



# Doubly compressed diffusion LMS over adaptive networks

Ibrahim El Khalil Harrane, Rémi Flamary, Cédric Richard

## ► To cite this version:

Ibrahim El Khalil Harrane, Rémi Flamary, Cédric Richard. Doubly compressed diffusion LMS over adaptive networks. 2016 50th Asilomar Conference on Signals, Systems and Computers, Nov 2016, Pacific Grove, France. pp.987-991, 10.1109/ACSSC.2016.7869515 . hal-03633804

**HAL Id: hal-03633804**

**<https://hal.science/hal-03633804>**

Submitted on 7 Apr 2022

**HAL** is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

# Doubly compressed diffusion LMS over adaptive networks

Ibrahim El Khalil Harrane, Rémi Flamary, Cédric Richard  
 Université Côte d'Azur, OCA, CNRS, France  
 ibrahim.harrane@oca.eu, remi.flamary@unice.fr, cedric.richard@unice.fr

**Abstract**—Diffusion LMS is an efficient strategy for solving distributed optimization problems with cooperating agents. Nodes are interested in estimating the same parameter vector and exchange information with their neighbors to improve their local estimates. Successful implementation of such applications relies on a substantial amount of communication resources. In this paper, we introduce diffusion LMS strategies that offer significantly reduced communication load without compromising performance. We perform analyses in the mean and mean-square sense of these algorithms. Simulations results are provided to confirm the theoretical findings.

## I. INTRODUCTION

Distributed adaptation over network has become an active research area with the increasing popularity of wireless sensor networks and, more recently, with the advent of the internet of things. An accessible overview of results in the field can be found in [1]–[5]. The interconnected agents continuously collect data, learn and adapt by carrying out data mining tasks on their own measurements as well as their neighbors. Albeit outperformed by centralized strategies, diffusion strategies are more robust and resilient to agents and links failures. Furthermore, their scalability and flexibility make them attractive for addressing inference problems in a collaborative manner. Among other possible strategies [6]–[8], diffusion LMS plays a central role with its enhanced efficiency and low complexity, and has been extensively studied in the literature in single task [1]–[3] and multitask frameworks [9]–[13]. It was also considered in other contexts such as nonlinear system identification [14] and dictionary learning [15].

Consider an interconnected network of  $N$  nodes. The aim of each node is to estimate an  $L \times 1$  unknown vector  $\mathbf{w}^o$  from collected measurements. Node  $k$  has access to local streaming measurements  $\{d_k(i), \mathbf{u}_{k,i}\}$  where  $d_k(i)$  is a scalar zero-mean reference signal, and  $\mathbf{u}_{k,i}$  is an  $L \times 1$  zero-mean regression vector with covariance matrix  $\mathbf{R}_{\mathbf{u}_k} = \mathbb{E}\{\mathbf{u}_{k,i}\mathbf{u}_{k,i}^\top\} > 0$ . The data at agent  $k$  and time  $i$  are assumed to be related via the linear regression model:

$$d_k(i) = \mathbf{u}_{k,i}^\top \mathbf{w}^o + v_k(i) \quad (1)$$

where  $\mathbf{w}^o$  is the unknown parameter vector to be estimated, and  $v_k(i)$  is a zero-mean i.i.d. noise with variance  $\sigma_{v,k}^2$ . The noise  $v_k(i)$  is assumed to be independent of any other signal. Let  $J_k(\mathbf{w})$  be a differentiable convex cost function at agent  $k$ . In this paper, we shall consider the mean-square-error criterion:

$$J_k(\mathbf{w}) = E\{|d_k(i) - \mathbf{u}_{k,i}^\top \mathbf{w}|^2\} \quad (2)$$

Diffusion LMS strategies seek the minimizer of the following aggregate cost function:

$$J^{\text{glob}}(\mathbf{w}) = \sum_{k=1}^N J_k(\mathbf{w}) \quad (3)$$

in a cooperative manner. Let  $\mathbf{w}_{k,i}$  denote the estimate of the minimizer of (3) at node  $k$  and time instant  $i$ . The general structure of diffusion LMS in its Adapt-then-Combine (ATC) form is given by:

$$\psi_{k,i} = \mathbf{w}_{k,i-1} - \mu_k \sum_{\ell \in \mathcal{N}_k} c_{\ell k} \hat{\nabla}_{\mathbf{w}} J_\ell(\mathbf{w}_{k,i-1}) \quad (4)$$

$$\mathbf{w}_{k,i} = \sum_{\ell \in \mathcal{N}_k} a_{\ell k} \psi_{\ell,i} \quad (5)$$

where  $\hat{\nabla}_{\mathbf{w}} J_\ell(\mathbf{w}_{k,i-1}) = -\mathbf{u}_{\ell,i}[d_\ell(i) - \mathbf{u}_{\ell,i}^\top \mathbf{w}_{k,i-1}]$ ,  $\mathcal{N}_k$  denotes the neighborhood of node  $k$  including  $k$ , and  $\mu_k$  is a positive step-size. Nonnegative coefficients  $\{a_{\ell k}\}$  and  $\{c_{\ell k}\}$  define a left-stochastic matrix  $\mathbf{A}$  and a right-stochastic matrix  $\mathbf{C}$ , respectively. Diffusion LMS provides a scalable estimation framework with high flexibility for ad-hoc deployment. Nevertheless, the requirement for all nodes to exchange their current estimates  $\mathbf{w}_{k,i-1}$ ,  $\hat{\nabla}_{\mathbf{w}} J_\ell(\mathbf{w}_{k,i-1})$  and  $\psi_{k,i}$  with their neighbors at each iteration imposes a substantial burden on communication and energy resources. Therefore, reducing the communication cost while maintaining the benefits of cooperation is of major importance for systems with limited energy budget such as wireless sensor networks.

