

EXACT CONVERGENCE RATE OF THE LAST ITERATE IN SUBGRADIENT METHODS*

MOSLEM ZAMANI[†] AND FRANÇOIS GLINEUR[‡]

Abstract. We study the convergence of the last iterate in subgradient methods applied to the minimization of a nonsmooth convex function with bounded subgradients. Using a novel proof technique that tracks distances to varying reference points, we derive tight worst-case convergence rates for the (projected) subgradient method with constant step sizes or step lengths. We identify the optimal constant step size for a given number of iterations N , yielding a last-iterate accuracy smaller than $\frac{BR}{\sqrt{N+1}}\sqrt{1 + \log(N)}/4$, where B bounds subgradient norms and R bounds the initial distance to a minimizer. We also propose a new optimal subgradient method based on a linearly decaying sequence of step sizes, which achieves a rate for the last iterate equal to $\frac{BR}{\sqrt{N+1}}$, matching the theoretical lower bound. Finally, we show that no universal step size sequence can attain this rate across all iterations, highlighting that the dependence of the step size sequence in N is unavoidable.

Key words. convex optimization, nonsmooth optimization, subgradient method, constant step size, constant step length, convergence rate, last iterate, optimal subgradient method

MSC codes. 90C25, 90C60, 49J52

1. Introduction.

1.1. Subgradient methods. Subgradient methods are iterative techniques for solving nonsmooth convex optimization problems, studied by N. Shor and others in the 1960s [20, 21]. They are both simple and widely used, and continue to be actively studied. New variants have been recently developed that can handle a wider range of optimization problems, see [2, 8] and the references therein.

Let $X \subseteq \mathbb{R}^n$ be a convex set and f be a convex function whose domain contains X . Consider the following convex optimization problem

$$(1.1) \quad \min_{x \in X} f(x).$$

The set of subgradients of function f at a point x , denoted as $\partial f(x)$, is given by

$$\partial f(x) = \{g \in \mathbb{R}^n \text{ such that } f(y) \geq f(x) + \langle g, y - x \rangle \text{ holds for all } y \in \text{dom} f\}.$$

The class of projected subgradient methods is given in Algorithm 1.1, where $P_X(\cdot)$ stands for the Euclidean projection on X . An instance of the method requires to define the sequence of N step sizes $\{h_k\}_{1 \leq k \leq N}$, where N is the number of iterations to perform. Moreover, if $x^k = x^{k+1}$ in Algorithm 1.1, then x^k will be an optimal solution., and the algorithm terminates.

For the method to be well-defined, we will assume the following throughout the paper:

ASSUMPTION 1.1.

1. The subdifferential $\partial f(x)$ is nonempty for every $x \in X$.

*Submitted to the editors June 30, 2025

Funding: This work was supported by the Fonds de la Recherche Scientifique-FNRS under Grant EOS 30468160.

[†]ICTEAM/INMA, Université catholique de Louvain, B-1348 Louvain-la-Neuve, Belgium (moslem.zamani@uclouvain.be).

[‡]ICTEAM/INMA & CORE, Université catholique de Louvain, B-1348 Louvain-la-Neuve, Belgium (Francois.Glineur@uclouvain.be).

Algorithm 1.1 Projected subgradient method with generic step sizes

Parameters: number of iterations N , sequence of positive step sizes $\{h_k\}_{1 \leq k \leq N}$.

Inputs: convex set X , convex function f defined on X , initial iterate $x^1 \in X$.

For $k = 1, 2, \dots, N$ perform the following steps:

1. Select a subgradient $g^k \in \partial f(x^k)$.
2. Compute $x^{k+1} = P_X(x^k - h_k g^k)$.

Output: last iterate x_{N+1}

2. The set X is closed, convex and nonempty.

The first assumption is necessary to compute a direction for the next iterate. It holds for example if set X is contained in the relative interior of the domain of function f . The second assumption ensures that the projection on X is well-defined and unique.

1.2. Convergence rates. Unlike the gradient method, the subgradient method is not a descent method, meaning that inequality $f(x_{k+1}) \leq f(x_k)$ does not necessarily hold at each iteration. For this reason, most convergence rates for the subgradient method describe the best iterate, or an average of the iterates, see for example [1, 12, 21, 14]. Convergence results typically require two more assumptions:

ASSUMPTION 1.2. Function f has B -bounded subgradients on set X , meaning

$$x \in X \text{ and } g \in \partial f(x) \Rightarrow \|g\| \leq B.$$

Note that a convex function f is Lipschitz continuous with modulus B if its subgradients are B -bounded on its domain.

ASSUMPTION 1.3. Function f admits a minimizer x^* , and the distance between initial iterate x^1 and x^* is bounded by a constant R , that is

$$\|x^1 - x^*\| \leq R.$$

For example, under these assumptions, the best iterate of the subgradient method with generic positive step sizes $\{h_k\}$ will satisfy (see e.g. [1, 12, 21])

$$(1.2) \quad \min_{1 \leq k \leq N+1} f(x^k) - f(x^*) \leq \frac{R^2 + B^2 \sum_{k=1}^{N+1} h_k^2}{2 \sum_{k=1}^{N+1} h_k}.$$

The same bound can be shown to also hold for the average iterate defined as

$$x^{\text{avg}} = \sum_{k=1}^{N+1} w_k x^k \quad \text{with } w_k = \frac{h_k}{\sum_{k=1}^{N+1} h_k}.$$

Note that these rates depend on step size h_{N+1} that is not used in the algorithm (i.e. they are valid for any value of $h_{N+1} \geq 0$). This already indicates that the simple bound (1.2) may not be optimal. If we know constants B and R , it can be shown that the optimal choice of step sizes, i.e. which minimizes the rate, consists of the following constant sequence

$$h_k = \frac{R}{B} \frac{1}{\sqrt{N+1}} \quad \text{which implies} \quad \min_{1 \leq k \leq N+1} f(x^k) - f(x^*) \leq \frac{BR}{\sqrt{N+1}}.$$

A final remark is that this last result cannot be improved. Indeed, it is known that no sequence of step sizes in Algorithm 1.1 can achieve a rate better than $\frac{BR}{\sqrt{N+1}}$ [7, Theorem 2]. This lower bound is actually valid for any black-box method that moves in a direction combining past subgradients at each step, see beginning of Section 5 for more details.

1.3. Rates on the final iterate. From the above it appears that subgradient methods are both simple and efficient, matching the best possible convergence rate. Nevertheless, we observe two drawbacks: first, the optimal sequence of step sizes $h_k = \frac{R}{B} \frac{1}{\sqrt{N+1}}$ requires knowledge of constants B and R and, more importantly, depends on the number of iterations N .

Second, these worst-case convergence rates only hold for the best or the average iterate, and nothing is guaranteed about the sequence of iterates, or in particular about the last iterate. As the final iterate is commonly selected in practice as the output of the subgradient method [11], analyzing its convergence rate may be of interest. In some situations the sequence of iterates computed by the subgradient method is itself under study, e.g. when the subgradient method is an adjustment strategy for approaching a stable state of some system; see [17].

The question of last-iterate convergence rates was previously raised in [19]. In [16] the authors introduce a modified subgradient method with double averaging for which the whole sequence of iterates converges with the rate $O(\frac{1}{\sqrt{N}})$. Moving back to standard subgradient methods as described by Algorithm 1.1, a lower bound of order $O(\frac{\log(N)}{\sqrt{N}})$ for the convergence rate for the last iterate is established in [9, 10] for a specific choice of step sizes $h_k = \frac{R}{B} \frac{1}{\sqrt{k}}$. The authors also prove a high probability upper bound with the same order $O(\frac{\log(N)}{\sqrt{N}})$ in the stochastic case. Finally, we can find in [11] a subgradient method with a different choice of step sizes for which a $O(\frac{1}{\sqrt{N}})$ convergence rate for the last iterate is obtained when the feasible set X is bounded. For more discussions on this topic, the interested reader can refer to [18].

In this paper, we continue to explore this question and contribute in two ways. First, we establish in Section 3 exact convergence rates for the last iterate of the subgradient method with either constant step sizes or constant step lengths. These results are based on a key lemma presented in Section 2, which generalizes the standard analysis of subgradient methods. Second, we present in Section 5 an optimal subgradient method for which the last-iterate convergence rate matches exactly the established lower bound for black-box nonsmooth convex optimization problems, namely for which we have $f(x^N) - f(x^*) \leq \frac{BR}{\sqrt{N+1}}$, improving the constant in the rate of [11] by an order of magnitude.

