates it table as $Q_{t+1}(s; a) = Q_t(s; a) + \alpha (R(s; a) + \max_{a'} Q_t(s; a') - Q_t(s; a))$, if $(s; a) = (S_t; A_t)$ and $Q_{t+1}(s; a) = Q_t(s; a)$ otherwise. The $|S| \times |A|$ elements of the vector Q_t are updated iteratively until it converges to Q^* , corresponding to an optimal policy. Using Q^* , one can obtain an optimal policy via greedy selection.

3.3. Stochastic Approximation and Finite Sample Bounds

Both TD-learning and Q-learning can be seen as variants of stochastic approximation (Chen et al., 2020b; 2019b;a; 2021d; Tsitsiklis, 1994). While generic stochastic approximation algorithms are studied under i.i.d. noise (Even-Dar & Mansour, 2004; Shah & Xie, 2018; Wainwright, 2019; Liu et al., 2015; Dalal et al., 2018), to apply them for studying RL we need to understand stochastic approximation under Markovian noise (Tsitsiklis, 1994; Qu & Wierman, 2020; Srikant & Ying, 2019; Chen et al., 2021c) which is significantly more challenging.

For a generic stochastic approximation (i.i.d. or Markovian noise) with constant step size α , parameter vector x_T , and convergent point x^* , it can be shown that the algorithm have the following convergence behaviour

$$E[\|x_T - x^*\|^2] \leq C_1(1 - C_0)^T + C_2; \quad (3)$$

where C_0 ; C_1 ; and C_2 are some problem dependent positive constants (Look at Appendix A for a discussion on a lower bound on the convergence of general stochastic approximation). The first term is denoted as the bias and the second term is called the variance. According to this bound, x_T geometrically converges to a ball around x^* with radius proportional to C_2 . Notice that we can always reduce the variance term by reducing the step size α , but this will lead to slower convergence in the bias term. In particular, in order to get $E[\|x_T - x^*\|^2] \leq \epsilon$, it is easy to see that we need $T \geq \frac{C_2}{\alpha^2} \log \frac{1}{\epsilon}$ sample complexity. Now suppose the constant C_2 is large. In this case, the variance term in the bound in (3) is large, and the sample complexity, which is proportional to C_2 will be poor. Notice that by the dis-

cussion in Appendix A, this bound is tight and cannot be improved.

This is where the FL can be employed in order to control the variance term by generating more data. For instance, in federated TD-learning, multiple agents work together to evaluate the value function simultaneously. Due to this collaboration, the agents can estimate the true value function with a lower variance. The same holds for estimating Q in Q-learning.

4. Federated On-policy RL

4.1. On-policy TD-learning with linear function approximation

In this section we describe the TD-learning with linear function approximation and online data samples in the single node setting. In this problem, we consider a full rank feature matrix $\Phi \in \mathbb{R}^{S \times d}$, and we denote s -th row of this matrix with $\phi(s)$; $s \in \mathcal{S}$. The goal is to find $v \in \mathbb{R}^d$ which solves the following fixed point equation:

$$v = (T)^n v \quad (4)$$

In equation (4), $(\cdot)$ is the projection with respect to the weighted 2-norm, i.e., $(V) = \arg \min_{v \in \mathbb{R}^d} \|v - V\|_P$. Here $\|v\|_P = \sqrt{v^T P v}$ and P is a diagonal matrix with diagonal entries corresponding to γ^s . In equation (4), $(T)^n$ denotes the n -step Bellman operator (Tsitsiklis & Van Roy, 1997). It is known (Tsitsiklis & Van Roy, 1997) that equation (4) has a unique solution v , and v is "close" to the true value function V . n -step TD-learning algorithm, which was shown in (2), is an iterative algorithm to obtain this unique fixed point using samples from the environment. Note that in this algorithm states and actions are sampled over a single trajectory, and hence the noise in updating v_t is Markovian. Furthermore, since the policy which samples the actions and the evaluating policy are both π , this algorithm is on-policy. As described in (Tsitsiklis & Van Roy, 1997; Bertsekas & Tsitsiklis, 1996a), the TD-learning algorithm can be studied under the umbrella of linear stochastic approximation with Markovian noise. More recently, the authors in (Bhandari et al., 2018; Srikant & Ying, 2019) have shown that the update parameter of TD-learning v_t converges to v in the form $E[\|v_t - v\|^2] = O((1 + C_0)^t + \epsilon)$. In the next section we show how FL can improve this result.

4.2. Federated TD-learning with linear function approximation

The federated version of on-policy n -step TD-learning with linear function approximation is shown in Algorithm 1. In this algorithm we consider N agents which collaboratively work together to evaluate v . For each agent i ; $i = 1; 2; \dots; N$, we initialize their corresponding pa-

rameters $v_0^i = 0$. Furthermore, each agent i samples its initial state S_0^i from some given distribution μ . In the next time steps, each agent follows a single Markovian trajectory generated by policy π , independently from other agents. At each time t , the parameter of each agent i is updated using this independently generated trajectory as $v_{t+1}^i = v_t^i + (S^i)^T E_t^i - v_t^i$. Finally, in order to ensure convergence to a global optimum, every K time steps all the agents send their parameters to a central server. The central server evaluates the average of these parameters and returns this average to each of the agents. Each agent then continues their update procedure using this average.

Notice that the averaging step is essential to ensure synchronization among the agents. Smaller K results in more frequent synchronization, and hence better convergence guarantees. However, setting smaller K is equivalent to more number of communications between the single agents and the central server, which incurs higher cost. Hence, an intermediate value for K has to be chosen to strike a balance between the communication cost and the accuracy. At the end, the algorithm samples a time step $T \sim qT^c$, where

$$q_T^c(t) = \frac{c^t}{\sum_{t=0}^{T-1} c^t} \quad \text{for } t = 0; 1; \dots; T-1 \quad (5)$$

and $c > 1$ is some constant. Since we have $q_T^c(t) > 0$ and $\sum_{t=0}^{T-1} q_T^c(t) = 1$, it is clear that $q_T^c(\cdot)$ is a probability distribution over the time interval $[0; T-1]$. In Theorem 4.1 we characterize the convergence of this algorithm as a function of ϵ , N , and K . Throughout the paper, $\mathcal{O}(\cdot)$ ignores the logarithmic terms.

Algorithm 1 Federated n -step TD (On-policy, Function Approx.)

```

1: Input: Policy  $\pi$ ;
2: Initialization:  $v_0^i = 0$  and  $S_0^i$  and  $f(A^i); S_1^i \sim \mu_1$  for
    $0 \leq i \leq N-1$  and all  $i$ 
3: for  $t = 0$  to  $T-1$  do
4:   for  $i = 1; \dots; N$  do
5:     Sample  $s_{t+n}^i \sim \pi(\cdot | S_{t+n}^i); S_{t+n+1}^i$ 
6:      $e_{t,i}^i = P(S_{t+n}^i | S_{t+n}^i) - v_{t+n}^i + (S_{t+n}^i)^T v_{t+n}^i - (S_{t+n}^i)^T v_t^i$  for
        $l = t; \dots; t+n-1$ 
7:      $E_{t,n}^i = \frac{1}{n} \sum_{l=t}^{t+n-1} e_{l,i}^i$ 
8:      $v_{t+1}^i = v_t^i + (S_{t+1}^i)^T E_{t,n}^i$ 
   end for
10:  if  $t+1 \bmod K = 0$  then
11:     $\bar{v}_{t+1} = \frac{1}{N} \sum_{j=1}^N v_{t+1}^j$ ;  $8 \leq i \leq N$ 
12:  end if
13: end for
14: Sample  $T \sim qT^c$ ;
Return:  $\frac{1}{N} \sum_{i=1}^N v_T^i$ 

```

Theorem 4.1. Let $\mathbf{v}_\wedge = \frac{1}{N} \sum_{i=1}^N \mathbf{v}^i$ denote the average of the parameters across agents at the random time $\mathbf{T}^\wedge$. For small enough step size α , and $T = O(\log \frac{1}{\alpha})$, there exist constant $c_{TDL} \in (0, 1)$ (see Section C.1 for precise statement), such that we have

$$\mathbb{E}[\|\mathbf{v}_{\mathbf{T}^\wedge} - \mathbf{v}^*\|_2^2] \leq c_{TDL} \left(\frac{1}{N} \sum_{i=1}^N \|\mathbf{v}^i - \mathbf{v}^*\|_2^2 + \frac{c_{TDL}}{N^2} \right) + c_{TDL} (K - 1)^2 + c_{TDL} \frac{32}{4},$$

where c_{TDL}^i , $i = 0; 1; 2; 3; 4$ are problem dependent constants, and $\frac{1}{N} \sum_{i=1}^N \|\mathbf{v}^i - \mathbf{v}^*\|_2^2 = O(\log(1/\alpha))$. By choosing $\alpha = O(\frac{\log(NT)}{T})$ and $K = T = N$, we achieve $\mathbb{E}[\|\mathbf{v}_{\mathbf{T}^\wedge} - \mathbf{v}^*\|_2^2] \leq \frac{1}{N}$ within $T = O(\frac{1}{\alpha})$ iterations.

For brevity purposes, here we did not show the exact dependence of the constants c_{TDL}^i , $i = 0; 1; 2; 3$ on the problem dependent constants. For a discussion on the detailed expression look at Section C.1 in the appendix.

Theorem 4.1 shows that federated TD-learning with linear function approximation enjoys a linear speedup with respect to the number of agents. Compared to the convergence bound of general stochastic approximation in (3), the bound in Theorem 4.1 has three differences. Firstly, the variance term which is proportional to the step size α is divided with the number of agents N . This will allow us to control the variance (and hence improve sample complexity) by employing more number of agents. Secondly, we have an extra term which is zero with perfect synchronization $K = 1$. Although this term is not divided with N , but it is proportional to α^2 , which is one order higher than the variance term in (3). Finally, the last term is of the order $O(\alpha^3)$, which can be handled by choosing small enough step size.

Furthermore, according to the choice of K in Theorem 4.1, after T iterations, the communication cost of federated TD is $T = K = N$. However, by employing federated TD-learning in the naive setting where all the agents communicate with the central server at every time step, the communication cost will be $O(T)$. Hence, we observe that by carefully tuning the hyper parameters of federated TD, we can significantly reduce the communication cost of the overall algorithm, while not losing performance in terms of the sample complexity.

Finally, federated TD-learning Algorithm 1 preserves the privacy of the agents. In particular, since the single agents only require to share their parameters $\mathbf{v}_{t+1}^i$, the central server will not be exposed to the state-action-reward trajectory generated by each agent. This can be essential in some applications where privacy is an issue (Mothukuri et al., 2021; Truex et al., 2019). Examples of such applications include autonomous driving (Liang et al., 2019; Zhao et al., 2021), Internet of Things (IoT) (Nguyen et al., 2021;

Ren et al., 2019; Wang et al., 2020b), and cloud robotics (Liu et al., 2019; Xu et al., 2021).

Remark. In algorithm 1, the randomness in choosing $\mathbf{T}^\wedge$ is independent of all the other randomness in the problem. Hence, in a practical setting, one can sample $\mathbf{T}^\wedge$ ahead of time, before running the algorithm, and stop the algorithm at time step $\mathbf{T}^\wedge$ and output $\mathbf{v}_{\mathbf{T}^\wedge}$. By this method, we require only a single data point to be saved, which results in the memory complexity of $O(1)$ for the algorithm.

5. Federated Off-Policy RL

On-policy TD-learning requires online sampling from the environment, which might be costly (e.g. robotics (Gu et al., 2017; Levine et al., 2020)), high risk (e.g. self-driving cars (Yurtsever et al., 2020; Maddern et al., 2017)), or unethical (e.g. in clinical trials (Gottesman et al., 2019; Liu et al., 2018; Gottesman et al., 2020)). Off-policy training in RL refers to the paradigm where we use data collected by a fixed behaviour policy to run the algorithm. When employed in federated setting, off-policy RL has privacy advantages as well (Foerster et al., 2016; Qi et al., 2021; Zhuo et al., 2019). In particular, suppose each single agent attains a unique sampling policy, and they do not wish to reveal these policies to the central server. In off-policy FL, agents only transmit sampled data, and hence the sampling policies remain private to each agent.

In Section 5.1 we will discuss off-policy TD-learning and in Section 5.2 we will discuss Q-learning, which is an off-policy control algorithm. For the off-policy algorithms, we only study the tabular setting. Notice that it has been observed that the combination of off-policy sampling and function approximation in RL (also known as deadly triad (Sutton & Barto, 2018)) can result in instability or even divergence (Baird, 1995). Recently there has been some work to overcome deadly triad (Chen et al., 2021b). Extension of our work to function approximation in the off-policy setting is a future research direction.

5.1. Federated Off-Policy TD-learning

In the following, we first discuss single-node off-policy TD-learning, and then we generalize it to the federated setting.

5.1.1. OFF-POLICY TD-LEARNING

In off-policy TD-learning the goal is to evaluate the value function $V = (V(s))_{s \in \mathcal{S}}$ corresponding to the policy using data sampled from some fixed behaviour policy b . In this setting, the evaluating policy π and the sampling policy b can be arbitrarily different, and we need to account for this difference while performing the evaluation. Although π and b can be different, notice that the value function V does not depend on b . In order to account for this

difference, we introduce the notion of importance sampling as $I^b(s; a) = \frac{(a|s)}{(a|s)}$ which is employed in the off-policy TD-learning.

Recently, several works studied the finite-time convergence of off-policy TD-learning. In particular, the authors in (Khadadian et al., 2021; Chen et al., 2021c; 2020b; 2021d) show that, similar to on-policy TD, off-policy TD-learning can be studied under the umbrella of stochastic approximation. Hence, this algorithm enjoys similar convergence behaviour as (3).

5.1.2. FEDERATED OFF-POLICY TD-LEARNING

The federated version of n -step off-policy TD-learning is shown in Algorithm 2. In this algorithm, each agent i attains a unique (and private) sampling policy i and follows an independent trajectory generated by this policy. Furthermore, at each time step t , each agent i attains a jSj -dimensional vector V_t^i and updates this vector using the samples generated by i . In order to account for the off-policy sampling, each agent utilizes $I^{(i)}(S^i; A^i) = \frac{(A^i|S^i)}{(A^i|S^i)}$ in the update of their algorithm. We further define $I_{\max} = \max_{s,a;i} I^{(i)}(s; a)$, which is a measure of discrepancy between the evaluating policy and sampling policy i of all the agents.

In order to ensure synchronization, all the agents transmit their parameter vectors to the central server every K time steps. The central server returns the average of these vectors to each agent and each agent follows this averaged vector afterwards. Notice that in federated off-policy TD-learning Algorithm 2, each agent share neither their sampled trajectory of state-action-rewards, nor their sampling policy with the central server. This provides two levels of privacy for the single agents. At the end, the algorithm samples a time step $T \sim q_T^{\text{TD}}$, where the distribution q_T^{TD} is defined in (5) and $C_{TD} = 1 - \frac{1}{2}$, where $\frac{1}{2} = \frac{0.5e^{1/4}(2 - \min(1, \frac{n+1}{n}))}{e^{1/4} + \frac{2 - \min(1, \frac{n+1}{n})}{2}}$. Here, we de-

note $\min = \min_{s,i} I^{(i)}(s)$. The constant C_{TD} is carefully chosen to ensure the convergence of Algorithm 2. Furthermore, for small enough step size γ , it can be shown that $0 < C_{TD} < 1$. Theorem 5.1 states the convergence of this Algorithm.

Theorem 5.1. Consider the federated n -step off-policy TD-learning Algorithm 2. Denote $V_T = \frac{1}{N} \sum_{i=1}^N V_T^i$. For small enough step size γ and large enough T , we have

$$E[kV_T - V_k^2] \leq C_1^{TD} \frac{1}{N} C_{TD}^2 + \frac{C_2^{TD}}{3} (K - 1)^2;$$

$$\text{where } C_1^{TD} = C_1^{TD} \cdot \frac{I_{\max}^2}{\min(1, \frac{n+1}{n})^2}, \quad C_2^{TD} =$$

Algorithm 2 Federated n -step TD (Off-policy Tabular Setting)

```

1: Input: Policy ;
2: Initialization:  $V_0^i = 0$  and  $S_0^i$  and  $fA^i$ ;  $S_1^i \sim \pi_1^i$  for  $0 \leq i \leq N-1$  and all  $i$ 
3: for  $t = 0$  to  $T-1$  do
4:   for  $i = 1; \dots; N$  do
5:     Sample  $A_{t+n}^i(jS^i); S_{t+n}^i \sim P(S_{t+n}^i | S_t^i, A_{t+n}^i)$ 
6:      $e_t^i = R(S_{t+n}^i; A_{t+n}^i) + V_{t+n}^i - V_t^i$ 
7:     Update  $V_{t+1}^i(s) = V_t^i(s) + \gamma I^{(i)}(S^i; A^i) e_t^i$  if  $s = S_t^i$  and  $V_{t+1}^i(s) = V_t^i(s)$  otherwise.
8:   end for
9:   if  $t+1 \bmod K = 0$  then
10:     $V_{t+1} = \frac{1}{N} \sum_{j=1}^N V_{t+1}^j$ ;  $8 \leq i \leq N$ 
11:   end if
12: end for
13: Sample  $T \sim q_T^{\text{TD}}$ 
14: Return:  $\frac{1}{N} \sum_{i=1}^N V_T^i$ 
    
```

$C_2^{TD} = \frac{I_{\max}^2}{\min(1, \frac{n+1}{n})^2} \log^2(jS)$, $C_3^{TD} = C_3^{TD} \cdot \frac{I_{\max}^2}{\min(1, \frac{n+1}{n})^2} \log^2(jS)$, and C_i^{TD} , $i = 1; 2; 3$ are universal problem independent constants. In addition, choosing $\gamma = \frac{8 \log(NT)}{T}$ and $K = T/N$, we have $E[kV_T - V_k^2] \leq O(\frac{1}{N} \cdot \frac{I_{\max}^2}{\min(1, \frac{n+1}{n})^2} \log^2(jS))$ after $T = O(\frac{1}{N} \cdot \frac{I_{\max}^2}{\min(1, \frac{n+1}{n})^2} \log^2(jS))$ iterations.

The proof is given in Section C.2 in the appendix.

Note that similar to on-policy TD-learning Algorithm 1, off-policy TD-learning also enjoys a linear speedup while maintaining a low communication cost. In addition, this algorithm preserves the privacy of the agents by holding both the data and the sampling policy private.

5.2. Federated Q-learning

So far we have discussed policy evaluation problem with on and off-policy samples. Next we aim at solving the control problem by employing the celebrated Q-learning algorithm (Watkins & Dayan, 1992; Tsitsiklis, 1994). In the next section we will explain the Q-learning algorithm. Further, in Section 5.2.2 we will provide a federated version of Q-learning along with its convergence result.

5.2.1. Q-LEARNING

The goal of Q-learning is to evaluate Q , which is the unique Q-function corresponding to the optimal policy. Knowing Q , one can obtain an optimal policy through a greedy selection (Puterman, 2014), and hence resolve the control problem.

Suppose $f_{S_t}; A_t g_{t,0}$ is generated by a fixed behaviour policy b . At each time step t , Q-learning preserves a $jS:jA$ table Q_t and updates the elements of this table as shown in Section 3.2. By assuming b to be an ergodic policy, the asymptotic convergence of Q_t to Q has been established in (Bertsekas & Tsitsiklis, 1996b). Furthermore, it can be shown that Q-learning is a special case of stochastic approximation and enjoys a convergence bound similar to (3) (Beck & Srikant, 2012a; Li et al., 2020; Qu & Wierman, 2020; Chen et al., 2021c).

Two points worth mentioning about the Q-learning algorithm. Firstly, Q-learning is an off-policy algorithm in the sense that only samples from a fixed ergodic policy is needed to perform the algorithm. Secondly, as opposed to the TD-learning, the update of the Q-learning is non-linear. This imposes a sharp contrast between the analysis of Q-learning and TD-learning (Chen et al., 2019a).

5.2.2. FEDERATED Q-LEARNING

Algorithm 3 provides the federated version of Q-learning. We characterize its convergence in the following theorem.

Algorithm 3 Federated Q-learning

```

1: Input: Sampling policy  $\mathbb{f}$  for  $i = 1; 2; \dots; N$ , initial distribution
2: Initialization:  $Q_0^i = 0$  and  $S_0^i$  for all  $i$  3:
   for  $t = 0$  to  $T - 1$  do
4:   for  $i = 1; \dots; N$  do
5:     Sample  $A_t^i(jS^i); S_{t+1}^i \sim P(jS^i; A_t^i)$ 
6:     Update  $Q_{t+1}^i(s; a) = Q_t^i(S_t^i; A_t^i) +$ 
        $R(S_t^i; A_t^i) + \max_a Q_t(S_{t+1}^i; a) - Q_t(S_t^i; A_t^i)$ , if
        $(s; a) = (S_t^i; A_t^i)$  and  $Q_{t+1}^i(s; a) = Q_t^i(s; a)$ 
       otherwise.
7:   end for
8:   if  $t + 1 \bmod K = 0$  then
9:      $Q_{t+1}^i = \frac{1}{N} \sum_{j=1}^N Q_{t+1}^j; \forall i \in [N]$ 
10:  end if
11: end for
12: Sample:  $\tau \sim p_{\tau}^{Q_T}$ 
13: Return:  $\frac{1}{N} \sum_{i=1}^N Q_T^i$ 
    
```

Theorem 5.2. Consider the federated Q-learning Algorithm 3 with $c_Q = \frac{1}{2} \frac{c_Q}{c_Q + 1} (0; 1)$, where $c_Q = \frac{1}{1 + \frac{0.5e^{1-4}(2 - \min(1))}{\frac{1}{2} \min(1)}^2}$ and we denote $\min = \min_{s,a,i} \mathbb{P}(s|a) \mathbb{P}(a|s)$. Denote $Q^* = \frac{1}{N} \sum_{i=1}^N Q^i$. For small enough step size and large enough T , we have

$$\mathbb{E}[kQ_T^* - Q^*k^2] \leq C_1 \frac{1}{T} + C_2 \frac{1}{N} (K - Q^*)^2;$$

where $C^Q = C_1^Q \frac{1}{\min(1)^3} + C_2^Q \frac{jSj^2 \log^2(jSj)}{\min(1)^4} + C_3^Q \frac{jSj^2 \log^2(jSj)}{\min(1)^8}$, and $C_i^Q, i = 1; 2; 3$ are universal problem

independent constants. In addition, choosing $\frac{8 \log(NT)}{T^2 c_Q}$

and $K = T = N$, we have $\mathbb{E}[kQ_T^* - Q^*k^2] \leq \frac{1}{N}$ within

$$T = \tilde{O}\left(\frac{1}{N} \frac{jSj^2 \log^2(jSj)}{\min(1)^9}\right) \text{ iterations.}$$

According to Theorem 5.2, federated Q-learning Algorithm 3, similar to federated off-policy TD-learning, enjoys linear speedup, communication efficiency as well as privacy guarantees. We would like to emphasize that the update of Q-learning is non-linear. Hence the result of Theorem 5.2 cannot be derived from Theorems 4.1 and 5.1.