In recent years, several strategies have been proposed to address this issue. They can be divided into two categories. On the one hand, some authors propose to restrict the number of active links between neighbors at each time instant [5]. On the other hand, there are authors that recommend to reduce the communication load by projecting parameter vectors onto lower dimensional spaces [16], or transmitting only partial parameter vectors [17]–[19]. In [17]–[19], the combination step (5) is redefined as:

$$\mathbf{w}_{k,i} = a_{kk} \psi_{k,i} + \sum_{\ell \in \mathcal{N}_k \setminus \{k\}} a_{\ell k} \left( \mathbf{H}_{\ell,i} \psi_{\ell,i} + [\mathbf{I} - \mathbf{H}_{\ell,i}] \psi_{k,i} \right) \quad (6)$$

where  $\mathbf{H}_{\ell,i}$  is a diagonal entry-selection matrix with  $M$  ones and  $L-M$  zeros on its diagonal. This means that the nodes can use the entries of their own intermediate estimates in lieu of the ones from the neighbors that have not been communicated. Matrix  $\mathbf{H}_{\ell,i}$  can be deterministic, or can randomly select  $M$  entries from all entries.

The literature has mainly focused on the reformulation (6) of the combination step (4) of the diffusion LMS. This means that, with the so-called *partial-diffusion LMS*, only the intermediate estimates  $\psi_{k,i}$  are partially shared by neighboring nodes. However, it is also of interest to consider the adaptation step as it accounts for a major part of information exchanges.

Our main contribution in this paper is to consider and compress every transmitted information over the network, namely, local estimates  $\mathbf{w}_{i,k}$  and estimated gradients  $\hat{\nabla}_w J_k(\mathbf{w}_{\ell,i-1})$ . In a preliminary step, we study the case where only the local estimates in (4) are partially transmitted. Next, we study the case where both the estimates and the estimated gradients are partially transmitted. This method allows a control of network communication flows by setting the numbers  $M$  and  $M_\nabla$  of selected entries in  $\mathbf{w}_{i,k}$  and  $\hat{\nabla}_w J_k(\mathbf{w}_{\ell,i-1})$ , respectively. Finally, we carry out a theoretical analysis of the algorithms in the mean and mean-square sense, and we perform numerical experiments to confirm the theoretical findings.

*Notation:* Boldface small letters denote vectors. All vectors are column vectors. Boldface capital letters denote matrices. The  $(k, \ell)$ -th entry of a matrix is denoted by  $(\cdot)_{k\ell}$ , and the  $(k, \ell)$ -th block of a block matrix is denoted by  $[\cdot]_{k\ell}$ . Matrix trace is denoted by  $\text{trace}\{\cdot\}$ . The expectation operator is denoted by  $\mathbb{E}\{\cdot\}$ . The identity matrix of size  $N$  is denoted by  $\mathbf{I}_N$ , and the all-one vector of length  $N$  is denoted by  $\mathbf{1}_N$ . We denote by  $\mathcal{N}_k$  the set of node indices in the neighborhood of node  $k$ , including  $k$  itself, and  $|\mathcal{N}_k|$  its cardinality. The operator  $\text{col}\{\cdot\}$  stacks its vector arguments on the top of each other to generate a connected vector. The notation  $\text{diag}\{a, b\}$  denotes a diagonal matrix with entries  $a$  and  $b$ . Likewise, the notation  $\text{diag}\{\mathbf{A}, \mathbf{B}\}$  denotes a block diagonal matrix with block entries  $\mathbf{A}$  and  $\mathbf{B}$ . The other symbols will be defined in the context where they are used.

## II. COMPRESSED DIFFUSION LMS

We shall now introduce and analyze the stochastic behavior of the compressed diffusion LMS (CD) as an intermediate step towards the doubly-compressed diffusion LMS (DCD). For the sake of simplicity, we shall consider that  $\mathbf{C}$  is doubly stochastic. We shall set  $\mathbf{A}$  to the identity matrix  $\mathbf{I}_N$ , with the purpose of limiting the network load. The compressed diffusion LMS algorithm is defined as:

$$\begin{aligned} \mathbf{w}_{k,i} = & \mathbf{w}_{k,i-1} + \mu_k \sum_{\ell \in \mathcal{N}_k} c_{\ell,k} \mathbf{u}_{\ell,i} [d_\ell(i) - \mathbf{u}_{\ell,i}^\top (\mathbf{H}_{k,i} \mathbf{w}_{k,i-1} \\ & + (\mathbf{I}_L - \mathbf{H}_{k,i}) \mathbf{w}_{\ell,i-1})] \end{aligned} \quad (7)$$

where  $\mathbf{H}_{\ell,i} = \text{diag}\{\mathbf{h}_{\ell,i}\}$ . We shall assume that  $\mathbf{h}_{\ell,i}$  is a  $L \times 1$  binary vector, generated by randomly setting  $M$  of its  $L$  entries to 1, and the other  $L - M$  entries to 0. We shall also assume that all possible outcomes for  $\mathbf{h}_{\ell,i}$  are equally likely, and i.i.d. over time and space. This leads to:

$$\mathbb{E}\{\mathbf{H}_{\ell,i}\} = \frac{M}{L} \mathbf{I}_L \quad (8)$$

and, given any  $L \times L$  matrix  $\Sigma$ ,

$$\mathbb{E}\{\mathbf{H}_{\ell,i} \Sigma \mathbf{H}_{k,i}\} = \begin{cases} \frac{M}{L} \left( \left(1 - \frac{M-1}{L-1}\right) \mathbf{I}_L \odot \Sigma + \frac{M-1}{L-1} \Sigma \right) & \text{if } \ell = k \\ \left(\frac{M}{L}\right)^2 \Sigma & \text{otherwise} \end{cases} \quad (9)$$

where  $\odot$  denotes the Hadamard entrywise product. Due to the constraint  $\mathbf{1}_L^\top \mathbf{h}_{\ell,i} = M$ , node  $\ell$  receives exactly  $M$  entries from its neighbors at each time instant  $i$ . The missing  $(L - M)$  entries are filled in with estimates that are available at node  $\ell$ .