2. Key lemma for convergence proofs. All convergence rates established in this paper will be derived from the following key lemma. Its proof is based on tracking the distance between the current iterate and a different reference point at each iteration ($\|x^k - z^k\|$ in the proof below, where z^k is constructed as a convex combination of the iterates and an arbitrary feasible point—typically x^*). In classical analyses of subgradient methods, convergence rates are often derived by bounding the distance between the current iterate and a fixed optimal solution, leveraging the subgradient inequality at the current iterate x^k and x^* . By introducing the auxiliary sequence $\{z^k\}$, our analysis gains additional flexibility, as it enables the application of the subgradient inequality at x^k and $x^1, \dots, x^{k-1}, x^*$. This refinement enables sharper convergence estimates by aggregating information across iterates. Note however that

a consequence of the above is that classical analyses hold for star-convex functions, while our improved version does not.

LEMMA 2.1. *Let f be a convex function and let X be a closed convex set. Suppose that $\hat{x} \in X$, $h_{N+1} \geq 0$ and $0 < v_0 \leq v_1 \leq \dots \leq v_N \leq v_{N+1}$. If Algorithm 1.1 with the starting point $x^1 \in X$ generates iterates and subgradients $\{(x^k, g^k)\}$ for $1 \leq k \leq N+1$, then*

$$(2.1) \quad \sum_{k=1}^{N+1} \left(h_k v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} h_i v_i \right) f(x^k) - v_0 \sum_{k=1}^{N+1} h_k v_k f(\hat{x}) \leq \frac{v_0^2}{2} \|x^1 - \hat{x}\|^2 + \frac{1}{2} \sum_{k=1}^{N+1} h_k^2 v_k^2 \|g^k\|^2.$$

Note that this inequality can also be equivalently written

$$\sum_{k=1}^{N+1} c_k (f(x_k) - f(\hat{x})) \leq \frac{v_0^2}{2} \|x^1 - \hat{x}\|^2 + \frac{1}{2} \sum_{k=1}^{N+1} h_k^2 v_k^2 \|g^k\|^2$$

with coefficients c_k defined as $c_k = h_k v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} h_i v_i$, since one can show that $\sum_{k=1}^{N+1} c_k = v_0 \sum_{k=1}^{N+1} h_k v_k$ using summation by parts.

Proof. Let $v_{N+2} = v_{N+1}$ and $z^0 = \hat{x}$. We define a dummy variable $x^{N+2} = P_X(x^{N+1} - h_{N+1}g^{N+1})$. Let z^k be defined as follows,

$$(2.2) \quad z^k = \left(1 - \frac{v_{k-1}}{v_k}\right) x^k + \frac{v_{k-1}}{v_k} z^{k-1}, \quad k \in \{1, \dots, N+2\}.$$

Due to the definition z^k may be written as a convex combination of $x^1, \dots, x^k, \hat{x}$. Indeed,

$$(2.3) \quad z^k = \frac{1}{v_k} \sum_{i=1}^k (v_i - v_{i-1}) x^i + \frac{v_0}{v_k} \hat{x}.$$

By (2.2), we have $v_{k+1} \|x^{k+1} - z^{k+1}\| = v_k \|x^{k+1} - z^k\|$. Using the non-expansiveness of the projection operator and the fact that $z^k \in X$, we obtain

$$\|P_X(x^k - h_k g^k) - z^k\| \leq \|x^k - h_k g^k - z^k\|.$$

Thus,

$$\begin{aligned} v_{k+1}^2 \|x^{k+1} - z^{k+1}\|^2 &\leq v_k^2 \|(x^k - h_k g^k) - z^k\|^2 \\ &= v_k^2 \left(\|x^k - z^k\|^2 + 2h_k \langle g^k, z^k - x^k \rangle + h_k^2 \|g^k\|^2 \right) \\ &\leq v_k^2 \|x^k - z^k\|^2 + 2h_k v_k^2 (f(z^k) - f(x^k)) + h_k^2 v_k^2 \|g^k\|^2, \end{aligned}$$

where the last inequality follows from the subgradient inequality. Hence,

$$h_k v_k^2 (f(x^k) - f(z^k)) \leq \frac{v_k^2}{2} \|x^k - z^k\|^2 - \frac{v_{k+1}^2}{2} \|x^{k+1} - z^{k+1}\|^2 + \frac{v_k^2 h_k^2}{2} \|g^k\|^2.$$

Summing up these inequalities over $k \in \{1, \dots, N+1\}$, we obtain

$$\sum_{k=1}^{N+1} h_k v_k^2 (f(x^k) - f(z^k)) \leq \frac{v_1^2}{2} \|x^1 - z^1\|^2 - \frac{v_{N+2}^2}{2} \|x^{N+2} - z^{N+2}\|^2 + \sum_{k=1}^{N+1} \frac{h_k^2 v_k^2}{2} \|g^k\|^2.$$

As $x^1 - z^1 = \frac{v_0}{v_1} (x^1 - \hat{x})$, we get

$$(2.4) \quad \sum_{k=1}^{N+1} h_k v_k^2 (f(x^k) - f(z^k)) \leq \frac{v_0^2}{2} \|x^1 - \hat{x}\|^2 + \sum_{k=1}^{N+1} \frac{h_k^2 v_k^2}{2} \|g^k\|^2.$$

On the other hand, by Jensen's inequality and (2.3), we get

$$(2.5) \quad \sum_{k=1}^{N+1} h_k v_k^2 (f(x^k) - f(z^k)) \geq \sum_{k=1}^{N+1} h_k v_k^2 f(x^k) - \sum_{k=1}^{N+1} \sum_{i=1}^k h_k v_k (v_i - v_{i-1}) f(x^i) - v_0 \sum_{k=1}^{N+1} h_k v_k f(\hat{x}).$$

Inequalities (2.4) and (2.5) imply the desired inequality and the proof is complete. $\square$

Compared to the standard analysis of subgradient methods, additional flexibility is provided by the sequence of weights $\{v_k\}$. Note that by setting $v_k = 1$ for all k and $\hat{x} = x^*$ in (2.1), we get

$$(2.6) \quad \sum_{k=1}^{N+1} h_k (f(x^k) - f^*) \leq \frac{1}{2} \|x^1 - x^*\|^2 + \frac{1}{2} \sum_{k=1}^{N+1} h_k^2 \|g^k\|^2.$$

from which it is straightforward to derive the standard convergence rate (1.2) with respect to the best objective value or the average of iterates [1, 12].

In order to establish a last-iterate convergence rate via Lemma 2.1, one should choose appropriate values for the $N+3$ parameters, $v_0, \dots, v_{N+1}$ and h_{N+1} , so that coefficients c_k are zero for all k except c_{N+1} . One can actually see with some algebra that all parameters are uniquely determined if we assign values to v_{N+1} , h_{N+1} and the coefficient c_{N+1} of $f(x^{N+1})$ in (2.1).

After the submission of this paper to arXiv, some authors extended Lemma 2.1 to broader settings; see [13]. In this interesting line of work, Defazio et al. [3, 4] investigate the convergence of the stochastic subgradient method with respect to the last iterate. They also establish the optimal step sizes presented in Section 5 for the stochastic setting, based on [3, Theorem 3]. In the sequel, we show that a deterministic counterpart of [3, Theorem 3] can be derived directly from Lemma 2.1.

PROPOSITION 2.2. *Let $w_1, \dots, w_{N+1}$ be a sequence of positive numbers and let $S = \sum_{i=1}^{N+1} w_i$. If Algorithm 1.1 with the starting point $x^1 \in X$ and step sizes $h_k = \frac{w_k}{S} \sum_{i=k+1}^{N+1} w_i$ generates iterates and subgradients $\{(x^k, g^k)\}$ for $1 \leq k \leq N+1$, then we have*

$$(2.7) \quad f(x^{N+1}) - f(x^*) \leq \frac{1}{2S} \|x^1 - x^*\|^2 + \frac{1}{2S} \sum_{k=1}^{N+1} w_k^2 \|g^k\|^2.$$

164 *Proof.* The proof follows from a direct application of Lemma 2.1. Let $S_k =$
 165 $\sum_{i=k+1}^{N+1} w_i$ for $k \in \{0, \dots, N\}$. We define

$$166 \quad v_k = \frac{\sqrt{S}}{S_k} \quad k \in \{0, \dots, N\},$$

167 $v_{N+1} = v_N$, $\hat{x} = x^*$ and $h_{N+1} = \frac{w_{N+1}^2}{S}$. It is readily seen that $0 < v_0 \leq v_1 \leq \dots \leq$
 168 $v_N \leq v_{N+1}$. For each $k \in \{1, \dots, N\}$, a direct computation yields

$$169 \quad h_k v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} h_i v_i = \frac{w_k}{S_k} - \left(\frac{w_k \sqrt{S}}{S_{k-1} S_k} \right) \left(\frac{S_{k-1}}{\sqrt{S}} \right) = 0.$$

170 In addition, $h_{N+1} v_{N+1}^2 - (v_{N+1} - v_N) \sum_{i=N+1}^{N+1} h_i v_i = 1$. On the other hand, we have
 171 $h_k^2 v_k^2 = \frac{w_k^2}{S}$ for $k \in \{1, \dots, N+1\}$. Substituting the above relations into (2.1), we
 172 obtain (2.7) and the proof is complete. $\square$

173 It is worth noting that for any step sizes $h_1, \dots, h_N$, there exist $w_1, \dots, w_{N+1}$
 174 that satisfy the assumptions of Proposition 2.2; see [3, Theorem 9].

