

Decentralized Gossip-Based Stochastic Bilevel Optimization over Communication Networks

Shuoguang Yang
IEDA, HKUST
yangsg@ust.hk

Xuezhou Zhang
Princeton University
xz7392@princeton.edu

Mengdi Wang
Princeton University
mengdiw@princeton.edu

Abstract

Bilevel optimization have gained growing interests, with numerous applications being found in meta learning, minimax games, reinforcement learning, and nested composition optimization. This paper studies the problem of decentralized distributed stochastic bilevel optimization over a network where each agent can only communicate with its neighbors, and gives examples from multi-task, multi-agent learning and federated learning. In this paper, we propose a gossip-based decentralized bilevel learning algorithm that allows networked agents to solve both the inner and outer optimization problems in a single timescale and share information through network propagation. We show that our algorithm enjoys the $\mathcal{O}(\frac{1}{\epsilon})$ per-agent sample complexity for general nonconvex bilevel optimization and $\mathcal{O}(\frac{1}{\epsilon})$ for Polyak-Łojasiewicz objectives, achieving a speedup that scales linearly with the network size K . The sample complexities are optimal in both ϵ and K . We test our algorithm on the examples of hyperparameter tuning and decentralized reinforcement learning. Simulated experiments confirmed that our algorithm achieves the state-of-the-art training efficiency and test accuracy.

1 Introduction

In recent years, stochastic bilevel optimization (SBO) has attracted increasing attention from the machine learning community. It has been found to provide favorable solutions to a variety of problems, such as meta learning and hyperparameter optimization (Franceschi et al., 2018; Snell et al., 2017; Bertinetto et al., 2018), composition optimization (Wang and Liu, 2016), two-player games (Von Stackelberg and Von, 1952), reinforcement learning and imitation learning (Arora et al., 2020; Hong et al., 2020). While the majority of the above mentioned work focuses on algorithm designs in the classic centralized setting, such problems often arise in distributed/federated applications, where agents are unwilling to share data but rather perform local updates and communicate with neighbors. Theories and algorithms for distributed stochastic bilevel optimization are less developed.

Consider the decentralized learning setting where the data are distributed over K agents $K = \{1, 2, \dots, K\}$ over a communication network. Each agent can only communicate with its neighbors over the network. One example is federated learning which is often concerned with a single-server-multi-user system, where agents communicate with a central server to solve a task cooperatively (Lan and Zhou, 2018; Ge et al., 2018). Another example is the sensor network, where sensors are fully decentralized and can only communicate with nearby neighbors (Taj and Cavallaro, 2011).

We consider the following decentralized stochastic bilevel optimization (DSBO) problem,

$$\min_{x \in \mathbb{R}^{d_x}} F(x) = \frac{1}{K} \sum_{k=1}^K f^k(x, y^*(x)) \quad , \quad \text{s.t. } y^*(x) = \operatorname{argmin}_{y \in \mathbb{R}^{d_y}} \frac{1}{K} \sum_{k=1}^K g^k(x, y) \quad , \quad (1)$$

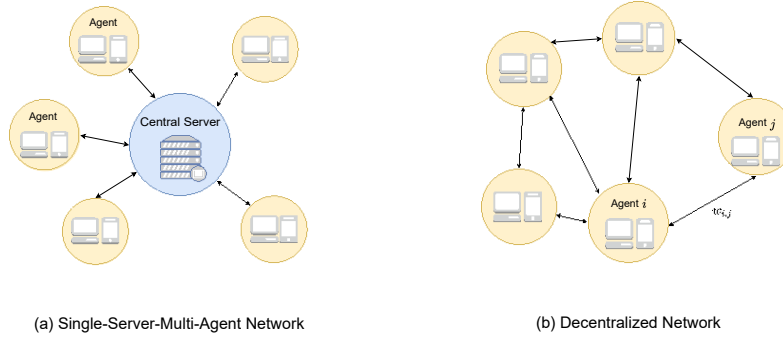


Figure 1: Illustration of Distributed Network Structures. In a centralized network, the agents may cooperate to solve a specific task by communicating with the central server, commonly seen in federated learning.

where $f^k(x, y) = \mathbb{E}_{\downarrow^k} [f^k(x, y; \downarrow^k)]$ and $g^k(x, y) = \mathbb{E}_{\leftarrow^k} [g^k(x, y; \leftarrow^k)]$ may vary between agents. The expectations $\mathbb{E}_{\downarrow^k} [\bullet]$ and $\mathbb{E}_{\leftarrow^k} [\bullet]$ are taken with respect to the random variables $\downarrow^k$ and $\leftarrow^k$, with heterogeneous distributions across agents. We consider the scenario where each $g^k(x, y)$ is strongly convex in y . We use the notation $F^* = \min_{x \in \mathbb{R}^{d_x}} F(x)$, $f_k(x, y) = \frac{1}{K} \sum_{k=1}^K f^k(x, y)$, and $g(x, y) = \frac{1}{K} \sum_{k=1}^K g^k(x, y)$ for convenience.

1.1 Example applications of SBO

SBO was first employed to formulate the resource allocation problem (Bracken and McGill, 1973) and has since found its applications in many classic operations research settings (Cramer et al., 1994; Sobieszczyński-Sobieski and Haftka, 1997; Livne, 1999; Tu et al., 2020; Tu, 2021), and more recently in machine learning problems (Franceschi et al., 2018; Snell et al., 2017; Bertinetto et al., 2018; Wang and Liu, 2016). In particular, we introduce two applications that have recently attracted lot of attention, namely hyperparameter optimization and compositional optimization.

Hyperparameter Optimization The problem of hyper-parameter tuning (Okuno et al., 2021) often takes the following form:

$$\min_{x \in \mathbb{R}^d} \left(\frac{1}{I} \sum_{i \in D_{\text{val}}} \ell_i(y^*(x)) \right), \text{ s.t. } y^*(x) \in \underset{y \in \mathbb{R}^{d_y}}{\text{argmin}} \left(\frac{1}{J} \sum_{j \in D_{\text{train}}} \ell_j(y) + R(x, y) \right), \quad (2)$$

where D_{train} and D_{val} are two datasets used for training and validation, respectively, $y \in \mathbb{R}^{d_y}$ is a vector of unknown parameters to optimize, $\ell_j(y)$ is a convex loss over data j , and $x \in \mathbb{R}^d$ is a vector of hyper-parameters for a strongly convex regularizer $R(x, y)$. For any hyper-parameter $x \in \mathbb{R}^d$, the inner-level problem solves for the best parameter $y^*(x)$ over the training set D_{train} under the regularized training loss $\frac{1}{J} \sum_{j \in D_{\text{train}}} \ell_j(y) + R(x, y)$. The goal is to find the hyper-parameter $x^* \in \mathbb{R}^d$ whose corresponding best response $y^*(x)$ yields the least loss over the validation set D_{val} . In practice, continuous hyperparameters are often tuned by a grid search which is exponentially expensive. An efficient SBO algorithm should find the optimal parameters in time increasingly polynomially with the dimension, rather than exponentially. When the training and validation sets are distributed across nodes, the problem becomes a distributed SBO.