**Assumption 1** The regression vectors  $\mathbf{u}_{k,i}$  arise from a zero-mean random process that is temporally white and spatially independent. A direct consequence of this assumption is that  $\mathbf{u}_{k,i}$  is independent of  $\mathbf{w}_{\ell,j}$  for all  $\ell$  and  $j < i$ .

**Assumption 2** The matrices  $\mathbf{H}_{i,k}$  arise from a random process that is temporally white, spatially independent, and independent of any other process.

We introduce the  $L \times 1$  error vectors:

$$\tilde{\mathbf{w}}_{k,i} = \mathbf{w}^o - \mathbf{w}_{k,i} \quad (10)$$

and we collect them from across all nodes into the vectors:

$$\tilde{\mathbf{w}}_i = \text{col}\{\tilde{\mathbf{w}}_{1,i}, \tilde{\mathbf{w}}_{2,i}, \dots, \tilde{\mathbf{w}}_{N,i}\} \quad (11)$$

Let  $\mathbf{R}_{u_{\ell,i}} = \mathbf{u}_{\ell,i} \mathbf{u}_{\ell,i}^\top$ . We also introduce:

$$\mathbf{M} = \text{diag}\{\mu_1 \mathbf{I}_L, \mu_2 \mathbf{I}_L, \dots, \mu_N \mathbf{I}_L\} \quad (12)$$

$$\mathbf{R}_i = \text{diag}\left\{\sum_{\ell \in \mathcal{N}_1} c_{\ell,1} \mathbf{R}_{u_{\ell,i}}, \dots, \sum_{\ell \in \mathcal{N}_N} c_{\ell,N} \mathbf{R}_{u_{\ell,i}}\right\} \quad (13)$$

$$\mathbf{R}_{u,i} = \text{diag}\{\mathbf{R}_{u_{1,i}}, \mathbf{R}_{u_{2,i}}, \dots, \mathbf{R}_{u_{N,i}}\} \quad (14)$$

$$\mathbf{H}_i = \text{diag}\{\mathbf{H}_{1,i}, \mathbf{H}_{2,i}, \dots, \mathbf{H}_{N,i}\} \quad (15)$$

$$\mathbf{C} = \mathbf{C} \otimes \mathbf{I}_L \quad (16)$$

where  $\otimes$  denotes the Kronecker product. Finally, we introduce the  $N \times N$  block matrix  $\mathbf{R}_{m,i}$  with each block defined as:

$$[\mathbf{R}_{I-H,i}]_{k\ell} = c_{\ell,k} \mathbf{R}_{u_{\ell,i}} (\mathbf{I}_L - \mathbf{H}_{k,i}) \quad (17)$$

Using recursion (7) and definitions (10), (11), we write:

$$\tilde{\mathbf{w}}_i = \mathbf{B}_i \tilde{\mathbf{w}}_{i-1} - \mathbf{G} \mathbf{s}_i \quad (18)$$

where

$$\mathbf{B}_i = \mathbf{I}_{NL} - \mathbf{M} \mathbf{R}_i \mathbf{H}_i - \mathbf{M} \mathbf{R}_{I-H,i} \quad (19)$$

$$\mathbf{G} = \mathbf{M} \mathbf{C}^\top \quad (20)$$

$$\mathbf{s}_i = \text{col}\{\mathbf{u}_{1,i} v_1(i), \mathbf{u}_{2,i} v_2(i), \dots, \mathbf{u}_{N,i} v_N(i)\} \quad (21)$$

### A. Mean convergence

Taking expectations of both sides of recursion (18), using Assumptions 1 and 2, and  $\mathbb{E}\{\mathbf{s}_i\} = \mathbf{0}$ , we find:

$$\begin{aligned} \mathbb{E}\{\tilde{\mathbf{w}}_i\} &= [\mathbf{I}_{NL} - \mathbb{E}\{\mathbf{M} \mathbf{R}_i \mathbf{H}_i\} - \mathbb{E}\{\mathbf{M} \mathbf{R}_{I-H,i}\}] \mathbb{E}\{\tilde{\mathbf{w}}_{i-1}\} \\ &\quad - \mathbb{E}\{\mathbf{M} \mathbf{C}^\top \mathbf{s}_i\} \\ &= \left[ \mathbf{I}_{NL} - \frac{M}{L} \mathbf{M} \mathbf{R} - \left(1 - \frac{M}{L}\right) \mathbf{M} \mathbf{C}^\top \mathbf{R}_u \right] \mathbb{E}\{\tilde{\mathbf{w}}_{i-1}\} \end{aligned} \quad (22)$$

where

$$\mathbf{R} = \mathbb{E}\{\mathbf{R}_i\} = \text{diag}\{\mathbf{R}_1, \dots, \mathbf{R}_N\} \quad (23)$$

$$\mathbf{R}_u = \mathbb{E}\{\mathbf{R}_{u,i}\} = \text{diag}\{\mathbf{R}_{u_1}, \mathbf{R}_{u_2}, \dots, \mathbf{R}_{u_N}\} \quad (24)$$

with

$$\mathbf{R}_k = \sum_{\ell \in \mathcal{N}_k} c_{\ell,k} \mathbf{R}_{u_\ell} \quad (25)$$

