



Variable Smoothing for Weakly Convex Composite Functions

Axel Böhm¹ · Stephen J. Wright²

Received: 17 March 2020 / Accepted: 19 December 2020
© The Author(s) 2021

Abstract

We study minimization of a structured objective function, being the sum of a smooth function and a composition of a weakly convex function with a linear operator. Applications include image reconstruction problems with regularizers that introduce less bias than the standard convex regularizers. We develop a variable smoothing algorithm, based on the Moreau envelope with a decreasing sequence of smoothing parameters, and prove a complexity of $\mathcal{O}(\epsilon^{-3})$ to achieve an ϵ -approximate solution. This bound interpolates between the $\mathcal{O}(\epsilon^{-2})$ bound for the smooth case and the $\mathcal{O}(\epsilon^{-4})$ bound for the subgradient method. Our complexity bound is in line with other works that deal with structured nonsmoothness of weakly convex functions.

Keywords Variable smoothing · Weakly convex · Composite minimization

1 Introduction

We study minimization of the sum of a smooth function and a nonsmooth, weakly convex function composed with a linear operator. The case in which the nonsmooth regularizer is convex has been studied extensively; see [1,2]. Weakly convex functions (which can be expressed as the difference between a convex function and a quadratic) share some properties with convex functions but include many interesting nonconvex

Communicated By Aris Daniilidis and Stephen J. Wright.

✉ Axel Böhm
axel.boehm@univie.ac.at

Stephen J. Wright
swright@cs.wisc.edu

¹ Faculty of Mathematics, University of Vienna, Oskar-Morgenstern-Platz 1, 1090 Vienna, Austria

² Computer Sciences Department and Wisconsin Institute for Discovery, University of Wisconsin-Madison, 1210 W. Dayton St, Madison, WI 53706, USA

cases, as we discuss in Sect. 2.1. For example, any smooth function with a uniformly Lipschitz continuous gradient is a weakly convex function.

Our approach makes use of a smooth approximation of the weakly convex function known as the *Moreau envelope*, parametrized by a positive scalar μ . Since evaluation of the gradient of the Moreau envelope is obtained by applying a proximal operator to the function, our method is suitable for problems where this proximal operator can be evaluated at reasonable cost. Our method requires $\mathcal{O}(\epsilon^{-3})$ iterations to obtain an ϵ -approximate stationary point.

The remainder of the paper is organized as follows. Section 2 is concerned with other problem formulations related to ours and describes specific problems with weakly convex regularizers. In Sect. 3, we give the necessary mathematical preliminaries including a detailed discussion about the notion of stationarity we use. Section 4 describes our approach and its convergence properties. In Sect. 5, we highlight the difference between the variable smoothing technique and a simple proximal gradient approach, for the case in which the linear operator is not present in the weakly smooth term.

2 Problem Class and Algorithmic Approach

The problem we address in this paper has the form

$$\min_{x \in \mathbb{R}^d} F(x) := h(x) + g(Ax), \quad (1)$$

for a smooth function $h : \mathbb{R}^d \rightarrow \mathbb{R}$, a weakly convex function $g : \mathbb{R}^n \rightarrow \mathbb{R}$ (generally nonsmooth) and a matrix $A \in \mathbb{R}^{n \times d}$. For some $\rho \geq 0$, we say that

$$g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}} \text{ is } \rho\text{-weakly convex if } g + \frac{\rho}{2} \|\cdot\|^2 \text{ is convex.}$$

When g is a smooth function with a uniformly Lipschitz continuous gradient, with Lipschitz constant L , then g is weakly convex with $\rho = L$. Other interesting weakly convex functions are discussed in Sect. 2.1.

The *Moreau envelope* g_μ is a smooth approximation of g , parametrized by a positive scalar μ . The Moreau envelope and the closely related proximal operator are defined as follows.

Definition 2.1 For a proper, ρ -weakly convex and lower semicontinuous function $g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$, the Moreau envelope of g with the parameter $\mu \in]0, \rho^{-1}[$ is the function from $\mathbb{R}^n$ to $\mathbb{R}$ defined by

$$g_\mu(y) := \inf_{z \in \mathbb{R}^n} \left\{ g(z) + \frac{1}{2\mu} \|z - y\|^2 \right\}.$$

The proximal operator of the function μg is the arg min of the right-hand side in this definition, that is,

$$\text{prox}_{\mu g}(y) := \arg \min_{z \in \mathbb{R}^n} \left\{ g(z) + \frac{1}{2\mu} \|z - y\|^2 \right\} = \arg \min_{z \in \mathbb{R}^n} \left\{ \mu g(z) + \frac{1}{2} \|z - y\|^2 \right\}.$$

Note that $\text{prox}_{\mu g}(y)$ is defined *uniquely* by this formula, because the function being minimized is strongly convex. We describe in Lemma 3.1 the relationship between $\nabla g_\mu(y)$ and $\text{prox}_{\mu g}(y)$, which is key to our algorithm.

Steps of our algorithm have the form

$$x \leftarrow x - \gamma \nabla(h + g_\mu \circ A)(x),$$

for some step length γ . Accelerated versions of these approaches have been proposed for convex problems in [3–5]. The use of acceleration makes the analysis more complicated than for the gradient case; see [6,7].

2.1 Composite Problems

We discuss several instances of problems of the form (1).

Regularization with $\|\cdot\|_1$ (LASSO). Functions that are “sharp” around zero have a long history as regularizers that induce sparsity in the solution vector x . Foremost among such functions is the vector norm $\|\cdot\|_1$, which is used for example in sparse least-squares regression (also known as LASSO [8]):

$$\min_x \frac{1}{2} \|Bx - b\|^2 + \|x\|_1. \quad (2)$$

This formulation is convex and forms a special case of (1) in which A is given by the identity. Regularization with the norm $\|\cdot\|_1$ is used also in logistic regression [9].

Other Convex Regularizers. The case of problems (1) in which g is nonsmooth and *convex* (with possible smooth and/or nonsmooth additive terms) has received a great deal of attention in the literature on convex optimization and applications; see for example [1,4,10–12]. The most notable applications are found in inverse problems involving images. In particular, discrete (an)isotropic *Total Variation (TV)* denoising has the form

$$\min_x \frac{1}{2} \|x - b\|^2 + \|\nabla x\|_1,$$

where b is the observed (noisy) image and ∇ denotes the discretized gradient in two or three dimensions. TV deblurring problems have the form

$$\min_x \frac{1}{2} \|Bx - b\|^2 + \|\nabla x\|_1,$$

where B is the blurring operator; see [1,13].

Other examples of convex problems of the form (1) include generalized convex feasibility [4] and support vector machine classification [5]. A typical formulation of the latter problem has $h(x) = (\lambda/2)\|x\|^2$ and $g(Ax) = \sum_{i=1}^n \phi(y_i a_i^T x)$, where $\phi(s) = \max\{-s, 0\}$ is the hinge loss and the rows of A are $y_i a_i^T$, $i = 1, 2, \dots, n$, where $(y_i, a_i) \in \{-1, 1\} \times \mathbb{R}^d$ are the training points and their labels.

