



Subgradient Methods for Sharp Weakly Convex Functions

Damek Davis¹ · Dmitriy Drusvyatskiy² · Kellie J. MacPhee² · Courtney Paquette³

Received: 8 March 2018 / Accepted: 4 August 2018 / Published online: 10 September 2018
© Springer Science+Business Media, LLC, part of Springer Nature 2018

Abstract

Subgradient methods converge linearly on a convex function that grows sharply away from its solution set. In this work, we show that the same is true for sharp functions that are only weakly convex, provided that the subgradient methods are initialized within a fixed tube around the solution set. A variety of statistical and signal processing tasks come equipped with good initialization and provably lead to formulations that are both weakly convex and sharp. Therefore, in such settings, subgradient methods can serve as inexpensive local search procedures. We illustrate the proposed techniques on phase retrieval and covariance estimation problems.

Keywords Subgradient · Weakly convex · Sharp · Error bound · Linear rate

Mathematics Subject Classification 65K05 · 65K10

Communicated by Evgeni A. Nurminski.

✉ Dmitriy Drusvyatskiy
ddrusv@uw.edu
<http://www.math.washington.edu/~ddrusv>

Damek Davis
<http://people.orie.cornell.edu/dsd95/>

Kellie J. MacPhee
<http://sites.math.washington.edu/~kmacphee>

Courtney Paquette
cop318@lehigh.edu

¹ School of Operations Research and Information Engineering, Cornell University, Ithaca, NY 14850, USA

² Department of Mathematics, University of Washington, Seattle, WA 98195, USA

³ Industrial and Systems Engineering Department, Lehigh University, Bethlehem, PA 18015, USA

1 Introduction

Typical methods for statistics and signal processing tasks follow the two-step strategy: (i) find a moderately accurate solution at a low sample complexity cost (e.g., using spectral initialization), and (ii) refine the approximate solution by an iterative “local search algorithm” that converges rapidly under natural statistical assumptions. For smooth problem formulations, the term “local search” almost universally refers to gradient descent or a close variant thereof; see, for example, [1–8]. For nonsmooth and nonconvex problems, the meaning of local search is much less clear. In this work, we ask the following question.

Is there a generic gradient-based local search procedure for nonsmooth and non-convex problems, which converges linearly under standard regularity conditions?

We answer this question positively for the class of weakly convex functions, i.e., those that become convex after a sufficiently large quadratic perturbation. We show that standard subgradient methods, originally designed for sharp convex problems, locally converge linearly when applied to sharp functions that are only weakly convex. We focus on three stepsize rules: Polyak stepsize [9,10], geometrically decaying stepsize [11,12], and constant stepsize [13–15]. As a proof of concept, we illustrate the resulting algorithms on phase retrieval and covariance estimation problems.

2 Discussion

Our approach is rooted in subgradient methods for convex optimization. To motivate the discussion, consider the constrained optimization problem

$$\min_{x \in \mathcal{X}} g(x), \quad (1)$$

where g is an L -Lipschitz convex function on $\mathbb{R}^d$ and $\mathcal{X}$ is a closed and convex set. Given a current iterate x_k , subgradient methods proceed as follows:

$$\left\{ \begin{array}{l} \text{Choose any } \zeta_k \in \partial g(x_k) \\ \text{Set } x_{k+1} = \text{proj}_{\mathcal{X}} \left(x_k - \alpha_k \cdot \frac{\zeta_k}{\|\zeta_k\|} \right) \end{array} \right\}.$$

Here, the symbol $\text{proj}_{\mathcal{X}}(y)$ denotes the nearest point of $\mathcal{X}$ to y and $\{\alpha_k\}$ is a specified stepsize sequence. The choice of the sequence $\{\alpha_k\}$ determines the behavior of the scheme and is the main distinguishing feature among subgradient methods. In this work, we will only be interested in subgradient methods that are linearly convergent. As usual, linear rates of convergence of iterative methods require some regularity conditions to hold. Here, the appropriate regularity condition is *sharpness* [16,17] (or equivalently a global error bound): There exists a real $\mu > 0$ satisfying

$$g(z) - \min_{x \in \mathcal{X}} g(x) \geq \mu \cdot \text{dist}(z; \mathcal{X}^*) \quad \text{for all } z \in \mathcal{X},$$

where $\mathcal{X}^*$ denotes the set of minimizers of (1), and $\text{dist}(z; \mathcal{X}^*)$ denotes the distance of z to $\mathcal{X}^*$. Assuming sharpness holds, subgradient methods with a judicious choice of $\{\alpha_k\}$ produce iterates that converge to $\mathcal{X}^*$ at the linear rate $\sqrt{1 - (\mu/L)^2}$. Results of this type date back to the 1960s and 1970s [9–12, 18], while some more recent approaches have appeared in [13–15].

Various contemporary problems lead to formulations that are indeed sharp, but are only *weakly convex* and *locally Lipschitz*. Recall that a function g is ρ -weakly convex [19] if the perturbed function $x \mapsto g(x) + \frac{\rho}{2} \|\cdot\|^2$ is convex for some $\rho \geq 0$.¹ Note that weakly convex functions need not be smooth nor convex. A quick computation (Lemma 3.1) shows that if g is μ -sharp and ρ -weakly convex, then there is a tube around the solution set $\mathcal{X}^*$ that contains no extraneous stationary points:

$$\mathcal{T} := \left\{ x \in \mathcal{X} : \text{dist}(x; \mathcal{X}^*) < \frac{2\mu}{\rho} \right\}.$$

In this work, we show that the standard linearly convergent subgradient methods, originally designed for convex problems, apply in this much greater generality, provided they are initialized within a slight contraction of the tube $\mathcal{T}$. The methods exhibit essentially the same linear rate of convergence as in the convex case, while the weak convexity constant ρ only determines the validity of the initialization. We focus on three stepsize rules: Polyak stepsize [9, 10], geometrically decaying stepsize [11, 12], and constant stepsize [13–15]. As a proof of concept, we illustrate the resulting algorithms on phase retrieval and covariance estimation problems.

Our current work sits within the broader scope of analyzing subgradient and proximal methods for weakly convex problems [19–26]; see also the recent survey [27]. In particular, the paper [21] proves a global sublinear rate of convergence, in terms of a natural stationarity measure, of a (stochastic) subgradient method on any weakly convex function. In contrast, here we are interested in subgradient methods that are locally linearly convergent under the additional sharpness assumption. The arguments we present are all quick modification of the proofs already available in the convex setting. Nonetheless, we believe that the drawn conclusions are interesting and powerful, opening the door to generic local search procedures for nonsmooth and nonconvex problems.

3 Notation

Throughout, we consider the Euclidean space $\mathbb{R}^d$, equipped with an inner product $\langle \cdot, \cdot \rangle$ and the induced norm $\|x\| := \sqrt{\langle x, x \rangle}$. The *distance* and the *projection* of any point

¹ To the best of our knowledge, the class of weakly convex functions was introduced in 1972 in a local publication of the Institute of Cybernetics in Kiev and then used in [19]. These are the functions that are lower bounded by a linear function up to a little-o term that is uniform in the base point. The term ρ -weak convexity, as we use it here, corresponds to the setting when the error term is a quadratic. A number of other names are used in the literature instead of ρ -weak convexity, such as semi-convex, prox-regular, and lower- $\mathcal{C}^2$.

$y \in \mathbb{R}^d$ onto a set $\mathcal{X}$ are defined by

$$\text{dist}(y; \mathcal{X}) := \inf_{x \in \mathcal{X}} \|y - x\| \quad \text{and} \quad \text{proj}_{\mathcal{X}}(y) := \underset{x \in \mathcal{X}}{\text{argmin}} \|y - x\|,$$

respectively. Note that $\text{proj}_{\mathcal{X}}(y)$ is nonempty as long as $\mathcal{X}$ is a closed set. The *indicator function* of a set $\mathcal{X}$, denoted by $\delta_{\mathcal{X}}$, is defined to be zero on $\mathcal{X}$ and $+\infty$ off it.

