



Globally convergent coderivative-based generalized Newton methods in nonsmooth optimization

Pham Duy Khanh¹ · Boris S. Mordukhovich² · Vo Thanh Phat² · Dat Ba Tran²

Received: 7 September 2021 / Accepted: 4 May 2023

© Springer-Verlag GmbH Germany, part of Springer Nature and Mathematical Optimization Society 2023

Abstract

This paper proposes and justifies two globally convergent Newton-type methods to solve unconstrained and constrained problems of nonsmooth optimization by using tools of variational analysis and generalized differentiation. Both methods are coderivative-based and employ generalized Hessians (coderivatives of subgradient mappings) associated with objective functions, which are either of class $\mathcal{C}^{1,1}$, or are represented in the form of convex composite optimization, where one of the terms may be extended-real-valued. The proposed globally convergent algorithms are of two types. The first one extends the damped Newton method and requires positive-definiteness of the generalized Hessians for its well-posedness and efficient performance, while the other algorithm is of the regularized Newton-type being well-defined when the generalized Hessians are merely positive-semidefinite. The obtained convergence rates for both methods are at least linear, but become superlinear under the semismooth*

Pham Duy Khanh: Research of this author is funded by the Ministry of Education and Training Research Funding under grant B2023-SPS-02, Boris S. Mordukhovich: Research of this author was partly supported by the US National Science Foundation under Grants DMS-1808978 and DMS-2204519, by the US Air Force Office of Scientific Research under Grant #15RT0462, and by the Australian Research Council under Discovery Project DP-190100555. Vo Thanh Phat: Research of this author was partly supported by the US National Science Foundation under Grants DMS-1808978 and DMS-2204519, and by the US Air Force Office of Scientific Research under Grant #15RT0462, Dat Ba Tran: Research of this author was partly supported by the US National Science Foundation under Grant DMS-1808978 and DMS-2204519.

✉ Pham Duy Khanh
pdkhanh182@gmail.com; khanhpd@hcmue.edu.vn

✉ Boris S. Mordukhovich
aa1086@wayne.edu

Vo Thanh Phat
phatvt@wayne.edu

Dat Ba Tran
tranbadat@wayne.edu

¹ Group of Analysis and Applied Mathematics, Department of Mathematics, Ho Chi Minh City University of Education, Ho Chi Minh City, Vietnam

² Department of Mathematics, Wayne State University, Detroit, MI, USA

property of subgradient mappings. Problems of convex composite optimization are investigated with and without the strong convexity assumption on smooth parts of objective functions by implementing the machinery of forward–backward envelopes. Numerical experiments are conducted for Lasso problems and for box constrained quadratic programs with providing performance comparisons of the new algorithms and some other first-order and second-order methods that are highly recognized in nonsmooth optimization.

Keywords Nonsmooth optimization · Variational analysis · Generalized Newton methods · Global convergence · Linear and superlinear convergence rates · Convex composite optimization · Lasso problems

Mathematics Subject Classification 90C31 · 49J52 · 49J53

1 Introduction

It has been well recognized that the classical Newton method furnishes a highly efficient algorithm to solve unconstrained optimization problems of the type

$$\text{minimize } \varphi(x) \quad \text{subject to } x \in \mathbb{R}^n \quad (1.1)$$

with $\mathcal{C}^2$ -smooth objective functions φ , provided that the Hessian matrix $\nabla^2\varphi(\bar{x})$ is positive-definite at the reference solution $\bar{x}$ and the starting point x^0 is chosen sufficiently close to $\bar{x}$. In this case, the Newton iterations exhibit the local convergence with a quadratic rate; see, e.g., [2, 26, 80].

To achieve the global convergence of the Newton method, various line search algorithms are implemented by using iterative procedures given in the form

$$x^{k+1} := x^k + \tau_k d^k \quad \text{for all } k \in \mathbb{N} := \{1, 2, \dots\} \quad (1.2)$$

with a stepsize $\tau_k \geq 0$ and a search direction $d^k \neq 0$. For Newton-type methods, the search directions are chosen by solving the linear equations

$$-\nabla\varphi(x^k) = H_k d^k, \quad (1.3)$$

where $H_k := \nabla^2\varphi(x^k)$ in the classical case, while H_k is an appropriate approximation of the Hessian for various *quasi-Newton methods*. An efficient way to choose H_k is provided by the *BFGS (Broyden–Fletcher–Goldfarb–Shanno) method*; see [16, 26, 44] for more details on this and related algorithms. If $H_k = \nabla^2\varphi(x^k)$ is positive-definite, algorithm (1.2) with the *backtracking line search* is called the *damped Newton method* [2, 9] to distinguish it from the *pure Newton method*, which uses a fixed stepsize. When $\nabla^2\varphi(x^k)$ is merely positive-semidefinite, H_k in (1.3) is often taken as the regularized Hessian $H_k := \nabla^2\varphi(x^k) + \mu_k I$ with the sequence $\{\mu_k\}$ being chosen as $\mu_k := c\|\nabla\varphi(x^k)\|$ for some constant $c > 0$. The corresponding algorithm is called

the *regularized Newton method*. We refer the reader to [15, 50, 94] for many interesting results in this direction.

Among the most popular Newton-type methods to solve problems of nonsmooth optimization (1.1) with objective functions of class $\mathcal{C}^{1,1}$ (i.e., continuously differentiable with Lipschitzian gradients) is the *semismooth Newton method*. The literature on this method and its modifications is enormous; the reader is referred to the books [26, 44, 47] and the bibliographies therein for various developments and historical remarks. In fact, most of the known results address solving the equations $f(x) = 0$ with Lipschitzian vector functions f , as well as their generalized versions, to which optimization problems are reduced via stationarity conditions (observe that this is not the case of our paper). The main idea behind the semismooth Newton method is the usage of Clarke's *generalized Jacobian* of Lipschitzian mappings. In this way, local convergence results, together with some globalization procedures, were obtained for this method under the *nonsingularity* of generalized Jacobians. The reader is referred to, e.g., [39, 93] for infinite-dimensional versions of the semismooth Newton method with applications to optimization and control problems governed by partial differential equations. Other versions of Newton-type methods to solve nonsmooth equations, generalized equations, optimization and variational problems can be found in [7, 18, 20, 26, 42, 44, 47, 71, 85] among other publications. We are not in a position here to review numerous contributions to Newtonian methods that are not directly related to our paper; see more commentaries below concerning publications related to our results.

This paper develops two *globally convergent* generalized Newton algorithms to solve optimization problems (1.1) starting with the case of $\mathcal{C}^{1,1}$ objective functions (i.e., those being second-order nonsmooth) and then considering problems of *convex composite optimization* with objectives represented as sums of two convex functions such that one of them is smooth, while the other one may be extended-real-valued, which allows us to include problems of constrained optimization. The developed algorithms constitute the coderivative-based *generalized damped Newton method* (GDNM) and the coderivative-based *generalized regularized Newton method* (GRNM) for the classes of problems under consideration.

Roughly speaking, the major feature of both generalized Newton methods developed here is the replacement of the classical Hessian $\nabla^2\varphi$ of $\mathcal{C}^2$ -smooth functions by the *generalized Hessian* (or *second-order subdifferential*) $\partial^2\varphi$ of extended-real-valued, lower semicontinuous ones. This construction was introduced by Mordukhovich [57] as the *coderivative* of the subgradient mapping, while enjoying nowadays comprehensive *calculus rules* and constructive *computations* for major classes of functions that naturally appear in variational analysis, optimization, and optimal control; see Sect. 2 for more details and references. Due to such massive developments in variational analysis and its applications, coderivative-based algorithms deserve a strong attention in numerical optimization. This largely motivates our current study.

Note that coderivatives have been recently employed by Gfrerer and Outrata [32] to design a *pure* Newton-type algorithm of solving *generalized equations* and—very differently—by Mordukhovich and Sarabi [68] to find local minimizers of (1.1) that were assumed to be *tilt-stable* in the sense of Poliquin and Rockafellar [79]. The results of [68] were obtained first for $\mathcal{C}^{1,1}$ objectives and then were propagated to a general

class of *prox-regular* functions by using Moreau envelopes. The *local* superlinear convergence of these Newtonian algorithms was established in [32, 68] under the *semismooth** assumption on the mapping in question, the property introduced in [32] as a less restrictive version of semismoothness. The paper by Khanh et al. [46] developed a coderivative-based algorithm of the pure Newton type to solve *subgradient inclusions* defined by prox-regular functions with justifying the local superlinear convergence of iterates under the *semismooth** property of the corresponding subgradient mappings. The question about how to achieve the global convergence of the coderivative-based Newton methods has not been investigated in aforementioned papers and/or any other publication. In particular, the new GDNM algorithms answer in the affirmative the open question formulated in [68] on whether coderivative-based Newton methods can be globalized via a damping strategy.

As mentioned, this paper addresses not only the design and justification of new globally convergent algorithms for unconstrained problems of $\mathcal{C}^{1,1}$ optimization, but also develops such algorithms for generally *constrained* problems of convex composite optimization. For the latter class, we employ the *forward-backward envelope* (FBE), the construction that has been recently introduced in variational analysis and optimization, while has been already proved useful in constrained optimization; see, e.g., [91] with the references therein.

A central assumption in our *GDNM algorithm* for problems of $\mathcal{C}^{1,1}$ optimization is the *positive-definiteness* of the generalized Hessian $\partial^2\varphi$, which is a direct extension of that for the classical Hessian $\nabla^2\varphi$ in the damped Newton method. This assumption alone ensures that GDNM for such problems is *well-defined* and converges globally to a *tilt-stable minimizer* of (1.1) with at least some *linear rate*. The *Q-superlinear* convergence rate of GDNM is guaranteed under the *semismooth** property of the gradient mapping $\nabla\varphi$ and some relationships between parameters of the algorithm and the problem data.

The proposed *GRNM algorithm* to solve problems of $\mathcal{C}^{1,1}$ optimization does not generally require the positive-definiteness of the generalized Hessian $\partial^2\varphi$: we construct it and verify its *well-posedness* and *global convergence* to stationary points of φ under merely *positive-semidefiniteness* of the generalized Hessian. To establish results on the linear and superlinear *convergence rates* of GRNM, the *metric regularity* of the gradient mappings is additionally imposed. The latter property has been well understood, characterized, and broadly applied in variational analysis, optimization, and related areas.

Considering further problems of *convex composite optimization* in the form

$$\text{minimize } \varphi(x) := f(x) + g(x) \quad \text{subject to } x \in \mathbb{R}^n, \quad (1.4)$$

where f is a convex smooth function and $g: \mathbb{R}^n \rightarrow \overline{\mathbb{R}} := (-\infty, \infty]$ is a lower semicontinuous (l.s.c.) extended-real-valued convex function, we reduce them to $\mathcal{C}^{1,1}$ optimization by using *FBE*. Employing *second-order calculus rules* allows us to express the generalized Hessian of FBE via the problem data and relate the metric regularity and tilt stability properties of φ from (1.4) to the corresponding ones for FBE. In this way, we establish constructive results on well-posedness, global convergence, and convergence rates for both GDNM and GRNM algorithms to solve (1.4)

with and without the *strong convexity* assumption on f in a highly important case of quadratic functions in (1.4); see below.

The results established by using both GDNM and GRNM for problems of convex composite optimization with quadratic functions f in (1.4) are employed to solve a basis class of *Lasso problems*, which can be written in this form. Such problems were introduced by Tibshirani [92] motivated by applications to statistics, and since that they have been largely investigated and applied to practical models in machine learning, image processing, etc. Computing all the ingredients of both algorithms in Lasso terms, we conduct MATLAB numerical experiments by using random data sets. To compare the performance of GDNM and GRNM with other quite popular and efficient algorithms of nonsmooth optimization, we conduct parallel numerical experiments with the same Lasso data for the recent second-order *Semismooth Newton Augmented Lagrangian Methods* (SSNAL) developed in [51] and the two well-recognized first-order methods: the *Alternating Direction Method of Multipliers* (ADMM) taken from [8] and the *Fast Iterative Shrinkage-Thresholding Algorithm* (FISTA) developed in [4]. In addition, we also provide numerical experiments to solve *convex quadratic programming* problems with *box constraints*, which arise in many applications as well as subproblems of more complex optimization problems [72]. The conducted numerical experiments for this part are compared with the *trust region reflexive algorithm* [2].