Weakly Convex Regularizers. The use of the ℓ_1 regularizer in (2) tends to depress the magnitude of nonzero elements of the solution, resulting in *bias*. This phenomenon is a consequence of the fact that the proximal operator of the 1-norm, often called the *soft thresholding operator*, does not approach the identity even for large arguments. For this reason, nonconvex alternatives to $\|\cdot\|_1$ are often used to reduce bias. These include ℓ_p -norms (with $0 < p < 1$) which are not weakly convex, and several weakly convex regularizers, which we now describe. The *minimax concave penalty (MCP)*, introduced in [14] and used in [15, 16], is a family of functions $r_{\lambda, \theta} : \mathbb{R} \rightarrow \mathbb{R}_+$ involving two positive parameters λ and θ , and defined by

$$r_{\lambda, \theta}(x) := \begin{cases} \lambda|x| - \frac{x^2}{2\theta}, & |x| \leq \theta\lambda, \\ \frac{\theta\lambda^2}{2}, & \text{otherwise.} \end{cases}$$

(Note that this function satisfies the definition of ρ -weak convexity with $\rho = \theta^{-1}$.) The proximal operator of this function (called *firm threshold* in [17]) can be written in the following closed form when $\theta > \gamma$:

$$\text{prox}_{\gamma r_{\lambda, \theta}}(x) = \begin{cases} 0, & |x| < \gamma\lambda, \\ \frac{x - \lambda\gamma \text{sgn}(x)}{1 - (\gamma/\theta)}, & \gamma\lambda \leq |x| \leq \theta\lambda, \\ x, & |x| > \theta\lambda. \end{cases}$$

The *fractional penalty function* (cf. [16, 18]) $\phi_a : \mathbb{R} \rightarrow \mathbb{R}_+$ (for parameter $a > 0$) is

$$\phi_a(x) := \frac{|x|}{1 + a|x|/2}.$$

The *smoothly clipped absolute deviation (SCAD)* [19] (cf. [16]) is defined for parameters $\lambda > 0$ and $\theta > 2$ as follows:

$$r_{\lambda, \theta}(x) = \begin{cases} \lambda|x|, & |x| \leq \lambda, \\ \frac{-x^2 + 2\theta\lambda|x| - \lambda^2}{2(\theta-1)}, & \lambda < |x| \leq \theta\lambda, \\ \frac{(\theta+1)\lambda^2}{2}, & |x| > \theta\lambda. \end{cases}$$

(This function is $(\theta - 1)^{-1}$ -weakly convex.)

Since these functions approach (or attain) a finite value as their argument grows in magnitude, they do not introduce as much bias in the solution as does the ℓ_1 norm, and their proximal operators approach the identity for large arguments.

The regularizers of this section, and the convex $\|\cdot\|$ regularizer, have been used mostly in the simple additive setting

$$\min_{x \in \mathbb{R}^d} h(x) + g(x)$$

for a smooth data fidelity term h and nonsmooth regularizer g , for example in least squares or logistic regression [15] and compressed sensing (cf. [17]).

Weakly convex composite losses The use of weakly convex functions composed with linear operators has been explored in the robust statistics literature. An early instance is the *Tukey biweight* function [20], in which $g(Ax)$ has the form

$$g(Ax) = \sum_{i=1}^n \phi(A_i \cdot x - b_i), \quad \text{where } \phi(\theta) = \frac{\theta^2}{1 + \theta^2}, \quad (3)$$

where $A_i \cdot$ denotes the i th row of A . This function behaves like the usual least-squares loss when $\theta^2 \ll 1$ but asymptotes at 1. It is ρ -weakly convex with $\rho = 6$.

A slightly different definition of the Tukey biweight function appears in [21, Section 2.1]. This same reference also mentions another nonconvex loss function, the *Cauchy loss*, which has the form (3) except that ϕ is defined by

$$\phi(\theta) = \frac{\xi^2}{2} \log \left(1 + \frac{\theta^2}{\xi^2} \right),$$

for some parameter ξ . This function is ρ -weakly convex with $\rho = 6$.

2.2 Complexity Bounds for Weakly Convex Problems

To put our results in perspective, we provide a review of the literature on complexity bounds for optimization problems related to our formulation (1), including weakly convex functions. In all cases, these are bounds on the number of iterations required to find an approximately stationary point, where our measure of stationarity is based the norm of the gradient of the Moreau envelope (a smooth proxy).

The best-known complexity for black box subgradient optimization for weakly convex functions is $\mathcal{O}(\epsilon^{-4})$. This result is proved for *stochastic* subgradient in [22], but as in the convex case, there is no known improvement in the deterministic setting. As in convex optimization, subgradient methods are quite general and implementable for weakly convex functions. However, when more structure is present in the function, algorithms that achieve better complexity can be devised. In particular, when the proximal operator of the nonsmooth weakly convex function can be calculated analytically, complexity bounds of $\mathcal{O}(\epsilon^{-2})$ can be proved (see Sect. 5), the same bounds as for steepest descent methods in the smooth nonconvex case. This means that the entire difficulty introduced by the nonsmoothness can be mitigated as long as the nonsmoothness can be treated by a proximal operator.

For convex optimization problems, bounds of $\mathcal{O}(\epsilon^{-1})$ can be obtained for gradient methods on smooth functions and $\mathcal{O}(\epsilon^{-1/2})$ for accelerated gradient methods. These

same bounds can also be obtained for nonsmooth problems provided that the nonsmooth term is handled by a proximal operator. When the explicit proximal operator is not available and subgradient methods have to be used, the complexity reverts to $\mathcal{O}(\epsilon^{-2})$.

It is possible to keep the $\mathcal{O}(\epsilon^{-2})$ rate when just a local model of the weakly convex part is evaluated by a convex operator. The paper [23] studies optimization problems of the type

$$\min_x h(x) + g(c(x))$$

where h is convex, proper and closed; g is convex and Lipschitz continuous; and c is smooth. (Under these assumptions, the composition $g \circ c$ is weakly convex.) The $\mathcal{O}(\epsilon^{-2})$ bound is proved for an algorithm in which the (convex) subproblem

$$\min_y h(y) + g(c(x) + \nabla c(x)(y - x)) + \frac{1}{2t} \|y - x\|^2 \quad (4)$$

is solved explicitly. In the more realistic case in which (4) must be solved by an iterative procedure, a bound of $\tilde{\mathcal{O}}(\epsilon^{-3})$ is obtained in [23]. (The symbol $\tilde{\mathcal{O}}$ hides logarithmic terms.)

Functions of the form $g(c(x))$ have also been studied in [24] for the case of a smooth nonlinear vector function c and a prox-regular g . This formulation is more general than those considered in this paper, both in the fact that all weakly convex functions are prox-regular, and in the nonlinearity of the inner map c . The subproblems in [24] have a form similar to (4), and while convergence results are proved in the latter paper, it does not contain rate-of-convergence results or complexity results.

A different weakly convex structure is explored in [25], in which the weak convexity stems from a smooth saddle point problem. This paper studies the problem

$$\min_x \max_{y \in Y} l(x, y),$$

for a compact set $Y \subset \mathbb{R}^m$, where $l(x, \cdot)$ is concave, $l(\cdot, y)$ is nonconvex, and $l(\cdot, \cdot)$ is smooth. An iteration bound of $\tilde{\mathcal{O}}(\epsilon^{-3})$ is proved for a method that uses only gradient evaluations.

In light of the considerations above, the complexity bound of $\mathcal{O}(\epsilon^{-3})$ for our algorithm seems almost inevitable. It interpolates between the setting without structural assumptions about the nonsmoothness (black box subgradient) and the perfect structural knowledge of the nonsmoothness (explicit knowledge of the proximal operator).

