

MARINA: Faster Non-Convex Distributed Learning with Compression

Eduard Gorbunov^{1,2,3} Konstantin Burlachenko³ Zhize Li³ Peter Richtárik³

¹ Moscow Institute of Physics and Technology, Russia

² Institute for Information Transmission Problems RAS, Russia

³ King Abdullah University of Science and Technology, Kingdom of Saudi Arabia

Abstract

We develop and analyze MARINA: a new communication efficient method for non-convex distributed learning over heterogeneous datasets. MARINA employs a novel communication compression strategy based on the compression of gradient differences which is reminiscent of but different from the strategy employed in the DIANA method of Mishchenko et al (2019). Unlike virtually all competing distributed first-order methods, including DIANA, ours is based on a carefully designed *biased* gradient estimator, which is the key to its superior theoretical and practical performance. To the best of our knowledge, the communication complexity bounds we prove for MARINA are strictly superior to those of all previous first order methods. Further, we develop and analyze two variants of MARINA: VR-MARINA and PP-MARINA. The first method is designed for the case when the local loss functions owned by clients are either of a finite sum or of an expectation form, and the second method allows for partial participation of clients – a feature important in federated learning. All our methods are superior to previous state-of-the-art methods in terms of the oracle/communication complexity. Finally, we provide convergence analysis of all methods for problems satisfying the Polyak-Łojasiewicz condition.

Contents

1	Introduction	2
1.1	Contributions	4
1.2	Related Work	6
1.3	Preliminaries	7
2	MARINA: Compressing Gradient Differences	7
2.1	Convergence Results for Generally Non-Convex Problems	8
2.2	Convergence Results Under Polyak-Łojasiewicz condition	9
3	MARINA and Variance Reduction	10
3.1	Finite Sum Case	10
3.2	Online Case	12
4	MARINA and Partial Participation	14
5	Numerical Experiments	15

A Basic Facts and Auxiliary Results	23
A.1 Useful Properties of Expectations	23
A.2 One Lemma	23
B Missing Proofs for MARINA	24
B.1 Generally Non-Convex Problems	24
B.2 Convergence Results Under Polyak-Łojasiewicz condition	26
C Missing Proofs for VR-MARINA	29
C.1 Finite Sum Case	29
C.1.1 Generally Non-Convex Problems	29
C.1.2 Convergence Results Under Polyak-Łojasiewicz condition	33
C.2 Online Case	36
C.2.1 Generally Non-Convex Problems	36
C.2.2 Convergence Results Under Polyak-Łojasiewicz condition	40
D Missing Proofs for PP-MARINA	44
D.1 Generally Non-Convex Problems	44
D.2 Convergence Results Under Polyak-Łojasiewicz condition	46

1 Introduction

Non-convex optimization problems appear in various applications of machine learning, such as training deep neural networks [Goodfellow et al., 2016] and matrix completion and recovery [Ma et al., 2018, Bhojanapalli et al., 2016]. Because of their practical importance, these problems gained a lot of attention in recent years that led to the rapid development of new efficient methods for non-convex optimization problems [Danilova et al., 2020], and especially the training of deep learning models [Sun, 2019].

Training deep neural networks is notoriously computationally challenging and time-consuming. In the quest to improve the generalization performance of modern deep learning models, practitioners are resorting to using larger and larger datasets in the training process, and to support such workloads, it is imperative to use advanced parallel and distributed hardware, systems, and algorithms. Distributed computing is often necessitated by the desire to train models from data that is naturally distributed across a number of edge devices, as is the case in federated learning [Konečný et al., 2016, McMahan et al., 2017]. However, even when this is not the case, distributed methods are often very efficient at reducing the training time [Goyal et al., 2017, You et al., 2020]. Due to these and other reasons, distributed optimization has gained immense popularity in recent years.

However, distributed methods almost invariably suffer from the so-called *communication bottleneck*: the cost of communication of the information necessary for the workers to jointly solve the problem at hand is often very high, and depending on the particular compute architecture, workload, and algorithm used, it can be orders of magnitude higher than the cost of computation. A popular technique for resolving this issue is *communication compression* [Seide et al., 2014, Konečný et al., 2016, Suresh et al., 2017, Konečný and Richtárik, 2018], which is based on applying a lossy transformation/compression to the models, gradients, or tensors to be sent over the network in order to save on communication. Since applying a lossy compression generally decreases the utility

Table 1: Summary of the state-of-the-art results for finding an ε -stationary point for the problem (1), i.e., such a point $\hat{x}$ that $\mathbf{E} [\|\nabla f(\hat{x})\|^2] \leq \varepsilon^2$. Dependences on the numerical constants, “quality” of the starting point, and smoothness constants are omitted in the complexity bounds. Abbreviations: “PP” = partial participation; “Communication complexity” = number of communications rounds needed to find ε -stationary point; “Oracle complexity” = number of (stochastic) first-order oracle calls needed to find ε -stationary point. Notation: ω = quantization parameter (see Def. 1.1); n = number of nodes; m = the size of the local dataset; r = (expected) number of clients sampled at each iteration; b' = batchsize for VR-MARINA at the iterations with compressed communication. To simplify the bounds we assume that the expected density ζ_Q of the quantization operator Q (see Def. 1.1) satisfies $\omega + 1 = \Theta(d/\zeta_Q)$ (e.g., this holds for RandK and ℓ_2 -quantization see [Beznosikov et al., 2020]). We notice, that [Haddadpour et al., 2020] and [Das et al., 2020] contain also better rates under different assumptions on clients’ similarity.

Setup	Method	Citation	Communication Complexity	Oracle Complexity
(1)	DIANA	[Mishchenko et al., 2019] [Horváth et al., 2019] [Li and Richtárik, 2020]	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2}$	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2}$
	FedCOMGATE (1)	[Haddadpour et al., 2020]	$\frac{1+\omega}{\varepsilon^2}$	$\frac{1+\omega}{n\varepsilon^4}$
	FedSTEPH, $r = n$	[Das et al., 2020]	$\frac{1+\omega/n}{\varepsilon^4}$	$\frac{1+\omega/n}{\varepsilon^4}$
	MARINA (Alg. 1)	Thm. 2.1 & Cor. 2.1 (NEW)	$\frac{1+\omega/\sqrt{n}}{\varepsilon^2}$	$\frac{1+\omega/\sqrt{n}}{\varepsilon^2}$
(1)+(5)	DIANA	[Li and Richtárik, 2020]	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$
	VR-DIANA	[Horváth et al., 2019]	$\frac{(m^{2/3}+\omega)\sqrt{1+\omega/n}}{\varepsilon^2}$	$\frac{(m^{2/3}+\omega)\sqrt{1+\omega/n}}{\varepsilon^2}$
	VR-MARINA (Alg. 2) $b' = 1$ (2)	Thm. 3.1 & Cor. 3.1 (NEW)	$\frac{1+\max\{\omega, \sqrt{(1+\omega)m}\}/\sqrt{n}}{\varepsilon^2}$	$\frac{1+\max\{\omega, \sqrt{(1+\omega)m}\}/\sqrt{n}}{\varepsilon^2}$
(1)+(6)	DIANA (3)	[Mishchenko et al., 2019] [Li and Richtárik, 2020]	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$
	FedCOMGATE (3)	[Haddadpour et al., 2020]	$\frac{1+\omega}{\varepsilon^2}$	$\frac{1+\omega}{n\varepsilon^4}$
	VR-MARINA (Alg. 2) $b' = 1$	Thm. 3.2 & Cor. 3.2 (NEW)	$\frac{1+\omega/\sqrt{n}}{\varepsilon^2} + \frac{\sqrt{1+\omega}}{n\varepsilon^3}$	$\frac{1+\omega/\sqrt{n}}{\varepsilon^2} + \frac{\sqrt{1+\omega}}{n\varepsilon^3}$
	VR-MARINA (Alg. 2) $b' = \Theta\left(\frac{1}{n\varepsilon^2}\right)$	Thm. 3.2 & Cor. 3.2 (NEW)	$\frac{1+\omega/\sqrt{n}}{\varepsilon^2}$	$\frac{1+\omega/\sqrt{n}}{n\varepsilon^4} + \frac{1+\omega}{n\varepsilon^3}$
PP, (1)	FedSTEPH	[Das et al., 2020]	$\frac{1+\omega/n}{r\varepsilon^4} + \frac{(1+\omega)(n-r)}{r(n-1)\varepsilon^4}$	$\frac{1+\omega/n}{r\varepsilon^4} + \frac{(1+\omega)(n-r)}{r(n-1)\varepsilon^4}$
	PP-MARINA (Alg. 4)	Thm. 4.1 & Cor. 4.1 (NEW)	$\frac{1+(1+\omega)\sqrt{n/r}}{\varepsilon^2}$	$\frac{1+(1+\omega)\sqrt{n/r}}{\varepsilon^2}$

(1) The results for FedCOMGATE are derived under assumption that for all vectors $x_1, \dots, x_n \in \mathbb{R}^d$ the quantization operator Q satisfies $\mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n Q(x_i) \right\|^2 - \left\| Q\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \right\|^2 \right] \leq G$ for some constant $G \geq 0$. In fact, this assumption does not hold for classical quantization operators like RandK and ℓ_2 -quantization on $\mathbb{R}^d$. The counterexample: $n = 2$ and $x_1 = -x_2 = (t, t, \dots, t)^\top$ with arbitrary large $t > 0$.

(2) One can even further improve the communication complexity by increasing b' .

(3) No assumptions on the smoothness of the stochastic realizations $f_\xi(x)$ are used.

of the exchanged messages, such an approach will typically lead to an increase in the number of communications, and the overall usefulness of this technique manifests itself in situations where the communication savings are larger compared to the increased need for the number of communication rounds [Horváth et al., 2019].

The optimization and machine learning communities have exerted considerable effort in recent years to design distributed methods supporting compressed communication. From the many methods proposed, we emphasize VR-DIANA [Horváth et al., 2019], FedCOMGATE [Haddadpour et al., 2020], and FedSTEPH [Das et al., 2020] because we do not assume any similarity between the local loss functions and these papers contain the state-of-the-art results in this setup.

1.1 Contributions

We propose several new distributed optimization methods supporting compressed communication, specifically focusing on smooth but nonconvex problems of the form

$$\min_{x \in \mathbb{R}^d} \left\{ f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x) \right\}, \quad (1)$$

where n workers/devices/clients/peers are connected in a centralized way with a parameter-server and client i has an access to the local loss function f_i only. We establish strong complexity rates for them and show that they are better than previous state-of-the-art results.

- MARINA.** The main contribution of our paper is a new distributed method supporting communication compression called MARINA (Algorithm 1). In this algorithm, workers apply an unbiased compression operator to the *gradient differences* at each iteration with some probability and send them to the server that performs aggregation by averaging. Unlike all known methods operating with unbiased compression operators, this procedure leads to a *biased* gradient estimator. We prove convergence guarantees for MARINA that are strictly better than previous state of the art methods (see Table 1). For example, MARINA's rate $\mathcal{O}(\frac{1+\omega/\sqrt{n}}{\varepsilon^2})$ is $\mathcal{O}(\sqrt{\omega})$ times better than that of the state-of-the-art method DIANA [Mishchenko et al., 2019], where ω is the variance parameter associated with the deployed compressor. For example, in the case of the Rand1 sparsification compressor we have $\omega = d - 1$, and hence we get an improvement by the factor $\mathcal{O}(\sqrt{d})$. Since the number d of features can be truly very large when training modern models, this is a substantial improved that can even amount to *several orders of magnitude*.
- Variance Reduction on Nodes.** We generalize MARINA to VR-MARINA, which can handle the situation when the local functions f_i have either a finite-sum (each f_i is an average of m functions) or an expectation form, and when it is more efficient to rely on local stochastic gradients rather than on local gradients. When compared with MARINA, VR-MARINA additionally performs *local variance reduction* on all nodes, progressively removing the variance coming from the stochastic approximation, leading to a better oracle complexity than previous state-of-the-art results (see Table 1). When no compression is used (i.e., $\omega = 0$), the rate of VR-MARINA is $\mathcal{O}(\frac{\sqrt{m}}{\sqrt{n\varepsilon^2}})$, while the rate of the state-of-the-art method VR-DIANA is $\mathcal{O}(\frac{m^{2/3}}{\varepsilon^2})$. This is an improvement by the factor $\mathcal{O}(\sqrt{nm}^{1/6})$. When a lot of compression is applied and ω is large, our method is faster by the factor $\mathcal{O}(\frac{m^{2/3}+\omega}{m^{1/2}+\omega^{1/2}})$. In the special case, when there is just a single node ($n = 1$) and no compression is used, VR-MARINA reduces to the PAGE method of [Li et al., 2020]. PAGE is an optimal first-order algorithm for smooth non-convex finite-sum/online optimization problems (see also [Horváth et al., 2020] for a similar method).
- Partial Participation.** We develop a modification of MARINA allowing for *partial participation* of the clients, which is a feature critical in federated learning. The resulting method, PP-MARINA, has superior communication complexity to existing methods developed for this settings (see Table 1).
- Convergence Under the Polyak-Łojasiewicz Condition.** We analyze all proposed methods for problems satisfying the Polyak-Łojasiewicz condition [Polyak, 1963, Łojasiewicz,

Table 2: Summary of the state-of-the-art results for finding an ε -solution for the problem (1) satisfying **Polyak-Łojasiewicz condition** (see As. 2.1), i.e., such a point $\hat{x}$ that $\mathbf{E}[f(\hat{x}) - f(x^*)] \leq \varepsilon$. Dependences on the numerical constants and $\log(1/\varepsilon)$ factors are omitted and all smoothness constanst are denoted by L in the complexity bounds. Abbreviations: “PP” = partial participation; “Communication complexity” = number of communications rounds needed to find ε -stationary point; “Oracle complexity” = number of (stochastic) first-order oracle calls needed to find ε -stationary point. Notation: ω = quantization parameter (see Def. 1.1); n = number of nodes; m = the size of the local dataset; r = (expected) number of clients sampled at each iteration; b' = batchsize for VR-MARINA at the iterations with compressed communication. To simplify the bounds we assume that the expected density ζ_Q of the quantization operator Q (see Def. 1.1) satisfies $\omega + 1 = \Theta(d/\zeta_Q)$ (e.g., this holds for RandK and ℓ_2 -quantization see [Beznosikov et al., 2020]). We notice, that [Haddadpour et al., 2020] and [Das et al., 2020] contain also better rates under different assumptions on clients’ similarity.