6. Generalized Federated Stochastic Approximation

In this section we study the convergence of a general federated stochastic approximation for contractive operators, FedSAM, which is presented in Algorithm 4. In this algorithm there are N agents $i = 1; 2; \dots; N$. At each time step $t = 0$, each agent i maintains the parameter $\theta_t^i \in \mathbb{R}^d$. At time $t = 0$, all agents initialize their parameters with $\theta_0^i = \theta_0$. Next, at time $t = 0$, each agent i updates its parameter as $\theta_{t+1}^i = \theta_t^i + \eta_t G^i(\theta_t^i; y_t^i) + b^i(\theta_t^i)$. Here η_t denotes the step size, and y^i is a noise which is Markovian along the time t , but is independent across the agents i . This notion is defined more concretely in Assumption 6.4. We note that functions $G^i(\cdot; \cdot)$ and $b^i(\cdot)$ are allowed to be dependent on the agent i . This allows us to employ the convergence bound of FedSAM in order to derive the convergence bound of off-policy TD-learning with different behaviour policies across agents. In order to avoid divergence, every K time steps we synchronize the parameters of all the agents as $\theta_{t+1}^i = \theta_{t+1}^j, \frac{1}{N} \sum_{j=1}^N \theta_{t+1}^j$ for all $i \in [N]$. Note that although smaller K corresponds to more frequent synchronization and hence more “accurate” updates, at the same time it results in a higher communication cost, which is not desirable. Hence, in order to determine the optimal choice of synchronization period, it is essential to characterize the dependence of the convergence on K . This is one of the results which we will derive in Theorem B.1. Finally, the algorithm samples $\tau \sim q_T$, where $q_T(\tau) = \frac{p_{\tau} c_{\tau}}{p_{\tau} c_{\tau} + 1}$ and outputs $\hat{\theta}_T$. This sampling scheme is essential for the convergence of overall algorithm. We further make some assumptions regarding the underlying process.

First, we assume that the expectation of $G^i(\cdot; y^i)_t$ geometrically converges to some function $G^i(\cdot)$ and the expectation of $b^i(y^i)$ geometrically converges to 0. In particular, we have the following assumption.

Assumption 6.1. For every agent i , there exist a function $G^i(\cdot)$ such that we have

$$\lim_{t \rightarrow \infty} \mathbb{E}[G^i(\cdot; y^i)_t] = G^i(\cdot)$$

$$\lim_{t \rightarrow \infty} \mathbb{E}[b^i(y^i_t)] = 0; \quad (6)$$

Algorithm 4 Federated Stochastic Approximation with Markovian Noise (FedSAM)

```

1: Input:  $C_{FSAM}; T; 0; K; 2: i$ 
   = 0 for all  $i = 1; \dots; N$ . 3: for  $t$ 
   = 0 to  $T - 1$  do
4:    $i_{t+1} = i_t + G^i(i_t; y_t^i) - i_t + b^i(y_t^i); 8: i \in [N]$ 
5:   if  $t + 1 \bmod K = 0$  then
6:      $i_{t+1} = t+1, \frac{1}{N} \sum_{j=1}^N i_{t+1}; 8: i \in [N]$ 
7:   end if
8: end for
9: Sample  $\mathbf{T} \sim q_{P_T}^{C_{FSAM}}()$ .
10: Return:  $\frac{1}{N} \sum_{i=1}^N i_T$ 
    
```

Furthermore, there exists $m_1, m_2 \geq 0$ and $\alpha \in [0, 1]$, such that for every $i = 1; 2; \dots; N$,

$$\|G^i(\cdot) - E[G^i(\cdot; y^i)]\|_{k_c} \leq m_1 k_c^t + k E[b^i(y^i)]_{k_c} \leq m_2 t; \quad (7)$$

where k, k_c is a given norm.

Next, we assume a contraction property on the expected operator $G^i(\cdot)$.

Assumption 6.2. We assume all expected operators $G^i(\cdot)$ are contraction mappings with respect to k, k_c with contraction factor $\alpha \in [0, 1]$. That is, for all $i = 1; 2; \dots; N$,

$$\|G^i(\cdot) - G^i(\cdot)k_c\|_{k_c} \leq \alpha k_c; \quad \forall i; 2 \leq R^d:$$

Next, we consider some Lipschitz and boundedness properties on $G^i(\cdot; y^i)$ and $b^i(\cdot)$.

Assumption 6.3. For all $i = 1; \dots; N$, there exist constants A_1, A_2 and B such that

1. $\|G^i(\cdot; y^i) - G^i(\cdot; y^i)k_c\|_{k_c} \leq A_1 k_c + \alpha k_c$, for all $i; 2 \leq R^d$.
2. $\|G^i(\cdot; y^i)k_c\|_{k_c} \leq A_2 k_c$ for all $i; y^i$.
3. $\|b^i(y^i)k_c\|_{k_c} \leq B$ for all y^i .

Remark. By Assumption 6.2 and due to the Banach fixed point theorem, $G^i(\cdot)$ has a unique fixed point for all $i = 1; 2; \dots; N$. Furthermore, by Assumption 6.3, we have $G^i(0; y) = 0$. Hence the point 0 is the unique fixed point of $G^i(\cdot)$.

Finally, we impose an assumption on the random data y_t^i .

Assumption 6.4. We assume that the Markovian noise y_t^i (Markovian with respect to time t) is independent across agents i . In other words, for all measurable functions $f(\cdot)$ and $g(\cdot)$, we assume the following

$$E_r[f(y_t^i)g(y_t^j)] = E_r[f(y_t^i)]E_r[g(y_t^j)];$$

for all $r, t; i = j$.

Theorem 6.1 states the convergence of Algorithm 4.

Theorem 6.1. Consider the federated stochastic approximation Algorithm 4 with $C_{FSAM} = 1 - \frac{2}{N} (0; 1)$ ($\frac{2}{N}$ is defined in Equation (14) in the appendix), and synchronization frequency K . Denote $\bar{x}_t = \frac{1}{N} \sum_{i=1}^N x_t^i$ and consider $\bar{x}_t$ as the output of this algorithm after T iterations. Assume $\alpha = d \log \frac{1}{\epsilon}$. For $T \geq \max\{K + 2g\}$ and small enough step size α , we have

$$E[\|k_T k_c^T\|_{k_c}^2] \leq C_1 \alpha^{T+1} + C_2 + C_3 (K^2 \alpha)^2 + C_4 \alpha^2; \quad (8)$$

where $C_i, i = 1; 2; 3; 4$ are some constants which are specified precisely in Appendix B, and are independent of $K; N$. Choosing $\alpha = \frac{8 \log(NT)}{T^2}$ and $K = T/N$,

we get $T = \tilde{O}(\frac{1}{\alpha})$ sample complexity for achieving $E[\|k_T k_c\|^2] \leq \epsilon$.

Theorem 6.1 establishes the convergence of $\bar{x}_t$ to zero in the expected mean-squared sense. The first term in (8) converges geometrically to zero as T grows. The second term is proportional to $\frac{1}{N}$ similar as (3). However, the number of agents N in the denominator ensures linear speedup, meaning that for small enough α (such that $\alpha = N$ is the dominant term), the sample complexity of each individual agent, relative to a centralized system, is reduced by a factor of N . The third term has quadratic dependence on α , and is zero when we have perfect synchronization, i.e. $K = 1$. The last term is proportional to α^3 , and has the weakest dependency on the step size α . For $K > 1$ we can merge the last two terms by upper bounding $\alpha^3 \leq \alpha^2$. The current upper bound, however, is tighter since with $K = 1$ (i.e. perfect synchronization) we have no term in the order α^2 .

Note that similar bounds (sans the last α^3 term) have been established for the simpler i.i.d. noise case in the federated setting (Khaled et al., 2020; Koloskova et al., 2020). Consequently, we achieve the same sample complexity results for the more general federated setting with Markov noise.

Remark. The bound in Theorem 6.1 holds only after $T > \max\{K + 2g\}$ and for all synchronization periods $K \geq 1$. At $K = 1$ the third term in the bound goes away, and we will be left only with the first order term, which is linearly decreasing with respect to the number of agents N , and the third order term $\tilde{O}(\alpha^3)$. The last term, however, is not tight and can be further improved to be of the order $\tilde{O}(\alpha^j); j > 3$. However, for that we need to assume larger α , which means the bound only hold after a longer waiting time. In particular, by choosing $\alpha = d \log \frac{1}{\epsilon}$, we can get $\tilde{O}(2^{T-1})$ for the last term (see the proof of Lemma B.2).

Acknowledgment

This work was partially supported by NSF awards CCF-1944993, CCF-2045694, CNS-2112471, CMMI-2112533, EPCN-2144316, and an award from Raytheon technologies.

References

- Akkaya, I., Andrychowicz, M., Chociej, M., Litwin, M., McGrew, B., Petron, A., Paino, A., Plappert, M., Powell, G., Ribas, R., Schneider, J., Tezak, N., Tworek, J., Welinder, P., Weng, L., Yuan, Q., Zaremba, W., and Zhang, L. Solving Rubik’s cube with a robot hand. Preprint arXiv:1910.07113, 2019.
- Baird, L. Residual algorithms: Reinforcement learning with function approximation. In *Machine Learning Proceedings 1995*, pp. 30–37. Elsevier, 1995.
- Banach, S. Sur les opérations dans les ensembles abstraits et leur application aux équations intégrales. *Fund. math.*, 3(1):133–181, 1922.
- Beck, A. First-order methods in optimization, volume 25. SIAM, 2017.
- Beck, C. L. and Srikant, R. Error bounds for constant step-size Q-learning. *Syst. Control. Lett.*, 61:1203–1208, 2012a.
- Beck, C. L. and Srikant, R. Error bounds for constant step-size Q-learning. *Systems & control letters*, 61(12): 1203–1208, 2012b.
- Beck, C. L. and Srikant, R. Improved upper bounds on the expected error in constant step-size Q-learning. In *2013 American Control Conference*, pp. 1926–1931. IEEE, 2013.
- Bertsekas, D. P. and Tsitsiklis, J. N. *Neuro-dynamic programming*. Athena Scientific, 1996a.
- Bertsekas, D. P. and Tsitsiklis, J. N. *Neuro-dynamic programming*. Athena Scientific, 1996b.
- Bertsekas, D. P., Bertsekas, D. P., Bertsekas, D. P., and Bertsekas, D. P. *Dynamic programming and optimal control*, volume 2. Athena scientific Belmont, MA, 1995.
- 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.
- Borkar, V. S. and Meyn, S. P. The ODE method for convergence of stochastic approximation and reinforcement learning. *SIAM Journal on Control and Optimization*, 38(2):447–469, 2000.
- Chen, T. and Giannakis, G. B. Bandit convex optimization for scalable and dynamic iot management. *IEEE Internet of Things Journal*, 6(1):1276–1286, 2018.
- Chen, T., Zhang, K., Giannakis, G. B., and Basar, T. Communication-efficient policy gradient methods for distributed reinforcement learning. *IEEE Transactions on Control of Network Systems*, 2021a.
- Chen, Z., Zhang, S., Doan, T. T., Maguluri, S. T., and Clarke, J.-P. Finite-sample analysis of nonlinear stochastic approximation with applications in reinforcement learning. Under review by Automatica, Preprint arXiv:1905.11425, 2019a.
- Chen, Z., Zhang, S., Doan, T. T., Maguluri, S. T., and Clarke, J.-P. Performance of Q-learning with linear function approximation: Stability and finite-time analysis. In *OptRL Workshop at NeurIPS 2019*, 2019b.
- Chen, Z., Maguluri, S. T., Shakkottai, S., and Shanmugam, K. Finite-sample analysis of stochastic approximation using smooth convex envelopes. Under Review at *Mathematics of Operations Research*, Preprint arXiv:2002.00874, 2020a.
- Chen, Z., Maguluri, S. T., Shakkottai, S., and Shanmugam, K. Finite-sample analysis of stochastic approximation using smooth convex envelopes. In *Advances in Neural Information Processing Systems*, 2020b. URL <https://proceedings.neurips.cc/paper/2020/file/5d44ee6f2c3f71b73125876103c8f6c4-Paper.pdf>.
- Chen, Z., Khodadadian, S., and Maguluri, S. T. Finite-Sample Analysis of Off-Policy Natural Actor-Critic with Linear Function Approximation. Preprint arXiv:2105.12540, 2021b. Submitted to NeurIPS 2021.
- Chen, Z., Maguluri, S. T., Shakkottai, S., and Shanmugam, K. A Lyapunov Theory for Finite-Sample Guarantees of Asynchronous Q-Learning and TD-Learning Variants. Under review by JMLR, Preprint arXiv:2102.01567, 2021c.
- Chen, Z., Maguluri, S. T., Shakkottai, S., and Shanmugam, K. Finite-Sample Analysis of Off-Policy TD-Learning via Generalized Bellman Operators. Preprint arXiv:2106.12729, 2021d.
- Chen, Z., Zhou, Y., and Chen, R. Multi-agent off-policy td learning: Finite-time analysis with near-optimal sample complexity and communication complexity. arXiv preprint arXiv:2103.13147, 2021e.
- Chen, Z., Zhou, Y., Chen, R., and Zou, S. Sample and communication-efficient decentralized actor-critic algorithms with finite-time analysis. arXiv preprint arXiv:2109.03699, 2021f.

- Dalal, G., Szörényi, B., Thoppe, G., and Mannor, S. Finite sample analysis for TD(0) with function approximation. In *Thirty-Second AAAI Conference on Artificial Intelligence*, 2018.
- Doan, T., Maguluri, S., and Romberg, J. Finite-Time Analysis of Distributed TD(0) with Linear Function Approximation on Multi-Agent Reinforcement Learning. In *International Conference on Machine Learning*, pp. 1626–1635, 2019.
- Duan, Y., Schulman, J., Chen, X., Bartlett, P. L., Sutskever, I., and Abbeel, P. RI^2 : Fast reinforcement learning via slow reinforcement learning. *arXiv preprint arXiv:1611.02779*, 2016.
- Espeholt, L., Soyer, H., Munos, R., Simonyan, K., Mnih, V., Ward, T., Doron, Y., Firoiu, V., Harley, T., Dunning, I., Legg, S., and Kavukcuoglu, K. IMPALA: Scalable Distributed Deep-RL with Importance Weighted Actor-Learner Architectures. In *International Conference on Machine Learning*, pp. 1406–1415, 2018.
- Even-Dar, E. and Mansour, Y. Learning Rates for Q-Learning. *J. Mach. Learn. Res.*, 5:1–25, 2004. ISSN 1532-4435.
- Foerster, J. N., Assael, Y. M., De Freitas, N., and Whiteson, S. Learning to communicate with deep multi-agent reinforcement learning. *arXiv preprint arXiv:1605.06676*, 2016.
- Gottesman, O., Johansson, F., Komorowski, M., Faisal, A., Sontag, D., Doshi-Velez, F., and Celi, L. A. Guidelines for reinforcement learning in healthcare. *Nature medicine*, 25(1):16–18, 2019.
- Gottesman, O., Futoma, J., Liu, Y., Parbhoo, S., Celi, L., Brunskill, E., and Doshi-Velez, F. Interpretable off-policy evaluation in reinforcement learning by highlighting influential transitions. In *International Conference on Machine Learning*, pp. 3658–3667. PMLR, 2020.
- Gu, S., Holly, E., Lillicrap, T., and Levine, S. Deep reinforcement learning for robotic manipulation with asynchronous off-policy updates. In *2017 IEEE international conference on robotics and automation (ICRA)*, pp. 3389–3396. IEEE, 2017.
- Hu, B. and Syed, U. A. Characterizing the exact behaviors of temporal difference learning algorithms using markov jump linear system theory. *arXiv preprint arXiv:1906.06781*, 2019.
- Islam, R., Henderson, P., Gomrokchi, M., and Precup, D. Reproducibility of benchmarked deep reinforcement learning tasks for continuous control. *arXiv preprint arXiv:1708.04133*, 2017.
- Jaakkola, T., Jordan, M. I., and Singh, S. P. Convergence of stochastic iterative dynamic programming algorithms. In *Advances in neural information processing systems*, pp. 703–710, 1994.
- Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., et al. Advances and open problems in federated learning. *Preprint arXiv:1912.04977*, 2019.
- Kalashnikov, D., Irpan, A., Pastor, P., Ibarz, J., Herzog, A., Jang, E., Quillen, D., Holly, E., Kalakrishnan, M., Vanhoucke, V., and Levine, S. Qt-opt: Scalable deep reinforcement learning for vision-based robotic manipulation. *Preprint arXiv:1806.10293*, 2018.
- Khaled, A., Mishchenko, K., and Richtárik, P. Tighter theory for local sgd on identical and heterogeneous data. In *International Conference on Artificial Intelligence and Statistics*, pp. 4519–4529. PMLR, 2020.
- Khodadadian, S., Chen, Z., and Maguluri, S. T. Finite-Sample Analysis of Off-Policy Natural Actor-Critic Algorithm. In *International Conference on Machine Learning*, 2021.
- Kober, J., Bagnell, J. A., and Peters, J. Reinforcement learning in robotics: A survey. *The International Journal of Robotics Research*, 32(11):1238–1274, 2013.
- Koloskova, A., Loizou, N., Boreiri, S., Jaggi, M., and Stich, S. A unified theory of decentralized sgd with changing topology and local updates. In *International Conference on Machine Learning*, pp. 5381–5393. PMLR, 2020.
- Konda, V. R. and Tsitsiklis, J. N. Actor-critic algorithms. In *Advances in neural information processing systems*, pp. 1008–1014, 2000.
- Konečný, J., McMahan, H. B., Ramage, D., and Richtárik, P. Federated optimization: Distributed machine learning for on-device intelligence. *arXiv preprint arXiv:1610.02527*, 2016.
- Lakshminarayanan, C. and Szepesvari, C. Linear stochastic approximation: How far does constant step-size and iterate averaging go? In *International Conference on Artificial Intelligence and Statistics*, pp. 1347–1355, 2018.
- Levine, S., Kumar, A., Tucker, G., and Fu, J. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *Preprint arXiv:2005.01643*, 2020.
- Li, G., Wei, Y., Chi, Y., Gu, Y., and Chen, Y. Sample Complexity of Asynchronous Q-Learning: Sharper Analysis and Variance Reduction. *Advances in neural information processing systems*, 2020.

- Li, X., Huang, K., Yang, W., Wang, S., and Zhang, Z. On the convergence of fedavg on non-iid data. arXiv preprint arXiv:1907.02189, 2019.
- Liang, X., Liu, Y., Chen, T., Liu, M., and Yang, Q. Federated transfer reinforcement learning for autonomous driving. arXiv preprint arXiv:1910.06001, 2019.
- Lim, H.-K., Kim, J.-B., Heo, J.-S., and Han, Y.-H. Federated reinforcement learning for training control policies on multiple iot devices. *Sensors*, 20(5):1359, 2020.
- Liu, B., Liu, J., Ghavamzadeh, M., Mahadevan, S., and Petrik, M. Finite-sample analysis of proximal gradient TD algorithms. In *Proceedings of the Thirty-First Conference on Uncertainty in Artificial Intelligence*, pp. 504–513, 2015.
- Liu, B., Wang, L., and Liu, M. Lifelong federated reinforcement learning: a learning architecture for navigation in cloud robotic systems. *IEEE Robotics and Automation Letters*, 4(4):4555–4562, 2019.
- Liu, R. and Olshesky, A. Distributed td (0) with almost no communication. arXiv preprint arXiv:2104.07855, 2021.
- Liu, Y., Gottesman, O., Raghu, A., Komorowski, M., Faisal, A. A., Doshi-Velez, F., and Brunskill, E. Representation Balancing MDPs for Off-policy Policy Evaluation. *Advances in Neural Information Processing Systems*, 31: 2644–2653, 2018.
- Maddern, W., Pascoe, G., Linegar, C., and Newman, P. 1 year, 1000 km: The oxford robotcar dataset. *The International Journal of Robotics Research*, 36(1):3–15, 2017.
- Maei, H. R. Convergent actor-critic algorithms under off-policy training and function approximation. Preprint arXiv:1802.07842, 2018.
- McMahan, B., Moore, E., Ramage, D., Hampson, S., and y Arcas, B. A. Communication-efficient learning of deep networks from decentralized data. In *Artificial Intelligence and Statistics*, pp. 1273–1282. PMLR, 2017.
- Mnih, V., Badia, A. P., Mirza, M., Graves, A., Lillicrap, T., Harley, T., Silver, D., and Kavukcuoglu, K. Asynchronous methods for deep reinforcement learning. In *International conference on machine learning*, pp. 1928–1937, 2016.
- Mothukuri, V., Parizi, R. M., Pouriyeh, S., Huang, Y., Dehghantanha, A., and Srivastava, G. A survey on security and privacy of federated learning. *Future Generation Computer Systems*, 115:619–640, 2021.
- Nadiger, C., Kumar, A., and Abdelhak, S. Federated reinforcement learning for fast personalization. In 2019 IEEE Second International Conference on Artificial Intelligence and Knowledge Engineering (AIKE), pp. 123–127. IEEE, 2019.
- Nair, A., Srinivasan, P., Blackwell, S., Alcicek, C., Fearon, R., De Maria, A., Panneershelvam, V., Suleyman, M., Beattie, C., Petersen, S., et al. Massively parallel methods for deep reinforcement learning. arXiv preprint arXiv:1507.04296, 2015.
- Nguyen, D. C., Ding, M., Pathirana, P. N., Seneviratne, A., Li, J., and Poor, H. V. Federated learning for internet of things: A comprehensive survey. arXiv preprint arXiv:2104.07914, 2021.
- Puterman, M. L. Markov decision processes: discrete stochastic dynamic programming. John Wiley & Sons, 2014.
- Qi, J., Zhou, Q., Lei, L., and Zheng, K. Federated reinforcement learning: Techniques, applications, and open challenges. arXiv preprint arXiv:2108.11887, 2021.
- Qu, G. and Wierman, A. Finite-Time Analysis of Asynchronous Stochastic Approximation and Q-Learning. In *Conference on Learning Theory*, pp. 3185–3205. PMLR, 2020.
- Qu, Z., Lin, K., Kalagnanam, J., Li, Z., Zhou, J., and Zhou, Z. Federated learning’s blessing: Fedavg has linear speedup. arXiv preprint arXiv:2007.05690, 2020.
- Rakhlin, A., Shamir, O., and Sridharan, K. Making gradient descent optimal for strongly convex stochastic optimization. In *Proceedings of the 29th International Conference on Machine Learning*, pp. 1571–1578, 2012.
- Recht, B., Re, C., Wright, S., and Niu, F. Hogwild!: A lock-free approach to parallelizing stochastic gradient descent. In Shawe-Taylor, J., Zemel, R., Bartlett, P., Pereira, F., and Weinberger, K. Q. (eds.), *Advances in Neural Information Processing Systems*, volume 24. Curran Associates, Inc., 2011. URL <https://proceedings.neurips.cc/paper/2011/file/218a0aefd1d1a4be65601cc6ddc1520e-Paper.pdf>.
- Ren, J., Wang, H., Hou, T., Zheng, S., and Tang, C. Federated learning-based computation offloading optimization in edge computing-supported internet of things. *IEEE Access*, 7:69194–69201, 2019.
- Shah, D. and Xie, Q. Q-learning with nearest neighbors. In *Advances in Neural Information Processing Systems*, pp. 3111–3121, 2018.