In Sect. 5, we treat the simpler setting in which the linear operator from (1) is the identity, so that $F(x) = h(x) + g(x)$. Similar problems have been analyzed before, for example, in [15, 17]. However, it is assumed in [17] that convexity in the data fidelity term h compensates for nonconvexity in the regularizer g such that the overall objective function F remains convex. (We make no such assumption here.) The paper [15] does not make such restrictive assumptions and proves convergence but not complexity bounds.

3 Preliminaries

The concept of subgradient of a convex function can be adapted to weakly convex functions via the following definition.

Definition 3.1 (Fréchet subdifferential) Let $g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ be a function and $\bar{y}$ a point such that $g(\bar{y})$ is finite. Then, the *Fréchet subdifferential* of g at $\bar{y}$, denoted by $\partial g(\bar{y})$, is the set of all vectors $v \in \mathbb{R}^n$ such that

$$g(y) \geq g(\bar{y}) + \langle v, y - \bar{y} \rangle + o(\|y - \bar{y}\|) \quad \text{as } y \rightarrow \bar{y}. \quad (5)$$

Modifying the convex case, in which subgradients are the slopes of linear functions that underestimate g but coincide with it at $\bar{y}$, Fréchet subgradients do so *up to first order*. This definition makes sense for arbitrary functions, but for lower semicontinuous ρ -weakly convex functions, more can be said. For example, for this class of function we know that subgradients satisfy the following stronger version of (5), for all $v \in \partial g(\bar{y})$,

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

Further, if we assume the weakly convex function to be continuous at a point y , then its subdifferential is nonempty at y . Both of these claims can be verified directly by adding $\frac{\rho}{2} \|\cdot\|^2$ to g and considering the convex subdifferential; see [26, Lemma 2.1].

Another nice property of weakly convex functions is that the definition of a Moreau envelope (see Definition 2.1) extends without modification to weakly convex functions, subject only to a restriction on the parameter μ . The proximal operator of Definition 2.1 also extends to this setting, and this operator and the Moreau envelope fulfil the same identity as in the convex setting.

Lemma 3.1 ([27, Corollary 3.4]) Let $g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ be a proper, ρ -weakly convex, and lower semicontinuous function, and let $\mu \in]0, \rho^{-1}[$. Then, the Moreau envelope $g_\mu(\cdot)$ is continuously differentiable on $\mathbb{R}^n$ with gradient

$$\nabla g_\mu(y) = \frac{1}{\mu} \left(y - \text{prox}_{\mu g}(y) \right), \quad \text{for all } y \in \mathbb{R}^n.$$

This gradient is $\max \left\{ \mu^{-1}, \frac{\rho}{1-\rho\mu} \right\}$ -Lipschitz continuous. In particular, a gradient step with respect to the Moreau envelope corresponds to a proximal step, that is,

$$y - \mu \nabla g_\mu(y) = \text{prox}_{\mu g}(y), \quad \text{for all } y \in \mathbb{R}^n. \quad (6)$$

Lemma 3.1 not only clarifies the smoothness of the Moreau envelope, but also gives a way of computing its gradient via the prox operator. Obviously, a smooth representation whose gradient could not be computed would be of only limited usefulness from an algorithmic standpoint. The only difference between the weakly convex and convex settings is that the Moreau envelope need not be convex in the former case.

3.1 Stationarity

We say that a point $\bar{x}$ is a stationary point for a function if the Fréchet subdifferential of the function contains 0 at $\bar{x}$. The concept of *nearly stationary* is not quite so straightforward. We motivate our approach by looking first at the simple additive composite problem, also discussed in Sect. 5, which corresponds to setting $A = I$ in (1), that is,

$$\min_x h(x) + g(x). \quad (7)$$

Stationarity for (7) means that $0 \in \partial(h + g)(\bar{x})$, that is, $-\nabla h(\bar{x}) \in \partial g(\bar{x})$. A natural definition for ϵ -approximate stationarity would thus be

$$\text{dist}(-\nabla h(x), \partial g(x)) \leq \epsilon, \quad (8)$$

where dist denotes the distance between two sets and is given for a point $x \in \mathbb{R}^d$ and a set $\mathcal{A} \subset \mathbb{R}^d$ by $\text{dist}(x, \mathcal{A}) := \inf_{y \in \mathcal{A}} \{\|x - y\|\}$. However, since we are running gradient descent on the *smoothed* problem, our algorithm will naturally compute and detect points with that satisfy a threshold condition of the form

$$\|\nabla h(x) + \nabla g_\mu(x)\| \leq \epsilon. \quad (9)$$

The next lemma helps to clarify relationship between these two conditions.

Lemma 3.2 *Let $g : \mathbb{R}^n \rightarrow \bar{\mathbb{R}}$ be a proper, ρ -weakly convex, and lower semicontinuous function; and let $\mu \in]0, \rho^{-1}[$. Then,*

$$\nabla g_\mu(x) \in \partial g(\text{prox}_{\mu g}(x)). \quad (10)$$

Proof From Definition 2.1, we have that

$$0 \in \partial g(\text{prox}_{\mu g}(x)) + \frac{1}{\mu}(\text{prox}_{\mu g}(x) - x),$$

from which the result follows when we use (6). $\square$

(This result is proved for the case of g convex in [23, Lemma 2.1], with essentially the same proof.)

This lemma tells us that when (9) holds, then (8) is nearly satisfied, except that in the argument of ∂g , x is replaced by $\text{prox}_{\mu g}(x)$. In general, however, $\text{prox}_{\mu g}(x)$ might be arbitrarily far away from x . We can remedy this issue by requiring g to be Lipschitz continuous also.

Lemma 3.3 *Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ be a ρ -weakly convex function that is L_g -Lipschitz continuous, and let $\mu \in]0, \rho^{-1}[$. Then, the Moreau envelope g_μ is Lipschitz continuous with*

$$\|\nabla g_\mu(x)\| \leq L_g \quad (11)$$

and

$$\|x - \text{prox}_{\mu g}(x)\| \leq \mu L_g, \quad \forall x \in \mathbb{R}^n. \quad (12)$$

Proof Lipschitz continuity is equivalent to bounded subgradients [28], so by (10), we have for all $x \in \mathbb{R}^n$

$$\|\nabla g_\mu(x)\| \leq \sup \left\{ \|v\| : v \in \partial g(\text{prox}_{\mu g}(x)) \right\} \leq L_g,$$

proving (11). The bound (12) follows immediately when we use the fact that $x - \text{prox}_{\mu g}(x) = \mu \nabla g_\mu(x)$ from Lemma 3.1. $\square$

When $x \in \mathbb{R}^n$ satisfies (9), ∇h is $L_{\nabla h}$ -Lipschitz continuous, g is L_g Lipschitz continuous, we have

$$\begin{aligned} & \text{dist}(-\nabla h(\text{prox}_{\mu g}(x)), \partial g(\text{prox}_{\mu g}(x))) \\ & \leq \|\nabla h(\text{prox}_{\mu g}(x)) - \nabla h(x)\| + \text{dist}(-\nabla h(x), \partial g(\text{prox}_{\mu g}(x))) \\ & \leq L_{\nabla h} \|x - \text{prox}_{\mu g}(x)\| + \epsilon \quad (\text{from (9) and (10)}) \\ & \leq L_{\nabla h} L_g \mu + \epsilon \quad (\text{from (12)}). \end{aligned}$$

Thus, if μ is sufficiently small and x satisfies (9), then $\text{prox}_{\mu g}(x)$ is near-stationary for (7).