Setup	Method	Citation	Communication Complexity	Oracle Complexity
(1)	DIANA	[Li and Richtárik, 2020]	$\frac{L(1+(1+\omega)\sqrt{\omega/n})}{\mu}$	$\frac{L(1+(1+\omega)\sqrt{\omega/n})}{\mu}$
	FedCOMGATE (1)	[Haddadpour et al., 2020]	$\frac{L(1+\omega)}{\mu}$	$\frac{L(1+\omega)}{\mu}$
	MARINA (Alg. 1)	Thm. 2.2 & Cor. B.2 (NEW)	$\omega + \frac{L(1+\omega/\sqrt{n})}{\mu}$	$\omega + \frac{L(1+\omega/\sqrt{n})}{\mu}$
(1)+(5)	DIANA	[Li and Richtárik, 2020]	$\frac{L(1+(1+\omega)\sqrt{\omega/n})}{\mu} + \frac{\frac{L(1+\omega)}{n\mu} (\frac{L}{\mu} + \frac{1}{\varepsilon})}{\mu}$	$\frac{L(1+(1+\omega)\sqrt{\omega/n})}{\mu} + \frac{\frac{L(1+\omega)}{n\mu} (\frac{L}{\mu} + \frac{1}{\varepsilon})}{\mu}$
	VR-DIANA	[Li and Richtárik, 2020]	$\frac{L(m^{2/3} + \omega)\sqrt{1+\omega/n}}{\mu}$	$\frac{L(m^{2/3} + \omega)\sqrt{1+\omega/n}}{\mu}$
	VR-MARINA (Alg. 2) $b' = 1$ (2)	Thm. C.2 & Cor. C.2 (NEW)	$\omega + m + \frac{L(1+\max\{\omega, \sqrt{(1+\omega)m}\}/\sqrt{n})}{\mu}$	$\omega + m + \frac{L(1+\max\{\omega, \sqrt{(1+\omega)m}\}/\sqrt{n})}{\mu}$
(1)+(6)	DIANA (3)	[Mishchenko et al., 2019]	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$	$\frac{1+(1+\omega)\sqrt{\omega/n}}{\varepsilon^2} + \frac{1+\omega}{n\varepsilon^4}$
	FedCOMGATE (3)	[Li and Richtárik, 2020]	$\frac{L(1+\omega)}{\mu}$	$\frac{L(1+\omega)}{n\mu\varepsilon}$
	VR-MARINA (Alg. 2) $b' = 1$	Thm. C.4 & Cor. C.4 (NEW)	$\omega + \frac{L(1+\omega/\sqrt{n})}{\mu} + \frac{L\sqrt{1+\omega}}{n\mu\varepsilon}$	$\omega + \frac{L(1+\omega/\sqrt{n})}{\mu} + \frac{L\sqrt{1+\omega}}{n\mu\varepsilon}$
	VR-MARINA (Alg. 2) $b' = \Theta(\frac{1}{n\mu\varepsilon})$	Thm. C.4 & Cor. C.4 (NEW)	$\omega + \frac{L(1+\omega/\sqrt{n})}{\mu}$	$\frac{1+\omega}{n\mu\varepsilon} + \frac{L(1+\omega/\sqrt{n})}{n\mu^2\varepsilon} + \frac{L(1+\omega)}{n\mu^2\sqrt{\varepsilon}}$
PP, (1)	FedSTEPH (4)	[Das et al., 2020]	$\left(\frac{L}{\mu}\right)^{3/2}$	$\left(\frac{L}{\mu}\right)^{3/2}$
	PP-MARINA (Alg. 4)	Thm. D.2 & Cor. D.2 (NEW)	$\frac{(\omega+1)n}{r} + \frac{L(1+(1+\omega)\sqrt{n/r})}{\mu}$	$\frac{(\omega+1)n}{r} + \frac{L(1+(1+\omega)\sqrt{n/r})}{\mu}$

(1) The results for FedCOMGATE are derived under assumption that for all vectors $x_1, \dots, x_n \in \mathbb{R}^d$ the quantization operator Q satisfies $\mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n Q(x_i) \right\|^2 - \left\| Q \left(\frac{1}{n} \sum_{i=1}^n x_i \right) \right\|^2 \right] \leq G$ for some constant $G \geq 0$. In fact, this assumption does not hold for classical quantization operators like RandK and ℓ_2 -quantization on $\mathbb{R}^d$. The counterexample: $n = 2$ and $x_1 = -x_2 = (t, t, \dots, t)^\top$ with arbitrary large $t > 0$.

(2) One can even further improve the communication complexity by increasing b' .

(3) No assumptions on the smoothness of the stochastic realizations $f_\xi(x)$ are used.

(4) The rate is derived under assumption that $r = \Omega((1+\omega)\sqrt{L/\mu} \log(1/\varepsilon))$.

1963]. Again, the obtained results are strictly better than previous state-of-the-art (see Table 2). Statements and proofs of all these results are in the appendix.

- **Simple Analysis.** The simplicity and flexibility of our analysis offer several extensions. For example, one can easily generalize our analysis to the case of different quantization operators and different batch sizes used by clients. Moreover, one can combine the ideas of VR-MARINA and PP-MARINA and obtain a single distributed algorithm with compressed communications, variance reduction on nodes, and clients’ sampling. We did not do this in order to keep the exposition simpler.

1.2 Related Work

Non-Convex Optimization. Since finding a global minimum of a non-convex function is, in general, an NP-hard problem [Murty and Kabadi, 1987], a lot of researchers in non-convex optimization focus on relaxed goals such as finding an ε -stationary point. The theory of stochastic first-order methods for finding ε -stationary points is well-developed: it contains lower bounds for expectation minimization without smoothness of stochastic realizations [Arjevani et al., 2019] and for finite-sum/expectation minimization [Fang et al., 2018, Li et al., 2020] as well as optimal methods matching the lower bounds (see [Danilova et al., 2020, Li et al., 2020] for the overview). Recently, distributed variants of such methods were proposed [Sun et al., 2020, Sharma et al., 2019, Khanduri et al., 2020].

Compressed Communications. Works on distributed methods supporting communication compression can be roughly split into two big groups: the first group focuses on methods using *unbiased* compression operators (which refer to as quantizations in this paper) such as RandK, and the second one studies methods using *biased* compressors such as TopK. One can find a detailed summary of the most popular compression operators in [Safaryan et al., 2021, Beznosikov et al., 2020].

Unbiased Compression. In this line of work, the first convergence result in the non-convex case was obtained by Alistarh et al. [2017] for QSGD, under assumptions that the local loss functions are the same for all workers and the stochastic gradient has uniformly bounded second moment. After that, Mishchenko et al. [2019] proposed DIANA (and its momentum version) and proved its convergence rate for non-convex problems without any assumption on the boundedness of the second moment of the stochastic gradient, but under the assumption that the dissimilarity between local loss functions is bounded. This restriction was later eliminated by Horváth et al. [2019] for the variance reduced version of DIANA called VR-DIANA, and the analysis was extended to a large class of unbiased compressors. Finally, the results for QSGD and DIANA were recently generalized and tightened by Li and Richtárik [2020] in a unifying framework which included many other methods as well.

Biased Compression. Biased compression operators are less “optimization-friendly” than unbiased ones. Indeed, one can construct a simple convex quadratic problem for which distributed SGD with Top1 compression diverges exponentially fast [Beznosikov et al., 2020]. However, this issue can be resolved using *error compensation* [Seide et al., 2014]. The first analysis of error-compensated SGD (EC-SGD) for non-convex problems was obtained by Karimireddy et al. [2019] for homogeneous problems under the assumption that the second moment of the stochastic gradient is uniformly bounded. The last assumption was recently removed from the analysis of EC-SGD by Stich and Karimireddy [2020], Beznosikov et al. [2020], while the first results without the homogeneity assumption were obtained by Koloskova et al. [2020a] for Choco-SGD, but still under the assumption that the second moment of the stochastic gradient is uniformly bounded. This issue was resolved by Beznosikov et al. [2020]. In general, the current understanding of optimization methods with biased compressors is far from complete: even in the strongly convex case, the first linearly converging [Gorbunov et al., 2020] and accelerated [Qian et al., 2020] error-compensated stochastic methods were proposed just recently.

Other Approaches. Besides communication compression, there are also different techniques aiming to reduce the overall communication cost of distributed methods. The most popular ones are based on decentralized communications and multiple local steps between communication rounds, where the second technique is very popular in Federated Learning [Konečný et al., 2016, Kairouz et al., 2019]. One can find the state-of-the-art distributed optimization methods using these techniques and their combinations in [Lian et al., 2017, Karimireddy et al., 2020, Li et al., 2019, Koloskova et al., 2020b]. Moreover, there exist results based on the combinations of communication compression with either decentralized communication, e.g., Choco-SGD [Koloskova et al., 2020a], or local updates, e.g., Qsparse-Local-SGD [Basu et al., 2019], FedCOMGATE [Haddadpour et al., 2020], FedSTEPH [Das et al., 2020], where in [Basu et al., 2019] the convergence rates were derived under assumption that the stochastic gradient has uniformly bounded second moment and the results for Choco-SGD, FedCOMGATE, FedSTEPH were described either earlier in the text, or in Table 1.

1.3 Preliminaries

We will rely on two key assumptions throughout the text.

Assumption 1.1 (Uniform lower bound). *There exists $f_* \in \mathbb{R}$ such that $f(x) \geq f_*$ for all $x \in \mathbb{R}^d$.*

Assumption 1.2 (L -smoothness). *We assume that f_i is L_i -smooth for all $i \in [n] = \{1, 2, \dots, n\}$ meaning that the following inequality holds $\forall x, y \in \mathbb{R}^d, \forall i \in [n]$:*

$$\|\nabla f_i(x) - \nabla f_i(y)\| \leq L_i \|x - y\|. \quad (2)$$

This assumption implies that f is L_f -smooth with $L_f^2 \leq L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$.

Finally, we describe a large class of unbiased compression operators satisfying a certain variance bound which we will refer to in this paper by the name *quantization*.

Definition 1.1 (Quantization). *We say that a stochastic mapping $\mathcal{Q} : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a quantization operator/quantization if there exists $\omega > 0$ such that for any $x \in \mathbb{R}^d$, we have*

$$\mathbf{E}[\mathcal{Q}(x)] = x, \quad \mathbf{E}[\|\mathcal{Q}(x) - x\|^2] \leq \omega \|x\|^2. \quad (3)$$

For the given quantization operator $\mathcal{Q}(x)$ we define the the expected density as

$$\zeta_{\mathcal{Q}} = \sup_{x \in \mathbb{R}^d} \mathbf{E}[\|\mathcal{Q}(x)\|_0],$$

where $\|y\|_0$ is the number of non-zero components of $y \in \mathbb{R}^d$.

Notice that the expected density is well-defined for any quantization operator since $\|\mathcal{Q}(x)\|_0 \leq d$.

2 MARINA: Compressing Gradient Differences

In this section, we describe the main algorithm of this work: MARINA (see Algorithm 1). At each iteration of MARINA, each worker i either sends to the server the dense vector $\nabla f_i(x^{k+1})$ with probability p , or it sends the quantized gradient difference $\mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k))$ with probability $1 - p$. In the first situation, the server just averages the vectors received from workers and gets $g^{k+1} = \nabla f(x^{k+1})$, while in the second case, the server averages the quantized differences

Algorithm 1 MARINA

```
1: Input: starting point  $x^0$ , stepsize  $\gamma$ , probability  $p \in (0, 1]$ , number of iterations  $K$ 
2: Initialize  $g^0 = \nabla f(x^0)$ 
3: for  $k = 0, 1, \dots, K - 1$  do
4:   Sample  $c_k \sim \text{Be}(p)$  (i.e.,  $c_k = 1$  with probability  $p$ , and  $c_k = 0$  with probability  $1 - p$ )
5:   Broadcast  $g^k$  to all workers
6:   for  $i = 1, \dots, n$  in parallel do
7:      $x^{k+1} = x^k - \gamma g^k$ 
8:     Define local gradient estimator
           
$$g_i^{k+1} = \begin{cases} \nabla f_i(x^{k+1}) & \text{if } c_k = 1 \\ g^k + \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) & \text{if } c_k = 0 \end{cases}$$

9:   end for
10:   $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$ 
11: end for
12: Return:  $\hat{x}^K$  chosen uniformly at random from  $\{x^k\}_{k=0}^{K-1}$ 
```

from all workers and then adds the result to g^k to get g^{k+1} . Moreover, if $\mathcal{Q}$ is identity quantization, i.e., $\mathcal{Q}(x) = x$, then MARINA reduces to Gradient Descent (GD).

However, for non-trivial quantizations we have $\mathbf{E}[g^{k+1} \mid x^{k+1}] \neq \nabla f(x^{k+1})$ unlike all other distributed methods using exclusively unbiased compressors we know of. That is, g^{k+1} is a *biased* stochastic estimator of $\nabla f(x^{k+1})$. However, MARINA is an example of a rare phenomenon in stochastic optimization when the *bias of the stochastic gradient helps to achieve better complexity*.

2.1 Convergence Results for Generally Non-Convex Problems

We start with the following result.

Theorem 2.1. *Let Assumptions 1.1 and 1.2 be satisfied. Then, after*

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right)$$

iterations with $\Delta_0 = f(x^0) - f_$, $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and the stepsize*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right)}$$

MARINA produces point $\hat{x}^K$ for which $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$.

One can find the full statement of the theorem together with its proof in Section B.1 of the appendix.

The following corollary provides the bounds on the number of iterations/communication rounds and estimates the total communication cost needed to achieve ε -stationary point in expectation. Moreover, for simplicity, throughout the paper we assume that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Corollary 2.1. *Let the assumptions of Theorem 2.1 hold and $p = \zeta_Q/d$. If*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_Q} - 1 \right)} \right)},$$

then MARINA requires

$$\mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_Q} - 1 \right)} \right) \right)$$

iterations/communication rounds in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is $\mathcal{O}(d + \zeta_Q K)$.

Let us clarify the obtained result. First of all, if $\omega = 0$ (no quantization), then $\zeta_Q = 0$ and the rate coincides with the rate of Gradient Descent (GD). Since GD is optimal among first-order methods in terms of reducing the norm of the gradient [Carmon et al., 2019], the dependence on ε in our bound cannot be improved in general. Next, if n is big enough, i.e., $n \geq \omega(d/\zeta_Q - 1)$, then¹ the iteration complexity of MARINA (method with compressed communications) and GD (method with dense communications) coincide. This means that in this regime, MARINA is able to reach a provably better communication complexity than GD!

2.2 Convergence Results Under Polyak-Łojasiewicz condition

In this section, we provide a complexity bounds for MARINA under the Polyak-Łojasiewicz (PL) condition.

Assumption 2.1 (PL condition). *Function f satisfies Polyak-Łojasiewicz (PL) condition with parameter μ , i.e.,*

$$\|\nabla f(x)\|^2 \geq 2\mu (f(x) - f(x^*)). \quad (4)$$

holds for $x^ = \arg \min_{x \in \mathbb{R}^d} f(x)$ and for all $x \in \mathbb{R}^d$.*

Under this and previously introduced assumptions we derive the following result.

Theorem 2.2. *Let Assumptions 1.1, 1.2 and 3.1 be satisfied. Then, after*

$$K = \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

iterations with $\Delta_0 = f(x^0) - f(x^)$, $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and the stepsize*

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2(1-p)\omega}{pn}} \right)}, \frac{p}{2\mu} \right\}$$

MARINA produces a point x^K for which $\mathbf{E}[f(x^K) - f(x^)] \leq \varepsilon$.*

One can find the full statement of the theorem together with its proof in Section B.2 of the appendix.

¹For ℓ_2 -quantization this requirement is satisfied when $n \geq d$.

3 MARINA and Variance Reduction

Throughout this section, we assume that the local loss on each node has either a finite-sum form (finite sum case),

$$f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x), \quad (5)$$

or an expectation form (online case),

$$f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i}[f_{\xi_i}(x)]. \quad (6)$$

3.1 Finite Sum Case

In this section, we generalize MARINA to problems of the form (1)+(5), obtaining VR-MARINA (see Algorithm 2). At each iteration of VR-MARINA, devices are to compute the full gradients $\nabla f_i(x^{k+1})$

Algorithm 2 VR-MARINA: finite sum case

- 1: **Input:** starting point x^0 , stepsize γ , minibatch size b' , probability $p \in (0, 1]$, number of iterations K
- 2: Initialize $g^0 = \nabla f(x^0)$
- 3: **for** $k = 0, 1, \dots, K - 1$ **do**
- 4: Sample $c_k \sim \text{Be}(p)$
- 5: Broadcast g^k to all workers
- 6: **for** $i = 1, \dots, n$ in parallel **do**
- 7: $x^{k+1} = x^k - \gamma g^k$
- 8: Define local gradient estimator

$$g_i^{k+1} = \begin{cases} \nabla f_i(x^{k+1}) & \text{if } c_k = 1, \\ g^k + \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) & \text{if } c_k = 0, \end{cases}$$

where $I'_{i,k}$ is the set of the indices in the minibatch, $|I'_{i,k}| = b'$

- 9: **end for**
 - 10: $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$
 - 11: **end for**
 - 12: **Return:** $\hat{x}^K$ chosen uniformly at random from $\{x^k\}_{k=0}^{K-1}$
-

and send them to the server with probability p . Typically $p \leq 1/m$ and m is large, meaning that workers compute full gradients rarely (once per $\geq m$ iterations in expectation). At other iterations workers compute minibatch stochastic gradients evaluated at the current and previous points, compress them using an unbiased compression operator, i.e., quantization/quantization operator, and send the resulting vectors $g_i^{k+1} - g^k$ to the server. Moreover, if $\mathcal{Q}$ is the identity quantization, i.e., $\mathcal{Q}(x) = x$, and $n = 1$, then MARINA reduces to the optimal method PAGE [Li et al., 2020].

In this part, we will rely on the following average smoothness assumption.

Assumption 3.1 (Average $\mathcal{L}$ -smoothness). *For all $k \geq 0$ and $i \in [n]$ the minibatch stochastic gradients difference $\tilde{\Delta}_i^k = \frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k))$ computed on the i -th worker satisfies $\mathbf{E} [\tilde{\Delta}_i^k \mid x^k, x^{k+1}] = \Delta_i^k$ and*

$$\mathbf{E} \left[\left\| \tilde{\Delta}_i^k - \Delta_i^k \right\|^2 \mid x^k, x^{k+1} \right] \leq \frac{\mathcal{L}_i^2}{b'} \|x^{k+1} - x^k\|^2 \quad (7)$$

with some $\mathcal{L}_i \geq 0$, where $\Delta_i^k = \nabla f_i(x^{k+1}) - \nabla f_i(x^k)$.