3.1 Weakly Convex Functions

Our main focus is on those functions that are convex up to an additive quadratic perturbation. Namely, a function $g: \mathbb{R}^d \rightarrow \mathbb{R} \cup \{+\infty\}$ is called ρ -weakly convex (with $\rho \geq 0$), if the assignment $x \mapsto g(x) + \frac{\rho}{2}\|x\|^2$ is a convex function. The algorithms we consider will all use generalized derivative constructions. Variational analytic literature highlights a number of distinct subdifferentials (e.g., [28–31]); for weakly convex functions, all of these constructions coincide. Consider a ρ -weakly convex function g . The *subdifferential* of g at x , denoted $\partial g(x)$, is the set of all vectors $v \in \mathbb{R}^d$ satisfying

$$g(y) \geq g(x) + \langle v, y - x \rangle + o(\|y - x\|) \quad \text{as } y \rightarrow x. \quad (2)$$

Though the condition (2) appears to lack uniformity with respect to the basepoint x , the subgradients of g automatically satisfy the much stronger property:

$$g(y) \geq g(x) + \langle v, y - x \rangle - \frac{\rho}{2}\|y - x\|^2, \quad \forall x, y \in \mathbb{R}^d, \quad v \in \partial g(x). \quad (3)$$

Indeed, this follows quickly by applying the subgradient inequality to the convex function $g + \frac{\rho}{2}\|\cdot\|^2$. Thus, we may use the two conditions, (2) and (3), interchangeably for weakly convex functions. We note in passing that localizing condition (3) leads to so-called prox-regular functions, introduced in [32].

Weakly convex functions are widespread in applications and are typically easy to recognize. One common source is the composite problem class:

$$\min_x F(x) := h(c(x)), \quad (4)$$

where $h: \mathbb{R}^m \rightarrow \mathbb{R}$ is convex and L -Lipschitz, and $c: \mathbb{R}^d \rightarrow \mathbb{R}^m$ is a C^1 -smooth map with β -Lipschitz gradient. An easy argument shows that F is $L\beta$ -weakly convex. This is a worst case estimate. In concrete circumstances, the composite function F may have a much more favorable weak convexity constant ρ . The elements of the subdifferential $\partial F(x)$ are straightforward to compute through the chain rule [28, Theorem 10.6, Corollary 10.9]:

$$\partial F(x) = \nabla c(x)^* \partial h(c(x)) \quad \text{for all } x \in \mathbb{R}^d.$$

For a discussion of some recent uses of weakly convex functions in optimization, see the short survey [27]. Throughout the paper, we will use the following two running examples to illustrate our results.

Example 3.1 (Phase retrieval) Phase retrieval is a common computational problem, with applications in diverse areas such as imaging, X-ray crystallography, and speech processing. For simplicity, we will focus on the version of the problem over the reals. The (real) phase retrieval problem seeks to determine a point x satisfying the magnitude conditions,

$$|\langle a_i, x \rangle|^2 \approx b_i \quad \text{for } i = 1, \dots, m,$$

where $a_i \in \mathbb{R}^d$ and $b_i \in \mathbb{R}$ are given. Note that we can only recover the optimal x up to a universal sign change, since $|\langle a_i, x \rangle| = |\langle a_i, -x \rangle|$. In this work, we will focus on the following optimization formulation of the problem [24,33,34]:

$$\min_x \frac{1}{m} \sum_{i=1}^m |\langle a_i, x \rangle|^2 - b_i|.$$

Clearly, this is an instance of (4). Indeed, under mild statistical assumptions on the way a_i are generated, the formulation is ρ -weakly convex, for some numerical constant ρ independent of d and m [24, Corollary 3.2]. Moreover, under an appropriate model of the noise in the measurements, the problem is sharp [24, Proposition 3]. It is worthwhile to mention that numerous other approaches to phase retrieval exist, based on different problem formulations; see, for example, [3,35–37].

Experiment setup All of the experiments on phase retrieval will be generated according to the following procedure. In the *exact setup*, we generate standard Gaussian measurements $a_i \sim N(0, I_{d \times d})$, for $i = 1, \dots, m$, and generate the target signal $\bar{x} \sim N(0, I_{d \times d})$. We then set $b_i = \langle a_i, \bar{x} \rangle^2$ for each $i = 1, \dots, m$. In the *corrupted setup*, we generate a_i and $\bar{x}$ as in the exact case. We then corrupt a proportion of the measurements with outliers. Namely, we set $b_i = (1 - z_i) \langle a_i, \bar{x} \rangle^2 + z_i |\zeta_i|$, where $z_i \sim \text{Bernoulli}(0.1)$ and $\zeta_i \sim \mathcal{N}(0, 100)$.

Example 3.2 (Covariance matrix estimation) The problem of covariance estimation from quadratic measurements, introduced in [38], is a higher rank variant of phase retrieval. Let $a_1, \dots, a_m \in \mathbb{R}^d$ be measurement vectors. The goal is to recover a low-rank decomposition of a covariance matrix $\bar{X} \bar{X}^T$, with $\bar{X} \in \mathbb{R}^{d \times r}$ for a given $0 \leq r \leq d$, from quadratic measurements

$$b_i \approx a_i^T \bar{X} \bar{X}^T a_i = \text{Tr}(\bar{X} \bar{X}^T a_i a_i^T).$$

Note that we can only recover $\bar{X}$ up to multiplication by an orthogonal matrix. This problem arises in a variety of contexts, such as covariance sketching for data streams and spectrum estimation of stochastic processes. We refer the reader to [38] for details.

In our examples, we will assume m is even and will focus on the potential function²

$$\min_{X \in \mathbb{R}^{d \times r}} \frac{1}{m} \sum_{i=1}^m \left| \left\langle X X^T, a_{2i} a_{2i}^T - a_{2i-1} a_{2i-1}^T \right\rangle - (b_{2i} - b_{2i-1}) \right|. \quad (5)$$

Under exact measurements, i.e., $b_i = a_i^T \bar{X} \bar{X}^T a_i$ and under appropriate statistical assumptions on how a_i are generated, the formulation (5) is ρ -weakly convex for a numerical constant ρ , independent of d or m , and is sharp. Indeed, this is a simple consequence of two results, namely [38, Corollary 1] and [2, Lemma 5.4]. It is possible to show that the objective is also sharp when the measurements are corrupted by gross outliers. This guarantee is beyond the scope of our current work and will appear in a different paper.

Experiment setup All of the experiments on covariance matrix estimation will be generated according to the following procedure. In the *exact setup*, we generate standard Gaussian measurements $a_i \sim N(0, I_{d \times d})$ for $i = 1, \dots, m$, and generate the target matrix $\bar{X} \in \mathbb{R}^{d \times r}$ as a standard Gaussian. We then set $b_i = \|\bar{X}^T a_i\|_F^2$ for each $i = 1, \dots, m$. In the *corrupted setup*, we generate a_i and $\bar{X}$ as in the exact case. We then corrupt a proportion of the measurements with outliers. Namely, we set $b_i = (1 - z_i) \|\bar{X}^T a_i\|_F^2 + z_i |\zeta_i|$, where $\zeta_i \sim \mathcal{N}(0, 100)$ and $z_i \sim \text{Bernoulli}(0.1)$. All plots will show iteration counter k versus the scaled Procrustes distance $\text{dist}(X_k, \mathcal{X}^*) / \|\bar{X}\| = \min_{\Omega^T \Omega = I} \|\Omega X_k - \bar{X}\|_F / \|\bar{X}\|_F$.

3.2 Setting of the Paper

Throughout the manuscript, we make the following assumption.

Assumption A Consider the optimization problem

$$\min_{x \in \mathcal{X}} g(x), \quad (6)$$

satisfying the following properties for some real $\rho \geq 0$ and $\mu > 0$.