- Shalev-Shwartz, S., Shammah, S., and Shashua, A. Safe, multi-agent, reinforcement learning for autonomous driving. Preprint arXiv:1610.03295, 2016.
- Shalev-Shwartz, S. et al. Online learning and online convex optimization. *Foundations and Trends® in Machine Learning*, 4(2):107–194, 2012.
- Shao, K., Tang, Z., Zhu, Y., Li, N., and Zhao, D. A survey of deep reinforcement learning in video games. arXiv preprint arXiv:1912.10944, 2019.
- Shen, H., Zhang, K., Hong, M., and Chen, T. Asynchronous advantage actor critic: Non-asymptotic analysis and linear speedup. arXiv preprint arXiv:2012.15511, 2020.
- Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., van den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., Dieleman, S., Grewe, D., Nham, J., Kalchbrenner, N., Sutskever, I., Lillicrap, T., Leach, M., Kavukcuoglu, K., Graepel, T., and Hassabis, D. Mastering the game of Go with deep neural networks and tree search. *Nature*, 529(7587):484, 2016.
- Spiridonoff, A., Olshevsky, A., and Paschalidis, I. C. Communication-efficient sgd: From local sgd to one-shot averaging. In *Advances in Neural Information Processing Systems*, volume 34, 2021.
- Srikant, R. and Ying, L. Finite-time error bounds for linear stochastic approximation and TD learning. In *Conference on Learning Theory*, pp. 2803–2830. PMLR, 2019.
- Stich, S. U. Local sgd converges fast and communicates little. In *International Conference on Learning Representations*, 2018.
- Sun, J., Wang, G., Giannakis, G. B., Yang, Q., and Yang, Z. Finite-time analysis of decentralized temporal-difference learning with linear function approximation. In *International Conference on Artificial Intelligence and Statistics*, pp. 4485–4495. PMLR, 2020.
- Sutton, R. S. Learning to predict by the methods of temporal differences. *Machine learning*, 3(1):9–44, 1988.
- Sutton, R. S. and Barto, A. G. Reinforcement learning: An introduction. MIT press, 2018.
- Tadić, V. On the convergence of temporal-difference learning with linear function approximation. *Machine learning*, 42(3):241–267, 2001.
- Truex, S., Baracaldo, N., Anwar, A., Steinke, T., Ludwig, H., Zhang, R., and Zhou, Y. A hybrid approach to privacy-preserving federated learning. In *Proceedings of the 12th ACM Workshop on Artificial Intelligence and Security*, pp. 1–11, 2019.
- Tsitsiklis, J. N. Asynchronous stochastic approximation and Q-learning. *Machine learning*, 16(3):185–202, 1994.
- Tsitsiklis, J. N. and Van Roy, B. Analysis of temporal-difference learning with function approximation. In *Advances in neural information processing systems*, pp. 1075–1081, 1997.
- Wai, H.-T. On the convergence of consensus algorithms with markovian noise and gradient bias. In *2020 59th IEEE Conference on Decision and Control (CDC)*, pp. 4897–4902. IEEE, 2020.
- Wainwright, M. J. Stochastic approximation with cone-contractive operators: Sharp ϵ_1 -bounds for Q-learning. Preprint arXiv:1905.06265, 2019.
- Wang, G., Lu, S., Giannakis, G., Tesauro, G., and Sun, J. Decentralized TD Tracking with Linear Function Approximation and its Finite-Time Analysis. In *Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M. F., and Lin, H. (eds.), Advances in Neural Information Processing Systems*, volume 33, pp. 13762–13772. Curran Associates, Inc., 2020a. URL <https://proceedings.neurips.cc/paper/2020/file/9ec51f6eb240fb631a35864e13737bca-Paper.pdf>.
- Wang, J. and Joshi, G. Cooperative sgd: A unified framework for the design and analysis of local-update sgd algorithms. *Journal of Machine Learning Research*, 22(213):1–50, 2021.
- Wang, X., Wang, C., Li, X., Leung, V. C., and Taleb, T. Federated deep reinforcement learning for internet of things with decentralized cooperative edge caching. *IEEE Internet of Things Journal*, 7(10):9441–9455, 2020b.
- Watkins, C. J. and Dayan, P. Q-learning. *Machine learning*, 8(3-4):279–292, 1992.
- Xu, M., Peng, J., Gupta, B., Kang, J., Xiong, Z., Li, Z., and Abd El-Latif, A. A. Multi-agent federated reinforcement learning for secure incentive mechanism in intelligent cyber-physical systems. *IEEE Internet of Things Journal*, 2021.
- Yang, Q., Liu, Y., Cheng, Y., Kang, Y., Chen, T., and Yu, H. Federated learning. *Synthesis Lectures on Artificial Intelligence and Machine Learning*, 13(3):1–207, 2019.
- Yun, W. J., Park, S., Kim, J., Shin, M., Jung, S., Mohaisen, A., and Kim, J.-H. Cooperative multi-agent deep reinforcement learning for reliable surveillance via autonomous multi-uav control. *IEEE Transactions on Industrial Informatics*, 2022.

- Yurtsever, E., Lambert, J., Carballo, A., and Takeda, K. A survey of autonomous driving: Common practices and emerging technologies. *IEEE Access*, 8:58443–58469, 2020.
- Zeng, S., Doan, T. T., and Romberg, J. Finite-time analysis of decentralized stochastic approximation with applications in multi-agent and multi-task learning. *arXiv preprint arXiv:2010.15088*, 2020.
- Zhang, S., Liu, B., Yao, H., and Whiteson, S. Provably convergent two-timescale off-policy actor-critic with function approximation. In *International Conference on Machine Learning*, pp. 11204–11213. PMLR, 2020.
- Zhang, S. Q., Lin, J., and Zhang, Q. A multi-agent reinforcement learning approach for efficient client selection in federated learning. *arXiv preprint arXiv:2201.02932*, 2022.
- Zhao, L., Ran, Y., Wang, H., Wang, J., and Luo, J. Towards cooperative caching for vehicular networks with multi-level federated reinforcement learning. In *ICC 2021-IEEE International Conference on Communications*, pp. 1–6. IEEE, 2021.
- Zhuo, H. H., Feng, W., Lin, Y., Xu, Q., and Yang, Q. Federated deep reinforcement learning. *arXiv preprint arXiv:1901.08277*, 2019.

Appendices

The appendices are organized as follows. In section A we discuss the lower bound on the convergence of general stochastic approximation. In Section B we derive the convergence bound of FedSAM algorithm. Next, we employ the results in Section B to derive the convergence bounds of federated TD-learning in Section C and federated Q-learning in Section D.

A. Lower Bound on the Convergence of General Stochastic Approximation

In this section we discuss the convergence of general stochastic approximation. In this discussion we provide a simple stochastic approximation with iid noise which can give insight into the general convergence bound in (3). In particular, we show that the convergence bound in (3) is tight and cannot be improved.

Consider a one dimensional random variable X with zero mean $E[X] = 0$ and bounded variance $E[X^2] = \sigma^2$. Consider the following update

$$x_{t+1} = x_t + (X_t - x_t); \quad t \geq 0; \quad (9)$$

where we start with some fixed deterministic x_0 and X_t is an iid sample of the random variable X . It is easy to see that the update (9) is a special case of the update of the general stochastic approximation with the fixed point $x^* = 0$.

By expanding the update (9), we have

$$x_t = (1 - \eta)^t x_0 + \sum_{k=0}^{t-1} (1 - \eta)^{t-k-1} \eta X_k$$

Hence, we have

$$x_t^2 = (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 X_k^2 + 2(1 - \eta)^t x_0 \sum_{k=0}^{t-1} (1 - \eta)^{t-k-1} \eta X_k$$

Taking expectation on both sides, and using the zero mean property of X_k , we have

$$\begin{aligned} E[x_t^2] &= (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 E[X_k^2] + 2(1 - \eta)^t x_0 \sum_{k=0}^{t-1} (1 - \eta)^{t-k-1} \eta E[X_k] \\ &= (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 \sigma^2 + 2(1 - \eta)^t x_0 \sum_{k=0}^{t-1} (1 - \eta)^{t-k-1} \eta E[X_k] \\ &= (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 \sigma^2 + 2(1 - \eta)^t x_0 \sum_{k=0}^{t-1} (1 - \eta)^{t-k-1} \eta E[X_k] \quad (\text{iid property}) \\ &= (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 \sigma^2 \quad (\text{zero mean}) \\ &= (1 - \eta)^{2t} x_0^2 + \sum_{k=0}^{t-1} (1 - \eta)^{2(t-k-1)} \eta^2 \sigma^2 \\ &= x_0^2 (1 - \eta)^{2t} + \frac{1 - (1 - \eta)^{2t}}{2} \sigma^2 \\ &= \underbrace{\left(x_0^2 \frac{1 - (1 - \eta)^{2t}}{2} \right)}_{T_1: \text{bias}} + \underbrace{\left(\frac{1 - (1 - \eta)^{2t}}{2} \right) \sigma^2}_{T_2: \text{variance}} \quad (10) \end{aligned}$$

It is clear that (10) has the same form as the bound in (3) with T_1 as the geometric term which converges to zero as $t \rightarrow \infty$, and T_2 term proportional to the step size η . In addition, note that in the above derivation, we did not use any inequality, and hence the bounds in (10) as well as (3) are tight.

B. Analysis of Federated Stochastic Approximation

First, we restate Theorem 6.1 with explicit expressions of the different constants.

Theorem B.1. Consider the FedSAM Algorithm 4 with $c_{\text{FSAM}} = 1 - \frac{1}{2}$ ($\frac{1}{2}$ is defined in (14)), and suppose Assumptions 6.1, 6.2, 6.3, 6.4 are satisfied. Consider small enough step size α which satisfies the assumptions in (21), (23), (32), (35). Furthermore, denote $\beta = d2 \log e$, and take large enough T such that $T > \max\{K + \frac{1}{\alpha}; 2g\}$. Then, the output of the FedSAM Algorithm 4, $\bar{\theta}_T = \frac{1}{N} \sum_{i=1}^N \theta_t^i$, satisfies

$$E[k_T^2] \leq C_1 \frac{1}{N} \frac{1}{1 - \frac{1}{2}} \frac{2+1}{2} + C_2 \frac{1}{N} + C_3 (K - 1)^2 + C_4 \beta^2; \quad (11)$$

where $C_1 = 16d_{\text{cm}} M_0 (\log \frac{1}{e} + \frac{1}{2})$, $M_0 = \frac{1}{l_{\text{cm}}^2} \frac{1}{C_1^2} B + (A_2 + 1) k_0 k_c + \frac{B}{2C_1} + k_0 k_c^2 C_2$
 $= 8u_{\text{cm}} C_8 \frac{1}{2} + C_{11} \frac{1}{2}$, and $C_3 = 80B^2 C_{17} u_{\text{cm}} \frac{1}{1 + \beta(1 - \frac{1}{2})} = \frac{1}{2} u_{\text{cm}}$ and $C_4 = 8u_{\text{cm}} C_7$
 $\frac{1}{2} + C_{11} + 0.5C_3 C_9 \frac{1}{2} C_3 C_{10} + 2C_1 C_3 C_{10} + 3A_1 C_3 + C_{13} = \frac{1}{2}$. Here the constants $u_{\text{cm}}, l_{\text{cm}}, \frac{1}{2}, C_1, C_3, C_7, C_8, C_9, C_{10}, C_{11}, C_{12}, C_{13}, C_{17}, B$ are problem dependent constants which are defined in the following proposition and lemmas.

Next, we characterize the sample complexity of the FedSAM Algorithm 4, where we establish a linear speedup in the convergence of the algorithm.

Corollary B.1.1. Consider FedSAM Algorithm 4 with fixed number of iterations T and step size $\alpha = \frac{8 \log(NT)}{1 - \frac{1}{2}}$. Suppose T is large enough, such that α satisfies the requirements of the step size in Theorem B.1 and $T > 4$. Also take $K = T = N$. Then we have $E[k_T^2]$ after $T = O(N)$ iterations.

Corollary B.1.1 establishes the sample and communication complexity of FedSAM Algorithm 4. The $O(1/N)$ sample complexity shows the linear speedup with increasing N . Another aspect of the cost is the number of communications required between the agents and the central server. According to Corollary B.1.1, we need $T = N = O(N)$ rounds of communications in order to reach an ϵ -optimal solution. Hence, even in the presence of Markov noise, the required number of communications is independent of the desired final accuracy ϵ , and grows linearly with the number of agents. Our result generalizes the existing result achieved for the simpler i.i.d. noise case in (Khaled et al., 2020; Spiridonoff et al., 2021).

In the following sections, we discuss the proof of Theorem B.1. In Section B.1, we introduce some notations and preliminary results to facilitate our Lyapunov-function based analysis. Next, in Section B.2, we state some primary propositions, which are then used to prove Theorem B.1. In Sections B.3, B.4, B.5, we prove the aforementioned propositions. Along the way, we state several intermediate results, which are stated and proved in Sections B.6 and B.7, respectively.

Throughout the appendix we have several sets of constants. The constants $C_i; i = 1; \dots; 17$ are problem dependent constants which we define recursively. The final constants which appears in the resulting bound in Theorem B.1 are shown as $C_i; i = 1; \dots; 4$. Finally, the constant c is used in the sampling of the time step T .

B.1. Preliminaries

We define the following notations:

- $\bar{\theta}_t = \frac{1}{N} \sum_{i=1}^N \theta_t^i$: virtual sequence of average (across agents) parameter.
- $\theta_t = \{\theta_t^i; i = 1; \dots, N\}$: set of local parameters at individual nodes.
- $Y_t = \{y_t^i; i = 1; \dots, N\}$: Markov chains at individual nodes.
- π^i : the stationary distribution of y^i as $t \rightarrow \infty$.
- $G(t; Y_t) = \frac{1}{N} \sum_{i=1}^N G^i(t; y_t^i)$: average of the noisy local operators at the individual local parameters.
- $G(t; \bar{\theta}_t) = \frac{1}{N} \sum_{i=1}^N G^i(t; \bar{\theta}_t)$: average of the noisy local operators at the average parameter.
- $b(Y_t) = \frac{1}{N} \sum_{i=1}^N b^i(y_t^i)$: average of Markovian noise.

$$E_{t-2} [M(t+1)]$$

$$\begin{aligned}
 & \frac{1}{2} \frac{2LA_2}{2} \frac{2L(A_1+1)}{2} \frac{2+26Lu^2_m}{2} \frac{m_1}{2} + \frac{(A_2+1)^2}{2} E_t \frac{2[M(t)]_{cs}}{2} \quad (15) \\
 & + \frac{2}{2} \frac{m_2 L}{2} E_t \frac{2[k_t - k_c]}{2} \quad (16) \\
 & + \frac{L}{2} \frac{L}{2} + A_2 \frac{1}{2} + \frac{u^2}{2} + (A_1+1) \frac{u^2}{2} + 3m_1^2 E_t \frac{2[k_t - k^2]}{2} \quad (17) \\
 & + \frac{L}{2} \frac{L}{2} + 3u^2(A_1+31) + A_2 + \frac{3}{2} \frac{L}{2} E_t \frac{2[A_1^2]}{2} \quad (18) \\
 & + \frac{L}{2} \frac{3u^{2m}}{2} + 3(A_1+1) + m_1 E_t \frac{2}{2} \quad (19) \\
 & + \frac{m_2^2 u_{cm}^2 L^2}{2} \frac{2}{2} \quad (20) \\
 & + \frac{1}{2} \frac{3L}{2} E_t \frac{2[kb(Y_t)k^2]}{2} \quad (21)
 \end{aligned}$$

where $1, 2, 3$, and 4 are arbitrary positive constants.

Proof. The proof of Proposition B.2 is presented in full detail in Section B.3. $\square$

Before discussing the bound in its full generality, we discuss a few special cases.

- Perfect synchronization ($K = 1$) with i.i.d. noise: since $t = 0$ for all t , $= 0$ (independence across time), and $C_7 = 0$ (see Lemmas B.2 and B.5 in Section B.6), the terms (15), (16) (17), (18), (19) will not appear in the bound, which is the form we get for centralized systems with i.i.d. noise (Rakhlin et al., 2012).
- Infrequent synchronization ($K > 1$) with i.i.d. noise: $= 0$, and $C_7 = 0$, the terms (15), (16), (18), (19) will not appear in the bound, which is the form we get for federated stochastic optimization with i.i.d. noise (Khaled et al., 2020).
- Perfect synchronization ($K = 1$) with Markov noise: since $t = 0$ for all t , the bound in Proposition B.2 generalizes the results in (Srikant & Ying, 2019; Chen et al., 2021c).

Next, we substitute the bound on $k_t - k_c$ (Lemma B.6) and $k_t - k^2$ (Lemma B.7) to further bound the one-step drift. Establishing a tight bound for these two quantities are essential to ensure linear speedup.

Proposition B.3 (One-step drift - II). Consider the update of the FedSAM Algorithm 4. Suppose Assumptions 6.1, 6.2, 6.3, and 6.4 are satisfied. Define $C_{15}() = \frac{6m_1 L u_{cm}^2}{2} + \frac{6L(A_2+1)^2 u^2}{2} + 144(A_2^2 + 1)C_{cm}^2$. For

$$\min \frac{1}{360(A_2+1)^2} \frac{2}{2C_{15}()} ; \quad (21)$$

$= d2 \log e$, and $t \geq 2$, we have

$$E_t \frac{2[M(t+1)]}{2} \leq (1 - \frac{1}{2}) E_t \frac{2[M(t)]}{2} + C_{14}()^4 + C_{16}()^2 + C_{17} \sum_{k=t}^T E_t \frac{2[k]}{2}$$

where $C_{14}() = C_7 + C_{11} + 0.5C_3^2C_9^2 + C_3C_{10} + 2C_1C_3C_{10} + 3A_1C_3 + C_{13} + C_8\frac{u_{c,D}^2}{l_{c,D}^2}2m_2^2$, $C_{16}() = C_8\frac{u_{c,D}^2}{l_{c,D}^2}B^2 + \frac{1}{2} + C_{12}$ and $C_{17} = (3A_1C_3 + 8A^2C_4 + C_5 + C_6)$. Here we define $C_9 = \frac{8}{l_{c,D}}\frac{u_{c,D}}{B}$, $C_{10} = \frac{8m_2u_{c,D}}{l_{c,D}(1-\gamma)}$, $C_{11} = \frac{8C_1^2C_3^2u_{c,D}^2}{l_{c,D}^2}$, $C_{12} = \frac{8C_4u_{c,D}^2B^2}{l_{c,D}^2}$, $C_{13} = \frac{14C_4u_{c,D}^2m_2^2}{l_{c,D}^2(1-\gamma)}$.

Proof. The proof of Proposition B.3 is presented in full detail in Section B.4. $\square$

Conditional expectation $E_{t-2}[\cdot]$. The conditional expectation $E_{t-2}[\cdot]$ used in Proposition B.3 is essential when dealing with Markovian noise. The idea of using conditional expectation to deal with Markovian noise is not novel per se. In the previous work (Bhandari et al., 2018; Srikant & Ying, 2019; Chen et al., 2021c), conditioning on t is sufficient to establish the convergence results. Due to the mixing property (Assumption 6.1), the Markov chain geometrically converges to its stationary distribution. Therefore, choosing “large enough”, and conditioning on t , one can ensure that the Markov chain at time t is “almost in steady state.” However, in federated setting, conditioning on t results in bounds that are too loose. In particular, consider the differences $k_t - k_c$ and $k_t - k^2$ in (15) and (16) respectively. In the centralized setting, as in (Bhandari et al., 2018; Srikant & Ying, 2019; Chen et al., 2021c), these terms can be bounded deterministically to yield 2 bound. However, in the federated setting, this crude bound does not result in linear speedup in N . In this work, to achieve a finer bound on $k_t - k_c$, we go steps further back in time. This ensures that the difference behaves almost like the difference of average of i.i.d. random variables, resulting in a tighter bound (see Lemma B.6). By exploiting the conditional expectation $E_{t-2}[\cdot]$, we derive a refined analysis to bound this term as $O(2=N+4)$, which guarantees a linear speedup (see Lemmas B.6 and B.7).

Taking total expectation in Proposition B.3 (using tower property), we get

$$E[M(t+1)] \leq (1-\gamma_2)E[M(t)] + C_{14}()^4 + C_{16}()^2 + C_{17}\sum_{k=t}^t E[k] \quad (22)$$

To understand the bound in Proposition B.3, consider the case of $K = 1$ (i.e. full synchronization). In this case we have $i = 0$ for all i , and the bound in Proposition B.3 simplifies to $E[M(t+1)] \leq (1-\gamma_2)E[M(t)] + C_{14}()^3 + C_{16}()^2$. This recursion is sufficient to achieve linear speedup. However, the bound in Proposition B.3 also include terms which are proportional to the error due to synchronization. In order to ensure convergence along with linear speedup, we need to further upper bound this term with terms which are of the order $O(3(K-1))$ and $M(i)$. The following lemma is the next important contribution of the paper where we establish such a bound for weighted sum of the synchronization error. Notice that the weights w_t are carefully chosen to ensure the best rate of convergence for the overall algorithm.

Proposition B.4 (Synchronization Error). Suppose $T > K + 1$. For t such that

$$\min\left\{\frac{1}{2}, \frac{\ln(5=4)}{-2(1+\tilde{A}_1)(K-1)^2}\right\} \leq \frac{2(\log() + 1)}{2K^2} \leq \frac{v}{2} \quad (23)$$

$$2(\log() + 1)^3 \exp\left\{-\frac{v}{2}\right\} \leq 2\log + 1 \leq 4K \exp\left\{-\frac{v}{2}\right\} \leq \frac{\ln(5=4)}{2} \leq \frac{1}{(1+\tilde{A}_1)}$$

where $v = \frac{2}{80\tilde{A}_2C_{17}u_{c,D}^2}$, the weighted consensus error satisfies

$$\frac{2C_{17}}{2W_T} \sum_{t=2}^T w_t X^T E \sim \frac{10C}{1+B(1-\gamma)} \frac{4m_2}{(1+\gamma)(K-1)+1} \frac{1}{2W_T} \sum_{t=0}^T w_t E M(t) \quad (24)$$

Proof. The proof is presented in Section B.5. $\square$

Finally, by incorporating the results in Propositions B.2, B.3, and B.4, we can establish the convergence of FedSAM in Theorem B.1.