Compositional Optimization Let $g(x, \leftarrow) : \mathbb{R}^d \rightarrow \mathbb{R}^{d_y}$ and $f(y, \downarrow) : \mathbb{R}^{d_y} \rightarrow \mathbb{R}$ be two stochastic mappings. Stochastic compositional optimization (SCO) (Wang et al., 2017a) takes the form

$$\min_{x \in \mathbb{X}} \mathbb{E}_{\downarrow} [f(\mathbb{E}_{\leftarrow} [h(x, \leftarrow)]; \downarrow)].$$

SCO applies to risk management (Yang et al., 2019; Ruszczyński, 2021), machine learning (Chen et al., 2021b), and reinforcement learning (Wang et al., 2017b). SCO was identified as a special case of SBO by (Chen et al., 2021c). To see this, we take the inner optimization objective to be $g(x, y) = \mathbb{E}_{\leftarrow} [(y - h(x, \leftarrow))^2]$. Thus $y^*(x) = \underset{y}{\text{argmin}} g(x, y)$ becomes $y^*(x) = \mathbb{E}_{\leftarrow} [h(x, \leftarrow)]$, arriving at an instance of SBO.

1.2 Challenges with Distributed SBO

Despite the recent rapid development of distributed single-level optimization, a method appropriate to DSBO remains elusive. The major hurdle to solving DSBO lies in the absence of explicit knowledge of $y^*(x)$, so that an unbiased gradient for $r_f(x, y^*(x))$ is not available. Recently, by applying the implicit function theorem, (Couellan and Wang, 2016; Ghadimi and Wang, 2018) showed that the gradient of a non-distributed SBO can be expressed as

$$r_x f(x, y^*(x)) = r_{xy}^2 g(x, y^*(x)) [r_y^2 g(x, y^*(x))]^{-1} r_y f(x, y^*(x)), \quad (3)$$

providing a connection between SBO and classical stochastic optimization. Since then, various algorithms have been proposed to obtain sharp estimators for $y^*(x)$ and reduce the bias of the constructed gradients (Chen et al., 2021a; Hong et al., 2020; Ji et al., 2021; Yang et al., 2021). These techniques result in tight convergence analysis and give rise to algorithms widely used in applications. However, no prior algorithms can be applied to the distributed setting.

Solving Problem (1) becomes challenging in the distributed setting in several aspects:

1. Even in single-agent SBO, the lack of $y^*(x)$ makes outer optimization nontrivial, and estimating $y^*(x)$ requires additional stochastic approximation, weighted averaging, and sophisticated calculation.
2. Calculating the outer gradient is highly nontrivial, even when we have an inner solution y . Note that

$$\frac{1}{K} \sum_{k=2}^K r_{xy}^2 g^k(x, y) [r_y^2 g^k(x, y)]^{-1} r_y f^k(x, y) = r_{xy}^2 g(x, y) [r_y^2 g(x, y)]^{-1} r_y f(x, y). \quad (4)$$

In other words, even if the inner problem is solved, the outer gradient requires a new estimation mechanism.

3. Now we move to SBO over distributed networks. In network learning, communication between agents can be limited by the network structure and communication protocol, so taking a simple average across agents may require multiple communication rounds.

Because of the above difficulties, it remains unclear how to estimate the outer gradient sharply under a distributed network. In an attempt to tackle this problem, this paper studies the convergence theory and sample complexity of gossip-based algorithms. In particular, we ask two theoretical questions:

- (i) How does the sample complexity of DSBO scale with the optimality gap and network size?
- (ii) How is the efficiency of DSBO affected by the network structure?

In this paper, we develop a gossip-based stochastic approximation scheme where each agent solves an optimization problem collaboratively by sampling stochastic first- and second-order information using its data and making gossip communications with its neighbors. In addition, we develop novel techniques for convergence analysis to characterize the convergence behavior of our algorithm. To the best of our knowledge, our work is the first to formulate DSBO mathematically and propose an algorithm with theoretical convergence guarantees. Moreover, we show that our algorithm enjoys an $\mathcal{O}(\frac{1}{K})$ sample complexity for finding ϵ -stationary points for nonconvex objectives, where $\mathcal{O}(\cdot)$ hides logarithmic factors, and enjoys an $\mathcal{O}(\frac{1}{K})$ sample complexity for Polyak-Łojasiewicz (PL) functions, subsuming strongly convex optimization. These results subsume the state-of-the-art results for non-federated stochastic bilevel optimization (Chen et al., 2021c) and central-server regimes (Tarzanagh et al., 2022), showing that almost no degradation is induced by network consensus. Further, the above results suggest that our algorithm exhibits a linear speed-up effect for decentralized settings; that is, the required per-agent sample complexities decrease linearly with the number of agents.

2 Related Works

Bilevel optimization was first formulated by (Bracken and McGill, 1973) for solving resource allocation problems. Later, a class of constrained-based algorithms was proposed by (Hansen et al., 1992; Shi et al., 2005), which treats the inner-level optimality condition as constraints to the out-level problem. Recently, (Couellan and Wang, 2016) examined the finite-sum case for unconstrained strongly convex lower-level problems and proposed a gradient-based algorithm that

Stochastic Bilevel Optimization					
	BSA	stocBio	ALSET	FEDNEST	This Work
Network	Single-Agent		Central-Server		Decentralized
Samples in $\downarrow$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2} + \frac{K}{(1-\alpha)^2})$
Samples in $\leftarrow$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^3})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$
	$O(\frac{1}{\epsilon^2})$				

Stochastic Compositional Optimization					
	SCGD	NASA	ALSET	FEDNEST	This Work
Network	Single-Agent		Central-Server		Decentralized
Samples in $\downarrow$	$O(\frac{1}{\epsilon^4})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2} + \frac{K}{(1-\alpha)^2})$
Samples in $\leftarrow$	$O(\frac{1}{\epsilon^4})$	$O(\frac{1}{\epsilon^4})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$	$O(\frac{1}{\epsilon^2})$
	$O(\frac{1}{\epsilon^2})$				

Table 1: Summary of per-agent sample complexities for nonconvex stochastic bilevel and compositional optimization in different settings: BSA (Ghadimi and Wang, 2018), ALSET (Chen et al., 2021c), stocBio (Ji et al., 2021), FEDNEST (Tarzanagh et al., 2022), SCGD (Wang and Liu, 2016), and NASA (Ghadimi et al., 2020). Given the weight matrix W , $\alpha = kW - \frac{1}{K} \mathbf{1}\mathbf{1}^T$.

exhibits asymptotic convergence under certain step-sizes. For SBO, (Ghadimi and Wang, 2018) developed a double-loop algorithm and established the first known complexity results. Subsequently, various methods have been employed to improve the sample complexity, including two-timescale stochastic approximation (Hong et al., 2020), acceleration (Chen et al., 2021a), momentum (Khanduri et al., 2021), and variance reduction (Guo et al., 2021; Ji et al., 2021; Yang et al., 2021).