1. (Weak convexity) The function $g: \mathbb{R}^d \rightarrow \mathbb{R}$ is ρ -weakly convex, and the set $\mathcal{X} \subset \mathbb{R}^d$ is closed and convex. The set of minimizers $\mathcal{X}^* := \underset{x \in \mathcal{X}}{\operatorname{argmin}} g(x)$ is nonempty.
2. (Sharpness) The inequality

$$g(z) - \min_{x \in \mathcal{X}} g(x) \geq \mu \cdot \text{dist}(z; \mathcal{X}^*) \quad \text{holds for all } z \in \mathcal{X}.$$

² Note this potential function is different from the direct generalization of the one used in phase retrieval, $\min_x \sum_i \left| \|X^T a_i\|^2 - b_i \right|$. The reason for the more exotic formulation is that the differences $a_{2i} a_{2i}^T - a_{2i-1} a_{2i-1}^T$ in (5) allow the authors of [38] to prove sharpness of the objective function by appealing to an ℓ_2/ℓ_1 restricted isometry property of the resulting linear map. See [38, Section III.B] for details. It is not clear whether the naive objective function is also sharp with high probability under reasonable statistical assumptions.

We will say that a point $\bar{x} \in \mathcal{X}$ is *stationary* for the target problem (6) if

$$g(x) - g(\bar{x}) \geq o(\|x - \bar{x}\|) \quad \text{as } x \rightarrow \bar{x} \text{ in } \mathcal{X}.$$

That is, $\bar{x}$ is *stationary* precisely when the zero vector is a subgradient of $g + \delta_{\mathcal{X}}$ at $\bar{x}$.

Shortly, we will discuss subgradient methods that converge linearly to $\mathcal{X}^*$ under appropriate initialization. As a first step, therefore, we must identify a neighborhood of $\mathcal{X}^*$ that is devoid of extraneous stationary points of (6). This is the content of the following lemma.

Lemma 3.1 (Neighborhood with no stationary points) *The problem (6) has no stationary points x satisfying*

$$0 < \text{dist}(x; \mathcal{X}^*) < \frac{2\mu}{\rho}. \quad (7)$$

Proof Fix a stationary point $x \in \mathcal{X} \setminus \mathcal{X}^*$ of (6). Choosing an arbitrary $\bar{x} \in \text{proj}_{\mathcal{X}^*}(x)$, observe

$$\mu \cdot \text{dist}(x; \mathcal{X}^*) \leq g(x) - g(\bar{x}) \leq \frac{\rho}{2} \|x - \bar{x}\|^2 = \frac{\rho}{2} \cdot \text{dist}^2(x; \mathcal{X}^*).$$

Dividing through by $\text{dist}(x; \mathcal{X}^*)$, the result follows. $\square$

In light of Lemma 3.1, for any $\gamma > 0$ we define the following tube

$$\mathcal{T}_\gamma := \left\{ x \in \mathcal{X} : \text{dist}(x; \mathcal{X}^*) < \gamma \cdot \frac{\mu}{\rho} \right\}$$

and the constant

$$L := \sup \{ \|\zeta\| : \zeta \in \partial g(x), x \in \mathcal{T}_1 \}. \quad (8)$$

Lemma 3.1 guarantees that the tubes $\mathcal{T}_\gamma$ contain no extraneous stationary points of the problem for any $0 < \gamma \leq 2$. Moreover, observe that μ and L play reciprocal roles; consequently, the ratio $\tau := \mu/L$ should serve as a measure of conditioning. The following lemma verifies the inclusion $\tau \in [0, 1]$.

Lemma 3.2 (Condition measure) *The inclusion $\tau \in [0, 1]$ holds.*

Proof Consider an arbitrary point $x \in \mathcal{T}_1 \setminus \mathcal{X}^*$ and choose any $\bar{x} \in \text{proj}_{\mathcal{X}^*}(x)$. By Lebourg's mean value theorem [39, Theorem 2.3.7], there exists a point z in the open segment $(x, \bar{x})$ and a vector $\zeta \in \partial g(z)$ satisfying

$$g(x) - g(\bar{x}) = \langle \zeta, x - \bar{x} \rangle. \quad (9)$$

Trivially z lies in $\mathcal{T}_1$, and therefore $\|\zeta\| \leq L$. Using this estimate and sharpness in (9) yields the guarantee

$$\mu \cdot \text{dist}(x; \mathcal{X}^*) \leq g(x) - g(\bar{x}) \leq \|\zeta\| \cdot \|x - \bar{x}\| \leq L \cdot \text{dist}(x; \mathcal{X}^*).$$

The result follows. $\square$

To summarize, we will use the following symbols to describe the parameters of the problem class (6): ρ is the weak convexity constant of g , μ is the sharpness constant of g , L is the maximal subgradient norm at points in the tube $\mathcal{T}_1$, and τ is the condition measure $\tau = \mu/L \in [0, 1]$.

4 Polyak Subgradient Method

In this section, we consider the Polyak subgradient method for problem (6). A preliminary version of this material applied to the phase retrieval problem appeared in [34]; we present the arguments here for the sake of completeness.

The Polyak subgradient method is summarized in Algorithm 1. The method requires knowing the optimal value $\min_{x \in \mathcal{X}} g(x)$. In a number of circumstances, this indeed is reasonable (e.g., exact penalty approach for solving nonlinear equations). The latter sections explore subgradient methods that do not require a known optimal value.

Algorithm 1: Polyak Subgradient Method

Data: Initial point $x_0 \in \mathbb{R}^d$
Step k : ($k \geq 0$)
 Choose $\zeta_k \in \partial g(x_k)$
if $\zeta_k \neq 0$ **then**
 Set $x_{k+1} = \text{proj}_{\mathcal{X}} \left(x_k - \frac{g(x_k) - \min_{x \in \mathcal{X}} g(x)}{\|\zeta_k\|^2} \zeta_k \right)$
else
 Set $x_{k+1} = x_k$
end

The following theorem shows that Algorithm 1, originally proposed for convex problems enjoys the same linear convergence guarantees for functions that are only weakly convex, provided it is initialized within a certain tube of the optimal solution set.

Theorem 4.1 (Linear rate) *Fix a real $0 < \gamma < 1$. Then, Algorithm 1 initialized at any point $x_0 \in \mathcal{T}_\gamma$ produces iterates that converge Q -linearly to $\mathcal{X}^*$. That is, for each index $k \geq 0$ we have*

$$\text{dist}^2(x_{k+1}; \mathcal{X}^*) \leq \left(1 - (1 - \gamma)\tau^2\right) \text{dist}^2(x_k; \mathcal{X}^*). \quad (10)$$

Proof We proceed by induction. Suppose that the theorem holds up to iteration $k - 1$. We will prove inequality (10). To this end, observe that the inductive hypothesis implies, $\text{dist}(x_k; \mathcal{X}^*) \leq \text{dist}(x_0; \mathcal{X}^*)$, and therefore x_k lies in $\mathcal{T}_\gamma$. Observe that in the case $\zeta_k = 0$, we have $x_{k+1} = x_k$ by the definition of the subgradient method and $x_k \in \mathcal{X}^*$ by Lemma 3.1; hence, (10) holds trivially. Therefore, we may suppose $\zeta_k \neq 0$ for the rest of the proof. Choose now a point $\bar{x} \in \text{proj}_{\mathcal{X}^*}(x_k)$. Taking into account that $\text{proj}_{\mathcal{X}}$ is nonexpansive, we successively deduce

$$\begin{aligned}
\|x_{k+1} - \bar{x}\|^2 &\leq \left\| (x_k - \bar{x}) - \frac{g(x_k) - g(\bar{x})}{\|\zeta_k\|^2} \zeta_k \right\|^2 \\
&= \|x_k - \bar{x}\|^2 + \frac{2(g(x_k) - g(\bar{x}))}{\|\zeta_k\|^2} \cdot \langle \zeta_k, \bar{x} - x_k \rangle + \frac{(g(x_k) - g(\bar{x}))^2}{\|\zeta_k\|^2} \\
&\leq \|x_k - \bar{x}\|^2 + \frac{2(g(x_k) - g(\bar{x}))}{\|\zeta_k\|^2} \left(g(\bar{x}) - g(x_k) + \frac{\rho}{2} \|x_k - \bar{x}\|^2 \right) \\
&\quad + \frac{(g(x_k) - g(\bar{x}))^2}{\|\zeta_k\|^2} \\
&= \|x_k - \bar{x}\|^2 + \frac{(g(x_k) - g(\bar{x}))}{\|\zeta_k\|^2} \left(\rho \|x_k - \bar{x}\|^2 - (g(x_k) - g(\bar{x})) \right) \\
&\leq \|x_k - \bar{x}\|^2 + \frac{(g(x_k) - g(\bar{x}))}{\|\zeta_k\|^2} \left(\rho \|x_k - \bar{x}\|^2 - \mu \|x_k - \bar{x}\| \right) \\
&= \|x_k - \bar{x}\|^2 + \frac{\rho(g(x_k) - g(\bar{x}))}{\|\zeta_k\|^2} \left(\|x_k - \bar{x}\| - \frac{\mu}{\rho} \right) \|x_k - \bar{x}\|.
\end{aligned}$$