3.2 Stationarity for the Composite Problem

It follows immediately from (10) in Lemma 3.2 that for $\mu \in]0, \rho^{-1}[$, we have for all $x \in \mathbb{R}^d$

$$\nabla(g_\mu \circ A)(x) = A^* \nabla g_\mu(Ax) \in A^* \partial g(\text{prox}_{\mu g}(Ax)). \quad (13)$$

Extending the results of the previous section to the case of a general linear operator A in (1) requires some work. Stationarity for (1) requires that $0 \in \nabla h(x) + A^* \partial g(Ax)$, so ϵ -near stationarity requires

$$\text{dist}(-\nabla h(x), A^* \partial g(Ax)) \leq \epsilon. \quad (14)$$

Our method can compute a point x such that

$$\|\nabla h(x) + \nabla(g_\mu \circ A)(x)\| \leq \epsilon$$

which by (13) implies that

$$\text{dist}(-\nabla h(x), A^* \partial g(z)) \leq \epsilon, \quad \text{where } z = \text{prox}_{\mu g}(Ax), \quad (15)$$

where provided that g is L_g -Lipschitz continuous, we have

$$\|Ax - z\| \leq L_g \mu. \quad (16)$$

The bound in (15) measures the criticality, while the bound in (16) concerns feasibility. The bounds (15), (16) are not a perfect match with (14), since the subdifferentials of h and $g \circ A$ are evaluated at different points.

Surjectivity of A . When A is surjective, we can perturb the x that satisfies (15), (16) to a nearby point x^* that satisfies a bound of the form (14). Since $z = \text{prox}_{\mu g}(Ax)$ is in the range of A , we can define

$$x^* := \arg \min_{x' \in \mathbb{R}^d} \{\|x - x'\|^2 : Ax' = z\}, \quad (17)$$

which is given explicitly by

$$x^* = x - A^*(AA^*)^{-1}(Ax - z) = x - A^\dagger(Ax - z)$$

where $A^\dagger := A^*(AA^*)^{-1}$ is the pseudoinverse of A . The operator norm of the pseudoinverse is bounded by the inverse of the smallest singular value $\sigma_{\min}(A)$ of A , so when g is L_g -Lipschitz continuous, we have from (16) that

$$\|x - x^*\| \leq \sigma_{\min}(A)^{-1} \|Ax - z\| \leq \sigma_{\min}(A)^{-1} L_g \mu. \quad (18)$$

The point x^* is approximately stationary in the sense of (14), for μ sufficiently small, because

$$\begin{aligned} & \text{dist}(-\nabla h(x^*), A^* \partial g(Ax^*)) \\ & \leq \|\nabla h(x^*) - \nabla h(x)\| + \text{dist}(-\nabla h(x), A^* \partial g(z)) \quad (\text{since } Ax^* = z) \\ & \leq L_{\nabla h} \|x - x^*\| + \epsilon \quad (\text{from (15)}) \\ & \leq L_{\nabla h} \sigma_{\min}(A)^{-1} L_g \mu + \epsilon \quad (\text{from (18)}). \end{aligned} \quad (19)$$

By choosing μ small, x^* will be an approximate solution in the stronger sense (14) and not just the weaker notion of (15), (16), which is the case if A is not surjective.

4 Variable Smoothing

We describe our variable smoothing approaches for the problem (1), where we assume that h is $L_{\nabla h}$ -smooth; g is possibly nonsmooth, ρ -weakly convex, and L_g -Lipschitz continuous; and A is a nonzero linear continuous operator. For convenience, we define the smoothed approximation $F_k : \mathbb{R}^d \rightarrow \mathbb{R}$ based on the Moreau envelope with parameter μ_k as follows:

$$F_k(x) := h(x) + g_{\mu_k}(Ax).$$

We note from Lemma 3.1 and the chain rule that

$$\nabla F_k(x) = \nabla h(x) + \frac{1}{\mu_k} A^*(Ax - \text{prox}_{\mu_k g}(Ax)). \quad (20)$$

The quantity L_k defined by

$$L_k := L_{\nabla h} + \|A\|^2 \max \left\{ \mu_k^{-1}, \frac{\rho}{1 - \rho \mu_k} \right\} \quad (21)$$

is a Lipschitz constant of the gradient of ∇F_k ; see Lemma 3.1. When $\rho\mu_k \leq 1/2$, the maximum in (21) is achieved by μ_k^{-1} , so in this case we can define

$$L_k := L_{\nabla h} + \|A\|^2/\mu_k. \quad (22)$$

4.1 An Elementary Approach

Our first algorithm takes gradient descent steps on the smoothed problem, that is,

$$x_{k+1} = x_k - \gamma_k \nabla F_k(x_k), \quad (23)$$

for certain values of the parameter μ_k and step size γ_k . From (20), the formula (23) is equivalent to

$$x_{k+1} = x_k - \frac{\gamma_k}{\mu_k} A^*(Ax_k - \text{prox}_{\mu_k g}(Ax_k)) - \gamma_k \nabla h(x_k).$$

Our basic algorithm is described next.

Algorithm 1 Variable Smoothing

Require: $x_1 \in \mathbb{R}^d$;
for $k = 1, 2, 3, \dots$ **do**
 Set $\mu_k \leftarrow (2\rho)^{-1}k^{-1/3}$, define L_k as in (22), set $\gamma_k \leftarrow 1/L_k$;
 Set $x_{k+1} \leftarrow x_k - \gamma_k \nabla F_k(x_k)$;
end for

We now state the convergence result for Algorithm 1. This result and later results make use of a quantity

$$F^* := \liminf_{k \rightarrow \infty} F_k(x_k), \quad (24)$$

which is finite if F is bounded below (and possibly in other circumstances too). (When $F^* = -\infty$, the claim of the theorem is vacuously true.) We also make use of the following quantity:

$$x_j^* := x_j - A^\dagger(Ax_j - \text{prox}_{\mu_j g}(Ax_j)). \quad (25)$$

Theorem 4.1 *Suppose that Algorithm 1 is applied to the problem (1), where g is ρ -weakly convex and ∇h and g are Lipschitz continuous with constants $L_{\nabla h}$ and L_g , respectively. We have*

$$\begin{aligned} \min_{1 \leq j \leq k} \text{dist}(-\nabla h(x_j), A^* \partial g(\text{prox}_{\mu_j g}(Ax_j))) \\ \leq k^{-1/3} \sqrt{L_{\nabla h} + 2\rho\|A\|^2} \sqrt{F_1(x_1) - F^* + (2\rho)^{-1}L_g^2}, \end{aligned}$$

where

$$\|Ax_j - \text{prox}_{\mu_j g}(Ax_j)\| \leq j^{-1/3}(2\rho)^{-1}L_g,$$

and F^* is defined as in (24). If A is also surjective, then for x_k^* defined as in (25), we have

$$\begin{aligned} & \min_{1 \leq j \leq k} \text{dist}(-\nabla h(x_j^*), A^* \partial g(Ax_j^*)) \\ & \leq k^{-1/3} \left(\sqrt{L_{\nabla h} + 2\rho\|A\|^2} \sqrt{F_1(x_1) - F^* + (2\rho)^{-1}L_g^2} + L_{\nabla h} \sigma_{\min}(A)^{-1}L_g \right) \end{aligned}$$

$$\text{and } \|x_j - x_j^*\| \leq \sigma_{\min}(A)^{-1}L_g\mu_j = \sigma_{\min}(A)^{-1}L_g(2\rho)^{-1}j^{-1/3}.$$

Before proving this theorem, we state and prove a lemma that relates the function values of two Moreau envelopes with two different smoothing parameters. In the convex case, such statements are well known, but in the nonconvex case this result is novel.