The subsequent parts of the paper are organized as follows. In Sect. 2, we briefly overview the tools of variational analysis and generalized differentiation used in our algorithmic developments. Section 3 describes and justifies the coderivative-based GDNM to solve problems of $\mathcal{C}^{1,1}$ optimization. In Sect. 4, we design the coderivative-based GRNM for the same class of problems with deriving well-posedness and global convergence results. Section 5 develops both GDNM and GRNM to solve problems of convex composite optimization. Section 6 is devoted to numerical experiments for our methods and their comparison with the standard semismooth Newton method in $\mathcal{C}^{1,1}$ optimization. Then we conduct numerical experiments to employ GDNM and GRNM for solving a basic class of Lasso problems and compare the achieved numerical results with those obtained by using SSNAL, ADMM, and FISTA. This section also contains applications and numerical experiments to solve box constrained problems of quadratic programming. The concluding Sect. 7 lists the main achievements of the paper and discusses some topics of our future research. For the reader's convenience, we place several technical lemmas in the Appendix.

2 Preliminaries from variational analysis

For the reader's convenience, this section presents some preliminaries from variational analysis and generalized differentiation that are broadly employed in what follows. More details can be found in the monographs [59, 60, 88] from which we borrow the standard notation used below. Recall that $\mathbb{N} := \{1, 2, \dots\}$.

Given a set-valued mapping (multifunction) $F: \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ between finite-dimensional spaces, its (sequential Painlevé-Kuratowski) *outer limit* at $\bar{x}$ is defined by

$$\operatorname{Lim}_{x \rightarrow \bar{x}} \sup F(x) := \{y \in \mathbb{R}^n \mid \exists \text{ sequences } x_k \rightarrow \bar{x}, y_k \rightarrow y \text{ with } y_k \in F(x_k), k \in \mathbb{N}\}.$$

Using the notation $z \xrightarrow{\Omega} \bar{z}$ meaning that $z \rightarrow \bar{z}$ with $z \in \Omega$ for a given nonempty set $\Omega \subset \mathbb{R}^s$, the (Fréchet) *regular normal cone* to Ω at $\bar{z} \in \Omega$ is

$$\widehat{N}_{\Omega}(\bar{z}) := \left\{ v \in \mathbb{R}^s \mid \limsup_{z \xrightarrow{\Omega} \bar{z}} \frac{\langle v, z - \bar{z} \rangle}{\|z - \bar{z}\|} \leq 0 \right\},$$

while the (Mordukhovich) *limiting normal cone* to Ω at $\bar{z} \in \Omega$ is defined by

$$\begin{aligned} N_{\Omega}(\bar{z}) &:= \operatorname{Lim}_{z \xrightarrow{\Omega} \bar{z}} \sup \widehat{N}_{\Omega}(z) \\ &= \left\{ v \in \mathbb{R}^s \mid \exists z_k \xrightarrow{\Omega} \bar{z}, v_k \rightarrow v \text{ as } k \rightarrow \infty \text{ with } v_k \in \widehat{N}_{\Omega}(z_k) \right\}. \end{aligned} \quad (2.1)$$

The corresponding limiting *coderivative* of $F: \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ at $(\bar{x}, \bar{y}) \in \operatorname{gph} F$ is defined by

$$D^*F(\bar{x}, \bar{y})(v) := \{u \in \mathbb{R}^n \mid (u, -v) \in N_{\operatorname{gph} F}(\bar{x}, \bar{y})\}, \quad v \in \mathbb{R}^m, \quad (2.2)$$

where $\operatorname{gph} F := \{(x, y) \in \mathbb{R}^n \times \mathbb{R}^m \mid y \in F(x)\}$, and where $\bar{y}$ is omitted in the coderivative notation if $F(\bar{x})$ is a singleton. If $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$ is a single-valued mapping which is $\mathcal{C}^1$ -smooth around $\bar{x}$, then

$$D^*F(\bar{x})(v) = \{\nabla F(\bar{x})^*v\} \text{ for all } v \in \mathbb{R}^m$$

via the adjoint/transpose Jacobian matrix $\nabla F(\bar{x})^*$. The defined coderivative of general multifunctions satisfies comprehensive *calculus rules* based on *variational/extremal principles* of variational analysis. Among the most impressive and useful advantages of the coderivative (2.2) are complete characterizations in its terms the fundamental well-posedness properties (metric regularity, linear openness, and Lipschitzian behavior) of general multifunctions that were developed in [58] and were labeled in [88] as the *Mordukhovich criteria*. In this paper we employ these characterizations for the property of *metric regularity* of $F: \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ around $(\bar{x}, \bar{y}) \in \operatorname{gph} F$ meaning that there exist a number $\mu > 0$ and neighborhoods U of $\bar{x}$ and V of $\bar{y}$ such that

$$\operatorname{dist}(x; F^{-1}(y)) \leq \mu \operatorname{dist}(y; F(x)) \text{ for all } (x, y) \in U \times V, \quad (2.3)$$

where $F^{-1}(y) := \{x \in \mathbb{R}^n \mid y \in F(x)\}$, and where ‘dist’ stands for the distance between a point and a set. If in addition F^{-1} has a single-valued localization around $(\bar{y}, \bar{x})$, i.e., there exist neighborhoods U of $\bar{x}$ and V of $\bar{y}$ together with a single-valued mapping $\vartheta: V \rightarrow U$ such that $\operatorname{gph} F^{-1} \cap (V \times U) = \operatorname{gph} \vartheta$, then F is *strongly metrically regular* around $(\bar{x}, \bar{y})$ with modulus $\mu > 0$. The aforementioned coderivative characterization from [58, Theorem 3.6] tells us that, whenever F is closed-graph

around $(\bar{x}, \bar{y}) \in \text{gph } F$, its metric regularity around this point is equivalent to the implication

$$[v \in \mathbb{R}^m, 0 \in D^*F(\bar{x}, \bar{y})(v)] \implies v = 0. \quad (2.4)$$

Moreover, the *exact regularity bound* of F at $(\bar{x}, \bar{y})$, i.e., the infimum of all $\mu > 0$ such that (2.3) holds for some neighborhoods U and V , is calculated by

$$\text{reg } F(\bar{x}, \bar{y}) = \|D^*F(\bar{x}, \bar{y})^{-1}\| = \|D^*F^{-1}(\bar{y}, \bar{x})\| \quad (2.5)$$

via the norm of the coderivatives of F and F^{-1} as positive homogeneous multifunctions; see [59, 60, 88] for more discussions, different proofs, and infinite-dimensional extensions.

Another notion in variational analysis used in what follows concerns a strong version of local monotonicity for set-valued mappings. We say that $T: \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ is *strongly locally monotone* with modulus $\kappa > 0$ around $(\bar{x}, \bar{y}) \in \text{gph } T$ if there exist neighborhoods U of $\bar{x}$ and V of $\bar{y}$ such that

$$\langle x - u, v - w \rangle \geq \kappa \|x - u\|^2 \quad \text{for all } (x, v), (u, w) \in \text{gph } T \cap (U \times V).$$

If in addition $\text{gph } T \cap (U \times V) = \text{gph } S \cap (U \times V)$ for any monotone operator $S: \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ satisfying $\text{gph } T \cap (U \times V) \subset \text{gph } S$, then T is *strongly locally maximal monotone* with modulus $\kappa > 0$ around $(\bar{x}, \bar{y})$. The reader is referred to [62] and [60, Sect. 5.2] for coderivative characterizations of the latter property, which is significantly more relaxed than the strong metric regularity of T around $(\bar{x}, \bar{y})$.

Next we consider an extended-real-valued function $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ with the domain and epigraph

$$\text{dom } \varphi := \{x \in \mathbb{R}^n \mid \varphi(x) < \infty\}, \quad \text{epi } \varphi := \{(x, \alpha) \in \mathbb{R}^{n+1} \mid \alpha \geq \varphi(x)\}.$$

Recall that an l.s.c. proper function $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ is *strongly convex* on a convex set $\Omega \subset \mathbb{R}^n$ with modulus $\kappa > 0$ if the quadratically shifted function $\varphi - (\kappa/2) \|\cdot\|^2$ is convex on Ω .

The (limiting) *subdifferential* of φ at $\bar{x} \in \text{dom } \varphi$ is defined geometrically by

$$\partial\varphi(\bar{x}) := \{v \in \mathbb{R}^n \mid (v, -1) \in N_{\text{epi } \varphi}(\bar{x}, \varphi(\bar{x}))\} \quad (2.6)$$

via the limiting normal cone (2.1), while admitting various analytic representations and satisfying comprehensive calculus rules that can be found in [59, 60, 88]. Observe the useful scalarization formula

$$D^*F(\bar{x})(v) = \partial\langle v, F \rangle(\bar{x}) \quad \text{for all } v \in \mathbb{R}^m \quad (2.7)$$

connecting the coderivative (2.2) of a locally Lipschitzian mapping $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$ and the subdifferential (2.6) of the function $x \mapsto \langle v, F \rangle(x)$ whenever $v \in \mathbb{R}^m$.

Following [57], we define the *second-order subdifferential*, or *generalized Hessian*, $\partial^2\varphi(\bar{x}, \bar{v}): \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ of $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ at $\bar{x} \in \text{dom } \varphi$ for $\bar{v} \in \partial\varphi(\bar{x})$ as the coderivative of the subgradient mapping

$$\partial^2\varphi(\bar{x}, \bar{v})(u) := (D^*\partial\varphi)(\bar{x}, \bar{v})(u) \quad \text{for all } u \in \mathbb{R}^n. \quad (2.8)$$

If φ is $\mathcal{C}^2$ -smooth around $\bar{x}$, then we have

$$\partial^2\varphi(\bar{x})(u) = \{\nabla^2\varphi(\bar{x})u\} \quad \text{for all } u \in \mathbb{R}^n, \quad (2.9)$$

while for φ of class $\mathcal{C}^{1,1}$ around $\bar{x}$, we get by the scalarization formula (2.7) that

$$\partial^2\varphi(\bar{x})(u) = \partial\langle u, \nabla\varphi \rangle(\bar{x}) \quad \text{for all } u \in \mathbb{R}^n. \quad (2.10)$$

As follows from (2.10), calculus rules and computations of the second-order subdifferential for $\mathcal{C}^{1,1}$ functions reduce in fact to those for the first-order construction (2.6). We also have well-developed second-order calculus rules for (2.8) for rather general classes of extended-real-valued functions; see, e.g., [59, 60, 65] with many additional references. Furthermore, the second-order subdifferential has been computed and analyzed in terms of the given data for broad classes of structural functional systems appearing in numerous aspects of variational analysis, optimization, stability, and optimal control among other areas, with subsequent applications to optimality conditions, sensitivity analysis, numerical algorithms, stochastic programming, electricity markets, etc. The reader can find more information in, e.g., [11, 13, 17, 19, 35, 36, 57, 59, 60, 63–66, 87, 95] along with other publications on such developments and related topics of second-order variational analysis. Some new results in this direction are presented in what follows.

In this paper, we use the fundamental notion of *tilt-stable local minimizers* and its second-order characterizations for the justification of the proposed Newton-type algorithms.

Definition 1 (Tilt-stable local minimizers) Given $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$, a point $\bar{x} \in \text{dom } \varphi$ is a TILT-STABLE LOCAL MINIMIZER of φ if there exists a number $\gamma > 0$ such that the mapping

$$M_\gamma: v \mapsto \operatorname{argmin}\{\varphi(x) - \langle v, x \rangle \mid x \in \mathbb{B}_\gamma(\bar{x})\}$$

is single-valued and Lipschitz continuous on some neighborhood of $0 \in \mathbb{R}^n$ with $M_\gamma(0) = \{\bar{x}\}$. By a MODULUS of tilt stability of φ at $\bar{x}$ we understand a Lipschitz constant of M_γ around the origin.

This notion was introduced by Poliquin and Rockafellar in [79] and characterized there via $\partial^2\varphi$ for a broad class of prox-regular functions $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ that are overwhelmingly involved in second-order variational analysis. More recently, developing second-order subdifferential calculus and second-order growth conditions made it possible to establish complete characterizations of tilt-stable local minimizers for

various classes of problems in constrained optimization including nonlinear programming, extended nonlinear programming, composite optimization, minimax problems, second-order cone programming, semidefinite programming, etc.; see, e.g., [12, 21, 22, 31, 60, 61, 65] among other publications on tilt stability in optimization.

Finally in this section, we recall for completeness the notions of convergence rates used for our algorithms.