Proof of Theorem B.1. Assume $w_0 = 1$, and consider the weights w_t generated by the recursion $w_t = w_{t-1} - \frac{1}{2} \nabla \ell(w_{t-1})$. Multiplying both sides of (22) by $\frac{1}{2} w_t$ and rearranging the terms, we get

$$\begin{aligned} & \frac{1}{2} w_t \mathbb{E} M(t) - \frac{1}{2} w_t \mathbb{E} M(t+1) + \frac{1}{2} w_t \mathbb{E} M(t+1) - \frac{1}{2} w_t \mathbb{E} M(t+1) \\ & = \frac{1}{2} [w_{t-1} \mathbb{E} M(t) - w_t \mathbb{E} M(t+1)] + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N + \frac{1}{2} w_t C_{17} X^t \mathbb{E} [\\ & \quad \cdot]; \end{aligned} \quad (25)$$

where we use $w_{t-1} = w_t - \frac{1}{2} \nabla \ell(w_{t-1})$. Summing (25) over $t = 2$ to T (define $W_T = \sum_{t=2}^T w_t$), we get

$$\begin{aligned} & \frac{1}{W_T} X^T \sum_{t=2}^T w_t \mathbb{E} M(t) \\ & = \frac{1}{2 W_T} [w_2 \mathbb{E} M(2) - w_T \mathbb{E} M(T+1)] + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N - \frac{1}{W_T} X^T \sum_{t=2}^T w_t \\ & \quad + \frac{1}{W_T} \sum_{t=2}^T w_t C_{17} X^t \mathbb{E} [\\ & \quad \frac{2 w_{t-1}}{2} \cdot \frac{1}{2} \mathbb{E} M(t) - \frac{1}{2} \mathbb{E} M(t+1) + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N + \frac{1}{2} w_t C_{17} X^t \mathbb{E} [\\ & \quad \frac{1}{2} \cdot \frac{1}{2} \mathbb{E} M(t) - \frac{1}{2} \mathbb{E} M(t+1) + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N + \frac{1}{2} w_t C_{17} X^t \mathbb{E} [\\ & \quad \cdot]; \end{aligned} \quad (26)$$

Substituting the bound on $\frac{1}{W_T} \sum_{t=2}^T w_t \mathbb{E} M(t)$ from Proposition B.4 into (26), we get

$$\begin{aligned} & \frac{1}{W_T} X^T \sum_{t=2}^T w_t \mathbb{E} M(t) \leq \frac{1}{2} \frac{2^{T-2+1}}{2} \mathbb{E} M(2) + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N \\ & \quad + \frac{1}{2} B^{\frac{1}{2}} \frac{O_G}{2} \left(1 + \frac{4m_2}{B(1)} \right) (K+1) + \frac{1}{2 W_T} \sum_{t=0}^T w_t \mathbb{E} M(t) \\ & \quad \frac{1}{W_T} X^T \mathbb{E} [M(t)] \leq \frac{1}{2} \frac{2^{T-2+1}}{2} M_0 + \frac{1}{2} [C_{14}()^3 + C_{16}()]_N \\ & \quad + \frac{1}{2} B^{\frac{1}{2}} \frac{O_G}{2} \left(1 + \frac{4m_2}{B(1)} \right) (K+1) + \frac{1}{2} \frac{2^{T-2+1}}{2} M_0; \end{aligned} \quad (27)$$

where M_0 is a problem dependent constant and is defined in Lemma B.6. To simplify (27), we define $C_{18}() = B^{\frac{1}{2}} \frac{40 C_{17}}{2} \left(1 + \frac{4m_2}{B(1)} \right)$, $C_{19}() = \frac{1}{2} C_{16}()$. We have

$$\frac{1}{W_T} X^T \sum_{t=2}^T w_t \mathbb{E} M(t) \leq \frac{1}{2} \frac{2^{T-2+1}}{2} M_0 + \frac{1}{2} [C_{14}()^3 + C_{18}() (K+1)^2 + C_{19}()]_N \quad (28)$$

Furthermore, define $\tilde{W}_T = \sum_{t=0}^T w_t$. By definition of T_Λ we have $\mathbb{E}[M(T_\Lambda)] = \frac{1}{\tilde{W}_T} \sum_{t=0}^{T_\Lambda} w_t \mathbb{E} M(t)$, and hence

$$\mathbb{E}[M(T_\Lambda)] = \frac{1}{\tilde{W}_T} \sum_{t=0}^{T_\Lambda} w_t \mathbb{E} M(t) + \frac{1}{\tilde{W}_T} X^T \sum_{t=2}^T w_t \mathbb{E} M(t)$$

where $C_{20}(\cdot) = 4M_0 + \frac{4M}{2}$. Furthermore, by Proposition B.1, we have $M(\cdot) = \frac{1}{2}k_m^2 + \frac{1}{2}k_T^2$, and hence

where $C_1() = 2u_{cm} \tilde{e}_{20}()$, $C_2() = 2u_{cm} C_{19}()$, $C_3() = 2u_{cm} C_{18}()$ and $C_4() = 2u_{cm} : {}^4C_{14}()$. $\quad \square$

Furthermore, we have $C_1(\cdot) = 2u_{cm} C_{19}(\cdot) = 2u_{cm', 2} C_{16}(\cdot) = \frac{8u_{cm} (C_8 + \frac{1}{2} + C_{12})}{2} \frac{8u_{cm} C_8 + \frac{1}{2} + C_{12}}{2} = C_2$, where we denote to emphasize the dependence of on , and $C_2 = \frac{8u_{cm} (C_8 + \frac{1}{2} + C_{12})}{2} \frac{8u_{cm} C_8 + \frac{1}{2} + C_{12}}{2}$. Note that we have $= O(\log(1))$.

In addition, $C_3() = 2u_{cm}C_{18}() = \frac{80B^2C_{17}u_{cm}(1+B(1))}{C_4() = \frac{80B^2C_{17}u_{cm}^2(1+B(1))}{2}}$: C_3 , where C_3 And lastly, $= \frac{80B^2C_{17}u_{cm}^2(1+B(1))}{2}$.
 $C_4() = \frac{8u_{cm}(C_7+C_{11}+0.5C_3C_9+C_3C_{10}+2C_1C_3C_{10}+3A_1C_3+C_{13})}{8u_{cm}C_7+C_{11}+0.5C^2C^2+C_3C_{10}+2C_1C_3C_{10}+3A_1C_3+C_{13}} = \frac{C_4^2}{2}$, where $C_4 = \frac{C_4^2}{2}$. $\square$

Next we will state the proof of Corollary B.1.1.

Proof of Corollary B.1.1.3 By this choice of step size, for large enough T , will be small enough and can satisfy the requirements of the step size in (21), (23), (32), (35). Furthermore, the first term in (29) will be

[illegible]

$$= \mathcal{O} \left(\frac{C_1' \cdot 2}{NT} \right) :$$

Furthermore, for the second term we have

$$C_2^N = \mathcal{O} \left(\frac{C_2'^2 \cdot 2}{NT} \right) :$$

Finally, for the third and the fourth terms we have

$$C_3(K-1)^2 + C_4^{32} = \mathcal{O} \left(\frac{C_3'^2 + C_4'^3}{N} \right) \frac{\log^2(NT)}{T^2} = \mathcal{O} \left(\frac{C_3'^2 + C_4'^3}{NT} \right) \cdot \frac{2}{NT}$$

Upper bounding (29) with , we get $\mathcal{O} \left(\frac{C_1' \cdot 2 + C_2'^2 \cdot 2 + C_3'^2 + C_4'^3}{N} \right) \cdot \frac{NT}{T^2}$. Hence, we need to have $T =$

$\mathcal{O} \left(\frac{C_1' \cdot 2 + C_2'^2 \cdot 2 + C_3'^2 + C_4'^3}{N} \right) \cdot \frac{NT}{T^2}$ number of iterations to get to a ball around the optimum with radius . $\square$

B.3. Proof of Proposition B.2

The update of the virtual parameter sequence $f_t g$ can be written as follows

$$f_{t+1} = f_t + (G(f_t; Y_t) - f_t + b(Y_t)) : \quad (30)$$

Using $\frac{p-1}{p}$ -smoothness of $M(\cdot)$ (Proposition B.1), we get

$$\begin{aligned} M(f_{t+1}) - M(f_t) &+ hrM(f_t); f_{t+1} - f_t + \frac{L}{2} \|k_{t+1} - k\|_S^2 \quad (\text{Smoothness of } M(\cdot)) \\ &= M(f_t) + hrM(f_t); G(f_t; Y_t) - f_t + b(Y_t) + \frac{L^2}{2} \|kG(f_t; Y_t) - f_t + b(Y_t)\|_S^2 \\ &\quad + \underbrace{rM(f_t); G(f_t) - f_t + hrM(f_t); b(Y_t)}_{T_1: \text{Expected update}} \\ &\quad + \underbrace{hrM(f_t); G(f_t; Y_t) - G(f_t)}_{T_2: \text{Error due to Markovian noise } b(Y_t)} + \underbrace{hrM(f_t); G(f_t) - G(f_t)}_{T_3: \text{Error due to Markovian noise } Y_k} \\ &\quad + \underbrace{\frac{L^2}{2} \|kG(f_t; Y_t) - f_t + b(Y_t)\|_S^2}_{T_4: \text{Error due to local updates}} \\ &\quad + \underbrace{\frac{L^2}{2} \|kG(f_t; Y_t) - f_t + b(Y_t)\|_S^2}_{T_5: \text{Error due to noise and discretization}} \end{aligned} \quad (31)$$

The inequality in (31) characterizes one step drift of the Lyapunov function $M(f_t)$. The term T_1 is responsible for negative drift of the overall recursion. T_2 and T_3 appear due to the presence of the Markovian noises $b^i(y^i)$ and $G^i(i; y^i)$ in the update of Algorithm 4. T_4 appears due to the mismatch between the parameters of the agents $i; i = 1; \dots; N$. Finally, T_5 appears due to discretization error in the smoothness upper bound. Next, we state bounds on $T_1; T_2; T_3; T_4; T_5$ in the following intermediate lemmas.

Lemma B.1. For all $2 \leq R^d$, the operator $G(\cdot)$ satisfies the following

$$T_1 = rM(f_t); G(f_t) - 2' \cdot 2 M(f_t) :$$

Lemma B.1 guarantees the negative drift in the one step recursion analysis of Proposition B.2. It follows from the Moreau envelope construction (Chen et al., 2021c) and the contraction property of the operators $G^i(\cdot); i = 1; \dots; N$ (Assumption 6.2).

Lemma B.2. Consider the iteration t of the Algorithm 4, and consider $\epsilon = d \log e$. We have

$$\mathbb{E}_t [T_2] = \mathbb{E}_t [hrM(f_t); b(Y_t)]$$

$$\begin{aligned} & \frac{L}{2} \frac{E_t^2}{l_{cs}^2} k_t^2 - k_c + \frac{E_t^2}{2} k_b(Y_t) k_c^2 \\ & + \frac{m_2 L}{l_{cs}^2} E_t k_t - k_c + \frac{1}{2} \frac{L m_2^2 u_{cm}^2}{l^2} + E_t [M(t)];_{cs} \end{aligned}$$

where ϵ_1 is an arbitrary positive constant.

In the i.i.d. noise setting, $E[T_2] = 0$. In Markov noise setting, going back steps (which introduces t) enables us to use Markov chain mixing property (Assumption 6.1).

Lemma B.3. For any $t \geq 0$, denote $T_3 = \text{tr} M(t); G(t; Y_t) - G(t)$. For any $\epsilon < t$, we have $E_t [T_3]$

$$\begin{aligned} & \frac{2LA_2}{2} + \frac{L(A_1 + 1)}{l_{cs}^2} + \frac{6m_1 L u_{cm}}{3 l_{cs}^2} E_t [M(t)]^2 \\ & + \frac{2LA_2}{l_{cs}^2} \frac{1}{2} + \frac{u_{cm}^2}{2^2} + \frac{L(A_1 + 1)}{l_{cs}^2} \frac{3u_{cm}^2}{2^2} + 2 + \frac{3m_1 L}{l_{cs}^2} E_t k_t^2 \\ & + \frac{3u_{cm}^2}{2^2} + \frac{3}{2} \frac{L(A_1 + 1)}{l_{cs}^2} + \frac{LA_2}{l_{cs}^2} E_t [\\ & + \frac{3u_{cm}^2}{2^2} + \frac{3}{2} \frac{L(A_1 + 1)}{l_{cs}^2} + \frac{m_1^2 L}{2 l_{cs}^2} E_t [\\ &] \end{aligned}$$

where ϵ_2 and ϵ_3 are arbitrary positive constants.

Lemma B.4. For any $t \geq 0$, we have

$$T_4 = \text{tr} M(t); G(t) - G(t) \leq 4 M(t) + \frac{L^2 u_{cm}^2}{2 l_{cs}^2} + \frac{L^2 u_{cm}^2}{4}$$

where ϵ_4 is an arbitrary positive constant.

Lemma B.5. For any $0 \leq t$, we have

$$T_5 = k G(t; Y_t) - k_t + b(Y_t) k_c^2 \leq \frac{6(A_2 + 1)^2 u_{cm}^2}{l_{cs}^2} M(t) + \frac{3A_1}{l_{cs}^2} + \frac{1}{l_{cs}^2}$$

Substituting the bounds in Lemmas B.1, B.2, B.3, B.4, B.5, and taking expectation, we get the final bound in Proposition B.2.

B.4. Proof of Proposition B.3

First we state the following two intermediate lemmas, which are proved in Section B.7.

Lemma B.6. Suppose $\epsilon = d \log e$ and

$$\min \left(\frac{1}{12(A_2 + 1)}; \frac{1}{8(A_2 + 1)} \right) : \quad (32)$$

For any $0 \leq t \leq 2$ we have the following

$$M(t) \leq \frac{1}{l_{cm}^2} + \frac{1}{C_1^2} B + (A_2 + 1) k_0 k_c + \frac{B}{2C_1^2} + k_0 k_c^2 - M_0 \quad (33)$$

Furthermore, for any $t \geq 2$, we have the following

$$\begin{aligned} E_t [k_t - k_c] & \leq 4C_1 E_t [k_t k_c] + 8 \frac{B}{l_{cd}} + \frac{u_{cd}}{l_{cd}} \left(\frac{1}{1} + \frac{u_{cd}}{2C_1} \right) \\ & + 6A_1 E_t [i]_{i=t} \end{aligned} \quad (34)$$

Lemma B.7. Suppose $\beta = d^2 \log e$ and

$$\min \frac{1}{C_1}; \frac{1}{4C_2}; \frac{1}{40C_1^2} ; \quad (35)$$

where $C_1 = 3 \frac{\beta}{A_2^2 + 1}$ and $C_2 = 3C_1 + 8$. We have the following

$$\begin{aligned} E_{t-2} [k_t - t - k^2] &\leq 8^2 C^2 E_{t-2} [k_t^2] + \frac{8 \beta C}{l^2} \frac{B^2}{C} \frac{u^2}{l^2} \frac{1}{N} \\ &\quad + 14 \frac{u^2}{l^2} \frac{m^2}{1} \frac{1}{2} + 8^2 A^2 \sum_{i=1}^t E_{t-2} [k_i^2] \end{aligned}$$

In Lemma B.6, we define $C_9 = \frac{8u_{CD}B}{l_{CD}}$; $C_{10} = \frac{8m_2 u_{CD}}{l_{CD}(1-\beta)}$. Hence, we can bound the term in (15) as

$$\begin{aligned} &2C_3 E_{t-2} [k_t - t - k_c] \\ &\leq 2C_3 \left(4C_1 E_{t-2} [k_t k_c] + C_9 \frac{1}{N} + C_{10}^2 (1 + 2C_1) + 6A_1 \sum_{k=t}^t E_{t-2} [k] \right) \\ &\leq 2C_3 \left(4C_1 u_{cm} E_{t-2} \frac{1}{2M(t)} + C_9 \frac{1}{N} + C_{10}^2 (1 + 2C_1) + 6A_1 \sum_{k=t}^t E_{t-2} [k] \right) \quad (\text{Proposition B.1}) \\ &= \frac{1}{6} 4C_1 C_3 u_{cm}^5 \frac{\beta}{6} \frac{1}{2} E_{t-2} [M(t)] \\ &\quad + 2C_3 \left(C_9 \frac{1}{N} + C_{10}^2 (1 + 2C_1) + 6A_1 \sum_{k=t}^t E_{t-2} [k] \right) \quad (\beta \text{ is concave, } \beta > 0) \\ &\leq \frac{8^2 C_1^2 u_{cm}^2}{C} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] + C_3 C_9 \frac{1}{N} + C_3 C_{10}^2 (1 + 2C_1) \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] \\ &\quad + 6A_1 C_3 \sum_{k=t}^t E_{t-2} \left[\frac{1}{2} + \frac{1}{2} k \right] \quad (\text{Young's inequality}) \\ &\leq \frac{8C_1^2}{2} \frac{C_3^2 u_{cm}^2}{6} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] + C_3 C_9 \frac{1}{N} + 4[C_3 C_{10}^2 (1 + 2C_1) + 3A_1 C_3 (\beta + 1)] \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] \\ &\quad + 3A_1 C_3 \sum_{k=t}^t E_{t-2} [k] \quad (\text{By (58)}) \\ &\leq C_{11} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] + \frac{1}{2} C^2 C^{24} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] + 4[C_3 C_{10}^2 (1 + 2C_1) + 3A_1 C_3 (\beta + 1)] + 3A_1 C_3 \sum_{k=t}^t E_{t-2} [k] \end{aligned} \quad (36)$$

Furthermore, using Lemma B.7, (58) and Proposition B.1, the term in (16) can be bounded as follows:

$$\begin{aligned} &C_4 E_{t-2} [k_t - t - k^2] \leq 16^2 C^2 C_4 u_m^2 E_{t-2} [M(t)] + 8^2 A^2 C_4 \sum_{k=0}^t E_{t-2} [k] \\ &\quad + \frac{8C_4 u_{CD}^2 B^2}{l_{CD}^2} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] + \frac{14C_4 u_{CD} m^2}{l_{CD}^2 (1-\beta)} \frac{1}{2} \frac{1}{6} \frac{1}{2} E_{t-2} [M(t)] \end{aligned} \quad (37)$$

$$\frac{2C_{17}}{W} \sum_{t=2}^T X_t^T W_t X_t^t E$$

$$\frac{2C_{17}}{2} \frac{X^T}{W_T} w_t X^t 5^2 (s' \cdot K) B^2 \hat{1} + \frac{4m_2}{B(1)} + 5^2 (s' \cdot K) A_2 \sim \frac{X^1}{E k_{t^0} K C^2} : \quad (42)$$

where $s' = b' = Kc$. Hence, $s' \cdot K$ denotes the last time instant before $'$ when synchronization happened. The first term in (42) can be upper bounded as follows.

$$\begin{aligned} & \frac{2C_{17}}{2} \frac{X^T}{W_T} w_t X^t 5^2 (s' \cdot K) B^2 \sim 1 + \frac{B}{(1)} \frac{4m_2}{(1)} \\ & = 2 \sim 2 \frac{1}{B} + \frac{4m_2}{B(1)} \frac{1}{2} \frac{1}{W_T} \frac{X^T}{t=2} w_t X^t (s' \cdot K) \frac{1}{2} B^2 \frac{1}{1} \\ & + \frac{B}{(1)} \frac{1}{2} \frac{4m_2}{W_T} \frac{1}{t=2} \frac{w_t}{C} \frac{(K)}{1} X^t \frac{1}{1} X^t \quad (s' \cdot K \quad K \quad 1) \\ & = 2 B^2 \sim 1 + \frac{1}{B(1)} \frac{1}{2} \frac{1}{W_T} \frac{X^T}{t=2} w_t \frac{1}{\#} + 1)(K \quad 1) \\ & 2 B^2 \sim 1 + \frac{4m_2}{B(1)} \frac{1}{2} \frac{1}{W_T} \frac{1}{t=2} \frac{1}{C} \frac{1}{(1)} (K \quad 1): \end{aligned} \quad (43)$$

Next, we compute the second term in (42).

$$\begin{aligned} & \frac{2C_{17}}{2} \frac{X^T}{W_T} w_t X^t 5^2 (s' \cdot K) A_2 \sim \frac{X^1}{k_{t^0} k_c^2} \\ & 2 A_2 \sim \frac{20C_{17} u_{cm}^2}{2} \frac{X^T}{W_T} w_t X^t (s' \cdot K) \frac{X^1}{M(t^0)^2} \quad (\text{Proposition B.1}) \\ & 2 A_2 \sim \frac{20C_{17} u_{cm}^2}{2} \frac{X^T}{W_T} w_t X^t (s' \cdot K) \frac{X^1}{M(t^0)^2} + \frac{X^T}{W_T} w_t X^t (s' \cdot K) \frac{X^1}{M(t^0)^2} \quad (44) \end{aligned}$$

where if $K < 2$, $l_1 = 0$. Next, we bound l_1 ; l_2 separately.

$$\begin{aligned} l_1 &= \frac{X^K}{t=2} w_t X^t (s' \cdot K) \frac{X^1}{M(t^0)} \\ & (K \quad 1) \frac{X^K}{t=2} w_t X^t \frac{X^1}{M(t^0)} \quad (\text{for } ' < K, s' = 0; \text{ for } ' = K, s' \cdot K = 0) \\ & (K \quad 1)(+1) \frac{X^K}{t=2} w_t X^t \frac{X^1}{M(t^0)} \\ & (K \quad 1)(K \quad 2 + 1)(+1) w_K \frac{X^K}{t=0} \frac{X^1}{M(t)}; \end{aligned} \quad (45)$$

where, (45) follows since $w_{t-1} = w_t$; $8 t$. Next, to bound l_2 in (44), we again split it into two terms.