Lemma 4.1 *Let $g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ be a proper, closed, and ρ -weakly convex function, and let μ_2 and μ_1 be parameters such that $0 < \mu_2 \leq \mu_1 < \rho^{-1}$. Then, we have*

$$g_{\mu_2}(y) \leq g_{\mu_1}(y) + \frac{1}{2} \frac{\mu_1 - \mu_2}{\mu_2} \mu_1 \|\nabla g_{\mu_1}(y)\|^2.$$

If, in addition, g is L_g -Lipschitz continuous, we have

$$g_{\mu_2}(y) \leq g_{\mu_1}(y) + \frac{1}{2} \frac{\mu_1 - \mu_2}{\mu_2} \mu_1 L_g^2.$$

Proof By using the definition of the Moreau envelope, together with Lemma 3.1, we obtain

$$\begin{aligned} g_{\mu_2}(y) &= \min_{u \in \mathbb{R}^n} \left\{ g(u) + \frac{1}{2\mu_2} \|y - u\|^2 \right\} \\ &= \min_{u \in \mathbb{R}^n} \left\{ g(u) + \frac{1}{2\mu_1} \|y - u\|^2 + \frac{1}{2} \left(\frac{1}{\mu_2} - \frac{1}{\mu_1} \right) \|y - u\|^2 \right\} \\ &\leq g(\text{prox}_{\mu_1 g}(y)) + \frac{1}{2\mu_1} \|y - \text{prox}_{\mu_1 g}(y)\|^2 + \frac{1}{2} \left(\frac{1}{\mu_2} - \frac{1}{\mu_1} \right) \|y - \text{prox}_{\mu_1 g}(y)\|^2 \\ &= g_{\mu_1}(y) + \frac{1}{2} \left(\frac{\mu_1 - \mu_2}{\mu_2} \right) \mu_1 \|\nabla g_{\mu_1}(y)\|^2, \end{aligned}$$

proving the first claim. The second claim follows immediately from (11). $\square$

Proof of Theorem 4.1 Since $L_k = 1/\gamma_k$ is the Lipschitz constant of ∇F_k , we have for any $k = 1, 2, \dots$ that

$$F_k(x_{k+1}) \leq F_k(x_k) + \langle \nabla F_k(x_k), x_{k+1} - x_k \rangle + \frac{1}{2\gamma_k} \|x_{k+1} - x_k\|^2.$$

By substituting the definition of x_{k+1} from (23), we have

$$F_k(x_{k+1}) \leq F_k(x_k) - \frac{\gamma_k}{2} \|\nabla F_k(x_k)\|^2. \quad (26)$$

From Lemma 4.1, we have for all $x \in \mathbb{R}^d$

$$F_{k+1}(x) \leq F_k(x) + \frac{1}{2}(\mu_k - \mu_{k+1}) \frac{\mu_k}{\mu_{k+1}} \|(\nabla g_{\mu_k})(Ax)\|^2 \leq F_k(x) + (\mu_k - \mu_{k+1}) L_g^2,$$

where we used in the second inequality that $\frac{\mu_k}{\mu_{k+1}} \leq 2$. We set $x = x_{k+1}$ and substitute into (26) to obtain

$$F_{k+1}(x_{k+1}) \leq F_k(x_k) - \frac{\gamma_k}{2} \|\nabla F_k(x_k)\|^2 + (\mu_k - \mu_{k+1}) L_g^2.$$

By summing both sides of this expression over $k = 1, 2, \dots, K$, and telescoping, we deduce that

$$\begin{aligned} \sum_{k=1}^K \frac{\gamma_k}{2} \|\nabla F_k(x_k)\|^2 &\leq F_1(x_1) - F_K(x_K) + (\mu_1 - \mu_K) L_g^2 \\ &\leq F_1(x_1) - F^* + \mu_1 L_g^2. \end{aligned} \quad (27)$$

Since

$$\gamma_k = \frac{\mu_k}{\mu_k L_{\nabla h} + \|A\|^2} \geq k^{-1/3} \frac{(2\rho)^{-1}}{(2\rho)^{-1} L_{\nabla h} + \|A\|^2} = k^{-1/3} \frac{1}{L_{\nabla h} + 2\rho \|A\|^2}.$$

we have from (27) that

$$\frac{1}{L_{\nabla h} + 2\rho \|A\|^2} \min_{1 \leq j \leq K} \|\nabla F_j(x_j)\|^2 \frac{1}{2} \sum_{k=1}^K k^{-1/3} \leq F_1(x_1) - F^* + (2\rho)^{-1} L_g^2. \quad (28)$$

Now we observe that

$$\begin{aligned} \sum_{k=1}^K k^{-1/3} &\geq \sum_{k=1}^K \int_k^{k+1} x^{-1/3} dx = \int_1^{K+1} x^{-1/3} dx = \frac{3}{2} \left((K+1)^{2/3} - 1 \right) \\ &\geq (K+1)^{2/3} - 1 \geq \frac{1}{2} K^{2/3}, \quad K = 1, 2, \dots \end{aligned}$$

where the final inequality can be checked numerically. Therefore, by substituting into (28), we have

$$\min_{1 \leq j \leq K} \|\nabla F_j(x_j)\|^2 \leq 4 \frac{L_{\nabla h} + (2\rho)\|A\|^2}{K^{2/3}} \left(F_1(x_1) - F^* + (2\rho)^{-1} L_g^2 \right),$$

and so

$$\min_{1 \leq j \leq K} \|\nabla F_j(x_j)\| \leq \frac{C}{K^{1/3}}, \quad (29)$$

where $C := 2\sqrt{L_{\nabla h} + (2\rho)\|A\|^2} \sqrt{F_1(x_1) - F^* + (2\rho)^{-1} L_g^2}$. By combining this bound with (15), and defining $z_j := \text{prox}_{\mu_j g}(Ax_j)$, we obtain

$$\min_{1 \leq j \leq k} \text{dist}(-\nabla h(x_j), A^* \partial g(z_j)) \leq \min_{1 \leq j \leq k} \|\nabla F_j(x_j)\| \leq \frac{C}{k^{1/3}}, \quad (30)$$

where we deduce from (12) that

$$\|Ax_j - z_j\| \leq \frac{(2\rho)^{-1} L_g}{j^{1/3}}, \quad \text{for all } j \geq 1. \quad (31)$$

The second statement concerning surjectivity of A follows from the consideration made in (17) to (19). $\square$

There is a mismatch between the two bounds in this theorem. The first bound (the criticality bound) indicates that during the first $k = O(\epsilon^{-3})$ iterations, we will encounter an iteration j at which the first-order optimality condition is satisfied to within a tolerance of ϵ . However, this bound could have been satisfied at an early iteration (that is, $j \ll \epsilon^{-3}$), for which value the second (feasibility) bound, on $\|Ax_j - \text{prox}_{\mu_j g}(Ax_j)\|$, may not be particularly small. The next section describes an algorithm that remedies this defect.

