

# Stochastic convergence of parallel asynchronous adaptive first-order methods

S. Gratton\* and Ph. L. Toint†

23 VI 2026

## Abstract

A new class of asynchronous adaptive first-order optimization methods is introduced, comprising asynchronous variants of several popular algorithms. Versions of these methods using momentum and/or inexact normalization are also considered. The convergence of methods in the class on non-convex functions is analyzed in a fully stochastic setting, and is shown to be of order  $\mathcal{O}(1/\sqrt{t})$  (up to logarithmic factors) under bounded-delay and cumulative variance assumptions. Numerical experiments suggest that such asynchronous adaptive algorithms are very relevant in heterogeneous large-scale machine learning systems.

**Keywords:** Unconstrained nonconvex optimization, first-order methods, global rate of convergence, asynchronous methods, parallel computing, heterogeneous deep learning problems.

## 1 Introduction

We consider the problem of finding a first-order critical point for the optimization problem

$$\min_{X \in \mathbf{R}^N} F(X) = \mathbb{E}[f(X, \xi)] \quad (1.1)$$

where  $f$  is a smooth, possibly nonconvex function of  $X$  and  $\xi$  is a random variable (whose distribution may depend on  $X$ ) defined on some probability space  $(\Omega, \mathcal{B}, \mathbb{P})$ . In other words, we are seeking  $X$  such that  $\nabla_X \mathbb{E}[f(X, \xi)] = 0$ . Motivated by the growing importance of very large models in deep learning, a substantial literature has been devoted to the study of asynchronous algorithms for solving (1.1). By far the most popular are stochastic asynchronous steepest descent methods (also called asynchronous gradient methods), whose modern efficient lock-free framework has been pioneered by the **Hogwild!** method [43]. This method, like most of the methods we discuss in this paper, maintains a vector  $\hat{X}$  of values for the problem's variables in a memory shared by a number of processes/processors, each of which sequentially

1. selects a subset or block of the variables from  $\hat{X}$ ,
2. computes the gradient of  $f$  at  $\hat{X}$  with respect to the variables of the block (this requires reading the vector  $\hat{X}$  from shared memory),
3. computes a gradient step (with method-specific stepsize) in the subspace spanned by this block using the block gradient, resulting in new values of the block variables,
4. updates the values of  $\hat{X}$  corresponding to the block with these new values.

---

\*Université de Toulouse, INP, IRIT, Toulouse, France. Email: serge.gratton@enseeiht.fr. Work partially supported by 3IA Artificial and Natural Intelligence Toulouse Institute (ANITI), French "Investing for the Future - PIA3" program under the Grant agreement ANR-19-PI3A-0004"

†Université de Toulouse, INP, IRIT, Toulouse, France, and NAXYS, University of Namur, Namur, Belgium. Email: philippe.toint@unamur.be

The whole operation is asynchronous because each process/processor does not wait for other processes to complete their tasks (it is *lock-free*). Of course, this description is extremely synthetic and methods of the type just described may still differ on a number of issues.

**Geometry.** The first is to define how a particular block of variables is selected in Step 1. This selection can be random or deterministic, and is characterized by the size of the allowed “overlap” between blocks selected by different processes, ranging from totally separable (the blocks associated with running processes are disjoint), partially separable (limited overlap between running blocks is allowed — for an introduction to partial separability, see [19, 20, 11]), or arbitrary (the size of the overlap is not constrained).

**Atomicity.** The next issue is the management of the read/write operations between local processes and shared memory. Because of the asynchronous nature of the whole algorithm, it may happen that different processes attempt to read or write the same component of  $\hat{X}$  at the same time. To cope with this problem, it is generally assumed that the reading and writing operations are *atomic* at some granularity level, meaning that reading/writing an atomic component of  $\hat{X}$  by a process cannot be interrupted by another process. We will see below that this concept also applies to other quantities than  $\hat{X}$ , and that, in addition to purely hardware and message passing considerations, it also can be limited by the internal structure of the object to read/write, or by the norm used to measure it (for instance, objects like positive definite matrices need to maintain their positive definite nature).

**Consistency.** In addition, it may well happen that a process is reading parts of  $\hat{X}$  while other parts are being updated by other processes. This blurs the notion of algorithm’s iterates, because it is possible that a block gradient is computed for a value of the variables corresponding to an incomplete, ongoing update of  $\hat{X}$ . Most algorithms (but not all) allow these *inconsistent* read/writes.

**Delays.** In addition, the asynchrony of the whole algorithm implies that when a block gradient is computed, it uses the values of  $\hat{X}$  read from shared memory, but these values may result from writes (by other blocks) somewhat in the past. The time difference between the moment  $\hat{X}$  is updated and the moment where it is used for computing a block gradient is called a *delay*. The analysis of all algorithms considered here imposes some kind of upper bound on such delays, either explicitly or in expectation.

**Choice of stepsize.** Finally, once the block gradient is computed, yielding a descent direction, a stepsize along this direction must be defined. Algorithms in the literature either assume that the stepsize is a “suitably small” constant, or that it is given by a (user) predefined decreasing sequence converging to zero.

Having brushed the landscape, we now review chronologically significant proposals in the literature. The pioneering **Hogwild!** method proposed in [43] was instrumental in introducing lock-free variants of asynchronous steepest-descent. It applies to strongly convex functions with partially separable geometry, for which block gradients are computable. Its stochastic convergence analysis assumes that the gradient oracles are bounded and that the constant stepsize can be chosen small enough compared to the gradient’s Lipschitz constant. Inconsistent reads are allowed. This is not the case for the **AsySG-con** [33] which, as the end of its name suggests, insists on coherent reads, but applies to possibly nonconvex functions with full-size gradients without any separability requirement. Its analysis assumes bounded gradients and predefined diminishing stepsizes. The **ARock** algorithm [44] is of a slightly different type, as it is better viewed as a fixed point method, but it faces similar issues. It uses full gradient vectors and requires that stepsizes decrease to zero. Its analysis holds for convex functions with bounded gradients but without any specific separability. The **KroMagnon** algorithm [36] is similar to **Hogwild!** in that it is a gradient-based method, allows block gradient evaluations and requires a partially separable geometry. Establishing its rate of convergence needs convex functions but no bound on the gradients is necessary. The requirement on partially separable geometry is relaxed for the versions of **Hogwild!** studied in [42, 41] and boundedness of the gradients is no longer requested. Instead of assuming a small enough stepsize as in the original version, the convergence analysis requires a stepsize decreasing to zero. **ASAGA** [32]

|                   | Problem |                    | Algorithm's characteristics |        |            | Assumptions  |                   |                 |
|-------------------|---------|--------------------|-----------------------------|--------|------------|--------------|-------------------|-----------------|
|                   | convex  | separable geometry | method type                 | stoch. | cons. read | step size    | bounded gradients | bounded delay   |
| Hogwild! [43]     | strong  | partially          | gradient                    | yes    | no         | small        | yes               | yes             |
| AsySG-con [33]    | no      | no                 | gradient                    | yes    | yes        | $\searrow 0$ | no                | yes             |
| "SGD" [9]         | yes     | no                 | gradient                    | yes    | no         | $\searrow 0$ | yes               | in $\mathbb{E}$ |
| ARock [44]        | yes     | no                 | fixed point                 | yes    | no         | small        | yes               | yes             |
| KroMagnon [36]    | yes     | partially          | gradient                    | yes    | no         | small        | no                | yes             |
| Hogwild! [42, 41] | strong  | relaxed            | gradient                    | yes    | no         | $\searrow 0$ | no                | yes             |
| ASAGA [32]        | strong  | partially          | gradient                    | yes    | no         | small        | no                | yes             |
| FDG [55]          | no      | yes                | gradient                    | yes    | no         | small        | yes               | yes             |
| AsyFLEXA [8]      | no      | no                 | convex approx.              | no     | no         | small        | no                | yes             |
| PASSM [29, 30]    | no      | yes                | gradient+mom                | yes    | no         | $\searrow 0$ | in $\mathbb{E}$   | in $\mathbb{E}$ |
| DADAO [38]        | yes     | no                 | gradient                    | yes    | no         | small        | no                | yes             |
| this paper        | no      | yes                | prec. grad.                 | yes    | no         | adapt.       | no                | yes             |
| this paper        | no      | yes                | prec. grad.+mom             | yes    | no         | adapt.       | no                | yes             |

Table 1: A summary of the relevant problem class, algorithm's characteristics and conditions of analysis for stochastic asynchronous gradient methods. Brackets around items in the "shared info" columns indicate the items are not mandatory.

is another gradient-based method requiring, as is the case for the original Hogwild!, partially separable geometry, strongly convex functions and a sufficiently small stepsize. However, there is no need to assume bounded gradients. The method occasionally evaluates the full gradient (in shared memory) in order to reduce the variance of the estimates. The FDG algorithm of [55] partitions the variables into blocks and interestingly proposes to exploit the structure of backpropagation in neural networks to compute the gradients with respect to the variables of a block independently from those of another, in turn leading to an asynchronous delayed-gradient method. The analysis stops with a descent lemma requiring bounded gradients and knowledge of the gradient's Lipschitz constant. The AsyFLEXA [8] deterministic algorithm is again slightly different in that it relies on convex local models (instead of purely linear ones as is the case for most gradient-based methods) and applies to possibly nonconvex functions with block gradient evaluations without separability requirement. Once more, the analysis is carried out for small enough stepsize, but without assuming bounded gradients. The reference [30] discusses several variants of a stochastic gradient-based method with momentum, as seen in the context of continuous differential equations. We report here on PASSM, a variant for potentially nonconvex totally separable functions with block gradient evaluations, which appears to perform best. Bounded gradients are required for its analysis and its stepsizes follow a geometrically decreasing sequence. Interestingly, a lock is required for writing the updated variables in shared memory. One should also mention contributions of [39] and [14]. The first introduces the concept of "elastic consistency", formalizing the conditions on delays for non-adaptive gradient methods and showing that such conditions are necessary for convergence of this type of methods. The second introduces general recurrence occurring in many non-adaptive asynchronous stochastic gradient methods and uses them to sharpen existing convergence results, in particular by considering average delays rather than maximum ones. Finally, the proposal of [38] considers an asynchronous method for strictly convex functions with known Lipschitz constant, putting special emphasis on communication costs.

Table 1 on this page provides an easy-to-read summary of the above discussion, where column "stoch." indicates whether the method is stochastic, and "cons. read" indicates whether the method requires consistent reading of shared memory. In the "method type" column, the abbreviations "convex approx", "prec. grad." and "mom" stand for "convex approximation", "preconditioned gradient" and "momentum", respectively. Finally, the abbreviation "adapt." in the "stepsize" column refers to adaptive stepsize, the technique we will develop below.

All the algorithms we have reviewed so far exhibit two (in our view, significant) limitations. The first and most important one is that pure gradient steps can lead to very slow convergence on (even moderately) ill-conditioned problems. This issue has long been addressed by the use of steepest-descent steps (a norm-dependent concept, at variance with gradient steps) and (adaptive)

preconditioning (see [7, Section 9.4], for instance). While these techniques have successfully been applied for solving problem (1.1) in the synchronous case, for instance in methods like **AdaGrad** [12, 37, 50], **Adam** [27], **Shampoo** [22, 48] or **Muon** [26, 46, 54] or in the distributed case [34], they haven't yet (as far as we know) been considered for asynchronous optimization. The second limitation is that theoretical analysis of many of the published algorithms ignores the use of momentum, despite its widespread nature (see [34], for instance).

The main contribution of the present paper is to propose and analyze a (to the authors' knowledge, first) class of algorithms containing several *stochastic asynchronous adaptively preconditioned gradient methods*, possibly *using momentum*, for the nonconvex case. In particular, **AdaNorm**, full and diagonal **AdaGrad** as well as adaptive variants of **Shampoo** and **Muon** are covered. Like for most recent proposals, bounded gradients are not necessary for its analysis. The proposed algorithm significantly extends the ADPREC framework for (synchronous) stochastic first-order methods [18], which is itself based on the use of general dual norms and covers a large class of stochastic adaptive gradient-based methods. The present proposal takes advantage of the powerful generality of this approach, and is defined in the geometric setting of totally separable problems, as is the case for **PASSM**. While this might be seen as theoretically restrictive, it still covers the very important case of layer-wise preconditioning in the training of large neural networks (see [53, 16, 45, 2] for instance), hence justifying our interest. As a side-benefit, our proposal also extends the theory presented in [18] by allowing *inexact step normalization strategies*, a clear advantage in practice for methods like **Muon** where the normalization is performed using a possibly inexact Newton–Schulz iteration [28, 5, 3, 48, 34]. The proposed addition of momentum to asynchronous steepest-descent methods (the second line for our paper in Table 1) is also useful to enhance numerical performance, but it may come with a (mostly theoretical) cost. For the common momentum definition using the gradient, assuming that the step-size is sufficiently small may be necessary to prove convergence, but the resulting degraded convergence rate may however be improved by adopting a suitable strategy for the momentum parameter. Remarkably, no assumption on the step-size is required for the momentum-less variant, or for a definition of the momentum using the momentum itself.

The paper is organized as follows. The momentum-less asynchronous framework is presented in Section 2, where it is shown to be a special case of an easier-to-analyze delayed framework, whose rate of convergence is studied in Section 3. Two different variants of momentum are introduced and the resulting rates of convergence examined in Section 4. Numerical experiments illustrating the use of asynchronous block-wise adaptive **Muon/AdaGrad** are presented in Section 6. Finally, some conclusions and perspectives are discussed in Section 7 while Appendix 1 and Appendix 2 give the details of the convergence proofs for the methods of Sections 2 and 3, respectively.

**Notations:** In what follows,  $\|\cdot\|_E$  denotes the Euclidean norm. If  $\|\cdot\|$  is a norm on some space  $\mathcal{S}$  endowed with an inner product  $\langle \cdot, \cdot \rangle$ , its dual norm  $\|\cdot\|_D$  is defined by  $\|x\|_D = \max_{y \in \mathcal{S}} \langle y, x \rangle / \|y\|$ . The symbol  $I_n$  denotes the identity matrix of dimension  $n$  and  $0_n$  the zero vector of size  $n$ . If  $x$  is a vector,  $[x]_i$  denotes its  $i$ -th component and  $[x]_{\mathcal{I}}$  the set of its components with indices in  $\mathcal{I}$ . If  $\mathcal{S}$  is a set,  $|\mathcal{S}|$  denotes its cardinal.

## 2 ASADPREC: a stochastic asynchronous first-order framework

The basis of our approach is to assume that the problem variables come in disjoint *blocks* of homogeneous nature. In deep learning applications, this occurs for instance when some variables are layer weights and other thresholds or biases of activation nodes. The first may be scalar and the second matrices [53, 16]. In line with [18], we assume that there are  $L$  blocks of variables, and problem (1.1) is therefore cast in the Cartesian product space

$$\mathbb{R}^N = \prod_{\ell=1}^L \mathbb{R}^{n_\ell \times m_\ell} = \prod_{\ell=1}^L \mathbb{R}^{d_\ell}.$$

We denote by  $\mathcal{I}_\ell$  the set of indices of the variables in block  $\ell$ . For  $X \in \mathbb{R}^N$ ,  $[X]_{\mathcal{I}_\ell}$  is then the subset of variables in  $X$  belonging to block  $\ell$ . Following [18], we then specify, for each block, a

particular norm  $\|\cdot\|_\ell$  and its dual  $\|\cdot\|_{\ell,*}$ . The inner product on  $\mathbb{R}^N$  is then defined as

$$\langle U, V \rangle = \sum_{\ell=1}^L \langle U_\ell, V_\ell \rangle_F \quad \text{where} \quad \langle A, B \rangle_F = \text{tr}(A^T B)$$

and the norm on  $\mathbb{R}^N$  given by

$$\|V\|^2 = \sum_{\ell=1}^L \|V_\ell\|_\ell^2 \quad (2.1)$$

whose dual norm is

$$\|U\|_*^2 = \sum_{\ell=1}^L \|U_\ell\|_{\ell,*}^2. \quad (2.2)$$

Because of the structure of the variables' space, it is natural to avoid mixing variables of different blocks, and this therefore defines the geometry of the asynchronous minimization method we are going to construct. Moreover, as is shown by the practical success of methods like **Muon** for matrix variables, different minimization methods may be appropriate for different types of variables. For each block, we therefore define a collection of methods applicable to variables in the block<sup>1</sup>, whose general format is

$$\Pi = \Pi^- + \mathcal{U}(G)^2, \quad Z = \Pi^{-1/2}G, \quad X^+ = X - \eta\|Z\|_*\mathcal{N}(Z),$$

where  $\Pi^-$  is a preconditioner inherited from previous iterations,  $G$  is the gradient,  $\mathcal{U}(\cdot)$  is a preconditioner update and  $\mathcal{N}(\cdot)$  a normalization operator in the chosen primal norm. An adaptive version of **Muon** [54] for block  $\ell$  is, for example, given by setting  $\|\cdot\|_\ell$  to the spectral matrix norm,  $\|\cdot\|_*$  to its dual, the nuclear norm, defining  $\mathcal{N}_\ell(G) = UV^T$  where the singular value decomposition of  $G$  is given by  $U\Sigma V^T$ , and using the associated preconditioner update  $\mathcal{U}(G) = (\|G\|_*/\sqrt{d_\ell})I_{d_\ell}$  with  $\|G\|_* = \text{tr}(\Sigma)$ . Thus, as was already the case in [18], each *method* for block  $\ell$  is specified by a preconditioner's update operator  $\mathcal{U} : \mathbb{R}^{d_\ell} \rightarrow \mathbb{R}^{d_\ell}$  and a normalization operator  $\mathcal{N} : \mathbb{R}^{d_\ell} \rightarrow \mathbb{R}^{d_\ell}$ . The collection of methods applicable for block  $\ell$  is then given by  $\mathcal{M}_\ell = \{(\mathcal{U}_1, \mathcal{N}_1), \dots, (\mathcal{U}_{\hat{m}_\ell}, \mathcal{N}_{\hat{m}_\ell})\}$ . While this way to describe first-order adaptively preconditioned gradient methods may seem abstract, it was shown in [18] that it covers a fairly representative set of popular methods, such as **AdaNorm**, full and diagonal **Adagrad**, **Shampoo** and **Muon** (see [18, Section 4.5] for details).

We now specify the **ASADPREC**<sup>2</sup> algorithmic *framework* on the next page, meaning that its description can be instantiated in a number of different specific algorithms depending on the choice of method for each block at each iteration. In what follows, we use the shorthand “**ASADPREC** algorithm” to refer to any algorithm belonging to the **ASADPREC** framework.

In this description, we assume that  $P$  processors are available and that all blocks have been *a priori* assigned to a unique processor, or, equivalently, that a processor  $p \in \{1, \dots, P\}$  can only work on a block  $\ell$  in its private block list  $\mathcal{L}_p$ , where  $\{\mathcal{L}_j\}_{j=1}^P$  is a partition of  $\{1, \dots, L\}$ . Because of this assumption, the preconditioners may be stored in local memory, avoiding read/write operations with shared storage. The detailed algorithm is presented on the following page.

<sup>1</sup>For instance, **AdaNorm** [50] and **Adagrad** [12, 37] are applicable methods for blocks with scalar variables. **Shampoo** [22] and **Muon** [26] are applicable for blocks with matrix variables.

<sup>2</sup>ASynchronous ADaptive PREConditioned gradient

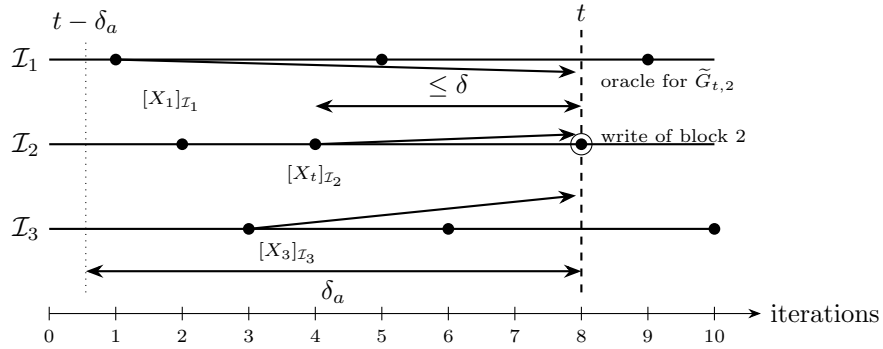
**Algorithm 2.1:** ASADPREC

 Given: a starting state  $\hat{X}_0$ , constants  $\eta, \varsigma > 0$  and the partition  $\mathcal{L}_1, \dots, \mathcal{L}_P$ .