This assumption is satisfied in many standard minibatch regimes. In particular, if $I'_{i,k} = \{1, \dots, m\}$, then $\mathcal{L}_i = 0$, and if $I'_{i,k}$ consists of b' i.i.d. samples from the uniform distributions on $\{1, \dots, m\}$ and f_{ij} are L_{ij} -smooth, then $\mathcal{L}_i \leq \max_{j \in [m]} L_{ij}$.

Under this and the previously introduced assumptions, we derive the following result.

Theorem 3.1. *Consider the finite sum case (1)+(5). Let Assumptions 1.1, 1.2 and 3.1 be satisfied. Then, after*

$$K = \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \right) \right)$$

iterations with $\Delta_0 = f(x^0) - f_*$, $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$, $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$ and the stepsize

$$\gamma \leq \frac{1}{L + \sqrt{\frac{1-p}{pn}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)},$$

VR-MARINA produces such a point $\hat{x}^K$ that $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$.

One can find the full statement of the theorem together with its proof in Section C.1.1 of the appendix.

Corollary 3.1. *Let the assumptions of Theorem 3.1 hold and $p = \min \{\zeta_{\mathcal{Q}}/d, b'/(m+b')\}$, where $b' \leq m$. If*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)},$$

then VR-MARINA requires

$$\mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{nb'}} \right) \right)$$

iterations/communication rounds and $\mathcal{O}(m + b'K)$ stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is $\mathcal{O}(d + \zeta_{\mathcal{Q}}K)$.

First of all, when workers quantize differences of the full gradients, then $I'_{i,k} = \{1, \dots, m\}$ for all $i \in [n]$ and $k \geq 0$, implying $\mathcal{L} = 0$. In this case, the complexity bounds for VR-MARINA recover the ones for MARINA. Next, when $\omega = 0$ (no quantization) and $n = 1$, our bounds for

iteration and oracle complexities for VR-MARINA recover the bounds for PAGE [Li and Richtárik, 2020], which is optimal for finite-sum smooth non-convex optimization. This observation implies that the dependence on ε and m in the complexity bounds for VR-MARINA cannot be improved in the class of first-order stochastic methods. Next, we notice that up to the differences in smoothness constants, the iteration and oracle complexities for VR-MARINA benefit from the number of workers n . Finally, as Table 1 shows, the rates for VR-MARINA are strictly better than ones for the previous state-of-the-art method VR-DIANA [Horváth et al., 2019].

We provide the convergence results for VR-MARINA in the finite-sum case under the Polyak-Łojasiewicz condition, together with complete proofs, in Section C.1.2 of the appendix.

3.2 Online Case

In this section, we focus on problems of type (1)+(6). For this type of problems we consider a slightly modified version of VR-MARINA (see Algorithm 3). That is, we replace line 8 in Algorithm 2 by the following update rule: $g_i^{k+1} = \frac{1}{b} \sum_{j \in I_{i,k}} \nabla f_{\xi_{ij}^k}(x^{k+1})$ if $c_k = 1$, and $g_i^{k+1} = g^k + \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f_{\xi_{ij}^k}(x^k)) \right)$ otherwise, where $I_{i,k}, I'_{i,k}$ are the sets of the indices in the minibatches, $|I_{i,k}| = b$, $|I'_{i,k}| = b'$, and ξ_{ij}^k is independently sampled from $\mathcal{D}_i$ for $i \in [n]$, $j \in [m]$.

Algorithm 3 VR-MARINA: online case

- 1: **Input:** starting point x^0 , stepsize γ , minibatch sizes $b, b' < b$, probability $p \in (0, 1]$, number of iterations K
- 2: Initialize $g^0 = \frac{1}{nb} \sum_{i=1}^n \sum_{j \in I_{i,0}} \nabla f_{\xi_{ij}^0}(x^{k+1})$, where $I_{i,0}$ is the set of the indices in the minibatch, $|I_{i,0}| = b$, and ξ_{ij}^0 is independently sampled from $\mathcal{D}_i$ for $i \in [n]$, $j \in [m]$
- 3: **for** $k = 0, 1, \dots, K - 1$ **do**
- 4: Sample $c_k \sim \text{Be}(p)$
- 5: Broadcast g^k to all workers
- 6: **for** $i = 1, \dots, n$ in parallel **do**
- 7: $x^{k+1} = x^k - \gamma g^k$
- 8: Define local gradient estimator

$$g_i^{k+1} = \begin{cases} \frac{1}{b} \sum_{j \in I_{i,k}} \nabla f_{\xi_{ij}^k}(x^{k+1}), & \text{if } c_k = 1, \\ g^k + \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f_{\xi_{ij}^k}(x^k)) \right), & \text{if } c_k = 0, \end{cases}$$

where $I_{i,k}, I'_{i,k}$ are the sets of the indices in the minibatches, $|I_{i,k}| = b$, $|I'_{i,k}| = b'$, and ξ_{ij}^k is independently sampled from $\mathcal{D}_i$ for $i \in [n]$, $j \in [m]$

- 9: **end for**
 - 10: $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$
 - 11: **end for**
 - 12: **Return:** $\hat{x}^K$ chosen uniformly at random from $\{x^k\}_{k=0}^{K-1}$
-

Before we provide our convergence results in this setup, we reformulate Assumption 3.1 for the

online case.

Assumption 3.2 (Average $\mathcal{L}$ -smoothness). *For all $k \geq 0$ and $i \in [n]$ the minibatch stochastic gradients difference $\tilde{\Delta}_i^k = \frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f_{\xi_{ij}^k}(x^k))$ computed on the i -th worker satisfies $\mathbf{E} [\tilde{\Delta}_i^k \mid x^k, x^{k+1}] = \Delta_i^k$ and*

$$\mathbf{E} \left[\left\| \tilde{\Delta}_i^k - \Delta_i^k \right\|^2 \mid x^k, x^{k+1} \right] \leq \frac{\mathcal{L}_i^2}{b'} \|x^{k+1} - x^k\|^2 \quad (8)$$

with some $\mathcal{L}_i \geq 0$, where $\Delta_i^k = \nabla f_i(x^{k+1}) - \nabla f_i(x^k)$.

Moreover, we assume that the variance of the stochastic gradients on all nodes is uniformly upper bounded.

Assumption 3.3. *We assume that for all $i \in [n]$ there exists such constant $\sigma_i \in [0, +\infty)$ that for all $x \in \mathbb{R}^d$*

$$\mathbf{E}_{\xi_i \sim \mathcal{D}_i} [\nabla f_{\xi_i}(x)] = \nabla f_i(x), \quad (9)$$

$$\mathbf{E}_{\xi_i \sim \mathcal{D}_i} [\|\nabla f_{\xi_i}(x) - \nabla f_i(x)\|^2] \leq \sigma_i^2. \quad (10)$$

Under these and previously introduced assumptions, we derive the following result.

Theorem 3.2. *Consider the online case (1)+(6). Let Assumptions 1.1, 1.2, 3.2 and 3.3 be satisfied. Then, after*

$$K = \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) \right)$$

iterations with $\Delta_0 = f(x^0) - f_*$, $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$, $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$, the stepsize

$$\gamma \leq \frac{1}{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}},$$

and $b = \Theta(\sigma^2/(n\varepsilon^2))$, $\sigma^2 = \frac{1}{n} \sum_{i=1}^n \sigma_i^2$, VR-MARINA produces a point $\hat{x}^K$ for which $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$.

One can find the full statement of the theorem, together with its proof, in Section C.2.1 of the appendix.

Corollary 3.2. *Let the assumptions of Theorem 3.2 hold and choose $p = \min \{\zeta_Q/d, b'/(b+b')\}$, where $b' \leq b$, $b = \Theta(\sigma^2/(n\varepsilon^2))$. If*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\{d/\zeta_Q - 1, b/b'\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}},$$

then VR-MARINA requires

$$\mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega}{n} \max \left\{ \frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega)}{nb'} \max \left\{ \frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) \right)$$

Algorithm 4 PP-MARINA

```
1: Input: starting point  $x^0$ , stepsize  $\gamma$ , probability  $p \in (0, 1]$ , number of iterations  $K$ , clients-  
   batchsize  $r \leq n$   
2: Initialize  $g^0 = \nabla f(x^0)$   
3: for  $k = 0, 1, \dots, K - 1$  do  
4:   Sample  $c_k \sim \text{Be}(p)$   
5:   Choose  $I'_k = \{1, \dots, n\}$  if  $c_k = 1$ , and choose  $I'_k$  as the set of  $r$  i.i.d. samples from the uniform  
   distribution over  $\{1, \dots, n\}$  otherwise  
6:   Broadcast  $g^k$  to all workers  
7:   for  $i = 1, \dots, n$  in parallel do  
8:      $x^{k+1} = x^k - \gamma g^k$   
9:     Set  $g_i^{k+1} = \begin{cases} \nabla f_i(x^{k+1}) & \text{if } c_k = 1, \\ g^k + \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) & \text{if } c_k = 0. \end{cases}$   
10:  end for  
11:  Set  $g^{k+1} = \begin{cases} \nabla f(x^{k+1}) & \text{if } c_k = 1, \\ g^k + \frac{1}{r} \sum_{i_k \in I'_k} \mathcal{Q}(\nabla f_{i_k}(x^{k+1}) - \nabla f_{i_k}(x^k)) & \text{if } c_k = 0. \end{cases}$   
12: end for  
13: Return:  $\hat{x}^K$  chosen uniformly at random from  $\{x^k\}_{k=0}^{K-1}$ 
```

iterations/communication rounds and $\mathcal{O}(\zeta_{\mathcal{Q}}K + \sigma^2/(n\varepsilon^2))$ stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is $\mathcal{O}(d + \zeta_{\mathcal{Q}}K)$.

Similarly to the finite-sum case, when $\omega = 0$ (no quantization) and $n = 1$, our bounds for iteration and oracle complexities for VR-MARINA recover the bounds for PAGE [Li and Richtárik, 2020], which is optimal for online smooth non-convex optimization as well. That is, the dependence on ε in the complexity bound for VR-MARINA cannot be improved in the class of first-order stochastic methods. As before, up to the differences in smoothness constants, the iteration and oracle complexities for VR-MARINA benefit from an increase in the number of workers n .

We provide the convergence results for VR-MARINA in the online case under the Polyak-Łojasiewicz condition, together with complete proofs, in Section C.2.2 of the appendix.

4 MARINA and Partial Participation

Finally, we propose another modification of MARINA. In particular, we prove an option for *partial participation* of the clients - a feature important in federated learning. The resulting method is called PP-MARINA (see Algorithm 4). At each iteration of PP-MARINA, the server receives the quantized gradient differences from r clients with probability $1 - p$, and aggregates full gradients from all clients with probability p , i.e., PP-MARINA coincides with MARINA up to the following difference: $g_i^{k+1} = \nabla f_i(x^{k+1})$, $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$ if $c_k = 1$, and $g_i^{k+1} = g^k + \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k))$, $g^{k+1} = \frac{1}{r} \sum_{i_k \in I'_k} g_i^{k+1}$ otherwise, where I'_k is the set of r i.i.d. samples from the uniform distribution over $\{1, \dots, n\}$. That is, if the probability p is chosen to be small enough, then with high probability the server receives only quantized vectors from a subset of clients at each iteration.

Below, we provide a convergence result for PP-MARINA for smooth non-convex problems.

Theorem 4.1. *Let Assumptions 1.1 and 1.2 be satisfied. Then, after*

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right)$$

iterations with $\Delta_0 = f(x^0) - f_$, $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and the stepsize*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right)}$$

PP-MARINA produces a point $\hat{x}^K$ for which $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$.

One can find the full statement of the theorem together with its proof in Section D.1 of the appendix.

Corollary 4.1. *Let the assumptions of Theorem 4.1 hold and choose $p = \zeta_{\mathcal{Q}}^r / (dn)$, where $r \leq n$. If*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{1+\omega}{b^r} \left(\frac{dn}{\zeta_{\mathcal{Q}}^r} - 1 \right)} \right)},$$

then PP-MARINA requires

$$\mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_{\mathcal{Q}}^r} - 1 \right)} \right) \right)$$

iterations/communication rounds in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost is $\mathcal{O}(dn + \zeta_{\mathcal{Q}} r K)$.

When $r = n$, i.e., all clients participate in communication with the server at each iteration, the rate for PP-MARINA recovers the rate for MARINA under the assumption that $(1 + \omega)(d/\zeta_{\mathcal{Q}} - 1) = \mathcal{O}(\omega(d/\zeta_{\mathcal{Q}} - 1))$ that holds for a wide class of quantization operators, e.g., for identical quantization, RandK, and ℓ_p -quantization. In general, the derived complexity is strictly better than previous state-of-the-art (see Table 1).

We provide the convergence results for PP-MARINA under the Polyak-Łojasiewicz condition, together with complete proofs, in Section D.2 of the appendix.

5 Numerical Experiments

We conduct several numerical experiments on binary classification problem involving non-convex loss [Zhao et al., 2010] (used for two-layer neural networks) with LibSVM data [Chang and Lin, 2011] to justify the theoretical claims of the paper. That is, we consider the following problem:

$$\min_{x \in \mathbb{R}^d} \left\{ f(x) = \frac{1}{N} \sum_{t=1}^N \ell(a_t^\top x, y_i) + \frac{\lambda}{2} \|x\|^2 \right\}, \quad (11)$$

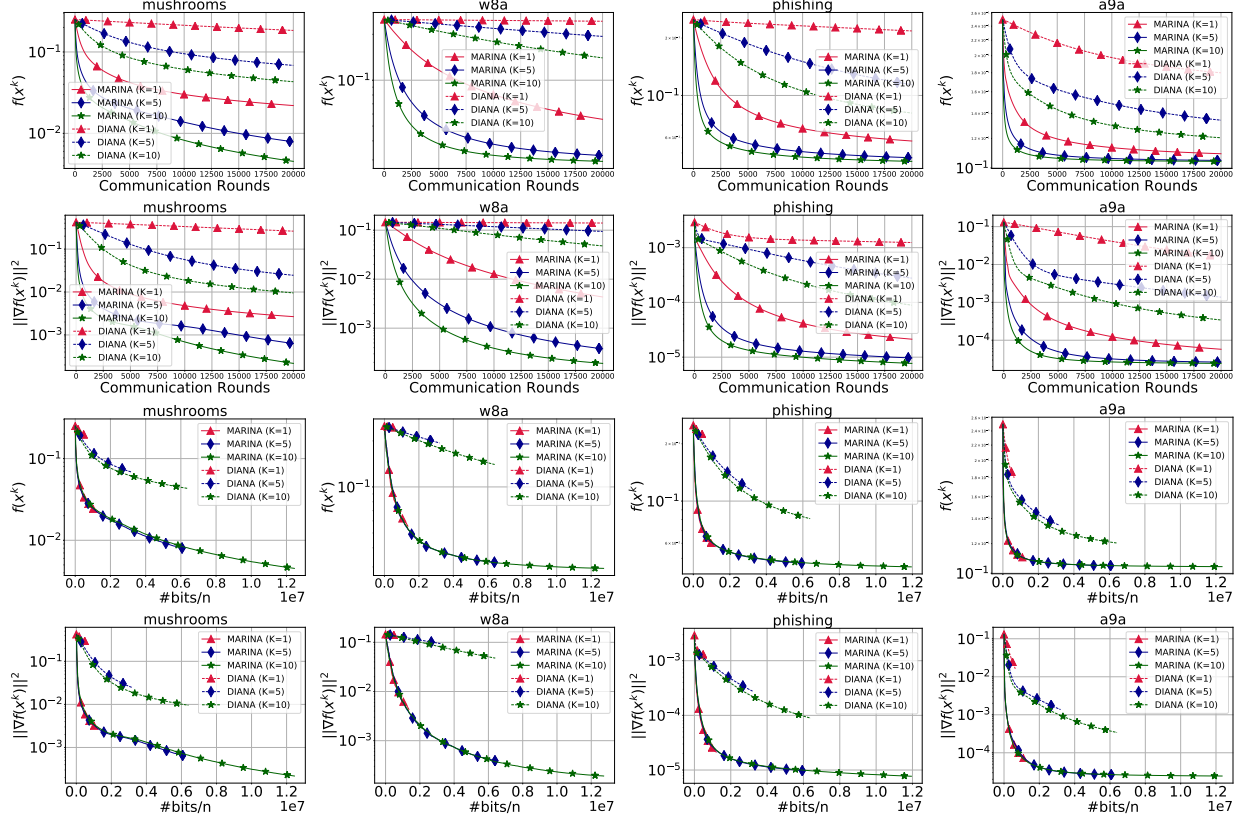


Figure 1: Comparison of MARINA with DIANA on binary classification problem involving non-convex loss (11) with LibSVM data [Chang and Lin, 2011]. Parameter n is chosen as per Tbl. 3 ($n = 5$). Stepsizes for the methods are chosen according to the theory. In all cases, we used the RandK sparsification operator with $K \in \{1, 5, 10\}$. No regularization is applied.

where $\{a_t\} \in \mathbb{R}^d$, $y_i \in \{-1, 1\}$ for all $t = 1, \dots, N$, $\lambda \geq 0$ is ℓ_2 -regularization parameter, and the function $\ell : \mathbb{R}^d \rightarrow \mathbb{R}$ is defined as

$$\ell(b, c) = \left(1 - \frac{1}{1 + \exp(-bc)}\right)^2.$$

The datasets were taken from LibSVM [Chang and Lin, 2011] and split into 5 equal parts among 5 clients (we excluded $N - 5 \cdot \lfloor N/5 \rfloor$ last datapoints from each dataset), see the summary in Table 3. The code was written in Python 3.8 using MPI4PY to emulate the distributed environment and then was executed on a machine with 48 cores, each is Intel(R) Xeon(R) Gold 6246 CPU 3.30GHz.