Combining the inclusion $x_k \in \mathcal{T}_\gamma$ with sharpness, we therefore deduce

$$\text{dist}^2(x_{k+1}; \mathcal{X}^*) \leq \|x_{k+1} - \bar{x}\|^2 \leq \left(1 - \frac{(1 - \gamma)\mu^2}{\|\zeta_k\|^2} \right) \|x_k - \bar{x}\|^2.$$

The result follows. $\square$

As a numerical illustration, we apply the Polyak subgradient method (Fig. 1) to our two running examples, phase retrieval and covariance matrix estimation. Notice that a linear rate of convergence is observed in all experiments except for two, with the rate improving monotonically with an increasing number of measurements m . In the two exceptional experiments, the number of measurements m is too small to guarantee that the initial point x_0 is within the basin of attraction, and the subgradient methods stagnate.

5 Subgradient Method with Constant Stepsize

Recall that the Polyak subgradient method (Algorithm 1) crucially relies on knowing the minimal value of the optimization problem (6). Henceforth, all the subgradient methods we consider are agnostic to this value. That being said, they will require some estimates on the problem parameters (μ, ρ, L) . We begin by analyzing a subgradient method with a constant stepsize (Algorithm 2). Constant stepsize schemes are often methods of choice in practice. We will show that when properly initialized, the subgradient method with constant stepsize generates iterates x_k such that $\text{dist}(x_k; \mathcal{X}^*)$ converges linearly up to a certain threshold.

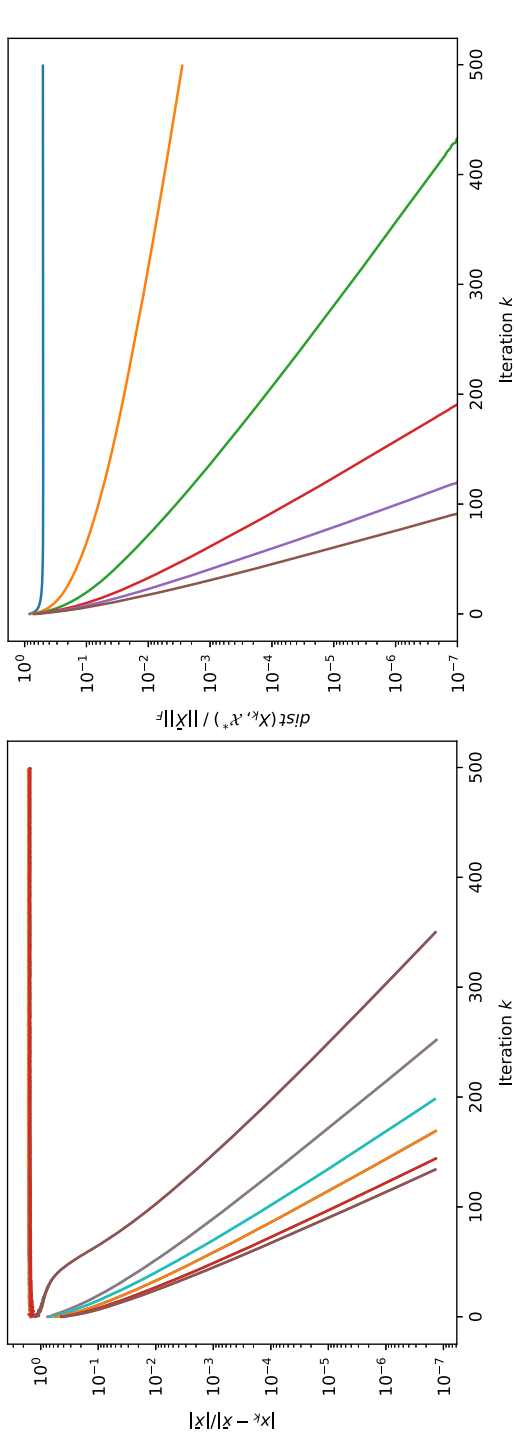


Fig. 1 Polyak subgradient method. (Left) Phase retrieval with the exact setup: $d = 5000$ and $m \in \{11,000, 12,225, 13,500, 14,750, 16,000, 17,250, 18,500\}$. (Right) Covariance matrix estimation with the exact setup: $d = 1000$, $r = 3$, $m \in \{5000, 8000, 11,000, 14,000, 17,000, 20,000\}$. In both experiments, convergence rates uniformly improve with increasing m

Algorithm 2: Subgradient method with constant stepsize**Data:** Initial point $x_0 \in \mathbb{R}^d$ and stepsize $\alpha > 0$ **Step k :** ($k \geq 0$)Choose $\zeta_k \in \partial g(x_k)$ **if** $\zeta_k \neq 0$ **then**Set $x_{k+1} = \text{proj}_{\mathcal{X}} \left(x_k - \alpha \cdot \frac{\zeta_k}{\|\zeta_k\|} \right)$ **else**Set $x_{k+1} = x_k$ **end**

The analysis we present fundamentally relies on the following estimate, often used in the analysis of subgradient methods. To simplify notation, for any point $x \in \mathbb{R}^d$, we set

$$E(x) := \text{dist}^2(x; \mathcal{X}^*).$$

Whenever x has an index k as a subscript, we will set $E_k := E(x_k)$. The following lemma will feature in both the constant and geometrically decaying stepsize schemes.

Lemma 5.1 (Basic recurrence) *Consider a point $x \in \mathcal{T}_1$ and a nonzero subgradient $\zeta \in \partial g(x)$, and define $x^+ := \text{proj}_{\mathcal{X}} \left(x - \alpha \frac{\zeta}{\|\zeta\|} \right)$ for some $\alpha > 0$. Then, the estimate holds:*

$$E(x^+) \leq \left(1 + \frac{\rho\alpha}{L} \right) E(x) - 2\alpha\tau\sqrt{E(x)} + \alpha^2. \quad (11)$$

Proof Choose an arbitrary point $\bar{x} \in \text{proj}_{\mathcal{X}^*}(x)$. Observe

$$\begin{aligned} \|x^+ - \bar{x}\|^2 &\leq \left\| (x - \bar{x}) - \alpha \frac{\zeta}{\|\zeta\|} \right\|^2 = \|x - \bar{x}\|^2 + \frac{2\alpha}{\|\zeta\|} \cdot \langle \zeta, \bar{x} - x \rangle + \alpha^2 \\ &\leq \|x - \bar{x}\|^2 + \frac{2\alpha}{\|\zeta\|} \cdot \left(g(\bar{x}) - g(x) + \frac{\rho}{2} \|x - \bar{x}\|^2 \right) + \alpha^2 \\ &\leq \left(1 + \frac{\alpha\rho}{\|\zeta\|} \right) \|x - \bar{x}\|^2 - \frac{2\alpha\mu}{\|\zeta\|} \cdot \|x - \bar{x}\| + \alpha^2. \end{aligned}$$

Thus, the inequality holds:

$$E(x^+) \leq \left(1 + \frac{\alpha\rho}{\|\zeta\|} \right) E(x) - \frac{2\alpha\mu}{\|\zeta\|} \cdot \sqrt{E(x)} + \alpha^2,$$

and consequently taking into account $\alpha/\|\zeta\| \geq \alpha/L$ we have

$$E(x^+) \leq \sup_{t \geq \alpha/L} \left\{ (1 + \rho t) E(x) - 2\mu t \cdot \sqrt{E(x)} + \alpha^2 \right\}.$$

Notice, the function inside the supremum is linear in t with slope given by $s := \rho E(x) - 2\mu\sqrt{E(x)}$. The inclusion $x \in \mathcal{T}_1$ directly implies $s \leq 0$. Therefore, the supremum on the right-hand side is attained at $t = \frac{\alpha}{L}$, yielding the claimed estimate (11). $\square$

In light of Lemma 5.1, we can now prove that the quantities $E_k = E(x_k)$ converge linearly below a certain fixed threshold. The proof is a modification of that in [15, Section 4].