 For  $\ell \in \{1, \dots, L\}$ , set  $\hat{\Pi}_\ell = \varsigma I_\ell$ .

 For each processor  $p \in \{1, \dots, P\}$  in parallel, continuously

1. select a block  $\ell \in \mathcal{L}_p$  and a method  $(\mathcal{U}_\ell, \mathcal{N}_\ell) \in \mathcal{M}_\ell$ .
2. read all blocks of  $\hat{X}$  block-atomically and form  $X_\ell^{\text{read}} \in \mathbb{R}^N$ , (inconsistent)
3. draw  $\xi_\ell$  and compute  $\tilde{G}_\ell = [\nabla_X f(X_\ell^{\text{read}}, \xi_\ell)]_{\mathcal{I}_\ell}$ ,
4.  $\Pi_\ell = \Pi_\ell^- + \mathcal{U}_\ell(\tilde{G}_\ell)^2$ ,
5.  $Z_\ell = \Pi_\ell^{-1/2} \tilde{G}_\ell$ ,
6. update  $[\hat{X}]_{\mathcal{I}_\ell} \leftarrow [\hat{X}]_{\mathcal{I}_\ell} - \eta \|Z_\ell\|_{*,\ell} \mathcal{N}_\ell(Z_\ell)$ . (block-atomic)



Oracle input at iteration  $t$ :  $\bar{X}_{\pi(t,2)} = ([X_1]_{\mathcal{I}_1}, [X_t]_{\mathcal{I}_2}, [X_3]_{\mathcal{I}_3})$ . The oracle evaluates the block gradient component  $[\nabla_X F(\bar{X}_{\pi(t,2)})]_{\mathcal{I}_2}$ , up to stochastic error, to produce  $\tilde{G}_{t,2}$ .

$\delta$ : maximal delay between two writes of the same block.

$\delta_a$ : maximal width of the reading window used to form the composite vector.

Some comments on this framework are useful at this stage.

1. The vector  $X_\ell^{\text{read}} \in \mathbb{R}^N$  denotes the full vector formed in Step 2 by block-atomic reads of  $\hat{X}$  by processor  $p$  for the update published in Step 6. It may be inconsistent across blocks in the sense that it needs not coincide with  $\hat{X}$  at any time.
2. Our assumption that blocks are assigned to fixed processors also implies that no two different processors can work on the same block. As a consequence,  $[\hat{X}]_{\mathcal{I}_\ell}$  is unmodified during the iteration on block  $\ell$  and, in particular,  $[X_\ell^{\text{read}}]_{\mathcal{I}_\ell} = [\hat{X}]_{\mathcal{I}_\ell}$  after Step 2. Should the assumption be abandoned, then a locking mechanism must be set up to guarantee this property and the preconditioners may need to be stored in and read from shared memory.
3. The computation of the block's gradient  $\tilde{G}_\ell$  in Step 3 may or may not require the computation of the complete gradient  $\nabla f(X, \xi_\ell)$ , depending on the particular form of the objective function and the computer architecture.
4. No writing conflict can occur in Step 6 because we have assumed a totally separable geometry, with no overlap between blocks.
5. The general framework of ASADPREC is reminiscent of the block-Jacobi [4, 15] setting for solving linear systems of equations.

Our analysis of ASADPREC's global rate of convergence is based on the view that, from an analytical standpoint, it can be viewed as a synchronous algorithmic framework with delays. To establish

this interpretation, let  $\widehat{X}$  denote the shared-memory vector manipulated by the asynchronous implementation, with initial value  $\widehat{X}_0$ . For the purpose of the analysis, we index iterations by block-atomic publications:  $X_t$  is the value of  $\widehat{X}$  immediately before the  $t$ -th block-atomic update in Step 6, and  $X_{t+1}$  is the value of  $\widehat{X}$  immediately after this update. If  $\ell$  is the block updated at Step 6 of iteration  $t$ , then

$$[X_{t+1}]_{\mathcal{I}_\ell} = [X_t]_{\mathcal{I}_\ell} - \eta \|Z_{t,\ell}\|_{*,\ell} \mathcal{N}_{t,\ell}(Z_{t,\ell}),$$

while, for every  $j \neq \ell_t$ ,

$$[X_{t+1}]_{\mathcal{I}_j} = [X_t]_{\mathcal{I}_j}.$$

In this iteration-based interpretation, we also identify  $\widetilde{G}_t$ , the approximate gradient at  $X_t$ , to

$$\widetilde{G}_t = \nabla_X f(X_{t,\ell}^{\text{read}}, \xi_{t,\ell}) \quad \text{and set} \quad \widetilde{G}_{t,\ell} = [\nabla_X f(X_{t,\ell}^{\text{read}}, \xi_{t,\ell})]_{\mathcal{I}_\ell}.$$

The computation of this block step was started at some time  $\pi(t, \ell)$  (where the nonnegative integer  $t - \pi(t, \ell)$  is called a *delay*) using  $[X_{\pi(t,\ell)}^{\text{read}}]_{\mathcal{I}_\ell}$  and the approximate block gradient  $\widetilde{G}_{\pi(t,\ell),\ell}$ . We may then (formally) consider that this step on block  $\ell$  was started at time  $t$  from  $X_t$ , but using the “delayed approximate block gradient”  $\widetilde{G}_{\pi(t,\ell),\ell}$ . This is exactly how the DADPREC algorithm described on the current page proceeds, at each iteration selecting a subset of blocks for which a complete step is computed using an approximate, *possibly delayed*, block gradient. We may thus see the sequence of iterates  $X_t$  generated by an ASADPREC algorithm as being instead generated by a DADPREC algorithm, where we also identify  $\Pi_\ell$  and  $Z_\ell$  in the former at time  $t$  with  $\Pi_{t,\ell}$  and  $Z_{t,\ell}$  in the latter.

**Algorithm 2.2:** DADPREC

Given: a starting point  $X_0$  and constants  $\eta, \varsigma > 0$ . Set  $\Pi_{-1,\ell} = \varsigma I_\ell$  for  $\ell \in \{1, \dots, L\}$ .

For  $t = 0, 1, \dots$

1. Draw  $\xi_t$  and compute an approximation  $\widetilde{G}_t$  of the gradient at  $X_t$ .
2. Select a nonempty block set  $\mathcal{A}_t \subseteq \{1, \dots, L\}$ .
3. For  $\ell \in \mathcal{A}_t$ , select  $(\mathcal{U}_{t,\ell}, \mathcal{N}_{t,\ell}) \in \mathcal{M}_\ell$ , define  $\widetilde{G}_{t,\ell} = [\widetilde{G}_t]_{\mathcal{I}_\ell}$  and compute

$$\Pi_{t,\ell} = \Pi_{t-1,\ell} + \mathcal{U}_{t,\ell}(\widetilde{G}_{t,\ell})^2, \tag{2.3}$$

$$Z_{t,\ell} = \Pi_{t,\ell}^{-1/2} \widetilde{G}_{t,\ell}, \tag{2.4}$$

$$X_{t+1,\ell} = X_{t,\ell} - \eta \|Z_{t,\ell}\|_{*,\ell} \mathcal{N}_{t,\ell}(Z_{t,\ell}). \tag{2.5}$$

4. For  $\ell \notin \mathcal{A}_t$ , (formally) set

$$\Pi_{t,\ell} = \Pi_{t-1,\ell}, \tag{2.6}$$

$$Z_{t,\ell} = 0_{n_\ell \times m_\ell} \tag{2.7}$$

$$X_{t+1,\ell} = X_{t,\ell} \tag{2.8}$$

Step 4 of DADPREC is purely formal, because it does not involve any computation. Note also that  $\widetilde{G}_{t,\ell}$  in DADPREC is a fairly general random approximation of the true block gradient, as it may include delays and other errors, such as resulting from sampling. Thus DADPREC generates a sequence of random approximate gradients  $\{\widetilde{G}_t = (\widetilde{G}_{t,1}^T, \dots, \widetilde{G}_{t,L}^T)^T\}$  and a resulting random process  $(\widetilde{G}_t, X_t, \{\Pi_{t,1}, \dots, \Pi_{t,L}\}, Z_t)$ , for which we define the conditional expectation at  $t$  as  $\mathbb{E}_t[\cdot] = \mathbb{E}[\cdot \mid \text{knowing}\{\widetilde{G}_0, \dots, \widetilde{G}_{t-1}\}]$ .

The reason for introducing the DADPREC framework is that its convergence analysis can be derived, with reasonable effort, from that of the proposal in [18]. We describe the necessary

assumptions and the results obtained (for DADPREC and thus for ASADPREC) in the next section, and postpone to the Appendix the detailed description of the modifications of the analysis of [18].

All assumptions and results will thus be formulated using the notation of the DADPREC framework, but may directly be interpreted in the context of ASADPREC.

As is clear from the algorithms' statements, this analysis does not include momentum. Variants including momentum will be considered in Section 4.

### 3 Convergence of the momentum-less DADPREC

Our analysis of the DADPREC framework starts by formulating requirements on the problem (1.1).

**Assumption 1 (Boundedness)** *There exists a constant  $f_{\text{low}}$  such that  $F(X) \geq f_{\text{low}}$  for all  $X \in \mathbb{R}^N$ .*

Defining  $G(X) = \nabla_X F(X)$ , we then require the following smoothness assumption.

**Assumption 2 (Smoothness)** *The objective function  $f$  is continuously differentiable and has a Lipschitz continuous gradient, that is there exists a constant  $L_G \geq 0$  such that, for all  $X, Y \in \mathbb{R}^N$ ,  $\|G(X) - G(Y)\|_* \leq L_G \|X - Y\|$ .*

We now introduce conditions that define the class of methods  $(\mathcal{U}_\ell, \mathcal{N}_\ell)$  for which we develop our theory. These are defined for a given block  $\ell \in \{1, \dots, L\}$  and depend on the choice of dual norm  $\|\cdot\|_{\ell,*}$  for this block. Again, we stress that *these abstract conditions do hold for several popular methods* for specific choices of the methods  $(\mathcal{U}, \mathcal{N}) \in \mathcal{M}_\ell$ , including **AdaNorm**, full and diagonal **Adagrad**, as well as adaptive variants of **Shampoo** and **Muon** [18, Section 4.5].

**Assumption 3 (Structural identities)** *For each  $t \geq 0$  and  $\ell \in \mathcal{A}_t$ , we have that, for all  $(\mathcal{U}, \mathcal{N}) \in \mathcal{M}_\ell$ ,*

$$\|Z_{t,\ell}\|_{*,\ell} \left\langle \tilde{G}_{t,\ell}, \mathcal{N}(Z_{t,\ell}) \right\rangle_F = \text{tr}(\Pi_{t,\ell}^{-1/2} \mathcal{U}(\tilde{G}_{t,\ell})^2), \quad (3.1)$$

and

$$\|Z_{t,\ell}\|_{*,\ell}^2 = \text{tr}(\Pi_{t,\ell}^{-1} \mathcal{U}(\tilde{G}_{t,\ell})^2). \quad (3.2)$$

**Assumption 4 (Gradient-preconditioner compatibility)** *There exists a constant  $\kappa_o > 0$  such that, for all  $\mathcal{U}$  in  $(\mathcal{U}, \mathcal{N}) \in \mathcal{M}_\ell$  and all  $G$ ,*

$$\|G\|_{*,\ell}^2 \leq \kappa_o^2 \text{tr}(\mathcal{U}(G)^2). \quad (3.3)$$

We also need to consider stochastic conditions on the quality of the gradient oracle.

**Assumption 5 (Cumulative variance)** *There exists a constant  $\omega \geq 0$  and a sequence  $\{\nu_t\}_{t \geq 0}$ , such that, for all  $t \geq 0$ ,*

$$\sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2] \leq \nu_t^2 + \omega^2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2]. \quad (3.4)$$

Now define

$$G_t \stackrel{\text{def}}{=} \nabla_X F(X_t) \quad \text{and} \quad G_{t,\ell} \stackrel{\text{def}}{=} [\nabla_X F(X_t)]_{\mathcal{I}_\ell}.$$

**Assumption 6 (Block unbiased)** *The block gradient oracle is unbiased, that is*

$$\mathbb{E}_t[\tilde{G}_{t,\ell}] = G_{t,\ell} \quad (3.5)$$

for all  $t \geq 0$  and all  $\ell \in \mathcal{A}_t$ .



Observe that Assumption 6 is made at the block level, and is weaker than assuming the more standard condition that  $\mathbb{E}_t[\tilde{G}_t] = G_t = G(X_t)$ . Similarly, Assumption 5 is also made at the block level. Moreover (3.4) requires a bound on the *cumulative* variance of all block gradient oracles over all past iterates, an approach also more general than assuming conditional variance at every iteration. In particular, it allows large variance at early iterations provided later iterations compensate. Trade-offs between different realizations and/or different blocks are also theoretically possible.

We finally impose a minimal condition on the regularity of the algorithm's choice of blocks.

**Assumption 7 (Bounded delays)** *There exists an integer  $\delta \geq 1$  such that, for every  $t > 0$  and every block  $\ell \notin \mathcal{A}_t$ , the last iteration before  $t$  at which block  $\ell$  was selected exists and satisfies*

$$\gamma(t, \ell) \stackrel{\text{def}}{=} \max\{i \in \{0, \dots, t-1\} \mid \ell \in \mathcal{A}_i\} \quad \text{and} \quad t - \gamma(t, \ell) \leq \delta.$$

In order to keep notation as simple as possible, we also assume, without loss of generality, that  $\mathcal{A}_0 = \{1, \dots, L\}$ , so that each block has been updated at least once before iteration  $t$  for  $t > 0$ . For the purpose of analysis, we also define, for  $t \geq 0$  and  $\ell \in \{1, \dots, L\}$ ,

$$\mathcal{J}_{t,\ell} = \{j \in \{0, \dots, t\} \mid \ell \in \mathcal{A}_j\},$$

the set of indices of iterations at which block  $\ell$  is effectively updated, and immediately note that  $|\mathcal{J}_{t,\ell}| \geq 1$  for all  $t \geq 0$  and all  $\ell \in \{1, \dots, L\}$ , and also that the summations  $\sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j}$  and  $\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}}$  are interchangeable (for instance in (3.4))

We now consider a useful sufficient condition ensuring Assumptions 5 and 7. In what follows,  $\bar{X}_{\pi(t,\ell)} \in \mathbb{R}^N$  denotes the possibly delayed global iterate read by the oracle producing  $\tilde{G}_{t,\ell}$ . For the ASADPREC algorithm, this vector coincides with the read iterate of Section 2, namely  $\bar{X}_{\pi(t,\ell)} = X_{\pi(t,\ell)}^{\text{read}}$ . The index  $\pi(t,\ell)$  is allowed to be delayed, but the delay is uniformly bounded by  $\delta_a$ , that is,  $t - \pi(t,\ell) \leq \delta_a$ . This is formalized by the following lemma.

**Lemma 3.1** Suppose that Assumption 2 holds. Suppose also that there exists a nonnegative integer  $\delta_a$  such that, for each  $(t, \ell)$ , there exist an iteration index  $\pi(t, \ell)$  and a vector  $\bar{X}_{\pi(t, \ell)} \in \mathbb{R}^N$  satisfying

$$0 \leq \pi(t, \ell) \leq t \quad \text{and} \quad t - \pi(t, \ell) \leq \delta_a, \quad (3.6)$$

and  $\pi(t, \ell)$  is the smallest index such that

$$[\bar{X}_{\pi(t, \ell)}]_{\mathcal{I}_\ell} = [X_t]_{\mathcal{I}_\ell} \quad \text{and} \quad \forall q \neq \ell, \exists s_{t, \ell, q} \in \{\pi(t, \ell), \dots, t\} \text{ with } [\bar{X}_{\pi(t, \ell)}]_{\mathcal{I}_q} = [X_{s_{t, \ell, q}}]_{\mathcal{I}_q}. \quad (3.7)$$

Then, Assumption 7 holds with  $\delta = \delta_a$ . If, in addition,

$$\mathbb{E}_{\pi(t, \ell)} \left[ \|\tilde{G}_{t, \ell} - \bar{G}_{\pi(t, \ell), \ell}\|_{*, \ell}^2 \right] \leq \tau_{\pi(t, \ell)}^2 + \omega_\pi^2 \sum_{r=0}^{\pi(t, \ell)} \sum_{q \in \mathcal{A}_r} \mathbb{E}[\|Z_{r, q}\|_{*, q}^2], \quad (3.8)$$

where

$$\bar{G}_{\pi(t, \ell)} = \nabla_X F(\bar{X}_{\pi(t, \ell)}) \quad \text{and} \quad \bar{G}_{\pi(t, \ell), \ell} = [\nabla_X F(\bar{X}_{\pi(t, \ell)})]_{\mathcal{I}_\ell}.$$

then, for all  $t \geq 0$ ,

$$\sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\tilde{G}_{j, \ell} - G_{j, \ell}\|_{*, \ell}^2] \leq 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \tau_{\pi(j, \ell)}^2 + 2(\omega_\pi^2 + \delta_a^2 L_G^2 \eta^2 L) \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|Z_{j, \ell}\|_{*, \ell}^2]. \quad (3.9)$$

In particular, Assumption 5 holds with

$$\nu_t^2 = 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \tau_{\pi(j, \ell)}^2 \quad \text{and} \quad \omega^2 = 2(\omega_\pi^2 + \delta_a^2 L_G^2 \eta^2 L).$$

**Proof.** The bounds (3.6) and the first part of (3.7) ensures that block  $\ell$  is activated at least every  $\delta_a$  iterations, yielding Assumption 7. Applying Assumption 2, we obtain, for every  $j \geq 0$  and every  $\ell \in \mathcal{A}_j$ ,

$$\|\bar{G}_{\pi(j, \ell)} - G_j\|_*^2 \leq L_G^2 \|\bar{X}_{\pi(j, \ell)} - X_j\|^2. \quad (3.10)$$

It remains to bound the distance between the read vector  $\bar{X}_{\pi(j, \ell)}$  and the current iterate  $X_j$ . By the bounded-delay condition (3.6), the vector  $\bar{X}_{\pi(j, \ell)}$  can differ from  $X_j$  only through the block publications that occurred between iterations  $\pi(j, \ell)$  and  $j - 1$ . Hence,

$$\bar{X}_{\pi(j, \ell)} - X_j = \sum_{i=\pi(j, \ell)}^{j-1} (X_i - X_{i+1}).$$

Taking norms and using the update formula (2.5) gives

$$\|\bar{X}_{\pi(j, \ell)} - X_j\|^2 \leq \left( \sum_{i=\pi(j, \ell)}^{j-1} \|X_{i+1} - X_i\| \right)^2 = \eta^2 \left( \sum_{i=\pi(j, \ell)}^{j-1} \|Z_i\|_* \right)^2.$$

If  $\pi(j, \ell) < j$ , then, since  $j - \pi(j, \ell) \leq \delta_a$ , the Cauchy's inequality gives that

$$\left( \sum_{i=\pi(j, \ell)}^{j-1} \|Z_i\|_* \right)^2 \leq (j - \pi(j, \ell)) \sum_{i=\pi(j, \ell)}^{j-1} \|Z_i\|_*^2 \leq \delta_a \sum_{i=\pi(j, \ell)}^{j-1} \|Z_i\|_*^2.$$

The inequality is trivial when  $\pi(j, \ell) = j$ . Thus,

$$\|\bar{X}_{\pi(j, \ell)} - X_j\|^2 \leq \delta_a \eta^2 \sum_{i=\pi(j, \ell)}^{j-1} \|Z_i\|_*^2,$$

and consequently, from (3.10),

$$\|\bar{G}_{\pi(j,\ell)} - G_j\|_*^2 \leq \delta_a L_G^2 \eta^2 \sum_{i=\pi(j,\ell)}^{j-1} \|Z_i\|_*^2.$$

We now sum this inequality over  $j$  and over the active blocks  $\ell \in \mathcal{A}_j$ . For a fixed index  $r$ , the term  $\|Z_r\|_*^2$  appears in the inner sum  $\sum_{i=\pi(j,\ell)}^{j-1} \|Z_i\|_*^2$  only if  $\pi(j,\ell) \leq r \leq j-1$ . Since (3.6) gives  $j - \pi(j,\ell) \leq \delta_a$ , this implies  $r+1 \leq j \leq r + \delta_a$ . Thus, for fixed  $r$ , the term  $\|Z_r\|_*^2$  can be counted for at most  $\delta_a$  different values of  $j$  for each active block. Hence, using  $|\mathcal{A}_j| \leq L$  and (2.2),

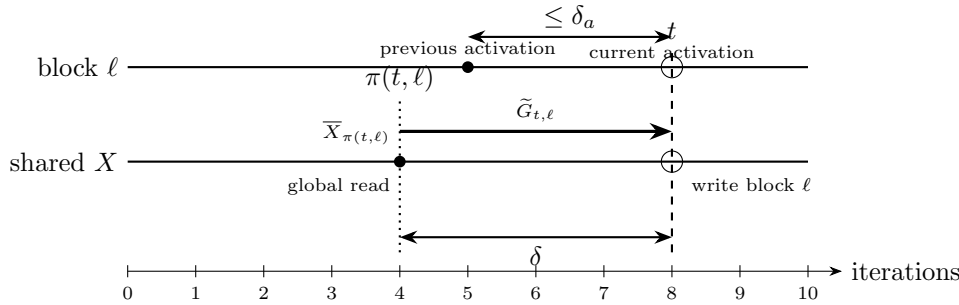
$$\begin{aligned} \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\bar{G}_{\pi(j,\ell),\ell} - G_{j,\ell}\|_{*,\ell}^2] &\leq \delta_a L_G^2 \eta^2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \sum_{i=\pi(j,\ell)}^{j-1} \mathbb{E}[\|Z_i\|_*^2] \\ &\leq \delta_a^2 L_G^2 \eta^2 L \sum_{i=0}^t \sum_{q \in \mathcal{A}_i} \mathbb{E}[\|Z_{i,q}\|_{*,q}^2]. \end{aligned}$$

Combining this bound with (3.8) and using  $\|a+b\|^2 \leq 2\|a\|^2 + 2\|b\|^2$ , we get

$$\begin{aligned} \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2] &\leq 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\tilde{G}_{j,\ell} - \bar{G}_{\pi(j,\ell),\ell}\|_{*,\ell}^2] + 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|\bar{G}_{\pi(j,\ell),\ell} - G_{j,\ell}\|_{*,\ell}^2] \\ &\leq 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \tau_{\pi(j,\ell)}^2 + 2(\omega_\pi^2 + \delta_a^2 L_G^2 \eta^2 L) \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2]. \end{aligned}$$

Thus (3.9) follows with  $\nu_t^2 = 2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \tau_{\pi(j,\ell)}^2$  and  $\omega^2 = 2(\omega_\pi^2 + \delta_a^2 L_G^2 \eta^2 L)$ .  $\square$

The following figure illustrates the various quantities occurring in the lemma (the top line describes the ASADPREC course of events, iterations refer to DADPREC iterations).