Distributed optimization was developed to handle real-world large-scale datasets (Dekel et al., 2012; Feyzmahdavian et al., 2016) and graph estimation (Wang et al., 2015). Centralized and decentralized systems are both important problems that have drawn significant attention. A centralized system considers the network topology where there is a central agent that communicates with the remaining agents (Lan and Zhou, 2018), while in a decentralized system (Gao and Huang, 2020; Koloskova et al., 2019; Lan et al., 2020; McMahan et al., 2017), each agent can only communicate with its neighbors by using gossip (Lian et al., 2017b) or gradient tracking (Pu and Nedic, 2021) communication strategies, with applications in multi-agent reinforcement learning (Xu et al., 2020). Variance reduction approaches (Xin et al., 2020, 2021; Lian et al., 2017a) have also been applied to improve the convergence rate of decentralized optimization. Random projection schemes have been studied to handle large sets of constraints (Wang and Bertsekas, 2015, 2016; Liu et al., 2015). All of the above trials were made on vanilla stochastic optimization problems. We are the first to study the decentralized bilevel optimization problem to our knowledge. Table 1 compares this work with prior arts under different settings.

3 Problem Setup

Assumption 3.1 (Sampling Oracle SO) Agent k may query the sampler and receive an independent locally sampled unbiased first- and second-order information $r_x f^k(x, y; \downarrow^k)$, $r_y f^k(x, y; \leftarrow^k)$, $r_y g^k(x, y; \leftarrow^k)$, $r_y^2 g_x^k(x, y; \leftarrow^k)$, and $r_y^2 g^k(x, y; \leftarrow^k)$.

Assumption 3.2 (Gossip Protocol) The network gossip protocol is specified by a $K \rightarrow K$ symmetric weight matrix W with nonnegative entries. Each agent k may receive information from its neighbors, e.g., z_j , $j \in N_k$, and aggregate them by a weighted sum $\sum_{j \in N_k} w_{k,j} z_j$. Further, matrix W satisfies

(i) W is doubly stochastic such that $\sum_i w_{i,j} = 1$ and $\sum_j w_{i,j} = 1$ for all $i, j \in [K]$.

(ii) There exists a constant $\alpha \in (0, 1)$ such that $kW - \frac{1}{K} \mathbf{1}\mathbf{1}^T \succeq \alpha I_K$, where kA denotes the spectral norm of $A \in \mathbb{R}^{K \times K}$.

These assumptions on the adjacency matrix are crucial to ensure the convergence of decentralized algorithms and are commonly made in the literature of decentralized optimization (Lian et al., 2017b).

Assumption 3.3 Let C_f, L_f be positive scalars. The outer level functions $\{f^k\}_{k \in [K]}$ satisfy the followings.

(i) There exists at least one optimal solution to Prob. (1).

(ii) Both $r_x f^k(x, y)$ and $r_y f^k(x, y)$ are L_f -Lipschitz continuous in (x, y) such that for all $x_1, x_2 \in \mathbb{R}^{d_x}$ and $y_1, y_2 \in \mathbb{R}^{d_y}$,

$$\begin{aligned} |r_x f^k(x_1, y_1) - r_x f^k(x_2, y_2)| &\leq L_f (kx_1 - x_2k + ky_1 - y_2k), \text{ and} \\ |r_y f^k(x_1, y_1) - r_y f^k(x_2, y_2)| &\leq L_f (kx_1 - x_2k + ky_1 - y_2k). \end{aligned}$$

(iii) For all $x \in \mathbb{R}^{d_x}$ and $y \in \mathbb{R}^{d_y}$,

$$E[kr_y f^k(x, y; \downarrow^k)k^2] \leq C_f^2 \text{ and } E[kr_x f^k(x, y; \downarrow^k)k^2] \leq C_f^2$$

Assumption 3.4 Let $C_g, L_g, \mathbb{E}_g, \mu_g, \mathbb{E}_g$ be positive scalars. The inner level functions $\{g^k\}_{k \in \mathcal{K}}$ satisfy the followings.

(i) For all $x \in \mathbb{R}^{d_x}$, $g(x, y)$ is μ_g -strongly convex in y .

(ii) For all $x \in \mathbb{R}^{d_x}$ and $y \in \mathbb{R}^{d_y}$, $g^k(x, y)$ is twice continuously differentiable in (x, y) .

(iii) $r_y g^k(x, y)$, $r_{yy}^2 g^k(x, y)$, and $r_{xy}^2 g^k(x, y)$ are Lipschitz continuous in (x, y) such that for all $x_1, x_2 \in \mathbb{R}^{d_x}$ and $y_1, y_2 \in \mathbb{R}^{d_y}$,

$$\begin{aligned} |r_y g^k(x_1, y_1) - r_y g^k(x_2, y_2)| &\leq L_g (kx_1 - x_2k + ky_1 - y_2k), \\ |r_{xy} g^k(x_1, y_1) - r_{xy} g^k(x_2, y_2)| &\leq L_g (kx_1 - x_2k + ky_1 - y_2k), \text{ and} \\ |r_{yy}^2 g^k(x_1, y_1) - r_{yy}^2 g^k(x_2, y_2)| &\leq \mathbb{E}_g (kx_1 - x_2k + ky_1 - y_2k). \end{aligned}$$

(iv) For all $x \in \mathbb{R}^{d_x}, y \in \mathbb{R}^{d_y}$, $r_y g^k(x, y; \leftarrow^k)$, $r_{yy}^2 g^k(x, y; \leftarrow^k)$, and $r_{xy}^2 g^k(x, y; \leftarrow^k)$ have bounded second-order moments such that $E[kr_y g^k(x, y; \leftarrow^k)k^2] \leq C_g^2$, $E[kr_{yy}^2 g^k(x, y; \leftarrow^k)k^2] \leq L_g^2$, and $E[kr_{xy}^2 g^k(x, y; \leftarrow^k)k^2] \leq L_g^2$.

(v) For all $x \in \mathbb{R}^{d_x}, y \in \mathbb{R}^{d_y}$, $\frac{1}{L_g} r_{yy}^2 g^k(x, y; \leftarrow^k)$ has bounded second moment such that $E[k\frac{1}{L_g} r_{yy}^2 g^k(x, y; \leftarrow^k)k^2] \leq (1 - \frac{\mu_g}{L_g})^2$, where $0 < \frac{\mu_g}{L_g} \leq 1$.

We defer the detailed assumptions to Section A.1 of the supplement.

Note that we denote by $kAk_2 = \max(A)$ the induced 2-norm for any matrix A . Here we point out that the above assumptions allow heterogeneity between functions f^k 's and g^k 's over the agents, and the smoothness and boundedness conditions are commonly adopted in SBO (Hong et al., 2020; Chen et al., 2021a).

4 Algorithm

As discussed in Section 1.2, the key challenge to solving DSBO is that each agent only has access to its own data but is required to construct estimators for the gradients and Hessian averaged across all agents. It is particularly challenging to construct such estimators when limited by the network's communication protocol.

To overcome this issue, we propose a gossip-based DSBO algorithm where each agent $k \in \mathcal{K}$ iteratively updates a solution pair (x_t^k, y_t^k) by using the combination of gossip communications and weighted-average stochastic approximation, where y_t^k serves as an estimator of the best response y^* : $\bar{y} = y^*(x^*)$ with $\bar{x}_t := \frac{1}{K} \sum_{k \in \mathcal{K}} x_t^k$. The full algorithm is given in Alg. 1.