In our experiments, we compare MARINA with the full-batch version of DIANA, and then VR-MARINA with VR-DIANA. We exclude FedCOMGATE and FedPATH from this comparison since they have significantly worse oracle complexities (see Table 1). Since one of the main goals of our experiments is to justify the theoretical findings of the paper, in the experiments, we used the stepsizes from the corresponding theoretical results for the methods (for DIANA and VR-DIANA the stepsizes were chosen according to [Horváth et al., 2019, Li and Richtárik, 2020]). Next, to compute the stochastic gradients we use batchsizes = $\max\{1, m/100\}$ for VR-MARINA and VR-DIANA.

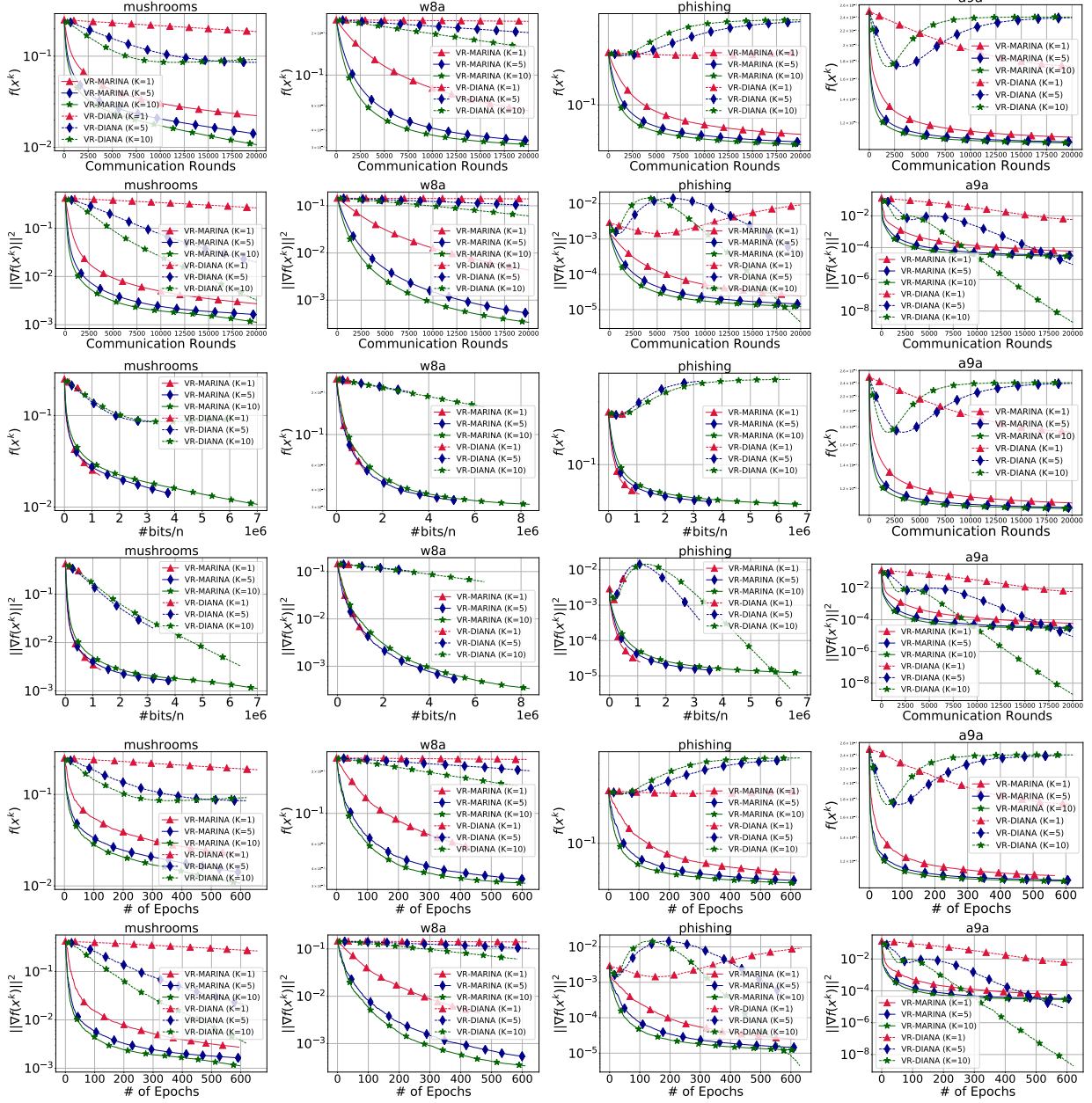


Figure 2: Comparison of VR-MARINA with VR-DIANA on binary classification problem involving non-convex loss (11) with LibSVM data [Chang and Lin, 2011]. Parameter n is chosen as per Tbl. 3 ($n = 5$). Stepsizes for the methods are chosen according to the theory and the batchsizes are $\sim m/100$. In all cases, we used the RandK sparsification operator with $K \in \{1, 5, 10\}$. No regularization is applied.

The results for the full-batched methods are reported in Figure 1, and the comparison of VR-MARINA and VR-DIANA is given in Figure 2. Clearly, in both cases, MARINA and VR-MARINA show faster convergence than the previous state-of-the-art methods, DIANA and VR-DIANA, for distributed

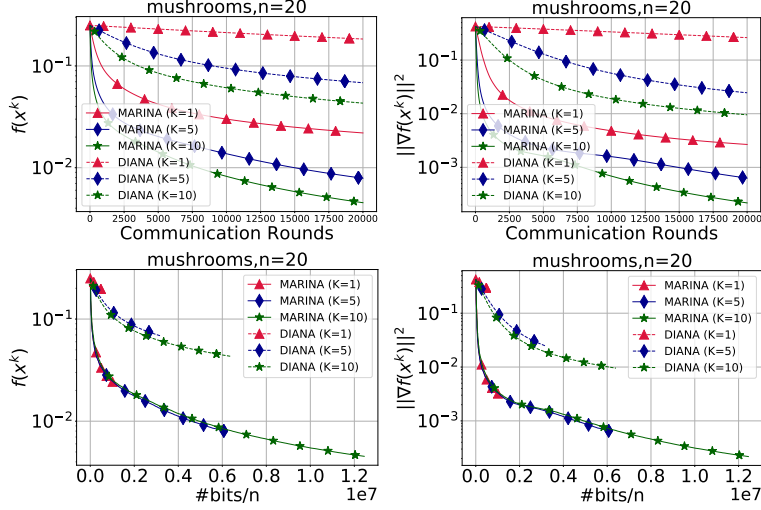


Figure 3: Comparison of MARINA with DIANA on binary classification problem involving non-convex loss (11) with `mushrooms` dataset and $n = 20$ workers. Stepsizes for the methods are chosen according to the theory. In all cases, we used the RandK sparsification operator with $K \in \{1, 5, 10\}$. No regularization is applied.

Table 3: Summary of the datasets and splitting of the data among clients.

Dataset	n	N (# of datapoints)	d (# of features)
<code>mushrooms</code>	5	8 120	112
<code>w8a</code>	5	49 745	300
<code>phishing</code>	5	11 055	69
<code>a9a</code>	5	32 560	124

non-convex optimization with compression in terms of $\|\nabla f(x^k)\|^2$ and $f(x^k)$ decrease w.r.t. the number of communication rounds, oracle calls per node and the total number of transferred bits from workers to the master.

We also tested MARINA and DIANA on `mushrooms` dataset with a bigger number of workers ($n = 20$). The results are reported in Figure 3. Similarly to the previous numerical tests, MARINA shows its superiority to DIANA with $n = 20$ as well.

Acknowledgments. The work of Peter Richtárik, Eduard Gorbunov, Konstantin Burlachenko and Zhize Li was supported by KAUST Baseline Research Fund. The paper was written while E. Gorbunov was a research intern at KAUST. The work of E. Gorbunov in Sections 1, 2, and B was also partially supported by the Ministry of Science and Higher Education of the Russian Federation (Goszadaniye) 075-00337-20-03, project No. 0714-2020-0005, and in Sections 3, 4, C, D – by RFBR, project number 19-31-51001. We would like to thank Konstantin Mishchenko (KAUST) for a suggestion related to the experiments.

References

- D. Alistarh, D. Grubic, J. Li, R. Tomioka, and M. Vojnovic. QSGD: Communication-efficient SGD via gradient quantization and encoding. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 1709–1720, 2017.
- Y. Arjevani, Y. Carmon, J. C. Duchi, D. J. Foster, N. Srebro, and B. Woodworth. Lower bounds for non-convex stochastic optimization. *arXiv preprint arXiv:1912.02365*, 2019.
- D. Basu, D. Data, C. Karakus, and S. Diggavi. Qsparse-local-SGD: Distributed SGD with quantization, sparsification and local computations. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 14668–14679, 2019.
- A. Beznosikov, S. Horváth, P. Richtárik, and M. Safaryan. On biased compression for distributed learning. *arXiv preprint arXiv:2002.12410*, 2020.
- S. Bhojanapalli, A. Kyrillidis, and S. Sanghavi. Dropping convexity for faster semi-definite optimization. In *Conference on Learning Theory (COLT)*, pages 530–582, 2016.
- Y. Carmon, J. C. Duchi, O. Hinder, and A. Sidford. Lower bounds for finding stationary points I. *Mathematical Programming*, pages 1–50, 2019.
- C.-C. Chang and C.-J. Lin. LIBSVM: A library for support vector machines. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 2(3):1–27, 2011.
- M. Danilova, P. Dvurechensky, A. Gasnikov, E. Gorbunov, S. Guminov, D. Kamzolov, and I. Shibaev. Recent theoretical advances in non-convex optimization. *arXiv preprint arXiv:2012.06188*, 2020.
- R. Das, A. Hashemi, S. Sanghavi, and I. S. Dhillon. Improved convergence rates for non-convex federated learning with compression. *arXiv preprint arXiv:2012.04061*, 2020.
- C. Fang, C. Li, Z. Lin, and T. Zhang. Near-optimal non-convex optimization via stochastic path integrated differential estimator. *Advances in Neural Information Processing Systems (NeurIPS)*, 31:689, 2018.
- I. Goodfellow, Y. Bengio, A. Courville, and Y. Bengio. *Deep learning*, volume 1. MIT press Cambridge, 2016.
- E. Gorbunov, D. Kovalev, D. Makarenko, and P. Richtárik. Linearly converging error compensated SGD. *Advances in Neural Information Processing Systems (NeurIPS)*, 33, 2020.
- P. Goyal, P. Dollár, R. Girshick, P. Noordhuis, L. Wesolowski, A. Kyrola, A. Tulloch, Y. Jia, and K. He. Accurate, large minibatch SGD: training imagenet in 1 hour. *arXiv preprint arXiv:1706.02677*, 2017.
- F. Haddadpour, M. M. Kamani, A. Mokhtari, and M. Mahdavi. Federated learning with compression: Unified analysis and sharp guarantees. *arXiv preprint arXiv:2007.01154*, 2020.
- S. Horváth, C.-Y. Ho, Ľudovít Horváth, A. N. Sahu, M. Canini, and P. Richtárik. Natural compression for distributed deep learning. *arXiv preprint arXiv:1905.10988*, 2019.

- S. Horváth, D. Kovalev, K. Mishchenko, S. Stich, and P. Richtárik. Stochastic distributed learning with gradient quantization and variance reduction. *arXiv preprint arXiv:1904.05115*, 2019.
- S. Horváth, L. Lei, P. Richtárik, and M. I. Jordan. Adaptivity of stochastic gradient methods for nonconvex optimization. *arXiv preprint arXiv:2002.05359*, 2020.
- P. Kairouz, H. B. McMahan, B. Avent, A. Bellet, M. Bennis, A. N. Bhagoji, K. Bonawitz, Z. Charles, G. Cormode, R. Cummings, et al. Advances and open problems in federated learning. *arXiv preprint arXiv:1912.04977*, 2019.
- S. P. Karimireddy, Q. Rebjock, S. Stich, and M. Jaggi. Error feedback fixes signSGD and other gradient compression schemes. In *International Conference on Machine Learning (ICML)*, pages 3252–3261, 2019.
- S. P. Karimireddy, S. Kale, M. Mohri, S. Reddi, S. Stich, and A. T. Suresh. Scaffold: Stochastic controlled averaging for federated learning. In *International Conference on Machine Learning (ICML)*, pages 5132–5143. PMLR, 2020.
- P. Khanduri, P. Sharma, S. Kalle, S. Bulusu, K. Rajawat, and P. K. Varshney. Distributed stochastic non-convex optimization: Momentum-based variance reduction. *arXiv preprint arXiv:2005.00224*, 2020.
- A. Koloskova, T. Lin, S. U. Stich, and M. Jaggi. Decentralized deep learning with arbitrary communication compression. In *International Conference on Learning Representations (ICLR)*, 2020a.
- A. Koloskova, N. Loizou, S. Boreiri, M. Jaggi, and S. Stich. A unified theory of decentralized SGD with changing topology and local updates. In *International Conference on Machine Learning (ICML)*, pages 5381–5393, 2020b.
- J. Konečný, H. B. McMahan, F. X. Yu, P. Richtárik, A. T. Suresh, and D. Bacon. Federated learning: Strategies for improving communication efficiency. *arXiv preprint arXiv:1610.05492*, 2016.
- J. Konečný and P. Richtárik. Randomized distributed mean estimation: accuracy vs communication. *Frontiers in Applied Mathematics and Statistics*, 4(62):1–11, 2018.
- J. Konečný, H. B. McMahan, F. Yu, P. Richtárik, A. T. Suresh, and D. Bacon. Federated learning: strategies for improving communication efficiency. In *NIPS Private Multi-Party Machine Learning Workshop*, 2016.
- X. Li, W. Yang, S. Wang, and Z. Zhang. Communication efficient decentralized training with multiple local updates. *arXiv preprint arXiv:1910.09126*, 5, 2019.
- Z. Li and P. Richtárik. A unified analysis of stochastic gradient methods for nonconvex federated optimization. *arXiv preprint arXiv:2006.07013*, 2020.
- Z. Li, H. Bao, X. Zhang, and P. Richtárik. PAGE: A simple and optimal probabilistic gradient estimator for nonconvex optimization. *arXiv preprint arXiv:2008.10898*, 2020.