The oracle uses the delayed global read iterate  $\bar{X}_{\pi(t,\ell)} = X_{\pi(t,\ell)}^{\text{read}}$  to form  $\tilde{G}_{t,\ell}$ , which is used to update block  $\ell$  at iteration  $t$ .

$\delta_a$ : read delay between the global read and the current write (ASADPREC)

$\delta$ : activation delay; block  $\ell$  is selected again within at most  $\delta$  iterations (DADPREC).

As a consequence of Lemma 3.1, not only the ASADPREC algorithm can be seen as a special case of DADPREC, but the theory developed for the latter (under Assumptions 5 and 7) is applicable to the former provided condition (3.6) (bounded delays) holds.

A general theorem stating the global rate of convergence of any DADPREC (and thus ASADPREC) algorithm may now be stated under the above assumptions.

**Theorem 3.2** Suppose that the DADPREC algorithm is applied to problem (1.1) and that Assumptions 1–5 and 7 hold. Then, for all  $t \geq 0$ ,

$$\min_{j \in \{0, \dots, t\}} \mathbb{E}[\|G_j\|_*] \leq \frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] \leq \frac{\kappa_a + \kappa_b \sqrt{\log(\Theta_t)} + \kappa_c \Theta_t}{\sqrt{t+1}} \quad (3.11)$$

with

$$\kappa_a = \delta L_G \eta \sqrt{L(L-1)\kappa_0 + (\delta+1)\sqrt{L}} (\nu_t + \omega \sqrt{\kappa_0}), \quad \kappa_b = \delta L_G \eta \sqrt{2NL(L-1) + (\delta+1)\omega \sqrt{2NL}}, \quad (3.12)$$

$$\kappa_c = (\delta+1)\kappa_o \sqrt{L}. \quad (3.13)$$

$$\Theta_t \stackrel{\text{def}}{=} \max \left[ \kappa_\Theta, 12\sqrt{N} \nu_t \sqrt{\max \left[ 1, \log \left( 12\sqrt{N} \nu_t \right) \right]} \right] \quad (3.14)$$

with

$$\kappa_\Theta = \max \left[ e^{\max[1, \frac{\kappa_o}{2N}]}, \frac{3(\mathbb{E}[f(X_0)] - f_{\text{low}} + \eta \varsigma N)}{\eta}, 24N \left( \omega + \frac{L_G \eta}{2} \right) \log \left( 24N \left( \omega + \frac{L_G \eta}{2} \right) \right) \right]$$

and  $\kappa_0 = \max[-\sum_{\ell=1}^L d_\ell \log(d_\ell) - N \log(\varsigma), 0]$ . Moreover, if Assumption 6 holds, then

$$\min_{j \in \{0, \dots, t\}} \mathbb{E}[\|G_j\|_*] \leq \frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] \leq \frac{\delta L_G \eta \sqrt{L(L-1)} (\sqrt{\kappa_0} + \sqrt{2N \log(\Theta_t)}) + \kappa_c \Theta_t}{\sqrt{t+1}}. \quad (3.15)$$

**Proof.** See Appendix 1. □

Since we prove (in Appendix) that each preconditioner  $\Pi_{t,\ell}$ , is bounded by the (in expectation) very slowly growing  $\Theta_t$  (see (3.14)), and since  $Z_t$  is the preconditioned approximate gradient, condition (3.4) is akin to a undelayed cumulative affine variance condition of the form

$$\sum_{j=0}^s \sum_{\ell=1}^L \mathbb{E}_j [\|\tilde{G}_{j,\ell} - G_{j,\ell}\|^2] \leq \nu_t^2 + \omega^2 \sum_{j=0}^s \sum_{\ell=1}^L \mathbb{E}_j [\|G_{j,\ell}\|_{*,\ell}^2],$$

whose iteration-wise variant has already been used in analysis of first-order methods (see [13, 49] or, for the stronger “affine- $\ast$ ” version, [1]). In particular, the “strong growth” assumption motivated by over-parametrized problems and used in [49] to derive an improved convergence rate for AdaGrad is essentially subsumed (at the cumulative level) for “preconditioned cumulative strong growth” (that is if  $\nu_t = 0$  is assumed in (3.4)).

As stated, the rate of convergence of DADPREC is dominated by the term in  $\kappa_c \Theta_t / \sqrt{t+1}$  in (3.11), where  $\Theta_t$  depends on  $\nu_t$ , the bound on the cumulative variance of the block gradient oracles, which may itself depend on  $t$ . The next corollary describes what can be said if one makes a more specific assumption on the oracle (block) variance at each iteration. It uses  $|\mathcal{J}_{t,\ell}|$  to count the number of updates to block  $\ell$  up to iteration  $t$ .

**Corollary 3.3** Suppose that the DADPREC algorithm is applied to problem (1.1) Assumptions 1 to 5 and 7 hold. Suppose also that, for all  $t \geq 0$  and  $\ell \in \mathcal{A}_t$ ,

$$\mathbb{E}_t \left[ \|\tilde{G}_{t,\ell} - G_{t,\ell}\|_*^2 \right] \leq \frac{\sigma_\ell^2}{|\mathcal{J}_{t,\ell}|^\alpha} + \omega^2 \mathbb{E}_t [\|Z_{t,\ell}\|_*^2] \quad (3.16)$$

for some  $\sigma_\ell \geq 0$ ,  $\ell \in \mathcal{A}_t$  and  $\alpha, \omega > 0$ . Then, if  $\psi_t = \max_{\ell \in \{1, \dots, L\}} |\mathcal{J}_{t,\ell}|$ ,

$$\frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] = \begin{cases} \mathcal{O} \left( \frac{\psi_t^{\frac{1}{2}(1-\alpha)} \sqrt{\log(\psi_t)}}{\sqrt{t+1}} \right) & \text{if } \alpha < 1, \\ \mathcal{O} \left( \frac{\sqrt{\log(\psi_t) \log(\log(t+1))}}{\sqrt{t+1}} \right) & \text{if } \alpha = 1, \\ \mathcal{O} \left( \frac{1}{\sqrt{t+1}} \right) & \text{if } \alpha > 1. \end{cases} \quad (3.17)$$

**Proof.** See Appendix 1. □

We have that  $\psi_t \leq t+1$  in general, but if, at each iteration  $t$ ,  $|\mathcal{A}_t| \ll L$  (a not unusual situation), then  $\psi_t \ll t+1$ .

Observe that the continuity of the bound expressed in Theorem 3.2 as a function of  $\nu_t$  is lost in the statement of Corollary 3.3, because the constants hidden in the  $\mathcal{O}(\cdot)$  notation depend on  $\alpha$ . In particular, this formulation does not support taking the limit for  $\alpha$  tending to one. Observe also that we could weaken (3.16) by requiring that

$$\mathbb{E}_t \left[ \|\tilde{G}_{t,\ell} - G_{t,\ell}\|_*^2 \right] \leq \frac{\sigma_\ell^2}{|\mathcal{J}_{t,\ell}|^\alpha} + \omega^2 \|Z_{t-1,\ell}\|_*^2 \quad (3.18)$$

for  $\ell \in \mathcal{A}_t$  with the convention that  $Z_{-1,\ell} = 0$ , since, obviously,

$$\sum_{\ell=1}^L \sum_{j=0}^{t-1} \mathbb{E}_j [\|Z_{j,\ell}\|_*^2] \leq \sum_{\ell=1}^L \sum_{j=0}^t \mathbb{E}_j [\|Z_{j,\ell}\|_*^2].$$

The advantage of (3.18) over (3.16) is that the right-hand-side of this new condition is now measurable at iteration  $t$ .

As was noted in [18], although the rates of convergence obtained in Corollary 3.3 under the bias and variance Assumptions 6 and 5 are reasonable, they do not quite match, for high-variance regimes, the best obtained so far for (momentum-less, synchronous) AdaGrad and AdaNorm [49]. However they recover (in order) the best  $\mathcal{O}(\sqrt{\log(k+1) \log(\log(k+1))/(k+1)})$  rate in the preconditioned cumulative strong growth context<sup>3</sup>. Importantly, they do so for a large class of asynchronous algorithms.

## 4 Convergence of DADPREC with momentum

Adding momentum to first-order methods has often been instrumental in improving numerical performance<sup>4</sup>. It is therefore of interest to consider asynchronous versions of such methods. We investigate two versions of this idea in this section, which share the common framework ASADPREC.M stated below for some given sequences  $\{\mu_k\}$  of momentum parameters with  $0 \leq \mu_k \leq \mu_{\max} < 1$  for all  $k \geq 0$ , with  $k_\ell$  indexing the number of updates to block  $\ell$ .

<sup>3</sup>That is not only in the case where  $\sigma_\ell = 0$  for each  $\ell$ , but, more generally, in the case where the “inaccuracy budget”  $\nu_t^2$  is finite.

<sup>4</sup>Although, as far as the authors are aware, this has not so far been translated in improved convergence theory.

**Algorithm 4.1: ASADPREC.M**

 Given: a starting state  $\widehat{X}_0$ , constants  $\eta, \varsigma > 0$  and the partition  $\mathcal{L}_1, \dots, \mathcal{L}_P$ .

 For  $\ell \in \{1, \dots, L\}$ ,  $\widehat{\Pi}_{-1,\ell} = \varsigma I_\ell$  and  $k_\ell = 0$ .

 For each processor  $p \in \{1, \dots, P\}$  in parallel, continuously

1. select a block  $\ell \in \mathcal{L}_p$  a method  $(\mathcal{U}_\ell, \mathcal{N}_\ell) \in \mathcal{M}_\ell$ .
2. read all blocks of  $\widehat{X}$  block-atomically and form  $X_\ell^{\text{read}} \in \mathbb{R}^N$ , (inconsistent)
3. draw  $\xi_\ell$  and compute  $\widetilde{G}_\ell = [\nabla_X f(X_\ell^{\text{read}}, \xi_\ell)]_{\mathcal{I}_\ell}$ ,
4. compute  $M_\ell$  and  $\Pi_\ell$ , using  $\mathcal{U}_\ell$  and  $k_\ell$ .
5.  $Z_\ell = \Pi_\ell^{-1/2} M_\ell$ ,
6.  $k_\ell \leftarrow k_\ell + 1$ ,
7. update  $[\widehat{X}]_{\mathcal{I}_\ell} \leftarrow [\widehat{X}]_{\mathcal{I}_\ell} - \eta \|Z_\ell\|_{*,\ell} \mathcal{N}_\ell(Z_\ell)$  (block-atomic)

The two variants are specified by detailing Step 4 either as

$$\left. \begin{aligned} \Pi_\ell &= \Pi_\ell^- + \mathcal{U}_\ell(\widetilde{G}_\ell)^2, \\ M_\ell &= \begin{cases} \mu_{k_\ell} M_\ell^- + (1 - \mu_{k_\ell}) \widetilde{G}_\ell & \text{if } k_\ell > 0, \\ \widetilde{G}_\ell & \text{if } k_\ell = 0, \end{cases} \end{aligned} \right\} \text{(ASADPREC.MG)}$$

or as

$$\left. \begin{aligned} M_\ell &= \begin{cases} \mu_{k_\ell} M_\ell^- + (1 - \mu_{k_\ell}) \widetilde{G}_\ell & \text{if } k_\ell > 0, \\ \widetilde{G}_\ell & \text{if } k_\ell = 0 \end{cases} \\ \Pi_\ell &= \Pi_\ell^- + \mathcal{U}_\ell(M_\ell)^2. \end{aligned} \right\} \text{(ASADPREC.MM)}$$

 In both cases, the momentum  $\widehat{M}_\ell$  and the counter  $k_\ell$  may also be stored in local memory.

Fortunately, the convergence theory of Section 3 can be adapted to cover both variants. This is achieved again by considering DADPREC.MM and DADPREC.MG, the obvious extensions of DADPREC corresponding to ASADPREC.MM and ASADPREC.MG, respectively, where (2.3) and (2.4) are now replaced by

$$\left. \begin{aligned} M_{t,\ell} &= \mu_{t,\ell} M_{t-1,\ell} + (1 - \mu_{t,\ell}) \widetilde{G}_{t,\ell}, \\ \Pi_{t,\ell} &= \Pi_{t-1,\ell} + \mathcal{U}_{t,\ell}(M_{t,\ell})^2, \\ Z_{t,\ell} &= \Pi_{t,\ell}^{-1/2} M_{t,\ell} \end{aligned} \right| \begin{aligned} \Pi_{t,\ell} &= \Pi_{t-1,\ell} + \mathcal{U}_{t,\ell}(\widetilde{G}_{t,\ell})^2, \\ M_{t,\ell} &= \mu_{t,\ell} M_{t-1,\ell} + (1 - \mu_{t,\ell}) \widetilde{G}_{t,\ell}, \\ Z_{t,\ell} &= \Pi_{t,\ell}^{-1/2} M_{t,\ell} \end{aligned}$$

for DADPREC.MM, for DADPREC.MG.

 We may also identify  $k_\ell$ , the counter of updates to block  $\ell$  for ASADPREC.MM and ASADPREC.MG, with  $|\mathcal{J}_{t,\ell}|$  for DADPREC.MM and DADPREC.MG.

 Because the update formula for  $Z_{t,\ell}$  above now involves the momentum  $M_{t,\ell}$  instead of  $\widetilde{G}_{t,\ell}$ , we reformulate Assumption 3 to reflect this change<sup>5</sup> as follows.

**Assumption 8 (Structural identities with momentum)** *For each  $t \geq 0$  and  $\ell \in \mathcal{A}_t$ , we have that, for all  $(\mathcal{U}, \mathcal{N}) \in \mathcal{M}_\ell$ ,*

$$\|Z_{t,\ell}\|_{*,\ell} \langle M_{t,\ell}, \mathcal{N}(Z_{t,\ell}) \rangle_F = \text{tr}(\Pi_{t,\ell}^{-1/2} \mathcal{U}(M_{t,\ell})^2), \quad (4.1)$$

and

$$\|Z_{t,\ell}\|_{*,\ell}^2 = \text{tr}(\Pi_{t,\ell}^{-1} \mathcal{U}(M_{t,\ell})^2). \quad (4.2)$$

Armed with this reformulated assumption, it is now possible to establish that the expected global rate of convergence of a DADPREC.MM algorithm is essentially identical to that of the version without momentum.

<sup>5</sup>The reformulation obviously still holds for AdaNorm, AdaGrad, Shampoo and Muon [18].

**Theorem 4.1** Suppose that an DADPREC.MM algorithm is applied to problem (1.1) and that Assumptions 1, 2, and 4, 5, 8 and 7 hold. Then

$$\min_{j \in \{0, \dots, t\}} \mathbb{E}[\|G_j\|_*] \leq \frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] \leq \frac{\kappa_a + \kappa_b \sqrt{\log(\Theta_t)} + \kappa_c \Theta_t}{\sqrt{t+1}} \quad (4.3)$$

with  $\Theta_t$  given by (3.14) and

$$\kappa_a = (\delta + 1)\sqrt{L} \left[ \nu_t + (L_G \eta \sqrt{L-1} + \omega) \sqrt{\kappa_0} \right], \quad \kappa_b = (\delta + 1)\sqrt{L} (L_G \eta \sqrt{L-1} + \omega) \sqrt{2N}$$

and

$$\kappa_c = (\delta + 1)\sqrt{L} \kappa_0, \quad \kappa_0 = \max \left[ - \sum_{\ell=1}^L d_\ell \log(d_\ell) - N \log(\varsigma), 0 \right].$$

Moreover, the conclusions of Corollary 3.3 do hold

**Proof.** See Appendix 2. □

We note at this point that our momentum definition and convergence analysis for the DADPREC.MM algorithm do not make the assumption that the stepsize parameter  $\eta$  is sufficiently small, in contrast with previous proofs of convergence where results assume that  $\eta$  is small enough, in particular smaller than a multiple of the (usually unknown) Lipschitz constant  $L_G$  [21, 25, 52]. Establishing convergence results for the DADPREC.MG variant unfortunately requires this assumption. This alternative theory hinges on the fact that  $\mathcal{U}(\tilde{G}_{k,\ell})$  can be viewed as a perturbation of  $\mathcal{U}(M_{k,\ell})$ , and therefore that the structural relations of Assumption 8 are themselves perturbed when the DADPREC.MG is used. As it turns out, it remains possible to bound these perturbations by a mix of variance-related and second-order terms, the first of which potentially affecting the resulting global convergence rate. The final outcome is given by the following theorem and corollary.