Definition 2 (*Rates of convergence*) Let $\{x^k\} \subset \mathbb{R}^n$ be a sequence of vectors converging to $\bar{x}$ as $k \rightarrow \infty$ with $\bar{x} \neq x^k$ for all $k \in \mathbb{N}$. The convergence rate is said to be:

(i) R- LINEAR if we have

$$0 < \limsup_{k \rightarrow \infty} \left(\|x^k - \bar{x}\| \right)^{1/k} < 1,$$

i.e., there exist $\mu \in (0, 1)$, $c > 0$, and $k_0 \in \mathbb{N}$ such that

$$\|x^k - \bar{x}\| \leq c\mu^k \quad \text{for all } k \geq k_0.$$

(ii) Q- LINEAR if we have

$$\limsup_{k \rightarrow \infty} \frac{\|x^{k+1} - \bar{x}\|}{\|x^k - \bar{x}\|} < 1,$$

i.e., there exist $\mu \in (0, 1)$ and $k_0 \in \mathbb{N}$ such that

$$\|x^{k+1} - \bar{x}\| \leq \mu \|x^k - \bar{x}\| \quad \text{for all } k \geq k_0.$$

(iii) Q- SUPERLINEAR if we have

$$\lim_{k \rightarrow \infty} \frac{\|x^{k+1} - \bar{x}\|}{\|x^k - \bar{x}\|} = 0.$$

3 Coderivative-based damped Newton method in $\mathcal{C}^{1,1}$ optimization

In this section, we concentrate on the unconstrained optimization problem (1.1), where the cost function $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}$ is of class $\mathcal{C}^{1,1}$. This kind of problem plays a crucial role not only in numerical optimization but also in applied areas including, e.g., machine learning. In particular, $\mathcal{C}^{1,1}$ optimization problems also arise frequently as *subproblems in augmented Lagrangian methods* [34, 86, 87]. Furthermore, L2-loss support vector regression problems are important classes of problems that are $\mathcal{C}^{1,1}$ unconstrained optimization problem [41]. More practical examples about $\mathcal{C}^{1,1}$ functions can be found in the paper [40]. A coderivative-based generalization of the pure Newton method to solve (1.1) *locally* was first suggested and investigated in [68] under the major assumption that a given point $\bar{x}$ is a *tilt-stable* local minimizer of (1.1). Then it

was extended in [46] to solve directly the gradient system $\nabla\varphi(x) = 0$ under certain assumptions on a given solution $\bar{x}$ of the gradient equation ensuring the well-posedness and local superlinear convergence of the algorithm. One of the serious disadvantages of the pure Newton method and its generalizations is that the corresponding sequence of iterates may not converge if the starting point is not sufficiently close to the solution. This motivates us to design and justify a *globally* convergent *damped Newton* counterpart of the generalized pure Newton algorithms from [46, 68] with *backtracking line search* to solve (1.1). Here is the algorithm.

Algorithm 1 Coderivative-based damped Newton algorithm for $\mathcal{C}^{1,1}$ functions

Input: $x^0 \in \mathbb{R}^n$, $\sigma \in (0, \frac{1}{2})$, $\beta \in (0, 1)$
1: **for** $k = 0, 1, \dots$ **do**
2: If $\nabla\varphi(x^k) = 0$, stop; otherwise go to the next step
3: Choose $d^k \in \mathbb{R}^n$ such that $-\nabla\varphi(x^k) \in \partial^2\varphi(x^k)(d^k)$
4: Set $\tau_k = 1$
5: **while** $\varphi(x^k + \tau_k d^k) > \varphi(x^k) + \sigma \tau_k \langle d^k, \nabla\varphi(x^k) \rangle$ **do**
6: set $\tau_k := \beta \tau_k$
7: **end while**
8: Set $x^{k+1} := x^k + \tau_k d^k$
9: **end for**

If φ is $\mathcal{C}^2$ -smooth, Algorithm 1 reduces to the standard damped Newton method (as, e.g., in [2, 9]) due to (2.9). In the general case of $\varphi \in \mathcal{C}^{1,1}$, it follows from (2.2) that the direction d^k in Step 3 of Algorithm 1 can be explicitly found from the inclusion

$$(-\nabla\varphi(x^k), -d^k) \in N((x^k, \nabla\varphi(x^k)); \text{gph } \nabla\varphi).$$

Note also that, due to the scalarization formula (2.10), the Newton equation in Step 3 of Algorithm 1 can be equivalently written in the form

$$-\nabla\varphi(x^k) \in \partial\langle d^k, \nabla\varphi \rangle(x^k), \quad (3.1)$$

which merely requires the first-order subdifferential computation.

We start justifying Algorithm 1 with the verification of its *well-posedness*. The following proposition establishes the existence of *descent* Newton directions under the *positive-definiteness* of $\partial^2\varphi(x)$.

Proposition 1 (existence of Newton directions and descent property) *Let $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}$ be of class $\mathcal{C}^{1,1}$ on $\mathbb{R}^n$. Suppose that $\nabla\varphi(x) \neq 0$ and that $\partial^2\varphi(x)$ is positive-definite, i.e.,*

$$\langle v, u \rangle > 0 \text{ for all } v \in \partial^2\varphi(x)(u) \text{ and } u \neq 0. \quad (3.2)$$

Then there exists a nonzero direction $d \in \mathbb{R}^n$ such that

$$-\nabla\varphi(x) \in \partial^2\varphi(x)(d). \quad (3.3)$$

Moreover, every such direction satisfies the inequality $\langle \nabla \varphi(x), d \rangle < 0$. Consequently, for each $\sigma \in (0, 1)$ and $d \in \mathbb{R}^n$ satisfying (3.3) we find $\delta > 0$ such that

$$\varphi(x + \tau d) \leq \varphi(x) + \sigma \tau \langle \nabla \varphi(x), d \rangle \text{ whenever } \tau \in (0, \delta). \quad (3.4)$$

Proof By the positive-definiteness of $\partial^2 \varphi(x)$, it follows from [60, Theorem 5.16] that $\nabla \varphi$ is strongly locally maximal monotone around $(x, \nabla \varphi(x))$. Thus $\nabla \varphi$ is strongly metrically regular around $(x, \nabla \varphi(x))$ by [60, Theorem 5.13]. Using [46, Corollary 4.2] yields the existence of $d \in \mathbb{R}^n$ with $-\nabla \varphi(x) \in \partial^2 \varphi(x)(d)$. To verify further that $d \neq 0$, suppose on the contrary that $d = 0$. Since $\nabla \varphi$ is locally Lipschitz around x , it follows from [59, Theorem 1.44] that

$$-\nabla \varphi(x) \in \partial^2 \varphi(x)(0) = (D^* \nabla \varphi)(x)(0) = \{0\},$$

which contradicts the assumption that $\nabla \varphi(x) \neq 0$. Employing again the positive-definiteness of $\partial^2 \varphi$ tells us that $\langle \nabla \varphi(x), d \rangle < 0$. Using [44, Lemmas 2.18 and 2.19], we arrive at (3.4) and thus complete the proof. $\square$

The next theorem establishes the *global linear convergence* of Algorithm 1 to a *tilt-stable minimizer* of (1.1) under the positive-definiteness assumption on the generalized Hessian $\partial^2 \varphi$.

Theorem 1 (global linear convergence of the coderivative-based damped Newton algorithm for $\mathcal{C}^{1,1}$ functions) *Let $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}$ be of class $\mathcal{C}^{1,1}$, and let $x^0 \in \mathbb{R}^n$ be an arbitrary point such that the generalized Hessian $\partial^2 \varphi(x)$ is positive-definite for all $x \in \Omega$, where*

$$\Omega := \{x \in \mathbb{R}^n \mid \varphi(x) \leq \varphi(x^0)\}. \quad (3.5)$$

Then Algorithm 1 either stops after finitely many iterations, or produces a sequence $\{x^k\} \subset \Omega$ such that $\{\varphi(x^k)\}$ is monotonically decreasing. Moreover, if the iterative sequence $\{x^k\}$ has an accumulation point $\bar{x}$ (in particular, when the level set Ω from (3.5) is bounded), then $\{x^k\}$ converges to $\bar{x}$, which is a tilt-stable local minimizer of φ . In this case, we have:

- (i) *The convergence rate of $\{\varphi(x^k)\}$ is at least Q -linear.*
- (ii) *The convergence rates of $\{x^k\}$ and $\{\|\nabla \varphi(x^k)\|\}$ are at least R -linear.*

Proof Proposition 1 easily ensures by induction that Algorithm 1 either stops after finitely many iterations, or produces a sequence $\{x^k\} \subset \Omega$ such that $\varphi(x^{k+1}) < \varphi(x^k)$ for all $k \in \mathbb{N}$. Suppose next that $\{x^k\}$ has an accumulation point $\bar{x}$. Since the set Ω is closed, we get $\bar{x} \in \Omega$, and hence have that $\partial^2 \varphi(\bar{x})$ is positive-definite. Then [11, Proposition 4.6] gives us positive numbers κ and δ such that

$$\langle z, w \rangle \geq \kappa \|w\|^2 \quad \text{for all } z \in \partial^2 \varphi(x)(w), \quad x \in \mathbb{B}_\delta(\bar{x}), \quad \text{and } w \in \mathbb{R}^n. \quad (3.6)$$

Since φ is of class $\mathcal{C}^{1,1}$ around $\bar{x}$, we get without loss of generality that $\nabla\varphi$ is Lipschitz continuous on $\mathbb{B}_\delta(\bar{x})$ with some constant $\ell > 0$. By [59, Theorem 1.44] we have

$$\|z\| \leq \ell\|w\| \quad \text{for all } z \in \partial^2\varphi(x)(w), \quad x \in \mathbb{B}_\delta(\bar{x}), \quad \text{and } w \in \mathbb{R}^n. \quad (3.7)$$

The rest of the proof is split into the following four claims.

Claim 1: *For any subsequence $\{x^{k_j}\}$ of $\{x^k\}$ such that $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, the corresponding sequence $\{\tau_{k_j}\}$ in Algorithm 1 is bounded from below by some $\gamma > 0$, the corresponding sequence $\{d^{k_j}\}$ is bounded, and we have*

$$\langle -\nabla\varphi(x^{k_j}), d^{k_j} \rangle \geq \kappa\|d^{k_j}\|^2, \quad (3.8)$$

$$\|\nabla\varphi(x^{k_j})\| \leq \ell\|d^{k_j}\|, \quad (3.9)$$

$$\varphi(x^{k_j}) - \varphi(x^{k_j+1}) \geq \sigma\gamma\kappa\|d^{k_j}\|^2, \quad (3.10)$$

for all large $j \in \mathbb{N}$. Since $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$ and $-\nabla\varphi(x^{k_j}) \in \partial^2\varphi(x^{k_j})(d^{k_j})$ for all $j \in \mathbb{N}$, we obtain (3.8) and (3.9) from (3.6) and (3.7), respectively. The Cauchy–Schwarz inequality yields $\|\nabla\varphi(x^{k_j})\| \geq \kappa\|d^{k_j}\|$ for such j . Since $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, the latter estimate verifies the boundedness of the sequence of directions $\{d^{k_j}\}$. It remains to show that $\{\tau_{k_j}\}_{j \in \mathbb{N}}$ is bounded from below by a positive number. Indeed, supposing on the contrary that the opposite holds and combining this with $\tau_k \geq 0$ give us a subsequence of $\{\tau_{k_j}\}$ that converges to 0. Assume without loss of generality that $\tau_{k_j} \rightarrow 0$ as $j \rightarrow \infty$. Thus $x^{k_j} + \tau_{k_j}d^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, and hence $x^{k_j} + \tau_{k_j}d^{k_j} \in \text{int } \mathbb{B}_\delta(\bar{x})$ whenever j is sufficiently large. Applying Lemma 3 from the Appendix, we have

$$\beta^{-1}\tau_{k_j} > \frac{2(\sigma - 1)\langle \nabla\varphi(x^{k_j}), d^{k_j} \rangle}{\ell\|d^{k_j}\|^2} \geq \frac{2(1 - \sigma)\kappa}{\ell},$$

where the second inequality follows from (3.8). Letting $j \rightarrow \infty$ gives us $\sigma \geq 1$, a contradiction due to the choice of σ . Hence there exists $\gamma > 0$ such that $\tau_{k_j} \geq \gamma$ for all $j \in \mathbb{N}$. Moreover, using the estimate in (3.8) allows us to find $j_0 \in \mathbb{N}$ such that

$$\varphi(x^{k_j}) - \varphi(x^{k_j+1}) \geq \sigma\tau_{k_j}\langle -\nabla\varphi(x^{k_j}), d^{k_j} \rangle \geq \sigma\gamma\kappa\|d^{k_j}\|^2 \quad \text{for all } j \geq j_0, \quad (3.11)$$

which therefore justifies Claim 1.