From (22), we observe that the algorithm (7) asymptotically converges in the mean toward  $\mathbf{w}^o$  if, and only if,

$$\rho\left(\mathbf{I}_{NL} - \frac{M}{L} \mathbf{M} \mathbf{R} - \left(1 - \frac{M}{L}\right) \mathbf{M} \mathbf{C}^\top \mathbf{R}_u\right) < 1 \quad (26)$$

where  $\rho(\cdot)$  denotes the spectral radius of its matrix argument. We know that  $\rho(\mathbf{X}) < \|\mathbf{X}\|$  for any induced norm. Then:

$$\begin{aligned} & \rho\left(\mathbf{I}_{NL} - \frac{M}{L} \mathbf{M} \mathbf{R} - \left(1 - \frac{M}{L}\right) \mathbf{M} \mathbf{C}^\top \mathbf{R}_u\right) \\ & \leq \|\mathbf{I}_{NL} - \frac{M}{L} \mathbf{M} \mathbf{R} - \left(1 - \frac{M}{L}\right) \mathbf{M} \mathbf{C}^\top \mathbf{R}_u\|_{b,\infty} \\ & \leq \max_{\ell,k} \|\mathbf{I}_L - \frac{M}{L} \mu_k \mathbf{R}_k - \left(1 - \frac{M}{L}\right) \mu_k c_{\ell k} \mathbf{R}_{u_\ell}\| \end{aligned} \quad (27)$$

where  $\|\cdot\|_{b,\infty}$  denotes the block maximum norm. This leads us to the following condition:

$$\mu_k < \frac{2}{\max_{\ell \in \mathcal{N}_k} \lambda_{\max}\left(\frac{M}{L} \mathbf{R}_k + \left(1 - \frac{M}{L}\right) c_{\ell k} \mathbf{R}_{u_\ell}\right)} \quad (28)$$

where  $\lambda_{\max}(\cdot)$  stands for the maximum eigenvalue of its matrix argument. Furthermore, by Weyl's theorem we have:

$$\mu_k < \frac{2}{\frac{M}{L} \lambda_{\max}(\mathbf{R}_k) + \left(1 - \frac{M}{L}\right) \max_{\ell \in \mathcal{N}_k} [c_{\ell k} \lambda_{\max}(\mathbf{R}_{u_\ell})]} \quad (29)$$

since  $\mathbf{R}_k$  and  $\mathbf{R}_{u_\ell}$  are Hermitian matrices.

### B. Mean-square stability

We shall now analyze the mean-square deviation  $\mathbb{E}\{\|\tilde{\mathbf{w}}_i\|_{\Sigma}^2\}$  where  $\Sigma$  denotes a nonnegative  $N \times N$  block diagonal matrix. The choice of matrix  $\Sigma$  determines the type of information extracted from the network and nodes. Using the independence assumption and (18), we find:

$$\mathbb{E}\|\tilde{\mathbf{w}}_i\|_{\Sigma}^2 = \mathbb{E}\{\tilde{\mathbf{w}}_{i-1}^\top \mathbf{B}_i^\top \Sigma \mathbf{B}_i \tilde{\mathbf{w}}_{i-1}\} + \mathbb{E}\{\mathbf{s}_i^\top \mathbf{G}^\top \Sigma \mathbf{G} \mathbf{s}_i\} \quad (30)$$

On the one hand, note that the second term on the RHS of (30) can be written as:

$$\mathbb{E}\{\mathbf{s}_i^\top \mathbf{G}^\top \Sigma \mathbf{G} \mathbf{s}_i\} = \text{trace}\{\mathbf{G} \Sigma \mathbf{G}^\top \Sigma\} \quad (31)$$

where

$$\mathbf{S} = \text{diag}(\sigma_{v,1}^2 \mathbf{R}_{u,1}, \dots, \sigma_{v,N}^2 \mathbf{R}_{u,N}) \quad (32)$$

On the other hand, the first term on the RHS of (30) can be expressed as:

$$\mathbb{E}\{\tilde{\mathbf{w}}_{i-1}^\top \mathbf{B}_i^\top \Sigma \mathbf{B}_i \tilde{\mathbf{w}}_{i-1}\} = \mathbb{E}\|\tilde{\mathbf{w}}_{i-1}\|_{\Sigma'}^2 \quad (33)$$

where the weighting matrix  $\Sigma'$  is defined as

$$\begin{aligned} \Sigma' &= \mathbb{E}\{\mathbf{B}_i^\top \Sigma \mathbf{B}_i\} \\ &= \Sigma - \frac{M}{L} \Sigma \mathbf{M} \mathbf{R} - \left(1 - \frac{M}{L}\right) \Sigma \mathbf{M} \mathbf{C}^\top \mathbf{R}_u \\ &\quad - \frac{M}{L} \mathbf{R}^\top \mathbf{M} \Sigma - \left(1 - \frac{M}{L}\right) \mathbf{R}_u \mathbf{C} \mathbf{M} \Sigma \\ &\quad + \mathbf{P}_1 + \mathbf{P}_2 + \mathbf{P}_2^\top + \mathbf{P}_3 \end{aligned} \quad (34)$$

with

$$\mathbf{P}_1 = \mathbb{E}\{\mathbf{H}_i \mathbf{R}_i^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R}_i \mathbf{H}_i\} \quad (35)$$

$$\mathbf{P}_2 = \mathbb{E}\{\mathbf{H}_i \mathbf{R}_i^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R}_{I-H,i}\} \quad (36)$$

$$\mathbf{P}_3 = \mathbb{E}\{\mathbf{R}_{I-H,i}^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R}_{I-H,i}\} \quad (37)$$