175 **3. Subgradient method with constant step sizes.** In this section, we inves-
 176 tigate the convergence rate of Algorithm 1.1 when the step size is constant for each
 177 iteration. Moreover, following the standard presentation of such convergence results,
 178 we assume that this constant step size is chosen proportionally to the ratio $\frac{R}{B}$, namely
 179 we define $h_k = \frac{hR}{B}$ ($k = 1, \dots, N$) for some $h > 0$. This normalization leads to slightly
 180 simpler expressions for the rates, which become proportional to a common factor BR .

181 **3.1. Increasing sequences $\{s_{\alpha,k}\}_{k \geq 1}$.** Before we prove our results we need to
 182 introduce a family of real sequences. Let $\alpha \geq 1$ be a real parameter. We define the
 183 sequence $\{s_{\alpha,k}\}_{k \geq 1}$ recursively as follows

$$184 \quad (3.1) \quad s_{\alpha,1} = \alpha, \quad s_{\alpha,k+1} = s_{\alpha,k} + \frac{1}{s_{\alpha,k}}.$$

185 The next proposition lists some properties of these sequences that will be used later.

186 **PROPOSITION 3.1.** *Any sequence $\{s_{\alpha,k}\}$ defined by (3.1) satisfies the following for*
 187 *all k :*

- 188 (a) $s_{\alpha,k+1} = \alpha + \sum_{i=1}^k \frac{1}{s_{\alpha,i}}$
- 189 (b) $s_{\alpha,k+1}^2 = \alpha^2 + 2k + \sum_{i=1}^k \frac{1}{s_{\alpha,i}^2}$
- 190 (c) $\beta > \alpha$ implies $s_{\beta,k} > s_{\alpha,k}$
- 191 (d) $\lim_{\alpha \rightarrow +\infty} s_{\alpha,k} = +\infty$.

192 *Proof.*

- 193 (a) This follows from telescoping in the sum of defining equalities $s_{\alpha,i+1} = s_{\alpha,i} + \frac{1}{s_{\alpha,i}}$
 194 for i ranging from 1 to k .
- 195 (b) Squaring the defining equality gives $s_{\alpha,i+1}^2 = (s_{\alpha,i} + \frac{1}{s_{\alpha,i}})^2 = s_{\alpha,i}^2 + 2 + \frac{1}{s_{\alpha,i}^2}$.
 196 Summing for i ranging from 1 to k and telescoping provides the result.
- 197 (c) This follows from the fact that $s \mapsto s + \frac{1}{s}$ is strictly increasing when $s \geq 1$.
- 198 (d) This follows from $\lim_{\alpha \rightarrow \infty} s_{\alpha,1} = +\infty$ and the fact that each sequence $\{s_{\alpha,k}\}$ is
 199 strictly increasing. $\square$

200 The sequence in the particular case $\alpha = 1$ will play a central role in our convergence
 201 rates. We denote $\{s_{1,k}\}$ by $\{s_k\}$ for convenience, meaning that

$$202 \quad s_1 = 1, \quad s_{k+1} = s_k + \frac{1}{s_k}$$

and provide the following estimate of its asymptotic behavior.

LEMMA 3.2. *For any $k \geq 2$ we have*

$$\sqrt{2k} \leq s_k \leq \sqrt{2k + \frac{1}{2} \log(k-1)}.$$

Proof. To prove the left inequality, since $s_{i+1}^2 = s_i^2 + \frac{1}{s_i^2} + 2 \geq s_i^2 + 2$, we have by induction that $s_k^2 \geq s_2^2 + 2(k-2) = 2k$ (using $s_2 = 2$). To prove the right inequality, we use (b) in Proposition 3.1 to obtain

$$s_k^2 = 1 + 2(k-1) + \sum_{i=1}^{k-1} \frac{1}{s_i^2} \leq 2k - 1 + \left(1 + \frac{1}{2} \sum_{i=2}^{k-1} \frac{1}{i}\right) \leq 2k + \frac{1}{2} \log(k-1),$$

where the first inequality follows from $s_k^2 \geq 2k$, and the second from the upper bound on harmonic numbers $\sum_{i=1}^k \frac{1}{i} \leq \log(k) + 1$. $\square$

It is worth noting that $\{s_k\}$ is employed for tuning the parameters of a primal-dual subgradient method in [15].

3.2. Convergence rate for the last iterate. We now turn to proving a convergence rate for the last iterate of the subgradient method with constant step sizes. With the choice of constant step size explained above $h_k = \frac{R}{B}h$ for some positive parameter h , we obtain Algorithm 3.1.

Algorithm 3.1 Projected subgradient method with constant step sizes

Parameters: number of iterations N , normalized step size parameter $h > 0$

Inputs: convex set X , convex function f defined on X with B -bounded subgradients, initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* .

For $k = 1, 2, \dots, N$ perform the following steps:

1. Select a subgradient $g^k \in \partial f(x^k)$.
2. Compute $x^{k+1} = P_X(x^k - h \frac{R}{B} g^k)$.

Output: last iterate x_{N+1}

Most of the effort in obtaining our convergence rate will be spent in obtaining the following lemma.

LEMMA 3.3. *Let f be a convex function with B -bounded subgradients on a convex set X , and let $\alpha \geq 1$. Let $\hat{x} \in X$ be a reference point. Consider N iterations of Algorithm 3.1 with step size parameter $h > 0$, starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - \hat{x}\| \leq R$. We have that the last iterate x_{N+1} satisfies*

$$(3.2) \quad f(x^{N+1}) - f(\hat{x}) \leq BR \left(\frac{1}{2} \left(s_{\alpha, N+1} \sqrt{h} - \frac{1}{s_{\alpha, N+1} \sqrt{h}} \right)^2 + 1 - Nh \right).$$

Proof. To prove inequality (3.2), we employ Lemma 2.1. Assume that

$$v_k = \frac{1}{s_{\alpha, N+1-k}}, \quad k \in \{0, 1, \dots, N\},$$

and $v_{N+1} = s_{\alpha,1}$. It is seen that $0 < v_0 \leq v_1 \leq \dots \leq v_{N+1}$. Suppose that $h_{N+1} = \frac{hR}{B}$.
By using Proposition 3.1 (a), one can verify that for $k \in \{1, \dots, N\}$,

$$\begin{aligned} v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} v_i &= \frac{1}{s_{\alpha,N+1-k}^2} - \left(\frac{1}{s_{\alpha,N+1-k}} - \frac{1}{s_{\alpha,N+2-k}} \right) \left(\sum_{i=1}^{N+1-k} \frac{1}{s_{\alpha,i}} + s_{\alpha,1} \right) \\ &= \frac{1}{s_{\alpha,N+1-k}^2} - \left(\frac{s_{\alpha,N+2-k}}{s_{\alpha,N+1-k}} - 1 \right) = 0. \end{aligned}$$

Furthermore,

$$\begin{aligned} v_{N+1}^2 - v_{N+1}(v_{N+1} - v_N) &= s_{\alpha,1} \left(\frac{1}{s_{\alpha,1}} \right) = 1, \\ v_0 \sum_{k=1}^{N+1} v_k &= \frac{1}{s_{\alpha,N+1}} \left(\sum_{k=1}^N \frac{1}{s_{\alpha,1}} + s_{\alpha,1} \right) = 1. \end{aligned}$$

Hence, by Lemma 2.1, we obtain

$$\begin{aligned} f(x^{N+1}) - f(\hat{x}) &\leq \frac{Rh}{2B} \sum_{i=1}^{N+1} v_k^2 \|g^i\|^2 + \frac{Bv_0^2}{2Rh} \|x^1 - x^*\|^2 \\ &\leq \frac{BRh}{2} \sum_{i=1}^N \frac{1}{s_{\alpha,N+1-k}^2} + \frac{BRh\alpha^2}{2} + \frac{BR}{2hs_{\alpha,N+1}^2} \\ &= BR \left(\frac{1}{2} \left(s_{\alpha,N+1} \sqrt{h} - \frac{1}{s_{\alpha,N+1} \sqrt{h}} \right)^2 + 1 - Nh \right), \end{aligned}$$

where the last equality follows from Proposition 3.1. Hence, we derive the desired inequality and the proof is complete. $\square$

The following theorem now uses Lemma 3.3 with the choice $\hat{x} = x^*$ to obtain a last-iterate convergence rate for the subgradient method with constant step sizes.