Here we briefly explain the idea of our DSBO algorithm. Suppose agent k would like to estimate $r_x f(x_t^k, y_t^k)$ by s_t^k , under Alg. 1 Step 6, it would query the stochastic first-order information using its own data, make gossip communications with neighbors, and update its estimators by taking the weighted average of its previous estimate s_{t-1}^k , neighbors j 's estimate s_{t-1}^j , and the newly sampled gradient $r_x f^k(x_t^k, y_t^k; \downarrow_t^k)$. Roughly speaking, this procedure can be viewed as taking a weighted average of the gradients sampled by all agents over the network, except that the effect of consensus should also be taken into account.

Algorithm 1 Gossip-Based Decentralized Stochastic Bilevel Optimization

input: Step-sizes $\{\epsilon_t\}$, $\{\eta_t\}$, $\{\gamma_t\}$, total iterations T , sampling oracle SO , weight matrix W , smoothness constant L_g , hessian sampling parameter b

$x_0^k = 0, y_0^k = 0, u_0^k = 0, s_0^k = 0, h_0^k = 0, v_{0,i}^k = \mu_g I$ for $i = 1, 2, \dots, b$

- 1: for $t = 0, 1, \dots, T-1$ do
- 2: for $k = 1, \dots, K$ do
- 3: **Local sampling**: Query SO at (x_t^k, y_t^k) to obtain $r_x f^k(x_t^k, y_t^k; \downarrow_t^k)$, $r_y f^k(x_t^k, y_t^k; \downarrow_t^k)$, $r_y g^k(x_t^k, y_t^k; \leftarrow_t^k)$, $r_{xy}^2 g^k(x_t^k, y_t^k; \leftarrow_t^k)$, and $\{r_{yy}^2 g^k(x_t^k, y_t^k; \leftarrow_{t,i}^k)\}_{i=1}^b$.
- 4: **Outer loop update**: $x_{t+1}^k = \frac{1}{K} \sum_{j=1}^K w_{k,j} x_t^j - \epsilon_t s_t^k + \eta_t q_t^k h_t^k$.
- 5: **Inner loop update**: $y_{t+1}^k = \frac{1}{K} \sum_{j=1}^K w_{k,j} y_t^j - \gamma_t r_y g^k(x_t^k, y_t^k; \leftarrow_t^k)$.
- 6: **Estimate** $r_x f(x_t, y_t)$: $s_{t+1}^k = (1 - \epsilon_t) s_t^k + \epsilon_t r_x f^k(x_t^k, y_t^k; \downarrow_t^k)$.
- 7: **Estimate** $r_y f(x_t, y_t)$: $h_{t+1}^k = (1 - \eta_t) h_t^k + \eta_t r_y f^k(x_t^k, y_t^k; \downarrow_t^k)$.
- 8: **Estimate** $r_{xy}^2 g(x_t, y_t)$: $u_{t+1}^k = (1 - \gamma_t) u_t^k + \gamma_t r_{xy}^2 g^k(x_t^k, y_t^k; \leftarrow_t^k)$.
- 9: **Estimate** $[r_{yy}^2 g(x_t, y_t)]^{-1}$: Set $Q_{t+1,0}^k = I$
- 10: for $i = 1, \dots, b$ do
- 11: $v_{t+1,i}^k = (1 - \gamma_t) v_{t,i}^k + \gamma_t r_{yy}^2 g^k(x_t^k, y_t^k; \leftarrow_{t,i}^k)$, 12: $Q_{t+1,i}^k = Q_{t,i}^k + (I - \frac{1}{L_g} v_{t+1,i}^k) Q_{t+1,i-1}^k$
- 13: end for
- 14: Set $q_{t+1}^k = \frac{1}{L_g} Q_{t+1,b}^k$
- 15: end for
- 16: end for

output: $\bar{x}_t = \frac{1}{K} \sum_{k=1}^K x_t^k$

The updates for $r_y f$ and $r_{xy}^2 g$ are conducted in a similar manner (Steps 7 & 8), but it requires extra efforts to evaluate $[r_{yy}^2 g]^{-1}$, because it has no unbiased estimator. To be specific, we note that $\frac{1}{K} \sum_{k=1}^K [r_{yy}^2 g^k]^{-1} \neq [\frac{1}{K} \sum_{k=1}^K r_{yy}^2 g^k]^{-1}$, making the unbiased estimator of the desired term unavailable even if each agent has an unbiased estimator for $[r_{yy}^2 g^k]^{-1}$. This is indeed a unique challenge for decentralized bilevel optimization. To overcome this issue, we propose a novel approach that each agent k constructs b independent estimators $\{v_{t,i}^k\}_{i=1}^b$ for $r_{yy}^2 g(x_t^k, y_t^k)$ using consensus and stochastic approximation. We then estimate $r_{yy}^2 g(x_t^k, y_t^k)$ by utilizing the following approximation.

$$\frac{1}{L_g} r_{yy}^2 g(x_t^k, y_t^k) \stackrel{?}{=} I + \frac{1}{K} \sum_{j=1}^K w_{k,j} \frac{1}{L_g} r_{yy}^2 g(x_t^j, y_t^j) \stackrel{?}{=} I + \frac{1}{K} \sum_{i=1}^b \frac{1}{L_g} v_{t,i}^k.$$

We provide details in Steps 9 - 14.

Finally, each agent will compute the gradient (3) using the estimators obtained in the above procedure and update the outer solution x_t^k by using the combination of gossip communication and stochastic gradient descent (Step 4). The inner loop solution y_t^k can be updated similarly by Step 5.

We highlight the following key features of our DSBO algorithm: (1) each agent only communicates its estimates instead of the raw data in the gossip-communication process, preserving data privacy. (2) agent k makes $O(|N_k|)$ communications with its neighbors in each round, which is much less than the total number of agents in a naive approach. (3) the algorithm is robust to contingencies in the network. If a communication channel fails, the agents can still jointly learn provided that the network is still connected. By contrast, a single-center-multi-user network would fail completely in case of a center failure.

5 Theory

In this section, we analyze the performance of our DSBO algorithm for both nonconvex and PL objectives and derive the convergence rates for both cases.

5.1 Nonconvex Objectives

We first consider the scenario where overall the objective function $F(x)$ is nonconvex. For nonconvex objectives, given the total number of iterations T , we employ the step-sizes in a constant form such that

$$\epsilon_t = C_0 \frac{q}{T}, \quad \eta_t = \frac{q}{T}, \quad \text{and } b = \frac{1}{T}(\log(T)) \text{ for all } t = 0, 1, \dots, T, \quad (5)$$

where $C_0 > 0$ is a small constant.

Compounded effect of consensus and SBO: As discussed earlier, to derive the convergence rate of SBO under decentralized federated setting, the key step is to quantify the compounded effect between the consensus errors induced by the network structure and the biases induced by estimating gradient within (3). Unlike the central-server or non-federated regimes, the consensus errors induced by the decentralized network structure must be handled carefully. We conduct a thorough analysis to derive the contraction of consensus errors, and further show that both the bias and variance of the averaged estimator diminish to zero, establishing a nontrivial convergence argument for the desired gradient and Hessian. In particular, the estimators preserve a concentration property so that their variances decrease proportionally to $1/K$, suggesting that the network consensus effect does not degrade the concentration of the generated stochastic samples. To achieve the best possible convergence rate, we carefully set the algorithm parameters, including the step-sizes ϵ_t , η_t and averaging weights η_t , to control the above consensus errors and biases.

