

A DECENTRALIZED VARIANCE-REDUCED METHOD FOR STOCHASTIC OPTIMIZATION OVER DIRECTED GRAPHS

Muhammad I. Qureshi[†], Ran Xin[‡], Soumya Kar[‡], Usman A. Khan[†]

[†]Tufts University, Medford, MA, USA, [‡]Carnegie Mellon University, Pittsburgh, PA, USA

ABSTRACT

In this paper, we propose a decentralized first-order stochastic optimization method **Push-SAGA** for finite-sum minimization over a strongly connected directed graph. This method features local variance reduction to remove the uncertainty caused by random sampling of the local gradients, global gradient tracking to address the distributed nature of the data, and push-sum consensus to handle the imbalance caused by the directed nature of the underlying graph. We show that, for a sufficiently small step-size, **Push-SAGA** linearly converges to the optimal solution for smooth and strongly convex problems, making it the first linearly-convergent stochastic algorithm over arbitrary strongly-connected directed graphs. We illustrate the behavior and convergence properties of **Push-SAGA** with the help of numerical experiments for strongly convex and non-convex problems.

Index Terms— Stochastic optimization, first-order methods, variance reduction, decentralized algorithms, directed graphs

1. INTRODUCTION AND RELATED WORK

Decentralized optimization has gained a significant interest recently because of its relevance to modern signal processing and machine learning tasks. In many such problems, the size of the data is very large and is available at nodes that are geographically distributed. Bringing the entire dataset to a central processor for learning an inference is computationally prohibitive and further requires extensive communication. It is therefore essential to design algorithms that are computationally efficient and requires minimal communication.

In this paper, we consider decentralized finite-sum minimization over a directed network of n nodes where each node i possesses a local batch of data, i.e.,

$$P : \min_{z \in \mathbb{R}^p} F(z) := \frac{1}{n} \sum_{i=1}^n f_i(z); \quad f_i(z) := \frac{1}{m_i} \sum_{j=1}^{m_i} f_{i,j}(z);$$

where each local cost function $f_i : \mathbb{R}^p \rightarrow \mathbb{R}$ is decomposed into m_i component functions $f_{i,j}$, and j indexes the local data

[†]The work of MIQ and UAK has been partially supported by NSF under awards #1903972 and #1935555. The work of SK and RX has been partially supported by NSF under award #1513936.

batch at node i ; this setup is illustrated in Fig.1. Clearly, an algorithm based on the local batch gradient ∇f_i requires m_i gradient evaluations which could be cost prohibitive. Our focus thus is on stochastic methods that randomly sample one data point per iteration.

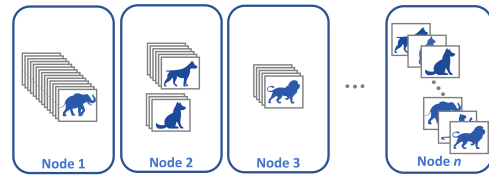


Fig. 1. Offline/Batch optimization problems.

A promising solution for Problem P is decentralized stochastic gradient descent (**DSGD**) [1, 2] that is further extended to directed graphs by **SGP** [3, 4] with the help of push-sum consensus [5, 6]. It is well known that **DSGD** (and **SGP**) incur a steady state error due to the variance of the stochastic gradients, and the difference in the local gradient ∇f_i and the global gradient ∇F [7]. The steady state error due to the local vs. global functions' dissimilarity is eliminated in [8] with the help of gradient tracking [9–14], but the performance still suffers from the variance of stochastic gradients. This variance is eliminated in [15] with the help of variance reduction [16, 17]. All these methods however require doubly stochastic network weight matrices and thus are not applicable to arbitrary directed graphs. Of relevance here are also **SADDOPT** [18] that adds gradient tracking to **SGP** and is a stochastic extension of the methods described in [13, 14].

In this paper, we propose **Push-SAGA** that systematically eliminates the steady state errors in **SGP** and **SADDOPT** with a combination of three ingredients: local variance reduction, global gradient tracking, and push-sum consensus. As a result, **Push-SAGA** converges linearly to the optimal solution for smooth and strongly convex problems over arbitrary directed graphs. To the best of our knowledge, **Push-SAGA** is the first stochastic method that provides a linear convergence guarantee over directed graphs. We now describe the rest of this paper. Section 2 develops and formally describes **Push-SAGA** while Section 3 provides the main result and convergence analysis. Finally, Section 4 presents numerical experiments on strongly convex

and non-convex problems.