To evaluate  $\mathbf{P}_1$ , we write:

$$\mathbf{P}_1 = \mathbb{E}_{\mathcal{H}}\left\{\mathbf{H}_i \mathbb{E}_{\mathcal{R}}\{\mathbf{R}_i^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R}_i\} \mathbf{H}_i\right\} \quad (38)$$

Consider first  $\mathbb{E}_{\mathcal{H}}\{\mathbf{H}_i \mathbf{\Pi} \mathbf{H}_i\}$  where  $\mathbf{\Pi}$  denotes any  $NL \times NL$  deterministic matrix. It can be shown that:

$$\begin{aligned} & \mathbb{E}\{\mathbf{H}_i \mathbf{\Pi} \mathbf{H}_i\} \\ &= \alpha_1 (\mathbf{I}_N \otimes \mathbf{1}_{LL}) \odot \mathbf{\Pi} + \alpha_2 \mathbf{I}_{NL} \odot \mathbf{\Pi} + \alpha_3 \mathbf{\Pi} \end{aligned} \quad (39)$$

where  $\mathbf{1}_{LL}$  denotes an all-one  $L \times L$  matrix, and

$$\alpha_1 = \frac{M}{L} \left(\frac{M-1}{L-1} - \frac{M}{L}\right) \quad \alpha_2 = \frac{M}{L} \left(1 - \frac{M-1}{L-1}\right) \quad \alpha_3 = \left(\frac{M}{L}\right)^2$$

Next, consider the inner expectation in the RHS of (38). The evaluation of this expectation depends on higher-order moments of the regression data. While we can continue with the analysis by calculating these terms, it is sufficient for the exposition to focus on the case of sufficiently small step-sizes where a reasonable approximation is [1]:

$$\mathbb{E}\{\mathbf{R}_i^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R}_i\} = \mathbf{R}^\top \mathbf{M} \Sigma \mathbf{M} \mathbf{R} \quad (40)$$

Substituting (39) and (40) into (38) leads to  $\mathbf{P}_1$ . Consider  $\mathbf{P}_2$ . Using the same higher-order approximation as above, we have:

$$\begin{aligned} [\mathbf{P}_2]_{k\ell} &= \frac{M}{L} \mathbf{R}_k [\mathbf{M} \Sigma \mathbf{M}]_{kk} c_{\ell k} \mathbf{R}_{u_\ell} \\ &\quad - \mathbb{E}\{\mathbf{H}_{k,i} \mathbf{R}_{k,i} [\mathbf{M} \Sigma \mathbf{M}]_{kk} c_{\ell k} \mathbf{R}_{u_\ell,i} \mathbf{H}_{k,i}\} \end{aligned} \quad (41)$$

The second term in RHS of (41) is of the form:

$$[\varphi(\mathbf{\Pi})]_{k\ell} = \mathbb{E}\{\mathbf{H}_{k,i} [\mathbf{\Pi}]_{k\ell} \mathbf{H}_{k,i}\} \quad (42)$$

with  $[\mathbf{\Pi}]_{k\ell} = \mathbf{R}_{k,i} [\mathbf{M} \Sigma \mathbf{M}]_{kk} c_{\ell k} \mathbf{R}_{u_\ell,i}$ . It can be shown that:

$$\varphi(\mathbf{\Pi}) = \alpha_2 (\mathbf{1}_{NN} \otimes \mathbf{I}_L) \odot \mathbf{\Pi} + (\alpha_1 + \alpha_3) \mathbf{\Pi} \quad (43)$$

This leads to:

$$\mathbf{P}_2 = \frac{M}{L} \mathbf{R} \mathbf{M} \Sigma \mathbf{M} \mathbf{C}^\top \mathbf{R}_u - \varphi(\mathbf{R} \mathbf{M} \Sigma \mathbf{M} \mathbf{C}^\top \mathbf{R}_u) \quad (44)$$

Finally, using the same higher-order approximation as above with  $\mathbf{P}_3$  leads to:

$$[\mathbf{P}_3]_{k\ell} = \sum_{m=1}^N \mathbb{E}\{[\mathbf{R}_{I-H,i}]_{km}^\top [\mathbf{M} \Sigma \mathbf{M}]_{mm} [\mathbf{R}_{I-H,i}]_{m\ell}\} \quad (45)$$

and

$$\begin{aligned} \mathbf{P}_3 &= \left(1 - 2\frac{M}{L}\right) \mathbf{R}_u \mathbf{C} \mathbf{M} \Sigma \mathbf{M} \mathbf{C}^\top \mathbf{R}_u \\ &\quad + \varphi(\mathbf{R}_u \mathbf{C} \mathbf{M} \Sigma \mathbf{M} \mathbf{C}^\top \mathbf{R}_u) \end{aligned} \quad (46)$$

Following the same reasoning as in [1] allows to express matrix  $\Sigma'$  in a vector form as

$$\boldsymbol{\sigma}' = \mathcal{F} \boldsymbol{\sigma} \quad (47)$$

where  $\sigma = \text{vec}(\Sigma)$ ,  $\sigma' = \text{vec}(\Sigma')$  and:

$$\begin{aligned} \mathcal{F} &= \mathbf{I}_{(NM)^2} - \frac{M}{L}(\mathcal{R}\mathcal{M} \otimes \mathbf{I}_{NL}) \\ &\quad - (1 - \frac{M}{L})(\mathcal{R}_u\mathcal{C}\mathcal{M} \otimes \mathbf{I}_{NL}) - \frac{M}{L}(\mathbf{I}_{NL} \otimes \mathcal{R}\mathcal{M}) \quad (48) \\ &\quad - (1 - \frac{M}{L})(\mathbf{I}_{NL} \otimes \mathcal{R}_u\mathcal{C}\mathcal{M}) + \mathbf{Z}_1 + \mathbf{Z}_2 + \mathbf{Z}_2^\top + \mathbf{Z}_4 \end{aligned}$$

where matrices  $\mathbf{Z}_j$  are obtained by applying the  $\text{vec}(\cdot)$  operator on  $\mathbf{P}_j$  and using the following property:

$$\text{vec}(\mathbf{ABC}) = (\mathbf{C}^\top \otimes \mathbf{A})\text{vec}(\mathbf{B}) \quad (49)$$