We provide the convergence rate of Alg. 1 in the following theorem and defer the detailed proof to Section B.7 of the supplementary material.

Theorem 5.1 Suppose Assumptions 3.1, 3.2, 3.3, and 3.4 hold. Letting $x_t = \frac{1}{K} \sum_{k=1}^K x_k^t$,

$$\frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[k r F(x_t) k^2] \leq O\left(\frac{1}{KT}\right) + O\left(\frac{K}{T(1-\alpha)^2}\right).$$

More details and proof are deferred to Section B.7 of the supplement.

Effect of consensus: In this result, the $O\left(\frac{K}{T(1-\alpha)^2}\right)$ term represents the errors induced by the consensus of the network. Despite the depending on the network structure, this term diminishes to zero in the order of $O(1/T)$, becoming a small order term when T is large. Consequently, given the network topology, the asymptotic convergence behavior of DSBO is independent of the network structure, answering question (ii) raised in Section 1.

Linear speedup: Because each agent queries $O(b)$ stochastic samples per round, clearly the required iteration and per-agent sample complexities for finding an ϵ -stationary point such that $\frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[k r F(x_t) k^2] \leq \epsilon$ are $O\left(\frac{1}{\epsilon}\right)$ and $O\left(\frac{1}{\epsilon}\right)$, respectively. This result implies that our algorithm achieves a linear-speed-up effect proportionate to the number of agents K . In other words,

in the presence of more agents, each agent needs to obtain fewer stochastic samples to achieve a specified accuracy. Meanwhile, our rate also matches the best-known $O(1/K\epsilon^2)$ iteration and per-agent sample complexities under the decentralized vanilla stochastic gradient descent settings (Lian et al., 2018). This is the first time such a result has been established for DSBO problems.

Single-center-multi-user-federated SBO: We point out that a simplified version of our algorithm solves single-center-multi-user-federated SBO, where the central server collects information directly from each agent and calculates the gradient by employing the weighted-average stochastic approximation scheme for the collected information. In such a scenario, the agents no longer communicate by gossip with neighbors but synchronously receive a common solution from the central server, so that the consensus effect disappears. In Theorem D.1, we show that a variant of our algorithm can indeed achieve the same $O\left(\frac{1}{KT}\right)$ convergence rate in this setting.

5.2 PL Objectives

Next we study the case where the objective function satisfies the following μ -PL condition.

Assumption 5.2 There exists a constant $\mu > 0$ such that the objective satisfies the PL condition:

$$2\mu(F(x) - F^*) \leq k r F(x) k^2.$$

Note that the class of strongly convex functions is a special case of PL functions. To utilize the μ -PL property and achieve fast convergence, unlike the nonconvex case (5) where the step-sizes are set as constants depending on the total number of iterations T , we employ step-sizes in a diminishing form that

$$\epsilon_t = \frac{2}{\mu(C_1 + t)}, \quad \tau_t = \tau = \frac{C_1}{C_1 + t}, \quad \text{and } b = \lceil \log(T) \rceil \text{ for } t \geq 1, \quad (6)$$

where $C_1 > 0$ is a large constant. By following an analytical process similar to that of the nonconvex scenario, in the next result, we derive the convergence rate of Alg. 1 for μ -PL objectives.

Theorem 5.3 Suppose Assumptions 3.1, 3.2, 3.3, and 3.4 hold and the function satisfies the μ -PL Assumption 5.2. Letting $\mathbf{x}_T^* = \frac{1}{K} \sum_{k=1}^K \mathbf{x}_T^k$, then

$$\mathbb{E}[F(\mathbf{x}_T^*)] - F^* \leq O\left(\frac{1}{KT}\right) + O\left(\frac{\ln T}{T^2(1 - \frac{1}{K})^2}\right).$$

The iteration and per-agent sample complexities for finding an ϵ -optimal point $\mathbb{E}[F(\mathbf{x}_T^*)] - F^* \leq \epsilon$ are $\tilde{O}(\frac{1}{\epsilon})$ and $\tilde{O}(\frac{1}{\epsilon})$, respectively.

Details and proof are deferred to Section C.2 of the supplement. This result shows that our algorithm achieves a faster convergence rate for functions satisfying the μ -PL condition in terms of both iteration and sample complexities. First, as in the nonconvex scenario, the consensus error decays in the order of $O(\frac{1}{T^2(1 - \frac{1}{K})^2})$. Dominated by $O(\frac{1}{KT})$, such an order of consensus decaying order indicates that the network structure will not affect Alg. 1's asymptotic convergence behavior under μ -PL objectives. Meanwhile, the above result implies that Alg. 1 speeds up linearly with the number of agents and matches the optimal $O(1/\epsilon)$ sample complexity for single-server vanilla strongly-convex stochastic optimization (Rakhlin et al., 2011). As a result, our algorithm achieves the optimal sample complexities for decentralized stochastic bilevel optimization, establishing the benchmark.

6 Numerical Experiments

In this section, we validate the practical performance of our algorithm in two examples: hyper-parameter optimization and policy evaluation in Markov Decision Processes (MDP), on artificially constructed decentralized ring networks. We run the experiments on a single server desktop computer. We provide the details of federate hyper-parameter optimization and federated policy evaluation.

Hyper-parameter Optimization We consider federated hyper-parameter optimization (2) for a handwriting recognition problem over the Australia handwriting dataset (Chang and Lin, 2011) consisting of data points $(\mathbf{w}, z)$, where $\mathbf{w} \in \mathbb{R}^{14}$ is the feature and $z \in \{0, 1\}$ indicates whether this data point belongs to category "1" or not. In our experiment, we consider the sigmoid loss function that $l(z) = 1/(1 + \exp(-z))$ and a strongly convex regularizer $R(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^d \frac{\alpha_i}{2} \|\mathbf{y}_i - \mathbf{x}_i\|^2$. We consider a ring network of K agents where each agent i preserves two neighbors ($i-1$) and $(i+1)$ and conducts a gossip communication strategy with adjacency matrix $w_{i,j} = \frac{1}{3}$ for $j \in \{i-1, i, i+1\}$.

Before testing Alg. 1, we first randomly split the dataset for training and validation, and then allocate both training and validation dataset over K agents. We then run Algorithm 1 for $T = 20000$ iterations, with $b = 200$, $\epsilon_t = 0.1/K/T$, and $\tau_t = \tau = 10/K/T$.

To provide a benchmark for comparison, we implement a baseline algorithm Decentralized Bilevel Stochastic Approximation (DBSA) algorithm, a naive extension of the double-loop BSA algorithm (Ghadimi and Wang, 2018) in the decentralized setting, formally stated in Section E.1 of the supplementary materials.

We first consider $K = 5$, test Alg. 1 for 5×10^4 iterations, and compare its performance with DBSA. We report the validation loss against total samples in Figure 2 and observe that DSBO exhibits better performance than DBSA. In particular, Further, we observe that our algorithm outperforms the baseline algorithm DBSA in that it requires fewer samples for DSBO to achieve a certain accuracy.