Lemma 5.2 (Contraction inequality) *Fix a constant $0 < \alpha < \frac{\tau\mu}{\rho}$, and let $\{x_k\}_{k \geq 0}$ be the iterates generated by Algorithm 2. Define the quantity*

$$E^* := \left(\frac{\alpha L}{\mu + \sqrt{\mu^2 - \alpha \rho L}} \right)^2. \quad (12)$$

Then, whenever an iterate x_k lies in $\mathcal{T}_1$, the estimate holds:

$$E_{k+1} - E^* \leq q_k (E_k - E^*),$$

where $q_k := 1 + \frac{\alpha}{L} \left(\rho - \frac{2\mu}{\sqrt{E_k} + \sqrt{E^}} \right)$ satisfies $q_k < 1$.*

Proof Looking back at estimate (11), consider the following equation in the variable e :

$$e = \left(1 + \frac{\alpha \rho}{L} \right) e - 2\alpha \tau \cdot \sqrt{e} + \alpha^2. \quad (13)$$

An easy computation shows that the minimal positive solution to (13) is exactly E^* , defined in (12). Note that E^* is well defined by the inequality $\alpha \leq \tau \cdot \frac{\mu}{\rho}$.

Subtracting (13) from (11) yields the estimate

$$\begin{aligned} E_{k+1} - E^* &\leq \left(1 + \frac{\alpha \rho}{L} \right) (E_k - E^*) - 2\alpha \tau (\sqrt{E_k} - \sqrt{E^*}) \\ &= \left(1 + \frac{\alpha}{L} \left(\rho - \frac{2\mu}{\sqrt{E_k} + \sqrt{E^*}} \right) \right) (E_k - E^*). \end{aligned}$$

Finally, noting

$$\sqrt{E^*} = \frac{\alpha L}{\mu + \sqrt{\mu^2 - \alpha \rho L}} = \frac{\mu - \sqrt{\mu^2 - \alpha \rho L}}{\rho} \leq \frac{\mu}{\rho},$$

we obtain

$$\rho - \frac{2\mu}{\sqrt{E_k} + \sqrt{E^*}} < \rho - \frac{2\mu}{2\mu/\rho} = 0.$$

This completes the proof of the lemma. $\square$

Iterating Lemma 5.2, we see that the quantities E_k decrease to a value at most E^* at a linear rate. Figure 2 illustrates this behavior on our two running examples. It is also clear from the figure that the linear rate of convergence improves as E_k tends to E^* . An explanation is immediate from the expression for q_k in Lemma 5.2. Indeed, as E_k decreases, so do the contraction factors q_k , and for $E_k \approx E^*$, we have $q_k \approx (1 - \tau^2) + \frac{\alpha \rho}{L}$. Thus, as the stepsize tends to zero, the limiting linear rate coincides with an ideal rate of $\approx 1 - \tau^2$.

Even though E_k could be smaller than E^* and q could be negative, it is apparent from Fig. 2 that even after E_k becomes smaller than E^* , all the following values E_k stay close to E^* . This is the content of the following theorem. The convex version of this theorem appears in [15, Theorem 1].

Theorem 5.1 (Convergence of fixed stepsize subgradient method) *Fix constants $0 < \gamma < 1$ and α satisfying*

$$0 < \alpha < \frac{\gamma\tau}{\sqrt{1+2\tau^2}} \cdot \frac{\mu}{\rho}. \quad (14)$$

Let x_k be the iterates generated by Algorithm 2 with stepsize α and initial point $x_0 \in T_\gamma$. Define the constants

$$E^* := \left(\frac{\alpha L}{\mu + \sqrt{\mu^2 - \alpha\rho L}} \right)^2 \quad \text{and} \quad D := \sqrt{\max\{E_0, 2\alpha^2 + E^*\}}.$$

Then, for each index k , the estimates hold:

$$\sqrt{E_k} \leq D < \frac{\gamma\mu}{\rho} \quad \text{and} \quad E_k - E^* \leq \max \left\{ q^k (E_0 - E^*), 2\alpha^2 \right\},$$

where the coefficient $q := 1 + \frac{\alpha}{L} \left(\rho - \frac{\mu}{D} \right)$ satisfies $0 < q < 1$.

Proof We first verify the claims that are independent of the iteration counter. To this end, observe that (14) directly implies $\alpha < \tau \cdot \frac{\mu}{\rho}$, and therefore E^* is well defined. Next, we show $D < \frac{\gamma\mu}{\rho}$. Indeed, noting $\sqrt{E^*} < \alpha\tau^{-1}$ and using (14), we deduce

$$D^2 = \max\{E_0, 2\alpha^2 + E^*\} \leq \max\{E_0, \alpha^2(2 + \tau^{-2})\} < \left(\frac{\gamma\mu}{\rho} \right)^2.$$

Next, we show the inequalities $0 < q < 1$. To this end, observe

$$-1 \leq -\frac{\tau\alpha}{D} < \frac{\alpha}{L} \left(\rho - \frac{\mu}{D} \right) \leq (1 - \gamma^{-1}) \cdot \frac{\rho\alpha}{L} < 0,$$

where the first inequality follows from the inequality, $\alpha \leq D$, and the third follows from the inequality, $D < \frac{\gamma\mu}{\rho}$. Thus, we conclude $0 < q < 1$, as claimed.

We now proceed by induction. Fix an index k and suppose as inductive hypothesis that for each index $i = 0, 1, \dots, k$, the estimates hold:

$$\sqrt{E_i} \leq D \quad \text{and} \quad E_i - E^* \leq \max \left\{ q^i (E_0 - E^*), 2\alpha^2 \right\}.$$