- X. Lian, C. Zhang, H. Zhang, C.-J. Hsieh, W. Zhang, and J. Liu. Can decentralized algorithms outperform centralized algorithms? a case study for decentralized parallel stochastic gradient descent. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 5330–5340, 2017.
- S. Łojasiewicz. A topological property of real analytic subsets. *Coll. du CNRS, Les équations aux dérivées partielles*, 117:87–89, 1963.
- C. Ma, K. Wang, Y. Chi, and Y. Chen. Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval and matrix completion. In *International Conference on Machine Learning (ICML)*, pages 3345–3354. PMLR, 2018.
- H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. Agüera y Arcas. Communication-efficient learning of deep networks from decentralized data. In *The 20th International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2017.
- K. Mishchenko, E. Gorbunov, M. Takáč, and P. Richtárik. Distributed learning with compressed gradient differences. *arXiv preprint arXiv:1901.09269*, 2019.
- K. Murty and S. Kabadi. Some np-complete problems in quadratic and nonlinear programming. *Mathematical programming*, 39(2):117–129, 1987.
- B. T. Polyak. Gradient methods for the minimisation of functionals. *USSR Computational Mathematics and Mathematical Physics*, 3(4):864–878, 1963.
- X. Qian, P. Richtárik, and T. Zhang. Error compensated distributed SGD can be accelerated. *arXiv preprint arXiv:2010.00091*, 2020.
- M. Safaryan, E. Shulgin, and P. Richtárik. Uncertainty principle for communication compression in distributed and federated learning and the search for an optimal compressor. *Information and Inference: A Journal of the IMA*, 2021.
- F. Seide, H. Fu, J. Droppo, G. Li, and D. Yu. 1-bit stochastic gradient descent and its application to data-parallel distributed training of speech dnns. In *Fifteenth Annual Conference of the International Speech Communication Association*, 2014.
- P. Sharma, S. Kalle, P. Khanduri, S. Bulusu, K. Rajawat, and P. K. Varshney. Parallel restarted spider—communication efficient distributed nonconvex optimization with optimal computation complexity. *arXiv preprint arXiv:1912.06036*, 2019.
- S. U. Stich and S. P. Karimireddy. The error-feedback framework: Better rates for SGD with delayed gradients and compressed updates. *Journal of Machine Learning Research*, 21:1–36, 2020.
- H. Sun, S. Lu, and M. Hong. Improving the sample and communication complexity for decentralized non-convex optimization: Joint gradient estimation and tracking. In *International Conference on Machine Learning (ICML)*, pages 9217–9228. PMLR, 2020.
- R. Sun. Optimization for deep learning: theory and algorithms. *arXiv preprint arXiv:1912.08957*, 2019.

- A. T. Suresh, F. X. Yu, S. Kumar, and H. B. McMahan. Distributed mean estimation with limited communication. In *International Conference on Machine Learning (ICML)*, 2017.
- Y. You, J. Li, S. Reddi, J. Hseu, S. Kumar, S. Bhojanapalli, X. Song, J. Demmel, K. Keutzer, and C.-J. Hsieh. Large batch optimization for deep learning: Training BERT in 76 minutes. In *International Conference on Learning Representations (ICLR)*, 2020.
- L. Zhao, M. Mammadov, and J. Yearwood. From convex to nonconvex: a loss function analysis for binary classification. In *2010 IEEE International Conference on Data Mining Workshops*, pages 1281–1288. IEEE, 2010.

Appendix

A Basic Facts and Auxiliary Results

A.1 Useful Properties of Expectations

Variance decomposition. For a random vector $\xi \in \mathbb{R}^d$ and any deterministic vector $x \in \mathbb{R}^d$ the variance can be decomposed as

$$\mathbf{E} \left[\|\xi - \mathbf{E}\xi\|^2 \right] = \mathbf{E} \left[\|\xi - x\|^2 \right] - \|\mathbf{E}\xi - x\|^2 \quad (12)$$

Tower property of mathematical expectation. For random variables $\xi, \eta \in \mathbb{R}^d$ we have

$$\mathbf{E} [\xi] = \mathbf{E} [\mathbf{E} [\xi \mid \eta]] \quad (13)$$

under assumption that all expectations in the expression above are well-defined.

A.2 One Lemma

In this section, we formulate a lemma from [Li et al., 2020] that hold in our settings as well. We omit the proof of this lemmas since it is identical to the one from [Li et al., 2020].

Lemma A.1 (Lemma 2 from [Li et al., 2020]). *Assume that function f is L -smooth and $x^{k+1} = x^k - \gamma g^k$. Then*

$$f(x^{k+1}) \leq f(x^k) - \frac{\gamma}{2} \|\nabla f(x^k)\|^2 - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2. \quad (14)$$

B Missing Proofs for MARINA

B.1 Generally Non-Convex Problems

In this section, we provide the full statement of Theorem 2.1 together with the proof of this result.

Theorem B.1 (Theorem 2.1). *Let Assumptions 1.1 and 1.2 be satisfied and*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right)}, \quad (15)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$. Then after K iterations of MARINA we have

$$\mathbf{E} \left[\|\nabla f(\hat{x}^K)\|^2 \right] \leq \frac{2\Delta_0}{\gamma K}, \quad (16)$$

where $\hat{x}^K$ is chosen uniformly at random from $x^0, \dots, x^{K-1}$ and $\Delta_0 = f(x^0) - f_*$. That is, after

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right) \quad (17)$$

iterations MARINA produces such a point $\hat{x}^K$ that $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$. Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right), \quad (18)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof of Theorem 2.1. The scheme of the proof is similar to the proof of Theorem 1 from [Li et al., 2020]. From Lemma A.1 we have

$$\begin{aligned} \mathbf{E}[f(x^{k+1})] &\leq \mathbf{E}[f(x^k)] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\ &\quad + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \end{aligned} \quad (19)$$

Next, we need to derive an upper bound for $\mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2]$. By definition of g^{k+1} we have

$$g^{k+1} = \begin{cases} \nabla f(x^{k+1}) & \text{with probability } p, \\ g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) & \text{with probability } 1-p. \end{cases}$$

Using this, variance decomposition (12) and tower property (13) we derive:

$$\begin{aligned} \mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\stackrel{(13)}{=} (1-p) \mathbf{E} \left[\left\| g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) - \nabla f(x^{k+1}) \right\|^2 \right] \\ &\stackrel{(13),(12)}{=} (1-p) \mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n \mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) - \nabla f(x^{k+1}) + \nabla f(x^k) \right\|^2 \right] \\ &\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \end{aligned}$$

Since $\mathcal{Q}(\nabla f_1(x^{k+1}) - \nabla f_1(x^k)), \dots, \mathcal{Q}(\nabla f_n(x^{k+1}) - \nabla f_n(x^k))$ are independent random vectors for fixed x^k and x^{k+1} we have

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &= (1-p)\mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n (\mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) - \nabla f_i(x^{k+1}) + \nabla f_i(x^k)) \right\|^2 \right] \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{1-p}{n^2} \sum_{i=1}^n \mathbf{E} [\|\mathcal{Q}(\nabla f_i(x^{k+1}) - \nabla f_i(x^k)) - \nabla f_i(x^{k+1}) + \nabla f_i(x^k)\|^2] \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(3)}{\leq} \frac{(1-p)\omega}{n^2} \sum_{i=1}^n \mathbf{E} [\|\nabla f_i(x^{k+1}) - \nabla f_i(x^k)\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2].
\end{aligned}$$

Using L -smoothness (2) of f_i together with the tower property (13) we get

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\leq \frac{(1-p)\omega}{n^2} \sum_{i=1}^n L_i^2 \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{(1-p)\omega L^2}{n} \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2]. \tag{20}
\end{aligned}$$

Next, we introduce new notation: $\Phi_k = f(x^k) - f_* + \frac{\gamma}{2p}\|g^k - \nabla f(x^k)\|^2$. Using this and inequalities (19) and (20) we establish the following inequality:

$$\begin{aligned}
\mathbf{E} [\Phi_{k+1}] &\leq \mathbf{E} \left[f(x^k) - f_* - \frac{\gamma}{2} \|\nabla f(x^k)\|^2 - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\
&\quad + \frac{\gamma}{2p} \mathbf{E} \left[\frac{(1-p)\omega L^2}{n} \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\
&= \mathbf{E} [\Phi_k] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2] + \left(\frac{\gamma(1-p)\omega L^2}{2pn} - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\
&\stackrel{(15)}{\leq} \mathbf{E} [\Phi_k] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2], \tag{21}
\end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)\omega L^2}{2pn} - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ following from (15). Summing up inequalities (21) for $k = 0, 1, \dots, K-1$ and rearranging the terms we derive

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbf{E} [\|\nabla f(x^k)\|^2] \leq \frac{2}{\gamma K} \sum_{k=0}^{K-1} (\mathbf{E} [\Phi_k] - \mathbf{E} [\Phi_{k+1}]) = \frac{2(\mathbf{E} [\Phi_0] - \mathbf{E} [\Phi_K])}{\gamma K} = \frac{2\Delta_0}{\gamma K},$$

since $g^0 = \nabla f(x^0)$ and $\Phi_{k+1} \geq 0$. Finally, using the tower property (13) and the definition of $\hat{x}^K$ we obtain (16) that implies (17) and (18). $\square$

Corollary B.1 (Corollary 2.1). *Let the assumptions of Theorem 2.1 hold and $p = \frac{\zeta_{\mathcal{Q}}}{d}$, where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right)},$$

then MARINA requires

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right) \right)$$

iterations/communication rounds to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is

$$\mathcal{O} \left(d + \frac{\Delta_0 L}{\varepsilon^2} \left(\zeta_{\mathcal{Q}} + \sqrt{\frac{\omega \zeta_{\mathcal{Q}}}{n} (d - \zeta_{\mathcal{Q}})} \right) \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof of Corollary 2.1. The choice of $p = \frac{\zeta_{\mathcal{Q}}}{d}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \frac{d}{\zeta_{\mathcal{Q}}} - 1, \\ pd + (1-p)\zeta_{\mathcal{Q}} &\leq \zeta_{\mathcal{Q}} + \left(1 - \frac{\zeta_{\mathcal{Q}}}{d}\right) \cdot \zeta_{\mathcal{Q}} \leq 2\zeta_{\mathcal{Q}}. \end{aligned}$$

Plugging these relations in (15), (17), and (18) we get that if

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right)},$$

then MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right) \right) \end{aligned}$$

iterations/communication rounds in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is

$$\begin{aligned} d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right) \\ &= \mathcal{O} \left(d + \frac{\Delta_0 L}{\varepsilon^2} \left(\zeta_{\mathcal{Q}} + \sqrt{\frac{\omega \zeta_{\mathcal{Q}}}{n} (d - \zeta_{\mathcal{Q}})} \right) \right) \end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

B.2 Convergence Results Under Polyak-Łojasiewicz condition

In this section, we provide the full statement of Theorem 2.2 together with the proof of this result.

Theorem B.2 (Theorem 2.2). *Let Assumptions 1.1, 1.2 and 2.1 be satisfied and*

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2(1-p)\omega}{pn}} \right)}, \frac{p}{2\mu} \right\}, \quad (22)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$. Then after K iterations of MARINA we have

$$\mathbf{E} [f(x^K) - f(x^*)] \leq (1 - \gamma\mu)^K \Delta_0, \quad (23)$$

where $\Delta_0 = f(x^0) - f(x^*)$. That is, after

$$K = \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \quad (24)$$

iterations MARINA produces such a point x^K that $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$. Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right), \quad (25)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof of Theorem 2.2. The proof is very similar to the proof of Theorem 2.1. From Lemma A.1 and PL condition we have

$$\begin{aligned} \mathbf{E}[f(x^{k+1}) - f(x^*)] &\leq \mathbf{E}[f(x^k) - f(x^*)] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\quad + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\ &\stackrel{(4)}{\leq} (1 - \gamma\mu) \mathbf{E} [f(x^k) - f(x^*)] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2]. \end{aligned}$$

Using the same arguments as in the proof of (20) we obtain

$$\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] \leq \frac{(1-p)\omega L^2}{n} \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p) \mathbf{E} [\|g^k - \nabla f(x^k)\|^2].$$

Putting all together we derive that the sequence $\Phi_k = f(x^k) - f(x^*) + \frac{\gamma}{p} \|g^k - \nabla f(x^k)\|^2$ satisfies

$$\begin{aligned} \mathbf{E} [\Phi_{k+1}] &\leq \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \frac{\gamma}{p} \mathbf{E} \left[\frac{(1-p)\omega L^2}{n} \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\ &= \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) + \left(\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \right) \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \left(\frac{\gamma(1-p)\omega L^2}{pn} - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\stackrel{(22)}{\leq} (1 - \gamma\mu) \mathbf{E} [\Phi_k], \end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)\omega L^2}{pn} - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ and $\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \leq (1 - \gamma\mu)\frac{\gamma}{p}$ following from (22). Unrolling the recurrence and using $g^0 = \nabla f(x^0)$ we obtain

$$\mathbf{E} [f(x^K) - f(x^*)] \leq \mathbf{E} [\Phi_K] \leq (1 - \gamma\mu)^K \Phi_0 = (1 - \gamma\mu)^K (f(x^0) - f(x^*))$$

that implies (24) and (25). $\square$

Corollary B.2. *Let the assumptions of Theorem 2.2 hold and $p = \frac{\zeta_{\mathcal{Q}}}{d}$, where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right)}, \frac{p}{2\mu} \right\},$$

then MARINA requires

$$K = \mathcal{O} \left(\max \left\{ \frac{d}{\zeta_{\mathcal{Q}}}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

iterations/communication rounds to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost per worker is

$$\mathcal{O} \left(d + \max \left\{ d, \frac{L}{\mu} \left(\zeta_{\mathcal{Q}} + \sqrt{\frac{\omega \zeta_{\mathcal{Q}}}{n} (d - \zeta_{\mathcal{Q}})} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof. The choice of $p = \frac{\zeta_{\mathcal{Q}}}{d}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \frac{d}{\zeta_{\mathcal{Q}}} - 1, \\ pd + (1-p)\zeta_{\mathcal{Q}} &\leq \zeta_{\mathcal{Q}} + \left(1 - \frac{\zeta_{\mathcal{Q}}}{d} \right) \cdot \zeta_{\mathcal{Q}} \leq 2\zeta_{\mathcal{Q}}. \end{aligned}$$

Plugging these relations in (22), (24), and (25) we get that if

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right)}, \frac{p}{2\mu} \right\},$$

then MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(\max \left\{ \frac{d}{\zeta_{\mathcal{Q}}}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega}{n} \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

iterations/communication rounds in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost per worker is

$$\begin{aligned} d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)\omega}{pn}} \right) \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(d + \max \left\{ d, \frac{L}{\mu} \left(\zeta_{\mathcal{Q}} + \sqrt{\frac{\omega \zeta_{\mathcal{Q}}}{n} (d - \zeta_{\mathcal{Q}})} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

C Missing Proofs for VR-MARINA

C.1 Finite Sum Case

C.1.1 Generally Non-Convex Problems

In this section, we provide the full statement of Theorem 3.1 together with the proof of this result.

Theorem C.1 (Theorem 3.1). *Consider the finite sum case (1)+(5). Let Assumptions 1.1, 1.2 and 3.1 be satisfied and*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \quad (26)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$. Then after K iterations of VR-MARINA we have

$$\mathbf{E} \left[\|\nabla f(\hat{x}^K)\|^2 \right] \leq \frac{2\Delta_0}{\gamma K}, \quad (27)$$

where $\hat{x}^K$ is chosen uniformly at random from $x^0, \dots, x^{K-1}$ and $\Delta_0 = f(x^0) - f_*$. That is, after

$$K = \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) \right) \quad (28)$$

iterations VR-MARINA produces such a point $\hat{x}^K$ that $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total number of stochastic oracle calls per node equals

$$m + K(pm + 2(1-p)b') = \mathcal{O} \left(m + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pm + (1-p)b') \right). \quad (29)$$

Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right), \quad (30)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof of Theorem 3.1. The proof of this theorem is a generalization of the proof of Theorem 2.1. From Lemma A.1 we have

$$\mathbf{E}[f(x^{k+1})] \leq \mathbf{E}[f(x^k)] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \quad (31)$$

Next, we need to derive an upper bound for $\mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2]$. Since $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$ we get the following representation of g^{k+1} :

$$g^{k+1} = \begin{cases} \nabla f(x^{k+1}) & \text{with probability } p, \\ g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) & \text{with probability } 1-p. \end{cases}$$

Using this, variance decomposition (12) and tower property (13) we derive:

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\stackrel{(13)}{=} (1-p)\mathbf{E} \left[\left\| g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) - \nabla f(x^{k+1}) \right\|^2 \right] \\
&\stackrel{(13),(12)}{=} (1-p)\mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) - \nabla f(x^{k+1}) + \nabla f(x^k) \right\|^2 \right] \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2].
\end{aligned}$$

Next, we use the notation: $\tilde{\Delta}_i^k = \frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k))$ and $\Delta_i^k = \nabla f_i(x^{k+1}) - \nabla f_i(x^k)$. These vectors satisfy $\mathbf{E} [\tilde{\Delta}_i^k \mid x^k, x^{k+1}] = \Delta_i^k$ for all $i \in [n]$. Moreover, $\mathcal{Q}(\tilde{\Delta}_1^k), \dots, \mathcal{Q}(\tilde{\Delta}_n^k)$ are independent random vectors for fixed x^k and x^{k+1} . These observations imply