**Theorem 4.2** Suppose that a DADPREC.MG algorithm is applied to problem (1.1). Suppose that Assumptions 1, 2, 4, 5, 7 and 8 hold. Suppose also that there exist constants  $\kappa_{\square}, \kappa_{\diamond} > 0$  such that

$$\mathcal{U}(\tilde{G}_a + \tilde{G}_b)^2 \preceq \kappa_{\square} \mathcal{U}(\tilde{G}_a)^2 + \kappa_{\square} \mathcal{U}(\tilde{G}_b)^2 \quad \text{and} \quad \text{tr}(\mathcal{U}(\tilde{G}_a)^2) \leq \kappa_{\diamond} \|\tilde{G}_a\|_{*,\ell}^2 \quad (4.4)$$

for all  $t \geq 0$  and  $\ell \in \mathcal{A}_t$  and all  $(\mathcal{U}, \mathcal{V}) \in \mathcal{M}_{\ell}$  and all  $\tilde{G}_a, \tilde{G}_b$  of suitable size, and that

$$\text{either } \eta \leq \frac{1 - \mu_{\max}}{L_G} \sqrt{\frac{\varsigma}{6\kappa_{\square}\kappa_{\diamond}\delta^2 L}} \quad \text{or} \quad \sum_{j=0}^{\infty} \|Z_j\|_*^2 \leq \kappa_Z \quad (4.5)$$

for some  $\kappa_Z \geq 0$ . Then

$$\min_{j \in \{0, \dots, t\}} \mathbb{E}[\|G_j\|_*] \leq \frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] \leq \frac{\kappa_a + \kappa_b \sqrt{\log(\Theta_t)} + \kappa_c \Theta_t}{\sqrt{t+1}} \quad (4.6)$$

where  $\kappa_a, \kappa_b$  and  $\kappa_c$  are of the form (3.12)–(3.13) using

$$\Theta_t = \max \left[ \kappa_{\Theta}, \frac{3\kappa_{\nu\nu}\theta_t^2}{\eta}, T_t \right], \quad (4.7)$$

with

$$\begin{aligned} \kappa_{\Theta} &= \max \left[ e^{\max[1, \frac{\kappa_0}{2N}]}, \frac{3\kappa_{\text{gap}}}{\eta}, 24N\kappa_{\Delta} \left( \omega^2 + \frac{L_G}{\eta} \right) \log \left( 24N\kappa_{\Delta} \left( \omega^2 + \frac{L_G}{\eta} \right) \right) \right], \\ \sigma(r, \ell) &= \min \{ i \in \mathcal{J}_{t,\ell} \mid i > r \}, \quad \bar{\mu}_{j,\ell} = \max \{ \mu_{j,\ell}, \mu_{\sigma(j,\ell),\ell} \}, \\ \theta_t^2 &= \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \bar{\mu}_{j,\ell}^2 \mathbb{E}[\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2], \quad T_t = 12\sqrt{N} \kappa_{\nu\Delta} \theta_t \sqrt{\max \left[ 1, \log \left( 12\sqrt{N} \kappa_{\nu\Delta} \theta_t \right) \right]}, \end{aligned} \quad (4.8)$$

$\kappa_{\text{gap}} = \mathbb{E}[f(X_0)] - f_{\text{low}} + \eta\varsigma N$ ,  $\kappa_{\circ}$  is defined in Assumption 4,  $\kappa_0$  in Theorem 3.2, and  $\kappa_{\nu\nu}$ ,  $\kappa_{\nu\Delta}$  and  $\kappa_{\Delta}$  are specified in (A.45)–(A.46). Moreover, if Assumption 6 holds, then (3.15) again holds.

**Proof.** See Appendix 2. □

Assuming (4.4) is not restrictive<sup>6</sup>, but supposing (4.5) is definitely less desirable because the value of the Lipschitz constant  $L_G$  or that of  $\kappa_Z$  (if it exists) is usually unknown. It is however common practice in the literature [21, 25, 52]. Note the term  $\kappa_{\nu\nu}\nu_t^2$  in the middle term of the expression of  $\Theta_t$  in (4.7), which is the first significant difference between (4.7) and (3.14) and induces a modified convergence rate. The second difference is the presence of the factor  $\bar{\mu}_{j,\ell}^2$  in the definition of the cumulative variance, which has the opposite effect of potentially improving the convergence rate. This is expressed by the following corollary.

<sup>6</sup>AdaNorm, AdaGrad, Shampoo and Muon continue to be covered, as verified in [18].



**Corollary 4.3** Suppose that a DADPREC.MG algorithm is applied to problem (1.1). Suppose that Assumptions 1, 2, 4, 7 and 8 hold. Suppose also that (4.4), (4.5) and (3.16) hold, and that, for each  $\ell \in \mathcal{A}_t$ ,

$$\mu_{t,\ell} \in \left[0, \frac{\mu_{\max}}{|\mathcal{J}_{t,\ell}|^\beta}\right] \quad (4.9)$$

for some  $\beta \geq 0$ . Then, if  $\psi_t = \max_{\ell \in \{1, \dots, L\}} |\mathcal{J}_{t,\ell}|$  and  $\alpha$  is given by (3.16),

$$\frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] = \begin{cases} \mathcal{O}\left(\frac{\psi_t^{1-\alpha-2\beta}}{\sqrt{t+1}}\right) & \text{if } \alpha + 2\beta < 1, \\ \mathcal{O}\left(\frac{\log(\psi_t)}{\sqrt{t+1}}\right) & \text{if } \alpha + 2\beta = 1, \\ \mathcal{O}\left(\frac{1}{\sqrt{t+1}}\right) & \text{if } \alpha + 2\beta > 1. \end{cases} \quad (4.10)$$

**Proof.** See Appendix 2. □

The rate of convergence now depends on the value of  $\alpha + 2\beta$ , indicating how *acting on the momentum parameter can partly alleviate the effect of high variance of the gradient oracle*. Indeed, choosing a small momentum parameter in effect reduces the propagation of larger errors in the gradient oracle across iterations. In particular, setting  $\mu_{t,\ell}$  to a multiple of  $|\mathcal{J}_{t,\ell}|^{-\frac{1}{4}(1-\alpha)}$  recovers the rate of the momentum-less variant also for the high-variance regime ( $\alpha < 1$ ). If however  $\mu_{t,\ell}$  is kept constant (or bounded away from zero), that is if  $\beta = 0$ , then the rate of convergence for  $\alpha < 1$  is now  $\mathcal{O}\left(\psi_t^{1-\alpha}/\sqrt{t+1}\right)$  instead of  $\mathcal{O}\left(\psi_t^{\frac{1}{2}(1-\alpha)}/\sqrt{t+1}\right)$ . Again, the constants hidden in the  $\mathcal{O}(\cdot)$  depend on  $\alpha$  and  $\beta$ , in particular preventing taking limits for  $\alpha + 2\beta$  tending to one.

## 5 Approximate normalization

We have assumed, in our theory, that normalization of the step is exact, in the sense that  $\|\mathcal{N}_{t,\ell}(Z)\|_\ell = 1$  for all  $t$ , all  $\ell \in \mathcal{A}_t$  and all  $Z \in \mathbb{R}^{n_\ell \times m_\ell}$ . This choice has been made for simplicity, but the reader can readily verify that it is not necessary, and that it is sufficient to assume the existence of constants  $0 < \kappa_{1,\mathcal{N}} \leq 1 \leq \kappa_{2,\mathcal{N}}$  such that  $\|\mathcal{N}_{t,\ell}(Z)\|_\ell \in [\kappa_{1,\mathcal{N}}, \kappa_{2,\mathcal{N}}]$  for all  $t$ ,  $\ell \in \mathcal{A}_t$  and all  $Z \in \mathbb{R}^{n_\ell \times m_\ell}$ . The constants  $\kappa_{1,\mathcal{N}}$  and  $\kappa_{2,\mathcal{N}}$  then percolate through all proofs, complicating the expression of relevant factors even more, but without affecting the order of global convergence. This observation is particularly useful for methods involving matrix blocks, such as Shampoo or Muon, where step normalization is typically performed by a possibly inexact Newton–Schulz iteration [28, 5, 3, 48, 34].

## 6 Numerical experiments

### 6.1 Experimental objectives

The experiments in this section illustrate the practical behavior of the asynchronous block-parallel adaptive preconditioning method (ASADPREC) introduced in the previous sections. The convergence analysis in Section 3 establishes that the asynchronous block-coordinate iteration converges under a bounded-delay condition, provided the block preconditioners capture sufficient local curvature to control the inter-block coupling. The experiments are designed to verify that these theoretical predictions hold on concrete machine learning problems, and to quantify the practical trade-offs involved.

The OFFO (Objective-Function-Free Optimization) framework of [18] provides the algorithmic foundation: the parameter vector is decomposed into  $L$  blocks, each equipped with its own

adaptive preconditioner, and the method requires only gradient evaluations—no function value computations. In the synchronous regime, this reduces to a preconditioned block-Jacobi iteration; in the asynchronous regime, each block is updated independently using a possibly stale snapshot of the other blocks. Classical results on asynchronous chaotic relaxation [10] and asynchronous iterations [4, 15] predict convergence when the spectral radius of the “asynchronous iteration matrix” is less than one—a condition that is favored by weak inter-block coupling and strong local preconditioning. In practice, the throughput gain from asynchronous execution depends on the computational heterogeneity between blocks. When all blocks have equal per-iteration cost, asynchronous execution offers no advantage over the synchronized variant (the barrier imposes no idle time). When block costs differ, the synchronization barrier forces fast blocks to idle while waiting for the slowest block, and the throughput gain from removing the barrier is bounded by an Amdahl-type ratio: the cost of the slowest worker relative to the average. In what follows, we numerically explore this relationship across a range of heterogeneity levels and address the following questions:

1. Does asynchronous block-parallel optimization preserve convergence quality (training loss, gradient norm, generalization metric) relative to its synchronous counterpart, as predicted by the bounded-delay convergence theory?
2. What throughput gain does asynchronous execution provide, and how does this gain relate to the measured computational heterogeneity between blocks (Amdahl-type bound)?
3. How does exponential moving average (EMA) momentum interact with the stale reads inherent to asynchronous execution?

The benchmarks are chosen to cover a range of architectural patterns (dense MLPs, sparse embedding models, hybrid architectures) and heterogeneity levels, so that the interplay between block structure, computational cost imbalance, and asynchronous scheduling can be studied in a controlled setting. The goal is *not* to achieve state-of-the-art predictive performance, but rather to verify the theoretical predictions and to identify the regimes in which asynchronous block parallelism is most beneficial. We consider three algorithmic variants. All three apply the same per-block preconditioned updates and differ only in scheduling and synchronization.

**Synchronous single-gradient (SyncSingle):** A single mini-batch is sampled and the stochastic gradient  $G_k = \nabla F_{\mathcal{B}_k}(x_k)$  is computed once. All  $L$  blocks are updated from the same snapshot, using Steps 4–6 of the ASADPREC framework. This variant minimizes gradient computation cost but assumes that a single gradient evaluation can be shared across all block solvers.

**Synchronous multi-gradient (SyncMulti):** Each of the  $L$  workers independently computes the full stochastic gradient on the same shared batch and snapshot, then extracts its own block component. Workers are synchronized by barriers: no block update is applied until every worker has completed its gradient computation. The mathematical iterate is identical to that of SyncSingle, using Steps 4–6 of ASADPREC, but the wall-clock cost per iteration includes the overhead of  $L$  gradient evaluations and barrier synchronization.

**Asynchronous block-owner (ASADPREC):** The synchronization barriers are removed. Each block  $X_\ell$  is permanently assigned to a dedicated worker. All workers share access to the full parameter vector  $\hat{X} = (\hat{X}_1, \dots, \hat{X}_L)$ . Worker  $\ell$  repeatedly:

- (i) reads the full vector  $\hat{X}$  by cloning each block  $\hat{X}_j$  under block  $j$ ’s lock, ensuring that all parameters within a given block come from the same publication epoch (intra-block consistency); different blocks may come from different epochs (inter-block asynchrony);
- (ii) samples a local mini-batch  $\mathcal{B}_{t,\ell}$ , computes the full stochastic gradient  $\nabla F_{\mathcal{B}_{t,\ell}}(\hat{X}_k)$ , and extracts the block component  $G_{t,\ell} = [\nabla F_{\mathcal{B}_{t,\ell}}(\hat{X}_k)]_\ell$ ;
- (iii) applies the preconditioned update to  $\hat{X}_\ell$  and publishes the result atomically under block  $\ell$ ’s lock.

This is the asynchronous block-Jacobi model [4, 15] used by ASADPREC. A bounded-delay throttle pauses any worker whose publication counter exceeds the slowest worker’s by more than  $\delta_a = 512$ .

## 6.2 Test problems and architectures

Table 2 summarizes the five benchmarks. Each model is decomposed so that one or two blocks use the matrix-valued adaptive **Muon** preconditioner (Newton–Schulz polar approximation [26] for SVHN, SVD for the others) while the remaining blocks use diagonal **AdaGrad**.

| Dataset        | Architecture                                               | $n_{\text{train}}$ | $p$        | $L$ | Muon block(s)                                                                                                                         |
|----------------|------------------------------------------------------------|--------------------|------------|-----|---------------------------------------------------------------------------------------------------------------------------------------|
| FashionMNIST   | MLP $784 \rightarrow 4096 \rightarrow 10$                  | 60,000             | 3,256,330  | 4   | $W_1 \in \mathbb{R}^{4096 \times 784}$                                                                                                |
| MoE-FMNIST     | 4 experts (see text)                                       | 60,000             | 3,887,950  | 6   | $W_{1,1} \in \mathbb{R}^{512 \times 2048}$ ,<br>$W_{2,1} \in \mathbb{R}^{512 \times 1024}$ ,<br>$W_2 \in \mathbb{R}^{512 \times 256}$ |
| CovType        | MLP $54 \rightarrow 512 \rightarrow 256 \rightarrow 7$     | 160,000            | 161,287    | 4   | $W_2 \in \mathbb{R}^{512 \times 256}$                                                                                                 |
| MovieLens 100K | NeuralMF, $d=128$ , $h=512$                                | 80,000             | 470,722    | 5   | $W_1 \in \mathbb{R}^{512 \times 256}$                                                                                                 |
| Criteo         | Hash-embed + MLP $429 \rightarrow 256 \rightarrow 1$       | 160,000            | 634,625    | 3   | $W_1 \in \mathbb{R}^{429 \times 256}$                                                                                                 |
| SVHN           | MLP $3072 \rightarrow 4096 \rightarrow 512 \rightarrow 10$ | 73,257             | 14,689,802 | 5   | $W_1 \in \mathbb{R}^{4096 \times 3072}$ ,<br>$W_2 \in \mathbb{R}^{512 \times 4096}$                                                   |

Table 2: Test problems and block decompositions.

**FashionMNIST [51]:** A one-hidden-layer MLP with  $3.26 \times 10^6$  parameters decomposed into 4 blocks:  $W_1$  (**Muon**),  $b_1$ ,  $W_2$ ,  $b_2$  (**AdaGrad**).

**MoE-FMNIST [51]:** A Mixture-of-Experts architecture applied to flattened  $28 \times 28$  Fashion-MNIST images, with a backbone MLP ( $784 \rightarrow 512784 \rightarrow 512$ ), a top-2/4 softmax router, four experts of varying width ( $512 \rightarrow h \rightarrow 512$ ) with  $h \in \{2048, 1024, 256, 64\}$ , and an output head ( $512 \rightarrow 10$ ), decomposed into 6 blocks:  $W_{1,1}$  and  $W_{2,1}$  (**Muon**, Newton–Schulz, 5 iterations), expert 3 and expert 4 full parameters, backbone and expert 1 and 2 biases and second-layer weights, router and output head (all **AdaGrad**).

**CovType [6]:** A two-hidden-layer MLP with  $1.61 \times 10^5$  parameters (200,000 samples, 80/20 split), decomposed into 4 blocks:  $(W_1, b_1)$ ,  $W_2$  (**Muon**),  $b_2$ ,  $(W_3, b_3)$  (**AdaGrad**).

**MovieLens 100K [24]:** A NeuralMF model with  $4.71 \times 10^5$  parameters (100,000 ratings, 943 users, 1,682 items), decomposed into 5 blocks:  $P$ ,  $Q$ ,  $(b_u, b_i)$  (sparse row-wise **AdaGrad**),  $W_1$  (**Muon**),  $(c_1, W_2, c_2)$  (**AdaGrad**).

**Criteo [31]:** A click-through-rate model with  $6.35 \times 10^5$  parameters, decomposed into 3 blocks: embedding table (**AdaGrad**),  $W_1$  (**Muon**),  $(b_1, W_2, b_2)$  (**AdaGrad**).

**SVHN [40]:** A two-hidden-layer MLP with  $14.7 \times 10^6$  parameters applied to flattened  $32 \times 32 \times 3$  images, decomposed into 5 blocks:  $W_1$  and  $W_2$  (**Muon**, Newton–Schulz, 5 iterations),  $b_1$ ,  $b_2$ ,  $(W_3, b_3)$  (**AdaGrad**). This benchmark features two Muon blocks of different sizes ( $12.6 \times 10^6$  and  $2.1 \times 10^6$  entries), creating a hierarchy of block costs.

### 6.2.1 Measured block heterogeneity

The raw block cost ratio ( $t_{\text{max}}/t_{\text{min}}$ ) exceeds 90 on all benchmarks, but the effective heterogeneity is considerably lower. Each worker’s per-iteration cost has two components: the gradient computation  $t_{\text{grad}}$  (identical for all workers, since each computes the full gradient  $\nabla F$ ) and the block update  $t_{\text{block}, \ell}$ . When  $t_{\text{grad}}$  is large relative to  $t_{\text{block}, \ell}$ , all workers run at similar speeds regardless of block update costs. Table 3 reports per-block computational costs measured on an Apple M3 Pro processor (single-threaded, median over 50 repetitions). The *effective heterogeneity ratio* is defined as  $(t_{\text{grad}} + t_{\text{max}})/(t_{\text{grad}} + t_{\text{min}})$ .

| Dataset      | Gradient | Slowest block | Fastest block | $t_{\max}/t_{\min}$ | Eff. ratio |
|--------------|----------|---------------|---------------|---------------------|------------|
| FashionMNIST | 5.6      | 105.2 (Muon)  | 0.004         | 25,000              | 19.6       |
| MoE-FMNIST   | 37.6     | 451.0 (Muon)  | 0.035         | 12,733              | 5.7        |
| CovType      | 0.8      | 0.42 (Muon)   | 0.004         | 96                  | 1.5        |
| MovieLens    | 1.0      | 5.61 (Muon)   | 0.011         | 526                 | 6.5        |
| Criteo       | 1.8      | 4.83 (Muon)   | 0.012         | 419                 | 3.7        |
| SVHN         | 23.4     | 580.2 (Muon)  | 0.004         | 134,000             | 25.8       |

Table 3: Per-block computational cost (ms) and effective heterogeneity.

On FashionMNIST, the Muon polar factorization (105 ms) far exceeds the gradient (5.6 ms), so the effective ratio remains high (19.6). On CovType, the gradient (0.8 ms) exceeds the slowest block update (0.42 ms), reducing the effective ratio to  $1.5\times$ . SVHN exhibits the highest effective ratio (25.8) because of the Newton–Schulz iteration on  $W_1 \in \mathbb{R}^{4096 \times 3072}$  (580 ms) relative to the gradient (23 ms).

### 6.3 Implementation and experimental protocol

All experiments use PyTorch on a shared-memory CPU architecture (Apple M3 Pro, 12 cores) with `torch.set_num_threads(1)`. The CPython GIL prevents byte-code level parallelism, but PyTorch C++ operations (matrix multiplications, SVD, tensor cloning) release the GIL, allowing concurrent execution across cores. The ASADPREC mode uses per-block locking: each block is read and written under its own lock, ensuring intra-block consistency while allowing inter-block asynchrony. A fixed wall-clock budget of 30 seconds is used for problems CovType, MovieLens and Criteo, a budget of 60 seconds for Fashion-MNIST, and MoE-FMNIST and a budget of 120 seconds for SVHN. All runs use 8 seeds (1234, ..., 1241), and identical initial parameters per seed. Throughput is reported in model-equivalent updates per second:  $\text{throughput} = N_{\text{block}}/(L \cdot T)$ .

### 6.4 Results without momentum

Tables 4 and 5 present the results obtained with the momentum-less ASADPREC in the setting just described. We start by reporting, for all benchmarks, the full-dataset training loss and the Euclidean norm of the full training gradient  $\|\nabla F(X)\|$  as convergence diagnostics.

| Dataset      | Final training loss gradient $\ \nabla F(X)\ $ |              |              | Final training loss value $F(X)$ |           |              |
|--------------|------------------------------------------------|--------------|--------------|----------------------------------|-----------|--------------|
|              | SyncSingle                                     | SyncMulti    | ASADPREC     | SyncSingle                       | SyncMulti | ASADPREC     |
| FashionMNIST | 0.421                                          | 0.343        | <b>0.277</b> | 0.321                            | 0.321     | <b>0.308</b> |
| MoE-FMNIST   | <b>0.632</b>                                   | 0.826        | 0.744        | 0.309                            | 0.314     | <b>0.279</b> |
| CovType      | 0.568                                          | <b>0.482</b> | 0.485        | 0.243                            | 0.260     | <b>0.240</b> |
| MovieLens    | 0.024                                          | 0.027        | <b>0.021</b> | <b>0.294</b>                     | 0.340     | 0.353        |
| Criteo       | 0.044                                          | 0.053        | <b>0.033</b> | <b>0.076</b>                     | 0.112     | 0.088        |
| SVHN         | 1.004                                          | <b>0.979</b> | 1.180        | 0.840                            | 0.833     | <b>0.819</b> |

Table 4: Convergence diagnostics for momentum-less variants

ASADPREC obtains the best gradient (in norm) for Criteo, FashionMNIST and MovieLens, while there is little difference between SyncSingle and ASADPREC for CovType. The gradient norm is higher for SVHN and MoE-FMNIST (due to delayed gradients). The best loss value is obtained by ASADPREC for 4/6 benchmarks. SyncSingle remains better for MovieLens and Criteo because it performs only one gradient evaluation per iteration.