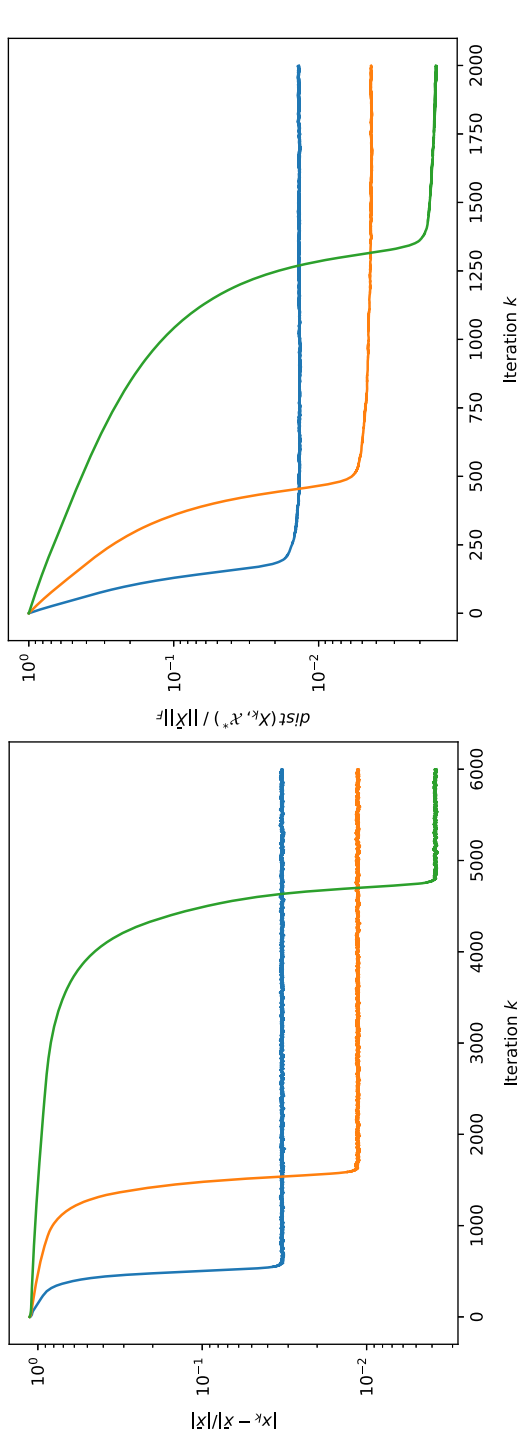


Fig. 2 Constant stepsize subgradient method. (Left) Phase retrieval with the corrupted setup; $d = 1000$, $m = 3$, $r = 10,000$, and $\alpha \in \{1, 1/3, 1/9\}$. (Right) Covariance matrix estimation with the corrupted setup; $d = 1000$, $r = 3$, $m = 10,000$, and $\alpha \in \{1, 1/3, 1/9\}$. The lower curves correspond to smaller step sizes in both experiments

Let us consider two cases. Suppose first $E_k \geq E^*$. Then, by applying Lemma 5.2, we deduce

$$\begin{aligned} E_{k+1} - E^* &\leq \left(1 + \frac{\alpha}{L} \left(\rho - \frac{2\mu}{\sqrt{E_k} + \sqrt{E^*}}\right)\right) (E_k - E^*) \\ &\leq \left(1 + \frac{\alpha}{L} \left(\rho - \frac{\mu}{D}\right)\right) (E_k - E^*) \\ &= q(E_k - E^*). \end{aligned}$$

Suppose now that the second case, $E_k < E^*$, holds. Then, Lemma 5.2 implies

$$\begin{aligned} E_{k+1} &\leq E_k + \frac{\alpha\rho}{L} (E_k - E^*) - 2\alpha\tau(\sqrt{E_k} - \sqrt{E^*}) \\ &\leq \max_{E \in [0, E^*]} \{E + \frac{\alpha\rho}{L} (E - E^*) - 2\alpha\tau(\sqrt{E} - \sqrt{E^*})\} \\ &= \max\{E^*, \frac{\alpha}{L} (2\mu\sqrt{E^*} - \rho E^*)\}. \end{aligned}$$

Subtracting E^* , we conclude

$$E_{k+1} - E^* \leq \max\{0, 2\tau\alpha\sqrt{E^*}\} \leq 2\alpha^2.$$

Thus, in both cases, we have the estimate

$$E_{k+1} - E^* \leq \max\left\{q(E_k - E^*), 2\alpha^2\right\}.$$

In particular, we immediately deduce $\sqrt{E_{k+1}} \leq D$. Applying the inductive hypothesis, we conclude

$$\begin{aligned} E_{k+1} - E^* &\leq \max\left\{q \cdot \max\{q^k(E_0 - E^*), 2\alpha^2\}, 2\alpha^2\right\} \\ &= \max\{q^{k+1}(E_0 - E^*), 2\alpha^2\}. \end{aligned}$$

The theorem is proved. $\square$

6 Geometrically Decaying Stepsize

In the previous section, we showed linear convergence of the constant stepsize scheme up to a fixed tolerance E^* . To obtain a linearly convergent method to the true solution set, we will allow the stepsize to decrease geometrically. The analogous strategy in the convex setting goes back to Goffin [12], and our argument is a direct generalization of [12] to the weakly convex setting. The intuition for why one may expect linear convergence under such stepsizes may be gleaned from the Polyak method under the optimal stepsize

$$\alpha_k = \frac{g(x_k) - \min_{x \in \mathcal{X}} g(x)}{\|\zeta_k\|}.$$

It is easy to verify that since the E_k tend to zero Q -linearly, the stepsizes α_k tend to zero R -linearly. We implement such a geometrically decaying stepsize in Algorithm 3 and prove linear convergence of the method in Theorem 6.1. In the proof, we use the notation $E_k := \text{dist}^2(x_k; \mathcal{X}^*)$ as in the previous section.

Algorithm 3: Subgradient method with geometrically decaying stepsize

Data: Constants $\lambda > 0$ and $0 < q < 1$.

Step k : ($k \geq 0$)

Choose $\zeta_k \in \partial g(x_k)$

Set stepsize $\alpha_k = \lambda \cdot q^k$

if $\zeta_k \neq 0$ **then**

Set $x_{k+1} = \text{proj}_{\mathcal{X}} \left(x_k - \alpha_k \frac{\zeta_k}{\|\zeta_k\|} \right)$

else

Set $x_{k+1} = x_k$

end

Theorem 6.1 Fix a real $0 < \gamma < 1$ and suppose $\tau \leq \sqrt{\frac{1}{2-\gamma}}$. Set

$$0 < \lambda \leq \frac{\gamma \mu^2}{\rho L} \quad \text{and} \quad q := \sqrt{1 - (1 - \gamma)\tau^2},$$

where λ is assumed finite. Then, the iterates x_k generated by Algorithm 3, initialized at some point x_0 with $\text{dist}(x_0; \mathcal{X}^*) \leq \frac{\lambda}{\tau - \sqrt{\tau^2 - (1 - \rho\lambda L - 1 - q^2)}}$, satisfy:

$$\text{dist}^2(x_k; \mathcal{X}^*) \leq \max \left\{ \frac{\lambda^2}{\tau^2}, \text{dist}(x_0; \mathcal{X}^*)^2 \right\} \left(1 - (1 - \gamma)\tau^2 \right)^k. \quad (15)$$

Proof We will prove the result by induction.

To this end, suppose the bound (15) holds for all $i = 0, \dots, k$. Appealing to Lemma 5.1 and using the relation $\alpha_k = \lambda q^k$, we obtain

$$E_{k+1} \leq \left(1 + \frac{\rho \lambda q^k}{L} \right) E_k - 2\lambda \tau q^k \sqrt{E_k} + \lambda^2 q^{2k}. \quad (16)$$

Define the constant $M := \max \left\{ \frac{\lambda}{\tau}, \text{dist}(x_0; \mathcal{X}^*) \right\}$. Recall the induction assumption guarantees $\sqrt{E_k} \leq M q^k$. Let us therefore fix some value $R \in [0, M]$ satisfying $\sqrt{E_k} = R q^k$. Inequality (16) then implies

$$E_{k+1} \leq \max_{R \in [0, M]} \left\{ R^2 q^{2k} + \frac{\rho \lambda R^2}{L} q^{3k} - 2\lambda \tau R q^{2k} + \lambda^2 q^{2k} \right\}.$$

Note that the expression inside the maximum is a convex quadratic in R and therefore the maximum must occur either at $R = 0$ or $R = M$. We therefore deduce

$$E_{k+1} \leq q^{2k} \cdot \max \left\{ \lambda^2, M^2 + \frac{\rho \lambda}{L} M^2 q^k - 2\lambda \tau M + \lambda^2 \right\}. \quad (17)$$

To complete the induction, it is therefore sufficient to show

$$\lambda^2 \leq M^2 q^2 \quad \text{and} \quad M^2 + \frac{\rho\lambda}{L} M^2 q^k - 2\lambda\tau M + \lambda^2 \leq M^2 q^2. \quad (18)$$

We begin by verifying the first property. Note $M \geq \frac{\lambda}{\tau}$. Hence, it suffices to show that $\tau \leq q$. Observe that the assumption $\tau \leq \sqrt{\frac{1}{2-\gamma}}$ directly implies

$$\tau^2 + (1 - \gamma)\tau^2 \leq 1.$$

Rearranging yields $\tau^2 \leq 1 - (1 - \gamma)\tau^2 = q^2$. Hence, the first condition in (18) holds.