To investigate the efficiency of Alg. 1 to the network structure, we test Alg. 1 over $K = 5, 10, 20$, and report the details of training and validation loss in Figure 2. Further, comparing the performances of Alg. 1 over different agents $K = 5, 10, 20$, we observe that Alg. 1 converges faster when using more agents. This observation suggests that Alg. 1 exhibits a speed-up effect when using more agents.

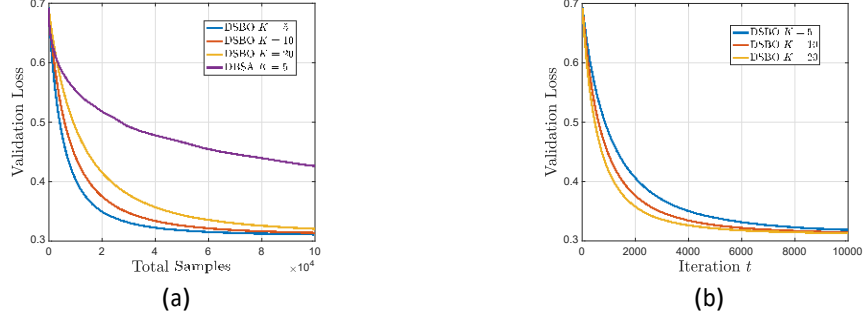


Figure 2: (a) Empirical averaged training loss against total samples for DSBO $K = 5, 10, 20$ and DBSA $K = 5$. (b) Empirical averaged validation loss against iteration for DSBO $K = 5, 10, 20$. All figures are generated through 10 independent simulations over the Australia handwriting dataset.

We provide additional experiments on networks of larger size ($K = 100$) and various topologies (fully-connected and randomly-connected) in Section E.1 of the supplement.

Distributed Policy Evaluation for Reinforcement Learning We consider a multi-agent MDP problem that arises in reinforcement learning. Let S be the state space. For any state $s \in S$, we denote the value function by $V(s)$. We consider the scenario where the value function can be approximated by a linear function such that $V(s) = \phi_s^\top x^*$, where $\phi_s \in \mathbb{R}^m$ is a feature and $x^* \in \mathbb{R}^m$ is an unknown parameter. To obtain the optimal x^* , we consider the following regularized Bellman minimization problem

$$\min_x F(x) = \frac{1}{2|S|} \sum_{s \in S} \phi_s^\top x - \mathbb{E}_{s^0} [r(s, s^0) + \gamma \phi_s^\top x | s]^2 + \frac{\kappa}{2} \|x\|^2,$$

where $r(s, s^0)$ is the random reward incurred by a transition s to s^0 , $\gamma \in (0, 1)$ is the discount factor, κ is the coefficient for the ℓ_2 -regularizer, and the expectation is taken over all random transitions from s to s^0 .

In the federated learning setting, we consider a ring network of K agents. Here each agent k has access to its own data with a heterogeneous random reward function r^k and can only communicate with its two neighbors $k+1$ and $k-1$. We denote by $y_s^k(x) = \phi_s^\top x - \mathbb{E}_{s^0} [\frac{1}{K} \sum_{k=1}^K r^k(s, s^0) + \phi_s^\top x | s]$ where $r^k(s, s^0)$ is the random reward function for agent k . The above problem can then be recast as a bilevel optimization problem

$$\min_{x \in \mathbb{R}^d} f(x, y^*(x)) = \frac{1}{2|S|} \sum_{s \in S} (y_s^*(x))^2 + \frac{\kappa}{2} \|x\|^2.$$

As pointed out by (Wang et al., 2016), the above problem is γ -strongly convex.

In our experiments, we simulate an environment with state space $|S| = 100$ and set the regularizer parameter $\kappa = 1$. We test the performance of Alg. 1 over three cases with $K = 5, 10, 20$ and conduct 10 independent simulations for each K . We implement a baseline double-loop algorithm DSGD that first estimates $y_s^k(x_t)$ with t samples in iteration t and then optimizes the solution x_t . We defer the implementation details of the environment and above algorithms to Section E.2 of the supplement.

We first consider $K = 5$, run Alg. 1 for 10^4 iterations and compare its performance with DSGD. We plot the empirical averaged mean square error $\kappa \|x_t - x^*\|^2$ against total samples generated by all agents in Figure 2. This empirical result suggests that Alg. 1 outperforms the DSGD. To investigate the convergence rate of DSBO, we compare the performance of DSBO over all three scenarios with $K = 5, 10, 20$ and plot the trajectory of the averaged log-error $\log(\kappa \|x_t - x^*\|^2)$ averaged, with one a straight line of slope -1 provided for comparison. We observe that for all four cases, the slopes of $\log(\kappa \|x_t - x^*\|^2)$ are close to -1, matching our theoretical claim in Theorem 5.3 that Alg. 1 converges at a rate of $O(1/t)$ for strongly convex objectives.

In the above experiment, we also note that Alg. 1 converges faster when using more agents. To further demonstrate the linear speedup effect, we compute the total generated samples for finding an ϵ -optimal solution $\kappa \|x_t - x^*\|^2 \leq \epsilon$ and plot the 75% confidence region of the log-sample against the number of agents $K = 5, 10, 20$ in Figure 3. We observe that it takes a roughly same number of samples to

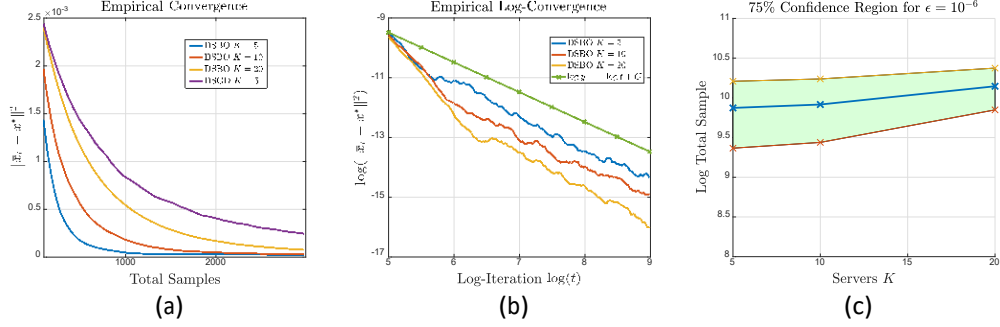


Figure 3: (a) Empirical averaged $\text{MSE} \mathbb{E} \|x^k - x^*\|^2$ against total samples for DSBO $K = 5, 10, 20$ and DSGB $K = 5$. (b) Empirical averaged log-MSE $\log(\mathbb{E} \|x^k - x^*\|^2)$ against log-iteration $\log(t)$ for DSBO $K = 5, 10, 20$. (c) 75% confidence region of log- total samples for achieving $\mathbb{E} \|x^k - x^*\|^2 \leq \epsilon$, with varying network sizes $K = 5, 10, 20$. All figures are generated through 10 independent simulations.

find a 10^{-6} -optimal solution despite different number of agents being involved. This suggests that the per-node sample complexity decreases linearly with K , validating the linear speedup claim in Theorem 5.3. We provide additional numerical results for other optimality level ϵ in Section E.2 of the supplementary material to further demonstrate the linear speedup effect.