THEOREM 3.4. *Let f be a convex function with B -bounded subgradients on a convex set X . Consider N iterations of Algorithm 3.1 with step size parameter $h > 0$, starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* . We have that the last iterate x_{N+1} satisfies*

$$f(x^{N+1}) - f^* \leq \begin{cases} BR(1 - Nh) & \text{when } h \leq \frac{1}{s_{N+1}^2}, \\ BR \left(\left(\frac{1}{2} s_{N+1}^2 - N \right) h + \frac{1}{2s_{N+1}^2 h} \right) & \text{when } h > \frac{1}{s_{N+1}^2}. \end{cases}$$

Proof. We prove the theorem by plugging a suitable value for α into inequality (3.2) written for the choice $\hat{x} = x^*$, since it holds for any $\alpha \geq 1$. In the first case, when $h \leq \frac{1}{s_{N+1}^2}$, we may select by Proposition 3.1 a value of $\alpha \geq 1$ such that

$$s_{\alpha,N+1} \sqrt{h} - \frac{1}{s_{\alpha,N+1} \sqrt{h}} = 0$$

which leads to the desired inequality. In the second case where $h > \frac{1}{s_{N+1}^2}$ one can choose $\alpha = 1$ to obtain the inequality. $\square$

It is interesting to compare this last-iterate convergence rate to the one holding for the best iterate. Plugging $h_k = \frac{R}{B}h$ into the rate (1.2), we obtain that

$$\min_{1 \leq k \leq N+1} f(x^k) - f(x^*) \leq BR \left(\frac{1}{2} \frac{1}{(N+1)h} + \frac{1}{2}h \right).$$

Using the bounds $2N + 2 \leq s_{N+1}^2 \leq 2N + 2 + \frac{1}{2} \log(N)$ from Lemma 3.2, we rewrite the rate for larger steps from Theorem 3.4 in the following slightly weaker but easier to interpret form:

$$f(x^{N+1}) - f(x^*) \leq BR \left(\left(1 + \frac{1}{4} \log(N)\right)h + \frac{1}{4(N+1)h} \right).$$

Finally, when using the constant step size $h = \frac{1}{\sqrt{N+1}}$, which gives the optimal rate $\frac{BR}{\sqrt{N+1}}$ for both the best and the average iterate, we obtain instead for the last iterate

$$f(x^{N+1}) - f(x^*) \leq \frac{BR}{\sqrt{N+1}} \left(\frac{5}{4} + \frac{1}{4} \log(N) \right),$$

which shows a logarithmic loss.

A last interesting remark is that Lemma 3.3 can be used with a reference $\hat{x}$ that is not a minimizer. In essence, it shows that subgradient methods converge to any sublevel set of the objective function with the same worst-case rate, provided the constant R in its numerator is taken as the distance from the initial iterate to that sublevel set.

3.3. Tightness of the convergence rate. A convergence rate is said to be exact (or tight) if there exists a problem instance achieving that rate. We now show that the convergence rates we obtained for the subgradient method with constant step sizes are exact.

In the case of shorter steps ($h \leq \frac{1}{s_{N+1}^2}$), it is readily seen that the convergence rate in Theorem 3.4 is exact. Indeed, it is enough to consider an unconstrained optimization univariate problem ($n = 1$, $X = \mathbb{R}$) with objective function $f(x) = B|x|$ and the initial point $x^1 = R$.

In the case of longer steps ($h > \frac{1}{s_{N+1}^2}$), the following theorem illustrates that the convergence rate in Theorem 3.4 is also exact. To establish exactness, we may assume without loss of generality that $R = B = 1$. Before we get to the theorem, we need to present a lemma.

LEMMA 3.5. *Let $\phi_1 : \mathbb{R} \rightarrow \mathbb{R}$ given by $\phi_1(x) = |x|$, and let $\phi_k : \mathbb{R}^k \rightarrow \mathbb{R}$ be defined as follows,*

$$(3.3) \quad \phi_k(x) = \max\{x_1, \sqrt{1 - \frac{1}{s_k^2}} \phi_{k-1}(\bar{x}_1) - \frac{1}{s_k^2} x_1\},$$

where $\bar{x}_1 = (x_2, \dots, x_k)^T$. Then, ϕ_k may be written as a piece-wise linear function $\phi_k(x) = \max_{1 \leq j \leq k+1} \langle g^j, x \rangle$, such that $g^1 = e^1$,

$$(3.4) \quad \|g^j\| = 1 \quad j \in \{1, \dots, k+1\}, \quad \langle g^i, g^j \rangle = \langle g^{\min(i,j)}, g^{\min(i,j)+1} \rangle \quad i \neq j,$$

and

$$(3.5) \quad \langle g^j, g^{j+1} \rangle \geq \langle g^{j+1}, g^{j+2} \rangle \quad j \in \{1, \dots, k+1\}.$$

Furthermore, $\sum_{j=1}^k \langle g^j, g^{j+1} \rangle = k - \frac{1}{2} s_{k+1}^2$ and

$$(3.6) \quad \sum_{j=1}^i \langle g^j, g^{j+1} \rangle \leq 0, \quad i \in \{1, \dots, k\}.$$

Proof. We establish the lemma by induction. By the definition of f_1 , it is readily seen that whole statements hold for $k = 1$. Suppose that $k \geq 2$ and $f_{k-1}(x) = \max_{1 \leq j \leq k} \langle g^j, x \rangle$. By the definition, we get $\phi_k(x) = \max_{1 \leq j \leq k+1} \langle \hat{g}^j, x \rangle$, where

$$\hat{g}^1 = e^1, \quad \hat{g}^{j+1} = \sqrt{1 - \frac{1}{s_k^4}} \begin{pmatrix} 0 \\ g^j \end{pmatrix} - \frac{1}{s_k^2} e^1, \quad j \in \{1, \dots, k\}.$$

One can check that we have (3.4) for $\hat{g}^1, \dots, \hat{g}^{k+1}$. Furthermore, it is seen that (3.5) holds for $j \in \{2, \dots, k+1\}$ due to the definition of $\hat{g}^j$. For $j = 1$, we get

$$\begin{aligned} \langle \hat{g}^1, \hat{g}^2 \rangle - \langle \hat{g}^2, \hat{g}^3 \rangle &= -\frac{1}{s_k^2} + \frac{1}{s_{k-1}^2} \left(1 - \frac{1}{s_k^4}\right) - \frac{1}{s_k^4} = -1 - \frac{1}{s_k^2} + \left(1 + \frac{1}{s_{k-1}^2}\right) \left(1 - \frac{1}{s_k^4}\right) \\ &= \left(1 + \frac{1}{s_k^2}\right) \left(\left(1 - \frac{1}{s_k^2}\right) \left(1 + \frac{1}{s_{k-1}^2}\right) - 1\right) = \left(1 + \frac{1}{s_k^2}\right) \frac{s_k^2 - s_{k-1}^2 - 1}{s_{k-1}^2 s_k^2} \geq 0, \end{aligned}$$

where the last inequality follows from $s_k^2 = s_{k-1}^2 + \frac{1}{s_{k-1}^2} + 2$.