4.2 An Epoch-Wise Approach with Improved Convergence Guarantees

We describe a variant of Algorithm 1 in which the steps are organized into a series of epochs, each of which is twice as long as the one before. We show that there is some iteration $j = O(\epsilon^{-3})$ such that both $\|Ax_j - \text{prox}_{\mu_j g}(Ax_j)\|$ and $\text{dist}(-\nabla h(x_j), A^* \partial g(\text{prox}_{\mu_j g}(Ax_j)))$ are smaller than the given tolerance ϵ .

Theorem 4.2 *Consider solving (1) using Algorithm 2, where h and g satisfy the assumptions of Theorem 4.1 and F^* defined in (24) is finite. For a given tolerance $\epsilon > 0$, Algorithm 2 generates an iterate x_j for some $j = O(\epsilon^{-3})$ such that*

$$\text{dist}(-\nabla h(x_j), A^* \partial g(z_j)) \leq \epsilon \quad \text{and} \quad \|Ax_j - z_j\| \leq \epsilon,$$

where $z_j = \text{prox}_{\mu_j g}(Ax_j)$.

Algorithm 2 Variable Smoothing with Epochs

Require: $x_1 \in \mathbb{R}^d$ and tolerance $\epsilon > 0$;
for $l = 0, 1, \dots$ **do**
 Set $S_l \leftarrow \infty$, Set $j_l \leftarrow 2^l$;
 for $k = 2^l, 2^l + 1, \dots, 2^{l+1} - 1$ **do**
 Set $\mu_k \leftarrow (2\rho)^{-1}k^{-1/3}$, define L_k as in (22), set $\gamma_k \leftarrow 1/L_k$;
 Set $x_{k+1} \leftarrow x_k - \gamma_k \nabla F_k(x_k)$;
 if $\|\nabla F_{k+1}(x_{k+1})\| \leq S_l$ **then**
 Set $S_l \leftarrow \|\nabla F_{k+1}(x_{k+1})\|$; Set $j_l \leftarrow k + 1$;
 if $S_l \leq \epsilon$ and $\|Ax_{k+1} - \text{prox}_{\mu_{k+1}g}(Ax_{k+1})\| \leq \epsilon$ **then**
 STOP;
 end if
 end if
 end for
end for

Proof As in (27), by using monotonicity of $\{F_k(x_k)\}$ and discarding nonnegative terms, we have that

$$\sum_{k=2^l}^{2^{l+1}-1} \frac{\gamma_k}{2} \|\nabla F_k(x_k)\|^2 \leq F_1(x_1) - F^* + (2\rho)^{-1} L_g^2.$$

With the same arguments as in the earlier proof, we obtain

$$\begin{aligned} \sum_{k=2^l}^{2^{l+1}-1} k^{-1/3} &\geq \sum_{k=2^l}^{2^{l+1}-1} \int_k^{k+1} x^{-1/3} dx = \int_{2^l}^{2^{l+1}} x^{-1/3} dx \\ &= \frac{3}{2} \left((2^{l+1})^{2/3} - (2^l)^{2/3} \right) = \frac{3}{2} \left(2^{2/3} - 1 \right) (2^l)^{2/3} \geq \frac{1}{2} (2^l)^{2/3}. \end{aligned}$$

Therefore, we have

$$\min_{2^l \leq j \leq 2^{l+1}-1} \|\nabla F_j(x_j)\| \leq \frac{C}{(2^l)^{1/3}},$$

with $C = 2\sqrt{L_{\nabla h} + 2\rho\|A\|^2} \sqrt{F_1(x_1) - F^* + (2\rho)^{-1} L_g^2}$ as before. Noting that $z_j := \text{prox}_{\mu_j g}(Ax_j)$, we have as in (30) that

$$\min_{2^l \leq j \leq 2^{l+1}-1} \text{dist}(-\nabla h(x_j), A^* \partial g(z_j)) \leq \frac{C}{(2^l)^{1/3}}, \quad (32)$$

as previously. Further, we have for $2^l \leq j \leq 2^{l+1} - 1$ that

$$\|Ax_j - z_j\| \leq L_g \mu \leq \frac{(2\rho)^{-1} L_g}{j^{1/3}} \leq \frac{(2\rho)^{-1} L_g}{(2^l)^{1/3}}. \quad (33)$$

From (32) and (33), we deduce that Algorithm 2 must terminate before the end of epoch l , that is, before 2^{l+1} iterations have been completed, where l is the first nonnegative

integer such that

$$2^l \geq \max\{C^3, (2\rho)^{-3}L_g^3\}\epsilon^{-3}.$$

Thus, termination occurs after at most $2 \max\{C^3, (2\rho)^{-3}L_g^3\}\epsilon^{-3}$ iterations. $\square$

For the case of A surjective, we have the following stronger result.

Corollary 4.1 *Suppose that the assumptions of Theorem 4.2 hold, that A is also surjective, and that the condition $\|Ax_{k+1} - \text{prox}_{\mu_{k+1}g}(Ax_{k+1})\| \leq \epsilon$ in Algorithm 2 is replaced by $\|x_{k+1} - x_{k+1}^*\| \leq \epsilon$, where x_{k+1}^* is defined in (25). Then, for some $j' = O(\epsilon^{-3})$, we have that*

$$\text{dist}(-\nabla h(x_{j'}^*), A^* \partial g(Ax_{j'}^*)) \leq \epsilon$$

$$\text{and } \|x_{j'} - x_{j'}^*\| \leq \epsilon.$$

Proof With the considerations made in the previous proof as well as the one made in (17) to (19), we can choose l to be the smallest positive integer such that

$$2^{l+1} \geq 2 \max\{C^3, \sigma_{\min}(A)^{-3}L_g^3(2\rho)^{-3}\}\epsilon^{-3}.$$

The claim then holds for some $j' \leq 2^{l+1}$. $\square$

Although Algorithm 2 seems more complicated than Algorithm 1, the steps are the same. The only difference is that for the second algorithm, we do not search for the iterate that minimizes criticality across *all* iterations but only across at most the last $k/2$ iterations, where k is the total number of iterations.

Remark 4.1 For both versions of our proposed method we use an explicit choice of smoothing parameters, choosing μ_k to be a multiple of $\mathcal{O}(k^{-1/3})$. This specific dependence on k achieves a balance between criticality and feasibility. As can be seen from (29) (criticality measure) and (31) (feasibility measure) both measures decrease like $k^{-1/3}$. A slower decrease in μ_k would result in a faster decrease in the criticality measure but a slower decrease in the feasibility measure—and vice versa.

Remark 4.2 Our technique does not adapt in an obvious way to the case in which g is actually convex. Typically, we know in advance whether or not h and g in (1) are convex, and if they are, we could choose one of the well-established methods that make use of gradients, proximal operators, and possibly acceleration. See, for example the proximal accelerated gradient approach of [3], which achieves a rate of $\mathcal{O}(k^{-1})$. A method in the spirit of [29], which automatically adapts to convexity and simultaneously achieves the optimal rates for both nonconvex and convex problems would be desirable, but is outside the scope of this work.

5 Proximal Gradient

Here we derive a complexity bound for the proximal gradient algorithm applied to the more elementary problem (7) studied in Sect. 3.1, that is,

$$\min_{x \in \mathbb{R}^d} F(x) := h(x) + g(x), \quad (34)$$

for $h : \mathbb{R}^d \rightarrow \mathbb{R}$ a $L_{\nabla h}$ -smooth function and $g : \mathbb{R}^d \rightarrow \overline{\mathbb{R}}$ a possibly nonsmooth, ρ -weakly convex function. Such a bound has not been made explicit before, to the authors' knowledge, though it is a fairly straightforward consequence of existing results. The bound makes an interesting comparison with the result in Sect. 4, where the nonsmoothness issue becomes more complicated due to the composition with a linear operator. In this section, we assume that a closed-form proximal operator is available for g , and we show that the complexity bound of $\mathcal{O}(\epsilon^{-2})$ is the same order as for gradient descent applied to smooth nonconvex functions.