7 Conclusion

In this paper, we propose a novel formulation for decentralized stochastic bilevel optimization. We develop a gossip-based stochastic approximation scheme to solve this problem in various settings. We show that our proposed algorithm finds a stationary point at a rate of $O(\frac{1}{K_T})$ for nonconvex objectives, and converges to the optimal solution at a rate of $O(\frac{1}{K_T})$ for PL objectives. Numerical experiments on hyper-parameter optimization and multi-agent federated MDP demonstrate the practical efficiency of our algorithm, exhibit the effect of speed-up in a decentralized setting, and validate our theoretical claims. In our future work, we wish to develop algorithms that achieve lower iteration complexities and enjoy lower communication costs.

Acknowledgments and Disclosure of Funding

Mengdi Wang acknowledges support by NSF grants DMS-1953686, IIS-2107304, CMMI-1653435, and ONR grant 1006977.

References

- Sanjeev Arora, Simon Du, Sham Kakade, Yuping Luo, and Nikunj Saunshi. Provable representation learning for imitation learning via bi-level optimization. In International Conference on Machine Learning, pages 367–376. PMLR, 2020.
- Luca Bertinetto, Joao F Henriques, Philip HS Torr, and Andrea Vedaldi. Meta-learning with differentiable closed-form solvers. arXiv preprint arXiv:1805.08136, 2018.
- Jerome Bracken and James T McGill. Mathematical programs with optimization problems in the constraints. Operations Research, 21(1):37–44, 1973.
- Chih-Chung Chang and Chih-Jen Lin. Libsvm: a library for support vector machines. ACM transactions on intelligent systems and technology (TIST), 2(3):1–27, 2011. URL <https://www.csie.ntu.edu.tw/~cjlin/libsvmtools/datasets/>.
- Tianyi Chen, Yuejiao Sun, and Wotao Yin. A single-timescale stochastic bilevel optimization method. arXiv preprint arXiv:2102.04671, 2021a.

- Tianyi Chen, Yuejiao Sun, and Wotao Yin. Solving stochastic compositional optimization is nearly as easy as solving stochastic optimization. *IEEE Transactions on Signal Processing*, 69:4937–4948, 2021b.
- Tianyi Chen, Yuejiao Sun, and Wotao Yin. Tighter analysis of alternating stochastic gradient method for stochastic nested problems. *arXiv preprint arXiv:2106.13781*, 2021c.
- Nicolas Couellan and Wenjuan Wang. On the convergence of stochastic bi-level gradient methods. *Optimization*, 2016.
- Evin J Cramer, John E Dennis, Jr, Paul D Frank, Robert Michael Lewis, and Gregory R Shubin. Problem formulation for multidisciplinary optimization. *SIAM Journal on Optimization*, 4(4): 754–776, 1994.
- Ofar Dekel, Ran Gilad-Bachrach, Ohad Shamir, and Lin Xiao. Optimal distributed online prediction using mini-batches. *Journal of Machine Learning Research*, 13(1), 2012.
- Hamid Reza Feyzmahdavian, Arda Aytakin, and Mikael Johansson. An asynchronous mini-batch algorithm for regularized stochastic optimization. *IEEE Transactions on Automatic Control*, 61(12):3740–3754, 2016.
- Luca Franceschi, Paolo Frasconi, Saverio Salzo, Riccardo Grazi, and Massimiliano Pontil. Bilevel programming for hyperparameter optimization and meta-learning. In *International Conference on Machine Learning*, pages 1568–1577. PMLR, 2018.
- Hongchang Gao and Heng Huang. Periodic stochastic gradient descent with momentum for decentralized training. *arXiv preprint arXiv:2008.10435*, 2020.
- Jason Ge, Zhaoran Wang, Mengdi Wang, and Han Liu. Minimax-optimal privacy-preserving sparse pca in distributed systems. In *International Conference on Artificial Intelligence and Statistics*, pages 1589–1598. PMLR, 2018.
- Saeed Ghadimi and Guanghui Lan. Accelerated gradient methods for nonconvex nonlinear and stochastic programming. *Mathematical Programming*, 156(1):59–99, 2016.
- Saeed Ghadimi and Mengdi Wang. Approximation methods for bilevel programming. *arXiv preprint arXiv:1802.02246*, 2018.
- Saeed Ghadimi, Andrzej Ruszczyński, and Mengdi Wang. A single timescale stochastic approximation method for nested stochastic optimization. *SIAM Journal on Optimization*, 30(1):960–979, 2020.
- Zhishuai Guo, Quanqi Hu, Lijun Zhang, and Tianbao Yang. Randomized stochastic variance-reduced methods for multi-task stochastic bilevel optimization. *arXiv preprint arXiv:2105.02266*, 2021.
- Pierre Hansen, Brigitte Jaumard, and Gilles Savard. New branch-and-bound rules for linear bilevel programming. *SIAM Journal on scientific and Statistical Computing*, 13(5):1194–1217, 1992.
- Mingyi Hong, Hoi-To Wai, Zhaoran Wang, and Zhuoran Yang. A two-timescale framework for bilevel optimization: Complexity analysis and application to actor-critic. *arXiv preprint arXiv:2007.05170*, 2020.
- Kaiyi Ji, Junjie Yang, and Yingbin Liang. Bilevel optimization: Convergence analysis and enhanced design. In *International Conference on Machine Learning*, pages 4882–4892. PMLR, 2021.
- Prashant Khanduri, Siliang Zeng, Mingyi Hong, Hoi-To Wai, Zhaoran Wang, and Zhuoran Yang. A near-optimal algorithm for stochastic bilevel optimization via double-momentum. *Advances in Neural Information Processing Systems*, 34, 2021.
- Anastasia Koloskova, Tao Lin, Sebastian U Stich, and Martin Jaggi. Decentralized deep learning with arbitrary communication compression. *arXiv preprint arXiv:1907.09356*, 2019.
- Guanghui Lan and Yi Zhou. Random gradient extrapolation for distributed and stochastic optimization. *SIAM Journal on Optimization*, 28(4):2753–2782, 2018.