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &= (1-p)\mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n (\mathcal{Q}(\tilde{\Delta}_i^k) - \Delta_i^k) \right\|^2 \right] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{1-p}{n^2} \sum_{i=1}^n \mathbf{E} [\|\mathcal{Q}(\tilde{\Delta}_i^k) - \tilde{\Delta}_i^k + \tilde{\Delta}_i^k - \Delta_i^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(12)}{=} \frac{1-p}{n^2} \sum_{i=1}^n \left(\mathbf{E} [\|\mathcal{Q}(\tilde{\Delta}_i^k) - \tilde{\Delta}_i^k\|^2] + \mathbf{E} [\|\tilde{\Delta}_i^k - \Delta_i^k\|^2] \right) \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(3)}{=} \frac{1-p}{n^2} \sum_{i=1}^n \left(\omega \mathbf{E} [\|\tilde{\Delta}_i^k\|^2] + \mathbf{E} [\|\tilde{\Delta}_i^k - \Delta_i^k\|^2] \right) + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(12)}{=} \frac{1-p}{n^2} \sum_{i=1}^n \left(\omega \mathbf{E} [\|\Delta_i^k\|^2] + (1+\omega) \mathbf{E} [\|\tilde{\Delta}_i^k - \Delta_i^k\|^2] \right) \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2].
\end{aligned}$$

Using L -smoothness (2) and average $\mathcal{L}$ -smoothness (7) of f_i together with the tower property (13) we get

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\leq \frac{1-p}{n^2} \sum_{i=1}^n \left(\omega L_i^2 + \frac{(1+\omega)\mathcal{L}_i^2}{b'} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2].
\end{aligned} \tag{32}$$

Next, we introduce new notation: $\Phi_k = f(x^k) - f_* + \frac{\gamma}{2p} \|g^k - \nabla f(x^k)\|^2$. Using this and inequalities (31)

and (32) we establish the following inequality:

$$\begin{aligned}
\mathbf{E}[\Phi_{k+1}] &\leq \mathbf{E} \left[f(x^k) - f_* - \frac{\gamma}{2} \|\nabla f(x^k)\|^2 - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\
&\quad + \frac{\gamma}{2p} \mathbf{E} \left[\frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\
&= \mathbf{E}[\Phi_k] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] + \left(\frac{\gamma(1-p)}{2pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\
&\stackrel{(26)}{\leq} \mathbf{E}[\Phi_k] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2], \tag{33}
\end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)}{2pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ following from (26). Summing up inequalities (33) for $k = 0, 1, \dots, K-1$ and rearranging the terms we derive

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbf{E}[\|\nabla f(x^k)\|^2] \leq \frac{2}{\gamma K} \sum_{k=0}^{K-1} (\mathbf{E}[\Phi_k] - \mathbf{E}[\Phi_{k+1}]) = \frac{2(\mathbf{E}[\Phi_0] - \mathbf{E}[\Phi_K])}{\gamma K} = \frac{2\Delta_0}{\gamma K},$$

since $g^0 = \nabla f(x^0)$ and $\Phi_{k+1} \geq 0$. Finally, using the tower property (13) and the definition of $\hat{x}^K$ we obtain (27) that implies (28), (29), and (30). $\square$

Remark C.1 (About batchsizes dissimilarity). *We notice that our analysis can be easily extended to handle the version of VR-MARINA with different batchsizes $b'_1, \dots, b'_n$ on different workers, i.e., when $|I'_{i,k}| = b'_i$ and $\tilde{\Delta}_i^k = \frac{1}{b'_i} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k))$. In this case, the statement of Theorem 3.1 remains the same with the small modification: instead of $\frac{\mathcal{L}^2}{b'}$ the complexity bounds will have $\frac{1}{n} \sum_{i=1}^n \frac{\mathcal{L}_i^2}{b'_i}$.*

Corollary C.1 (Corollary 3.1). *Let the assumptions of Theorem 3.1 hold and $p = \min \left\{ \frac{\zeta_Q}{d}, \frac{b'}{m+b'} \right\}$, where $b' \leq m$ and ζ_Q is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\{d/\zeta_Q - 1, m/b'\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}},$$

then VR-MARINA requires

$$\mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_Q - 1, m/b'\}}{n}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_Q - 1, m/b'\}}{nb'}} \right) \right)$$

iterations/communication rounds,

$$\mathcal{O} \left(m + \frac{\Delta_0}{\varepsilon^2} \left(L \left(b' + \sqrt{\frac{\omega \max\{(d/\zeta_Q - 1)(b')^2, mb'\}}{n}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{(d/\zeta_Q - 1)b', m\}}{n}} \right) \right)$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is

$$\mathcal{O} \left(d + \frac{\Delta_0 \zeta_Q}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_Q - 1, m/b'\}}{n}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_Q - 1, m/b'\}}{nb'}} \right) \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof of Corollary 3.1. The choice of $p = \min \left\{ \frac{\zeta_{\mathcal{Q}}}{d}, \frac{b'}{m+b'} \right\}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{m}{b'} \right\}, \\ pm + (1-p)b' &\leq \frac{2mb'}{m+b'} \leq 2b', \\ pd + (1-p)\zeta_{\mathcal{Q}} &\leq \frac{\zeta_{\mathcal{Q}}}{d} \cdot d + \left(1 - \frac{\zeta_{\mathcal{Q}}}{d}\right) \cdot \zeta_{\mathcal{Q}} \leq 2\zeta_{\mathcal{Q}}. \end{aligned}$$

Plugging these relations in (26), (28), (29) and (30) and using $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ we get that if

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}},$$

then VR-MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{L^2 \frac{\omega \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n} + \mathcal{L}^2 \frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{nb'}} \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{nb'}} \right) \right) \end{aligned}$$

iterations/communication rounds and

$$\begin{aligned} m + K(pm + 2(1-p)b') &= \mathcal{O} \left(m + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pm + (1-p)b') \right) \\ &= \mathcal{O} \left(m + \frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{nb'}} \right) b' \right) \\ &= \mathcal{O} \left(m + \frac{\Delta_0}{\varepsilon^2} \left(L \left(b' + \sqrt{\frac{\omega \max\{(d/\zeta_{\mathcal{Q}}-1)(b')^2, mb'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{(d/\zeta_{\mathcal{Q}}-1)b', m\}}{n}} \right) \right) \end{aligned}$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is

$$\begin{aligned} d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right) \\ &= \mathcal{O} \left(d + \frac{\Delta_0 \zeta_{\mathcal{Q}}}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}}-1, m/b'\}}{nb'}} \right) \right) \end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

C.1.2 Convergence Results Under Polyak-Łojasiewicz condition

In this section, we provide an analysis of VR-MARINA under Polyak-Łojasiewicz condition in the finite sum case.

Theorem C.2. *Consider the finite sum case (1)+(5). Let Assumptions 1.1, 1.2, 3.1 and 2.1 be satisfied and*

$$\gamma \leq \min \left\{ \frac{1}{L + \sqrt{\frac{2(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \frac{p}{2\mu} \right\}, \quad (34)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$. Then after K iterations of VR-MARINA we have

$$\mathbf{E} [f(x^K) - f(x^*)] \leq (1 - \gamma\mu)^K \Delta_0, \quad (35)$$

where $\Delta_0 = f(x^0) - f(x^*)$. That is, after

$$K = \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \quad (36)$$

iterations VR-MARINA produces such a point x^K that $\mathbf{E} [f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total number of stochastic oracle calls per node equals

$$m + K(pm + 2(1-p)b') = \mathcal{O} \left(m + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pm + (1-p)b') \log \frac{\Delta_0}{\varepsilon} \right). \quad (37)$$

Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right), \quad (38)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof. The proof is very similar to the proof of Theorem 3.1. From Lemma A.1 and PL condition we have

$$\begin{aligned} \mathbf{E}[f(x^{k+1}) - f(x^*)] &\leq \mathbf{E}[f(x^k) - f(x^*)] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\quad + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\ &\stackrel{(4)}{\leq} (1 - \gamma\mu) \mathbf{E} [f(x^k) - f(x^*)] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2]. \end{aligned}$$

Using the same arguments as in the proof of (32) we obtain

$$\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] \leq \frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p) \mathbf{E} [\|g^k - \nabla f(x^k)\|^2].$$

Putting all together we derive that the sequence $\Phi_k = f(x^k) - f(x^*) + \frac{\gamma}{p} \|g^k - \nabla f(x^k)\|^2$ satisfies

$$\begin{aligned} \mathbf{E} [\Phi_{k+1}] &\leq \mathbf{E} \left[(1-\gamma\mu)(f(x^k) - f(x^*)) - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \frac{\gamma}{p} \mathbf{E} \left[\frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\ &= \mathbf{E} \left[(1-\gamma\mu)(f(x^k) - f(x^*)) + \left(\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \right) \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \left(\frac{\gamma(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\stackrel{(34)}{\leq} (1-\gamma\mu) \mathbf{E} [\Phi_k], \end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ and $\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \leq (1-\gamma\mu)\frac{\gamma}{p}$ following from (34). Unrolling the recurrence and using $g^0 = \nabla f(x^0)$ we obtain

$$\mathbf{E} [f(x^{k+1}) - f(x^*)] \leq \mathbf{E} [\Phi_{k+1}] \leq (1-\gamma\mu)^{k+1} \Phi_0 = (1-\gamma\mu)^{k+1} (f(x^0) - f(x^*))$$

that implies (36), (37), and (38). $\square$

Corollary C.2. *Let the assumptions of Theorem C.2 hold and $p = \min \left\{ \frac{\zeta_{\mathcal{Q}}}{d}, \frac{b'}{m+b'} \right\}$, where $b' \leq m$ and $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \min \left\{ \frac{1}{L + \sqrt{\frac{2 \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \frac{p}{2\mu} \right\},$$

then VR-MARINA requires

$$\mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n}} \right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

iterations/communication rounds,

$$\mathcal{O} \left(m + \max \left\{ \frac{b'}{p}, \frac{L}{\mu} \left(b' + \sqrt{\frac{\omega \max\{(d/\zeta_{\mathcal{Q}} - 1)(b')^2, mb'\}}{n}} \right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{(d/\zeta_{\mathcal{Q}} - 1)b', m\}}{n}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost per worker is

$$\mathcal{O} \left(d + \zeta_{\mathcal{Q}} \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n}} \right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof. The choice of $p = \min \left\{ \frac{\zeta_{\mathcal{Q}}}{d}, \frac{b'}{m+b'} \right\}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{m}{b'} \right\}, \\ pm + (1-p)b' &\leq \frac{2mb'}{m+b'} \leq 2b', \\ pd + (1-p)\zeta_{\mathcal{Q}} &\leq \frac{\zeta_{\mathcal{Q}}}{d} \cdot d + \left(1 - \frac{\zeta_{\mathcal{Q}}}{d}\right) \cdot \zeta_{\mathcal{Q}} \leq 2\zeta_{\mathcal{Q}}. \end{aligned}$$

Plugging these relations in (34), (36), (37) and (38) and using $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ we get that if

$$\gamma \leq \min \left\{ \frac{1}{L + \sqrt{\frac{2 \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \frac{p}{2\mu} \right\},$$

then VR-MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}}{\mu} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L + \sqrt{L^2 \frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n} + \mathcal{L}^2 \frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}}}{\mu} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n}} \right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

iterations/communication rounds and

$$\begin{aligned} m + K(pm + 2(1-p)b') &= \mathcal{O} \left(m + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}}{\mu} \right\} (pm + (1-p)b') \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(m + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}} \right\} b' \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(m + \max \left\{ \frac{b'}{p}, \frac{L}{\mu} \left(b' + \sqrt{\frac{\omega \max\{(d/\zeta_{\mathcal{Q}} - 1)(b')^2, mb'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max\{(d/\zeta_{\mathcal{Q}} - 1)b', m\}}{n}} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected

total communication cost per worker is

$$\begin{aligned}
d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right) \\
&= \mathcal{O} \left(d + \zeta_{\mathcal{Q}} \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max \{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{n}} \right) \right. \right. \\
&\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max \{d/\zeta_{\mathcal{Q}} - 1, m/b'\}}{nb'}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)
\end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

C.2 Online Case

C.2.1 Generally Non-Convex Problems

In this section, we provide the full statement of Theorem 3.2 together with the proof of this result.

Theorem C.3 (Theorem 3.2). *Consider the finite sum case (1)+(6). Let Assumptions 1.1, 1.2 and 3.1 be satisfied and*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \quad (39)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$. Then after K iterations of VR-MARINA we have

$$\mathbf{E} \left[\|\nabla f(\hat{x}^K)\|^2 \right] \leq \frac{2\Delta_0}{\gamma K} + \frac{\sigma^2}{nb}, \quad (40)$$

where $\hat{x}^K$ is chosen uniformly at random from $x^0, \dots, x^{K-1}$ and $\Delta_0 = f(x^0) - f_*$. That is, after

$$K = \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) \right) \quad (41)$$

iterations with $b = \Theta(\frac{\sigma^2}{n\varepsilon^2})$ VR-MARINA produces such a point $\hat{x}^K$ that $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total number of stochastic oracle calls per node equals

$$b + K(pb + 2(1-p)b') = \mathcal{O} \left(\frac{\sigma^2}{n\varepsilon^2} + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) \left(p \frac{\sigma^2}{n\varepsilon^2} + (1-p)b' \right) \right). \quad (42)$$

Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right), \quad (43)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof of Theorem 3.2. The proof follows the same steps as the proof of Theorem 3.1. From Lemma A.1 we have

$$\mathbf{E}[f(x^{k+1})] \leq \mathbf{E}[f(x^k)] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2}\right) \mathbf{E}[\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \quad (44)$$

Next, we need to derive an upper bound for $\mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2]$. Since $g^{k+1} = \frac{1}{n} \sum_{i=1}^n g_i^{k+1}$ we get the following representation of g^{k+1} :

$$g^{k+1} = \begin{cases} \frac{1}{nb} \sum_{i=1}^n \sum_{j \in I_{i,k}} \nabla f_{\xi_{ij}^k}(x^{k+1}) & \text{with probability } p, \\ g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f_{\xi_{ij}^k}(x^k)) \right) & \text{with probability } 1 - p. \end{cases}$$

Using this, variance decomposition (12), tower property (13), and independence of ξ_{ij}^k for $i \in [n]$, $j \in I_{i,k}$ we derive:

$$\begin{aligned} \mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\stackrel{(13)}{=} (1-p) \mathbf{E} \left[\left\| g^k + \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) - \nabla f(x^{k+1}) \right\|^2 \right] \\ &\quad + \frac{p}{n^2 b^2} \mathbf{E} \left[\left\| \sum_{i=1}^n \sum_{j \in I_{i,k}} (\nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f(x^{k+1})) \right\|^2 \right] \\ &\stackrel{(13),(12)}{=} (1-p) \mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) - \nabla f(x^{k+1}) + \nabla f(x^k) \right\|^2 \right] \\ &\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2] + \frac{p}{n^2 b^2} \sum_{i=1}^n \sum_{j \in I_{i,k}} \mathbf{E} \left[\left\| \nabla f_{\xi_{ij}^k}(x^{k+1}) - \nabla f(x^{k+1}) \right\|^2 \right] \\ &\stackrel{(13),(10)}{=} (1-p) \mathbf{E} \left[\left\| \frac{1}{n} \sum_{i=1}^n \mathcal{Q} \left(\frac{1}{b'} \sum_{j \in I'_{i,k}} (\nabla f_{ij}(x^{k+1}) - \nabla f_{ij}(x^k)) \right) - \nabla f(x^{k+1}) + \nabla f(x^k) \right\|^2 \right] \\ &\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2] + \frac{p\sigma^2}{nb}, \end{aligned}$$

where $\sigma^2 = \frac{1}{n} \sum_{i=1}^n \sigma_i^2$. Applying the same arguments as in the proof of inequality (32) we obtain

$$\begin{aligned} \mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\leq \frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\ &\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2] + \frac{p\sigma^2}{nb}. \end{aligned} \quad (45)$$