Standard proximal gradient applied to problem (34), for a given step size $\lambda \in]0, \min\{\rho^{-1}/2, L_{\nabla h}^{-1}\}]$ and initial point x_1 , is as follows:

$$\begin{aligned} x_{k+1} &:= \arg \min_{x \in \mathbb{R}^d} \left\{ g(x) + \langle \nabla h(x_k), x - x_k \rangle + \frac{1}{2\lambda} \|x - x_k\|^2 \right\}, \\ &= \text{prox}_{\lambda g}(x_k - \lambda \nabla h(x_k)), \quad k = 1, 2, \dots \end{aligned} \quad (35)$$

where the choice of λ ensures that the function to be minimized in (35) is $(\lambda^{-1} - \rho)$ -strongly convex, so that x_{k+1} is uniquely defined.

We have the following convergence result.

Theorem 5.1 *Consider the algorithm defined by (35) applied to problem (34), where we assume that g is proper, lower semicontinuous and ρ -weakly convex and that ∇h is Lipschitz continuous with constant $L_{\nabla h}$. Supposing that the step size $\lambda \in]0, \min\{\rho^{-1}/2, L_{\nabla h}^{-1}\}]$, we have for all $k \geq 1$ that*

$$\min_{2 \leq j \leq k+1} \text{dist}(0, \partial(h + g)(x_j)) \leq k^{-1/2} \sqrt{2(F(x_1) - F^*)} \frac{\lambda^{-1} + L_{\nabla h}}{\sqrt{\lambda^{-1} - \rho}},$$

where F^* is defined in (24).

Proof Note first that the result is vacuous if $F^* = -\infty$, so we assume henceforth that F^* is finite. We have for every $x \in \mathbb{R}^d$ that

$$\begin{aligned} g(x_{k+1}) + h(x_k) + \langle \nabla h(x_k), x_{k+1} - x_k \rangle + \frac{1}{2\lambda} \|x_{k+1} - x_k\|^2 + \frac{1}{2} (\lambda^{-1} - \rho) \|x - x_{k+1}\|^2 \\ \leq g(x) + h(x_k) + \langle \nabla h(x_k), x - x_k \rangle + \frac{1}{2\lambda} \|x - x_k\|^2. \end{aligned}$$

By applying the inequality

$$h(x_{k+1}) \leq h(x_k) + \langle \nabla h(x_k), x_{k+1} - x_k \rangle + \frac{1}{2\lambda} \|x_{k+1} - x_k\|^2, \quad \text{for all } x \in \mathbb{R}^d,$$

obtained from the Lipschitz continuity of ∇h and the fact that $\lambda \leq L_{\nabla h}^{-1}$, we deduce that

$$F(x_{k+1}) + \frac{1}{2}(\lambda^{-1} - \rho) \|x - x_{k+1}\|^2 \leq g(x) + h(x_k) + \langle \nabla h(x_k), x - x_k \rangle + \frac{1}{2\lambda} \|x - x_k\|^2,$$

for every $x \in \mathbb{R}^d$. By setting $x = x_k$, we obtain

$$F(x_{k+1}) + \frac{1}{2}(\lambda^{-1} - \rho) \|x_k - x_{k+1}\|^2 \leq F(x_k),$$

which shows, together with the definition (24), that

$$\sum_{k=1}^{\infty} \|x_k - x_{k+1}\|^2 \leq \frac{2(F(x_1) - F^*)}{\lambda^{-1} - \rho}. \quad (36)$$

From the optimality conditions for (35), we obtain

$$0 \in \nabla h(x_k) + \partial g(x_{k+1}) + \lambda^{-1}(x_{k+1} - x_k)$$

which also shows that

$$w_{k+1} := \frac{1}{\lambda}(x_k - x_{k+1}) + \nabla h(x_{k+1}) - \nabla h(x_k) \in \partial(h + g)(x_{k+1}), \quad (37)$$

so that

$$\|w_{k+1}\|^2 \leq (\lambda^{-1} + L_{\nabla h})^2 \|x_k - x_{k+1}\|^2.$$

By combining this bound with (36), we obtain

$$\sum_{k=1}^{\infty} \|w_{k+1}\|^2 \leq 2(F(x_1) - F^*) \frac{(\lambda^{-1} + L_{\nabla h})^2}{\lambda^{-1} - \rho}.$$

from which it follows that

$$\min_{1 \leq j \leq k} \|w_{j+1}\| \leq \sqrt{2(F(x_1) - F^*)} \frac{(\lambda^{-1} + L_{\nabla h})}{\sqrt{\lambda^{-1} - \rho}}.$$

The result now follows from (37), when we note that

$$\min_{1 \leq j \leq k} \text{dist}(0, \partial(h + g)(x_{j+1})) \leq \min_{1 \leq j \leq k} \|w_{j+1}\|. \quad \square$$

This theorem indicates that the proximal gradient algorithm requires at most $\mathcal{O}(\epsilon^{-2})$ to find an iterate with ϵ -approximate stationarity. This bound contrasts with the bound $\mathcal{O}(\epsilon^{-3})$ of Sect. 4 for the case of general A . Moreover, the $\mathcal{O}(\epsilon^{-2})$ bound has the same order as the bound for gradient descent applied to general smooth nonconvex optimization.

6 Conclusions

We consider a standard problem formulation in which a linear transformation of the input variables is composed with a nonsmooth regularizer and added to a smooth function. In most works, the regularizer is assumed to be convex, but we extend here to the case in which it is only weakly convex. This extension allows for functions which introduce desired properties, such as sparsity, without causing a bias. [Two examples from robust statistics are minimax concave penalty (MCP) and smoothly clipped absolute deviation (SCAD)], We propose a novel method based on the variable smoothing framework and show a complexity of $\mathcal{O}(\epsilon^{-3})$ to obtain an ϵ -approximate solution. This iteration complexity falls strictly between the iteration complexity of smooth (nonconvex) problems ($\mathcal{O}(\epsilon^{-2})$) and that of the black box subgradient method for weakly convex function ($\mathcal{O}(\epsilon^{-4})$) which assumes no knowledge of the structure of the nonsmoothness.

A performance comparison between our smoothed approach and the black-box subgradient algorithm on an image denoising problem that uses an MCP total variation regularizer is shown in Figs. 1.