Claim 2: $\bar{x}$ is a tilt-stable local minimizer of φ . To verify this, we only need to show that $\bar{x}$ is a stationary point of φ , by taking into account the positive-definiteness of $\partial^2\varphi(\bar{x})$ and the second-order characterization of tilt-stability from [79, Theorem 1.3]. Since $\bar{x}$ is an accumulation point of $\{x^k\}$, there exists a subsequence $\{x^{k_j}\}_{j \in \mathbb{N}}$ of $\{x^k\}$ such that $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$. Due to Claim 1, we find $\gamma > 0$ such that (3.10) is satisfied. Since the sequence $\{\varphi(x^k)\}$ is nonincreasing and since $\varphi(\bar{x})$ is an accumulation point of $\{\varphi(x^k)\}$, the sequence $\{\varphi(x^k)\}$ must converge to $\varphi(\bar{x})$ as $k \rightarrow \infty$. Letting $j \rightarrow \infty$ in the inequality (3.10), we have $\|d^{k_j}\| \rightarrow 0$ as $j \rightarrow \infty$. Passing to the limit as

$j \rightarrow \infty$ in the inequality $\|\nabla\varphi(x^{k_j})\| \leq \ell\|d^{k_j}\|$ for all large j from (3.9) tells us that $\nabla\varphi(\bar{x}) = 0$, which readily justifies Claim 2.

Claim 3: *The iterative sequence $\{x^k\}$ is convergent.* To verify this, we use Ostrowski's condition from [26, Proposition 8.3.10]. First we show that no other accumulation point of $\{x^k\}$ exists in $\mathbb{B}_\delta(\bar{x})$. Assuming the contrary, find $\tilde{x} \in \mathbb{B}_\delta(\bar{x})$ such that $\tilde{x} \neq \bar{x}$ and that $\tilde{x}$ is an accumulation point of $\{x^k\}$. Arguing similarly to Claim 2 tells us that $\tilde{x}$ is a tilt-stable local minimizer of φ , which contradicts the strong convexity of φ on $\mathbb{B}_\delta(\bar{x})$. Supposing next that $\{x^{k_j}\}$ is an arbitrary subsequence of $\{x^k\}$ with $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, we need to check that

$$\lim_{j \rightarrow \infty} \|x^{k_j+1} - x^{k_j}\| = 0. \quad (3.12)$$

Indeed, Claim 1 gives us $\gamma > 0$ such that (3.10) holds, which implies in turn that

$$\|x^{k_j+1} - x^{k_j}\|^2 = \tau_{k_j}^2 \|d^{k_j}\|^2 \leq \|d^{k_j}\|^2 \leq \frac{1}{\sigma\gamma\kappa} \left(\varphi(x^{k_j}) - \varphi(x^{k_j+1}) \right) \rightarrow 0$$

and hence verifies (3.12). Finally, it follows from [26, Proposition 8.3.10] that the sequence $\{x^k\}$ converges to $\bar{x}$ as $k \rightarrow \infty$, which therefore completes the proof of Claim 3.

Claim 4: *The convergence rate of $\{\varphi(x^k)\}$ is at least Q -linear, while the convergence rates of $\{x^k\}$ and $\{\|\nabla\varphi(x^k)\|\}$ are at least R -linear.* Indeed, the strong convexity of φ on $\mathbb{B}_\delta(\bar{x})$ shows that

$$\langle \nabla\varphi(x) - \nabla\varphi(u), x - u \rangle \geq \kappa \|x - u\|^2 \quad \text{for all } x, u \in \mathbb{B}_\delta(\bar{x}). \quad (3.13)$$

Since $x^k \rightarrow \bar{x}$ as $k \rightarrow \infty$, we have that $x^k \in U$ for all $k \in \mathbb{N}$ sufficiently large. Substituting $x := x^k$ and $u := \bar{x}$ into (3.13) and then using the Cauchy–Schwarz inequality together with the stationarity condition $\nabla\varphi(\bar{x}) = 0$ give us the lower estimate

$$\|\nabla\varphi(x^k)\| \geq \kappa \|x^k - \bar{x}\| \quad (3.14)$$

for large k . The local Lipschitz continuity of $\nabla\varphi$ around $\bar{x}$ and the result of [44, Lemma A.11] ensure the existence of $\ell > 0$ such that

$$\varphi(x^k) - \varphi(\bar{x}) = |\varphi(x^k) - \varphi(\bar{x}) - \langle \nabla\varphi(\bar{x}), x^k - \bar{x} \rangle| \leq \frac{\ell}{2} \|x^k - \bar{x}\|^2 \quad \text{for large } k. \quad (3.15)$$

Furthermore, estimate (3.7) together with the inclusion $-\nabla\varphi(x^k) \in \partial^2\varphi(x^k)(d^k)$ implies that

$$\|\nabla\varphi(x^k)\| \leq \ell\|d^k\| \quad \text{for large } k. \quad (3.16)$$

Claim 1 tells us that the sequence $\{\tau_k\}$ is bounded from below by some constant $\gamma > 0$ such that

$$\varphi(x^k) - \varphi(x^{k+1}) \geq \sigma \gamma \kappa \|d^k\|^2 \quad \text{for large } k.$$

Combining the above inequality with (3.16) yields the estimates

$$\varphi(x^k) - \varphi(x^{k+1}) \geq \sigma \gamma \kappa \ell^{-2} \|\nabla \varphi(x^k)\|^2 \quad (3.17)$$

if k is large. Using (3.14), (3.15), (3.17) and then applying Lemma 4 from the Appendix with

$$\alpha_k := \varphi(x^k) - \varphi(\bar{x}), \quad \beta_k := \|\nabla \varphi(x^k)\|, \quad \gamma_k := \|x^k - \bar{x}\|$$

$c_1 := \sigma \gamma \kappa \ell^{-2}$, $c_2 := \kappa$, and $c_3 = \ell/2$ verify Claim 4 and thus completes the proof of the theorem. $\square$

Remark 1 (on proof of Theorem 1) Following the suggestion of the referee, we present an alternative proof of Claim 2 in Theorem 1 by assuming the contrary and using the result of *gradient related property* introduced in [6]. Indeed, suppose that $\nabla \varphi(\bar{x}) \neq 0$, we check that the sequence $\{d^k\}$ is *gradient related* to $\{x^k\}$ in the sense of [6], i.e., for any subsequence $\{x^{k_j}\}$ that converges to $\bar{x}$ as $j \rightarrow \infty$, the corresponding subsequence $\{d^{k_j}\}$ is bounded and satisfies the inequality

$$\limsup_{j \rightarrow \infty} \langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle < 0. \quad (3.18)$$

The boundedness of $\{d^{k_j}\}$ is clarified in Claim 1. Combining (3.8) and (3.9), we deduce that

$$\langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle \leq -\kappa \|d^{k_j}\|^2 \leq -\kappa \ell^{-2} \|\nabla \varphi(x^{k_j})\|^2, \quad \text{for large } j \in \mathbb{N}.$$

Taking the the upper limit above and using $\nabla \varphi(\bar{x}) \neq 0$ justifies (3.18). Therefore, it follows from [6, Proposition 1.2.1] that $\nabla \varphi(\bar{x}) = 0$. This is a contradiction, which verifies Claim 2.

Our next goal is to establish the *Q-superlinear convergence* of Algorithm 1. First recall additional notions of variational analysis and generalized differentiation needed for these developments. A highly recognized concept used for Newton-type methods addressing single-valued Lipschitz continuous mappings is semismoothness. A mapping $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ is *semismooth* at $\bar{x}$ if it is locally Lipschitzian around this point and the limit

$$\lim_{\substack{A \in \text{co} \overline{\nabla} f(\bar{x} + tu') \\ u' \rightarrow u, t \downarrow 0}} Au' \quad (3.19)$$

exists for all $u \in \mathbb{R}^n$, where ‘co’ stands for the convex hull of a set, and where $\overline{\nabla} f$ is defined by

$$\overline{\nabla} f(x) := \left\{ A \in \mathbb{R}^{m \times n} \mid \exists x_k \xrightarrow{\Omega_f} x \text{ such that } \nabla f(x_k) \rightarrow A \right\}, \quad x \in \mathbb{R}^n$$

with $\Omega_f := \{x \in \mathbb{R}^n \mid f \text{ is differentiable at } x\}$; see [26, 44, 47, 83] for further discussions. Note that any semismooth mapping automatically admits the classical directional derivative at the reference point.

The next proposition about the acceptance of the unit stepsize under semismoothness follows from a close look at the proof of [25, Theorem 3.3].

Proposition 2 (acceptance of unit stepsize under semismoothness) *Suppose that a function $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}$ is $\mathcal{C}^1$ -smooth around $\bar{x}$ with $\nabla \varphi(\bar{x}) = 0$, and that $\nabla \varphi$ is semismooth at this point. Let a sequence $\{x^k\}$ converge to $\bar{x}$ with $x^k \neq \bar{x}$ as $k \in \mathbb{N}$, and let a sequence $\{d^k\}$ satisfy the condition*

$$\|x^k + d^k - \bar{x}\| = o(\|x^k - \bar{x}\|). \quad (3.20)$$

Suppose further that there exists $\kappa > 0$ such that $\langle \nabla \varphi(x^k), d^k \rangle \leq \kappa^{-1} \|d^k\|^2$ whenever $k \in \mathbb{N}$ is sufficiently large. Then for any $\sigma \in (0, 1/2)$ we have the estimate

$$\varphi(x^k + d^k) \leq \varphi(x^k) + \sigma \langle \nabla \varphi(x^k), d^k \rangle \text{ for all large } k. \quad (3.21)$$

Quite recently [32], the concept of semismoothness has been improved and extended to set-valued mappings. To formulate the latter notion, recall first the construction of the *directional limiting normal cone* to a set $\Omega \subset \mathbb{R}^s$ at $\bar{z} \in \Omega$ in the direction $d \in \mathbb{R}^s$ introduced in [33] by

$$N_\Omega(\bar{z}; d) := \{v \in \mathbb{R}^s \mid \exists t_k \downarrow 0, d_k \rightarrow d, v_k \rightarrow v \text{ with } v_k \in \widehat{N}_\Omega(\bar{z} + t_k d_k)\}. \quad (3.22)$$

It is obvious that (3.22) agrees with the limiting normal cone (2.1) for $d = 0$. The *directional limiting coderivative* of $F : \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ at $(\bar{x}, \bar{y}) \in \text{gph } F$ in the direction $(u, v) \in \mathbb{R}^n \times \mathbb{R}^m$ is defined in [30] by

$$\begin{aligned} D^* F((\bar{x}, \bar{y}); (u, v))(v^*) \\ := \{u^* \in \mathbb{R}^n \mid (u^*, -v^*) \in N_{\text{gph } F}((\bar{x}, \bar{y}); (u, v))\} \text{ for all } v^* \in \mathbb{R}^m. \end{aligned} \quad (3.23)$$

Using (3.23), we come to the aforementioned property of set-valued mappings introduced in [32].

Definition 3 (emismooth* property of set-valued mappings) A mapping $F : \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ is SEMISMOOTH* at $(\bar{x}, \bar{y}) \in \text{gph } F$ if whenever $(u, v) \in \mathbb{R}^n \times \mathbb{R}^m$ we have

$$\langle u^*, u \rangle = \langle v^*, v \rangle \text{ for all } (v^*, u^*) \in \text{gph } D^* F((\bar{x}, \bar{y}); (u, v)).$$

Among various properties of semismooth* mappings obtained in [32], recall that this property holds if the graph of $F: \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ is represented as a union of finitely many closed and convex sets, as well as for the normal cone mappings generated by convex polyhedral sets. Note also that the semismooth* property of single-valued locally Lipschitzian mappings $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ around $\bar{x}$ agrees with the semismooth property (3.19) at this point provided that f is *directionally differentiable* at $\bar{x}$; see [32, Corollary 3.8]. Although the standard semismooth property of locally Lipschitzian and directionally differentiable mappings has been conventionally used in the Newton method literature, some important results were obtained without the directional differentiability assumption; see, e.g., Meng et al. [53]. Such a relaxed semismooth property of single-valued locally Lipschitzian mappings is known as *G-semismoothness*. Note that, in contrast to *G-semismoothness*, the semismooth* property is defined for arbitrary set-valued mappings, and it is used for subgradient ones in this paper; see Sect. 5. But even for single-valued Lipschitzian mappings, the semismooth* definition based on coderivatives may have some advantages in comparison with the *G-semismooth* one due to comprehensive coderivative calculus rules. More recent results on the semismooth* property can be found in [27].