$$l_2 = \frac{X^{K+1}}{t=K+1} w_t X^t (s' \cdot K) \frac{X^1}{M(t^0)} + \frac{X_T}{t=K+1} w_t X_t (s' \cdot K) \frac{X^1}{M(t^0)} :$$

$$\begin{aligned}
 & (K-1)(t+1) \mathbb{E} \left[\sum_{s=0}^{t-1} \gamma^s \mathbf{X}^T \mathbf{X}^s \mathbf{w}_{t+K+1} \mathbf{M}(\cdot) + \mathbf{w}_T^T \mathbf{X}^{t-1} \mathbf{M}(\cdot) \right] \\
 & = (K-1)(t+1) \mathbb{E} \left[\sum_{s=0}^{t-1} \gamma^s \mathbf{X}^T \mathbf{X}^s \frac{\mathbf{w}}{1 - \frac{\gamma}{2}^{K+1}} \mathbf{M}(\cdot) + \frac{\mathbf{w}}{1 - \frac{\gamma}{2}^{t-K+1}} \mathbf{M}(\cdot) \right] \\
 & \quad (w_t = w_{t+1} - \frac{\gamma}{2}) \\
 & (K-1)(t+1) \mathbb{E} \left[\sum_{s=0}^{t-1} \gamma^s \mathbf{X}^T \mathbf{X}^s \mathbf{w}_t \mathbf{M}(t) \right] \tag{47}
 \end{aligned}$$

Using the bounds on I_3 ; I_4 from (46) and (47) respectively, we can bound I_2 .

$$I_2 \leq w_{K+K^2} \mathbb{E} \left[\sum_{t=0}^{K-1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + 3w_{K+} \mathbb{E} \left[\sum_{t=K}^{K+1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + (K-1)(t+1) \mathbb{E} \left[\sum_{t=K}^{t-1} \gamma^s \mathbf{X}^T \mathbf{X}^s \mathbf{w}_t \mathbf{M}(t) \right] \tag{48}$$

Substituting the bounds on I_1 ; I_2 from (45), (48) respectively, into (44), we get

$$\begin{aligned}
 & \frac{2C_{17}}{W_T} \mathbb{E} \left[\sum_{t=2}^T \gamma^t \mathbf{X}^T \mathbf{X}^t \mathbf{w}_t \mathbf{M}(t) \right] \leq 5^2 (\gamma^s K) A_2 \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + k_0 k^2 \frac{2}{c} \\
 & \quad + 2A_2 \frac{20C_{17} u_{cm}^2}{W} (K-1)(K+2+1)(t+1) w_K \mathbb{E} \left[\sum_{t=0}^{K-1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + K^2 w_{K+} \mathbb{E} \left[\sum_{t=0}^{K-1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \\
 & \quad + 2A_2 \frac{20C_{17} u_{cm}^2}{W_T} \mathbb{E} \left[\sum_{t=K}^{K+1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + (K-1)(t+1) \mathbb{E} \left[\sum_{t=K}^{t-1} \gamma^s \mathbf{X}^T \mathbf{X}^s \mathbf{w}_t \mathbf{M}(t) \right] \tag{49}
 \end{aligned}$$

We analyze the terms in (49) separately. First, for the terms with $\mathbb{E} \left[\sum_{t=0}^{K-1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right]$,

$$\begin{aligned}
 & \frac{1}{W_T} 2A_2 \frac{20C_{17} u_{cm}^2}{W} K(t+1) \left[(K+2+1)w_K + w_{K+} \right] \mathbb{E} \left[\sum_{t=0}^{K-1} \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \\
 & = \frac{1}{W_T} 2A_2 \frac{20C_{17} u_{cm}^2}{W} K(t+1) \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \\
 & = \frac{1}{W_T} 2A_2 \frac{20C_{17} u_{cm}^2}{W} K(t+1) \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + \frac{1}{W_T} \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \\
 & \quad + 2A_2 \frac{20C_{17} u_{cm}^2}{W} K(t+1) \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] + \frac{1}{W_T} \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \\
 & \quad + \frac{1}{2W_T} \mathbb{E} \left[\sum_{t=0}^{K-1} \gamma^s \mathbf{X}^T \mathbf{X}^t \mathbf{M}(t) \right] \tag{50}
 \end{aligned}$$

where (50) holds since we choose γ small enough such that

$$2A_2 \frac{20C_{17} u_{cm}^2}{W} (K+2+1) + \frac{1}{1 - \frac{\gamma}{2}^{K+1}} < \frac{1}{2}$$

To get this, we use the inequality $1 - x \leq \exp(-x)$ for $x < 1$, $\frac{1}{2} < 2 \log + 1$, we get $\frac{1}{(1 - \frac{\gamma}{2}^2)}$ $\exp(-\frac{\gamma}{2} 2 \log + 1)$. For (50) to hold, it is sufficient that

$$\frac{1}{80C_{17} u_{cm} A_2 K^2 (\log(t) + 1)} \min \left\{ \frac{1}{K}; \frac{1}{2(\log(t) + 1) \exp(-\frac{\gamma}{2} 2 \log + 1)} \right\} \tag{51}$$

$$\frac{r M_f^g(x)}{D} \stackrel{?}{=} k x k_m ;$$

B.7. Proof of Lemmas

Proof of Lemma B.1. Using Cauchy-Schwarz inequality, we have

$$\left| \frac{1}{m} \sum_{i=1}^m \langle \mathbf{z}_i, \mathbf{G}(\mathbf{t}) \rangle \right| \leq \sqrt{\frac{1}{m} \sum_{i=1}^m \|\mathbf{z}_i\|^2} \sqrt{\frac{1}{m} \sum_{i=1}^m \|\mathbf{G}(\mathbf{t})\|^2} \quad (59)$$

where $\|\cdot\|_m$ denotes the dual of the norm $\|\cdot\|_m$. Furthermore, by Proposition B.1 we have (Lemma B.11)

$$\begin{aligned} (a) \quad T_{11} &= \|\mathbf{G}(\mathbf{t})\|_m^2 \\ &= \left\| \frac{1}{N} \sum_{i=1}^N \mathbf{G}^i(\mathbf{t}) \right\|_m^2 \\ &\leq \frac{1}{N} \sum_{i=1}^N \|\mathbf{G}^i(\mathbf{t})\|_m^2 \quad (\text{triangle inequality}) \\ &= \frac{1}{N} \sum_{i=1}^N \|\mathbf{G}^i(\mathbf{t})\|_m^2 \\ &= \frac{1}{N} \sum_{i=1}^N \|\mathbf{G}^i(\mathbf{t})\|_m^2 \quad (\text{Assumption 6.2, Proposition B.1}) \end{aligned}$$

Furthermore, by the convexity of the $\|\cdot\|_m$ norm (Lemma B.11), we have (60)

$$T_{12} \leq \|\mathbf{G}(\mathbf{t})\|_m^2$$

Hence, using (60), (61) in (59) we get (61)

$$\left| \frac{1}{m} \sum_{i=1}^m \langle \mathbf{z}_i, \mathbf{G}(\mathbf{t}) \rangle \right| \leq \frac{1}{\sqrt{m}} \|\mathbf{G}(\mathbf{t})\|_m \quad (62)$$

where $\|\cdot\|_2$ is defined in (14). $\square$

Proof of Lemma B.2. Given some $t < T$, $T_2 = \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle$ can be written as follows:

$$\begin{aligned} T_2 &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \\ &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \\ &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \quad (\text{Cauchy-Schwarz}) \\ &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \quad (\text{Proposition B.1}) \end{aligned} \quad (63)$$

Taking expectation on both sides, we have

$$\begin{aligned} \mathbb{E}_t [T_2] &= \frac{1}{m} \sum_{i=1}^m \mathbb{E}_t [\langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle] \\ &= \frac{1}{m} \sum_{i=1}^m \mathbb{E}_t [\langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle] \quad (64) \end{aligned}$$

For T_{21} , we have

$$\begin{aligned} T_{21} &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \\ &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \quad (\text{Assumption 6.1}) \\ &= \frac{1}{m} \sum_{i=1}^m \langle \mathbf{b}(\mathbf{Y}_t), \mathbf{G}(\mathbf{t}) \rangle \quad (\text{assumption on } \mathbf{G}) \end{aligned}$$

$$\frac{m_2^2 L}{l_{cs}^2} \left[k_t - \frac{1}{l_{cs}} \left(k_t k_s + k_r M(t) - r M(0) k_s \right) \right] \quad (64)$$

$$\frac{m_2^2 L}{l_{cs}^2} \left[k_t - \frac{1}{l_{cs}} \left(k_t k_c + k_t k_c \right) \right] \quad (65)$$

$$\frac{m_2^2 L}{l_{cs}^2} \left[k_t - \frac{1}{l_{cs}} \left(k_t k_c + \frac{m_2^2 u_{cm} L}{l_{cs}^2} \left(\frac{k_t k_{cm}}{l_{cs}^2} - \frac{1}{2} \right) \right) \right] \quad (66)$$

$$\frac{m_2^2 L}{l_{cs}^2} \left[k_t - \frac{1}{l_{cs}} \left(k_t k_c + \frac{1}{2} \left(\frac{m_2^2 u_{cm} L}{l_{cs}^2} - \frac{1}{2} \right) + \frac{1}{2} M(t) \right) \right] \quad (67)$$

Substituting (65) in (64), we get

$$E_t [T_2] \leq \frac{L}{2} \left(\frac{1}{l_{cs}^2} E_t [k_t^2] + \frac{1}{2} E_t [k_b(Y_t) k^2] + \frac{L}{l_{cs}} E_t \left[k_t \frac{m_2}{2} k_c + \frac{1}{2} \left(\frac{m_2^2 u_{cm} L}{l_{cs}^2} - \frac{1}{2} \right) + \frac{1}{2} E_t [M(t)] \right] \right)$$

Proof of Lemma B.3.

□

$$\begin{aligned} T_3 &= \frac{1}{N} \sum_{i=1}^N \left(r M(t) - r M(t); G(t; Y_t) - G(t) \right) \\ &= \frac{1}{N} \sum_{i=1}^N \left(\underbrace{r M(t) - r M(t); G(t; Y_t) - G(t)}_{T_{31}} \right) \\ &\quad + \frac{1}{N} \sum_{i=1}^N \left(\underbrace{r M(t); G(t; Y_t) - G(t; Y_t) + G(t) - G(t)}_{T_{32}} \right) \\ &\quad + \frac{1}{N} \sum_{i=1}^N \left(\underbrace{r M(t); G(t; Y_t) - G(t) - G(t)}_{T_{33}} \right) \end{aligned} \quad (68)$$

Next, we bound all three terms individually.

I. Bound on T_{31} :

$$\begin{aligned} T_{31} &= \frac{1}{N} \sum_{i=1}^N \left(r M(t) - r M(t); G^i(t; y^i) - G^i(t) \right) \\ &= \frac{1}{N} \sum_{i=1}^N \left(\underbrace{r M(t) - r M(t); G^i(t; y^i) - G^i(t)}_{T_{311}} \right) + \frac{1}{N} \sum_{i=1}^N \left(\underbrace{r M(t); G^i(t; y^i) - G^i(t)}_{T_{312}} \right) \end{aligned} \quad (69)$$

For T_{311} , we have

$$T_{311} = \frac{1}{N} \sum_{i=1}^N \left(k_r M(t) - k_r M(t); k_s \right) \leq \frac{L}{l_{cs}} \left(k_t - k_c \right) \quad (70)$$

where the inequality follows from Proposition B.1 and (13). For T_{312} , we have,

$$\begin{aligned} T_{312} &= \frac{1}{N} \sum_{i=1}^N \left(k G^i(t; y^i) - E_{y^i} G^i(t; y) \right) k_{\xi} \\ &\leq \frac{1}{N} \sum_{i=1}^N \left(\underbrace{G^i(t; y^i) - E_{y^i} G^i(t; y)}_{T_{3121}} \right) k_{\xi} \quad (\text{Triangle and Jensen's inequality}) \\ &\leq \frac{1}{N} \sum_{i=1}^N \left(\underbrace{A_2 k_t k_c + A_2 k_t k_c}_{2A_2} \right) \quad (\text{Assumption 6.3}) \\ &\leq \frac{1}{N} \sum_{i=1}^N \left(\underbrace{k_t k_c + k_t k_c}_{k_t k_c} \right) \quad (\text{triangle inequality}) \\ &= \frac{2A_2}{N} \sum_{i=1}^N k_t k_c : cs \end{aligned} \quad (71)$$

[illegible]

(74)

III. Bound on T_{33} : Taking expectation on both sides of T_{33} , we have

(75)

Substituting the bounds on $T_{31}; T_{32}; T_{33}$ from (70), (74), (75), in (66), we get the result.

Proof of Lemma B.4. Denote $T_4 = \text{hrM}(t); G(t) - G(t)$. By the Cauchy–Schwarz inequality, we have

(76)

For T_{41} we have

(77)

For T_{42} we have

$$\begin{aligned} kG(t) - G(t)k_s &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t) = \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (\text{Jensen's inequality}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} = \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (78) \end{aligned}$$

where (78) follows from Assumption 6.2. Combining (77) and (78), for an arbitrary $\epsilon > 0$, we get

$$\begin{aligned} T_4 &\leq \frac{L u_{cm}}{l} \frac{1}{2} M(t) + \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (c_1) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + \frac{L u_{cm}}{l} \frac{1}{2} M(t) + \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + \frac{L u_{cm}}{l} \frac{1}{2} M(t) + \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (\text{Young's inequality}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + \frac{L u_{cm}}{l} \frac{1}{2} M(t) + \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (79) \end{aligned}$$

□

Proof of Lemma B.5. Denote $T_5 = kG(t; Y_t) - G(t; Y_t)k_s + b(Y_t)k_s$. We have

$$\begin{aligned} T_5 &= kG(t; Y_t) - G(t; Y_t)k_s + b(Y_t)k_s = \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + b(Y_t)k_s \quad (80) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + b(Y_t)k_s = \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} + b(Y_t)k_s \end{aligned}$$

For T_{51} we have

$$\begin{aligned} T_{51} &\leq l_{cs} (kG(t; Y_t)k_c + k_t k_c)^2 l_{cs} \quad (\text{By (13)}) \\ &= (A_2 k_t k_c + k_t k_c)^2 = 2(A_2 + 1)^2 u_{cm}^2 M(t) \quad (\text{Assumption 6.3}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t - G^i(t)_{s_i=1} \quad (81) \end{aligned}$$

For T_{52} we have

$$\begin{aligned} T_{52} &= \frac{1}{N} \sum_{i=1}^N (G^i(i)_t - G^i(t; y^i))_t^2 = \frac{1}{N} \sum_{i=1}^N G^i(i)_t^2 - 2 G^i(i)_t G^i(t; y^i) + G^i(t; y^i)^2 \quad (k_{s_i}^2 \text{ is convex}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t^2 - 2 G^i(i)_t G^i(t; y^i) + G^i(t; y^i)^2 \quad (\text{By (13)}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t^2 - 2 G^i(i)_t G^i(t; y^i) + G^i(t; y^i)^2 \quad (\text{Assumption 6.3}) \\ &= \frac{1}{N} \sum_{i=1}^N G^i(i)_t^2 - 2 G^i(i)_t G^i(t; y^i) + G^i(t; y^i)^2 \quad (82) \end{aligned}$$

Using the bounds in (81), (82), we get

$$T_5 \leq \frac{6(A_2 + 1)^2 u_{cm}^2}{l^2} M(t) + \frac{3A_1}{l^2} \quad (83)$$

where the last inequality is by the assumption on .

□

Proof of Lemma B.6. Define $\bar{l} = \frac{1}{N} \sum_{i=1}^N l_i$. By the update rule of Algorithm 4, if $l + 1 \bmod K = 0$, we have

$$\begin{aligned} l_{l+1} &= \frac{1}{N} \sum_{i=1}^N l_{l+1}^i = \frac{1}{N} \sum_{i=1}^N \left(G^i(l_i; y_i) - l_i + b^i(y_i) \right)_{i=1}^N \\ &= \frac{1}{N} \sum_{i=1}^N \left(k G^i(l_i; y_i) k_{c_i} + k l_i k_c + k b^i(y_i) k_{c_i} \right)_{i=1}^N \\ &= \frac{1}{N} \sum_{i=1}^N \left((A_2 + 1) k l_i k_c + B^i \right)_{i=1}^N \\ &= (1 + (A_2 + 1)) \bar{l} + B; \end{aligned} \quad \begin{array}{l} \text{(triangle inequality)} \\ \text{(Assumption 6.3)} \end{array} \quad (84)$$

Furthermore, if $l + 1 \bmod K \neq 0$, we have $l_{l+1} = \frac{1}{N} \sum_{j=1}^N \left(l_j + (G^j(l_j; y_j) - l_j + b^j(y_j)) \right)$, and hence

$$\begin{aligned} l_{l+1} &= \frac{1}{N} \sum_{i=1}^N l_{l+1}^i = \frac{1}{N} \sum_{i=1}^N \left(l_i + \frac{1}{N} \sum_{j=1}^N (G^j(l_j; y_j) - l_j + b^j(y_j)) \right)_{i=1}^N \\ &= \frac{1}{N} \sum_{i=1}^N \left(k G^i(l_i; y_i) k_{c_i} + k l_i k_c + k b^i(y_i) k_{c_i} \right)_{i=1}^N \end{aligned} \quad \begin{array}{l} \text{(triangle inequality)} \end{array}$$

and the same bound as in (84) holds. By recursive application of (84), we get

$$\begin{aligned} l_{l+1} &= (1 + (A_2 + 1))^{l+1} l_0 + B \sum_{i=0}^l (1 + (A_2 + 1))^i \\ &= (1 + (A_2 + 1))^{l+1} l_0 + B \frac{(1 + (A_2 + 1))^{l+1} - 1}{A_2 + 1}. \end{aligned} \quad (85)$$

Notice that for $x \leq \frac{\log 2}{2}$, we have $(1 + x)^{l+1} \leq 1 + 2x(l + 1)$. If $0 \leq l \leq 1$, by the assumption on ϵ , we have $(1 + (A_2 + 1))^{l+1} \leq (1 + (A_2 + 1))^2 \leq 1 + 4(A_2 + 1) \leq 2$. Hence, we have

$$l_{l+1} \leq 2 l_0 + B \frac{4(A_2 + 1)}{2} = 2 l_0 + 4B \leq 2 \quad (86)$$

Furthermore, we have

$$\begin{aligned} k_{l+1} - l_k &= k G(l_l; y_l) - l_l + b(y_l) k_{c_l} - k G(l_l; y_l) k_{c_l} \\ &= \frac{1}{N} \sum_{i=1}^N \left(G^i(l_i; y_i) - l_i + b^i(y_i) \right)_{i=1}^N + k b(y_l) k_{c_l} \\ &= \frac{1}{N} \sum_{i=1}^N \left(G^i(l_i; y_i) - l_i + b^i(y_i) \right)_{i=1}^N + B^i \\ &= \frac{1}{N} \sum_{i=1}^N \left((A_2 + 1) l_i k_c + B^i \right)_{i=1}^N \\ &= (A_2 + 1) \bar{l} + B; \end{aligned} \quad \begin{array}{l} \text{(triangle inequality)} \\ \text{(convexity of norm)} \\ \text{(Assumption 6.3)} \end{array} \quad (87)$$

Suppose $0 \leq t \leq 2$. We have

$$k_t - l_0 k_c \leq \sum_{k=0}^{t-1} (k_{k+1} - l_k k_c) \quad \text{(triangle inequality)}$$

$$\frac{1}{C_1} B + (A_2 + 1) \cdot 0 + \frac{B}{2C_1} : \quad (t \ 2)$$

$$M(t) = \frac{1}{2l_{cm}^2} (k_t \cdot 0 + 0 \cdot k_c)^2_{cm} = \frac{1}{2l_{cm}^2} (k_t \cdot 0 + k_0 k_c)^2_{cm} \quad (\text{triangle inequality})$$
[illegible]
$$k_{l+1} - k_c \leq \frac{q}{A_2 + 1} k_c + 2kb(Y_l)k_c + 2A_{11} \quad (90)$$
$$k_{l+1}k_c \left(1 + 3 \frac{q}{A_2 + 1} k_l k_c + 2kb(Y_l)k_c + 2A_{1l} \right)$$

$$= (1 + C_1)k_l k_c + 2kb(Y_l)k_c + 2A_1 l; \quad (91)$$

where we denote $C_1 = 3 \frac{p}{A_2^2 + 1}$. Assuming $t \geq l$, and taking expectation on both sides, we have

$$\begin{aligned} & E_{t-2}[k_{l+1}k_c] + (1 + C_1)E_{t-2}[k_l k_c] + 2E_{t-2}[kb(Y_l)k_c] + 2A_1 E_{t-2}[l] \\ &= (1 + C_1)E_{t-2}[k_l k_c] + 2 \frac{u_{cD}}{B} + 2m_2^{l-t+2} + 2A_1 E_{t-2}[l] \quad (\text{Lemma B.9}) \\ &= (1 + C_1)E_{t-2}[k_l k_c] + 2 \frac{u_{cD}}{B} + 2m_2^{l-t+1} + 2A_1 E_{t-2}[l] \quad (\text{Assumption on } \rho) \\ &= (1 + C_1)E_{t-2}[k_l k_c] + c_t(l) + 2A_1 E_{t-2}[l]; \end{aligned} \quad (92)$$

where $c_t(l) = 2 \frac{u_{cD}}{B} + 2m_2^{l-t+1}$. By applying this inequality recursively, we have

$$\begin{aligned} & E_{t-2}[k_{l+1}k_c] + (1 + C_1)E_{t-2}[k_l k_c] + c_t(l) + 2A_1 E_{t-2}[l] \\ &= (1 + C_1)[(1 + C_1)E_{t-2}[k_{l-1}k_c] + c_t(l-1) + 2A_1 E_{t-2}[l-1]] + c_t(l) + 2A_1 E_{t-2}[l] \\ &= (1 + C_1)^2 E_{t-2}[k_{l-2}k_c] + (1 + C_1)c_t(l-1) + c_t(l) + 2A_1 E_{t-2}[l] \\ &= \dots \\ &= (1 + C_1)^{t-l+1} E_{t-2}[k_t k_c] + (1 + C_1)^{t-l} c_t(l) + c_t(l) + 2A_1 E_{t-2}[l] \\ &= (1 + C_1)^{t-l+1} E_{t-2}[k_t k_c] + (1 + C_1)^{t-l} c_t(l) + c_t(l) + 2A_1 E_{t-2}[l] \end{aligned} \quad (93)$$

We study T_1 and T_2 in (93) separately. For T_1 we have

$$\begin{aligned} T_1 &= 2 \frac{u_{cD}}{B} \sum_{k=0}^{t-l+1} (1 + C_1)^k + 2m_2^{t-l+1} \\ &= 2 \frac{u_{cD}}{B} \frac{(1 + C_1)^{t-l+1} - 1}{C_1} + 2m_2^{t-l+1} \\ &= 2 \frac{u_{cD}}{B} \frac{(1 + C_1)^{t-l+1} - 1}{C_1} + 2m_2^{t-l+1} \quad (\text{due to } C_1 > 0) \\ &= 2 \frac{u_{cD}}{B} \frac{(1 + C_1)^{t-l+1} - 1}{C_1} + 2m_2^{t-l+1} \end{aligned} \quad (94)$$