Next, we show that M satisfies the second property in (18). Rearranging the expression, we must establish

$$\left(1 + \frac{\rho\lambda}{L} q^k - q^2\right) M^2 - 2\lambda\tau M + \lambda^2 \leq 0. \quad (19)$$

We will show that the quadratic on the left-hand side in M has two real positive roots. To this end, a quick computation shows that the two roots are

$$\frac{\lambda\tau \pm \sqrt{\lambda^2\tau^2 - \lambda^2\left(1 + \frac{\rho\lambda}{L} q^k - q^2\right)}}{1 + \frac{\rho\lambda}{L} q^k - q^2} = \frac{\lambda}{\tau \mp \sqrt{\tau^2 - \left(1 + \frac{\rho\lambda}{L} q^k - q^2\right)}}.$$

To see that the discriminant is nonnegative, observe

$$\tau^2 - \left(1 + \frac{\rho\lambda}{L} q^k - q^2\right) \geq \tau^2 - \left(1 + \frac{\rho\lambda}{L} - q^2\right) \geq \tau^2 - (1 + \gamma\tau^2 - q^2) = 0.$$

Thus, the convex quadratic in (19) has two real roots. Observe $M \geq \lambda/\tau$, and therefore M is greater than the minimal root. It suffices therefore to argue that M is smaller than the largest root. If $M = \lambda/\tau$, then this is clearly true. Otherwise, it must be that $M = \sqrt{E_0}$. Our assumption on the initial distance $\sqrt{E_0}$ then guarantees that $\sqrt{E_0}$ is smaller than the largest positive root of (19), as claimed. Hence, the condition (19) holds, and the induction is complete. $\square$

In the convex setting, $\rho = 0$, Theorem 6.1 recovers the guarantees in Goffin [12]. Theorem 6.1 demonstrates a precise relationship between the initial distance, the constant λ which controls the rate, and the constants τ and ρ . In particular, the larger the initial distance, the larger λ needs to be to guarantee convergence. It is now straightforward to see how one should set the parameters q and λ to guarantee linear convergence within the tube $\mathcal{T}_\gamma$.

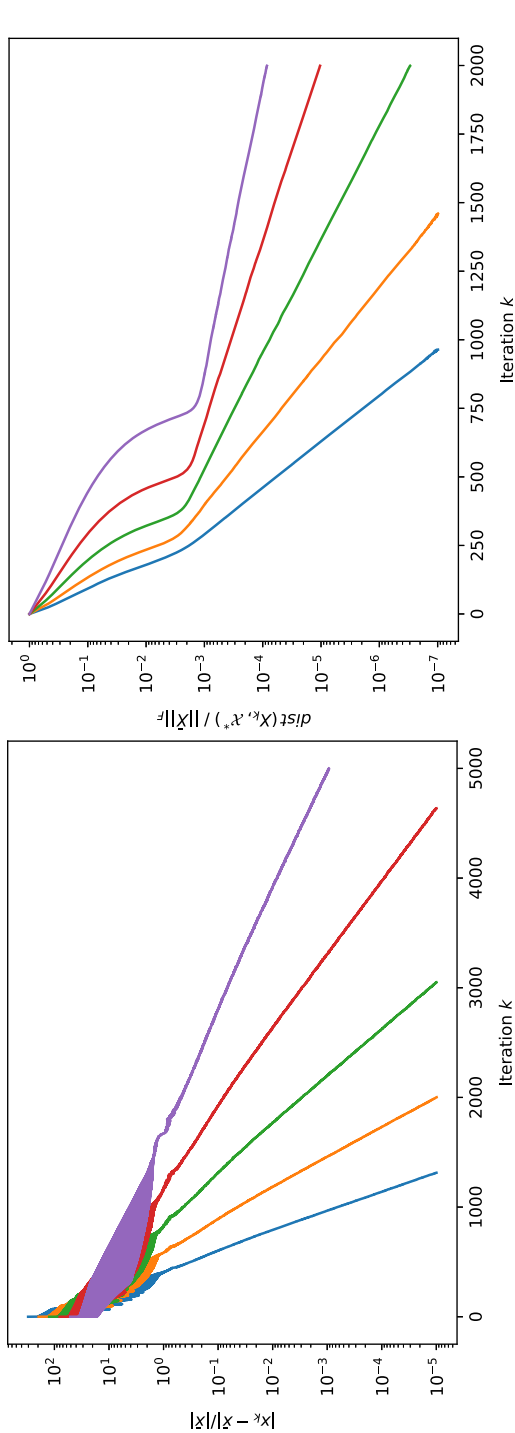


Fig. 3 Geometrically decaying stepsize. (Left) Phase retrieval with the corrupted setup; $d = 1000$, $m = 3000$, $q \in \{0.983, 0.989, 0.993, 0.996, 0.997\}$; (Right) Covariance matrix estimation with the corrupted setup; $d = 1000$, $r = 3$, $m = 10,000$, $q \in \{0.986, 0.991, 0.994, 0.996, 0.998\}$. The depicted rates uniformly improve with lower values of q , in both figures

Corollary 6.1 Fix a real $0 < \gamma < 1$ and suppose $\tau \leq \sqrt{\frac{1}{2-\gamma}}$ and $\rho > 0$. Set

$$\lambda := \frac{\gamma\mu^2}{\rho L} \quad \text{and} \quad q := \sqrt{1 - (1 - \gamma)\tau^2}.$$

Then, the iterates x_k generated by Algorithm 3, initialized at any point $x_0 \in \mathcal{T}_\gamma$, satisfy:

$$\text{dist}^2(x_k; \mathcal{X}^*) \leq \frac{\gamma^2\mu^2}{\rho^2} \left(1 - (1 - \gamma)\tau^2\right)^k. \quad (20)$$

Proof Looking back at Theorem 6.1, we see that $\text{dist}(x_0; \mathcal{X}^*)$ satisfies the assumed upper bound:

$$\text{dist}(x_0; \mathcal{X}^*) < \frac{\gamma\mu}{\rho} = \frac{\lambda}{\tau} \leq \frac{\lambda}{\tau - \sqrt{\tau^2 - (1 - \rho\lambda L^{-1} - q^2)}}.$$

Since $\frac{\lambda}{\tau} = \frac{\gamma\mu}{\rho}$, we conclude that $\max\{\text{dist}(x_0; \mathcal{X}^*), \frac{\lambda}{\tau}\} \leq \frac{\gamma\mu}{\rho}$. The result immediately follows after applying Theorem 6.1. $\square$

We now illustrate the performance of Algorithm 3 on our two running examples in Fig. 3. Empirically, we observed that $\lambda > 0$ and $0 < q < 1$ must be tuned for performance, which is what we did in the experiments. We observe linear convergence in all cases, and the convergence rate of the method improves monotonically as the chosen rate q is decreased. The Polyak scheme pictured in Fig. 1 clearly outperforms all other methods, while the geometrically decaying stepsize scheme performs much better than the constant stepsize scheme in Fig. 2. That being said, we should stress that the Polyak subgradient method requires knowledge of the optimal value of the problem, which is often unavailable. We also note that the work [15] provides schemes which can indeed automatically adapt to the unknown parameter μ for convex problems. We leave extensions of such algorithms to weakly convex problems for future work.

Acknowledgements Research of Drusvyatskiy and MacPhee was partially supported by the AFOSR YIP Award FA9550-15-1-0237 and by the NSF DMS 1651851 and CCF 1740551 Awards. Research of Paquette was supported by NSF CCF 1740796.