The next lemma discusses the acceptance of the unit stepsize for $\mathcal{C}^{1,1}$ functions with semismooth* gradients. The obtained estimates are of their own interest, while are instrumental to establish major superlinear convergence results in this and subsequent sections.

Lemma 1 (acceptance of unit stepsize under semismoothness*) *Let $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}$ be a $\mathcal{C}^1$ -smooth function around $\bar{x} \in \mathbb{R}^n$ with $\nabla\varphi(\bar{x}) = 0$. Suppose that $\nabla\varphi$ is locally Lipschitzian around $\bar{x}$ with modulus $\ell > 0$, and that $\nabla\varphi$ is semismooth* at this point. Take a sequence $\{x^k\}$ converging to $\bar{x}$ with $x^k \neq \bar{x}$ as $k \in \mathbb{N}$, and let a sequence $\{d^k\}$ satisfy condition (3.20). Assume also that there exists $\kappa > 0$ such that*

$$\varphi(x^k + d^k) - \varphi(x^k) \leq \langle \nabla\varphi(x^k + d^k), d^k \rangle - \frac{1}{2\kappa} \|d^k\|^2 \quad (3.24)$$

whenever k is sufficiently large. Then for any $\sigma \in (0, 1/(2\ell\kappa))$ we have estimate (3.21).

Proof Having (3.24) with $\sigma \in (0, 1/(2\ell\kappa))$ and using (3.20) together with [26, Lemma 7.5.7], we get

$$\lim_{k \rightarrow \infty} \|x^k - \bar{x}\|/\|d^k\| = 1, \quad (3.25)$$

which also yields the limiting relationship

$$\|x^k + d^k - \bar{x}\| = o(\|d^k\|) \quad \text{as } k \rightarrow \infty. \quad (3.26)$$

Then the assumed estimate (3.24) in (ii) leads us to the inequalities

$$\varphi(x^k + d^k) - \varphi(x^k) - \sigma \langle \nabla\varphi(x^k), d^k \rangle$$

$$\begin{aligned}
 &\leq \langle \nabla \varphi(x^k + d^k), d^k \rangle - \frac{1}{2\kappa} \|d^k\|^2 - \sigma \langle \nabla \varphi(x^k), d^k \rangle \\
 &\leq \|\nabla \varphi(x^k + d^k)\| \cdot \|d^k\| - \frac{1}{2\kappa} \|d^k\|^2 + \sigma \|\nabla \varphi(x^k)\| \cdot \|d^k\| \\
 &\leq \ell \|x^k + d^k - \bar{x}\| \cdot \|d^k\| - \frac{1}{2\kappa} \|d^k\|^2 + \sigma \ell \|x^k - \bar{x}\| \cdot \|d^k\| \\
 &\leq \|d^k\|^2 \left(\ell \frac{\|x^k + d^k - \bar{x}\|}{\|d^k\|} - \frac{1}{2\kappa} + \sigma \ell \frac{\|x^k - \bar{x}\|}{\|d^k\|} \right)
 \end{aligned}$$

for all large $k \in \mathbb{N}$. Finally, it follows from $\sigma < 1/(2\ell\kappa)$, (3.25), and (3.26) that

$$\varphi(x^k + d^k) - \varphi(x^k) - \sigma \langle \nabla \varphi(x^k), d^k \rangle \leq 0 \quad \text{when } k \text{ is sufficiently large.}$$

This verifies (3.21) and thus completes the proof of the lemma. $\square$

Remark 2 (on acceptance of unit stepsize) Proposition 2 provides a sufficient condition to ensure the *asymptotic acceptance* of the unit stepsize. The key assumption here is the semismoothness of $\nabla \varphi$ at the reference point $\bar{x}$, which always includes the directional differentiability of $\nabla \varphi$ at $\bar{x}$. Lemma 1 introduces an alternative condition without the latter property to attain the acceptance of the unit stepsize, which depends on the modulus κ in (3.24) and the modulus of the Lipschitz continuity of $\nabla \varphi$. The given proof of this result requires the technical condition $\sigma \in (0, 1/(2\ell\kappa))$. It is not clear to us whether this condition can be either removed, or replaced by the more simple one $\sigma \in (0, 1/2)$. In the case where φ is twice differentiable around the critical point $\bar{x}$ and $\nabla^2 \varphi$ is continuous at $\bar{x}$, effective sufficient conditions for the acceptance of the unit stepsize can be found in [45, Sect. 5.2].

Now we are ready to justify the Q -superlinear rate of convergence of iterates in Algorithm 1 under some additional assumptions and relationships between parameters of the problem and the algorithm.

Theorem 2 (superlinear convergence of the coderivative-based damped Newton algorithm in $C^{1,1}$ optimization) *In the setting of Theorem 1 ensuring the convergence of $\{x^k\}$ to a tilt-stable minimizer $\bar{x}$ of φ as $k \rightarrow \infty$, suppose that $\nabla \varphi$ is locally Lipschitzian around $\bar{x}$ with some constant $\ell > 0$ being also semismooth* at this point. Then the rate of the convergence of $\{x^k\}$ is at least Q -superlinear if either one of the following two conditions is satisfied:*

- (i) $\nabla \varphi$ is directionally differentiable at $\bar{x}$.
- (ii) $\sigma \in (0, 1/(2\ell\kappa))$, where $\kappa > 0$ is a modulus of tilt stability of $\bar{x}$.

Moreover, in both cases (i) and (ii) the sequence $\{\varphi(x^k)\}$ converges Q -superlinearly to $\varphi(\bar{x})$, and the sequence $\{\nabla \varphi(x^k)\}$ converges Q -superlinearly to 0 as $k \rightarrow \infty$.

Proof Fixing a tilt-stable minimizer $\bar{x}$ with modulus $\kappa > 0$ from the assertions of Theorem 1, we split the proof of this theorem into the three claims.

Claim 1: *The sequence of directions $\{d^k\}$ satisfies condition (3.20).* Indeed, by the characterization of tilt-stable minimizers via the combined second-order subdifferential taken from [61, Theorem 3.5] and [11, Proposition 4.6], we find a positive number

δ such that the inequality

$$\langle z, w \rangle \geq \frac{1}{\kappa} \|w\|^2 \quad \text{for all } z \in \partial^2 \varphi(x)(w), \quad x \in \mathbb{B}_\delta(\bar{x}), \quad \text{and } w \in \mathbb{R}^n \quad (3.27)$$

is satisfied. Employing the subadditivity property of coderivatives taken from [46, Lemma 5.6] gives us

$$\partial^2 \varphi(x^k)(d^k) \subset \partial^2 \varphi(x^k)(x^k + d^k - \bar{x}) + \partial^2 \varphi(x^k)(-x^k + \bar{x}).$$

Since $-\nabla \varphi(x^k) \in \partial^2 \varphi(x^k)(d^k)$, for all $k \in \mathbb{N}$ there exists $v^k \in \partial^2 \varphi(x^k)(-x^k + \bar{x})$ such that

$$-\nabla \varphi(x^k) - v^k \in \partial^2 \varphi(x^k)(x^k + d^k - \bar{x}).$$

Using further (3.27) and the Cauchy–Schwarz inequality, we get

$$\|x^k + d^k - \bar{x}\| \leq \kappa \|\nabla \varphi(x^k) + v^k\| \quad \text{for sufficiently large } k \in \mathbb{N}. \quad (3.28)$$

The semismoothness* of $\nabla \varphi$ at $\bar{x}$ together with $\nabla \varphi(\bar{x}) = 0$ implies by [46, Lemma 5.5] that

$$\|\nabla \varphi(x^k) + v^k\| = \|\nabla \varphi(x^k) - \nabla \varphi(\bar{x}) + v^k\| = o(\|x^k - \bar{x}\|). \quad (3.29)$$

Then it follows from (3.28) and (3.29) that $\|x^k + d^k - \bar{x}\| = o(\|x^k - \bar{x}\|)$, which justifies the claim.

Claim 2: We have $\tau_k = 1$ for all $k \in \mathbb{N}$ sufficiently large provided that either condition (i), or condition (ii) of this theorem is satisfied. To verify the claim, let us show that (3.21) holds for large k under the imposed assumptions. Suppose first that (i) is satisfied. The directional differentiability and semismoothness* of $\nabla \varphi$ at $\bar{x}$ ensure by [32, Corollary 3.8] that $\nabla \varphi$ is semismooth at $\bar{x}$. Due to the inclusion $-\nabla \varphi(x^k) \in \partial^2 \varphi(x^k)(d^k)$ and estimate (3.27), we have that $\langle \nabla \varphi(x^k), d^k \rangle \leq -\kappa^{-1} \|d^k\|^2$ if k is large enough. Then the fulfillment of (3.21) in case (i) follows directly from Proposition 2. In case (ii), we know from Claim 1 that $\{d^k\}$ converges to 0 and that $x^k + d^k \rightarrow \bar{x}$ as $k \rightarrow \infty$. Employing the uniform second-order growth condition for tilt-stable minimizers from [61, Theorem 3.2] gives us a neighborhood U of $\bar{x}$ such that

$$\varphi(x) \geq \varphi(u) + \langle \nabla \varphi(u), x - u \rangle + \frac{1}{2\kappa} \|x - u\|^2 \quad \text{for all } x, u \in U,$$

and thus verifies (3.24). Using Lemma 1 brings us to (3.21), which verifies the claim.

Claim 3: The conclusions on the Q -superlinear convergence in the theorem hold in both cases (i) and (ii). We see from Claim 2 that $\tau_k = 1$ for all k sufficiently large, and thus Algorithm 1 eventually becomes the generalized pure Newton algorithm from [46, Algorithm 5.3]. Hence the claimed Q -superlinear convergence results follow from [46, Theorems 5.7 and 5.12]. $\square$

Remark 3 (comparing Algorithm 1 with other Newton-type methods) There exist several generalized Newton-type methods providing superlinearly convergent iterates under appropriate assumptions; see [26, 44, 47] and the bibliographies therein. We briefly compare Algorithm 1 with the two most popular globalized Newtonian methods. The first one is known as the *semismooth Newton method* initiated independently by Kummer [48] and by Qi and Sun [83]. The second method was introduced by Pang [74] under the name of the *B-differential Newton method*. Both methods address solving Lipschitzian equations $f(x) = 0$, which reduce in the setting of (1.1) to finding solutions of the gradient equation

$$\nabla\varphi(x) = 0, \quad x \in \mathbb{R}^n. \quad (3.30)$$

Regarding the aforementioned methods to solve the gradient equation (3.30), observe the following:

(i) The semismooth Newton method and its globalizations for solving (3.30) are based on Clarke's generalized Jacobian $\partial_C \nabla\varphi$ of $\nabla\varphi$ (see below) to find Newton directions d^k as solutions to the system of linear equations

$$-\nabla\varphi(x^k) = A^k d, \quad (3.31)$$

where A^k is an element of $\partial_C \nabla\varphi(x^k)$ as the convex hull of the *Bouligand's Jacobian*

$$\partial_B \nabla\varphi(x) := \left\{ \lim_{m \rightarrow \infty} \nabla^2\varphi(u_m) \mid u_m \rightarrow x, u_m \in Q_\varphi \right\}, \quad x \in \mathbb{R}^n, \quad (3.32)$$

with Q_φ standing for the set on which φ is twice differentiable. A strong feature of (3.31) is the *linearity* of equations therein, although the solvability of these equations requires the *nonsingularity* of all the matrices A^k . Moreover, computation cost to solve the system of linear equations (3.31) is known to be expensive. Although system (3.1) for finding directions in Algorithm 1 is not linear, we get from it a smaller set of algorithm directions in comparison with (3.31). The detailed numerical experiment to compare our approach and semismooth Newton method in a specific $\mathcal{C}^{1,1}$ optimization problem can be found in Sect. 6. Another advantage of Algorithm 1 over the semismooth Newton method is well-developed second-order subdifferential calculus that is not available for (3.32) and its convexification. Theoretical comparisons of local convergence between *coderivative-based Newtonian methods* and semismooth Newton methods can be found in [46, 68], where the regularity assumptions and the nonsingularity of all matrices in the generalized Jacobian have been discussed in detail. Regarding the global convergence of our approach and the semismooth Newton method, observe from [88, Theorem 13.52] that

$$\text{co } \partial^2\varphi(x)(w) = \text{co } \{Aw \mid A \in \partial_B \nabla\varphi(x)\}, \quad x \in \mathbb{R}^n, w \in \mathbb{R}^n, \quad (3.33)$$

which tells us that the positive-definiteness of $\partial^2\varphi(x)$ is equivalent to the positive-definiteness of $\partial_C \nabla\varphi(x)$ and $\partial_B \nabla\varphi(x)$. This positive-definiteness is required for both global convergence of Algorithm 1 and globalization of semismooth Newton method.