We now turn to primary generalization metrics, depending on the benchmark. For classification problems (FashionMNIST, CovType, SVHN), we report the *test accuracy* (fraction of correctly classified test samples). For MovieLens, we report the *root mean square error* (RMSE), defined as  $\text{RMSE} = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{r}_i - r_i)^2}$ , where  $\hat{r}_i$  and  $r_i$  are the predicted and true ratings. For Criteo, we report the *area under the ROC curve* (AUC), which measures the probability that a randomly

chosen positive example (clicked ad) is ranked higher than a randomly chosen negative example by the model’s predicted score;  $\text{AUC} = 0.5$  corresponds to random ranking and  $\text{AUC} = 1$  to perfect ranking. Generalization metrics are reported in Table 5, along with the throughput of ASADPREC and the speedup of ASADPREC compared with SyncMulti. As can be seen in this table, ASADPREC is best on all benchmarks except Criteo.

| Dataset      | metric  | SyncSingle | SyncMulti    | ASADPREC     | Throughput | Speedup |
|--------------|---------|------------|--------------|--------------|------------|---------|
| FashionMNIST | acc.(%) | 87.00      | 86.94        | <b>87.34</b> | 10         | 1.43    |
| MoE-FMNIST   | acc.(%) | 86.82      | 86.52        | <b>87.39</b> | 42         | 1.27    |
| CovType      | acc.(%) | 89.49      | 88.82        | <b>89.56</b> | 105        | 1.27    |
| MovieLens    | RMSE    | 0.930      | 0.919        | <b>0.916</b> | 89         | 1.14    |
| Criteo       | AUC     | 0.715      | <b>0.722</b> | 0.717        | 89         | 1.17    |
| SVHN         | acc.(%) | 69.44      | 69.28        | <b>69.65</b> | 1          | 1.33    |

Table 5: Generalization metrics, throughput and speed-up for momentum-less variants

Globally, Tables 4 and 5 suggest that stale (delayed) gradients do not degrade the quality of the training and that the additional iterations made possible by asynchrony translate into better generalization.

## 6.5 Results with momentum

The approaches yielding DADPREC.MM/ASADPREC.MM and DADPREC.MG/ASADPREC.MG not only differ in theory, but their practical requirements are also significantly different, in particular for the Shampoo and Muon case. Indeed, if the SVD of  $M_\ell$  is given by  $M_\ell = U_\ell \Sigma_\ell V_\ell^T$ , the Muon recursions<sup>7</sup> are written, in this case, as

$$\begin{array}{ll}
 M_\ell &= \mu_k M_\ell^- + (1 - \mu_k) \tilde{G}_\ell \\
 \Pi_\ell &= \Pi_\ell^- + \frac{\|M_\ell\|_{*,\ell}^2}{d_\ell}, \\
 X_\ell^+ &= X_\ell - \eta \frac{\|M_\ell\|_{*,\ell}}{\sqrt{\Pi_\ell}} U_\ell V_\ell^T \\
 &\text{for Muon.MM,}
 \end{array}
 \quad
 \begin{array}{ll}
 \Pi_\ell &= \Pi_\ell^- + \frac{\|\tilde{G}_\ell\|_{*,\ell}^2}{d_\ell}, \\
 M_\ell &= \mu_k M_\ell^- + (1 - \mu_k) \tilde{G}_\ell \\
 X_\ell^+ &= X_\ell - \eta \frac{\|M_\ell\|_{*,\ell}}{\sqrt{\Pi_\ell}} U_\ell V_\ell^T \\
 &\text{for Muon.MG.}
 \end{array}$$

The main difference is that Muon.MM only requires  $\|M_\ell\|_{*,\ell}$ , which can be computed with a single polar decomposition (using SVD or Newton–Schulz iterations), while Muon.MG also requires  $\|\tilde{G}_\ell\|_{*,\ell}$ . This potentially requires another decomposition, which can significantly slow down the algorithm. It is therefore of interest to consider faster techniques for computing  $\|\tilde{G}_\ell\|_{*,\ell}$ . We have considered both the stochastic Lanczos quadrature [47] and a warm-started power iteration for this task. As it turns out, the first, although theoretically suitable, appears to be too slow for practical use. By contrast, the second, which can only provide a deterministic lower bound, nevertheless works remarkably well. The details of this technique are given in Appendix 3.

When applied to ASADPREC, these considerations lead us to a family of momentum variants, of which we have retained six, whose characteristics are given in Table 6. In this table,  $\beta$  is the rate of momentum decay appearing in (4.9) (with  $\mu_{\max} = 0.9$ ) and the last column indicates whether a warm-started power iteration is used to approximate  $\|\tilde{G}_\ell\|_{*,\ell}$  in the Muon updates.

To keep the paper as concise as possible, we only report in Table 7 the values of the final gradients (in norm) for these six variants, and refer the reader to Appendix 3 for more results. As can be seen from this table, no specific momentum variant appears to dominate, but the momentum-less variant (first column) is never the best.

## 6.6 Discussion

Globally, our results suggest the following conclusions.

<sup>7</sup> $\Pi_\ell$  is a scalar in this case.

| Variant  | type | $\beta$ | approx. $\ \tilde{G}_\ell\ _{*,\ell}$ |
|----------|------|---------|---------------------------------------|
| MMconst  | MM   | 0       | no                                    |
| MGconst  | MG   | 0       | no                                    |
| MGconstA | MG   | 0       | yes                                   |
| MMdecay  | MM   | 0.5     | no                                    |
| MGdecay  | MG   | 0.5     | no                                    |
| MGdecayA | MG   | 0.5     | yes                                   |

Table 6: Asynchronous momentum variants and their characteristics

| Dataset       | No Mom. | MMconst      | MGconst      | MGconstA | MMdecay      | MGdecay      | MGdecayA |
|---------------|---------|--------------|--------------|----------|--------------|--------------|----------|
| Fashion-MNIST | 0.277   | 0.551        | 0.500        | 0.526    | 0.255        | <b>0.246</b> | 0.312    |
| MoE-FMNIST    | 0.744   | 0.858        | <b>0.656</b> | 1.212    | 0.663        | 0.884        | 0.881    |
| CovType       | 0.485   | <b>0.246</b> | 0.279        | 0.265    | 0.345        | 0.386        | 0.755    |
| MovieLens     | 0.021   | 0.021        | 0.036        | 0.032    | <b>0.015</b> | 0.019        | 0.025    |
| Criteo        | 0.033   | 0.036        | <b>0.008</b> | 0.018    | 0.046        | 0.035        | 0.028    |
| SVHN          | 1.180   | 1.430        | 0.988        | 1.138    | 1.103        | <b>0.880</b> | 1.131    |

Table 7: Final gradient norm for momentum variants

1. *Asynchronous parallelism preserves and improves convergence quality.* DADPREC is the best momentum-less variant for 5/6 benchmarks, and is also best in many cases where momentum is used.
2. *Heterogeneity drives the speedup.* The per-block locking protocol—guaranteeing intra-block consistency while allowing inter-block asynchrony—matches the classical asynchronous block-Jacobi model [4, 15]. The effective heterogeneity ratio, which accounts for the shared gradient computation cost, is a better predictor of the asynchronous speedup than the raw block cost ratio.
3. *The cost of momentum is critical for matrix updates* like Muon. This makes the MM approach attractive from this point of view, because its lower cost per iteration allows more iterations in the time budget, resulting in gains in accuracy between +1.8 and +3.6 points compared with the MG approach.
4. The estimation of the nuclear norm of  $\tilde{G}_\ell$  using the *warm-started power iteration* is an *efficient proxy*, and does produce the best results for some problems.
5. As expected, the effect of decaying momentum (i.e.  $\beta > 0$  in (4.9)) is problem-dependent, because of its direct interaction with the variance of the problem’s oracle.
6. Although globally promising, our tests did not allow us to select a clear winner among the studied variants.

It should however be stressed that all our experiments were conducted assuming the “full gradient evaluation model”, that is the case where  $G_{t,\ell} = [\nabla F(X_t, \xi)]_{\mathcal{I}_\ell}$  cannot be computed on its own, but requires the evaluation of  $\nabla F(X_t, \xi)$ . Should the evaluation of  $G_{t,\ell}$  be possible at a lower cost<sup>8</sup>, the heterogeneity ratios would then be dramatically increased, which is likely to give the asynchronous option used in ASADPREC a substantial advantage.

## 7 Conclusions

We have presented a new class of asynchronous adaptive first-order methods including asynchronous variants of AdaNorm, AdaGrad, Shampoo and Muon, as well as variants thereof using momentum and/or inexact step normalization. We have shown that, under bounded-delay and cumulative variance assumptions, all methods in the class have a expected global convergence rate essentially of order  $\mathcal{O}(1/\sqrt{t})$ .

<sup>8</sup>Among other possibilities, we think here of structured networks such as mixture of experts [23], DeepONets [35] or multi-frequency networks [17]. The approach of delayed gradients suggested by [55], in some sense close to DADPREC, may also be worth investigating.

We also described numerical experiments with these methods, providing motivation for asynchronous adaptive optimization where different blocks use different geometries, different blocks have drastically different computational costs, and synchronization barriers become prohibitively expensive. Such situations naturally appear in modern large-scale machine learning systems mixing diagonal adaptive methods, with matrix-based preconditioned methods. Our results, although limited, provide empirical support for the claim that asynchronous adaptive block-coordinate optimization methods may be relevant in heterogeneous large-scale machine learning systems.

## References

- [1] A. Attia and T. Koren. SGD with AdaGrad stepsizes: Full adaptivity with high probability to unknown parameters, unbounded gradients and affine variance. *arXiv:2302.08783*, 2023.
- [2] J. Bernstein and L. Newhouse. Modular duality in deep learning. *arXiv:2410.21265v2*, 2024.
- [3] J. Bernstein and L. Newhouse. Old optimizer, new norm: an anthology. *arXiv:2409.20325v2*, 2024.
- [4] D.P. Bertsekas and J. N. Tsitsiklis. *Parallel and Distributed Computation: Numerical Methods*. Prentice Hall, 1989.
- [5] Å. Björck and C. Bowie. An iterative algorithm for computing the best estimate of an orthogonal matrix. *SIAM Journal on Numerical Analysis*, 8(2):358–364, 1971.
- [6] J. A. Blackard and D. J. Dean. Comparative accuracies of artificial neural networks and discriminant analysis in predicting forest cover types from cartographic variables. *Comput. Electron. Agric.*, 24(3):131–151, 1999.
- [7] S. Boyd and L. Vandenberghe. *Convex Optimization*. Cambridge University Press, Cambridge, England, 2004.
- [8] L. Cannelli, F. Facchinei, V. Kungurtev, and G. Scutari. Asynchronous parallel algorithms for nonconvex optimization. *Mathematical Programming, Series A*, 184:121–154, 2020.
- [9] S. C. Chaturapruek, J. C. Duchi, and Ch. Ré. Asynchronous stochastic convex optimization. In *NIPS’15: Proceedings of the 29th International Conference on Neural Information Processing Systems*, volume 1, pages 1531–1539, 2015.
- [10] D. Chazan and W. Miranker. Chaotic relaxation. *Linear Algebra and its Applications*, 2(2):199–222, 1969.
- [11] A. R. Conn, N. I. M. Gould, and Ph. L. Toint. *LANCELOT: a Fortran package for large-scale nonlinear optimization (Release A)*. Number 17 in Springer Series in Computational Mathematics. Springer Verlag, Heidelberg, Berlin, New York, 1992.
- [12] J. Duchi, E. Hazan, and Y. Singer. Adaptive subgradient methods for online learning and stochastic optimization. *Journal of Machine Learning Research*, 12, July 2011.
- [13] M. Faw, I. Tziotis, C. Caramanis, A. Mokhtari, S. Shakkottai, and R. Ward. The power of adaptivity in SGD: Self-tuning step sizes with unbounded gradients and affine variance. In *Proceedings of 35th Conference on Learning Theory*, volume 178, pages 313–355, 2022.
- [14] H. R. Feyzmahdavian and M. Johansson. Asynchronous iterations in optimization: New sequence results and sharper algorithmic guarantees. *Journal of Machine Learning Research*, 24:1–75, 2023.
- [15] A. Frommer and D. B. Szyld. On asynchronous iterations. *Journal of Computational and Applied Mathematics*, 123(1–2):201–216, 2000.
- [16] B. Ginsburg, P. Castonguay, O. Hrinchuk, O. Kuchaiev, R. Leary, V. Lavrukhin, J. Li, H. Nguyen, Y. Zhang, and J. M. Cohen. Training deep networks with stochastic gradient normalized by layerwise adaptive second moments. *arXiv:1905.11286v3*, 2019.
- [17] S. Gratton, V. Mercier, E. Riccietti, and Ph. L. Toint. A block-coordinate approach of multi-level optimization with an application to physics-informed neural networks. *Computational Optimization and Applications*, 89(2):385–417, 2024.
- [18] S. Gratton and Ph. L. Toint. A unified convergence theory for adaptive first-order methods in the nonconvex case, including AdaNorm, full and diagonal AdaGrad, Shampoo and Muon. *arXiv:2604.17423*, 2026.
- [19] A. Griewank and Ph. L. Toint. Local convergence analysis for partitioned quasi-Newton updates. *Numerische Mathematik*, 39:429–448, 1982.
- [20] A. Griewank and Ph. L. Toint. On the existence of convex decomposition of partially separable functions. *Mathematical Programming*, 28:25–49, 1984.
- [21] Z. Guo, W. Yin, R. Jin, and T. Yang. A novel convergence analysis for algorithms of the Adam family. In *Proceedings of OPT2021: 13th Annual Workshop on Optimization for Machine Learning*, 2021.
- [22] V. Gupta, T. Koren, and Y. Singer. Shampoo: Preconditioned stochastic tensor optimization. In *Proceedings of the International Conference on Machine Learning*, pages 1842–1850, 2018.

- [23] J. B. Hampshire and A. Waibel. The Meta-Pi network: building distributed knowledge representations for robust multisource pattern recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 14(7):751–769, 1992.
- [24] F. M. Harper and J. A. Konstan. The MovieLens datasets: History and context. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 5(4), 2025.
- [25] Y. Hong and J. Lin. Revisting convergence of Adagrad with relaxed assumptions. arXiv:2403.13794v2, 2024.
- [26] K. Jordan, Y. Jin, V. Boza, Y. Jiacheng, F. Cecista, L. Newhouse, and J. Bernstein. URL:https://kellerjordan.github.io.posts/muon, 2024.
- [27] D. Kingma and J. Ba. Adam: A method for stochastic optimization. In *Proceedings in the International Conference on Learning Representations (ICLR)*, 2015.
- [28] Z. Kovarik. Some iterative methods for improving orthonormality. *SIAM Journal on Numerical Analysis*, 7(3):386–389, 1970.
- [29] V. Kungurtsev, M. Egan, B. Chatterjee, and D Alistarh. Asynchronous optimization methods for efficient training of deep neural networks with guarantees. arXiv:1905.11845v2, 2019.
- [30] V. Kungurtsev, M. Egan, B. Chatterjee, and D Alistarh. Asynchronous optimization methods for efficient training of deep neural networks with guarantees. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 8209–8216, 2021.
- [31] Criteo Labs. Display advertising challenge. Kaggle, 2014.
- [32] R. Leblond, F. Pedregosa, and S. Lacoste-Julien. Improved asynchronous parallel optimization analysis for stochastic incremental methods. *Journal of Machine Learning Research*, 19(81):1–68, 2018.
- [33] X. Lian, Y. Huang, Y. Li, and J. Liu. Asynchronous parallel stochastic gradient for nonconvex optimization. In C. Cortes, N. Lawrence, D. Lee, M. Sugiyama, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc., 2015.
- [34] J. Liu, J. Lin, X. Yao, Z. Jiang, G. Lai, Y. Du, Y. Qin, W. Xu, E. Lu, J. Yan, Y. Chen and. H. Zheng, Y. Liu, S. Liu, B. Yin, W. He, H. Zhu, Y. Wang, J. Wang, M. Dong, Z. Zhang, Y. Kang, H. Zhang, X. Xu, Y. Zhang, Y. Wu, W. Zhou, and Z. Yang. MUON is scalable for LLM training. arXiv:2502.16982, 2025.
- [35] L. Lu, P. Jin, and G. Karniadakis. DeepONet: Learning nonlinear operators for identifying differential equations based on the universal approximation theorem of operators.
- [36] H. Mania, X. Pan, D. Papailiopoulos, B. Recht, K. Ramchandran, and M. I. Jordan. Perturbed iterate analysis for asynchronous stochastic optimization. *SIAM Journal on Optimization*, 27(4):2202–2229, 2017.
- [37] B. McMahan and M. Streeter. Adaptive bound optimization for online convex optimization. In *Conference on Learning Theory*, page 244sq, 2010.
- [38] A. Nabli and E. Oyallon. DADAO: Decoupled accelerated decentralized asynchronous optimization. arXiv:2208.00779v3, 2022.
- [39] G. Nadiradze, I. Markov, B. Chatterjee, and V. Kungurtsev. Elastic consistency: a practical consistency model for distributed stochastic gradient descent. *ACM SIGACT News*, 53(2):64–82, 2022.
- [40] Y. Netzer, T. Wang, A. Coates, A. Bissacco, B. Wu, and A.Y. Ng. Reading digits in natural images with unsupervised feature learning, 2011.
- [41] L. M. Nguyen, P. H. Nguyen, P. Richtarik, K. Scheinberg, M. Takáč, and M. van Dijk. New convergence aspects of stochastic gradient algorithms. *Journal of Machine Learning Research*, 20(176):1–49, 2019.
- [42] L. M. Nguyen, P. H. Nguyen, M. van Dijk, P. Richtarik, K. Scheinberg, and M. Takáč. SGD and Hogwild! convergence without the bounded gradients assumption. *Proceedings of Machine Learning Research*, 80:3750–3758, 2018.
- [43] F. Niu, B. Recht, Ch. Ré, and S. J. Wright. HOGWILD!: A lock-free approach to parallelizing stochastic gradient descent. *Advances in Neural Information Processing Systems*, 4(1):693–701, 2011.
- [44] Z. Peng, Y. Xu, M. Yan, and W. Yin. ARock: An algorithmic framework for asynchronous parallel coordinate updates. *SIAM Journal on Scientific Computing*, 38(5):2851–2879, 2016.
- [45] Th. Pethick, W. Xie, K. Antonakopoulos, Z. Zhu, A. Silveti-Falls, and V. Cevher. Training deep learning models with norm-constrained LMOs. arXiv:2502.07529v2, 2025.
- [46] C. Si, D. Zhang, and W. Shen. AdaMuon: Adaptive Muon optimizer. arXiv:2507.11005, 2025.
- [47] S. Ubaru, J. Chen, and Y. Saad. Fast estimation of  $\text{Str}(f(a))$  via stochastic Lanczos quadrature. *SIAM Journal on Matrix Analysis and Applications*, 38(4):1075–1099, 2017.
- [48] N. Vyas, D. Morwani, R. Zhao, M. Kwun, I. Shapira, B. Brandfonbrener, L. Janson, and S. Kakade. SOAP: Improving and stabilizing Shampoo using Adam. arXiv:2409.11321, 2024.
- [49] B. Wang, H. Zhang, Z. Ma, and W. Chen. Convergence of Adagrad for non-convex objectives: simple proofs and relaxed assumptions. In *Proceedings of the Thirty-Sixth Conference on Learning Theory*, pages 161–190, 2023.