- Guanghui Lan, Soomin Lee, and Yi Zhou. Communication-efficient algorithms for decentralized and stochastic optimization. *Mathematical Programming*, 180(1):237–284, 2020.
- Xiangru Lian, Mengdi Wang, and Ji Liu. Finite-sum composition optimization via variance reduced gradient descent. In *Artificial Intelligence and Statistics*, pages 1159–1167. PMLR, 2017a.
- Xiangru Lian, Ce Zhang, Huan Zhang, Cho-Jui Hsieh, Wei Zhang, and Ji Liu. Can decentralized algorithms outperform centralized algorithms? a case study for decentralized parallel stochastic gradient descent. *Advances in Neural Information Processing Systems*, 30, 2017b.
- Xiangru Lian, Wei Zhang, Ce Zhang, and Ji Liu. Asynchronous decentralized parallel stochastic gradient descent. In *International Conference on Machine Learning*, pages 3043–3052. PMLR, 2018.
- Jialin Liu, Yuantao Gu, and Mengdi Wang. Averaging random projection: A fast online solution for large-scale constrained stochastic optimization. In *2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 3586–3590, 2015.
- Eli Livne. Integrated aeroservoelastic optimization: status and direction. *Journal of Aircraft*, 36(1): 122–145, 1999.
- Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics*, pages 1273–1282. PMLR, 2017.
- Takayuki Okuno, Akiko Takeda, Akihiro Kawana, and Motokazu Watanabe. On lp-hyperparameter learning via bilevel nonsmooth optimization. *Journal of Machine Learning Research*, 22(245): 1–47, 2021.
- Shi Pu and Angelia Nedic. Distributed stochastic gradient tracking methods. *Mathematical Programming*, 187(1):409–457, 2021.
- Alexander Rakhlin, Ohad Shamir, and Karthik Sridharan. Making gradient descent optimal for strongly convex stochastic optimization. *arXiv preprint arXiv:1109.5647*, 2011.
- A. Ruszczyński and A. Shapiro. Optimization of convex risk functions. *Mathematics of Operations Research*, 31(3):433–452, 2006.
- Andrzej Ruszczyński. A stochastic subgradient method for nonsmooth nonconvex multilevel composition optimization. *SIAM Journal on Control and Optimization*, 59(3):2301–2320, 2021.
- Chenggen Shi, Jie Lu, and Guangquan Zhang. An extended kuhn–tucker approach for linear bilevel programming. *Applied Mathematics and Computation*, 162(1):51–63, 2005.
- Jake Snell, Kevin Swersky, and Richard Zemel. Prototypical networks for few-shot learning. *Advances in neural information processing systems*, 30, 2017.
- Jaroslav Sobieszczanski-Sobieski and Raphael T Haftka. Multidisciplinary aerospace design optimization: survey of recent developments. *Structural optimization*, 14(1):1–23, 1997.
- Murtaza Taj and Andrea Cavallaro. Distributed and decentralized multicamera tracking. *IEEE Signal Processing Magazine*, 28(3):46–58, 2011.
- Davoud Ataee Tarzanagh, Mingchen Li, Christos Thrampoulidis, and Samet Oymak. Fednest: Federated bilevel, minimax, and compositional optimization. *arXiv preprint arXiv:2205.02215*, 2022.
- Shenyinying Tu. Two-Stage Decomposition Algorithms and Their Application to Optimal Power Flow Problems. PhD thesis, Northwestern University, 2021.
- Shenyinying Tu, Andreas Wächter, and Ermin Wei. A two-stage decomposition approach for ac optimal power flow. *IEEE Transactions on Power Systems*, 36(1):303–312, 2020.
- Heinrich Von Stackelberg and Stackelberg Heinrich Von. The theory of the market economy. Oxford University Press, 1952.

- M. Wang and J. Liu. A stochastic compositional gradient method using markov samples. In Proceedings of the 2016 Winter Simulation Conference, pages 702–713. IEEE Press, 2016.
- M. Wang, E. X. Fang, and H. Liu. Stochastic compositional gradient descent: Algorithms for minimizing compositions of expected-value functions. *Mathematical Programming*, 161(1-2): 419–449, 2017a.
- M. Wang, J. Liu, and E. X. Fang. Accelerating stochastic composition optimization. *The Journal of Machine Learning Research*, 18(1):3721–3743, 2017b.
- Mengdi Wang and Dimitri P Bertsekas. Incremental constraint projection methods for variational inequalities. *Mathematical Programming*, 150(2):321–363, 2015.
- Mengdi Wang and Dimitri P Bertsekas. Stochastic first-order methods with random constraint projection. *SIAM Journal on Optimization*, 26(1):681–717, 2016.
- Mengdi Wang, Ji Liu, and Ethan Fang. Accelerating stochastic composition optimization. In *Advances in Neural Information Processing Systems*, pages 1714–1722, 2016.
- Xiaohan Wang, Mengdi Wang, and Y uantao Gu. A distributed tracking algorithm for reconstruction of graph signals. *IEEE Journal of Selected Topics in Signal Processing*, 9(4):728–740, 2015.
- Ran Xin, Usman A Khan, and Soumya Kar. Variance-reduced decentralized stochastic optimization with accelerated convergence. *IEEE Transactions on Signal Processing*, 68:6255–6271, 2020.
- Ran Xin, Usman A Khan, and Soumya Kar. An improved convergence analysis for decentralized online stochastic non-convex optimization. *IEEE Transactions on Signal Processing*, 69:1842–1858, 2021.
- Yue Xu, Zengde Deng, Mengdi Wang, Wenjun Xu, Anthony Man-Cho So, and Shuguang Cui. Voting-based multiagent reinforcement learning for intelligent iot. *IEEE Internet of Things Journal*, 8(4):2681–2693, 2020.
- Junjie Yang, Kaiyi Ji, and Yingbin Liang. Provably faster algorithms for bilevel optimization. *Advances in Neural Information Processing Systems*, 34, 2021.
- S. Yang, M. Wang, and E. X. Fang. Multilevel stochastic gradient methods for nested composition optimization. *SIAM Journal on Optimization*, 29(1):616–659, 2019.

Checklist

The checklist follows the references. Please read the checklist guidelines carefully for information on how to answer these questions. For each question, change the default **[TODO]** to **[Yes]** , **[No]** , or **[N/A]** . You are strongly encouraged to include a justification to your answer, either by referencing the appropriate section of your paper or providing a brief inline description. For example:

- Did you include the license to the code and datasets? **[Yes]**
- Did you include the license to the code and datasets? **[No]** The code and the data are proprietary.
- Did you include the license to the code and datasets? **[N/A]**

Please do not modify the questions and only use the provided macros for your answers. Note that the Checklist section does not count towards the page limit. In your paper, please delete this instructions block and only keep the Checklist section heading above along with the questions/answers below.

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope? **[Yes]**
 - (b) Did you describe the limitations of your work? **[Yes]**
 - (c) Did you discuss any potential negative societal impacts of your work? **[N/A]** This is an algorithmic and theoretical work justified by numerical experiments. We are not aware of any negative societal impacts.
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? **[Yes]**
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? **[Yes]**
 - (b) Did you include complete proofs of all theoretical results? **[Yes]** Complete proofs are provided in the supplement.
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? **[No]** We will provide the code via github in the camera-ready version stage.
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? **[Yes]**
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? **[No]** We provided averaged results over different random seeds and the obtained results are stable.
 - (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? **[Yes]**
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
 - (a) If your work uses existing assets, did you cite the creators? **[Yes]**
 - (b) Did you mention the license of the assets? **[N/A]** The dataset is publicly available and no license is required.
 - (c) Did you include any new assets either in the supplemental material or as a URL? **[N/A]** We are not releasing new assets.
 - (d) Did you discuss whether and how consent was obtained from people whose data you're using/curating? **[N/A]** We cited the website link where data is presented.
 - (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? **[No]** The data does not contain personally identifiable information or offensive content.
5. If you used crowdsourcing or conducted research with human subjects...
 - (a) Did you include the full text of instructions given to participants and screenshots, if applicable? **[N/A]** We did not involve human subjects.

- (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
- (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [N/A]