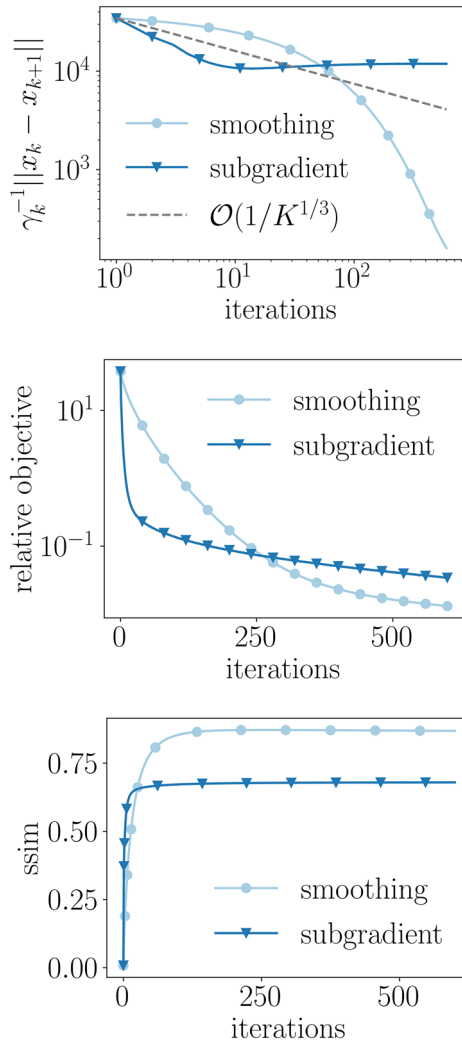


Fig. 1 Progress of our smoothing algorithm and a naive subgradient algorithm on an image denoising problem, where minimax concave penalty is used instead of the 1-norm in the anisotropic total variation, showing better performance by the smoothing approach. **Top:** The difference of consecutive iterates scaled by the inverse of the stepsize, representing the norm of the (sub)gradient used at each iteration. **Middle:** Relative difference between the objective function at the current iterate and an approximate minimum. **Bottom:** The quality of the resulting reconstruction measured via the *structural similarity index measure*, see [30]

Acknowledgements Research of the first author was supported by the doctoral program *Vienna Graduate School on Computational Optimization (VGSCO)*, FWF (Austrian Science Fund), project W 1260. Research of the second author was supported by NSF Awards 1628384, 1634597, and 1740707; Subcontract 8F-30039 from Argonne National Laboratory; and Award N660011824020 from the DARPA Lagrange Program.

Funding Open Access funding provided by University of Vienna.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

1. Chambolle, A., Pock, T.: A first-order primal-dual algorithm for convex problems with applications to imaging. *J. Math. Imaging Vis.* **40**(1), 120–145 (2011)
2. Rudin, L.I., Osher, S., Fatemi, E.: Nonlinear total variation based noise removal algorithms. *Phys. D* **60**(1–4), 259–268 (1992)
3. Boş, R.I., Böhm, A.: Variable smoothing for convex optimization problems using stochastic gradients. *J. Sci. Comput.* (2020). <https://doi.org/10.1007/s10915-020-01332-8>
4. Tran-Dinh, Q., Fercoq, O., Cevher, V.: A smooth primal-dual optimization framework for nonsmooth composite convex minimization. *SIAM J. Optim.* **28**(1), 96–134 (2018)
5. Boş, R.I., Hendrich, C.: A variable smoothing algorithm for solving convex optimization problems. *TOP* **23**(1), 124–150 (2015)
6. Beck, A., Teboulle, M.: A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM J. Imaging Sci.* **2**(1), 183–202 (2009)
7. Chambolle, A., Dossal, C.: On the convergence of the iterates of the “fast iterative shrinkage/thresholding algorithm”. *J. Optim. Theory Appl.* **166**(3), 968–982 (2015)
8. Tibshirani, R.: Regression shrinkage and selection via the LASSO. *J. R. Stat. Soc. B* **58**, 267–288 (1996)
9. Shi, W., Wahba, G., Wright, S.J., Lee, K., Klein, R., Klein, B.: LASSO-Patternsearch algorithm with application to ophthalmology data. *Stat. Interface* **1**, 137–153 (2008)
10. Vũ, B.C.: A splitting algorithm for dual monotone inclusions involving cocoercive operators. *Adv. Comput. Math.* **38**(3), 667–681 (2013)
11. Boş, R.I., Csetnek, E.R.: On the convergence rate of a forward-backward type primal-dual splitting algorithm for convex optimization problems. *Optimization* **64**(1), 5–23 (2015)
12. Boş, R.I., Hendrich, C.: Convergence analysis for a primal-dual monotone + skew splitting algorithm with applications to total variation minimization. *J. Math. Imaging Vis.* **49**(3), 551–568 (2014)
13. Chambolle, A., Ehrhardt, M.J., Richtárik, P., Schönlieb, C.B.: Stochastic primal-dual hybrid gradient algorithm with arbitrary sampling and imaging applications. *SIAM J. Optim.* **28**(4), 2783–2808 (2018)
14. Zhang, C.H.: Nearly unbiased variable selection under minimax concave penalty. *Ann. Stat.* **38**(2), 894–942 (2010)
15. Shen, X., Gu, Y.: Nonconvex sparse logistic regression with weakly convex regularization. *IEEE Trans. Signal Process.* **66**(12), 3199–3211 (2018)
16. Li, G., Pong, T.K.: Calculus of the exponent of kurdyka–łojasiewicz inequality and its applications to linear convergence of first-order methods. *Found. Comput. Math.* **18**(5), 1199–1232 (2018)
17. Bayram, I.: On the convergence of the iterative shrinkage/thresholding algorithm with a weakly convex penalty. *IEEE Trans. Signal Process.* **64**(6), 1597–1608 (2015)
18. Parekh, A., Selesnick, I.W.: Convex denoising using non-convex tight frame regularization. *IEEE Signal Process. Lett.* **22**(10), 1786–1790 (2015)
19. Fan, J.: Comments on wavelets in statistics: a review by A. Antoniadis. *J. Ital. Stat. Soc. Rev.* **6**(2), 131 (1997)
20. Beaton, A.E., Tukey, J.W.: The fitting of power series, meaning polynomials, illustrated on band-spectroscopic data. *Technometrics* **16**(2), 147–185 (1974)
21. Loh, P.L.: Statistical consistency and asymptotic normality for high-dimensional robust m -estimators. *Ann. Stat.* **45**(3), 866–896 (2017)
22. Davis, D., Drusvyatskiy, D.: Stochastic subgradient method converges at the rate $O(k^{-1/4})$ on weakly convex functions. arXiv preprint [arXiv:1802.02988](https://arxiv.org/abs/1802.02988) (2018)

23. Drusvyatskiy, D., Paquette, C.: Efficiency of minimizing compositions of convex functions and smooth maps. *Math. Program.* **178**, 503–558 (2019)
24. Lewis, A.S., Wright, S.J.: A proximal method for composite minimization. *Math. Program.* **158**(1–2), 501–546 (2016)
25. Thekumparampil, K.K., Jain, P., Netrapalli, P., Oh, S.: Efficient algorithms for smooth minimax optimization. In: *Advances in Neural Processing Systems*, pp. 12680–12691. Curran Associates (2019)
26. Davis, D., Drusvyatskiy, D.: Stochastic model-based minimization of weakly convex functions. *SIAM J. Optim.* **29**(1), 207–239 (2019)
27. Hoheisel, T., Laborde, M., Oberman, A.: On proximal point-type algorithms for weakly convex functions and their connection to the backward Euler method (2018)
28. Mordukhovich, B.S.: *Variational Analysis and Generalized Differentiation I: Basic Theory*, vol. 330. Springer, Berlin (2006)
29. Ghadimi, S., Lan, G.: Accelerated gradient methods for nonconvex nonlinear and stochastic programming. *Math. Program.* **156**(1–2), 59–99 (2016)
30. Horé, A., Ziou, D.: Image quality metrics: PSNR vs. SSIM. In: *2010 20th International Conference on Pattern Recognition*, pp. 2366–2369. IEEE (2010)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.