- [50] R. Ward, X. Wu, and L. Bottou. Adagrad stepsizes: sharp convergence over nonconvex landscapes. In *Proceedings of the International Conference on Machine Learning (ICML2019)*, 2019.
- [51] H. Xiao, K. Rasul, and R. Vollgraf. Fashion-MNIST, a novel image dataset for benchmarking machine learning algorithms. arXiv:1708.07747, 2017.
- [52] N. Xiao, X. Hu, X. Lin, and K.-C. Toh. Adam family methods for nonsmooth optimization with convergence guarantees. *Journal of Machine Learning Research*, 25, 2024.
- [53] A. W. Yu, L. Huang, Q. Lin, R. Salakhutdinov, and J. Carbonell. Block-normalized gradient method; an empirical study for training deep neural network. arXiv:1707.04822v2, 2017.
- [54] M. Zhang, Y. Liu, and H. Schaeffer. AdaGrad meets Muon: Adaptive stepsizes for orthogonal updates. arXiv:2509.02981v2, 2025.
- [55] H. Zhuang, Y. Wang, Q. Liu, S. Zhang, and Z. Lin. Fully decoupled neural network learning using delayed gradients. arXiv:1906.09108v3, 2019.

## Appendix 1: Detailed analysis for momentum-less DADPREC

The proof of Theorem 3.2 follows the broad lines of [18], but significant modifications are needed to handle the delayed block framework considered here. The main differences are that the descent and preconditioning estimates must be written on the active block set  $\mathcal{A}_t$ , and that the random block gradients  $\tilde{G}_{t,\ell}$  satisfy the cumulative approximation bound of Assumption 5. We give below the parts of the argument where these features enter explicitly. The elementary estimates on the normalization operator, on the growth of the block preconditioners, and on the associated telescopic sums are used in the same form as in [18].<sup>9</sup>

Our first step is to analyze the (expected) descent at iteration  $t$ .

**Lemma A.1** Suppose that Assumption 2 holds. Then

$$\begin{aligned} \mathbb{E}_t[f(X_{t+1})] &\leq f(X_t) - \eta \sum_{\ell \in \mathcal{A}_t} \mathbb{E}_t \left[ \|Z_{t,\ell}\|_{*,\ell} \left\langle \tilde{G}_{t,\ell}, \mathcal{N}_{t,\ell}(Z_{t,\ell}) \right\rangle_F \right] \\ &\quad + \eta \sum_{\ell \in \mathcal{A}_t} \mathbb{E}_t \left[ \|Z_{t,\ell}\|_{*,\ell} \left| \left\langle \tilde{G}_{t,\ell} - G_{t,\ell}, \mathcal{N}_{t,\ell}(Z_{t,\ell}) \right\rangle_F \right| \right] + \frac{L_G \eta^2}{2} \sum_{\ell \in \mathcal{A}_t} \mathbb{E}_t [\|Z_{t,\ell}\|_{*,\ell}^2]. \end{aligned} \quad (\text{A.1})$$

**Proof.** See [18, Lemma 3.1]. □

The following trace inequalities may then be proved.

**Lemma A.2** Suppose that a momentum-less DADPREC algorithm is applied to problem (1.1). Then, for all  $t \geq 0$ ,

$$\sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) - \sum_{\ell=1}^L \text{tr}(\Pi_{-1,\ell}^{1/2}) \leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \text{tr} \left( \Pi_{j,\ell}^{-1/2} \mathcal{U}_{j,\ell} (\tilde{G}_{j,\ell})^2 \right), \quad (\text{A.2})$$

and

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \text{tr} \left( \Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell} (\tilde{G}_{j,\ell})^2 \right) \leq \sum_{\ell=1}^L \text{tr}(\log(\Pi_{t,\ell})) - \sum_{\ell=1}^L \text{tr}(\log(\Pi_{-1,\ell})). \quad (\text{A.3})$$

<sup>9</sup>Notation has changed: in [18], the iteration index was  $k$ , the normalization operator was denoted by  $S_\ell(\cdot)$ , and the preconditioner update operator by  $\mathcal{L}_{k,\ell}(\cdot)$ .

**Proof.** See [18, Lemmas 3.2–3.4]. In the last of these lemmas, the sums over all iterations and all blocks are replaced here by  $\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}}$ , because only the iterations at which block  $\ell$  is active contribute to the update of  $\Pi_{t,\ell}$ .  $\square$

Because not every block is necessarily updated at iteration  $t$  in the DADPREC algorithm, we now need to re-prove the suitable variant of the classical telescoping argument.

**Theorem A.3** Suppose that Assumption 1, 2, 3 and 5 hold. Define

$$\Delta_t \stackrel{\text{def}}{=} \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\log \Pi_{t,\ell}) - \sum_{\ell=1}^L \text{tr}(\log \Pi_{-1,\ell}) \right]. \quad (\text{A.4})$$

Then, for every  $t \geq 0$ ,

$$\eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) \right] \leq \kappa_{\text{gap}} + \eta \nu_t \sqrt{\Delta_t} + \left( \eta \omega + \frac{L_G \eta^2}{2} \right) \Delta_t, \quad (\text{A.5})$$

where  $\kappa_{\text{gap}} \stackrel{\text{def}}{=} \mathbb{E}[f(X_0)] - f_{\text{low}} + \eta \varsigma N$ .

**Proof.** Taking the full expectation in the conditional descent inequality (A.1) and summing for  $j = 0$  to  $t$  gives

$$\begin{aligned} & \eta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \left\langle \tilde{G}_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] \\ & \leq \mathbb{E}[f(X_0)] - f_{\text{low}} + \eta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] \\ & \quad + \frac{L_G \eta^2}{2} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2 \|\mathcal{N}_{j,\ell}(Z_{j,\ell})\|_{\ell}^2]. \end{aligned} \quad (\text{A.6})$$

Using successively (3.1) and (A.2), we obtain that

$$\begin{aligned} \eta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \left\langle \tilde{G}_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] &= \eta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} \left[ \text{tr}(\Pi_{j,\ell}^{-1/2} \mathcal{U}_{j,\ell}(\tilde{G}_{j,\ell})^2) \right] \\ &\geq \eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) - \sum_{\ell=1}^L \text{tr}(\Pi_{-1,\ell}^{1/2}) \right]. \end{aligned} \quad (\text{A.7})$$

Now observe that the Cauchy-Schwarz inequality and the definition of the normalization operator  $\mathcal{N}_{j,\ell}$  give that

$$\left| \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right| \leq \|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell} \|\mathcal{N}_{j,\ell}(Z_{j,\ell})\|_{\ell} = \|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}.$$

Therefore, for  $j \in \mathcal{J}_{t,\ell}$ ,

$$\mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] \leq \sqrt{\mathbb{E} [\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2]} \sqrt{\mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2]}.$$

The Cauchy-Schwarz inequality also implies that, for any vectors  $a$  and  $b$  with nonnegative components,

$$\sum_j \sqrt{a_j} \sqrt{b_j} \leq \|\sqrt{a}\|_E \|\sqrt{b}\|_E = \sqrt{\sum_j a_j} \sqrt{\sum_j b_j}. \quad (\text{A.8})$$

Hence, summing over  $\ell \in \mathcal{A}_j$  and using (2.2) and (2.7), we obtain that

$$\sum_{\ell \in \mathcal{A}_j} \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \mid \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] \leq \sqrt{\sum_{\ell \in \mathcal{A}_j} \mathbb{E} [\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_*^2]} \sqrt{\sum_{\ell \in \mathcal{A}_j} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2]}.$$

Summing now over  $j \in \{0, \dots, t\}$ , noting that  $\sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} = \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}}$  and using (A.8) again, we deduce that

$$\begin{aligned} & \eta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \mid \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(z_{j,\ell}) \right\rangle_F \right] \\ & \leq \eta \sqrt{\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} [\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_*^2]} \sqrt{\sum_{\ell=1}^L \sum_{j=0}^t \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2]} \\ & \leq \eta \nu_t \sqrt{\sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2]} + \eta \omega \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2] \end{aligned}$$

where we used (3.4) to obtain the last inequality. But, by (3.2), (A.3) and (A.4),

$$\sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2] = \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2] = \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} [\text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell} (\tilde{G}_{j,\ell})^2)] \leq \Delta_t, \quad (\text{A.9})$$

so that

$$\eta \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} \left[ \|Z_{j,\ell}\|_{*,\ell} \mid \left\langle \tilde{G}_{j,\ell} - G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \right\rangle_F \right] \leq \eta \nu_t \sqrt{\Delta_t} + \eta \omega \Delta_t. \quad (\text{A.10})$$

Similarly, using the definition of the normalization operator  $\mathcal{N}_{j,\ell}$  and (A.9), the quadratic term satisfies

$$\frac{L_G \eta^2}{2} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2 \|\mathcal{N}_{j,\ell}(Z_{j,\ell})\|_\ell^2] = \frac{L_G \eta^2}{2} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2] \leq \frac{L_G \eta^2}{2} \Delta_t. \quad (\text{A.11})$$

Substituting (A.7), (A.10) and (A.11) into (A.6) then yields that

$$\eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) - \sum_{\ell=1}^L \text{tr}(\Pi_{-1,\ell}^{1/2}) \right] \leq f(X_0) - f_{\text{low}} + \eta \nu_t \sqrt{\Delta_t} + \left( \eta \omega + \frac{L_G \eta^2}{2} \right) \Delta_t, \quad (\text{A.12})$$

which is exactly (A.5) after taking into account that  $\Pi_{-1,\ell} = \varsigma I_\ell$  for all  $\ell \in \{1, \dots, L\}$ .  $\square$

We now recall from [18] a crucial consequence of the previous theorem, which provides a bound  $\Theta_t$  on the expected trace of all preconditioners at iteration  $t$ .

**Theorem A.4** Suppose that Assumptions 1, 2, 3 and 5 hold. Then, for all  $t \geq 0$ ,

$$\mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) \right] \leq \Theta_t \stackrel{\text{def}}{=} \max [\kappa_\Theta, T_t], \quad (\text{A.13})$$

where

$$\kappa_\Theta = \max \left[ e^{\max[1, \frac{\kappa_0}{2N}]}, \frac{3\kappa_{\text{gap}}}{\eta}, 24N \left( \omega + \frac{L_G \eta}{2} \right) \log \left( \max \left[ e, 24N \left( \omega + \frac{L_G \eta}{2} \right) \right] \right) \right],$$

with  $\kappa_{\text{gap}}$  defined in Theorem A.3,

$$T_t = 12\sqrt{N} \nu_t \sqrt{\log \left( \max \left[ e, 12\sqrt{N} \nu_t \right] \right)}, \quad (\text{A.14})$$

$\nu_t$  being defined in Assumption 5 and

$$\kappa_0 = \max \left[ - \sum_{\ell=1}^L d_\ell \log(d_\ell) - N \log(\varsigma), 0 \right]. \quad (\text{A.15})$$

Moreover,

$$\sum_{j=0}^t \mathbb{E} [\|Z_j\|_*^2] = \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E} [\|Z_{j,\ell}\|_{*,\ell}^2] \leq \kappa_0 + 2N \log(\Theta_t). \quad (\text{A.16})$$

**Proof.** Set

$$Y_t \stackrel{\text{def}}{=} \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) \right].$$

By Theorem A.3, we have, after division by  $\eta$ , that

$$Y_t \leq \frac{\kappa_{\text{gap}}}{\eta} + \nu_t \sqrt{\Delta_t} + \left( \omega + \frac{L_G \eta}{2} \right) \Delta_t.$$

Moreover, by [18, Lemma 3.6], applied here with the present block notation,

$$\Delta_t \leq \kappa_0 + 2N \log(Y_t), \quad (\text{A.17})$$

where  $\kappa_0$  is given by (A.15). The preceding two inequalities are exactly the scalar inequalities used in the proof of [18, Theorem 3.9], with the coefficient of  $\Delta_t$  now equal to

$$\omega + \frac{L_G \eta}{2}.$$

Therefore, the same argument gives

$$Y_t \leq \Theta_t \stackrel{\text{def}}{=} \max [\kappa_\Theta, T_t], \quad (\text{A.18})$$

where  $\kappa_\Theta$  and  $T_t$  are defined in (A.13)–(A.14).

It remains to prove (A.16). For every active block  $\ell \in \mathcal{A}_j$ , the definition of the normalized direction  $Z_{j,\ell}$  and Assumption 3 give

$$\|Z_{j,\ell}\|_{*,\ell}^2 = \text{tr} \left( \Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell} (\tilde{G}_{j,\ell})^2 \right).$$

Moreover,  $Z_{j,\ell} = 0$  when  $\ell \notin \mathcal{A}_j$ . Hence,

$$\sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] = \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2] = \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell} (\tilde{G}_{j,\ell})^2)].$$

Applying the logarithmic potential estimate (A.3) of Lemma A.2, and then taking expectations, yields

$$\sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \leq \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\log \Pi_{t,\ell}) - \sum_{\ell=1}^L \text{tr}(\log \Pi_{-1,\ell}) \right] = \Delta_t.$$

Combining this inequality with (A.17) and (A.18) and using the monotonicity of the logarithm, we obtain that

$$\sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \leq \Delta_t \leq \kappa_0 + 2N \log(Y_t) \leq \kappa_0 + 2N \log(\Theta_t),$$

which proves (A.16).  $\square$

We now prove two additional technical results. The first estimates the discrepancy between block gradients for iterations at which the block is or isn't selected.

**Lemma A.5** Suppose that Assumptions 2 and 7 hold. Define, for  $t > 0$  with  $j \notin \mathcal{A}_t$ ,

$$\gamma(t, j) = \max[i \in \{0, \dots, t-1\} \mid j \in \mathcal{A}_i], \quad (\text{A.19})$$

the index of the last iteration before iteration  $t$  where block  $j$  was selected. Then

$$\mathbb{E}[\|G_{t,j} - G_{\gamma(t,j),j}\|_{*,j}^2] \leq \delta L_G^2 \eta^2 \sum_{s=\gamma(t,j)}^{t-1} \mathbb{E}[\|Z_s\|_*^2] \quad (\text{A.20})$$

**Proof.** Since  $t > 1$ , block  $j$  has been selected at least once before iteration  $t$ ,  $\gamma(t, j)$  is well defined. Using (2.5) and the definition of the normalization operator  $\mathcal{N}(\cdot)$ , we have that

$$\|X_{s+1,\ell} - X_{s,\ell}\|_\ell = \eta \|Z_s\|_{*,\ell} \|\mathcal{N}(Z_s)\|_\ell = \eta \|Z_s\|_{*,\ell},$$

Therefore, using also (2.2), Assumption 2 and Assumption 7, we derive that

$$\begin{aligned} \mathbb{E}[\|G_{t,j} - G_{\gamma(t,j),j}\|_{*,j}^2] &\leq \mathbb{E}[\|G_t - G_{\gamma(t,j)}\|_*^2] \\ &\leq L_G^2 \mathbb{E}[\|X_t - X_{\gamma(t,j)}\|^2] \\ &= L_G^2 \mathbb{E} \left[ \left\| \sum_{s=\gamma(t,j)}^{t-1} (X_{s+1} - X_s) \right\|^2 \right] \\ &\leq \delta L_G^2 \sum_{s=\gamma(t,j)}^{t-1} \mathbb{E}[\|X_{s+1} - X_s\|^2] \\ &= \delta L_G^2 \eta^2 \sum_{s=\gamma(t,j)}^{t-1} \mathbb{E}[\|Z_s\|_*^2], \end{aligned}$$

where the factor  $\delta$  follows from Cauchy's inequality and  $t - \gamma(t, j) \leq \delta$ .  $\square$

The second shows that the sum of full gradient norms can be bounded by an expression involving the sum of block-gradient norms at active iterations.

**Lemma A.6** Suppose that Assumptions 1, 2, 3, 5 and 7 hold. Then, for all  $t > \delta$ ,

$$\sum_{i=0}^t \mathbb{E}[\|G_i\|_*] \leq \delta L_G \eta \sqrt{L(L-1)(t+1)} \sqrt{\kappa_0 + 2N \log(\Theta_t)} + (\delta + 1) \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}]. \quad (\text{A.21})$$

**Proof.**

For every  $i \in \{0, \dots, t\}$ , we split the full gradient into its active and inactive block components. For inactive blocks, using the definition of  $\gamma(i, \ell)$ , we have

$$\mathbb{E}[\|G_i\|_*] \leq \sum_{\ell \in \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell}\|_{*,\ell}] + \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}] + \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{\gamma(i,\ell),\ell}\|_{*,\ell}].$$

Summing over  $i = 0, \dots, t$  gives that

$$\sum_{i=0}^t \mathbb{E}[\|G_i\|_*] \leq \sum_{i=0}^t \sum_{\ell \in \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell}\|_{*,\ell}] + \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}] + \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{\gamma(i,\ell),\ell}\|_{*,\ell}]. \quad (\text{A.22})$$

We first bound the last term of this inequality. Let  $\mathcal{S}_t = \{(i, \ell) \mid 0 \leq i \leq t, \ell \notin \mathcal{A}_i\}$ . Since, for each  $i$ , at most  $L - 1$  blocks are inactive,

$$|\mathcal{S}_t| = \sum_{i=0}^t |\{1, \dots, L\} \setminus \mathcal{A}_i| \leq (L - 1)(t + 1).$$

For  $(i, \ell) \in \mathcal{S}_t$ , define  $Y_{i,\ell} = \|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}$ . Applying the Cauchy-Schwarz inequality in to the random vectors  $(1)_{(i,\ell) \in \mathcal{S}_t}$  and  $(Y_{i,\ell})_{(i,\ell) \in \mathcal{S}_t}$ , we obtain that

$$\begin{aligned} \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[Y_{i,\ell}] &= \mathbb{E} \left[ \sum_{(i,\ell) \in \mathcal{S}_t} Y_{i,\ell} \right] \leq |\mathcal{S}_t|^{1/2} \left( \sum_{(i,\ell) \in \mathcal{S}_t} \mathbb{E}[Y_{i,\ell}^2] \right)^{1/2} \\ &\leq \sqrt{(L-1)(t+1)} \left( \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}^2] \right)^{1/2}. \end{aligned} \quad (\text{A.23})$$

Using (A.20), we further deduce that

$$\sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}^2] \leq \delta L_G^2 \eta^2 \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \sum_{s=\gamma(i,\ell)}^{i-1} \mathbb{E}[\|Z_s\|_*^2].$$

We now bound the triple sum in the right-hand side of this inequality. Fix  $\ell$  and  $s$ . The term  $\mathbb{E}[\|Z_s\|_*^2]$  can occur in the inner sum associated with an index  $i$  only if  $\gamma(i, \ell) \leq s \leq i - 1$ . In particular,  $s + 1 \leq i$ . Moreover, Assumption 7 gives that  $i - \gamma(i, \ell) \leq \delta$ , and therefore  $i - \delta \leq \gamma(i, \ell) \leq s$ , which implies  $i \leq s + \delta$ . Hence, for fixed  $\ell$  and  $s$ , the possible indices  $i$  satisfy  $s + 1 \leq i \leq s + \delta$ , and there are at most  $\delta$  such indices. Consequently,

$$\sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \sum_{s=\gamma(i,\ell)}^{i-1} \mathbb{E}[\|Z_s\|_*^2] \leq \delta L \sum_{s=0}^t \mathbb{E}[\|Z_s\|_*^2].$$

Therefore,

$$\sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}^2] \leq \delta^2 L_G^2 \eta^2 L \sum_{s=0}^t \mathbb{E}[\|Z_s\|_*^2].$$

Combining this estimate with (A.23) gives

$$\begin{aligned} \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}] &\leq \sqrt{(L-1)(t+1)} \left( \delta^2 L_G^2 \eta^2 L \sum_{s=0}^t \mathbb{E}[\|Z_s\|_*^2] \right)^{1/2} \\ &= \delta L_G \eta \sqrt{L(L-1)(t+1)} \left( \sum_{s=0}^t \mathbb{E}[\|Z_s\|_*^2] \right)^{1/2}. \end{aligned}$$

Using now (A.16), we obtain that

$$\sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell} - G_{\gamma(i,\ell),\ell}\|_{*,\ell}] \leq \delta L_G \eta \sqrt{L(L-1)(t+1)} \sqrt{\kappa_0 + 2N \log(\Theta_t)}. \quad (\text{A.24})$$

It remains to bound the second term of (A.22) that involves  $G_{\gamma(i,\ell),\ell}$ . Fix a block  $\ell$  and an index  $j \in \mathcal{J}_{t,\ell}$ , and define

$$\mathcal{N}_{j,\ell} = \{i \in \{0, \dots, t\} \mid \ell \notin \mathcal{A}_i \text{ and } \gamma(i, \ell) = j\}.$$

Thus,  $\mathcal{N}_{j,\ell}$  is the set of iterations  $i$  for which the term  $G_{j,\ell}$  appears in the sum through the identity  $G_{\gamma(i,\ell),\ell} = G_{j,\ell}$ . If  $i \in \mathcal{N}_{j,\ell}$ , then necessarily  $j < i$ , because  $\gamma(i, \ell)$  is the last activation of block  $\ell$  strictly before iteration  $i$ . Moreover, Assumption 7 implies  $i - \gamma(i, \ell) \leq \delta$ , and therefore  $1 \leq i - j \leq \delta$ . Hence

$$\mathcal{N}_{j,\ell} \subseteq \{j+1, \dots, j+\delta\} \quad \text{and} \quad |\mathcal{N}_{j,\ell}| \leq \delta.$$

Consequently,

$$\sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{\gamma(i,\ell),\ell}\|_{*,\ell}] = \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \sum_{i \in \mathcal{N}_{j,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] \leq \delta \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}].$$

Adding the active-block contribution yields that

$$\sum_{i=0}^t \sum_{\ell \in \mathcal{A}_i} \mathbb{E}[\|G_{i,\ell}\|_{*,\ell}] + \sum_{i=0}^t \sum_{\ell \notin \mathcal{A}_i} \mathbb{E}[\|G_{\gamma(i,\ell),\ell}\|_{*,\ell}] \leq (\delta + 1) \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}].$$

Substituting this inequality and (A.24) in (A.22) proves (A.21).  $\square$

This last theorem finally gives us all the ingredients to derive the convergence and complexity results of Theorem 3.2. The bound (A.13), together with Assumption 4, indeed implies convergence of the criticality measure on each block with a *uniform rate-of-convergence bound*. As shown by the next theorem, the expected global rate of convergence is (in order) proportional to the value of  $\Theta_t$ . The proof of this statement is substantially more involved than [18, Theorem 3.10] because of the introduction of delays.