Basic Notation: We use lowercase bold letters to represent vectors and uppercase letters for matrices. The column vector of n ones is denoted by $\mathbf{1}_n$ and I_n denotes the $n \times n$ identity matrix. For a primitive column stochastic matrix $\underline{B} \in \mathbb{R}^{n \times n}$, from Perron Frobenius theorem [19], we have that $\underline{B}^{-1} := \lim_{k \rightarrow \infty} \underline{B}^k = \mathbf{1}\mathbf{f}^T$, where $\mathbf{f}$ is the right eigenvector corresponding to the unique eigen-value 1. We use $\rho(\underline{B})$ to denote the spectral radius of $\underline{B}$. We denote $\|\cdot\|_2$ and $\|\cdot\|$ as the Euclidean vector norm and the spectral (matrix) norm, respectively. Since $\rho(\underline{B}) < 1$, it can be shown that there exists a matrix norm $\|\cdot\|$, formally defined in [20], such that $\|\underline{B}\| < 1$.

2. ALGORITHM DEVELOPMENT

Recall Problem P and the fact that the underlying communication graph is not necessarily undirected. **SGP** proposed in [3, 4] is a decentralized optimization algorithm that can be implemented over arbitrary directed graphs since it only requires a column stochastic weight matrix $\underline{B} = \text{fb}_{ir} \in \mathbb{R}^{n \times n}$. **SGP** is formally described as follows: at each node i ,

$$\begin{aligned} \mathbf{y}_i^{k+1} &= \sum_{r=1}^n b_{ir} \mathbf{y}_r^k; \\ \mathbf{x}_i^{k+1} &= \sum_{r=1}^n b_{ir} \mathbf{x}_r^k - \text{rf}_{i;s_i^{k+1}}(\mathbf{z}_i^k); \\ \mathbf{z}_i^{k+1} &= \frac{\mathbf{x}_i^{k+1}}{\mathbf{y}_i^{k+1}}; \end{aligned}$$

where $\mathbf{x}_i^k$ and $\mathbf{z}_i^k$, both in $\mathbb{R}^p$, estimate the global optimizer $\mathbf{z}$ at node i and time k , s_i^{k+1} is drawn uniformly at random from the index set $\{1; \dots; m_i\}$, $\text{rf}_{i;s_i^{k+1}}(\mathbf{z}_i^k)$ is the gradient of a randomly drawn component function from the local batch, and $\mathbf{y}_i^k \in \mathbb{R}^n$ estimates the right eigenvector of $\underline{B}$. At each node i , **SGP** is initialized with an arbitrary $\mathbf{x}_i^0 \in \mathbb{R}^p$, $\mathbf{z}_i^0 = \mathbf{x}_i^0$, and $\mathbf{y}_i^0 = \mathbf{1}$. It can be shown [7] that the **SGP** iterate $\mathbf{z}_i^k$, at each node i , converges to an error ball that is a function of two components, i.e., the variance introduced by stochastic gradient $\text{rf}_{i;s_i^{k+1}}$ and the difference between the local batch gradient rf_i and the global gradient $\mathbf{r}\mathbf{f}$.

The proposed algorithm **Push-SAGA** eliminates the variance of the stochastic gradient, introduced due to the sampling of the local batch at each node, with the help of a variance reduction technique known as **SAGA** [16,17]. In particular, each node i maintains an iterate $\mathbf{g}_i^k$ that estimates the entire local batch gradient rf_i from randomly drawn samples of the local batch, and thus, loosely speaking, $\mathbf{g}_i^k \approx \text{rf}_i(\mathbf{z}_i^k)$. The **SAGA** estimator $\mathbf{g}_i^k$ requires a gradient table at each node i and thus results in extra storage costs. At each iteration of the algorithm, the s_i^{k+1} -th element of this gradient table is updated with the component gradient $\text{rf}_{i;s_i^{k+1}}$ evaluated at the

current iterate $\mathbf{z}_i^{k+1}$. Subsequently, these variance-reduced estimators $\mathbf{f}\mathbf{g}_i^k \mathbf{g}_i^k$'s are used to update the gradient tracking iterate $\mathbf{w}_i^k$ such that, loosely speaking, $\mathbf{w}_i^k \approx \mathbf{g}_i^k$, and thus the descent direction $\mathbf{w}_i^k$ asymptotically tracks the global gradient $\mathbf{r}\mathbf{f}$. The gradient tracking update comes from the dynamic average consensus protocol [21].