Next, we introduce new notation: $\Phi_k = f(x^k) - f_* + \frac{\gamma}{2p} \|g^k - \nabla f(x^k)\|^2$. Using this and inequalities (44)

and (45) we establish the following inequality:

$$\begin{aligned}
\mathbf{E}[\Phi_{k+1}] &\leq \mathbf{E}\left[f(x^k) - f_* - \frac{\gamma}{2}\|\nabla f(x^k)\|^2 - \left(\frac{1}{2\gamma} - \frac{L}{2}\right)\|x^{k+1} - x^k\|^2 + \frac{\gamma}{2}\|g^k - \nabla f(x^k)\|^2\right] \\
&\quad + \frac{\gamma}{2p}\mathbf{E}\left[\frac{1-p}{n}\left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'}\right)\|x^{k+1} - x^k\|^2 + (1-p)\|g^k - \nabla f(x^k)\|^2 + \frac{p\sigma^2}{nb}\right] \\
&= \mathbf{E}[\Phi_k] - \frac{\gamma}{2}\mathbf{E}[\|\nabla f(x^k)\|^2] + \left(\frac{\gamma(1-p)}{2pn}\left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'}\right) - \frac{1}{2\gamma} + \frac{L}{2}\right)\mathbf{E}[\|x^{k+1} - x^k\|^2] \\
&\quad + \frac{\gamma\sigma^2}{2nb} \\
&\stackrel{(39)}{\leq} \mathbf{E}[\Phi_k] - \frac{\gamma}{2}\mathbf{E}[\|\nabla f(x^k)\|^2] + \frac{\gamma\sigma^2}{2nb},
\end{aligned} \tag{46}$$

where in the last inequality we use $\frac{\gamma(1-p)}{2pn}\left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'}\right) - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ following from (39). Summing up inequalities (46) for $k = 0, 1, \dots, K-1$ and rearranging the terms we derive

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbf{E}[\|\nabla f(x^k)\|^2] \leq \frac{2}{\gamma K} \sum_{k=0}^{K-1} (\mathbf{E}[\Phi_k] - \mathbf{E}[\Phi_{k+1}]) + \frac{\sigma^2}{nb} = \frac{2(\mathbf{E}[\Phi_0] - \mathbf{E}[\Phi_K])}{\gamma K} + \frac{\sigma^2}{nb} = \frac{2\Delta_0}{\gamma K} + \frac{\sigma^2}{nb},$$

since $g^0 = \nabla f(x^0)$ and $\Phi_{k+1} \geq 0$. Finally, using the tower property (13) and the definition of $\hat{x}^K$ we obtain (40) that implies (41), (42), and (43). $\square$

Remark C.2 (About batchsizes dissimilarity). *Similarly to the finite sum case, our analysis can be easily extended to handle the version of VR-MARINA with different batchsizes $b_1, \dots, b_n$ and $b'_1, \dots, b'_n$ on different workers, i.e., when $|I_{i,k}| = b_i$, $|I'_{i,k}| = b'_i$ for $i \in [n]$. In this case, the statement of Theorem 3.2 remains the same with the small modification: instead of $\frac{\mathcal{L}^2}{b'}$ the complexity bounds will have $\frac{1}{n} \sum_{i=1}^n \frac{\mathcal{L}_i^2}{b'_i}$, and instead of the requirement $b = \Theta\left(\frac{\sigma^2}{n\varepsilon}\right)$ it will have $\frac{1}{n^2} \sum_{i=1}^n \frac{\sigma_i^2}{b_i} = \Theta(\varepsilon^2)$.*

Corollary C.3 (Corollary 3.2). *Let the assumptions of Theorem 3.2 hold and $p = \min\left\{\frac{\zeta_Q}{d}, \frac{b'}{b+b'}\right\}$, where $b' \leq b$, $b = \Theta\left(\frac{\sigma^2}{n\varepsilon^2}\right)$ and ζ_Q is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\left\{\frac{d}{\zeta_Q} - 1, \frac{b'}{b+b'}\right\}}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'}\right)}},$$

then VR-MARINA requires

$$\mathcal{O}\left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega}{n} \max\left\{\frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2}\right\}}\right) + \mathcal{L} \sqrt{\frac{(1+\omega)}{nb'} \max\left\{\frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2}\right\}}\right)\right)$$

iterations/communication rounds and

$$\mathcal{O}\left(\frac{\sigma^2}{n\varepsilon^2} + \frac{\Delta_0 L b'}{\varepsilon^2} + \frac{\Delta_0 L}{\varepsilon^2} \sqrt{\frac{\omega b'}{n} \max\left\{\left(\frac{d}{\zeta_Q} - 1\right) b', \frac{\sigma^2}{n\varepsilon^2}\right\}} + \frac{\Delta_0 \mathcal{L}}{\varepsilon^2} \sqrt{\frac{1+\omega}{n} \max\left\{\left(\frac{d}{\zeta_Q} - 1\right) b', \frac{\sigma^2}{n\varepsilon^2}\right\}}\right)$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost per worker is

$$\mathcal{O}\left(d + \frac{\Delta_0 \zeta_Q}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega}{n} \max\left\{\frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2}\right\}}\right) + \mathcal{L} \sqrt{\frac{1+\omega}{nb'} \max\left\{\frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2}\right\}}\right)\right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof of Corollary 3.1. The choice of $p = \min \left\{ \frac{\zeta_Q}{d}, \frac{b'}{b+b'} \right\}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \max \left\{ \frac{d}{\zeta_Q} - 1, \frac{b}{b'} \right\}, \\ pm + (1-p)b' &\leq \frac{2mb'}{m+b'} \leq 2b', \\ pd + (1-p)\zeta_Q &\leq \frac{\zeta_Q}{d} \cdot d + \left(1 - \frac{\zeta_Q}{d}\right) \cdot \zeta_Q \leq 2\zeta_Q. \end{aligned}$$

Plugging these relations in (39), (41), (42) and (43) and using $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ we get that if

$$\gamma \leq \frac{1}{L + \sqrt{\frac{\max\{d/\zeta_Q - 1, b/b'\}}{n}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)},$$

then VR-MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{L^2 \frac{\omega \max\{d/\zeta_Q - 1, b/b'\}}{n} + \mathcal{L}^2 \frac{(1+\omega) \max\{d/\zeta_Q - 1, b/b'\}}{nb'}} \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega}{n} \max \left\{ \frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) + \mathcal{L} \sqrt{\frac{(1+\omega)}{nb'} \max \left\{ \frac{d}{\zeta_Q} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) \right) \end{aligned}$$

iterations/communication rounds and

$$\begin{aligned} b + K(pb + 2(1-p)b') &= \mathcal{O} \left(b + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \right) (pb + (1-p)b') \right) \\ &= \mathcal{O} \left(b + \frac{\Delta_0}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega \max\{d/\zeta_Q - 1, b/b'\}}{n}} \right) \right. \right. \\ &\quad \left. \left. + \mathcal{L} \sqrt{\frac{(1+\omega) \max\{d/\zeta_Q - 1, b/b'\}}{nb'}} \right) b' \right) \\ &= \mathcal{O} \left(\frac{\sigma^2}{n\varepsilon^2} + \frac{\Delta_0}{\varepsilon^2} \left(L \left(b' + \sqrt{\frac{\omega b'}{n} \max \left\{ \left(\frac{d}{\zeta_Q} - 1 \right) b', \frac{\sigma^2}{n\varepsilon^2} \right\}} \right) \right. \right. \\ &\quad \left. \left. + \mathcal{L} \sqrt{\frac{1+\omega}{n} \max \left\{ \left(\frac{d}{\zeta_Q} - 1 \right) b', \frac{\sigma^2}{n\varepsilon^2} \right\}} \right) \right) \end{aligned}$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected

total communication cost per worker is

$$\begin{aligned}
d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \frac{\Delta_0}{\varepsilon^2} \left(L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)} \right) (pd + (1-p)\zeta_{\mathcal{Q}}) \right) \\
&= \mathcal{O} \left(d + \frac{\Delta_0 \zeta_{\mathcal{Q}}}{\varepsilon^2} \left(L \left(1 + \sqrt{\frac{\omega}{n} \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) \right. \right. \\
&\quad \left. \left. + \mathcal{L} \sqrt{\frac{1+\omega}{nb'} \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\varepsilon^2} \right\}} \right) \right)
\end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

C.2.2 Convergence Results Under Polyak-Łojasiewicz condition

In this section, we provide an analysis of VR-MARINA under Polyak-Łojasiewicz condition in the online case.

Theorem C.4. *Consider the finite sum case (1)+(6). Let Assumptions 1.1, 1.2, 3.1, 2.1 and 3.3 be satisfied and*

$$\gamma \leq \min \left\{ \frac{1}{L + \sqrt{\frac{2(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}, \frac{p}{2\mu} \right\}, \quad (47)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$ and $\mathcal{L}^2 = \frac{1}{n} \sum_{i=1}^n \mathcal{L}_i^2$. Then after K iterations of VR-MARINA we have

$$\mathbf{E} [f(x^K) - f(x^*)] \leq (1 - \gamma\mu)^K \Delta_0 + \frac{\sigma^2}{nb\mu}, \quad (48)$$

where $\Delta_0 = f(x^0) - f(x^*)$. That is, after

$$K = \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} \log \frac{\Delta_0}{\varepsilon} \right) \quad (49)$$

iterations with $b = \Theta \left(\frac{\sigma^2}{n\mu\varepsilon} \right)$ VR-MARINA produces such a point x^K that $\mathbf{E} [f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total number of stochastic oracle calls per node equals

$$b + K(pb + 2(1-p)b') = \mathcal{O} \left(m + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pb + (1-p)b') \log \frac{\Delta_0}{\varepsilon} \right). \quad (50)$$

Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost per worker equals

$$d + K(pd + (1-p)\zeta_{\mathcal{Q}}) = \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right), \quad (51)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof. The proof is very similar to the proof of Theorem 3.2. From Lemma A.1 and PL condition we have

$$\begin{aligned}
\mathbf{E}[f(x^{k+1}) - f(x^*)] &\leq \mathbf{E}[f(x^k) - f(x^*)] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2}\right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\
&\quad + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(4)}{\leq} (1 - \gamma\mu) \mathbf{E}[f(x^k) - f(x^*)] - \left(\frac{1}{2\gamma} - \frac{L}{2}\right) \mathbf{E}[\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2].
\end{aligned}$$

Using the same arguments as in the proof of (45) we obtain

$$\begin{aligned}
\mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\leq \frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\
&\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2] + \frac{p\sigma^2}{nb}.
\end{aligned} \tag{52}$$

Putting all together we derive that the sequence $\Phi_k = f(x^k) - f(x^*) + \frac{\gamma}{p} \|g^k - \nabla f(x^k)\|^2$ satisfies

$$\begin{aligned}
\mathbf{E}[\Phi_{k+1}] &\leq \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) - \left(\frac{1}{2\gamma} - \frac{L}{2}\right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\
&\quad + \frac{\gamma}{p} \mathbf{E} \left[\frac{1-p}{n} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 + \frac{p\sigma^2}{nb} \right] \\
&= \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) + \left(\frac{\gamma}{2} + \frac{\gamma}{p}(1-p)\right) \|g^k - \nabla f(x^k)\|^2 \right] + \frac{\gamma\sigma^2}{nb} \\
&\quad + \left(\frac{\gamma(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] \\
&\stackrel{(34)}{\leq} (1 - \gamma\mu) \mathbf{E}[\Phi_k] + \frac{\gamma\sigma^2}{nb},
\end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right) - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ and $\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \leq (1 - \gamma\mu)\frac{\gamma}{p}$ following from (47). Unrolling the recurrence and using $g^0 = \nabla f(x^0)$ we obtain

$$\begin{aligned}
\mathbf{E}[f(x^K) - f(x^*)] &\leq \mathbf{E}[\Phi_K] \leq (1 - \gamma\mu)^K \Phi_0 + \frac{\gamma\sigma^2}{nb} \sum_{k=0}^{K-1} (1 - \gamma\mu)^k \\
&\leq (1 - \gamma\mu)^K (f(x^0) - f(x^*)) + \frac{\gamma\sigma^2}{nb} \sum_{k=0}^{\infty} (1 - \gamma\mu)^k \\
&\leq (1 - \gamma\mu)^K (f(x^0) - f(x^*)) + \frac{\sigma^2}{nb\mu}.
\end{aligned}$$

Together with $b = \Theta\left(\frac{\sigma^2}{n\mu\varepsilon}\right)$ it implies (49), (50), and (51). $\square$

Corollary C.4. *Let the assumptions of Theorem C.4 hold and $p = \min\left\{\frac{\zeta_{\mathcal{Q}}}{d}, \frac{b'}{b+b'}\right\}$, where $b' \leq b$ and $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \min \left\{ \frac{1}{L + \sqrt{\frac{2 \max\{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{n}} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}, \frac{p}{2\mu} \right\}$$

and

$$b = \Theta\left(\frac{\sigma^2}{n\mu\varepsilon}\right), \quad \sigma^2 = \frac{1}{n} \sum_{i=1}^n \sigma_i^2,$$

then VR-MARINA requires

$$\mathcal{O}\left(\max\left\{\frac{1}{p}, \frac{L}{\mu}\left(1 + \sqrt{\frac{\omega}{n} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{nb'} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right\} \log \frac{\Delta_0}{\varepsilon}\right)$$

iterations/communication rounds,

$$\begin{aligned} \mathcal{O}\left(\frac{\sigma^2}{n\mu\varepsilon} + \max\left\{\frac{b'}{p}, \frac{L}{\mu}\left(b' + \sqrt{\frac{\omega b'}{n} \max\left\{\left(\frac{d}{\zeta_{\mathcal{Q}}} - 1\right) b', \frac{\sigma^2}{n\mu\varepsilon}\right\}}\right) \right. \\ \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{n} \max\left\{\left(\frac{d}{\zeta_{\mathcal{Q}}} - 1\right) b', \frac{\sigma^2}{n\mu\varepsilon}\right\}}\right\} \log \frac{\Delta_0}{\varepsilon}\right) \end{aligned}$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost per worker is

$$\mathcal{O}\left(d + \zeta_{\mathcal{Q}} \max\left\{\frac{1}{p}, \frac{L}{\mu}\left(1 + \sqrt{\frac{\omega}{n} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{nb'} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right\} \log \frac{\Delta_0}{\varepsilon}\right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof. The choice of $p = \min\left\{\frac{\zeta_{\mathcal{Q}}}{d}, \frac{b'}{b+b'}\right\}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{b}{b'}\right\}, \\ pm + (1-p)b' &\leq \frac{2bb'}{b+b'} \leq 2b', \\ pd + (1-p)\zeta_{\mathcal{Q}} &\leq \frac{\zeta_{\mathcal{Q}}}{d} \cdot d + \left(1 - \frac{\zeta_{\mathcal{Q}}}{d}\right) \cdot \zeta_{\mathcal{Q}} \leq 2\zeta_{\mathcal{Q}}. \end{aligned}$$

Plugging these relations in (47), (49), (50) and (51) and using $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$ we get that if

$$\gamma \leq \min\left\{\frac{1}{L + \sqrt{\frac{2 \max\{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{n} (\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'})}}, \frac{p}{2\mu}\right\},$$

then VR-MARINA requires

$$\begin{aligned} K &= \mathcal{O}\left(\max\left\{\frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} (\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'})}}{\mu}\right\} \log \frac{\Delta_0}{\varepsilon}\right) \\ &= \mathcal{O}\left(\max\left\{\frac{1}{p}, \frac{L + \sqrt{L^2 \frac{\omega \max\{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{n} + \mathcal{L}^2 \frac{(1+\omega) \max\{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{nb'}}}{\mu}\right\} \log \frac{\Delta_0}{\varepsilon}\right) \\ &= \mathcal{O}\left(\max\left\{\frac{1}{p}, \frac{L}{\mu}\left(1 + \sqrt{\frac{\omega}{n} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right) + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{nb'} \max\left\{\frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu}\right\}}\right\} \log \frac{\Delta_0}{\varepsilon}\right) \end{aligned}$$

iterations/communication rounds and

$$\begin{aligned}
b + K(pb + 2(1-p)b') &= \mathcal{O} \left(b + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pb + (1-p)b') \log \frac{\Delta_0}{\varepsilon} \right) \\
&= \mathcal{O} \left(b + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega \max \{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{n}} \right) \right. \right. \\
&\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{(1+\omega) \max \{d/\zeta_{\mathcal{Q}} - 1, b/b'\}}{nb'}} \right\} b' \log \frac{\Delta_0}{\varepsilon} \right) \\
&= \mathcal{O} \left(\frac{\sigma^2}{n\mu\varepsilon} + \max \left\{ \frac{b'}{p}, \frac{L}{\mu} \left(b' + \sqrt{\frac{\omega b'}{n} \max \left\{ \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right) b', \frac{\sigma^2}{n\mu\varepsilon} \right\}} \right) \right. \right. \\
&\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{n} \max \left\{ \left(\frac{d}{\zeta_{\mathcal{Q}}} - 1 \right) b', \frac{\sigma^2}{n\mu\varepsilon} \right\}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)
\end{aligned}$$

stochastic oracle calls per node in expectation in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost per worker is

$$\begin{aligned}
d + K(pd + (1-p)\zeta_{\mathcal{Q}}) &= \mathcal{O} \left(d + \max \left\{ \frac{1}{p}, \frac{L + \sqrt{\frac{1-p}{pn} \left(\omega L^2 + \frac{(1+\omega)\mathcal{L}^2}{b'} \right)}}{\mu} \right\} (pd + (1-p)\zeta_{\mathcal{Q}}) \log \frac{\Delta_0}{\varepsilon} \right) \\
&= \mathcal{O} \left(d + \zeta_{\mathcal{Q}} \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{\omega}{n} \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu} \right\}} \right) \right. \right. \\
&\quad \left. \left. + \frac{\mathcal{L}}{\mu} \sqrt{\frac{1+\omega}{nb'} \max \left\{ \frac{d}{\zeta_{\mathcal{Q}}} - 1, \frac{\sigma^2}{nb'\mu} \right\}} \right\} \log \frac{\Delta_0}{\varepsilon} \right)
\end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

D Missing Proofs for PP-MARINA

D.1 Generally Non-Convex Problems

In this section, we provide the full statement of Theorem 4.1 together with the proof of this result.