## Proof of Theorem 3.2

Assumption 4 gives that, for each  $\ell \in \{1, \dots, L\}$ ,

$$\sum_{j \in \mathcal{J}_{t,\ell}} \|\tilde{G}_{j,\ell}\|_{*,\ell}^2 \leq \kappa_o^2 \sum_{j \in \mathcal{J}_{t,\ell}} \text{tr}(\mathcal{U}_{j,\ell}(\tilde{G}_{j,\ell})^2) = \kappa_o^2 \text{tr}\left(\sum_{j \in \mathcal{J}_{t,\ell}} \mathcal{U}_{j,\ell}(\tilde{G}_{j,\ell})^2\right) \leq \kappa_o^2 \text{tr}(\Pi_{t,\ell})$$

But

$$\mathrm{tr}(\Pi_{t,\ell}) = \sum_{i=1}^{d_\ell} \left( \sqrt{\lambda_i[\Pi_{t,\ell}]} \right)^2 \leq \left( \sum_{i=1}^{d_\ell} \sqrt{\lambda_i[\Pi_{t,\ell}]} \right)^2 = \mathrm{tr}(\Pi_{t,\ell}^{1/2})^2,$$

and thus, summing over  $\ell \in \{1, \dots, L\}$  and using the fact that  $\sum_i a_i^2 \leq (\sum_i a_i)^2$ ,

$$\sqrt{\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \|\tilde{G}_{j,\ell}\|_{*,\ell}^2} \leq \kappa_o \sqrt{\sum_{\ell=1}^L \mathrm{tr}(\Pi_{t,\ell}^{1/2})^2} \leq \kappa_o \sqrt{\left( \sum_{\ell=1}^L \mathrm{tr}(\Pi_{t,\ell}^{1/2}) \right)^2} = \kappa_o \sum_{\ell=1}^L \mathrm{tr}(\Pi_{t,\ell}^{1/2}).$$

Using now the Cauchy-Schwarz and Jensen's inequalities with this last bound, we deduce that

$$\begin{aligned} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|\tilde{G}_{j,\ell}\|_{*,\ell}] &= \mathbb{E} \left[ \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \|\tilde{G}_{j,\ell}\|_{*,\ell} \right] \\ &\leq \sqrt{L \max_{\ell \in \{1, \dots, L\}} |\mathcal{J}_{t,\ell}|} \mathbb{E} \left[ \sqrt{\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \|\tilde{G}_{j,\ell}\|_{*,\ell}^2} \right] \\ &\leq \kappa_o \sqrt{L(t+1)} \mathbb{E} \left[ \sum_{\ell=1}^L \mathrm{tr}(\Pi_{t,\ell}^{1/2}) \right] \\ &\leq \kappa_o \sqrt{L(t+1)} \Theta_t \end{aligned} \tag{A.25}$$

where the last inequality results from (A.13). As a consequence, we have that

$$\begin{aligned} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] &\leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell} - \tilde{G}_{j,\ell}\|_{*,\ell}] + \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|\tilde{G}_{j,\ell}\|_{*,\ell}] \\ &\leq \sqrt{L(t+1)} \sqrt{\nu_t^2 + \omega^2 \sum_{j=0}^t \sum_{\ell \in \mathcal{A}_j} \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2]} + \kappa_o \sqrt{L(t+1)} \Theta_t. \end{aligned}$$

Using now (A.16), we obtain that

$$\begin{aligned} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] &\leq \sqrt{L(t+1)} \left( \sqrt{\nu_t^2 + \omega^2(\kappa_0 + 2N \log(\Theta_t))} + \kappa_o \Theta_t \right) \\ &\leq \sqrt{L(t+1)} \left( \nu_t + \omega \sqrt{\kappa_0} + \omega \sqrt{2N \log(\Theta_t)} + \kappa_o \Theta_t \right). \end{aligned} \tag{A.26}$$

Combining this inequality with (A.21) then ensures that

$$\begin{aligned} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] &\leq \delta L_G \eta \sqrt{L(L-1)(t+1)} \left( \sqrt{\kappa_0} + \sqrt{2N \log(\Theta_t)} \right) \\ &\quad + (\delta + 1) \sqrt{L(t+1)} \left( \nu_t + \omega \sqrt{\kappa_0} + \omega \sqrt{2N \log(\Theta_t)} + \kappa_o \Theta_t \right). \end{aligned}$$

Dividing by  $t+1$  yields (3.11).

Suppose now that Assumption 6 also holds. Then

$$\mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] \leq \mathbb{E}[\|\tilde{G}_{j,\ell}\|_{*,\ell}],$$

and therefore (A.25), combined with (A.21), gives (3.15).



### Proof of Corollary 3.3

Observe that, when  $\Theta_t$  (as defined in Theorem A.4) grows with  $t$ , the bound (3.11) is dominated by the term  $\kappa_c \Theta_t / \sqrt{t+1}$ . Since Corollary 3.3 solely aims at describing dominating terms for growing  $t$ , we may ignore the other terms and merely consider that (3.11) reduces to

$$\frac{1}{t+1} \sum_{j=0}^t \mathbb{E}[\|G_j\|_*] \leq \frac{\kappa_c \Theta_t}{\sqrt{t+1}}, \quad (\text{A.27})$$

which is identical to the result (equation (3.42)) of Theorem 3.10 in [18]. The rest of the proof then follows that of [18, Corollary 3.11] with  $\nu_k^2$  replaced by

$$\max_{\ell \in \{1, \dots, L\}} \sum_{j \in \mathcal{J}_{t,\ell}} \frac{\sigma_\ell^2}{|\mathcal{J}_{j,\ell}|^\alpha} = \max_{\ell \in \{1, \dots, L\}} \sigma_\ell^2 \sum_{j=0}^{|\mathcal{J}_{t,\ell}|-1} \frac{1}{(j+1)^\alpha} \leq \begin{cases} \frac{\sigma_{\max}^2}{1-\alpha} \psi_t^{1-\alpha} & \text{if } \alpha < 1, \\ \sigma_{\max}^2 \log(\psi_t) & \text{if } \alpha = 1, \\ \sigma_{\max}^2 \zeta(\alpha) & \text{if } \alpha > 1, \end{cases}$$

where  $\sigma_{\max} = \max_{\ell \in \{1, \dots, L\}} \sigma_\ell$ . This then yields (3.17).

## Appendix 2: Detailed analysis for DADPREC with momentum

We start by establishing a bound on the norm of the difference between the block-wise momentum  $M_{t,\ell}$  and the approximate gradient  $\tilde{G}_{t,\ell}$ .

**Lemma A.7** Consider Algorithms DADPREC.MM and DADPREC.MG. Suppose that Assumptions 2 and 7 hold. Define

$$E_{t,\ell} = M_{t,\ell} - \tilde{G}_{t,\ell} \quad \text{for } t \geq 0 \text{ and } \ell \in \mathcal{A}_t. \quad (\text{A.28})$$

Then

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|E_{j,\ell}\|_{*,\ell}^2] \leq \frac{3}{(1-\mu_{\max})^2} \left[ 2 \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \bar{\mu}_{j,\ell}^2 \mathbb{E}[\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2] + \delta^2 L L_G^2 \eta^2 \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \right] \quad (\text{A.29})$$

where  $\bar{\mu}_{r,\ell} = \max\{\mu_{r,\ell}, \mu_{\varphi(r,\ell),\ell}\}$  with  $\varphi(r,\ell) = \min[j \in \mathcal{J}_{t,\ell} \mid j > r]$ .

**Proof.** Suppose that  $\ell \in \mathcal{A}_t$  and  $t > 1$  (so that  $\gamma(t,\ell)$  is well defined). The definition of the momentum, (A.19) and (A.28) then give that, for  $\ell \in \mathcal{A}_t$ ,

$$\begin{aligned} E_{t,\ell} &= \mu_{t,\ell} M_{\gamma(t,\ell),\ell} + (1 - \mu_{t,\ell}) \tilde{G}_{t,\ell} - \tilde{G}_{t,\ell} \\ &= \mu_{t,\ell} (M_{\gamma(t,\ell),\ell} - \tilde{G}_{\gamma(t,\ell),\ell}) + \mu_{t,\ell} (\tilde{G}_{\gamma(t,\ell),\ell} - G_{\gamma(t,\ell),\ell}) + \mu_{t,\ell} (G_{\gamma(t,\ell),\ell} - \tilde{G}_{t,\ell}) \\ &= \mu_{t,\ell} E_{\gamma(t,\ell),\ell} + \mu_{t,\ell} (\tilde{G}_{\gamma(t,\ell),\ell} - G_{\gamma(t,\ell),\ell}) + \mu_{t,\ell} (G_{\gamma(t,\ell),\ell} - \tilde{G}_{t,\ell}). \end{aligned}$$

Thus, using the facts that  $\mu_{t,\ell} \leq \mu_{\max} < 1$ , that  $(a+b)^2 \leq (1+\rho)a^2 + (1+1/\rho)b^2$  for any  $\rho > 0$  and that  $(a+b+c)^2 \leq 3a^2 + 3b^2 + 3c^2$ , we obtain that

$$\begin{aligned} \|E_{t,\ell}\|_{*,\ell}^2 &\leq (1+\rho) \mu_{\max}^2 \|E_{\gamma(t,\ell),\ell}\|_{*,\ell}^2 \\ &\quad + 3 \left(1 + \frac{1}{\rho}\right) \mu_{t,\ell}^2 \left[ \|G_{t,\ell} - G_{\gamma(t,\ell),\ell}\|_{*,\ell}^2 + \|\tilde{G}_{\gamma(t,\ell),\ell} - G_{\gamma(t,\ell),\ell}\|_{*,\ell}^2 + \|\tilde{G}_{t,\ell} - G_{t,\ell}\|_{*,\ell}^2 \right]. \end{aligned}$$

Now let  $\rho = (1 - \mu_{\max})/\mu_{\max}$ . Then  $(1+\rho)\mu_{\max}^2 = \mu_{\max} < 1$  and  $(1+1/\rho) = 1/(1 - \mu_{\max})$ .

Hence, summing over  $j \in \mathcal{J}_{t,\ell}$ , we obtain that

$$\begin{aligned} \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \|E_{j,\ell}\|_{*,\ell}^2 &\leq \mu_{\max} \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \|E_{\gamma(j,\ell),\ell}\|_{*,\ell}^2 \\ &+ \frac{3}{1 - \mu_{\max}} \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{j,\ell}^2 \left[ \|G_{j,\ell} - G_{\gamma(j,\ell),\ell}\|_{*,\ell}^2 + \|\tilde{G}_{\gamma(j,\ell),\ell} - G_{\gamma(j,\ell),\ell}\|_{*,\ell}^2 + \|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2 \right] \end{aligned}$$

Observe now that (2.2), Assumptions 2 and 7 with (2.5) and the definition of the normalization operator imply (analogously to the derivation in Lemma A.5) that

$$\begin{aligned} \|G_{j,\ell} - G_{\gamma(j,\ell),\ell}\|_{*,\ell}^2 &\leq \|G_j - G_{\gamma(j,\ell)}\|_*^2 \\ &\leq L_G^2 \sum_{q=1}^L \|X_{j,q} - X_{\gamma(j,\ell),q}\|_q^2 \\ &\leq \delta L_G^2 \sum_{q=1}^L \sum_{s=\gamma(j,\ell)}^{j-1} \|X_{s+1,q} - X_{s,q}\|_q^2 \\ &= \delta L_G^2 \eta^2 \sum_{q=1}^L \sum_{s=\gamma(j,\ell)}^{j-1} \|Z_{s,q}\|_{*,q}^2. \end{aligned} \tag{A.30}$$

But Assumption 7 implies that the sum over  $s$  in the last right-hand side contains at most  $\delta$  terms. Thus, using also (2.2) and the fact that  $\mu_{j,\ell} < 1$ ,

$$\sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{j,\ell}^2 \|G_{j,\ell} - G_{\gamma(j,\ell),\ell}\|_{*,\ell}^2 \leq \delta^2 L_G^2 \eta^2 \sum_{q=1}^L \sum_{s=0}^{t-1} \|Z_{s,q}\|_{*,q}^2 = \delta^2 L_G^2 \eta^2 \sum_{s=0}^{t-1} \|Z_s\|_*^2.$$

Summing now over  $\ell \in \{1, \dots, L\}$ , taking full expectation and defining  $\vartheta_{t,\ell}^2 = \mathbb{E}[\|\tilde{G}_{t,\ell} - G_{t,\ell}\|_{*,\ell}^2]$ , we deduce that

$$\begin{aligned} \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mathbb{E}[\|E_{j,\ell}\|_{*,\ell}^2] &\leq \mu_{\max} \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mathbb{E}[\|E_{\gamma(j,\ell),\ell}\|_{*,\ell}^2] \\ &+ \frac{3}{1 - \mu_{\max}} \left[ \delta L_G^2 \eta^2 \sum_{\ell=1}^L \sum_{s=0}^{t-1} \mathbb{E}[\|Z_s\|_*^2] + \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{j,\ell}^2 (\vartheta_{\gamma(j,\ell)}^2 + \vartheta_{j,\ell}^2) \right] \end{aligned}$$

Now  $\varphi(\gamma(j,\ell),\ell) = j$  for  $j \leq t$ . Therefore

$$\sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{j,\ell}^2 \vartheta_{\gamma(j,\ell),\ell}^2 = \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{\varphi(\gamma(j,\ell),\ell)}^2 \vartheta_{\gamma(j,\ell),\ell}^2,$$

yielding

$$\begin{aligned} \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mathbb{E}[\|E_{j,\ell}\|_{*,\ell}^2] &\leq \mu_{\max} \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mathbb{E}[\|E_{j,\ell}\|_{*,\ell}^2] \\ &+ \frac{3}{1 - \mu_{\max}} \left[ \delta^2 L L_G^2 \eta^2 \sum_{s=0}^{t-1} \mathbb{E}[\|Z_s\|_*^2] + 2 \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mu_{j,\ell}^2 \vartheta_{j,\ell}^2 \right]. \end{aligned}$$

Hence,

$$\sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \mathbb{E}[\|E_{j,\ell}\|_{*,\ell}^2] \leq \frac{3}{(1-\mu_{\max})^2} \left[ \delta^2 L L_G^2 \eta^2 \sum_{s=0}^{t-1} \mathbb{E}[\|Z_s\|_*^2] + 2 \sum_{\ell=1}^L \sum_{\substack{j \in \mathcal{J}_{t,\ell} \\ t > 0}} \bar{\mu}_{j,\ell}^2 \vartheta_{j,\ell}^2 \right]$$

which, with the fact that  $E_{0,\ell} = 0$ , concludes the proof of (A.29).  $\square$

### Proof of Theorem 4.1

Observe now that the DADPREC.MM algorithm is nothing but Algorithm DADPREC where the momentum  $M_{t,\ell}$  plays the role of the approximate gradient  $\tilde{G}_{t,\ell}$ . Moreover Assumption 8 exactly reflects this change of perspective. Also note that, for all  $\ell \in \{1, \dots, L\}$  and  $j \in \mathcal{J}_{t,\ell}$ ,

$$\|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2 \leq 2\|E_{j,\ell}\|_{*,\ell}^2 + 2\|\tilde{G}_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2.$$

Hence, using Lemma A.7 and given that  $\bar{\mu}_{j,\ell} < 1$ , we have that

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2] < \left( \frac{12}{(1-\mu_{\max})^2} + 2 \right) \sum_{j=0}^t \mathbb{E}[\|\tilde{G}_j - G_j\|_*^2] + \frac{6\delta^2 L L_G^2 \eta^2}{(1-\mu_{\max})^2} \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2]. \quad (\text{A.31})$$

This provides an adapted formulation to replace (3.4) in Assumption 5, which thus continues to hold with

$$\nu_t^2 = \left( \frac{3}{(1-\mu_{\max})^2} + 2 \right) \sum_{j=0}^t \mathbb{E}[\|\tilde{G}_j - G_j\|_*^2] \quad \text{and} \quad \omega^2 = \frac{2\delta^2 L L_G^2 \eta^2}{(1-\mu_{\max})^2}. \quad (\text{A.32})$$

As a consequence, the theory developed for the DADPREC algorithm also holds for its DADPREC.MM variant and we obtain from (A.25) that

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] \leq \kappa_0 \sqrt{L(t+1)} \Theta_t. \quad (\text{A.33})$$

Now, using the triangle inequality, the Cauchy-Schwarz and Jensen inequalities, we obtain that

$$\begin{aligned} \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] &\leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] + \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}] \\ &\leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] + \sqrt{L(t+1)} \sqrt{\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}]^2} \\ &\leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] + \sqrt{L(t+1)} \sqrt{\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2]} \\ &= \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] + \sqrt{L(t+1)} \sqrt{\nu_t^2 + \omega^2 \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2]} \end{aligned} \quad (\text{A.34})$$

where we have inserted (A.32) in (A.31) to obtain the last equality. Hence, using (A.16), we deduce that

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|G_{j,\ell}\|_{*,\ell}] \leq \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|M_{j,\ell}\|_{*,\ell}] + \sqrt{L(t+1)} \sqrt{\nu_t^2 + \omega^2 (\kappa_0 + 2N \log(\Theta_t))}.$$

Finally, (4.3) follows by substitution (A.33) into this last bound, using  $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ , applying Lemma A.6 and dividing by  $t+1$ . The conclusions of Corollary 3.3 continue to hold because the order of convergence (as a function of  $t$ ) is unmodified.

## Proof of Theorem 4.2

This section considers the convergence of the DADPREC.MG algorithm. We first recall results from [18] on the impact of using this algorithm on the structural identities of Assumption 8.