The analytical expressions of  $\mathbf{Z}_i$  are not provided in this paper due to the lake of space. They will be made available in an extended version of this paper. Substituting (31) and (33) in (30), and applying the  $\text{vec}$  operation to both sides, we get:

$$\mathbb{E}\|\tilde{\mathbf{w}}_i\|_\sigma^2 = \mathbb{E}\|\tilde{\mathbf{w}}_{i-1}\|_{\mathcal{F}\sigma}^2 + [\text{vec}(\mathcal{Y}^\top)]^\top \sigma \quad (50)$$

where

$$\mathcal{Y} = \mathcal{G}\mathcal{S}\mathcal{G}^\top \quad (51)$$

### III. DOUBLY-COMPRESSED DIFFUSION LMS

The compressed diffusion LMS studied before only partially transmit local estimates while the estimated gradient vectors are fully transmitted through the network. We shall now introduce the doubly-compressed diffusion LMS algorithm, which offers an additional compression layer by transmitting partial estimated gradient vectors. Node  $k$  receives  $M_\nabla$  entries of  $\hat{\nabla}_w J_k(\mathbf{w}_{\ell,i-1})$  from its neighbors at each time instant  $i$ . The missing  $(L - M_\nabla)$  entries are filled in with the gradient estimates that are available at node  $k$ . The algorithm can be formulated as:

$$\begin{aligned} \mathbf{w}_{k,i} &= \mathbf{w}_{k,i-1} + \mu_k \sum_{\ell \in \mathcal{N}_k} c_{\ell,k} \mathbf{Q}_{\ell,i} \mathbf{u}_{\ell,i} [d_\ell(i) \\ &\quad - \mathbf{u}_{\ell,i}^\top (\mathbf{H}_{k,i} \mathbf{w}_{k,i-1} + (\mathbf{I}_L - \mathbf{H}_{k,i}) \mathbf{w}_{\ell,i-1})] \quad (52) \\ &\quad + \mu_k \sum_{\ell \in \mathcal{N}_k} c_{\ell,k} (\mathbf{I}_L - \mathbf{Q}_{\ell,i}) \mathbf{u}_{k,i} [d_k(i) \\ &\quad - \mathbf{u}_{k,i}^\top (\mathbf{H}_{k,i} \mathbf{w}_{k,i-1} + (\mathbf{I}_L - \mathbf{H}_{k,i}) \mathbf{w}_{\ell,i-1})] \end{aligned}$$

where  $\mathbf{Q}_{\ell,i}$  is a stochastic selection matrix that has the same properties as  $\mathbf{H}_{k,i}$  defined in the previous section.

**Assumption 3** The matrices  $\mathbf{Q}_{i,\ell}$  arise from a random process that is temporally white, spatially independent, and independent of any other process.

Proceeding as in the previous section, we find that:

$$\tilde{\mathbf{w}}_i = \mathcal{B}_i \tilde{\mathbf{w}}_{i-1} - \mathcal{G}_i \mathbf{s}_i \quad (53)$$

where

$$\begin{aligned} \mathcal{B}_i &= \mathbf{I}_{NL} - \mathcal{M}\mathcal{R}_{Q,i}\mathcal{H}_i - \mathcal{M}\mathcal{Q}'_i\mathcal{R}_{u,i}\mathcal{H}_i \\ &\quad - \mathcal{M}\mathcal{R}_{Q(I-H),i} - \mathcal{M}\mathcal{R}_{(I-Q)(I-H),i} \quad (54) \end{aligned}$$

$$\mathcal{G}_i = \mathcal{M}\mathcal{Q}_i\mathcal{C}^\top + \mathcal{M}\mathcal{Q}'_i \quad (55)$$

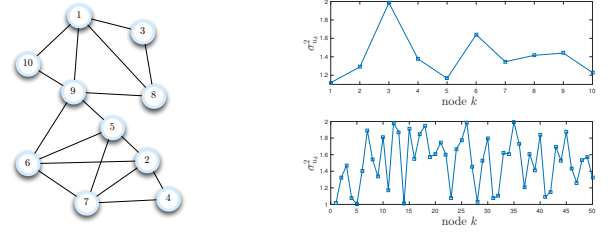


Fig. 1: (left) Network topology. (right) Variance  $\sigma_{u_k}^2$  of regressors in Experiment 1 (top) and in Experiment 2 (bottom).

with

$$\begin{aligned} \mathcal{R}_{Q,i} &= \text{diag} \left\{ \sum_{\ell \in \mathcal{N}_1} c_{\ell 1} \mathbf{Q}_{\ell,i} \mathbf{R}_{u_{\ell,i}}, \dots, \sum_{\ell \in \mathcal{N}_N} c_{\ell N} \mathbf{Q}_{\ell,i} \mathbf{R}_{u_{\ell,i}} \right\} \\ \mathcal{Q}_i &= \text{diag} \{ \mathbf{Q}_{1,i}, \mathbf{Q}_{2,i}, \dots, \mathbf{Q}_{N,i} \} \\ \mathcal{Q}'_i &= \text{diag} \left\{ \sum_{\ell \in \mathcal{N}_1} c_{\ell 1} (\mathbf{I}_L - \mathbf{Q}_{\ell,i}), \dots, \sum_{\ell \in \mathcal{N}_N} c_{\ell N} (\mathbf{I}_L - \mathbf{Q}_{\ell,i}) \right\} \\ [\mathcal{R}_{Q(I-H),i}]_{k\ell} &= c_{\ell k} \mathbf{Q}_{\ell,i} \mathbf{R}_{u_{\ell,i}} (\mathbf{I}_L - \mathbf{H}_{k,i}) \\ [\mathcal{R}_{(I-Q)(I-H),i}]_{k\ell} &= c_{\ell k} (\mathbf{I}_L - \mathbf{Q}_{\ell,i}) \mathbf{R}_{u_{\ell,i}} (\mathbf{I}_L - \mathbf{H}_{k,i}) \end{aligned}$$

The performance analysis of this algorithm is similar to the analysis of compressed diffusion LMS, although more tedious. The following result can be obtained following the same steps.