References

1. Jane, P., Netrapalli, P.: Fast exact matrix completion with finite samples. In: Grnwald, P., Hazan, E., Kale, S. (eds.) Proceedings of The 28th Conference on Learning Theory, Proceedings of Machine Learning Research, vol. 40, pp. 1007–1034. PMLR, Paris, France (2015). <http://proceedings.mlr.press/v40/Jain15.html>
2. Tu, S., Boczar, R., Simchowitz, M., Soltanolkotabi, M., Recht, B.: Low-rank solutions of linear matrix equations via procrustes flow. In: Proceedings of the 33rd International Conference on International Conference on Machine Learning—Volume 48, ICML'16, pp. 964–973. JMLR.org (2016). <http://dl.acm.org/citation.cfm?id=3045390.3045493>
3. Candès, E., Li, X., Soltanolkotabi, M.: Phase retrieval via Wirtinger flow: theory and algorithms. IEEE Trans. Inf. Theory **61**(4), 1985–2007 (2015). <https://doi.org/10.1109/TIT.2015.2399924>

4. Jain, P., Jin, C., Kakade, S., Netrapalli, P.: Global convergence of non-convex gradient descent for computing matrix squareroot. In: Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (AISTATS), pp. 479–488 (2017). <http://proceedings.mlr.press/v54/jain17a.html>
5. Chen, Y., Wainwright, M.: Fast low-rank estimation by projected gradient descent: general statistical and algorithmic guarantees (2015). [arXiv:1509.03025](https://arxiv.org/abs/1509.03025)
6. Meka, R., Jain, P., Dhillon, I.: Guaranteed rank minimization via singular value projection (2009). [arXiv:0909.5457](https://arxiv.org/abs/0909.5457)
7. Boumal, N.: Nonconvex phase synchronization. *SIAM J. Optim.* **26**(4), 2355–2377 (2016)
8. Brutzkus, A., Globerson, A.: Globally optimal gradient descent for a ConvNet with Gaussian inputs (2017). [arXiv:1702.07966](https://arxiv.org/abs/1702.07966)
9. Eremin, I.: The relaxation method of solving systems of inequalities with convex functions on the left-hand side. *Dokl. Akad. Nauk SSSR* **160**, 994–996 (1965)
10. Poljak, B.: Minimization of unsmooth functionals. *USSR Comput. Math. Math. Phys.* **9**, 14–29 (1969)
11. Shor, N.: The rate of convergence of the method of the generalized gradient descent with expansion of space. *Kibernetika (Kiev)* **2**, 80–85 (1970)
12. Goffin, J.: On convergence rates of subgradient optimization methods. *Math. Program.* **13**(3), 329–347 (1977). <https://doi.org/10.1007/BF01584346>
13. Supittayapornpong, S., Neely, M.: Staggered time average algorithm for stochastic non-smooth optimization with $O(1/t)$ convergence (2016). [arXiv:1607.02842](https://arxiv.org/abs/1607.02842)
14. Yang, T., Lin, Q.: RSG: beating subgradient method without smoothness and strong convexity (2016). [arXiv:1512.03107](https://arxiv.org/abs/1512.03107)
15. Johnstone, P., Moulin, P.: Faster subgradient methods for functions with Hölderian growth (2017). [arXiv:1704.00196](https://arxiv.org/abs/1704.00196)
16. Polyak, B.: Sharp minima. Institute of Control Sciences Lecture Notes, Moscow, USSR; Presented at the IIASA Workshop on Generalized Lagrangians and Their Applications, IIASA, Laxenburg, Austria, vol. 3, pp. 369–380 (1979)
17. Burke, J., Ferris, M.: Weak sharp minima in mathematical programming. *SIAM J. Control Optim.* **31**(5), 1340–1359 (1993). <https://doi.org/10.1137/0331063>
18. Poljak, B.: Subgradient methods: a survey of Soviet research. In: Nonsmooth Optimization (Proc. IIASA Workshop, Laxenburg, 1977), IIASA Proc. Ser., vol. 3, pp. 5–29. Pergamon, Oxford (1978)
19. Nurminkii, E.A.: The quasigradient method for the solving of the nonlinear programming problems. *Cybernetics* **9**(1), 145–150 (1973). <https://doi.org/10.1007/BF01068677>
20. Nurminkii, E.A.: Minimization of nondifferentiable functions in the presence of noise. *Cybernetics* **10**(4), 619–621 (1974). <https://doi.org/10.1007/BF01071541>
21. Davis, D., Drusvyatskiy, D.: Stochastic subgradient method converges at the rate $O(k^{-1/4})$ on weakly convex functions (2018). [arXiv:1802.02988](https://arxiv.org/abs/1802.02988)
22. Drusvyatskiy, D., Lewis, A.: Error bounds, quadratic growth, and linear convergence of proximal methods. To appear in *Mathematics of Operations Research* (2016). [arXiv:1602.06661](https://arxiv.org/abs/1602.06661)
23. Duchi, J., Ruan, F.: Stochastic methods for composite optimization problems (2017). Preprint [arXiv:1703.08570](https://arxiv.org/abs/1703.08570)
24. Duchi, J., Ruan, F.: Solving (most) of a set of quadratic equalities: composite optimization for robust phase retrieval (2017). [arXiv:1705.02356](https://arxiv.org/abs/1705.02356)
25. Drusvyatskiy, D., Paquette, C.: Efficiency of minimizing compositions of convex functions and smooth maps (2016). Preprint [arXiv:1605.00125](https://arxiv.org/abs/1605.00125)
26. Davis, D., Grimmer, B.: Proximally guided stochastic method for nonsmooth, nonconvex problems (2017). Preprint [arXiv:1707.03505](https://arxiv.org/abs/1707.03505)
27. Drusvyatskiy, D.: The proximal point method revisited. To appear in *SIAG/OPT Views and News* (2018). [arXiv:1712.06038](https://arxiv.org/abs/1712.06038)
28. Rockafellar, R., Wets, R.B.: Variational Analysis. Grundlehren der mathematischen Wissenschaften, vol. 317. Springer, Berlin (1998)
29. Mordukhovich, B.: Variational Analysis and Generalized Differentiation I: Basic Theory. Grundlehren der mathematischen Wissenschaften, vol. 330. Springer, Berlin (2006)
30. Penot, J.P.: Calculus Without Derivatives, Graduate Texts in Mathematics, vol. 266. Springer, New York (2013). <https://doi.org/10.1007/978-1-4614-4538-8>
31. Ioffe, A.: Variational Analysis of Regular Mappings. Springer Monographs in Mathematics. Springer, Cham (2017). <https://doi-org.offcampus.lib.washington.edu/10.1007/978-3-319-64277-2>. Theory and Applications

32. Poliquin, R., Rockafellar, R.: Prox-regular functions in variational analysis. *Trans. Am. Math. Soc.* **348**, 1805–1838 (1996)
33. Eldar, Y., Mendelson, S.: Phase retrieval: stability and recovery guarantees. *Appl. Comput. Harmon. Anal.* **36**(3), 473–494 (2014). <https://doi.org/10.1016/j.acha.2013.08.003>
34. Davis, D., Drusvyatskiy, D., Paquette, C.: The nonsmooth landscape of phase retrieval (2017). [arXiv:1711.03247](https://arxiv.org/abs/1711.03247)
35. Chen, Y., Candès, E.: Solving random quadratic systems of equations is nearly as easy as solving linear systems. *Commun. Pure Appl. Math.* **70**(5), 822–883 (2017). <https://doi-org.offcampus.lib.washington.edu/10.1002/cpa.21638>
36. Sun, J., Qu, Q., Wright, J.: A geometric analysis of phase retrieval. To appear in *Foundations of Computational Mathematics* (2017). [arXiv:1602.06664](https://arxiv.org/abs/1602.06664)
37. Tan, Y., Vershynin, R.: Phase retrieval via randomized Kaczmarz: theoretical guarantees (2017). [arXiv:1605.08285](https://arxiv.org/abs/1605.08285)
38. Chen, Y., Chi, Y., Goldsmith, A.J.: Exact and stable covariance estimation from quadratic sampling via convex programming. *IEEE Trans. Inf. Theory* **61**(7), 4034–4059 (2015)
39. Clarke, F.: *Optimization and Nonsmooth Analysis*. Wiley, New York (1983)