It happens because we not only need the existence of generalized Newton directions in inclusion (3.1) and in the system of linear equations (3.31), but the descent property of these directions is also needed to achieve the desired global convergence.

(ii) The B-differential Newton method for solving equation (3.30) developed in [74] and [82] is based on Robinson's B-derivative, which reduces for Lipschitzian mappings to the classical directional derivative. For this method, we need to solve the subproblem given below to find Newton directions d^k as solutions to

$$-f(x^k) = f'(x^k, d^k), \quad \text{where } f := \nabla\varphi, \quad (3.34)$$

with requiring the directional differentiability of $\nabla\varphi$. As shown in [74], the main assumptions to guarantee the solvability of (3.34) are the strict Fréchet differentiability of f and the nonsingularity of Jacobian ∇f at the solution point, which are rather restrictive requirements in comparison with our assumptions.

4 Coderivative-based regularized Newton method

Observe that the positive-definiteness of the generalized Hessian $\partial^2\varphi(x)$ in Algorithm 1 cannot be replaced by the less demanding positive-semidefiniteness of $\partial^2\varphi(x)$ to ensure the existence of descent Newton direction in Algorithm 1 as in Proposition 1. Indeed, consider the simplest linear function $\varphi(x) := x$ on $\mathbb{R}$. Then we obviously have that $\varphi''(x) \geq 0$ for all $x \in \mathbb{R}$, while there are no Newton directions $d \in \mathbb{R}$ such that the backtracking line search condition (3.4) holds. This means that Algorithm 1 cannot be even constructed without the positive-definiteness of $\partial^2\varphi(x)$. Here we propose the following globally convergent coderivative-based *generalized regularized Newton algorithm* to solve problems of $\mathcal{C}^{1,1}$ optimization that is *well-posed* and exhibits the *convergence* of its subsequences to *stationary points* of φ under merely the *positive-semidefiniteness* of the generalized Hessian. Linear and superlinear *convergence rates* are achieved under some additional assumptions.

Algorithm 2 Coderivative-based regularized Newton algorithm for $\mathcal{C}^{1,1}$ functions

Input: $x^0 \in \mathbb{R}^n$, $c > 0$, $\sigma \in (0, \frac{1}{2})$, $\beta \in (0, 1)$

```

1: for  $k = 0, 1, \dots$  do
2:   If  $\nabla\varphi(x^k) = 0$ , stop; otherwise let  $\mu_k := c\|\nabla\varphi(x^k)\|$  and go to next step
3:   Choose  $d^k \in \mathbb{R}^n$  such that  $-\nabla\varphi(x^k) \in \partial^2\varphi(x^k)(d^k) + \mu_k d^k$ 
4:   Set  $\tau_k = 1$ 
5:   while  $\varphi(x^k + \tau_k d^k) > \varphi(x^k) + \sigma\tau_k(\nabla\varphi(x^k), d^k)$  do
6:     set  $\tau_k := \beta\tau_k$ 
7:   end while
8:   Set  $x^{k+1} := x^k + \tau_k d^k$ 
9: end for
```

Our first major result in this section establishes the well-posedness and global convergence of iterates generated by Algorithm 2 to stationary points of φ under only the positive-semidefiniteness assumption on the generalized Hessian $\partial^2\varphi(x)$.

Theorem 3 (well-posedness and convergence of the coderivative-based regularized Newton algorithm) *For a function $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}$ of class $\mathcal{C}^{1,1}$ on $\mathbb{R}^n$, the following assertions hold:*

(i) *Let $x \in \mathbb{R}^n$ be such that $\nabla\varphi(x) \neq 0$ and $\partial^2\varphi(x)$ is positive-semidefinite, i.e.,*

$$\langle z, w \rangle \geq 0 \text{ for all } z \in \partial^2\varphi(x)(w) \text{ and } w \in \mathbb{R}^n. \quad (4.1)$$

Then for any $\varepsilon > 0$, there exists a nonzero direction $d \in \mathbb{R}^n$ with

$$-\nabla\varphi(x) \in \partial^2\varphi(x)(d) + \varepsilon d. \quad (4.2)$$

Moreover, every such direction satisfies the inequality $\langle \nabla\varphi(x), d \rangle < 0$. Consequently, for each $\sigma \in (0, 1)$ and $d \in \mathbb{R}^n$ satisfying (4.2) we have $\delta > 0$ such that

$$\varphi(x + \tau d) \leq \varphi(x) + \sigma \tau \langle \nabla\varphi(x), d \rangle \text{ whenever } \tau \in (0, \delta). \quad (4.3)$$

(ii) *Picking any starting point $x^0 \in \mathbb{R}^n$ such that $\partial^2\varphi(x)$ is positive-semidefinite for all $x \in \Omega$ from (3.5), we have that Algorithm 2 either stops after finitely many iterations, or produces a sequence of iterates $\{x^k\} \subset \Omega$ such that the sequence of values $\{\varphi(x^k)\}$ is monotonically decreasing. Moreover, all the accumulation points of $\{x^k\}$ satisfy the stationarity condition.*

Proof To justify (i), fix x satisfying the assumptions therein and consider the function

$$\varphi_\varepsilon(\cdot) = \varphi(\cdot) + \frac{\varepsilon}{2} \|\cdot\|^2 \text{ for any } \varepsilon > 0 \text{ on } \mathbb{R}^n.$$

It follows from the second-order subdifferential sum rule in [59, Proposition 1.121] that

$$\partial^2\varphi_\varepsilon(x)(w) = \partial^2\varphi(x)(w) + \varepsilon w \text{ for all } w \in \mathbb{R}^n. \quad (4.4)$$

Thus $z - \varepsilon w \in \partial^2\varphi(x)(w)$ whenever $z \in \partial^2\varphi_\varepsilon(x)(w)$. Due to the positive-semidefiniteness of $\partial^2\varphi(x)(w)$, we get $\langle z, w \rangle \geq \varepsilon \|w\|^2$, which implies that $\partial^2\varphi_\varepsilon(x)$ is positive-definite. It follows from [60, Theorem 5.16] that $\nabla\varphi_\varepsilon$ is locally strongly maximally monotone around $(x, \nabla\varphi_\varepsilon(x))$. Hence the gradient mapping $\nabla\varphi_\varepsilon$ is strongly metrically regular around this point due to [60, Theorem 5.13] telling us that the inverse mapping $\nabla\varphi_\varepsilon^{-1}: \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ admits a single-valued localization $\vartheta: V \rightarrow U$ around $(\nabla\varphi_\varepsilon(x), x)$, which is locally Lipschitzian around the point $\nabla\varphi_\varepsilon(x)$. Combining this with the scalarization formula (2.7) yields the representations

$$D^*\nabla\varphi_\varepsilon^{-1}(\nabla\varphi_\varepsilon(x), x)(\nabla\varphi(x)) = D^*\vartheta(\nabla\varphi_\varepsilon(x))(\nabla\varphi(x)) = \partial\langle \nabla\varphi(x), \vartheta \rangle(\nabla\varphi_\varepsilon(x)). \quad (4.5)$$

Since ϑ is locally Lipschitzian, we deduce from [59, Theorem 1.22] that $\partial\langle \nabla\varphi_\varepsilon(x), \vartheta \rangle(\nabla\varphi_\varepsilon(x)) \neq \emptyset$. Thus it follows from (4.5) that

$$D^*\nabla\varphi_\varepsilon^{-1}(\nabla\varphi_\varepsilon(x), x)(\nabla\varphi(x)) \neq \emptyset. \quad (4.6)$$

Picking any $-d \in D^*\nabla\varphi_\varepsilon^{-1}(\nabla\varphi_\varepsilon(x), x)(\nabla\varphi(x))$ and easily representing the coderivative of the inverse mapping $\nabla\varphi_\varepsilon^{-1}$ via that of $\nabla\varphi_\varepsilon$, we deduce from (4.6) the inclusion

$$-\nabla\varphi(x) \in D^*\nabla\varphi_\varepsilon(x)(d) = \partial^2\varphi_\varepsilon(x)(d).$$

Due to (4.4), it follows from the above that $-\nabla\varphi(x) \in \partial^2\varphi(x)(d) + \varepsilon d$.

To verify (i), it remains to show that $d \neq 0$. Supposing the contrary and using the local Lipschitz continuity of $\nabla\varphi$, we obtain from [59, Theorem 1.44] that

$$-\nabla\varphi(x) \in \partial^2\varphi(x)(0) = (D^*\nabla\varphi)(x)(0) = \{0\},$$

which contradicts the imposed assumption $\nabla\varphi(x) \neq 0$. The positive-definiteness of $\partial^2\varphi_\varepsilon(x)$ yields $\langle \nabla\varphi(x), d \rangle < 0$ that ensures in turn the fulfillment of (4.3) due to [44, Lemmas 2.18 and 2.19].

Next we proceed with the proof of (ii). It follows from (i) while arguing by induction that Algorithm 2 either stops after finitely many iterations, or produces a sequence of iterates $\{x^k\} \subset \Omega$ such that $\varphi(x^{k+1}) < \varphi(x^k)$ for all $k \in \mathbb{N}$. Let us first show that the sequence $\{d^k\}$ is bounded. Indeed, it follows from the construction that

$$-\nabla\varphi(x^k) - \mu_k d^k \in \partial^2\varphi(x^k)(d^k) \quad \text{for all } k \in \mathbb{N}, \quad (4.7)$$

and thus we get that $\langle -\nabla\varphi(x^k) - \mu_k d^k, d^k \rangle \geq 0$, i.e.,

$$\langle \nabla\varphi(x^k), -d^k \rangle \geq \mu_k \|d^k\|^2, \quad k \in \mathbb{N}. \quad (4.8)$$

Employing the Cauchy–Schwarz inequality and replacing μ_k with $c\|\nabla\varphi(x^k)\|$ lead us to

$$c\|\nabla\varphi(x^k)\| \cdot \|d^k\| = \mu_k \|d^k\| \leq \|\nabla\varphi(x^k)\|$$

and readily implies that $\|d^k\| \leq 1/c$ for all k .