Now, we prove the second part. We have

$$\begin{aligned} \sum_{j=1}^k \langle \hat{g}^j, \hat{g}^{j+1} \rangle &= -\frac{1}{s_k^2} + \left(1 - \frac{1}{s_k^4}\right) \sum_{j=1}^{k-1} \langle g^j, g^{j+1} \rangle + \frac{k-1}{s_k^4} = \\ &= \left(1 - \frac{1}{s_k^4}\right) \left(k - 1 - \frac{1}{2} s_k^2\right) + \frac{k-1}{s_k^4} - \frac{1}{s_k^2} = k - \frac{1}{2} s_{k+1}^2, \end{aligned}$$

where the last equality follows from $s_{k+1}^2 = s_k^2 + \frac{1}{s_k^2} + 2$. It is readily seen that (3.6) holds when $i = 1$. For the case $i \geq 2$, the inequality follows from Proposition 3.1, (b), as $s_k^2 \geq 2k$ for $k \geq 2$, which completes the proof. $\square$

THEOREM 3.6. *Let N be a number of iterations and $h > \frac{1}{s_{N+1}^2}$. Then there exists a convex function $f : \mathbb{R}^{N+1} \rightarrow \mathbb{R}$ and $x^1 \in \mathbb{R}^{N+1}$ such that*

- i) The function f is Lipschitz continuous with modulus one;*
- ii) $0 \in \operatorname{argmin} f(x)$;*
- iii) Algorithm 1.1 with initial point x^1 with $\|x^1\| = 1$ and step size h may generate x^{N+1} with*

$$f(x^{N+1}) = \left(\frac{1}{2} s_{N+1}^2 - N\right)h + \frac{1}{2s_{N+1}^2 h}.$$

Proof. We construct a piece-wise linear function that satisfies all given conditions. Let $f : \mathbb{R}^{N+1} \rightarrow \mathbb{R}$ given by

$$f(x) = \max \left(0, \sqrt{1 - \frac{1}{h^2 s_k^4}} \phi_N(\bar{x}_1) + \frac{1}{h s_k^2} x_1 \right),$$

where ϕ_N given in Lemma 3.5. By Lemma 3.5, f can be expressed in the form $f(x) = \max_{0 \leq k \leq N+1} \langle \hat{g}^k, x \rangle$, where $\hat{g}^0 = 0$ and

$$\hat{g}^{j+1} = \sqrt{1 - \frac{1}{h^2 s_k^4}} \begin{pmatrix} 0 \\ g^j \end{pmatrix} + \frac{1}{h s_k^2} e^1, \quad j \in \{1, \dots, N+1\}.$$

Furthermore, we have $\|\hat{g}^i\| = 1$,

$$(3.7) \quad \langle \hat{g}^i, \hat{g}^j \rangle = \langle \hat{g}^{\min(i,j)}, \hat{g}^{\min(i,j)+1} \rangle = \left(1 - \frac{1}{h^2 s_k^4}\right) \langle g^{\min(i,j)}, g^{\min(i,j)+1} \rangle + \frac{1}{h^2 s_k^4}.$$

and

$$(3.8) \quad \langle \hat{g}^j, \hat{g}^{j+1} \rangle \geq \langle \hat{g}^{j+1}, \hat{g}^{j+2} \rangle \quad j \in \{1, \dots, N+1\}.$$

for $i, j \in \{1, \dots, N+1\}$ with $i \neq j$. As $\|\hat{g}^k\| = 1$ for $k \in \{1, \dots, N+1\}$, the function f is Lipschitz continuous with modulus one. We have $f(0) = 0$. Hence, $0 \in \partial f(0)$, which implies that $0 \in \operatorname{argmin} f(x)$. Suppose that $x^1 = e^1$. In what follows, we establish that

$$(3.9) \quad \hat{g}^k \in \partial f \left(x^1 - h \sum_{i=1}^{k-1} \hat{g}^i \right), \quad k \in \{1, \dots, N+1\}.$$

To this end, we just need to show that for $k \in \{1, \dots, N+1\}$

$$\langle \hat{g}^i, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle \leq \langle \hat{g}^k, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle, \quad i \in \{0, \dots, N+1\}.$$

Assume that $i > k$. By (3.7), we have

$$\begin{aligned} \langle \hat{g}^i, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle &= \frac{1}{hs_k^2} - h \sum_{j=1}^{k-1} \langle \hat{g}^{\min(i,j)}, \hat{g}^{\min(i,j)+1} \rangle = \\ &= \frac{1}{hs_k^2} - h \sum_{j=1}^{k-1} \langle \hat{g}^{\min(k,j)}, \hat{g}^{\min(k,j)+1} \rangle = \langle \hat{g}^k, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle. \end{aligned}$$

For $i < k$, we have

$$\begin{aligned} \langle \hat{g}^i, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle &= \frac{1}{hs_k^2} - h \sum_{j=1}^{i-1} \langle \hat{g}^{\min(i,j)}, \hat{g}^{\min(i,j)+1} \rangle - h - h(k-i-1) \langle \hat{g}^i, \hat{g}^{i+1} \rangle \leq \\ &= \frac{1}{hs_k^2} - h \sum_{j=1}^{k-1} \langle \hat{g}^{\min(k,j)}, \hat{g}^{\min(k,j)+1} \rangle = \langle \hat{g}^k, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle, \end{aligned}$$

where the last inequality follows from (3.8) and $|\langle \hat{g}^{\min(k,i)}, \hat{g}^{\min(k,i)+1} \rangle| \leq 1$. Moreover, it follows from (3.6),

$$\langle \hat{g}^k, x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \rangle \geq 0.$$

The above equations imply (3.9). By (3.9), it is seen that Algorithm 1.1 with initial point x^1 and step size h may generate the following points

$$x^k = x^1 - h \sum_{j=1}^{k-1} \hat{g}^j \quad k \in \{2, \dots, N+1\}.$$

Moreover, we get

$$\begin{aligned} f(x^{N+1}) &= \langle \hat{g}^{N+1}, x^{N+1} \rangle = \langle \hat{g}^{N+1}, x^1 - h \sum_{j=1}^N \hat{g}^j \rangle = \frac{1}{hs_{N+1}^2} - \frac{N}{hs_{N+1}^4} - h(1 - \frac{1}{h^2 s_{N+1}^4}) \\ &= \sum_{j=1}^N \langle \hat{g}^{N+1}, \hat{g}^j \rangle = (\frac{1}{2} s_{N+1}^2 - N)h + \frac{1}{2s_{N+1}^4 h}, \end{aligned}$$

where the last equality follows from Lemma 3.5, and the proof is complete. $\square$

3.4. Optimal constant step size. In what follows, we determine the optimal constant step size for Algorithm 1.1, i.e. the value of h that minimizes the rate established in Theorem 3.4, and derive the resulting optimal last-iterate convergence rate for this class of subgradient methods.

THEOREM 3.7. *Let f be a convex function with B -bounded subgradients on a convex set X . Consider N iterations of Algorithm 3.1 starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* . The optimal value of the step size parameter h is given by*

$$h^* = \frac{1}{s_{N+1} \sqrt{s_{N+1}^2 - 2N}},$$

and with that choice the last iterate x_{N+1} satisfies

$$f(x^{N+1}) - f(x^*) \leq BR \sqrt{1 - \frac{2N}{s_{N+1}^2}}.$$

Proof. To get the optimal step size suffices to minimize the function $H : \mathbb{R}_+ \rightarrow \mathbb{R}$ given by

$$H(h) = \begin{cases} 1 - Nh & h \in [0, \frac{1}{s_{N+1}^2}) \\ (\frac{1}{2}s_{N+1}^2 - N)h + \frac{1}{2s_{N+1}^2 h} & h \in [\frac{1}{s_{N+1}^2}, \infty). \end{cases}$$

It is readily seen that H is a differentiable convex function, decreasing on its first piece, and the above minimizer h^* can be found by solving $H'(h^*) = 0$ on the second piece.

Plugging this optimal h^* in the rate Theorem 3.4 completes the proof. $\square$

A simpler, slightly weaker bound is obtained using $1 - \frac{2N}{s_{N+1}^2} = \frac{s_{N+1}^2 - 2N}{s_{N+1}^2} \leq \frac{2 + \frac{1}{2} \log(N)}{2N+2}$, leading to

$$f(x^{N+1}) - f(x^*) \leq \frac{BR}{\sqrt{N+1}} \sqrt{1 + \frac{1}{4} \log(N)} = O\left(\sqrt{\frac{\log N}{N}}\right)$$

which emphasizes the logarithmic loss compared to the best-iterate convergence rate.

4. Subgradient method with constant step lengths. To compute steps in Algorithm 3.1 requires identifying a bound B on the maximum norm of any subgradient, which may be difficult. Alternatively, one can adapt the subgradient method to remove the need to know such a bound, leading to an algorithm using constant step lengths (also sometimes said to use normalized steps or directions [14]). We still need to express this constant step length as a fraction t of the initial distance R (removing the requirement to know R appears to be much more difficult). Since the length of a step is equal to $\|h_k g_k\|$, this implies the choice of step sizes $h_k = \frac{tR}{\|g_k\|}$ for each k . Algorithm 4.1 below presents the subgradient method with constant step length. Note that when $g^k = 0$ we set $\frac{tR}{\|g^k\|} g^k = 0$ in Algorithm 4.1.

We give below a convergence rate for Algorithm 4.1 by using Lemma 2.1.