Notice that for $x \geq 0$, we have $(1 + x)^{t-l+1} \geq 1 + 2x(t-l+1)$. By the assumption on ρ , we have $(1 + C_1)^{t-l+1} \geq 1 + 2C_1(t-l+1)$ and $(1 + C_1) \geq 1 + 2C_1$. Hence, we have

$$T_1 \geq 2 \frac{u_{cD}}{B} \frac{1 + 2C_1(t-l+1)}{C_1} + 2m_2^{t-l+1} \quad (95)$$

Furthermore, for the term T_2 we have

$$\begin{aligned} T_2 &= \sum_{k=0}^{t-l+1} (1 + C_1)^{t-l+k} + \sum_{k=0}^{t-l+1} (1 + C_1)^{t-l+k} \\ &= \sum_{k=0}^{t-l+1} (1 + 2C_1)^{t-l+k} + \sum_{k=0}^{t-l+1} (1 + C_1)^{t-l+k} \end{aligned} \quad (\text{due to } C_1 > 0) \quad (96)$$

Substituting (95), (96) in (93), for every $t \geq l$, we get

$$E_{t-2}[k_l k_c] \leq 2E_{t-2}[k_t k_c] + \frac{2u_{cD}}{B} + \frac{4m_2}{1} + \frac{4A_1}{1} \sum_{k=0}^{t-l} \dots \quad (97)$$

But we have

$$\begin{aligned}
 & \mathbb{E}_{t-2} k_t - k_c \leq \mathbb{E}_{t-2} \sum_{i=t}^t \# k_{i+1} - k_c \quad (\text{triangle inequality}) \\
 & \leq \mathbb{E}_{t-2} [C_1 k_i k_c + 2kb(Y_i)k_c + 2A_{1i}] \quad (\text{by (90)}) \\
 & \leq \sum_{i=t}^{t-1} C_1 4\mathbb{E}_{t-2} [k_t - k_c] + 2u_{cD} \frac{4B}{l} + \frac{4m_2}{p} + 4A_1 \frac{\mathbb{E}_{t-2} [t-j]}{1} \sum_{j=0}^3 X_{t-j}^{cD} \quad (\text{by (97)}) \\
 & \quad + 2 \sum_{i=t}^t \frac{u_{cD}}{p} + 2m_2 \frac{1}{N} + 2A_1 \mathbb{E}_{t-2} [i] \quad (\text{Lemma B.9}) \\
 & \leq C_1 4\mathbb{E}_{t-2} [k_t - k_c] + 2u_{cD} \frac{4B}{l} + \frac{4m_2}{p} + 4A_1 \sum_{j=0}^3 \mathbb{E}_{t-2} [t-j] X_{t-j}^{cD} \quad (\text{assumption on }) \\
 & \quad + 2 \frac{l_{cD} u_{cD}}{B} + 4 \frac{l_{cD} m_2}{1} + 2A_1 \mathbb{E}_{t-2} [i] \\
 & \leq 2C_1 \mathbb{E}_{t-2} [k_t - k_c] + u_{cD} \frac{B}{l_{cD}} + \frac{8C_1^2}{l_{cD} p} + 2 + \frac{8m_2}{l_{cD} C} + \frac{u_{cD} 4m_2^2}{1} + \frac{1}{1} \\
 & \quad + (4A_1^2 C_1 + 2A_1) \sum_{i=t}^t \mathbb{E}_{t-2} [i] \quad (i=0) \\
 & \leq 2C_1 \mathbb{E}_{t-2} [k_t - k_c] + 4 \frac{u_{cD}}{B} + \frac{u_{cD}}{l} \frac{4m_2}{p} + \frac{2(1+2C_1)}{1} + 3A_1 \sum_{i=t}^t \mathbb{E}_{t-2} [i] \quad (98)
 \end{aligned}$$

Furthermore, by triangle inequality, we have $k_t - k_c \leq k_t - k_{t-1} + k_{t-1} - k_c + k_t k_c$. By assumption on , we have $2C_1 k_t - k_c \leq 2C_1 k_t - k_c + 2C_1 k_t k_c \leq 5k_t - k_c + 2C_1 k_t k_c$. By taking expectation on both sides, and substituting it in (98), we get (34).

Proof of Lemma B.7. From (91) we have

$$\begin{aligned}
 & k_{l+1} k^2 \left((1 + C_1)^2 k_l k^2 + 4^2 kb(Y_l) k^2 + 4^2 A_c^{22} \right) \\
 & \quad + \frac{4(1 + C_1) k_l k_c kb(Y_l) k_c + 4A_1 (1 + C_1) k_l k_c l + 8^2 A_{1l} kb(Y_l) k_c}{T_1} : \left\{ \frac{z}{T_2} \right\} \left\{ \frac{z}{T_3} \right\} \quad (99)
 \end{aligned}$$

For T_1 we have

$$\begin{aligned}
 T_1 &= 2 \frac{p}{(1 + C_1) k_l k_c} - 2 \frac{p}{(1 + C_1) kb(Y_l) k_c} \\
 & \leq 2(1 + C_1) k_l k_c + 2(1 + C_1) kb(Y_l) k_c + 4k_l k_c + 2 \quad (\text{ab } \frac{1}{2}a^2 + \frac{1}{2}b^2) \\
 & \leq 4kb(Y_l) k_c; \quad (100)
 \end{aligned}$$

where the last inequality is by the assumption on . Analogously for T_2 we have

$$\begin{aligned}
 T_2 &= 2 \frac{p}{(1 + C_1) k_l k_c} - 2A_1 \frac{p}{(1 + C_1) l} \\
 & \leq 4k_l k_c + 2A_{1l} : \quad (101)
 \end{aligned}$$

For T_3 we have

$$T_3 = 2kb(Y_I)k_c + 4A_{11} + 2^2kb(Y_I)k^2 + 8^2A^{22} \quad (102)$$

Combining the bounds in (100), (101), and (102), and noting that $(1 + C_1)^2 \leq 1 + 3C_1$, from (99) we have

$$\begin{aligned} k_{l+1}k_c &\leq (1 + 3C_1)k_l k_c + 4^2kb(Y_I)k_c + 4^2A_{11}^2 \\ &\quad + 4k_l k^2 + 4kb(Y_I)k^2 + 4k_l k^2 + 4A_c^{22} + 2^2kb(Y_I)k^2 + 8^2A^{22} = (1 + (3C_1 + 1 \\ &\quad 8))k_l k^2 + (6^2 + 4)kb(Y_I)k^2 + A^2(12^2 + 4)^2 \\ &\quad (1 + C_2)k_l k_c + 16kb(Y_I)k_c + 16A_{11} \end{aligned} \quad (103)$$

where $C_2 = 3C_1 + 8$. Taking expectation on both sides, we have

$$\begin{aligned} E_{t-2} k_{l+1} k^2 &\leq (1 + C_2) E_{t-2} k_l k^2 + 10E_{t-2} kb(Y_I)k^2 + 16A^2 E_{t-2} [^2] \\ &\quad (1 + C_2) E_{t-2} k_l k_c + 16 \frac{u_{cD}^2}{l^2} + 2m^{22(l-t+2)} + 16A_1 E_{t-2} [l] \quad \text{(Lemma B.9)} \\ &= (1 + C_2) E_{t-2} k_l k^2 + d(l) + 16A^2 E_{t-2} [^2] \end{aligned} \quad (104)$$

where $d(l) = 10 \frac{u_{cD}^2}{l^2} + 2m^{22(l-t+2)}$. Hence, for any $t \leq l$, we have

$$\begin{aligned} E_{t-2} k_{l+1} k^2 &\leq (1 + C_2) (1 + C_2) E_{t-2} k_l k_c + d(l-1) + 16A_1 E_{t-2} [l-1] \\ &\quad + d(l) + 16A_1 E_{t-2} [l] \\ &\quad (1 + C_2)^{l+1-t} E_{t-2} k_t k^2 + \sum_{i=t}^l (1 + C_2)^{l-i} d(i) + 16A^2 E_{t-2} \sum_{i=t}^l (1 + C_2)^{l-i} \\ &\quad (1 + C_2)^{+1} E_{t-2} k_t k^2 + \sum_{i=t}^l (1 + C_2)^{l-i} d(i) + 16A_1 E_{t-2} \sum_{i=t}^l (1 + C_2)^{l-i} \quad (105) \end{aligned}$$

For the term T_4 we have

$$\begin{aligned} T_4 &= 10 \frac{u_{cD}^2}{l^2} \sum_{i=0}^X (1 + C_2)^i B^2 + 2m^{22(i-cD)} \\ &= 10 \frac{u_{cD}^2}{l^2} \frac{B^2 (1 + C_2)^{+1}}{C_2} + 20 \frac{u_{cD}^2}{l^2} m^{22} (1 + C_2) \\ &\quad 10 \frac{u_{cD}^2}{l^2} \frac{B^2 (1 + C_2)^{+1}}{C_2} + 20 \frac{u_{cD}^2}{l^2} m^{22} (1 + C_2) \sum_{i=0}^{2i} 10 \frac{u_{cD}^2}{l^2} B^2 \\ &\quad (1 + C_2)^{+1} \frac{1}{l^2} + 20 \frac{u_{cD}^2}{l^2} m^{22} (1 + C_2) \frac{1}{l^2} \quad (C_2 \geq 0) \end{aligned} \quad (106)$$

By the same argument as in Lemma B.6, and by the assumption on , we have $(1 + C_2)^{+1} \leq 1 + 2C_2(+ 1) \leq 1 + 4C_2 - 2$ and $(1 + C_2) \leq 1 + 2C_2 - 1 = 2 - 2$. Hence, we have

$$T_4 \leq 40 \frac{u_{cD}^2}{l^2} \frac{B^2}{N} + \frac{m_2^2}{1 - 2} \quad (107)$$

Furthermore, for T_5 , we have

$$T_5 = \sum_{i=0}^X (1 + C_2)^i E_{t-2} [t+i]^2$$

$(C_2 > 0)$

(assumption on)

(108)

Combining the bounds on T_4 (107) and T_5 (108) in (105) we have for any $t \geq t_0$,

(109)

Furthermore, we have

(triangle inequality)

(from (89))

Taking expectation on both sides, and using the bounds in Lemma B.9 and (109), we get

(110)

Furthermore, by triangle inequality, we have $k_t - k_c \leq k_t - t - k_c + k_t k_c$. Squaring both sides, we have $k_t - k^2 \leq (k_t - t - k_c + k_t k_c)^2 \leq 2k_t - t - k^2 + 2k_t k^2$. By assumption on t , we have $2^{22} C^2 k_t - k^2 \leq 4^{22} C^2 k_{t_1} - k^2 + 4^{22} C^2 k_t k^2 \leq 0.5 k_{t_1} - t - k^2 + 4^{22} C^2 k_t k^2$.

Proof of Lemma B.8. For $s \leq K + 1 - t \leq (s + 1)K - 1$, where $s = bt = Kc$,

$$(* \ s^i K = {}_s K \text{ and Young's inequality})$$

$$\begin{aligned}
 & \frac{2}{N} \sum_{i=1}^N \frac{u_{c2}^2}{4} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \quad \text{(Triangle inequality, and } \mathbb{E}[X^2] \leq \mathbb{E}[X^2] \text{)} \\
 & \frac{2}{N} \sum_{i=1}^N \frac{u_{c2}^2}{4} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \quad \text{(} \mathbb{E}[X^2] \leq \mathbb{E}[X^2] \text{ and } \text{Var}(X) \leq \mathbb{E}[X^2] \text{)}
 \end{aligned}$$

Taking expectation

$$\begin{aligned}
 & \mathbb{E} \left[\frac{2}{N} \sum_{i=1}^N \frac{u_{c2}^2}{4} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \quad \text{(using Assumption 6.3)} \\
 & \frac{2^2 u_{c2}^2}{l_{c2}^2} (t - sK) B^2 + \frac{4^2 u_{c2}^2}{l_{c2}^2} \sum_{i=1}^N \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \quad \text{(Assumption 6.3)} \\
 & \text{(a) } \frac{2^2 u_{c2}^2}{l_{c2}^2} (t - sK) B^2 + \frac{4^2 u_{c2}^2}{l_{c2}^2} \sum_{i=1}^N \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \text{(b) } \frac{2^2 u_{c2}^2}{l_{c2}^2} (t - sK) B^2 + \frac{4^2 u_{c2}^2}{l_{c2}^2} \sum_{i=1}^N \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \\
 & + \frac{4^2}{N} (t - sK) \sum_{i=1}^N \mathbb{E} \left[\frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} G^i(i_0; y_{t^0}^{i_0})^2 \right] \\
 & \quad \text{(111) } \frac{1}{l_{c2}^2} \sum_{i=1}^N \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right]
 \end{aligned}$$

where (a) follows since

$$\begin{aligned}
 & \mathbb{E} \left[\frac{1}{4} \sum_{i=1}^N \frac{u_{c2}^2}{l_{c2}^2} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \right] \\
 & \leq \mathbb{E} \left[\frac{1}{4} \sum_{i=1}^N \frac{u_{c2}^2}{l_{c2}^2} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \right] \\
 & \quad \text{(Jensen's inequality)} \\
 & \leq \mathbb{E} \left[\frac{1}{4} \sum_{i=1}^N \frac{u_{c2}^2}{l_{c2}^2} \mathbb{E} \left[\frac{1}{6} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 + \frac{1}{2} \sum_{t=0}^{t^0} b^i(y_{t^0}^{i_0})^2 \right] \right] \\
 & \quad \text{(Assumption 6.1, 6.3)}
 \end{aligned}$$

$$\sim \gamma^{1/h} \quad i^t X^{k-1} \quad \#$$

$$\frac{4^2(t-sK)(1+A_1)}{t^0+E} = \frac{1+4^2(1+A_1)(t'-sK)}{t-k'=1} \quad E$$

$$t^0=sK$$

$$\begin{aligned}
 & 4^2(t-sK)(1+A_1) \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & + 4^2(t-sK)(1+A_1) \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & 4^2(t-sK) B^2 \sim 1 + \frac{4m_2}{B(1)} + \sum_{t'=sK}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} E \sim \frac{1}{2} \quad \# \\
 & t_0 + A_2 E k_{t_0} k_C \quad \text{(using (112))} \\
 & 4^2(t-sK)(1+A_1) \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & + 4^2(t-sK) \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & 4^2(1+A_1)(t-sK) B^2 \sim 1 + \frac{4m_2}{B(1)} + \sum_{t'=sK}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} E \sim \frac{1}{2} \quad \# \\
 & t_0 + A_2 E k_{t_0} k_C \quad : (115)
 \end{aligned}$$

Substituting (115) into (114), we see that the induction hypothesis in (114) holds. We can go back as far as the last instant of synchronization, $j = t - sK$. For $j = t - sK$, we get

$$\begin{aligned}
 E & \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & B^2 \sim 1 + \frac{4m_2}{B(1)} + \sum_{t'=sK}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} E k_{t_0} k_C^2 + (1+A_1)E \\
 & sK : (116)
 \end{aligned}$$

Next, using $1+x \leq e^x$ for $x \geq 0$, we get

$$\begin{aligned}
 & \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \leq \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & \exp \left(-2^2(1+A_1)(t-sK)^2 \right) \\
 & \frac{5}{4} \quad (117)
 \end{aligned}$$

if δ is small enough such that $2^2(1+A_1)(K-1)^2 \ln \frac{5}{4}$, which holds true by the assumption on the step size. Using (117) in (116), we get

$$\begin{aligned}
 E & \sim \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & \sim \frac{1}{(1+A_1)E} \sum_{t'=1}^{t-sK} \frac{1}{1+4^2(1+A_1)(t'-sK)} \sum_{i=t'-sK}^{t'-1} E \\
 & = 5^2(t-sK) B^2 \sim 1 + \frac{4m_2}{B(1)} + 5^2(t-sK) A_2 \sum_{t'=sK}^{t-sK} E k_{t_0} k_C ; \quad (118)
 \end{aligned}$$

which concludes the proof. $\square$

Proof of Lemma B.9. We have

$$\begin{aligned}
 & E_t r[kb(Y_t)k_C] - u_{CD} E_t r[kb(Y_t)k_D] \\
 & = u_{CD} E_t r \left[\frac{q}{b(Y_t) > D b(Y_t)} \right] \\
 & q \frac{1}{b(Y_t) > D b(Y_t)}
 \end{aligned}$$

$$u_{cD} = E_t \left[r \left[b(Y_t) > D b(Y_t) \right] \right]$$

(concavity of square root)

$$\begin{aligned}
 &= u_{cD} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 \\
 &= u_{cD} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 + \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \\
 &= u_{cD} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 + \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \quad (\text{Assumption 6.4}) \\
 &= u_{cD} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 + \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \quad (119)
 \end{aligned}$$

For the term T_1 we have

$$\begin{aligned}
 T_1 &= \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 \\
 &= \frac{B}{l_{cD}} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 \quad (\text{Assumption 6.3}) \\
 &= \frac{B}{l_{cD}} \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right)^2 \quad (120)
 \end{aligned}$$

For the term T_2 we have

$$\begin{aligned}
 T_2 &= \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \\
 &= \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \quad (\text{Cauchy-Schwarz}) \\
 &= \frac{2}{N} \sum_{i < j} \left(\frac{1}{N} \sum_{i=1}^N b^i(y_t) - \frac{1}{N} \sum_{i=1}^N b^i(y_t^i) \right) \left(\frac{1}{N} \sum_{i=1}^N b^j(y_t) - \frac{1}{N} \sum_{i=1}^N b^j(y_t^i) \right) \quad (\text{Assumption 6.1}) \\
 &= \frac{2m_2^r}{l_{cD}} \quad (121)
 \end{aligned}$$

Substituting (120) and (121) in (119), we get the result in (56).

The proof of (57) follows analogously. $\square$

Proof of Lemma B.10. By definition, we have

$$\begin{aligned}
 \|G(t; Y_t) - G(t; Y_t^c)\|_c^2 &= \frac{1}{N} \sum_{i=1}^N \left(G^i(i; y_t^i) - G^i(i; y_t^c) \right)^2 \\
 &= \frac{1}{N} \sum_{i=1}^N \left(G^i(i; y_t^i) - G^i(i; y_t^c) \right)^2 \quad (\text{convexity of norm}) \\
 &= \frac{1}{N} \sum_{i=1}^N \left(A_1 k_t^i - k_c \right)^2 \quad (\text{Assumption 6.3}) \\
 &= (A_1 k_t)^2 \\
 &= \frac{1}{N} \sum_{i=1}^N A_1 k_t^2 - k_c^2 \quad (\text{convexity of square})
 \end{aligned}$$

$$= A_1^2$$

$$t:$$

$$t)$$

(By definition of

Furthermore, by the convexity of $(\cdot)^2$, we have

$$t \geq \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{N} \sum_{i=1}^N \right)^2 = t$$

□

Proof of Lemma B.11. Since $M_f^B(\cdot)$ is convex, and there exists a norm, $\|\cdot\|_m$, such that $M_f^B(x) = \frac{1}{2} \|\cdot\|_m^2$ (see Proposition B.1), using the chain rule of subdifferential calculus,

$$r M_f^B(x) = \|\cdot\|_m u_x;$$

where $u_x \in \partial \|\cdot\|_m$ is a subgradient of $\|\cdot\|_m$ at x . Hence,

$$r M_f^B(x)^2 = \|\cdot\|_m^2 \langle u_x, u_x \rangle;$$

where $\|\cdot\|_m^2$ is the dual norm of $\|\cdot\|_m$. Since $\|\cdot\|_m$ is convex and, as a function of x , is 1-Lipschitz w.r.t. $\|\cdot\|_m$ we have $\|\cdot\|_m^2 \leq 1$ (see Lemma 2.6 in (Shalev-Shwartz et al., 2012)).

Further, by convexity of $\|\cdot\|_m$ norm, $\|\cdot\|_m \leq \|\cdot\|_m + \langle u_x, \cdot \rangle$. Therefore,

$$r M_f^B(x); x = \|\cdot\|_m \langle u_x, x \rangle \leq \|\cdot\|_m^2 = 2 M_f^B(x):$$

□

C. Federated TD-learning

C.1. On-policy Function Approximation

Proposition C.1. On-policy TD-learning with linear function approximation Algorithm 1 satisfies the following:

1. $\bar{v} = v^i_t - v$
2. $S_t = (S_t^1; \dots; S_t^N)$ and $A_t = (A_t^1; \dots; A_t^N)$
3. $y_t^i = (S_t^i; A_t^i; \dots; S_{t+n-1}^i; A_{t+n-1}^i; S_{t+n}^i)$ and $Y_t = (S_t; A_t; \dots; S_{t+n-1}; A_{t+n-1}; S_{t+n})$
4. π : Stationary distribution of the policy.

Furthermore, choose some arbitrary positive constant $\gamma > 0$. The corresponding $G^i(\cdot; y_t^i)$ and $b^i(y_t^i)$ in Algorithm 1 for On-policy TD-learning with linear function approximation is as follows

1. $G^i(\cdot; y_t^i)_t = \bar{v} + \frac{1}{t} \sum_{l=t}^{t+n-1} (S_l^i)^{\top} (S_{l+1}^i)^{\top} - \bar{v}$
2. $b^i(y_t^i) = \frac{1}{t} \sum_{l=t}^{t+n-1} R(S_l^i; A_l^i) + \frac{1}{t} (S_{t+n}^i)^{\top} \bar{v} - \bar{v}$

where $\bar{v}$ solves the projected bellman equation $\bar{v} = ((T)^n \bar{v})$. Furthermore, the corresponding step size α_t in Algorithm 4 is $\frac{1}{t}$.

Lemma C.1. Consider the federated on-policy TD-learning Algorithm 1 as a special case of FedSAM Algorithm 4 (see Proposition C.1). Suppose the trajectory $\{S_t^i\}_{t=0;1;\dots}$ converges geometrically fast to its stationary distribution as follows $d_{TV}(P(S_t^i = j | S_0^i = i)) \leq \rho^t$ for all $i, j = 1; 2; \dots; N$. The corresponding $G^i(\cdot)$ in Assumption 6.1 for the federated TD-learning is as follows

$$G^i() = + \gamma \left(\frac{1}{n} (P)^n \right); \quad (122)$$

where $\gamma > 0$ is an arbitrary constant introduced in Proposition C.1. In addition, (6) holds. Furthermore, for $t \geq n + 1$, we have $m = \max\{ \frac{2A_2 m}{n}, 2Bn \}$, where A_2 and B are specified in Lemma C.3 and $\gamma = \frac{1}{m}$.