Fix now an accumulation point $\bar{x}$ of the sequence of iterates $\{x^k\}$ and find a subsequence $\{x^{k_j}\}$ of $\{x^k\}$ such that $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$. Since the sequence $\{\varphi(x^k)\}$ is nonincreasing and $\varphi(\bar{x})$ is an accumulation point of $\{\varphi(x^k)\}$, this sequence converges to $\varphi(\bar{x})$ as $k \rightarrow \infty$. Moreover, we have

$$\varphi(x^{k+1}) - \varphi(x^k) \leq \sigma \tau_k \langle \nabla\varphi(x^k), d^k \rangle < 0 \quad \text{for all } k \in \mathbb{N},$$

which yields the equality

$$\lim_{k \rightarrow \infty} \tau_k \langle \nabla\varphi(x^k), d^k \rangle = 0. \quad (4.9)$$

Using the boundedness of $\{d^k\}$, we find a subsequence $\{d^{k_j}\}$ converging to some $\bar{d} \in \mathbb{R}^n$. Let us verify that

$$\langle \nabla\varphi(\bar{x}), \bar{d} \rangle = 0. \quad (4.10)$$

Indeed, if $\limsup_{j \rightarrow \infty} \tau_{k_j} > 0$, then (4.10) follows immediately from (4.9). Otherwise, we have $\lim_{j \rightarrow \infty} \tau_{k_j} = 0$, and the exit condition of the backtracking line search in Step 5 of Algorithm 2 brings us to

$$\varphi(x^{k_j} + \tau'_{k_j} d^{k_j}) > \varphi(x^{k_j}) + \sigma \tau'_{k_j} \langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle \quad (4.11)$$

for all $j \in \mathbb{N}$, where $\tau'_{k_j} := \tau_{k_j}/\beta$. Dividing now both sides of (4.11) by τ'_{k_j} and letting $j \rightarrow \infty$ implies that

$$\langle \nabla \varphi(\bar{x}), \bar{d} \rangle = \lim_{j \rightarrow \infty} \frac{\varphi(x^{k_j} + \tau'_{k_j} d^{k_j}) - \varphi(x^{k_j})}{\tau'_{k_j}} \geq \sigma \langle \nabla \varphi(\bar{x}), \bar{d} \rangle.$$

This tells us that $\langle \nabla \varphi(\bar{x}), \bar{d} \rangle \geq 0$ by $\sigma < 1$. Letting $j \rightarrow \infty$ in $\langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle \leq 0$, we get $\langle \nabla \varphi(\bar{x}), \bar{d} \rangle \leq 0$ and arrive at (4.10). Combining (4.8) and (4.10) verifies that $\mu_{k_j} \|d^{k_j}\|^2 \rightarrow 0$ as $j \rightarrow \infty$. By the definition of $\{\mu_k\}$ and the convergence $x^{k_j} \rightarrow \bar{x}$, we have $\mu_{k_j} = c \|\nabla \varphi(x^{k_j})\| \rightarrow c \|\nabla \varphi(\bar{x})\|$ as $j \rightarrow \infty$, which ensures that

$$c \|\nabla \varphi(\bar{x})\| \cdot \|\bar{d}\|^2 = \lim_{j \rightarrow \infty} \mu_{k_j} \|d^{k_j}\|^2 = 0. \quad (4.12)$$

Since φ is of class $\mathcal{C}^{1,1}$ around $\bar{x}$, it follows from [59, Theorem 1.44] and (4.7) that there exists $\ell > 0$ such that $\|\nabla \varphi(x^{k_j}) + \mu_{k_j} d^{k_j}\| \leq \ell \|d^{k_j}\|$ for all j sufficiently large. This yields

$$\|\nabla \varphi(x^{k_j})\|^2 + 2\mu_{k_j} \langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle + \mu_{k_j}^2 \|d^{k_j}\|^2 \leq \ell^2 \|d^{k_j}\|^2$$

for such j . Letting $j \rightarrow \infty$ in the above inequality, we arrive at $\|\nabla \varphi(\bar{x})\|^2 \leq \ell^2 \|\bar{d}\|^2$ due to (4.10) and the second equality in (4.12). Using the obtained estimate together with the first part of (4.12) gives us $\nabla \varphi(\bar{x}) = 0$ and thus completes the proof of the theorem. $\square$

Remark 4 (on proof of Theorem 3(ii)) Following the suggestion of the referee, we provide an alternative proof of the assertions in (ii) of Theorem 3 by assuming the contrary and using the *gradient related property* taken from [6]. Indeed, suppose that $\nabla \varphi(\bar{x}) \neq 0$ and then show that $\{d^k\}$ is gradient related to $\{x^k\}$, i.e., for any subsequence $\{x^{k_j}\}$ converging to $\bar{x}$, the corresponding subsequence $\{d^{k_j}\}$ is bounded and satisfies

$$\limsup_{j \rightarrow \infty} \langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle < 0. \quad (4.13)$$

Indeed, the boundedness of $\{d^k\}$ is verified by the same proof as in Theorem 3(ii). Due to (4.8), we have

$$\langle \nabla \varphi(x^{k_j}), d^{k_j} \rangle \leq -\mu_{k_j} \|d^{k_j}\|^2 = -c \|\nabla \varphi(x^{k_j})\| \cdot \|d^{k_j}\|^2 \text{ for all } j \in \mathbb{N}. \quad (4.14)$$

Since φ is of class $\mathcal{C}^{1,1}$ around $\bar{x}$, it follows from [59, Theorem 1.44] and (4.7) that there exists $\ell > 0$ such that $\|\nabla\varphi(x^{k_j}) + \mu_{k_j}d^{k_j}\| \leq \ell\|d^{k_j}\|$ for all j sufficiently large. By the definition of $\{\mu_k\}$ and the convergence $x^{k_j} \rightarrow \bar{x}$, we have that $\mu_{k_j} = c\|\nabla\varphi(x^{k_j})\| \rightarrow c\|\nabla\varphi(\bar{x})\|$ as $j \rightarrow \infty$, which ensures that μ_{k_j} is bounded from above by some $m > 0$, and thus arrive at the estimate

$$\|\nabla\varphi(x^{k_j})\| \leq \|\nabla\varphi(x^{k_j}) + \mu_{k_j}d^{k_j}\| + \mu_{k_j}\|d^{k_j}\| \leq (\ell + m)\|d^{k_j}\| \quad \text{for all } j \in \mathbb{N}. \quad (4.15)$$

Combining (4.14) and (4.15) tells us that

$$\langle \nabla\varphi(x^{k_j}), d^{k_j} \rangle \leq -c(\ell + m)^{-2}\|\nabla\varphi(x^{k_j})\|^3 \quad \text{for all } j \in \mathbb{N}.$$

By taking the upper limit in both sides above and using $\nabla\varphi(\bar{x}) \neq 0$ justifies (4.13). Therefore, it follows from [6, Proposition 1.2.1] that $\nabla\varphi(\bar{x}) = 0$, a contradiction that verifies (ii).

The next theorem establishes the linear and superlinear *convergence rates* of iterates in Algorithm 2 to *tilt-stable minimizers* under the *metric regularity assumption* on $\nabla\varphi$ at the solution point $\bar{x}$. Note that the latter property is constructively characterized by (2.4) and (2.10) as

$$\{u \in \mathbb{R}^n \mid 0 \in \partial\langle u, \nabla\varphi \rangle(\bar{x})\} = \{0\}.$$

Theorem 4 (linear and superlinear global convergence of coderivative-based regularized Newton algorithm) *In the setting of Theorem 3, let $\bar{x}$ be an accumulation point of $\{x^k\}$ such that $\nabla\varphi$ is metrically regular around this point. Then $\bar{x}$ is a tilt-stable local minimizer of φ and Algorithm 2 converges to $\bar{x}$ with the convergence rates as follows:*

- (i) *The sequence of values $\{\varphi(x^k)\}$ converges to $\varphi(\bar{x})$ at least Q -linearly.*
- (ii) *The sequences $\{x^k\}$ and $\{\nabla\varphi(x^k)\}$ converge at least R -linearly to $\bar{x}$ and 0, respectively.*
- (iii) *The convergence rates of $\{x^k\}$, $\{\varphi(x^k)\}$, and $\{\nabla\varphi(x^k)\}$ are at least Q -superlinear if $\nabla\varphi$ is semismooth* at $\bar{x}$ and either one of the following two conditions holds:*
 - (a) $\nabla\varphi$ is directionally differentiable at $\bar{x}$.
 - (b) $\sigma \in (0, 1/(2\ell\kappa))$, where $\kappa > 0$ and $\ell > 0$ are moduli of metric regularity and Lipschitz continuity of $\nabla\varphi$ around $\bar{x}$, respectively.

Proof We split the proof into the seven major claims of their own interest.

Claim 1 $\bar{x}$ is a tilt-stable local minimizer of φ . Due to Theorem 3, $\bar{x}$ is a stationary point of φ and $\bar{x} \in \Omega$, which implies that $\partial^2\varphi(\bar{x})$ is positive-semidefinite. This property and the imposed metric regularity of $\nabla\varphi$ around $\bar{x}$ with modulus κ allow us to conclude by using [22, Theorem 4.13] that $\bar{x}$ is a tilt-stable local minimizer of φ with the same modulus κ .

Claim 2 For any subsequence $\{x^{k_j}\}$ of $\{x^k\}$ with $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, the corresponding sequence $\{\tau_{k_j}\}$ in Algorithm 2 is bounded from below by a positive number

γ , and we have

$$\varphi(x^{k_j}) - \varphi(x^{k_j+1}) \geq \frac{\sigma\gamma}{\kappa} \|d^{k_j}\|^2 \quad \text{for all large } j \in \mathbb{N}. \quad (4.16)$$

Indeed, supposing on the contrary that $\{\tau_{k_j}\}$ is not bounded from below by a positive number and combining this with $\tau_k \geq 0$ give us a subsequence of $\{\tau_{k_j}\}$ that converges to 0. Let $\tau_{k_j} \rightarrow 0$ as $j \rightarrow \infty$ without loss of generality. Then using the characterization of tilt-stable minimizers via the combined second-order subdifferential taken from [61, Theorem 3.5] and [11, Proposition 4.6], we find $\delta > 0$ with

$$\langle z, w \rangle \geq \frac{1}{\kappa} \|w\|^2 \quad \text{for all } z \in \partial^2\varphi(x)(w), \quad x \in \mathbb{B}_\delta(\bar{x}), \quad \text{and } w \in \mathbb{R}^n. \quad (4.17)$$

Since $-\nabla\varphi(x^{k_j}) - \mu_{k_j}d^{k_j} \in \partial^2\varphi(x^{k_j})(d^{k_j})$ for all $j \in \mathbb{N}$, it follows from (4.17) that

$$\langle -\nabla\varphi(x^{k_j}), d^{k_j} \rangle \geq \left(\mu_{k_j} + \frac{1}{\kappa} \right) \|d^{k_j}\|^2 \geq \frac{1}{\kappa} \|d^{k_j}\|^2 \quad \text{for all } j \text{ sufficiently large}. \quad (4.18)$$

Then Claim 1 of Theorem 3 tells us that the sequence $\{d^k\}$ is bounded. Hence $x^{k_j} + \tau_{k_j}d^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$ and $x^{k_j} + \tau_{k_j}d^{k_j} \in \text{int } \mathbb{B}_\delta(\bar{x})$ whenever j is sufficiently large. Applying Lemma 3 from the Appendix, we get

$$\beta^{-1}\tau_{k_j} > \frac{2(\sigma - 1)\langle \nabla\varphi(x^{k_j}), d^{k_j} \rangle}{\ell\|d^{k_j}\|^2} \geq \frac{2(1 - \sigma)}{\kappa\ell},$$

where the second inequality follows from (4.18).

Letting $j \rightarrow \infty$ gives us $\sigma \geq 1$, a contradiction due to $\sigma < 1$. This verifies the existence of $\gamma > 0$ such that $\tau_{k_j} \geq \gamma$ for all $j \in \mathbb{N}$. Using estimate (4.18), we find $j_0 \in \mathbb{N}$ with

$$\varphi(x^{k_j}) - \varphi(x^{k_j+1}) \geq \sigma\tau_{k_j}\langle -\nabla\varphi(x^{k_j}), d^{k_j} \rangle \geq \frac{\sigma\gamma}{\kappa} \|d^{k_j}\|^2 \quad \text{for all } j \geq j_0, \quad (4.19)$$

which therefore justifies Claim 2.

Claim 3 The iterative sequence $\{x^k\}$ converges to $\bar{x}$. To verify this, we are based on Ostrowski's condition from [26, Proposition 8.3.10]. Let us first check that there is no other accumulation point of $\{x^k\}$ in $\mathbb{B}_\delta(\bar{x})$. On the contrary, suppose that there exists $\tilde{x} \in \mathbb{B}_\delta(\bar{x})$ such that $\tilde{x} \neq \bar{x}$ and $\tilde{x}$ is an accumulation point of $\{x^k\}$. It follows from Theorem 3 that $\tilde{x}$ is a stationary point of φ , which contradicts the strong convexity of φ on $\mathbb{B}_\delta(\bar{x})$. Supposing next that $\{x^{k_j}\}$ is an arbitrary subsequence of $\{x^k\}$ with $x^{k_j} \rightarrow \bar{x}$ as $j \rightarrow \infty$, we check that

$$\lim_{j \rightarrow \infty} \|x^{k_j+1} - x^{k_j}\| = 0. \quad (4.20)$$

Indeed, find by Claim 2 such $\gamma > 0$ that (4.16) holds, which implies that

$$\|x^{k_j+1} - x^{k_j}\|^2 = \tau_{k_j}^2 \|d^{k_j}\|^2 \leq \|d^{k_j}\|^2 \leq \frac{\kappa}{\sigma\gamma} (\varphi(x^{k_j}) - \varphi(x^{k_j+1})) \rightarrow 0$$

as $j \rightarrow \infty$ and thus verifies (4.20). Employing now [26, Proposition 8.3.10] ensures the convergence of $\{x^k\}$ to $\bar{x}$ as $k \rightarrow \infty$ and therefore completes the proof of this claim.