#### A. Convergence in mean

Taking expectations of both sides of (53), we obtain:

$$\begin{aligned} \mathbb{E}\{\tilde{\mathbf{w}}_i\} &= \left( \mathbf{I}_{NL} - \frac{MM_\nabla}{L^2} \mathcal{M}\mathcal{R} - \frac{M}{L} \left( 1 - \frac{M_\nabla}{L} \right) \mathcal{M}\mathcal{R}_u \right. \\ &\quad \left. - \frac{M_\nabla}{L} \left( 1 - \frac{M}{L} \right) \mathcal{M}\mathcal{C}^\top \mathcal{R}_u \right. \quad (56) \\ &\quad \left. - \left( 1 - \frac{M_\nabla}{L} \right) \left( 1 - \frac{M}{L} \right) \mathcal{M}\mathcal{R}_u \mathcal{C}^\top \right) \mathbb{E}\{\tilde{\mathbf{w}}_{i-1}\} \end{aligned}$$

From (56), the algorithm is stable as long as the following condition is met:

$$\mu_k < \frac{2}{\lambda_{\max,k}} \quad (57)$$

where

$$\begin{aligned} \lambda_k &= \frac{MM_\nabla}{L^2} \lambda_{\max}(\mathbf{R}_k) + \frac{M}{L} \left( 1 - \frac{M_\nabla}{L} \right) \lambda_{\max}(\mathbf{R}_{u_k}) \\ &\quad + \frac{M_\nabla}{L} \left( 1 - \frac{M}{L} \right) \max_{\ell \in \mathcal{N}_k} [c_{\ell k} \lambda_{\max}(\mathbf{R}_{u_{\ell}})] \quad (58) \\ &\quad + \left( 1 - \frac{M_\nabla}{L} \right) \left( 1 - \frac{M}{L} \right) \max_{\ell \in \mathcal{N}_k} [c_{\ell k} \lambda_{\max}(\mathbf{R}_{u_k})] \end{aligned}$$

Due to lake of space and of the complexity of the analysis in the mean-square sense, this analysis is not reported here. It will be made available in an extended version of this article.

### IV. SIMULATION RESULTS

We shall now evaluate the accuracy of the theoretical models with simulations. We performed two experiments to determine the performance and efficiency of both algorithms. First, we considered a small network in order to validate the theoretical models. Then, we considered a larger network and high dimensional measurements in order to test the algorithms. For both experiments, parameter vectors  $\mathbf{w}^o$  were generated from

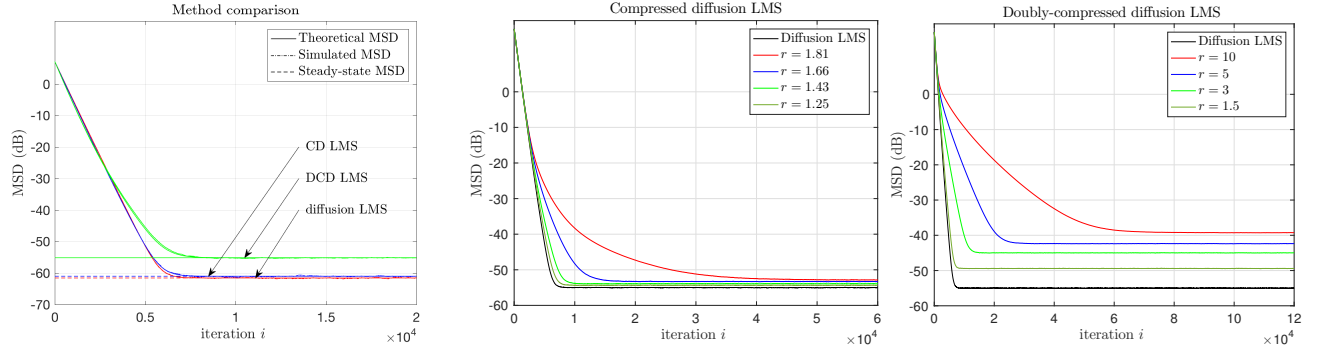


Fig. 2: Theoretical model and simulated curves (left), evolution of the MSD as a function of the compression ratio for compressed diffusion LMS (center), and doubly-compressed diffusion LMS (right).

a zero-mean Gaussian distribution. The input data  $\mathbf{u}_{k,i}$  were drawn from zero-mean Gaussian distributions with covariance  $\mathbf{R}_{u,k} = \sigma_{u,k}^2 \mathbf{I}_L$  reported in Fig. 1. The weighting matrices  $\mathbf{C}$  were generated using the Metropolis rule [1]. Noises  $v_k(i)$  were zero-mean, i.i.d. and Gaussian distributed with variance  $\sigma_{v,k}^2 = 10^{-3}$ . The step-sizes were set to  $\mu_k = 10^{-3}$ . The simulation results were averaged over 100 Monte Carlo runs.