THEOREM 4.1. *Let f be a convex function with B -bounded subgradients on a convex set X . Consider N iterations of Algorithm 4.1 with step length parameter $t > 0$, starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* . We have that the last iterate x_{N+1} satisfies the following rate*

Algorithm 4.1 Projected subgradient method with constant step lengths**Parameters:** number of iterations N , step length parameter $t > 0$ **Inputs:** convex set X , convex function f defined on X with B -bounded subgradients, initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* .For $k = 1, 2, \dots, N$ perform the following steps:

1. Select a subgradient $g^k \in \partial f(x^k)$.
2. Compute $x^{k+1} = P_X \left(x^k - t \frac{R}{\|g^k\|} g^k \right)$.

Output: last iterate x_{N+1}

348 *i) If $t \in (0, \frac{1}{s_{N+1}^2}]$, then*

349
$$f(x^{N+1}) - f(x^*) \leq BR(1 - Nt).$$

350 *ii) If $t \in [\frac{1}{s_{N+1}^2}, \infty)$, then*

351
$$f(x^{N+1}) - f(x^*) \leq BR \left(\left(\frac{1}{2} s_{N+1}^2 - N \right) t + \frac{1}{2 s_{N+1}^2 t} \right).$$

352 *Proof.* We employ Lemma 2.1 to establish the theorem. Let $h_{N+1} = \frac{tR}{B}$ and
 353 $v_{N+1} = \alpha$ for some $\alpha \geq 1$. Suppose that $h_k = \frac{tR}{\|g^k\|}$, $k \in \{1, \dots, N\}$. We define v^k
 354 recursively as follows,

355 (4.1)
$$v_k = \frac{h_{N+1}}{\sum_{i=k+1}^{N+1} h_i v_i}, \quad k \in \{N, \dots, 1, 0\}.$$

356 It is seen that $0 < v_0 \leq v_1 \leq \dots \leq v_N \leq v_{N+1}$. For $k \in \{1, \dots, N\}$, we have

357
$$h_k v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} h_i v_i = v_{k-1} \sum_{i=k}^{N+1} h_i v_i - v_k \sum_{i=k+1}^{N+1} h_i v_i = 0.$$

Furthermore,

$$h_{N+1} v_{N+1}^2 - (v_{N+1} - v_N) h_{N+1} v_{N+1} = h_{N+1}, \quad v_0 \sum_{k=1}^{N+1} h_k v_k = h_{N+1}.$$

358 On the other hand, by (4.1), we get $v_N = \frac{1}{\alpha}$ and

359
$$\frac{1}{v_k} = \frac{h_{k+1}}{h_{N+1}} v_{k+1} + \frac{1}{v_{k+1}} \geq v_{k+1} + \frac{1}{v_{k+1}}, \quad k \in \{0, \dots, N-1\},$$

360 where the last inequality follows from $\|g^{k+1}\| \leq B$. One can infer by induction that

361
$$v_k \leq \frac{1}{s_{\alpha, N+1-k}}, \quad k \in \{0, \dots, N\}.$$

362 Hence, by using Lemma 2.1, we obtain

363
$$f(x^{N+1}) - f(x^*) \leq \frac{v_0^2}{2h_{N+1}} \|x^1 - x^*\|^2 + \frac{t^2 R^2}{2h_{N+1}} \sum_{k=1}^N v_k^2 + \frac{h_{N+1} B^2 v_{N+1}^2}{2}$$

 364
$$\leq BR \left(\frac{1}{2} \left(s_{\alpha, N+1} \sqrt{t} - \frac{1}{s_{\alpha, N+1} \sqrt{t}} \right)^2 + 1 - Nh \right),$$

where the last inequality resulted from Proposition 3.1. The rest of the proof is analogous to Theorem 3.4. $\square$

5. Optimal subgradient methods.

5.1. Lower bound on last-iterate convergence rate. The convergence rate of any black-box method that relies on subgradients cannot be arbitrarily small. More precisely, for any method that moves at each iteration in a direction belonging to the span of the current and past subgradients, it is known that the accuracy of last iterate must obey a lower bound of the order $\Omega(\frac{1}{\sqrt{N}})$. Nesterov [14, Theorem 3.2.1] proposes a Lipschitz continuous function f with modulus $B > 0$, for which any subgradient method that calls the first-order oracle N times satisfies

$$f(x^{N+1}) - f(x^*) \geq \frac{BR}{2(2 + \sqrt{N+1})},$$

where $\|x^1 - x^*\| \leq R$. Drori and Teboulle [7] improved the above-mentioned lower bound and obtain the following

$$(5.1) \quad f(x^{N+1}) - f(x^*) \geq \frac{BR}{\sqrt{N+1}}.$$

Note that this exact optimal convergence rate has been established for variants of the subgradient method with momentum terms, see [5, Corollary 3] and [22, 23].

Since this lower bound plays a central role in this work, we provide a short self-contained proof. Given a number of iterations N , consider function

$$f(x) = \max\{0, x_1, x_2, \dots, x_{N+1}\} = \left[\max_{1 \leq k \leq N+1} x_k \right]_+$$

with $B = 1$ and $x_* = 0$, and pick starting point $x^1 = (1, 1, \dots, 1)$ with $R = \sqrt{N+1}$. As long as $f(x^k) > 0$, subgradient $g^k \in \partial f(x^k)$ can be chosen as a basis vector e^i for some $1 \leq i \leq N+1$ (and $\|g^k\| = B$). Define the following induction hypothesis, which is easy to check for $k = 1$,

$$(H_k) \quad x^k \text{ contains at least } N+2-k \text{ components equal to 1.}$$

Assume (H_k) holds for some $k \leq N$. Then $f(x^k) \geq 1$, so subgradient g^k can be chosen as some basis vector e^i , and x^{k+1} can differ only by at most one component from x^k , implying (H_{k+1}) holds. So we have proved that $f(x^k) \geq 1 = \frac{BR}{\sqrt{N+1}}$ for all $1 \leq k \leq N+1$.

A subgradient method for which the last-iterate accuracy would match this lower bound, or at least its order $\Omega(\frac{1}{\sqrt{N}})$ is certainly desirable [19]. Recently, Jain et al. [11] introduced such a subgradient method when the feasible set X is bounded. Indeed, they derive the following convergence rate for their proposed algorithm [11, Theorem 2.6]

$$f(x^{N+1}) - f(x^*) \leq \frac{15BD}{\sqrt{N+1}},$$

where $D = \max_{x, y \in X} \|x - y\|$.

5.2. Optimal subgradient method. In this section, we introduce Algorithm 5.1, a subgradient method based on a new sequence of step sizes for which the last-iterate convergence rate matches the lower bound (5.1).

We establish the optimal rate of convergence for Algorithm 5.1 using Lemma 2.1.

Algorithm 5.1 Optimal projected subgradient method**Parameters:** number of iterations N **Inputs:** convex set X , convex function f defined on X with B -bounded subgradients, initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* .For $k = 1, 2, \dots, N$ perform the following steps:

1. Select a subgradient $g^k \in \partial f(x^k)$.
2. Compute $x^{k+1} = P_X(x^k - h_k g^k)$ using step size $h_k = \frac{R}{B} \frac{(N+1-k)}{(N+1)^{3/2}}$.

Output: last iterate x_{N+1}

THEOREM 5.1. *Let f be a convex function with B -bounded subgradients on a convex set X . Consider N iterations of Algorithm 5.1 starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* . We have that the last iterate x_{N+1} satisfies*

$$(5.2) \quad f(x^{N+1}) - f(x^*) \leq \frac{BR}{\sqrt{N+1}}.$$

Proof. Suppose that v_k 's are given as follows,

$$v_k = \frac{(N+1)^{\frac{3}{4}}}{N+1-k} \sqrt{\frac{B}{R}}, \quad k \in \{0, \dots, N\},$$

and $v_{N+1} = v_N$. It is seen that $0 < v_0 \leq v_1 \leq \dots \leq v_N \leq v_{N+1}$. In addition, let $h_{N+1} = \frac{R}{B} \frac{1}{(N+1)^{3/2}}$. For $k \in \{1, \dots, N\}$, we have

$$h_k v_k^2 - (v_k - v_{k-1}) \sum_{i=k}^{N+1} h_i v_i = \frac{1}{N+1-k} - (N+2-k) \left(\frac{1}{N+1-k} - \frac{1}{N+2-k} \right) = 0.$$

In addition,

$$v_0 \sum_{k=1}^{N+1} h_k v_k = \frac{1}{N+1} \sum_{k=1}^{N+1} 1 = 1, \quad h_{N+1} v_{N+1}^2 = 1.$$

By Lemma 2.1, we get

$$\begin{aligned} f(x^{N+1}) - f(x^*) &\leq \frac{R}{2B(N+1)^{3/2}} \sum_{i=1}^{N+1} \|g^i\|^2 + \frac{B}{2R\sqrt{N+1}} \|x^1 - x^*\|^2 \\ &\leq \frac{R}{2B(N+1)^{3/2}} \sum_{i=1}^{N+1} B^2 + \frac{B}{2R\sqrt{N+1}} R^2 = \frac{BR}{\sqrt{N+1}}, \end{aligned}$$

and the proof is complete. $\square$

5.3. Absence of optimal subgradient method with universal sequence of step sizes. It is seen that the step sizes in Algorithm 5.1 are dependent on the number of iterations, N . As conjectured in [11], it is not possible to introduce a universal infinite sequence of step sizes $\{h_k\}_{k \in \mathbb{N}}$ for which (5.2) would hold for all iterations, i.e. simultaneously for all N (this is sometimes called *anytime* convergence). Before we show this point, we need to present a lemma.