Lemma C.2. Consider the federated on-policy TD-learning 1 as a special case of FedSAM (as specified in Proposition C.1). Consider the $jSj \times jSj$ matrix $U = \gamma \left(\frac{1}{n} (P)^n - I \right)$ with eigenvalues $f_1; \dots; f_{jSj}$. Define $\rho_{\max} = \max_i |f_i|$ and $\rho = \max_i \text{Re}[f_i] > 0$, where $\text{Re}[\cdot]$ evaluates the real part. By choosing γ large enough in the linear function (122), there exist a weighted 2-norm $\| \cdot \|_c = \| G^i(\cdot) - G^i(z) \|_c \leq k_1 \| \cdot - z \|_c$ for $c = 1, 2$ for $k_1 = \frac{1}{1 - \rho_{\max}}$.

Lemma C.3. Consider the federated on-policy TD-learning Algorithm 1 as a special case of FedSAM (as specified in Proposition C.1). There exist some constants A_1, A_2 , and B such that the properties of Assumption 6.3 are satisfied.

Lemma C.4. Consider the federated on-policy TD-learning Algorithm 1 as a special case of FedSAM (as specified in Proposition C.1). Assumption 6.4 holds for this algorithm.

C.1.1. PROOFS

Proof of Proposition C.1. Items 1-4 are by definition. Subtracting v from both sides of the update of the TD-learning, we have

$$\begin{aligned} v_{t+1}^i - v &= v_t^i - v + (S_t)^i - (S_{t+1})^i + R(S_t^i; A^i) - (S_{t+1})^i v - (S_t^i)^i v \\ &= \frac{1}{n} \sum_{l=1}^n (S_t^l)^i - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^i + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^i v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^i v \\ &= \frac{1}{n} \sum_{l=1}^n (S_t^l)^i - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^i + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^i v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^i v \\ &= \frac{1}{n} \sum_{l=1}^n (S_t^l)^i - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^i + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^i v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^i v \\ &= \frac{1}{n} \sum_{l=1}^n (S_t^l)^i - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^i + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^i v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^i v \end{aligned}$$

which proves items 1 and 2. Furthermore, for the synchronization part of TD-learning, we have

$$\begin{aligned} v_t^i &= \frac{1}{N} \sum_{j=1}^N v_t^j \\ &= \frac{1}{N} \sum_{j=1}^N \left(\frac{1}{n} \sum_{l=1}^n (S_t^l)^j - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^j + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^j v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^j v \right) \end{aligned}$$

which is equivalent to the synchronization step in FedSAM Algorithm 4. Notice that here we used the fact that all agents have the same fixed point v .

Proof of Lemma C.1. It is easy to observe that

$$G^i(\cdot; y_t)_i = \frac{1}{n} \sum_{l=1}^n (S_t^l)^i - \frac{1}{n} \sum_{l=1}^n (S_{t+1}^l)^i + \frac{1}{n} \sum_{l=1}^n (R(S_t^l; A^l) - (S_{t+1}^l)^i v) - \frac{1}{n} \sum_{l=1}^n (S_t^l)^i v$$

□

Taking expectation with respect to the stationary distribution, we have

$$\begin{aligned}
 G^i() &= E_{S_t} \left[\frac{1}{n} (S_t^i)^n (S_{t+n}^i)^n - (S_t^i)^n \right] \\
 &= E_{S_t} \left[\frac{1}{n} (S_t^i)^n (S_{t+n}^i)^n - (S_t^i)^n \right] \quad (\text{tower property of expectation}) \\
 &= E_{S_t} \left[\frac{1}{n} (S_t^i)^n E[(S_{t+n}^i)^n | S_t^i] - (S_t^i)^n \right] = E_{S_t} \left[\frac{1}{n} (S_t^i)^n (P^n)_{ii} - (S_t^i)^n \right] \\
 &= E_{S_t} \left[\frac{1}{n} (S_t^i)^n (P^n)_{ii} - (S_t^i)^n \right] = E_{S_t} \left[\frac{1}{n} (S_t^i)^n (P^n)_{ii} - (S_t^i)^n \right]
 \end{aligned}$$

where P is the transition probability matrix corresponding to the policy, and $\mathbf{I}$ is a diagonal matrix with diagonal entries corresponding to elements of $\mathbf{I}$.

As explained in (Tsitsiklis & Van Roy, 1997), the projection operator is a linear operator and can be written as $\mathbf{P} = (\mathbf{I} - \gamma \mathbf{T})^{-1} \mathbf{T}$, where $\mathbf{T}$ is a diagonal matrix with diagonal entries corresponding to the stationary distribution of the policy. Hence, the fixed point equation is as follows $\mathbf{v} = (\mathbf{I} - \gamma \mathbf{T})^{-1} \mathbf{T} \mathbf{v}$. Since $\mathbf{T}$ is a full column matrix, we can eliminate it from both sides of the equality, and further multiply both sides with $\mathbf{T}$. We have $\mathbf{T} \mathbf{v} = \mathbf{T} (\mathbf{I} - \gamma \mathbf{T})^{-1} \mathbf{T} \mathbf{v}$, and hence $\mathbf{T} (\mathbf{I} - \gamma \mathbf{T}) \mathbf{v} = \mathbf{0}$, which is equivalent to $E_{S_t} [\gamma (S_t^i) ((\mathbf{T}^n \mathbf{v})(S_t^i) - \mathbf{v}(S_t^i))] = 0$. By expanding $(\mathbf{T}^n)$, we have $E_{S_t} [\gamma (S_t^i) \sum_{l=0}^{n-1} (R(S_t^i; A_t^i) + \mathbf{v}(S_{t+l+1}^i) - \mathbf{v}(S_t^i))] = 0$, which means

$$E_{\mathbf{y}} \mathbf{b}^i(\mathbf{y}) = 0; \quad (123)$$

and proves (6).

Moreover, we have

$$\begin{aligned}
 \mathbf{k} G^i() &= E[G^i(; \mathbf{y}^i)] \mathbf{k}_c = E_{\mathbf{y}^i} [G^i(; \mathbf{y}^i)] = E[G^i(; \mathbf{y}^i)^c] = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) G^i(; \mathbf{y}^i)
 \end{aligned} \quad (\text{Assumption 6.3})$$

For brevity, we denote $P(S_t^i = s_t^i) = P(s_t^i)$. We have

$$\begin{aligned}
 &\sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) \\
 &= \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i) = \sum_{\mathbf{y}^i} P(\mathbf{y}^i = \mathbf{y}^i | \mathbf{y}^i)
 \end{aligned} \quad (t = n + 1)$$

$$\begin{aligned}
 &= \sum_{i=1}^n \mathbb{E} \left[\left(s^i \right)_t^\top P \left(s^i j S_t^i \right) s_h \right] \\
 &= 2d_{TV} \left(P \left(S_t^i = j S_n^i \right) \right) 2m \\
 &= (2m^n) : t
 \end{aligned}$$

$$\mathbb{E}[b^i(y^i)]^c = \mathbb{E}[b^i(y^i)] - \mathbb{E}_{y_t} [b^i(y^i)]_c \quad (\text{By (123)})$$

$$\begin{aligned}
 &= \sum_{i=1}^n \mathbb{E} \left[\left(y^i \right)_t^\top P \left(y^i = y_t^i j y_0^i \right) b^i(y^i) \right] \\
 &= \sum_{i=1}^n \mathbb{E} \left[\left(y^i \right)_t^\top P \left(y^i = y_t^i j y_0^i \right) b^i(y_0^i) \right]_c \\
 &= \sum_{i=1}^n \mathbb{E} \left[\left(y^i \right)_t^\top P \left(y_t = y_t^i j y_0 \right) \dot{B}^i \right] \\
 &= 2Bm : t \quad (\text{Assumption 6.3})
 \end{aligned}$$

Proof of Lemma C.2. Consider the $|S| \times |S|$ matrix $U = \gamma^n (P - I)$ with eigenvalues $f_1; \dots; f_{|S|}$. As shown in (Tsitsiklis & Van Roy, 1997), since P is a full rank matrix, the real part of f_i is strictly negative for all $i = 1; \dots; |S|$. $\square$

Furthermore, define $\gamma_{\max} = \max_i \gamma_i$ and $\gamma = \max_i \text{Re}[f_i] > 0$, where $\text{Re}[\cdot]$ evaluates the real part. Consider the matrix $U^0 = I + \frac{1}{2\gamma_{\max}} U$. It is easy to show that the eigenvalues of U^0 are $f_1 + \frac{1}{2\gamma_{\max}}$; $\dots$; $f_{|S|} + \frac{1}{2\gamma_{\max}}$. For an arbitrary i , the norm of the i 'th eigenvalue satisfies

$$\begin{aligned}
 1 + \frac{\gamma}{2\gamma_{\max}} &= 1 + \frac{\text{Re}[f_i]}{2\gamma_{\max}} + \frac{\text{Im}[f_i]^2}{2\gamma_{\max}^2} \\
 &= 1 + \frac{\text{Re}[f_i]}{2\gamma_{\max}} + \frac{\text{Im}[f_i]^2}{2\gamma_{\max}^2} \\
 &= 1 + \frac{\gamma}{2\gamma_{\max}^2} + \frac{\gamma_{\max}}{2\gamma_{\max}^2} \\
 &= 1 + \frac{\gamma}{4\gamma_{\max}} :
 \end{aligned} \quad (\text{Re}[f_i] < 0)$$

Hence, all the eigenvalues of U^0 are in the unit circle. By (Bertsekas et al., 1995), Page 46 footnote, we can find a weighted 2-norm as $\| \cdot \|_{\gamma}$ such that U^0 is contraction with respect to this norm with some contraction factor c . In particular, there exist a choice of γ such that we have $c = 1 - \frac{\gamma}{8\gamma_{\max}}$.

Proof of Lemma C.3. The existence of A_1 and A_2 immediately follows after observing that $G^i(i; y^i)_t$ is a linear function of y^i . Furthermore, the result on B follows due to v being bounded as shown in (Chen et al., 2021b). $\square$

Proof of Lemma C.4. For the sake of brevity, we write $S_t^i = s_t^i$ simply as s_t^i , and similarly for other random variables. We have

$$\mathbb{E}_t \left[r(y_t^i) - g(y_t^i) \right]$$

$$\begin{aligned}
 &= \sum_{\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r} P(\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r) f(y_t) g(y_t) \\
 &= \sum_{\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r} P(\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r) f(y_t) g(y_t) \\
 &= \sum_{\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r} P(\mathbf{s}_t^i, \mathbf{a}_t^i, \dots, \mathbf{s}_{t+n-1}^i, \mathbf{a}_{t+n-1}^i, \mathbf{s}_{t+n}^i, \mathbf{s}_t^j, \mathbf{a}_t^j, \dots, \mathbf{s}_{t+n-1}^j, \mathbf{a}_{t+n-1}^j, \mathbf{s}_{t+n}^j, \mathbf{F}_t, r) f(y_t) g(y_t) \\
 &= E_{\mathbf{r}}[f(y_t)] E_{\mathbf{r}}[g(y_t)]:
 \end{aligned}$$

□

Proof of Theorem 4.1. By Proposition C.1 and Lemmas C.1, C.2, C.3, and C.4, it is clear that the federated TD-learning with linear function approximation Algorithm 1 satisfies all the Assumptions 6.1, 6.2, 6.3, and 6.4 on the FedSAM Algorithm 4. Furthermore, by the proof of Theorem B.1, we have $w_t = (1 - \frac{\gamma}{2})^t$, and the constant c_{TDL} in the sampling distribution q_{TDL} in Algorithm 1 is $c_{\text{TDL}} = (1 - \frac{\gamma}{2})^{-1}$. Furthermore, by choosing the step size small enough, we can satisfy the requirements in (21), (23), (32), (35). By choosing K large enough, we can satisfy $K > \frac{1}{\gamma}$. Hence, the result of Theorem B.1 holds for this algorithm with some $c_{\text{TDL}} > 1$. Also, it is easy to see that $(1 - \frac{\gamma}{2})^t = O(1)$, which is a constant that can be absorbed in C_1^{TDL} . Finally, for the sample complexity result, we simply employ Corollary B.1.1.

Next, we derive the constant c_{TDL} . Since $k_{\text{c}} = k_{\text{c}}$, which is smooth, we choose $g() = \frac{1}{k_{\text{c}}} k_{\text{c}}^2$. By taking $\gamma = 1$, we have

$$l_{\text{CS}} = u_{\text{CS}} = 1. \text{ Therefore, we have } \gamma_1 = 1, \text{ and } \gamma_2 = 1 - c, \text{ and } c_{\text{TDL}} = \frac{1 - (1 - c)}{2}, \text{ where } c \text{ is defined in}$$

Lemma C.2.

□

C.2. Off-policy Tabular Setting

In this subsection, we verify that the Off-policy federated TD-learning Algorithm 2 satisfies the properties of the FedSAM Algorithm 4. In the following, V^* is the solution to the Bellman equation (124).

$$V(s) = \sum_a (a_j s) R(s; a) + \sum_{s^0} P(s^0 | s; a) V(s^0) \quad (124)$$

Note that V^* is independent of the sampling policy of the agent. Furthermore, we take $k_{\text{c}} = k_{\text{c}}$.

Proposition C.2. Off-policy n -step federated TD-learning is equivalent to the FedSAM Algorithm 4 with the following parameters.

1. $\bar{V}^i = V^i - V$
2. $S_t = (S_t^1, \dots, S_t^N)$ and $A_t = (A_t^1, \dots, A_t^N)$
3. $y_t^i = (S_t^i, A_t^i, \dots, S_{t+n-1}^i, A_{t+n-1}^i, S_{t+n}^i)$ and $Y_t = (S_t, A_t, \dots, S_{t+n-1}, A_{t+n-1}, S_{t+n})$
4. π^i : Stationary distribution of the sampling policy of the i -th agent.
5. $G_t^i(y_t^i) = V_t^i(s_t^i) + \frac{1}{n} \sum_{l=t}^{t+n-1} \pi^i(s_{t+l}^i) (V_t^i(s_{t+l}^i) - V_{t+l}^i(s_{t+l}^i))$
6. $b_t^i(y_t^i) = \frac{1}{n} \sum_{l=t}^{t+n-1} \pi^i(s_{t+l}^i) (V_t^i(s_{t+l}^i) - V_{t+l}^i(s_{t+l}^i))$

Lemma C.5. Consider the federated off-policy TD-learning Algorithm 2 as a special case of FedSAM (as specified in Proposition C.2). Suppose the trajectory $f_{S_t^i, t=0,1,\dots}$ converges geometrically fast to its stationary distribution as follows

$$= \frac{1}{N} \sum_{j=1}^N \left(\frac{V_t^j}{\{z^V\}} \right)$$

which is equivalent to the synchronization step in FedSAM Algorithm 4. Notice that here we used the fact that all agents have the same fixed point V . $\square$

Proof of Lemma C.5. By taking expectation of $G^i(\cdot; y_t)_s$, we have

$$\begin{aligned} G^i(\cdot)_s &= E_{S^{i:t}}(s) + \mathbf{1}_{f_s=S^i g} \sum_{l=t}^{t+n-1} J^{(i)}(S^i; A^i)_{j_l} (S^{i-1}_{l+1}) (S^i)^{l=t} \\ &= (s) + \sum_{l=t}^{t+n-1} \underbrace{E_{S^{i:t}} \mathbf{1}_{f_s=S^i g} J^{(i)}(S^i; A^i)_{j_l} (S^{i-1}_{l+1}) (S^i)^{l=t}}_{\frac{1}{T_l}} : \end{aligned}$$

Denote $E_k^i[\cdot] = E[J f S^i; A^i g_{rk-1}; S^i]$. For T_l , we have

$$\begin{aligned} T_l &= E_{S^{i:t}} E_{\mathbf{1}_{f_s=S^i g}} \sum_{j=t}^{t+n-1} J^{(i)}(S^i; A^i)_{j_l} (S^{i-1}_{l+1}) (S^i)^{l=t} \\ &= E_{S^{i:t}} \mathbf{1}_{f_s=S^i g} \sum_{j=t}^{t+n-1} J^{(i)}(S^i; A^i)_{j_l} E_{\mathbf{1}_{f_s=S^i g}} J^{(i)}(S^i; A^i)_{j_l} (S^{i-1}_{l+1}) (S^i)^{l=t} : \end{aligned}$$

Here,

$$\begin{aligned} E_{\mathbf{1}_{f_s=S^i g}} J^{(i)}(S^i; A^i)_{j_l} (S^{i-1}_{l+1}) &= \sum_{s;a} P(S_{l+1}^i = s; A^i = ajS^i) J^{(i)}(S^i; A^i)_{j_l} (s)_{s;a} \\ &= \sum_{s;a} P(S_{l+1}^i = s; A^i = ajS^i) \frac{(ajS^i)_{j_l}}{(ajS^i)_{j_l}} (s) \\ &= \sum_{s;a} (ajS^i)_{j_l} P(sjS^i; a) \frac{(ajS^i)_{j_l}}{(ajS^i)_{j_l}} \\ &= P(sjS^i; a) (ajS^i)_{j_l} [P](S^i)_{j_l}; s;a \end{aligned}$$

where $[P]_{s_0; s_1} = \sum_a P(s_1 | s_0; a) (aj s_0)$. Hence, we have

$$\begin{aligned} T_l &= E_{S^{i:t}} \mathbf{1}_{f_s=S^i g} \sum_{j=t}^{t+n-1} J^{(i)}(S^i; A^i)_{j_l} E_{\mathbf{1}_{f_s=S^i g}} J^{(i)}(S^i; A^i)_{j_l} (P)(S^i)_{j_l} (S^i)^{l=t} \\ &= E_{S^{i:t}} \mathbf{1}_{f_s=S^i g} \sum_{j=t}^{t+n-1} J^{(i)}(S^i; A^i)_{j_l} ((P)^2)(S^i)_{j_l} (P_l)(S_{l-1}^i) \\ &= \dots \\ &= E_{S^{i:t}} \mathbf{1}_{f_s=S^i g} \sum_{j=t}^{t+n-1} ((P)^{l-t+1})(S^i)_{j_l} ((P)^{l-t})(S^i)_{j_l} = (s)^{l-t} \\ &= ((P)^{l-t+1})(s) ((P)^{l-t})(s) \\ &= (P)^{l-t+1} (P)^{l-t} (s) = (P)^{l-t+1} (P)^{l-t} (s); \end{aligned}$$

where we denote i as diagonal matrix with diagonal entries corresponding to the stationary distribution i . Hence, in total we have

$$\begin{aligned} G^i(\cdot)_s &= (s) + \sum_{l=t}^{t+n-1} (P)^{l-t+1} (P)^{l-t} (s) \\ &= (s) + \sum_{l=t}^{t+n-1} (P)^{l-t+1} (P)^{l-t} (s); \end{aligned}$$

Furthermore, by the same argument as in the proof of Lemma C.1, we have

$$kG^i(\cdot) - E[G^i(\cdot; y_t)] k_C \leq \sum_i P(y_t = y_{tj} y_0) A_2 k k_C; y_t$$

[illegible]

$$\begin{aligned}
 & \max_s j_1(s) - j_2(s) + \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [k_1 - j_2 k_1 + k_1 - j_2 k_1] \quad (\text{definition of } k - k_1) \\
 & \max_s j_1(s) - j_2(s) + \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [k_1 - j_2 k_1 + k_1 - j_2 k_1] \quad (\text{definition of } l_{\max}) \\
 & = k_1 - j_2 k_1 + [k_1 - j_2 k_1 + k_1 - j_2 k_1] l_{\max} \\
 & = k_1 - j_2 k_1 + [k_1 - j_2 k_1 + k_1 - j_2 k_1] l_{\max} \quad \text{if } l_{\max} = 1 \\
 & \quad \frac{1 - (l_{\max})^{n+1}}{1 - l_{\max}} \quad \text{o.w.}
 \end{aligned}$$

Furthermore, we have

$$\begin{aligned}
 kb^l(y_t^i)k_c &= \max_{S_t^i; A_t^i; \dots; S_t^i} \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [R(S_t^i; A_t^i) + V(S_{t+1}^i) - V(S_t^i)] \quad (\text{triangle inequality}) \\
 &= \max_{S_t^i; A_t^i; \dots; S_{t+n}^i} \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [R(S_t^i; A_t^i) + V(S_{t+1}^i) - V(S_t^i)] \quad (\text{triangle inequality}) \\
 &= \max_{S_t^i; A_t^i; \dots; S_{t+n}^i} \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [R(S_t^i; A_t^i) + V(S_{t+1}^i) - V(S_t^i)] \quad (\text{triangle inequality}) \\
 &= \max_{S_t^i; A_t^i; \dots; S_{t+n}^i} \sum_{l=t}^{t+X-1} l^{(i)}(S^i; A^i) [R(S_t^i; A_t^i) + V(S_{t+1}^i) - V(S_t^i)] \quad (\text{triangle inequality}) \\
 &= \frac{2l_{\max}}{1} \sum_{l=t}^{t+X-1} (l_{\max})^{l-t} \\
 &= \frac{2l_{\max}}{1 - (l_{\max})^X} \sum_{l=t}^{t+X-1} (l_{\max})^{l-t} \quad \text{if } l_{\max} = 1 \\
 & \quad \frac{1 - (l_{\max})^{n+1}}{1 - l_{\max}} \quad \text{o.w.}
 \end{aligned}$$

□

Proof of Lemma C.8. The proof follows similar to Lemma C.4.

□

Proof of Theorem 5.1. By Proposition C.2 and Lemmas C.5, C.6, C.7, and C.8, it is clear that the federated off-policy TD-learning Algorithm 2 satisfies all the Assumptions 6.1, 6.2, 6.3, and 6.4 of the FedSAM Algorithm 4. Furthermore, by the proof of Theorem B.1, we have $w_t = (1 - \frac{c}{2})^t$, and the constant c in the sampling distribution q^c in Algorithm 1 is $c = (1 - \frac{c}{2})^{-1}$. In equation (125) we evaluate the exact value of w_t .

Furthermore, by choosing step size small enough, we can satisfy the requirements in (21), (23), (32), (35). By choosing K large enough, we can satisfy $K > \frac{1}{\epsilon}$, and by choosing T large enough we can satisfy $T > K + \frac{1}{\epsilon}$. Hence, the result of Theorem B.1 holds for this algorithm.

Next, we derive the constants involved in Theorem B.1 step by step. After deriving the constants C_1, C_2, C_3 , and C_4 in Theorem B.1, we can directly get the constants $C_i^T D^T$ for $i = 1; 2; 3; 4$.