Theorem D.1 (Theorem 4.1). *Let Assumptions 1.1 and 1.2 be satisfied and*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right)}, \quad (53)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$. Then after K iterations of PP-MARINA we have

$$\mathbf{E} \left[\|\nabla f(\hat{x}^K)\|^2 \right] \leq \frac{2\Delta_0}{\gamma K}, \quad (54)$$

where $\hat{x}^K$ is chosen uniformly at random from $x^0, \dots, x^{K-1}$ and $\Delta_0 = f(x^0) - f_*$. That is, after

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right) \quad (55)$$

iterations PP-MARINA produces such a point $\hat{x}^K$ that $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$. Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost (for all workers) equals

$$dn + K(pdn + (1-p)\zeta_{\mathcal{Q}}r) = \mathcal{O} \left(dn + \frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) (pdn + (1-p)\zeta_{\mathcal{Q}}r) \right), \quad (56)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof of Theorem 4.1. The proof is very similar to the proof of Theorem 3.1. From Lemma A.1 we have

$$\mathbf{E}[f(x^{k+1})] \leq \mathbf{E}[f(x^k)] - \frac{\gamma}{2} \mathbf{E}[\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E}[\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \quad (57)$$

Next, we need to derive an upper bound for $\mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2]$. By definition of g^{k+1} we have

$$g^{k+1} = \begin{cases} \nabla f(x^{k+1}) & \text{with probability } p, \\ g^k + \frac{1}{r} \sum_{i_k \in I'_k} \mathcal{Q}(\nabla f_{i_k}(x^{k+1}) - \nabla f_{i_k}(x^k)) & \text{with probability } 1-p. \end{cases}$$

Using this, variance decomposition (12) and tower property (13) we derive:

$$\begin{aligned} \mathbf{E}[\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\stackrel{(13)}{=} (1-p) \mathbf{E} \left[\left\| g^k + \frac{1}{r} \sum_{i_k \in I'_k} \mathcal{Q}(\nabla f_{i_k}(x^{k+1}) - \nabla f_{i_k}(x^k)) - \nabla f(x^{k+1}) \right\|^2 \right] \\ &\stackrel{(13),(12)}{=} (1-p) \mathbf{E} \left[\left\| \frac{1}{r} \sum_{i_k \in I'_k} \mathcal{Q}(\nabla f_{i_k}(x^{k+1}) - \nabla f_{i_k}(x^k)) - \nabla f(x^{k+1}) + \nabla f(x^k) \right\|^2 \right] \\ &\quad + (1-p) \mathbf{E}[\|g^k - \nabla f(x^k)\|^2]. \end{aligned}$$

Next, we use the notation: $\Delta_i^k = \nabla f_i(x^{k+1}) - \nabla f_i(x^k)$ for $i \in [n]$ and $\Delta^k = \nabla f(x^{k+1}) - \nabla f(x^k)$. These vectors satisfy $\mathbf{E} [\Delta_{i_k}^k \mid x^k, x^{k+1}] = \Delta^k$ for all $i_k \in I'_k$. Moreover, $\mathcal{Q}(\Delta_{i_k}^k)$ for $i_k \in I'_k$ are independent random vectors for fixed x^k and x^{k+1} . These observations imply

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &= (1-p)\mathbf{E} \left[\left\| \frac{1}{r} \sum_{i_k \in I'_k} (\mathcal{Q}(\Delta_{i_k}^k) - \Delta^k) \right\|^2 \right] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{1-p}{r} \mathbf{E} [\|\mathcal{Q}(\Delta_{i_k}^k) - \Delta_{i_k}^k + \Delta_{i_k}^k - \Delta^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(12)}{=} \frac{1-p}{r} \left(\mathbf{E} [\|\mathcal{Q}(\Delta_{i_k}^k) - \Delta_{i_k}^k\|^2] + \mathbf{E} [\|\Delta_{i_k}^k - \Delta^k\|^2] \right) \\
&\quad + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(3)}{=} \frac{1-p}{r} \left(\omega \mathbf{E} [\|\Delta_{i_k}^k\|^2] + \mathbf{E} [\|\Delta_{i_k}^k - \Delta^k\|^2] \right) + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&\stackrel{(13),(12)}{=} \frac{(1-p)(1+\omega)}{r} \mathbf{E} [\|\Delta_{i_k}^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2].
\end{aligned}$$

Using L -smoothness (2) of f_i together with the tower property (13) we get

$$\begin{aligned}
\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] &\leq \frac{(1-p)(1+\omega)}{nr} \sum_{i=1}^n L_i^2 \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\
&= \frac{(1-p)(1+\omega)L^2}{r} \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p)\mathbf{E} [\|g^k - \nabla f(x^k)\|^2]. \quad (58)
\end{aligned}$$

Next, we introduce new notation: $\Phi_k = f(x^k) - f_* + \frac{\gamma}{2p} \|g^k - \nabla f(x^k)\|^2$. Using this and inequalities (57) and (58) we establish the following inequality:

$$\begin{aligned}
\mathbf{E} [\Phi_{k+1}] &\leq \mathbf{E} \left[f(x^k) - f_* - \frac{\gamma}{2} \|\nabla f(x^k)\|^2 - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\
&\quad + \frac{\gamma}{2p} \mathbf{E} \left[\frac{(1-p)(1+\omega)L^2}{r} \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\
&= \mathbf{E} [\Phi_k] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2] + \left(\frac{\gamma(1-p)(1+\omega)L^2}{2pn} - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\
&\stackrel{(53)}{\leq} \mathbf{E} [\Phi_k] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2], \quad (59)
\end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)(1+\omega)L^2}{2pn} - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ following from (53). Summing up inequalities (33) for $k = 0, 1, \dots, K-1$ and rearranging the terms we derive

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbf{E} [\|\nabla f(x^k)\|^2] \leq \frac{2}{\gamma K} \sum_{k=0}^{K-1} (\mathbf{E} [\Phi_k] - \mathbf{E} [\Phi_{k+1}]) = \frac{2(\mathbf{E} [\Phi_0] - \mathbf{E} [\Phi_K])}{\gamma K} = \frac{2\Delta_0}{\gamma K},$$

since $g^0 = \nabla f(x^0)$ and $\Phi_{k+1} \geq 0$. Finally, using the tower property (13) and the definition of $\hat{x}^K$ we obtain (54) that implies (55) and (56). $\square$

Corollary D.1 (Corollary 4.1). *Let the assumptions of Theorem 4.1 hold and $p = \frac{\zeta_Q r}{dn}$, where $r \leq n$ and ζ_Q is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_Q r} - 1 \right)} \right)},$$

then PP-MARINA requires

$$K = \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_{\mathcal{Q}r}} - 1 \right)} \right) \right)$$

iterations/communication rounds to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost is

$$\mathcal{O} \left(dn + \frac{\Delta_0 L}{\varepsilon^2} \left(\zeta_{\mathcal{Q}r} + \sqrt{(1+\omega)\zeta_{\mathcal{Q}}(dn - \zeta_{\mathcal{Q}r})} \right) \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof of Corollary 4.1. The choice of $p = \frac{\zeta_{\mathcal{Q}r}}{dn}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \frac{dn}{\zeta_{\mathcal{Q}r}} - 1, \\ pdn + (1-p)\zeta_{\mathcal{Q}r} &\leq \zeta_{\mathcal{Q}r} + \left(1 - \frac{\zeta_{\mathcal{Q}r}}{dn}\right) \cdot \zeta_{\mathcal{Q}r} \leq 2\zeta_{\mathcal{Q}r}. \end{aligned}$$

Plugging these relations in (53), (55), and (56) we get that if

$$\gamma \leq \frac{1}{L \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_{\mathcal{Q}r}} - 1 \right)} \right)},$$

then PP-MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right) \\ &= \mathcal{O} \left(\frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_{\mathcal{Q}r}} - 1 \right)} \right) \right) \end{aligned}$$

iterations/communication rounds in order to achieve $\mathbf{E}[\|\nabla f(\hat{x}^K)\|^2] \leq \varepsilon^2$, and the expected total communication cost is

$$\begin{aligned} dn + K(pd n + (1-p)\zeta_{\mathcal{Q}r}) &= \mathcal{O} \left(dn + \frac{\Delta_0 L}{\varepsilon^2} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) (pd n + (1-p)\zeta_{\mathcal{Q}r}) \right) \\ &= \mathcal{O} \left(dn + \frac{\Delta_0 L}{\varepsilon^2} \left(\zeta_{\mathcal{Q}r} + \sqrt{(1+\omega)\zeta_{\mathcal{Q}}(dn - \zeta_{\mathcal{Q}r})} \right) \right) \end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$

D.2 Convergence Results Under Polyak-Łojasiewicz condition

In this section, we provide an analysis of PP-MARINA under Polyak-Łojasiewicz condition.

Theorem D.2. *Let Assumptions 1.1, 1.2 and 2.1 be satisfied and*

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2(1-p)(1+\omega)}{pr}} \right)}, \frac{p}{2\mu} \right\}, \quad (60)$$

where $L^2 = \frac{1}{n} \sum_{i=1}^n L_i^2$. Then after K iterations of PP-MARINA we have

$$\mathbf{E} [f(x^K) - f(x^*)] \leq (1 - \gamma\mu)^K \Delta_0, \quad (61)$$

where $\Delta_0 = f(x^0) - f(x^*)$. That is, after

$$K = \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \quad (62)$$

iterations PP-MARINA produces such a point x^K that $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$. Moreover, under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server we have that the expected total communication cost (for all workers) equals

$$dn + K(pdn + (1-p)\zeta_{\mathcal{Q}}r) = \mathcal{O} \left(dn + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right\} (pdn + (1-p)\zeta_{\mathcal{Q}}r) \log \frac{\Delta_0}{\varepsilon} \right), \quad (63)$$

where $\zeta_{\mathcal{Q}}$ is the expected density of the quantization (see Def. 1.1).

Proof. The proof is very similar to the proof of Theorem 4.1. From Lemma A.1 and PL condition we have

$$\begin{aligned} \mathbf{E}[f(x^{k+1}) - f(x^*)] &\leq \mathbf{E}[f(x^k) - f(x^*)] - \frac{\gamma}{2} \mathbf{E} [\|\nabla f(x^k)\|^2] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\quad + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2] \\ &\stackrel{(4)}{\leq} (1 - \gamma\mu) \mathbf{E} [f(x^k) - f(x^*)] - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] + \frac{\gamma}{2} \mathbf{E} [\|g^k - \nabla f(x^k)\|^2]. \end{aligned}$$

Using the same arguments as in the proof of (58) we obtain

$$\mathbf{E} [\|g^{k+1} - \nabla f(x^{k+1})\|^2] \leq \frac{(1-p)(1+\omega)L^2}{r} \mathbf{E} [\|x^{k+1} - x^k\|^2] + (1-p) \mathbf{E} [\|g^k - \nabla f(x^k)\|^2].$$

Putting all together we derive that the sequence $\Phi_k = f(x^k) - f(x^*) + \frac{\gamma}{p} \|g^k - \nabla f(x^k)\|^2$ satisfies

$$\begin{aligned} \mathbf{E} [\Phi_{k+1}] &\leq \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) - \left(\frac{1}{2\gamma} - \frac{L}{2} \right) \|x^{k+1} - x^k\|^2 + \frac{\gamma}{2} \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \frac{\gamma}{p} \mathbf{E} \left[\frac{(1-p)(1+\omega)L^2}{r} \|x^{k+1} - x^k\|^2 + (1-p) \|g^k - \nabla f(x^k)\|^2 \right] \\ &= \mathbf{E} \left[(1 - \gamma\mu)(f(x^k) - f(x^*)) + \left(\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \right) \|g^k - \nabla f(x^k)\|^2 \right] \\ &\quad + \left(\frac{\gamma(1-p)(1+\omega)L^2}{pr} - \frac{1}{2\gamma} + \frac{L}{2} \right) \mathbf{E} [\|x^{k+1} - x^k\|^2] \\ &\stackrel{(60)}{\leq} (1 - \gamma\mu) \mathbf{E} [\Phi_k], \end{aligned}$$

where in the last inequality we use $\frac{\gamma(1-p)(1+\omega)L^2}{pr} - \frac{1}{2\gamma} + \frac{L}{2} \leq 0$ and $\frac{\gamma}{2} + \frac{\gamma}{p}(1-p) \leq (1 - \gamma\mu)\frac{\gamma}{p}$ following from (60). Unrolling the recurrence and using $g^0 = \nabla f(x^0)$ we obtain

$$\mathbf{E} [f(x^K) - f(x^*)] \leq \mathbf{E} [\Phi_K] \leq (1 - \gamma\mu)^K \Phi_0 = (1 - \gamma\mu)^K (f(x^0) - f(x^*))$$

that implies (62) and (63). $\square$

Corollary D.2. *Let the assumptions of Theorem D.2 hold and $p = \frac{\zeta_Q r}{dn}$, where $r \leq n$ and ζ_Q is the expected density of the quantization (see Def. 1.1). If*

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2(1+\omega)}{r} \left(\frac{dn}{\zeta_Q r} - 1 \right)} \right)}, \frac{p}{2\mu} \right\},$$

then PP-MARINA requires

$$K = \mathcal{O} \left(\max \left\{ \frac{dn}{\zeta_Q r} \frac{L}{\mu} \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_Q r} - 1 \right)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

iterations/communication rounds to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost is

$$\mathcal{O} \left(dn + \max \left\{ dn, \frac{L}{\mu} \left(\zeta_Q r + \sqrt{(1+\omega)\zeta_Q (dn - \zeta_Q r)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right)$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server.

Proof. The choice of $p = \frac{\zeta_Q r}{dn}$ implies

$$\begin{aligned} \frac{1-p}{p} &= \frac{dn}{\zeta_Q r} - 1, \\ pdn + (1-p)\zeta_Q r &\leq \zeta_Q r + \left(1 - \frac{\zeta_Q r}{dn} \right) \cdot \zeta_Q r \leq 2\zeta_Q r. \end{aligned}$$

Plugging these relations in (60), (62), and (63) we get that if

$$\gamma \leq \min \left\{ \frac{1}{L \left(1 + \sqrt{\frac{2(1+\omega)}{r} \left(\frac{dn}{\zeta_Q r} - 1 \right)} \right)}, \frac{p}{2\mu} \right\},$$

then PP-MARINA requires

$$\begin{aligned} K &= \mathcal{O} \left(\max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(\max \left\{ \frac{dn}{\zeta_Q r} \frac{L}{\mu} \left(1 + \sqrt{\frac{1+\omega}{r} \left(\frac{dn}{\zeta_Q r} - 1 \right)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

iterations/communication rounds in order to achieve $\mathbf{E}[f(x^K) - f(x^*)] \leq \varepsilon$, and the expected total communication cost is

$$\begin{aligned} dn + K(pd n + (1-p)\zeta_Q r) &= \mathcal{O} \left(dn + \max \left\{ \frac{1}{p}, \frac{L}{\mu} \left(1 + \sqrt{\frac{(1-p)(1+\omega)}{pr}} \right) \right\} (pd n + (1-p)\zeta_Q r) \log \frac{\Delta_0}{\varepsilon} \right) \\ &= \mathcal{O} \left(dn + \max \left\{ dn, \frac{L}{\mu} \left(\zeta_Q r + \sqrt{(1+\omega)\zeta_Q (dn - \zeta_Q r)} \right) \right\} \log \frac{\Delta_0}{\varepsilon} \right) \end{aligned}$$

under assumption that the communication cost is proportional to the number of non-zero components of transmitted vectors from workers to the server. $\square$