1) *Experiment 1:* Due to the high dimensionality of the theoretical models, we considered the network of  $N = 10$  nodes shown in Fig. 1 (left). We set  $L = 5$ ,  $M = 3$  and  $M_{\nabla} = 1$ . This resulted in compression ratios of  $\frac{10}{8}$  and  $\frac{5}{2}$  for the compressed and the doubly-compressed diffusion LMS, respectively. It can be observed in Fig. 2 (left) that the theoretical models fit accurately the simulated results for both algorithms. Unsurprisingly, diffusion LMS outperformed its compressed counterparts at the expense of a higher communication load.

2) *Experiment 2:* Since compression is relevant for relatively large data flows, we considered a network with  $N = 50$  agents. Data dimension was set to  $L = 50$ . Figure 2 illustrates the performance of the algorithms for different compression ratios. Note that maximum compression ratio that can be reached with the compressed diffusion LMS is  $\frac{100}{55}$  since estimated gradient vectors are fully transmitted. Doubly-compressed diffusion LMS is more flexible since the compression ratio can be controlled via  $M$  and  $M_{\nabla}$ .

## V. CONCLUSION

In order to preserve energy resources, we investigated compression techniques for diffusion LMS that consist of exchanging partial local estimates. We carried out an analysis of the stochastic behavior of the proposed algorithms in the mean and mean-square sense. Finally we provided simulation results to illustrate the accuracy of the theoretical models and the performance of the compression layer.

## REFERENCES

- [1] A. H. Sayed, "Diffusion adaptation over networks," in *Academic Press Library in Signal Processing*, R. Chellapa and S. Theodoridis, Eds. Elsevier, 2014. Also available as arXiv:1205.4220 [cs.MA], May 2012., pp. 322–454.
- [2] A. H. Sayed, S.-Y. Tu, J. Chen, X. Zhao, and Z. J. Towfic, "Diffusion strategies for adaptation and learning over networks: an examination of distributed strategies and network behavior," *IEEE Signal Processing Magazine*, vol. 30, no. 3, pp. 155–171, 2013.
- [3] A. H. Sayed, "Adaptive networks," *Proceedings of the IEEE*, vol. 102, no. 4, pp. 460–497, 2014.
- [4] X. Zhao, S.-Y. Tu, and A. H. Sayed, "Diffusion adaptation over networks under imperfect information exchange and non-stationary data," *IEEE Transactions on Signal Processing*, vol. 60, no. 7, pp. 3460–3475, 2012.
- [5] R. Arablouei, S. Werner, K. Dogancay, and Y.-F. Huang, "Analysis of a reduced-communication diffusion LMS algorithm," *Signal Processing*, vol. 117, pp. 355–361, 2015.
- [6] A. Nedic and A. Ozdaglar, "Distributed subgradient methods for multi-agent optimization," *IEEE Transactions on Automatic Control*, vol. 54, no. 1, pp. 48–61, 2009.
- [7] M. G. Rabbat and R. D. Nowak, "Quantized incremental algorithms for distributed optimization," *IEEE J. Sel. Areas Commun.*, vol. 23, no. 4, pp. 798–808, Apr. 2005.
- [8] C. G. Lopes and A. H. Sayed, "Incremental adaptive strategies over distributed networks," *IEEE Transactions on Signal Processing*, vol. 55, no. 8, pp. 4064–4077, 2007.
- [9] J. Chen, C. Richard, A. O. Hero, and A. H. Sayed, "Diffusion lms for multitask problems with overlapping hypothesis subspaces," in *Proc. IEEE MLSP'14*, Reims, France, 2014, pp. 1–6.
- [10] J. Chen, C. Richard, and A. H. Sayed, "Multitask diffusion adaptation over networks," *IEEE Transactions on Signal Processing*, vol. 62, no. 16, pp. 4129–4144, 2014.
- [11] —, "Diffusion LMS over multitask networks," *IEEE Transactions on Signal Processing*, vol. 63, no. 11, pp. 2733–2748, 2015.
- [12] R. Nassif, C. Richard, A. Ferrari, and A. H. Sayed, "Multitask diffusion adaptation over asynchronous networks," *IEEE Transactions on Signal Processing*, vol. 64, no. 11, pp. 2835–2850, 2016.
- [13] —, "Proximal multitask learning over networks with sparsity-inducing coregularization," *IEEE Transactions on Signal Processing*, vol. 64, no. 23, pp. 6329–6344, 2016.
- [14] W. Gao, J. Chen, C. Richard, and J. Huang, "Diffusion adaptation over networks with kernel least-mean-square," in *Proc. IEEE CAMSAP'15*, Cancun, Mexico, 2015.
- [15] P. Chainais and C. Richard, "Learning a common dictionary over a sensor network," in *Proc. IEEE CAMSAP'13*, Saint Martin, French West Indies, 2013.
- [16] M. O. Sayin and S. S. Kozat, "Compressive diffusion strategies over distributed networks for reduced communication load," *IEEE Transactions on Signal Processing*, vol. 62, no. 20, pp. 5308–5323, 2014.
- [17] R. Arablouei, S. Werner, Y.-F. Huang, and K. Dogancay, "Distributed least mean-square estimation with partial diffusion," *IEEE Transactions on Signal Processing*, vol. 62, no. 2, pp. 472–484, 2014.
- [18] R. Arablouei, K. Dogancay, S. Werner, and Y.-F. Huang, "Adaptive distributed estimation based on recursive least-squares and partial diffusion," *IEEE Transactions on Signal Processing*, vol. 62, no. 14, pp. 3510–3522, 2014.
- [19] V. Vadidpour, A. Rastegarnia, A. Khalili, and S. Sane'i, "Partial-diffusion least mean-square estimation over networks under noisy information exchange," *arXiv preprint arXiv:1511.09044*, 2015.