Algorithm 1 **Push-SAGA** at each node i

Require: $\mathbf{x}_i^0 \in \mathbb{R}^p$; $\mathbf{z}_i^0 = \mathbf{x}_i^0$; $\mathbf{w}_i^0 = \mathbf{g}_i^0 = \text{rf}_i(\mathbf{z}_i^0)$; $\mathbf{v}_{i;j}^1 = \mathbf{z}_i^0$;
 $\delta_j \in [0, 1]$; m_i ; $\mathbf{y}_i^0 = \mathbf{1}$; $\gamma > 0$; $\text{fb}_{ir} \in \mathbb{R}^{n \times n}$, Gradient
table: $\text{rf}_{i;j}(\mathbf{v}_{i;j}^0) \in \mathbb{R}^{p \times m_i}$

- 1: for $k = 0; 1; 2; \dots$ do
- 2: $\mathbf{x}_i^{k+1} = \sum_{r=1}^n b_{ir} \mathbf{x}_r^k - \mathbf{w}_i^k$;
- 3: $\mathbf{y}_i^{k+1} = \sum_{r=1}^n b_{ir} \mathbf{y}_r^k$;
- 4: $\mathbf{z}_i^{k+1} = \frac{\mathbf{x}_i^{k+1}}{\mathbf{y}_i^{k+1}}$;
- 5: Select s_i^{k+1} uniformly at random from $\{1; \dots; m_i\}$
- 6: $\mathbf{g}_i^{k+1} = \text{rf}_{i;s_i^{k+1}}(\mathbf{z}_i^{k+1}) - \text{rf}_{i;s_i^{k+1}}(\mathbf{v}_{i;s_i^{k+1}}^{k+1}) + \frac{1}{m_i} \sum_{j=1}^{m_i} \text{rf}_{i;j}(\mathbf{v}_{i;j}^{k+1})$
- 7: Replace $\text{rf}_{i;s_i^{k+1}}(\mathbf{v}_{i;s_i^{k+1}}^{k+1})$ by $\text{rf}_{i;s_i^{k+1}}(\mathbf{z}_i^{k+1})$ in the gradient table
- 8: $\mathbf{w}_i^{k+1} = \sum_{r=1}^n b_{ir} \mathbf{w}_r^k + \mathbf{g}_i^{k+1} - \mathbf{g}_i^k$
- 9: if $j = s_i^{k+1}$, then $\mathbf{v}_{i;j}^{k+2} = \mathbf{z}_i^{k+1}$, else $\mathbf{v}_{i;j}^{k+2} = \mathbf{v}_{i;j}^{k+1}$
- 10: end if
- 11: end for

3. LINEAR CONVERGENCE OF PUSH-SAGA

This section formally describes the convergence analysis of **Push-SAGA**. We start with the following assumptions.

Assumption 1 (Strongly connected directed graphs). Each node communicates over a strongly connected directed graph. The underlying weight matrix $\underline{B} = \text{fb}_{ir}$ is such that it follows the sparsity of the graph. In addition, $\underline{B}$ is primitive and column stochastic, i.e., $\mathbf{1}_n^T \underline{B} = \mathbf{1}_n^T$ and $\rho(\underline{B}) < 1$.

Assumption 2 (Smooth and strongly convex cost functions). Each local cost function f_i is μ -strongly convex and each component cost $f_{i;j}$ is L -smooth.

The above assumptions are standard in the literature. Assumption 1 for example can be fulfilled when each node chooses a suitable weight for its outgoing neighbors [13]. A valid choice at node i is $b_{ri} = 1/d_i^{\text{out}}$, for each outgoing neighbor r , that requires each node to know its out-degree denoted by d_i^{out} . Following Assumption 2, it can be verified that the global cost F is L -smooth and μ -strongly convex and thus have a unique global minimizer that is denoted by $\mathbf{z}$.