**Lemma A.8** Suppose that Assumptions 2, 7 and 8 and (4.4) hold. Then

$$\|Z_{t,\ell}\|_{*,\ell} \langle G_{t,\ell}, \mathcal{N}_{t,\ell}(Z_{t,\ell}) \rangle_F \geq \frac{1}{\kappa_\square} \text{tr}(\Pi_{t,\ell}^{-1/2} \mathcal{U}_{t,\ell}(\tilde{G}_{t,\ell})^2) - \text{tr}(\Pi_{t,\ell}^{-1/2} \mathcal{U}_{t,\ell}(E_{t,\ell})^2) - \|Z_{t,\ell}\|_{*,\ell} \|E_{t,\ell}\|_{*,\ell} \quad (\text{A.35})$$

and

$$\|Z_{t,\ell}\|_{*,\ell}^2 \leq \kappa_\square \text{tr}(\Pi_{t,\ell}^{-1} \mathcal{U}_{t,\ell}(G_{t,\ell})^2) + \kappa_\square \text{tr}(\Pi_{t,\ell}^{-1} \mathcal{U}_{t,\ell}(E_{t,\ell})^2) \quad (\text{A.36})$$

**Proof.** See [18, Lemmas A.1 and A.2]. □

**Lemma A.9** Suppose that Assumption 2, 7 and 8 and (4.4) hold. Then

$$\sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell} \|E_{j,\ell}\|_{*,\ell}] \leq \frac{12\theta_t^2}{(1-\mu_{\max})^2} + \left( \frac{6\delta^2 L L_G^2 \eta^2}{(1-\mu_{\max})^2} + 2 \right) \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2], \quad (\text{A.37})$$

$$\sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1/2} \mathcal{U}_{j,\ell}(E_{j,\ell})^2)] \leq \frac{6\kappa_\diamond \theta_t^2}{(1-\mu_{\max})^2 \sqrt{\varsigma}} + \frac{3\kappa_\diamond \delta^2 L L_G^2 \eta^2}{(1-\mu_{\max})^2 \sqrt{\varsigma}} \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \quad (\text{A.38})$$

and

$$\sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell}(E_{j,\ell})^2)] \leq \frac{6\kappa_\diamond \theta_t^2}{(1-\mu_{\max})^2 \varsigma} + \frac{3\kappa_\diamond \delta^2 L L_G^2 \eta^2}{(1-\mu_{\max})^2 \varsigma} \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2]. \quad (\text{A.39})$$

**Proof.** See [18, Lemma A.3]. □

**Lemma A.10** Suppose that Assumption 2 and 8 and (4.4) hold. Suppose in addition that (4.5) is satisfied. Then

$$\begin{aligned} \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell} \langle G_{j,\ell}, \mathcal{N}_{j,\ell}(Z_{j,\ell}) \rangle_F] &\geq \frac{1}{\kappa_{\square}} \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1/2} \mathcal{U}_{j,\ell}(\tilde{G}_{j,\ell})^2)] \\ &\quad - \kappa_{1,\nu} \theta_t^2 - \kappa_{1,z} \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \end{aligned} \quad (\text{A.40})$$

with

$$\kappa_{1,\nu} = \frac{6\kappa_{\diamond}}{(1-\mu_{\max})^2\sqrt{\varsigma}} + \frac{12}{(1-\mu_{\max})^2} \quad \text{and} \quad \kappa_{1,z} = \frac{3\kappa_{\diamond}\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2\sqrt{\varsigma}} + \frac{6\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2} + 2, \quad (\text{A.41})$$

and

$$\sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \leq 2 \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell}(G_{j,\ell})^2)] + \kappa_{2,\nu} \theta_t^2 + \kappa_{2,z}, \quad (\text{A.42})$$

where

$$\kappa_{2,\nu} = \frac{12\kappa_{\square}\kappa_{\diamond}}{(1-\mu_{\max})^2\varsigma} \quad \text{and} \quad \kappa_{2,z} = \begin{cases} 0 & \text{if the first part of (4.5) holds,} \\ \frac{6\kappa_{\square}\kappa_{\diamond}\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2\varsigma} \kappa_{\mu Z} & \text{if the second part of (4.5) holds.} \end{cases} \quad (\text{A.43})$$

**Proof.** The inequality (A.40) is obtained by substituting (A.37) and (A.38) into (A.35). Moreover, taking the expectation in (A.36), summing for  $j \in \{0, \dots, t\}$  and  $\ell \in \{1, \dots, L\}$  and substituting (A.39) gives that

$$\begin{aligned} \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2] &\leq \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell}(G_{j,\ell})^2)] + \frac{6\kappa_{\square}\kappa_{\diamond}}{(1-\mu_{\max})^2\varsigma} \theta_t^2 \\ &\quad + \frac{3\kappa_{\square}\kappa_{\diamond}\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2\varsigma} \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2]. \end{aligned}$$

Suppose first that the first part of (4.5) holds. Then

$$\frac{3\kappa_{\square}\kappa_{\diamond}\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2\varsigma} \leq \frac{1}{2}$$

and (A.42) follows with  $\kappa_{2,z} = 0$ . Alternatively, if the second part of (4.5) holds, then (A.42) follows with

$$\kappa_{2,z} = \frac{6\kappa_{\square}\kappa_{\diamond}\delta^2 LL_G^2 \eta^2}{(1-\mu_{\max})^2\varsigma} \kappa_{\mu Z}.$$

□

From this point on, the analysis follows the lines of the argument of Appendix 1 with Lemma A.10 providing an alternate set of structural inequalities. Lemma A.2 is unchanged. The modifications to Theorem A.3 are minor. It is restated as follows.

**Theorem A.11** Suppose that Assumption 1, 2, 7 and 8 hold and that (4.4) and (4.5) hold. Define  $\Delta_j$  as in (A.4). Then, for every  $k \geq 0$ ,

$$\eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) \right] \leq \kappa_{\text{gap}} + \kappa_{\nu\nu} \theta_t^2 + \eta \kappa_{\nu\Delta} \theta_t \sqrt{\Delta_t} + \eta \kappa_{\Delta} \Delta_t, \quad (\text{A.44})$$

where

$$\kappa_{\text{gap}} \stackrel{\text{def}}{=} \mathbb{E}[f(X_0)] - f_{\text{low}} + \eta \varsigma N + \eta \sqrt{\kappa_{2,z}} + \left( \eta \kappa_{1,z} + \frac{L_G \eta^2}{2} \right) \kappa_{2,z} \quad (\text{A.45})$$

$$\kappa_{\nu\nu} = \eta \left( \kappa_{1,\nu} + \sqrt{\kappa_{2,\nu}} + \kappa_{1,z} \kappa_{2,\nu} + \frac{L_G \eta}{2} \right), \quad \kappa_{\nu\Delta} = \sqrt{2} \quad \text{and} \quad \kappa_{\Delta} = 2\kappa_{1,z} + L_G \eta. \quad (\text{A.46})$$

**Proof.** In the proof of Theorem A.3, a perturbation  $\eta \kappa_{1,\nu} \theta_t^2 + \eta \kappa_{1,z} \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\|Z_{t,\ell}\|_{*,\ell}^2]$  is subtracted from the right-hand side of (A.7) in order to reflect (A.40). The proof then goes on unmodified, up to the substitution leading to (A.12), in which the perturbed (A.7) then gives that

$$\eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) - \sum_{\ell=1}^L \text{tr}(\Pi_{-1,\ell}^{1/2}) \right] \leq f(X_0) - f_{\text{low}} + \eta \kappa_{1,\nu} \theta_t^2 + \eta \theta_t \sqrt{\zeta_t} + \left( \eta \kappa_{1,z} + \frac{L_G \eta^2}{2} \right) \zeta_t.$$

where  $\zeta_t = \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E}[\|Z_{j,\ell}\|_{*,\ell}^2]$ . But this definition, (A.42) and (A.3) give that

$$\zeta_t \leq \kappa_{2,z} + \kappa_{2,\nu} \theta_t^2 + 2 \sum_{j=0}^t \sum_{\ell=1}^L \mathbb{E} \left[ \text{tr}(\Pi_{j,\ell}^{-1} \mathcal{U}_{j,\ell} (G_{j,\ell})^2) \right] \leq \kappa_{2,z} + \kappa_{2,\nu} \theta_t^2 + 2\Delta_t,$$

and thus, using  $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ ,

$$\begin{aligned} \eta \mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) - \sum_{\ell=1}^L \text{tr}(\Pi_{-1,\ell}^{1/2}) \right] &\leq f(X_0) - f_{\text{low}} + \eta \sqrt{\kappa_{2,z}} + \eta \left( \kappa_{1,\nu} + \sqrt{\kappa_{2,\nu}} + \kappa_{1,z} \kappa_{2,\nu} + \frac{L_G \eta}{2} \right) \theta_t^2 \\ &\quad + \sqrt{2} \eta \theta_t \sqrt{\Delta_t} + 2 \left( \eta \kappa_{1,z} + \frac{L_G \eta^2}{2} \right) \Delta_t + \left( \eta \kappa_{1,z} + \frac{L_G \eta^2}{2} \right) \kappa_{2,z} \end{aligned}$$

yielding (A.44)-(A.46) after taking into account that  $\Pi_{-1,\ell} = \varsigma I_\ell$  for all  $\ell \in \{1, \dots, L\}$ .  $\square$

The form of bound (A.44) differs very little from that of (A.5): besides using different constants, (A.44) now involves a term in  $\kappa_{\nu\nu} \theta_t^2$ . This term then percolates through the proof of Theorem A.4, so that (A.13) now becomes

$$\mathbb{E} \left[ \sum_{\ell=1}^L \text{tr}(\Pi_{t,\ell}^{1/2}) \right] \leq \Theta_t \stackrel{\text{def}}{=} \max \left[ \kappa_{\Theta}, \frac{3\kappa_{\nu\nu} \theta_t^2}{\eta}, T_t \right]. \quad (\text{A.47})$$

The proof of Theorem 3.2 then stands without modification other than using this updated value of  $\Theta_t$ , finally yielding Theorem 4.2.

### Proof of Corollary 4.3

Finally, the proof of Corollary 4.3 is similar in spirit to that of Corollary 3.11 in [18], with two twists. The first is that the formula for the block-wise variance is now given by (3.16), which,

together with (4.9), implies<sup>10</sup> that

$$\begin{aligned}
 \sum_{\ell=1}^L \sum_{j \in \mathcal{J}_{t,\ell}} \|M_{j,\ell} - G_{j,\ell}\|_{*,\ell}^2 &\leq \mu_{\max} \sum_{\ell=1}^L \sigma_{\ell}^2 \sum_{j \in \mathcal{J}_{t,\ell}} \frac{1}{\min[|\mathcal{J}_{j,\ell}|, |\mathcal{J}_{j-1,\ell}|]^{\alpha+2\beta}} + \omega^2 \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \\
 &\leq \mu_{\max} \sum_{\ell=1}^L \sigma_{\ell}^2 \sum_{j \in \mathcal{J}_{t,\ell}} \frac{1}{|\mathcal{J}_{j-1,\ell}|^{\alpha+2\beta}} + \omega^2 \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2] \\
 &\leq \mu_{\max} \sigma_{\max}^2 \max_{\ell \in \{1, \dots, L\}} \left[ \sum_{j=1}^t \frac{1}{|\mathcal{J}_{j-1,\ell}|^{\alpha+2\beta}} \right] + \omega^2 \sum_{j=0}^t \mathbb{E}[\|Z_j\|_*^2].
 \end{aligned}$$

Therefore, using  $\psi_{t-1} \leq \psi_t$ , we may set the bound

$$\theta_t^2 = \begin{cases} \frac{\sigma_{\max}^2}{1-\alpha-2\beta} \psi_t^{1-\alpha-2\beta} & \text{if } \alpha + 2\beta < 1, \\ \sigma_{\max}^2 \log(\psi_t) & \text{if } \alpha + 2\beta = 1, \\ \sigma_{\max}^2 \zeta(\alpha + 2\beta) & \text{if } \alpha + 2\beta > 1. \end{cases} \quad (\text{A.48})$$

The second is the presence of term  $3\kappa_{\nu\nu}\theta_t^2$  in (4.7). Since the dominant term in the convergence bound (4.6) is the last, the final order of convergence is determined by the maximal power of  $\theta_t$  occurring in the expression of  $\Theta_t$ , which is now  $\theta_t^2$  because of this new term. Thus substituting (A.48) in (4.6) and only keeping the dominant term in  $t$  gives (4.10).  $\square$

## Appendix 3: Implementation and additional results for asynchronous momentum variants

We first describe our warm-started power iteration to compute an approximation of the nuclear norm of  $G$ . It is based on the proxy

$$\|G\|_* \approx \frac{\|G\|_F^2}{\sigma_{\max}[G]}$$

where  $\sigma_{\max}[G]$  is the largest singular value of  $G$ . The norm  $\|G\|_F^2$  is easily computable, and we estimate  $\sigma_{\max}[G]$  using two iterations of the power iterations on  $G$ . This does not ensure a convergent estimator, but gives a deterministic lower bound

$$\frac{\|G\|_F^2}{\sigma_{\max}[G]} = \frac{\sum_{i=1}^n \sigma_i[G]^2}{\sigma_{\max}[G]} \leq \sum_{i=1}^n \sigma_i[G] = \|G\|_*,$$

where equality holds if and only if all singular values are identical. Our algorithm is warm-started: we initialize the power iteration with the result obtained at the previous optimization iteration. As the gradient  $G$  changes progressively, so does its dominant singular vector, and two iterations are typically enough to obtain an accurate estimate  $\text{est}(\|G\|_*)$ . The detailed algorithm is given on the following page.

The cost of using **WSPower** for an  $m \times n$  matrix amounts to  $\mathcal{O}(4mn)$  flops for the four matrix-vector products, plus  $\mathcal{O}(mn)$  flops for  $\|G\|_F^2$ , giving a total of  $\mathcal{O}(5mn)$  flops, compared to  $\mathcal{O}(mn \min(m, n))$  flops for the SVD. In our experiments with  $256 \times 2048$  matrices, using **WSPower** is  $\min(m, n)/5 \approx 50$  times faster.

Although a mere proxy of  $\|G\|_*$ , **WSPower** works remarkably well, but its warm-start feature is crucial. Without it, accuracy results diverge on larger models (SVHN: 56%, MoE-FMNIST:

<sup>10</sup>With the convention that  $|\mathcal{J}_{-1,\ell}| = 1$  for all  $\ell \in \{1, \dots, L\}$ .

**Algorithm A.1: WSPower**

Given: a starting vector  $v$  resulting from the previous optimization iteration (a random vector at the first).

1. Compute  $u = Gv$ ,  $u = u/\|u\|$ ,  $v = G^T u$ ,  $\sigma_{\max} = \|v\|$ ,  $v = v/\|v\|$  (repeat twice)
2. Set  $\text{est}(\|G\|_*) = \frac{1}{\sigma_{\max}} \sum_{i=1}^m \sum_{j=1}^m G_{i,j}^2$ ,
3. Save  $v$  for the next call.

77%, to compare with SVHN  $70.44\% \pm 0.43$ , MoE  $88.63\% \pm 0.34$  with warm start, this last result being the best achieved for this benchmark in all our runs) The throughput of our asynchronous MG optimization variants using WSPower then become comparable to that of MM variants, as shown in Table A.10 below.

We now complete the results mentioned in Section 6.5. Table A.8 gives the final loss values obtained when using different asynchronous momentum variants.

| Dataset       | No Mom. | MMconst      | MGconst | MGconstA     | MMdecay      | MGdecay | MGdecayA |
|---------------|---------|--------------|---------|--------------|--------------|---------|----------|
| Fashion-MNIST | 0.308   | <b>0.239</b> | 0.336   | 0.259        | 0.267        | 0.316   | 0.273    |
| MoE-FMNIST    | 0.279   | 0.240        | 0.245   | <b>0.186</b> | 0.290        | 0.296   | 0.288    |
| CovType       | 0.240   | <b>0.149</b> | 0.226   | 0.173        | 0.225        | 0.256   | 0.229    |
| MovieLens     | 0.353   | <b>0.146</b> | 0.347   | 0.209        | 0.344        | 0.397   | 0.305    |
| Criteo        | 0.088   | <b>0.008</b> | 0.107   | 0.027        | 0.081        | 0.126   | 0.041    |
| SVHN          | 0.819   | 0.856        | 0.954   | 0.798        | <b>0.762</b> | 0.932   | 0.809    |

Table A.8: Final loss value for asynchronous momentum variants

We see that the MMconst variant dominates on CovType (0.149), Criteo (0.008!), Fashion-MNIST (0.239) and MovieLens (0.146), while MGconstA is best for MoE-FMNIST and MGdecay for SVHN.

The generalization metrics for the different variants with momentum are given in Table A.9.

| Dataset       | Metric  | No Mom. | MMconst      | MGconst      | MGconstA     | MMdecay      | MGdecay | MGdecayA |
|---------------|---------|---------|--------------|--------------|--------------|--------------|---------|----------|
| Fashion-MNIST | acc.(%) | 87.34   | 88.16        | 86.31        | 87.86        | <b>88.19</b> | 87.00   | 87.92    |
| MoE-FMNIST    | acc.(%) | 87.39   | 87.97        | 88.11        | <b>88.63</b> | 86.93        | 86.99   | 87.11    |
| CovType       | acc.(%) | 89.56   | <b>92.71</b> | 90.01        | 91.92        | 90.13        | 88.88   | 90.07    |
| MovieLens     | RMSE    | 0.916   | 0.992        | <b>0.914</b> | 0.956        | 0.915        | 0.923   | 0.926    |
| Criteo        | AUC     | 0.717   | 0.720        | <b>0.728</b> | 0.716        | 0.716        | 0.723   | 0.708    |
| SVHN          | acc.(%) | 69.39   | 68.69        | 65.10        | 70.44        | <b>71.55</b> | 66.19   | 69.45    |

Table A.9: Generalization metrics for asynchronous momentum variants

Care should, however, be taken: compared to the momentum-less option, the loss for MovieLens improves from 0.353 to 0.008, but the RMSE worsens from 0.916 to 0.992, illustrating a possible overfitting in that case.

The throughputs per problem are given in Table A.10, which shows that the introduction of the MM momentum has a limited effect on throughput, but that the MG approach has a clear negative impact, due to the double decomposition that we pinpointed in Section 6.5. Fortunately, the use of the warm-started power iteration to approximate the nuclear norm of  $\tilde{G}_\ell$  essentially corrects this deficiency.

Finally focussing on the classification problems, we report, in Table A.11, the variance of the accuracy obtained for the 8 independent runs. This table suggests that the variance on accuracy is nearly always better for the MM momentum than for the MG approach. In particular, a comparison of MMconst with MGconst shows that the former improves the accuracy of all classification problems by as much as 3.58% (for SVHN).



| Dataset       | No Mom. | MMconst | MGconst | MGconstA | MMdecay | MGdecay | MGdecayA |
|---------------|---------|---------|---------|----------|---------|---------|----------|
| Fashion-MNIST | 10      | 9       | 5       | 9        | 10      | 6       | 8        |
| MoE-FMNIST    | 42      | 34      | 18      | 34       | 28      | 21      | 27       |
| CovType       | 105     | 95      | 56      | 99       | 105     | 67      | 101      |
| MovieLens     | 89      | 90      | 57      | 84       | 79      | 50      | 87       |
| Criteo        | 89      | 84      | 72      | 93       | 82      | 58      | 96       |
| SVHN          | 1       | 1       | 1       | 1        | 1       | 1       | 1        |

Table A.10: Throughput for asynchronous momentum variants

| Dataset       | No Mom.             | MMconst    | MGconst    | MGconstA            | MMdecay             | MGdecay    | MGdecayA   |
|---------------|---------------------|------------|------------|---------------------|---------------------|------------|------------|
| Fashion-MNIST | $\pm 0.23$          | $\pm 0.23$ | $\pm 0.26$ | $\pm 0.30$          | $\pm \mathbf{0.09}$ | $\pm 0.18$ | $\pm 0.22$ |
| MoE-FMNIST    | $\pm \mathbf{0.30}$ | $\pm 0.33$ | $\pm 0.38$ | $\pm 0.34$          | $\pm 0.79$          | $\pm 0.38$ | $\pm 0.53$ |
| CovType       | $\pm \mathbf{0.11}$ | $\pm 0.14$ | $\pm 0.28$ | $\pm 0.17$          | $\pm 0.18$          | $\pm 0.24$ | $\pm 0.20$ |
| SVHN          | $\pm 2.78$          | $\pm 1.15$ | $\pm 1.54$ | $\pm \mathbf{0.43}$ | $\pm 1.27$          | $\pm 2.10$ | $\pm 1.22$ |

Table A.11: Accuracy variance for asynchronous momentum variants on classification problems