In this analysis we only consider the terms involving $jS_j, jA_j, \frac{1}{\sqrt{\tau_{\max}}}$, and $\min$. Since $k - k_c = k - k_1$, we choose $g() = \frac{1}{k} k^2$, i.e. the p -norm with $p = 2 \log(jS_j)$. It is known that $g()$ is $(p-1)$ smooth with respect to k k_p norm (Beck, 2017), and hence $L = (\log(jS_j))$. Hence, we have $l_{cs} = jS_j^{1-p} = \frac{1}{p} = (1)$ and $u_{cs} = \frac{1}{p} = 1$. Therefore, we

have $\gamma_1 = \frac{1 + \frac{u_{cs}}{c}}{1 + \frac{u_{cs}}{c}} = \frac{1 + \frac{1}{c}}{1 + \frac{1}{c}}$. By choosing $c = (\frac{1+c}{2})^2 - 1 = \frac{1+2\frac{3}{c}}{4\frac{3}{c}}(1 - c) = \min(1 - n+1) =$

$$\begin{aligned}
 (\min(1, \frac{1}{2})), \text{ which is } &= O(1), \text{ we have } \gamma_1 = \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and} \\
 \gamma_2 = 1 - \frac{1}{1 + \frac{1}{2}} &= \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and} \\
 &> 1 - \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and} \\
 (\min(1, \frac{1}{2})) &= O(1), \text{ we have } \gamma_3 = \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and} \\
 \gamma_3 = \frac{1}{1 + \frac{1}{2}} &= \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and}
 \end{aligned}$$

Using γ_2 , we have

$$w_t = 1 - \frac{\gamma_2}{2} = \frac{1}{2} = O(1) \quad (125)$$

Further, we have

$$\begin{aligned}
 l_{cm} &= (1 + l_{cs}^2)^{1/2} = (1) \\
 u_{cm} &= (1 + u_{cs}^2)^{1/2} = (1)
 \end{aligned}$$

Since TV-divergence is upper bounded with 1, we have $m = O(1)$. By Lemma C.7, we have

$$A_1 = A_2 = 1 + (1 + \frac{1}{l_{max}})^n \quad \text{if } l_{max} = 1 \quad \text{o.w.} \quad = O(l_{max}^n)$$

and $A_1 = A_2 =$

$$(1), \quad \text{if } l_{max} = 1 \quad \text{o.w.} \quad = O(1)$$

and $B =$

(1). Hence $m_1 = \frac{2A}{n} = O(l_{max})$. Also, we have $m_2 = 0$.

We choose the D-norm in Lemma B.9 as the 2-norm k_2 . Hence, by primary norm equivalence, we have $l_{cd} = \frac{1}{\sqrt{2}}$, and $u_{cd} = 1$, and hence $\frac{u_{cd}}{l_{cd}} = \sqrt{2}$.

We can evaluate the rest of the constants as follows

$$\begin{aligned}
 \gamma_1 &= \gamma_2 = \gamma_3 = \frac{1}{1 + \frac{1}{2}} = \frac{2}{3} = O(1), \text{ and} \\
 \gamma_2 &= 1 - \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and} \\
 \gamma_3 &= \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and}
 \end{aligned}$$

and similarly

$$\gamma_3 = \frac{1}{1 + \frac{1}{2}} = \frac{1}{3/2} = \frac{2}{3} = O(1), \text{ and}$$

$$C_1 = 3^q \frac{1}{A_2^2 + 1} = O \left(\frac{1}{I_{\max}^n} \right);$$

and

$$C_1 = \frac{1}{(1)};$$

$$C_2 = 3C_1 + 8 = O(\frac{1}{\eta_{\max}});$$

$$C_3 = \frac{m_2 L}{l_{cs}^2} = O(1);$$

$$\begin{aligned} C_4 &= \frac{L^2}{2} \frac{1}{\log^2(jSj)} + \frac{2LA_2}{2} \frac{1}{\log(jSj)} + \frac{L(A_1+1)}{2} \frac{1}{\log(jSj)} + \frac{3m_1 L^2}{2} \frac{1}{\log(jSj)} + \frac{cs}{2} \frac{1}{\log(jSj)} \\ &= O\left(\frac{1}{m_2 n} (1 - \frac{1}{2})^2 + \frac{1}{\min(1 - \frac{1}{2})} \frac{1}{m_1 n} (1 - \frac{1}{2})^2 + \frac{1}{\min(1 - \frac{1}{2})} \frac{1}{m_1 n} (1 - \frac{1}{2})^2 + \frac{1}{\min(1 - \frac{1}{2})} \log(jSj)\right) \\ &= O\left(\frac{1}{2m_1 n} (1 - \frac{1}{2})^2\right); \end{aligned}$$

$$\begin{aligned} C_5 &= \frac{3u_{2m}}{2} \frac{1}{\log^3(jSj)} + \frac{3L(A_1+1)}{2} \frac{1}{\log^2(jSj)} + \frac{LA_2}{2} \frac{1}{\log(jSj)} + \frac{L^2 u_{2m}}{2} \frac{1}{\log^2(jSj)} + \frac{3A^2 L^2}{2} \frac{1}{\log(jSj)} \\ &= O\left(\frac{1}{m_2 n} (1 - \frac{1}{2})^2 \frac{1}{\min(1 - \frac{1}{2})} + \frac{1}{\min(1 - \frac{1}{2})} \frac{1}{m_1 n} (1 - \frac{1}{2})^2 + \frac{1}{\min(1 - \frac{1}{2})} \log(jSj)\right) \\ &= O\left(\frac{1}{3m_1 n} (1 - \frac{1}{2})^2\right); \end{aligned}$$

$$\begin{aligned} C_6 &= \frac{3u_{2m}}{2} \frac{1}{\log^3(jSj)} + \frac{3L(A_1+1)}{2} \frac{1}{\log^2(jSj)} + \frac{m_1 L}{2} \frac{1}{\log(jSj)} = O\left(\frac{1}{m_2 n} (1 - \frac{1}{2})^2 \frac{1}{\min(1 - \frac{1}{2})} + \frac{1}{\min(1 - \frac{1}{2})} \log(jSj)\right) \\ &= O\left(\frac{1}{3m_1 n} (1 - \frac{1}{2})^2\right); \end{aligned}$$

$$C_7 = \frac{m_2^2 u_{cm}^2 L^2}{2l_{cs}^4} = O(1);$$

$$C_8 = \frac{1}{2} \frac{3L}{\log(jSj)} = O\left(\frac{1}{\log(jSj)}\right);$$

$$C_9 = \frac{8u_{cd} B}{l_{cd}} = O\left(\frac{1}{l_{cd}}\right);$$

$$C_{10} = \frac{8m_2 u_{cd}}{l_{cd} (1 - \frac{1}{2})} = O(1);$$

$$C_{11} = \frac{8C_1^2 C_3^2 u_{cm}^2}{6} = O(1);$$

$$C_{12} = \frac{8C_4 u_{2D} B^2}{l_{cd}^2} = O\left(\frac{1}{l_{cd}^2}\right) = O\left(\frac{1}{(1 - \frac{1}{2})^2}\right);$$

$$C_{13} = \frac{14C_4 u_{cd}^2 m_2^2}{l_{cd}^2 (1 - \frac{1}{2})} = O(1);$$

$$C_{14}() = C_7 + C_{11} + 0.5C_2^2C_9^2 + C_3C_{10} + 2C_1C_3C_{10} + 3A_1C_3 + C_{13} + C_8 \frac{u_D^2}{l^2} 2m_2^{22} = 0; c_D$$

$$C_{16}() = C_8 \frac{u_{\frac{B^2}{c_D}} + 1}{l_{\frac{c_D}{2}}^2} C_{12}^2 = O \frac{jSj \log(jSj) l_{\max}^{2n} + 1 + \frac{jSj \log(jSj) l_{\max}^{3n}}{2} \frac{1}{2}}{(1)_{\min}^3} \frac{1}{\min(1)_{\max}^4}$$

$$= O \frac{jSj \log_2(jSj) l_{\max}^{2n}}{(1)_{\min}^4} ;$$

$$C_{17} = (3A_1C_3 + 8A_1^2C_4 + C_5 + C_6)$$

$$= O \quad 0 + l_{\max} : \frac{2n^2}{(jSj)_{\max}^2} \frac{\log^1(jSj) l_{\max}^2}{m_{\min}^3} \frac{\log^2(jSj) l_{\max}^2}{m_{\min}^3} \frac{\log^1}{m_{\min}^3}$$

$$= O \quad \frac{l_{\max}^3 \log^2(jSj)}{m_{\min}^3} ;$$

Similar to $k_D = k_{k_2}$, we have $l_{c2} = \frac{1}{S}$ and $u_{c2} = 1$.

$$A_1 = \frac{2A_1^2 u_{c2}^2}{l_{c2}^2} = O \quad l_{\max}^{2n} jSj;$$

$$A_2 = \frac{2A_2 u_{c2}}{l_{c2}^2} = O \quad l_{\max}^{2n} jSj;$$

$$\sim \frac{1}{l^2} \frac{\max}{B} \frac{jSj l_{\max}^{2n}}{(1)_{\min}^2};$$

$$M_0 = \frac{1}{l_{cm}^2} \frac{1}{C_1^2} B + (A_2 + 1) k_0 k_c + \frac{2}{2C_1} \frac{B}{k_0 k^2} = O \quad \frac{l_{\max}^n + (l_{\max}^n)^2}{1} \frac{1}{1} + 1$$

$$= O \quad \frac{1}{(1)_{\min}^2}$$

$$C_1 = 16u_{cm}^2 M_0 \log \left(\frac{1}{e} + \frac{1}{2} \right) = O \quad \frac{l_{\max}^{4n}}{(1)_{\min}^2} \frac{1}{\min(1)} = O \quad \frac{l_{\max}^{4n}}{(1)_{\min}^3}$$

$$C_2 = \frac{8u_{cm} C_8 + \frac{1}{2} + C_{12}}{\min(1)_{\min}^2} \frac{1}{\min(1)} \frac{\log(jSj)}{m_{\min}^4} + 1 + \frac{jSj \log_2(jSj) l_{\max}^{3n}}{2} \frac{1}{\max}$$

$$= O \quad \frac{jSj \log_2(jSj) l_{\max}^{3n}}{m_{\min}^4} \frac{1}{\max}$$

$$C_3 = \frac{80B^2 C_1 u_{cm}^{1+B(14m_2)}}{\min(1)_{\min}^4} \frac{1}{(1)_{\min}^4} \frac{jSj^2 l_{\max}^{4n}}{m_{\min}^3} \frac{l_{\max}^{3n}}{2} \frac{\log(jSj)}{2}$$

$$= O \quad \frac{l_{\max}^{7n}}{m_{\min}^3} \frac{jSj^2 \log_2(jSj)}{m_{\min}^8} ;$$

$$C_4 = 8u_{cm} C_7 + C_{11} + 0.5C_2^2C_9^2 + C_3C_{10} + 2C_1C_3C_{10} + 3A_1C_3 + C_{13} = 0;$$

Finally, for the sample complexity result, we simply employ Corollary [B.1.1](#).

□

D. Federated Q-learning

In this section, we verify that the federated Q-learning algorithm 3 satisfies the properties of the FedSAM Algorithm 4. In the following, Q is the solution to the Bellman optimality equation (126)

$$Q(s; a) = R(s; a) + E_{S^0} P(j|s; a) \max_{a^0} Q(S^0; a^0) \quad (126)$$

Note that Q is independent of the sampling policy of the agent. Furthermore, $k_c = k_1$.

Proposition D.1. Federated Q-learning algorithm 3 is equivalent to the FedSAM Algorithm 4 with the following parameters.

1. $\bar{Q}_t^i = Q$
2. $S_t = (S_t^1; \dots; S_t^N)$ and $A_t = (A_t^1; \dots; A_t^N)$
3. $y_t^i = (S_t^i; A_t^i; S_{t+1}^i; A_{t+1}^i)$ and $Y_t = (S_t; A_t; S_{t+1}; A_{t+1})$
4. $\bar{\mu}^i$: Stationary distribution of the sampling policy of the i 'th agent.
5. $G^i(y_t^i)_{(s;a)} = \mathbb{1}_{f(S_t^i=s; A_t^i=a)} \max_{a^0} \{R(S_{t+1}^i; a^0) + Q(S_{t+1}^i; a^0) - Q(S_t^i; A_t^i)\} + \mathbb{1}_{f(S_t^i=s; A_t^i=a)} R(S_t^i; A_t^i) + \max_{a^0} Q(S_{t+1}^i; a^0) - Q(S_t^i; A_t^i)$

where $\mathbb{1}_A$ is the indicator function corresponding to set A , such that $\mathbb{1}_A = 1$ if A is true, and 0 otherwise.

Lemma D.1. Consider the federated Q-learning Algorithm 3 as a special case of FedSAM (as specified in Proposition D.1). Suppose the trajectory $fS_t^i; A_t^i; g_{t=0;1;\dots}$ converges geometrically fast to its stationary distribution as follows $d_{TV}(P(S_t^i = j; A_t^i = j) | \bar{\mu}^i) \leq \rho^t$ for all $i = 1, 2, \dots, N$. The corresponding $G^i()$ in Assumption 6.1 for the federated Q-learning is as follows

$$G^i(y_t^i)_{(s;a)} = R(s; a) + E_{S^0} P(j|s; a) \max_{a^0} \{R(S_{t+1}^i; a^0) + Q(S_{t+1}^i; a^0) - Q(S_t^i; A_t^i)\} + \mathbb{1}_{f(S_t^i=s; A_t^i=a)} R(S_t^i; A_t^i) + \max_{a^0} Q(S_{t+1}^i; a^0) - Q(S_t^i; A_t^i)$$

Furthermore, we have $m_1 = 2A_2\rho$ where A_2 is specified in Lemma D.3, $m_2 = 0$, and $\rho = \rho$.

Lemma D.2. Consider the federated Q-learning as a special case of FedSAM (as specified in Proposition D.1). The corresponding contraction factor c in Assumption 6.2 for this algorithm is $c = (1 - \rho)_{\min}$, where $\rho_{\min} = \min_{s,a,i} \rho^i(s; a)$

Lemma D.3. Consider the federated Q-learning as a special case of FedSAM (as specified in Proposition D.1). The constants A_1 , A_2 , and B in Assumption 6.2 are as follows: $A_1 = A_2 = 2$ and $B = \frac{1}{1-\rho}$.

Lemma D.4. Consider the federated Q-learning as a special case of FedSAM (as specified in Proposition D.1). Assumption 6.4 holds for this algorithm.

D.1. Proofs

Proof of Proposition D.1. Items 1-4 are by definition. Furthermore, by the update of the Q-learning, and subtracting Q from both sides, we have

$$\begin{aligned} & \left| \frac{Q_{t+1}^i(s; a) - Q(s; a)}{\mathbb{1}_{f(S_t^i=s; A_t^i=a)}} \right| = \left| \frac{Q_t^i(s; a) - Q(s; a)}{\mathbb{1}_{f(S_t^i=s; A_t^i=a)}} \right| \\ & + \left| \frac{\mathbb{1}_{f(S_t^i=s; A_t^i=a)} R(S_{t+1}^i; a) + \max_{a^0} Q_t(S_{t+1}^i; a^0) - Q_t(S_t^i; A_t^i) - Q(s; a)}{\mathbb{1}_{f(S_t^i=s; A_t^i=a)}} \right| \\ & = \left| \frac{\mathbb{1}_{f(S_t^i=s; A_t^i=a)} R(S_{t+1}^i; a) + \max_{a^0} Q_t(S_{t+1}^i; a^0) - Q_t(S_t^i; A_t^i) - Q(s; a)}{\mathbb{1}_{f(S_t^i=s; A_t^i=a)}} \right| \\ & \leq \left| \frac{\mathbb{1}_{f(S_t^i=s; A_t^i=a)} R(S_{t+1}^i; a) + \max_{a^0} Q_t(S_{t+1}^i; a^0) - Q_t(S_t^i; A_t^i) - Q(s; a)}{\mathbb{1}_{f(S_t^i=s; A_t^i=a)}} \right| \end{aligned}$$

$$Q_t^i \frac{1}{N} \sum_{j=1}^N Q_t^j = \frac{1}{N} \sum_{j=1}^N (Q_t^j \frac{1}{N} \sum_{i=1}^N Q_t^i)$$
$$\begin{aligned}
& \mathbb{E}[G^i(\cdot; \mathbf{y}^i)] \mathbf{k}_c \\
&= \mathbb{E}_{\mathbf{y}^i} [G^i(\cdot; \mathbf{y}^i)]^T \mathbb{E}[G^i(\cdot; \mathbf{y}^i)]^C = \mathbb{X} \\
&= \mathbb{E}_{\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0} [G^i(\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0) P(S^i = \mathbf{s}; A^i = \mathbf{a}; S_{t+1} = \mathbf{s}^0; A_{t+1} = \mathbf{a}^0 | S^i; A^i) G^i(\cdot; \mathbf{y}^i)] \\
&= \mathbb{E}_{\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0} [G^i(\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0) P(S^i = \mathbf{s}; A^i = \mathbf{a}; S_{t+1} = \mathbf{s}^0; A_{t+1} = \mathbf{a}^0 | S^i; A^i): G^i(\cdot; \mathbf{y}^i)] \mathbf{k}_c \quad (\text{kaxk}_c = \text{jajkxk}_c) \\
&= \mathbb{E}_{\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0} [G^i(\mathbf{s}; \mathbf{a}; \mathbf{s}^0; \mathbf{a}^0) P(S^i = \mathbf{s}; A^i = \mathbf{a}; S_{t+1} = \mathbf{s}^0; A_{t+1} = \mathbf{a}^0 | S^i; A_0): A_2 \mathbf{k}_c] \quad (\text{Assumption D.3}) \\
&= \mathbb{X} [G^i(\mathbf{s}; \mathbf{a}) P(\mathbf{s}^0 | \mathbf{s}; \mathbf{a}) (A^0 | \mathbf{s}^0)] P(S^i = \mathbf{s}; A^i = \mathbf{a} | S_0; A_0): A_2 \mathbf{k}_c \\
&= \mathbb{X} [G^i(\mathbf{s}; \mathbf{a}) P(S_t = \mathbf{s}; A_t = \mathbf{a} | S_0; A_0): A_2 \mathbf{k}_c] \\
&= 2d_{TV}(G^i(\cdot; \cdot); P(S^i = \cdot; A^i = \cdot | S_0; A_0)): A_2 \mathbf{k}_c \\
&2A_2 \mathbf{k}_c \mathbf{m}^T
\end{aligned}$$
$$\begin{aligned} kE[b^i(y_t^i)]k_c &= \max_{s,a} P(S_t^i = s; A_t^i = a | S_0^i, A_0^i) R(s; a) + E_{S^0 P(j_s; a)} [\max_{a^0} Q(s^0; a^0)] - Q(s; a) \\ &= 0: \quad \text{(Bellman optimality equation (126))} \end{aligned}$$

$$kG^i(1; y) - G^i(2; y)k_c$$

$$\begin{aligned}
 &= \max_{s,a} \left(\mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \mathbb{1}_{(S;A)} \max_{a^0} Q(S^0; a^0) \right. \\
 &\quad \left. + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (2 + Q(S^0; a^0)) \mathbb{1}_{(S;A)} \max_{a^0} Q(S^0; a^0) \right) \\
 &\quad + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \max_{a^0} (2 + Q(S^0; a^0)) \quad (\text{triangle inequality}) \\
 &\max_{s,a} k_1 \mathbb{1}_{(S;A)} + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \max_{a^0} (2 + Q(S^0; a^0)) \quad (\text{definition of } k_1) \\
 &\max_{s,a} k_1 \mathbb{1}_{(S;A)} + \mathbb{1}_{f_{S=s;A=a}} k_1 \mathbb{1}_{(S;A)} \quad (\text{By (127)}) \\
 &2k_1 \mathbb{1}_{(S;A)} \\
 &= 2k_1 \mathbb{1}_{(S;A)}
 \end{aligned}$$

Second, for A_2 , we have

$$\begin{aligned}
 kG^i(\gamma)k_c &= \max_{s,a} \left(\mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \mathbb{1}_{(S;A)} \max_{a^0} Q(S^0; a^0) \right. \\
 &\quad \left. + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (2 + Q(S^0; a^0)) \mathbb{1}_{(S;A)} \max_{a^0} Q(S^0; a^0) \right) \\
 &\quad + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \max_{a^0} (2 + Q(S^0; a^0)) \quad (\text{triangle inequality}) \\
 &\max_{s,a} k_1 \mathbb{1}_{(S;A)} + \mathbb{1}_{f_{S=s;A=a}} \max_{a^0} (1 + Q(S^0; a^0)) \max_{a^0} Q(S^0; a^0) \quad (\text{definition of } k_1) \\
 &\max_{s,a} k_1 \mathbb{1}_{(S;A)} + \max_{a^0} jQ(S^0; a^0)j \\
 &2k_1 \mathbb{1}_{(S;A)} \\
 &= 2k_1 \mathbb{1}_{(S;A)}
 \end{aligned}$$

Lastly, for B , we have

$$\begin{aligned}
 kb^i(y^i)k_c &= \max_{s,a} \left(\mathbb{1}_{f_{S=s;A=a}} \max_{a^0} R(S; A) + \max_{a^0} Q(S^0; a^0) \mathbb{1}_{(S;A)} \right. \\
 &\quad \left. + \max_{a^0} jQ(S^0; a^0)j + jQ(S; A)j \right) \quad (\text{triangle inequality}) \\
 &\max_{s,a} \left(1 + \frac{1}{1} + \frac{1}{1} \right) \\
 &= \frac{2}{1}
 \end{aligned}$$

Proof of Lemma D.4. The proof follows similar to Lemma C.4. $\square$

Proof of Theorem 5.2. By Proposition D.1 and Lemmas D.1, D.2, D.3, and D.4, it is clear that the federated Q Algorithm 3 satisfies all the Assumptions 6.1, 6.2, 6.3, and 6.4 of the FedSAM Algorithm 4. Furthermore, by the proof of Theorem B.1, we have $w_t = (1 - \frac{c}{2})^t$, and the constant c in the sampling distribution q^c_t in Algorithm 1 is $c = (1 - \frac{c}{2})^{-1}$. In equation (128) we evaluate the exact value of w_t .

Furthermore, by choosing step size small enough, we can satisfy the requirements in (21), (23), (32), (35). By choosing K large enough, we can satisfy $K > \frac{1}{\epsilon}$, and by choosing T large enough we can satisfy $T > K + \frac{1}{\epsilon}$. Hence, the result of Theorem B.1 holds for this algorithm.

Next, we derive the constants involved in Theorem B.1 step by step. In this analysis we only consider the terms involving jSj , jAj , $\frac{1}{1-\gamma}$, $l_{\max}$, and $\min$. Since $k_1 k_c = k_1 k_1$, we choose $g(\cdot) = \frac{1}{2} k_1 k_p^1$, i.e. the p -norm with $p = 2 \log(jSj)$.

ve

y

e

nd

ed

☐