Based on the above assumptions and **Push-SAGA** in Algorithm 1, we now provide the main result of this paper.

Theorem 1. Consider Problem P under Assumptions 1 and 2. For a sufficiently small step-size $\gamma > 0$, **Push-SAGA** linearly converges to the optimal solution $\mathbf{z}$ at each node.

In the following, we provide the convergence analysis and the auxiliary results that are needed to prove Theorem 1. To this aim, we write **Push-SAGA** in a vector-matrix format. Let $x^k; z^k; g^k; w^k$, all in $\mathbb{R}^{pn}$ be the global vectors that stack the local variables $x_i^k; z_i^k; g_i^k; w_i^k \in \mathbb{R}^p$, respectively, and $y^k \in \mathbb{R}^n$ concatenate y_i^k 's. Similarly, the global matrices can be defined as $B := \frac{B}{I_p}$ and $Y_k := \text{diag}(y^k)$. It can be verified that **Push-SAGA** in Algorithm 1 can be compactly written as

$$x^{k+1} = Bx^k - w^k; \quad (1a)$$

$$y^{k+1} = By^k; \quad (1b)$$

$$z^{k+1} = Y_k^{-1}x^{k+1}; \quad (1c)$$

$$w^{k+1} = Bw^k + g^{k+1} - g^k; \quad (1d)$$

We next introduce four error quantities that will aid in the convergence analysis:

- (i) Network agreement error: $E\|x^k - B^{-1}x^k\|^2$;
- (ii) Optimality gap: $E\|x^k - z^k\|^2$;
- (iii) Mean auxiliary gap: $E[t^k]$;
- (iv) Gradient tracking error: $E\|w^k - B^{-1}w^k\|^2$;

where $\bar{x}^k := \frac{1}{n} \sum_i x_i^k$ and

$$t^k := \sum_{i=1}^n \left(\frac{1}{m_i} \sum_{j=1}^{m_i} kv_{i,j}^k - zk^2 \right).$$

In order to show that **Push-SAGA** converges linearly to the global minimum, we establish that each of these four errors converges to zero linearly, which subsequently leads to $z_i^k \rightarrow z^*$. These properties are derived with the help of an LTI system that captures the evolution of these error quantities. We start with a standard result that comes from the literature on directed graphs.

Lemma 1. [22] Let Assumption 1 hold and $Y^1 := \lim_{k \rightarrow \infty} Y_k$, then $\|Y_k - Y^1\| \leq \frac{1}{k} T$, where $T := \frac{1}{\min_i h_i} \max_i k_i$, $h_i := \sum_j a_{ji}$ and $k_i := \max_j a_{ij}$.

Lemma 2. Consider **Push-SAGA** under Assumptions 1, 2, and for all $k \geq 0$, define $u^k; s^k \in \mathbb{R}^4$ and $G; H_k \in \mathbb{R}^{4 \times 4}$ as

$$u^k := \begin{bmatrix} E[\|x^k - B^{-1}x^k\|^2] \\ E[\|x^k - z^k\|^2] \\ E[t^k] \\ E[\|L^{-2}kw^k - B^{-1}w^k\|^2] \end{bmatrix}, \quad s^k := \begin{bmatrix} E[\|x^k - z^k\|^2] \\ 0 \\ 0 \\ 0 \end{bmatrix};$$

$$G := \begin{bmatrix} \frac{1}{2} & 0 & 0 & \frac{1}{2} \\ \frac{2L^2}{4} - \frac{1}{2} & \frac{1}{2} & \frac{m}{5M} & 0 \\ \frac{188}{4} & \frac{169}{2} & \frac{1}{38} & 0 \\ \frac{1}{2} & \frac{1}{2} & \frac{1}{4} & \frac{3+2}{4} \end{bmatrix};$$

$$H_k := \begin{bmatrix} \frac{2L^2}{4} & 0 & 0 & \frac{1}{2} \\ \frac{2}{4} & 0 & 0 & \frac{1}{5} \\ \frac{188}{4} & 0 & 0 & 0 \\ \frac{1}{2} & 0 & 0 & 0 \end{bmatrix} T^k;$$