LEMMA 5.2. Consider Algorithm 1.1 with $h_1 = \frac{1}{2\sqrt{2}}$ and $N = 2$.

i) If $h_2 \in (0, \frac{1}{8\sqrt{2}}]$, then there exists $f \in \mathcal{F}(\mathbb{R}^2)$ with 1-bounded subgradients and x^1 and such that

$$f(x^3) - f^* = \frac{1}{\sqrt{2}} - h_2,$$

where $\|x^1 - x^*\| \leq 1$.

ii) If $h_2 \in (\frac{1}{8\sqrt{2}}, \infty)$, then there exists $f \in \mathcal{F}(\mathbb{R}^3)$ with 1-bounded subgradients and x^1 and such that

$$f(x^3) - f^* = h_2 + \frac{1}{64h_2} + \frac{16h_2}{(1+8\sqrt{2}h_2)^2},$$

where $\|x^1 - x^*\| \leq 1$.

Proof. First we establish i). Let $f \in \mathcal{F}(\mathbb{R}^2)$ be given by

$$f(x) = \max \{x_1 - 1, x_2 - 1, -1\}.$$

It is readily seen that $x^* = 0$ is an optimal solution to problem $\min f(x)$. Algorithm 1.1 with initial point $x^1 = \frac{1}{\sqrt{2}}(1, 1)^T$ may generate the following points

$$x^2 = x^1 - \frac{1}{2\sqrt{2}}e_1, \quad x^3 = x^1 - \frac{1}{2\sqrt{2}}e_1 - h_2e_2.$$

In addition, $f(x^3) - f^* = \frac{1}{\sqrt{2}} - h_2$ and we introduce an optimization problem with the desired properties.

Now, we prove ii). Let $\gamma = \frac{32h_2}{(1+8\sqrt{2}h_2)^2}$. One can show that $\gamma \in [0, 1]$. Suppose that

$$\xi^1 = \begin{pmatrix} -\sqrt{\frac{\gamma}{1-\gamma^2}} \\ 0 \end{pmatrix}, \xi^2 = \begin{pmatrix} \frac{\gamma}{\sqrt{1-\gamma^2}} \\ -\sqrt{(1-\gamma^2)(1-\frac{1}{128h_2^2})} \end{pmatrix}, \xi^3 = \begin{pmatrix} \frac{\gamma}{\sqrt{1-\gamma^2}} \\ \sqrt{(1-\gamma^2)(1-\frac{1}{128h_2^2})} \end{pmatrix}.$$

It is readily seen that $\|\xi^k\| = 1$, $i \in \{1, 2, 3\}$. Let $z^1 = e_1$ and

$$z^2 = e_1 - \frac{1}{2\sqrt{2}}\xi^1, \quad z^3 = e_1 - \frac{1}{2\sqrt{2}}\xi^1 - h_2\xi^2.$$

Consider the following linear functions

$$\alpha_1(x) = \gamma + \langle \xi^1, x - z^1 \rangle, \quad \alpha_2(x) = \gamma + \frac{1-\gamma^2}{32h_2} - \frac{\gamma^2}{2\sqrt{2}} + \langle \xi^2, x - z^2 \rangle,$$

$$\alpha_3(x) = h_2 + \frac{1}{64h_2} + \frac{16h_2}{(1+8\sqrt{2}h_2)^2} + \langle \xi^3, x - z^3 \rangle.$$

We define $f : \mathbb{R}^3 \rightarrow \mathbb{R}$ by

$$f(x) = \max \{0, \alpha_1(x), \alpha_2(x), \alpha_3(x)\}.$$

One can show that $x^* = 0$ is an optimal solution to problem $\min f(x)$. By doing some algebra, one can check that Algorithm 1.1 with initial point $x^1 = e_1$ may generate the following points $x^2 = z^2$ and $x^3 = z^3$. It is seen that $f(x^3) - f(x^*) = h_2 + \frac{1}{64h_2} + \frac{16h_2}{(1+8\sqrt{2}h_2)^2}$, and the proof is complete. $\square$

Now, we present an argument why there does not exist a sequence $\{h_k\}$ that satisfies the convergence rate (5.2) for any arbitrary N . By contradiction, assume that there exists such a $\{h_k\}$. For the convenience suppose that $B = R = 1$. Due to

the exactness of rates given in Theorem 3.4, we have $h_1 = \frac{1}{2\sqrt{2}}$. By Lemma 5.2, one can infer that a convergence rate of Algorithm 1.1 with $h_1 = \frac{1}{2\sqrt{2}}$ and $N = 2$ cannot be lower than $o = 0.5785$. Note that o is computed by solving $\min_{h \geq 0} H(h)$, where H is given by

$$H(h) = \begin{cases} \frac{1}{\sqrt{2}} - h & h \in [0, \frac{1}{8\sqrt{2}}] \\ h + \frac{1}{64h} + \frac{16h}{(1+8\sqrt{2}h)^2} & h \in (\frac{1}{8\sqrt{2}}, \infty) \end{cases}.$$

On the other hand, $o > 0.5775 > \frac{1}{\sqrt{3}}$. Hence, it is not possible to have a universal sequence $\{h_k\}$ for which (5.2) holds for any N in the setting of Algorithm 1.1, and actually not even for the first two steps.

In the setting of Algorithm 1.1, step sizes guaranteeing an optimal last-iterate convergence rate must depend on the total number of iterations. Note that Zhu et al. [26] recently proved that the universal sequence of step sizes $h_k = \frac{R}{B\sqrt{k}}$ achieves the convergence rate of $\frac{3}{2} \frac{BR}{\sqrt{N+1}}$ for the average of any first N iterates, and [5, Corollary 4] achieves the optimal $\frac{BR}{\sqrt{N+1}}$ rate for all iterates of a variant of the subgradient method relying on line-search. We conjecture that the incorporation of suitable momentum terms in the subgradient method may lead to a universal optimal algorithm whose last-iterate convergence rate in $O\left(\frac{BR}{\sqrt{N+1}}\right)$ would hold for all iterations.

5.4. Optimal projected subgradient method using step lengths. In the last part of this section, we present Algorithm 5.2, an optimal subgradient method based on step lengths.

Algorithm 5.2 Optimal projected subgradient method (step lengths)

Parameters: number of iterations N

Inputs: convex set X , convex function f defined on X with B -bounded subgradients, initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* .

For $k = 1, 2, \dots, N$ perform the following steps:

1. Select a subgradient $g^k \in \partial f(x^k)$.
2. Compute $x^{k+1} = P_X\left(x^k - t_k \frac{g^k}{\|g^k\|}\right)$ with $t_k = \frac{R(N+1-k)}{(N+1)^{3/2}}$.

Output: last iterate x_{N+1}

In the forthcoming theorem, we provide a convergence rate for Algorithm 5.2.