where $m := \min_i m_i$; $M := \max_i m_i$; $\gamma := \sup_k \|Y_k\|$; $\gamma^{-1} := \inf_k \|Y_k^{-1}\|$, and $\gamma := \gamma^2(1+T)h$. Then, for $\alpha := \frac{1}{L}$ and the step-size $0 < \alpha < \frac{1}{28L}$, we have

$$\|u^k - Gu^{k-1} + H_k s^{k-1}\| \leq \rho \|u^{k-1}\|. \quad (2)$$

The proof of Lemma 2 can be found in [23] and follows the procedure in [15, 18]. Next, we show that $\|u^k\|$ goes to zero and further characterize its convergence rate. To this aim, note that H_k linearly decays to a zero matrix at $O(k)$. The evolution of (2) further requires the spectral radius ($\rho(G)$) of G . We characterize this spectral radius in Lemma 4 with the help of the matrix perturbation argument in Lemma 3.

Lemma 3. [19] Consider $X; Y \in \mathbb{R}^{n \times n}$. Let λ be a simple eigenvalue of X , and w^L and v^R be left and right eigenvectors of X corresponding to λ . Let (t) be an eigenvalue of $X + tY$. Then (t) is unique and differentiable at $t = 0$, and

$$\frac{d(t)}{dt} \Big|_{t=0} = \frac{w^H Y v}{w^H v}.$$

Lemma 4. Consider G defined in Lemma 2. If the stepsize α is sufficiently small, then $\rho(G) < 1$.

Proof. We start by noting that G can be decomposed into two components: one that is independent of the step-size α and another perturbation term that is controlled by α . In particular, we have that $G := G_0 + P$; such that

$$G_0 := \begin{bmatrix} \frac{1}{2} & 0 & 0 & \frac{1}{2} \\ \frac{2}{4} - \frac{1}{2} & \frac{1}{2} & \frac{1}{38} & 0 \\ \frac{188}{4} & \frac{169}{2} & \frac{1}{38} & 0 \\ \frac{1}{2} & \frac{1}{2} & \frac{1}{4} & \frac{3+2}{4} \end{bmatrix};$$

$$P := \begin{bmatrix} \frac{2L^2}{4} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix};$$

Since G_0 is lower triangular, its eigenvalues appear on the diagonal and it can be verified that $\rho(G_0) = 1$ since $\frac{1}{2} \in [0; 1)$ and $M \geq 1$. We next note that the left eigenvector of G_0 , corresponding to the eigenvalue 1, is $w^L = [0; 1; 0; 0]$, whereas, by choosing $r = 169M + 76m$, it can be shown that the corresponding right eigenvector is

$$v^R = \begin{bmatrix} (1 - \frac{1}{2})^2 \\ 0 \\ -\frac{1}{4Mr} \\ -\frac{1}{2mr} \end{bmatrix}; \quad 1$$

Denoting the eigenvalues of G by λ_i , Lemma 3 leads to

$$\lambda_i = \frac{w^H v}{w^H v} = 1 + \alpha \frac{w^H P v}{w^H v} < 1;$$

because 1 is a simple eigenvalue of G_0 . Finally, since eigenvalues are continuous functions of the parameters of a matrix, we have that $\rho(G)$ decreases from 1 as α slightly increases from zero and the proof follows. $\square$

The following corollary derives the linear convergence of u_k .