THEOREM 5.3. *Let f be a convex function with B -bounded subgradients on a convex set X . Consider N iterations of Algorithm 5.2 starting from an initial iterate $x^1 \in X$ satisfying $\|x^1 - x^*\| \leq R$ for some minimizer x^* . We have that the last iterate x_{N+1} satisfies*

$$(5.3) \quad f(x^{N+1}) - f(x^*) \leq \frac{BR}{\sqrt{N+1}}.$$

Proof. The proof is analogous to that of Theorem 4.1. Let

$$u_{N+1} = (N+1)^{\frac{3}{4}} \sqrt{\frac{B}{R}}, \quad h_{N+1} = \frac{R}{B(N+1)^{3/2}},$$

and let $h_k = \frac{t_k}{\|g^k\|}$, $k \in \{1, \dots, N\}$. Let us define u^k recursively in the following manner,

$$(5.4) \quad u_k = \frac{1}{\sum_{i=k+1}^{N+1} h_i u_i}, \quad k \in \{N, \dots, 1, 0\}.$$

It is readily seen that $0 < u_0 \leq u_1 \leq \dots \leq u_N \leq u_{N+1}$, and

$$h_k u_k^2 - (u_k - u_{k-1}) \sum_{i=k}^{N+1} h_i u_i = u_{k-1} \sum_{i=k}^{N+1} h_i u_i - u_k \sum_{i=k+1}^{N+1} h_i u_i = 0, \quad k \in \{1, \dots, N\}.$$

In addition,

$$h_{N+1} u_{N+1}^2 - (u_{N+1} - u_N) h_{N+1} u_{N+1} = 1, \quad u_0 \sum_{k=1}^{N+1} h_k u_k = 1.$$

Consider v_k given in the proof of Theorem 5.1. As

$$v_k = \frac{1}{\sum_{i=k+1}^{N+1} v_i h_i \frac{\|g^k\|}{B}}, \quad k \in \{0, 1, \dots, N\},$$

one can infer by induction that $u_k \leq v_k$, $k \in \{0, 1, \dots, N+1\}$. Thus, by Lemma 2.1, we get

$$\begin{aligned} f(x^{N+1}) - f(x^*) &\leq \frac{u_0^2}{2} \|x^1 - x^*\|^2 + \frac{1}{2} \sum_{k=1}^N h_k^2 u_k^2 \|g^k\|^2 + \frac{1}{2} h_{N+1}^2 u_{N+1}^2 B^2 \\ &= \frac{u_0^2}{2} \|x^1 - x^*\|^2 + \frac{1}{2} \sum_{k=1}^N \left(\frac{t_k}{B}\right)^2 u_k^2 B^2 + \frac{1}{2} h_{N+1}^2 u_{N+1}^2 B^2 \\ &\leq \frac{B}{2R\sqrt{N+1}} R^2 + \frac{R}{2B(N+1)^{3/2}} \sum_{i=1}^{N+1} B^2 = \frac{BR}{\sqrt{N+1}}, \end{aligned}$$

where the last inequality follows from $u_k \leq v_k$, $k \in \{0, 1, \dots, N+1\}$, and the proof is complete. $\square$

6. Conclusion. Before concluding, we briefly explain how most of the theorems in this paper were initially discovered. We used the performance estimation problem (PEP) methodology [6, 25], which allowed us to compute numerically the exact last-iterate convergence rate of subgradient methods applied to convex functions with bounded subgradients [24]. Assuming $B = R = 1$ without loss of generality, the values of the worst case accuracy for several choices of N and h were matched with explicit analytical expressions, which required some guesswork including the introduction of the sequence $\{s_k\}$. The next step was to use the numerical values of the dual multipliers to identify PEP-style proofs of those convergence rates, then to guess analytical expressions for those multipliers. Finally, we observed that large parts of the obtained proofs could be simplified by rewriting them as Jensen-type inequalities. After further simplifications, this ultimately leads to the proof technique that was exposed in Section 2.

To summarize, we have provided in this paper new convergence rates for the subgradient method with constant step sizes and constant step lengths, and proved their

tightness. Additionally, we have presented two optimal subgradient methods that attains the most favorable convergence rate achievable among subgradient algorithms. As avenues for future research, it would be valuable to investigate the convergence analysis of the (stochastic) proximal subgradient method with respect to the last iterate by employing some result analogous to Lemma 2.1. Moreover, deriving tighter convergence rates for the stochastic subgradient method may be of interest.

REFERENCES

- [1] S. BOYD, L. XIAO, AND A. MUTAPCIC, *Subgradient methods*, lecture notes of EE392o, Stanford University, Autumn Quarter, 2004 (2003), pp. 2004–2005.
- [2] D. DAVIS AND D. DRUSVYATSKIY, *Stochastic model-based minimization of weakly convex functions*, SIAM Journal on Optimization, 29 (2019), pp. 207–239.
- [3] A. DEFAZIO, A. CUTKOSKY, H. MEHTA, AND K. MISHCHENKO, *When, why and how much? adaptive learning rate scheduling by refinement*, arXiv preprint arXiv:2310.07831, (2023).
- [4] A. DEFAZIO, X. YANG, A. KHALED, K. MISHCHENKO, H. MEHTA, AND A. CUTKOSKY, *The road less scheduled*, Advances in Neural Information Processing Systems, 37 (2024), pp. 9974–10007.
- [5] Y. DRORI AND A. B. TAYLOR, *Efficient first-order methods for convex minimization: a constructive approach*, Mathematical Programming, 184 (2020), pp. 183–220.
- [6] Y. DRORI AND M. TEOULLE, *Performance of first-order methods for smooth convex minimization: a novel approach*, Mathematical Programming, 145 (2014), pp. 451–482.
- [7] Y. DRORI AND M. TEOULLE, *An optimal variant of Kelley’s cutting-plane method*, Mathematical Programming, 160 (2016), pp. 321–351.
- [8] B. GRIMMER AND D. LI, *Some primal-dual theory for subgradient methods for strongly convex optimization*, arXiv preprint arXiv:2305.17323, (2023).
- [9] N. J. HARVEY, C. LIAW, Y. PLAN, AND S. RANDHAWA, *Tight analyses for non-smooth stochastic gradient descent*, in Conference on Learning Theory, PMLR, 2019, pp. 1579–1613.
- [10] N. J. HARVEY, C. LIAW, AND S. RANDHAWA, *Tight analyses for subgradient descent I: Lower bounds*, Open Journal of Mathematical Optimization, 5 (2024), pp. 1–17.
- [11] P. JAIN, D. M. NAGARAJ, AND P. NETRAPALLI, *Making the last iterate of SGD information theoretically optimal*, SIAM Journal on Optimization, 31 (2021), pp. 1108–1130.
- [12] G. LAN, *First-order and stochastic optimization methods for machine learning*, vol. 1, Springer, 2020.
- [13] Z. LIU AND Z. ZHOU, *Revisiting the last-iterate convergence of stochastic gradient methods*, International Conference on Learning Representations, (2024).
- [14] Y. NESTEROV, *Introductory lectures on convex optimization: A basic course*, vol. 87, Springer Science & Business Media, 2003.
- [15] Y. NESTEROV, *Primal-dual subgradient methods for convex problems*, Mathematical programming, 120 (2009), pp. 221–259.
- [16] Y. NESTEROV AND V. SHIKHMAN, *Quasi-monotone subgradient methods for nonsmooth convex minimization*, Journal of Optimization Theory and Applications, 165 (2015), pp. 917–940.
- [17] Y. NESTEROV AND V. SHIKHMAN, *Algorithmic Principle of Least Revenue for Finding Market Equilibria*, Springer International Publishing, Cham, 2016, pp. 381–435.
- [18] F. ORABONA, *Last iterate of SGD converges even in unbounded domains*, 2020, <https://parameterfree.com/2020/08/07/last-iterate-of-sgd-converges-even-in-unbounded-domains/>. Accessed: 2024-06-21.
- [19] O. SHAMIR, *Open problem: Is averaging needed for strongly convex stochastic gradient descent?*, in Conference on Learning Theory, JMLR Workshop and Conference Proceedings, 2012, pp. 47–1.
- [20] N. Z. SHOR, *An application of the method of gradient descent to the solution of the network transportation problem*, Notes of Scientific Seminar on Theory and Applications of Cybernetics and Operations Research, Ukrainian Academy of Sciences, 1 (1962), pp. 9–17.
- [21] N. Z. SHOR, *Minimization Methods for Non-Differentiable Functions*, vol. 3 of Springer Series in Computational Mathematics, Springer-Verlag, Berlin, Heidelberg, 1985.
- [22] W. TAO, Z. PAN, G. WU, AND Q. TAO, *Primal averaging: A new gradient evaluation step to attain the optimal individual convergence*, IEEE transactions on cybernetics, 50 (2018), pp. 835–845.
- [23] W. TAO, Z. PAN, G. WU, AND Q. TAO, *The strength of Nesterov’s extrapolation in the individual convergence of nonsmooth optimization*, IEEE Transactions on Neural Networks and

- 527 Learning Systems, 31 (2019), pp. 2557–2568.
- 528 [24] A. B. TAYLOR, J. M. HENDRICKX, AND F. GLINEUR, *Exact worst-case performance of first-order*
529 *methods for composite convex optimization*, SIAM Journal on Optimization, 27 (2017),
530 pp. 1283–1313.
- 531 [25] A. B. TAYLOR, J. M. HENDRICKX, AND F. GLINEUR, *Smooth strongly convex interpolation*
532 *and exact worst-case performance of first-order methods*, Mathematical Programming, 161
533 (2017), pp. 307–345.
- 534 [26] Z. ZHU, Y. ZHANG, AND Y. XIA, *Convergence rate of projected subgradient method with time-*
535 *varying step-sizes*, Optimization Letters, (2024), pp. 1–5.