6. REFERENCES

- [1] S. S. Ram, A. Nedić, and V. V. Veeravalli, "Distributed stochastic subgradient projection algorithms for convex optimization," *Journal of Optimization Theory and Applications*, vol. 147, no. 3, pp. 516–545, 2010.
- [2] J. Chen and A. H. Sayed, "Diffusion adaptation strategies for distributed optimization and learning over networks," *IEEE Transactions on Signal Processing*, vol. 60, no. 8, pp. 4289–4305, 2012.
- [3] A. Nedić and A. Olshevsky, "Stochastic gradient-push for strongly convex functions on time-varying directed graphs," *IEEE Transactions on Automatic Control*, vol. 61, no. 12, pp. 3936–3947, 2016.
- [4] M. Assran, N. Loizou, N. Ballas, and M. G. Rabbat, "Stochastic gradient-push for distributed deep learning," in *36th International Conference on Machine Learning*, Jun. 2019, vol. 97, pp. 344–353.
- [5] D. Kempe, A. Dobra, and J. Gehrke, "Gossip-based computation of aggregate information," in *44th Annual IEEE Symposium on Foundations of Computer Science*, Oct. 2003, pp. 482–491.
- [6] C. N. Hadjicostis and T. Charalambous, "Average consensus in the presence of delays in directed graph topologies," *IEEE Transactions on Automatic Control*, vol. 59, no. 3, pp. 763–768, 2014.
- [7] R. Xin, S. Kar, and U. A. Khan, "Decentralized stochastic optimization and machine learning: A unified variance-reduction framework for robust performance and fast convergence," *IEEE Signal Processing Magazine*, vol. 37, no. 3, pp. 102–113, 2020.
- [8] S. Pu and A. Nedić, "Distributed stochastic gradient tracking methods," *Mathematical Programming*, 2020.
- [9] J. Xu, S. Zhu, Y. C. Soh, and L. Xie, "Augmented distributed gradient methods for multi-agent optimization under uncoordinated constant stepsizes," in *54th IEEE Conference on Decision and Control*, 2015, pp. 2055–2060.
- [10] P. D. Lorenzo and G. Scutari, "NEXT: in-network non-convex optimization," *IEEE Transactions on Signal and Information Processing over Networks*, vol. 2, no. 2, pp. 120–136, 2016.
- [11] Y. Tian, Y. Sun, and G. Scutari, "Achieving linear convergence in distributed asynchronous multi-agent optimization," *IEEE Transactions on Automatic Control*, pp. 1–1, 2020.
- [12] G. Qu and N. Li, "Harnessing smoothness to accelerate distributed optimization," *IEEE Transactions on Control of Network Systems*, vol. 5, no. 3, pp. 1245–1260, 2017.
- [13] C. Xi, R. Xin, and U. A. Khan, "ADD-OPT: Accelerated distributed directed optimization," *IEEE Transactions on Automatic Control*, vol. 63, no. 5, pp. 1329–1339, 2017.
- [14] A. Nedić, A. Olshevsky, and W. Shi, "Achieving geometric convergence for distributed optimization over time-varying graphs," *SIAM Journal on Optimization*, vol. 27, no. 4, pp. 2597–2633, 2017.
- [15] R. Xin, U. A. Khan, and S. Kar, "Variance-reduced decentralized stochastic optimization with accelerated convergence," *arXiv:1912.04230*, Dec. 2019.
- [16] M. Schmidt, N. Le Roux, and F. Bach, "Minimizing finite sums with the stochastic average gradient," *Mathematical Programming*, vol. 162, no. 1–2, pp. 83–112, 2017.
- [17] A. Defazio, F. Bach, and S. Lacoste-Julien, "SAGA: A fast incremental gradient method with support for non-strongly convex composite objectives," in *Advances in Neural Information Processing Systems*, 2014, pp. 1646–1654.
- [18] M. I. Qureshi, R. Xin, S. Kar, and U. A. Khan, "S-ADDOPT: Decentralized stochastic first-order optimization over directed graphs," *IEEE Control Systems Letters*, vol. 5, no. 3, pp. 953–958, 2021.
- [19] R. A. Horn and C. R. Johnson, *Matrix Analysis*, Cambridge University Press, Cambridge, 2nd edition, 2012.
- [20] R. Xin, A. K. Sahu, U. A. Khan, and S. Kar, "Distributed stochastic optimization with gradient tracking over strongly-connected networks," in *58th IEEE Conference of Decision and Control*, Nice, France, 2019.
- [21] M. Zhu and S. Martínez, "Discrete-time dynamic average consensus," *Automatica*, vol. 46, no. 2, pp. 322–329, 2010.
- [22] A. Nedić and A. Olshevsky, "Distributed optimization over time-varying directed graphs," *IEEE Transactions on Automatic Control*, vol. 60, no. 3, pp. 601–615, 2014.
- [23] M. I. Qureshi, R. Xin, S. Kar, and U. A. Khan, "Push-SAGA: A decentralized stochastic algorithm with variance reduction over directed graphs," *arXiv:2008.06082*, 2020.
- [24] B. Polyak, *Introduction to optimization*, Optimization Software, 1987.