Claim 4 The convergence rate of $\{\varphi(x^k)\}$ is at least Q-linear, while the convergence rates of $\{x^k\}$ and $\{\|\nabla\varphi(x^k)\|\}$ are at least R-linear. Indeed, the strong convexity of φ on $\mathbb{B}_\delta(\bar{x})$ implies that

$$\begin{aligned} \varphi(x) &\geq \varphi(u) + \langle \nabla\varphi(u), x - u \rangle + \frac{1}{2\kappa} \|x - u\|^2 \quad \text{and} \\ \langle \nabla\varphi(x) - \nabla\varphi(u), x - u \rangle &\geq \frac{1}{\kappa} \|x - u\|^2 \end{aligned} \quad (4.21)$$

for all $x, u \in \mathbb{B}_\delta(\bar{x})$. By the convergence $x^k \rightarrow \bar{x}$ we have that $x^k \in U$ for all k sufficiently large, which we assumed from now on. Substituting $x := x^k$ and $u := \bar{x}$ into (4.21) and then using the Cauchy–Schwarz inequality together with $\nabla\varphi(\bar{x}) = 0$ yield the estimates

$$\varphi(x^k) \geq \varphi(\bar{x}) + \frac{1}{2\kappa} \|x^k - \bar{x}\|^2 \quad \text{and} \quad (4.22)$$

$$\|\nabla\varphi(x^k)\| \geq \frac{1}{\kappa} \|x^k - \bar{x}\|. \quad (4.23)$$

The local Lipschitz continuity of $\nabla\varphi$ around $\bar{x}$ and the result of [44, Lemma A.11] ensure the existence of a positive number ℓ such that

$$\varphi(x^k) - \varphi(\bar{x}) = |\varphi(x^k) - \varphi(\bar{x}) - \langle \nabla\varphi(\bar{x}), x^k - \bar{x} \rangle| \leq \frac{\ell}{2} \|x^k - \bar{x}\|^2. \quad (4.24)$$

Moreover, since $-\nabla\varphi(x^k) - \mu_k d^k \in \partial^2\varphi(x^k)(d^k)$, by using [59, Theorem 1.44] we have

$$\|\nabla\varphi(x^k) + \mu_k d^k\| \leq \ell \|d^k\|. \quad (4.25)$$

It follows from the convergence $x^k \rightarrow \bar{x}$ and $\nabla\varphi(\bar{x}) = 0$ that $\mu_k = c\|\nabla\varphi(x^k)\| \rightarrow 0$ as $k \rightarrow \infty$, which implies that $\mu_k \leq \ell$. Combining the latter with (4.25) gives us the estimates

$$\|\nabla\varphi(x^k)\| \leq \|\nabla\varphi(x^k) + \mu_k d^k\| + \mu_k \|d^k\| \leq 2\ell \|d^k\|. \quad (4.26)$$

By Claim 2 we have that $\{\tau_k\}$ is bounded from below by some constant $\gamma > 0$ and that

$$\varphi(x^k) - \varphi(x^{k+1}) \geq \frac{\sigma\gamma}{\kappa} \|d^k\|^2,$$

which together with (4.26) yields the inequality

$$\varphi(x^k) - \varphi(x^{k+1}) \geq \frac{\sigma\gamma}{4\kappa\ell^2} \|\nabla\varphi(x^k)\|^2. \quad (4.27)$$

Combining finally (4.23), (4.24), and (4.27) and then applying Lemma 4 with the sequences $\alpha_k := \varphi(x^k) - \varphi(\bar{x})$, $\beta_k := \|\nabla\varphi(x^k)\|$, $\gamma_k := \|x^k - \bar{x}\|$ and positive numbers $c_1 := (\sigma\gamma)/(4\kappa\ell^2)$, $c_2 := 1/\kappa$, and $c_3 := \ell/2$, we complete the verification of all the conclusions of this claim.

Claim 5 $\|x^k + d^k - \bar{x}\| = o(\|x^k - \bar{x}\|)$ provided that $\nabla\varphi$ is semismooth* at $\bar{x}$. Indeed, the subadditivity property of coderivatives taken from [46, Lemma 5.6] tells us that

$$\partial^2\varphi(x^k)(d^k) \subset \partial^2\varphi(x^k)(x^k + d^k - \bar{x}) + \partial^2\varphi(x^k)(-x^k + \bar{x}).$$

Since $-\nabla\varphi(x^k) - \mu_k d^k \in \partial^2\varphi(x^k)(d^k)$, there exists $v^k \in \partial^2\varphi(x^k)(-x^k + \bar{x})$ such that

$$-\nabla\varphi(x^k) - \mu_k d^k - v^k \in \partial^2\varphi(x^k)(x^k + d^k - \bar{x}).$$

For large $k \in \mathbb{N}$ with $x^k \in \mathbb{B}_\delta(\bar{x})$, applying (4.17) yields

$$\langle \nabla\varphi(x^k), -d^k \rangle \geq (\kappa^{-1} + \mu_k) \|d^k\|^2 \geq \kappa^{-1} \|d^k\|^2, \quad (4.28)$$

$$\langle -\nabla\varphi(x^k) - \mu_k d^k - v^k, x^k + d^k - \bar{x} \rangle \geq \frac{1}{\kappa} \|x^k + d^k - \bar{x}\|^2.$$

Using again the Cauchy–Schwarz inequality together with $\nabla\varphi(\bar{x}) = 0$ ensures that

$$\|d^k\| \leq \kappa \|\nabla\varphi(x^k)\| = \kappa \|\nabla\varphi(x^k) - \nabla\varphi(\bar{x})\| \leq \kappa\ell \|x^k - \bar{x}\|, \quad (4.29)$$

$$\|x^k + d^k - \bar{x}\| \leq \kappa \|\nabla\varphi(x^k) + v^k + \mu_k d^k\| \leq \kappa \left(\|\nabla\varphi(x^k) + v^k\| + \mu_k \|d^k\| \right). \quad (4.30)$$

Furthermore, it follows from the Lipschitz continuity of $\nabla\varphi$ on $\mathbb{B}_\delta(\bar{x})$ and from $\nabla\varphi(\bar{x}) = 0$ that

$$\mu_k = c \|\nabla\varphi(x^k)\| = c \|\nabla\varphi(x^k) - \nabla\varphi(\bar{x})\| \leq \ell \|x^k - \bar{x}\|. \quad (4.31)$$

Combining (4.29), (4.30), (4.31) tells us that

$$\|x^k + d^k - \bar{x}\| \leq \kappa \|\nabla\varphi(x^k) + v^k\| + \kappa^2 \ell^2 \|x^k - \bar{x}\|^2. \quad (4.32)$$

Using now the semismooth* property of the gradient mapping $\nabla\varphi$ at $\bar{x}$ together with the stationarity condition $\nabla\varphi(\bar{x}) = 0$ implies by [46, Lemma 5.5] that

$$\|\nabla\varphi(x^k) + v^k\| = \|\nabla\varphi(x^k) - \nabla\varphi(\bar{x}) + v^k\| = o(\|x^k - \bar{x}\|). \quad (4.33)$$

It follows from (4.32) and (4.33) that $\|x^k + d^k - \bar{x}\| = o(\|x^k - \bar{x}\|)$ as $k \rightarrow \infty$, which verifies this claim.

Claim 6 We have $\tau_k = 1$ for all large $k \in \mathbb{N}$ provided that $\nabla\varphi$ is semismooth* at $\bar{x}$ and that either condition (a), or condition (b) of the theorem holds. To proceed, it suffices to verify the estimate in (3.21) under both conditions (a) and (b). If (a) is satisfied, then $\nabla\varphi$ is semismooth at $\bar{x}$. Then this estimate and the assertion of the claim follows directly by (4.28) and Proposition 2. Assuming now the condition in (b) and using Claim 5, we easily see that $\{d^k\}$ converges to 0 and that $x^k + d^k \rightarrow \bar{x}$ as $k \rightarrow \infty$. Employing (4.21) justifies (3.24). Then (3.21) follows from Lemma 1 and thus completes the verification of the claim.

Claim 7 The Q-superlinear convergence holds in both cases (a) and (b) of (iii). Indeed, we get from Claim 6 that $\tau_k = 1$ for all k sufficiently large. It follows from Claim 5 that

$$\|x^{k+1} - \bar{x}\| = \|x^k + \tau_k d^k - \bar{x}\| = \|x^k + d^k - \bar{x}\| = o(\|x^k - \bar{x}\|) \text{ as } k \rightarrow \infty,$$

which justifies the Q-superlinear convergence of $\{x^k\}$ in both cases. The Q-superlinear convergence of $\{\varphi(x^k)\}$ follows immediately from (4.22) and (4.24), while the Q-superlinear convergence of $\{\nabla\varphi(x^k)\}$ is a consequence of (4.23) and (4.31). This completes the proof of the theorem. $\square$

Remark 5 (comparison with related globalized Newton-type algorithms) Observe the following: (i) In contrast to the generalized damped Newton algorithm (Algorithm 1), the solvability of subproblems and the behavior of accumulation points in Algorithm 2 are guaranteed merely under the *positive-semidefiniteness* of $\partial^2\varphi(x)$ in Theorem 3. To achieve the convergence of the iterative sequence $\{x^k\}$ in Algorithm 2, we only need the metric regularity of $\nabla\varphi$ around the accumulation point $\bar{x}$ instead of the positive-definiteness of $\partial^2\varphi(x)$ for all $x \in \Omega$, which is the key assumption of the convergence in Algorithm 1. Note also that this is just a sufficient condition to ensure the convergence of Algorithm 2. The crucial open question we will pursue in our future research is whether it is possible to replace the metric regularity of $\nabla\varphi$ in Theorem 4 by a weaker assumption. One of the natural assumptions of this type is that $\|\nabla\varphi(x)\|$ provides a local error bound near the accumulation point, which was investigated in, e.g., [15, 50, 94] in different settings.

(ii) The generalized regularized Newton methods via the Bouligand Jacobian can be found in Pang and Qi [75] and in the book [26, Sect. 8.3.3]. Their method aims to solve the optimization problem

$$\text{minimize } \varphi(x) \text{ subject to } x \in P, \quad (4.34)$$

where P is a nonempty polyhedron in $\mathbb{R}^n$, φ is a convex $\mathcal{C}^{1,1}$ function defined on an open convex set $\Omega \subset \mathbb{R}^n$ containing P . In the case of unconstrained minimization (1.1) with $P = \mathbb{R}^n$, the generalized regularized Newton method via the Bouligand Jacobian requires to find direction d^k as solutions to the equation

$$-\nabla\varphi(x^k) = (A^k + \varepsilon_k I)d^k, \quad \text{where } A^k \in \partial_B \nabla\varphi(x^k), \text{ and } \varepsilon_k > 0.$$

The key assumption to guarantee the convergence of their methods is that all the matrices in $\partial_B \nabla\varphi(\bar{x})$ are nonsingular, where $\bar{x}$ is the accumulation point of the iterative sequence $\{x^k\}$ generated by their method; see, e.g., [26, Theorem 8.3.19]. This assumption is weaker than our assumption that $\nabla\varphi$ is metrically regular around $(\bar{x}, 0)$ due to the coderivative criterion (2.4) and the inclusion

$$\partial_B \nabla\varphi(\bar{x})w \subset \partial^2\varphi(\bar{x})(w) \quad \text{for all } w \in \mathbb{R}^n.$$

Observe to this end that the metric regularity property is defined for arbitrary set-valued mappings, and it is used for subgradient ones in this paper to guarantee the convergence; see Sect. 5. Note also the calculus rules developed for the $\partial_B \nabla\varphi$ are more limited in comparison with full calculus available for $\partial^2\varphi$.

5 Coderivative-based Newton methods in composite optimization

In this section, we consider a broad and highly important class of optimization problems given by

$$\text{minimize } \varphi(x) := f(x) + g(x), \quad x \in \mathbb{R}^n, \quad (5.1)$$

where $f: \mathbb{R}^n \rightarrow \mathbb{R}$ is a convex and smooth function, while the regularizer $g: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ is a convex and extended-real-valued one. This class is known as problems of *convex composite optimization*.

Problems written in format (5.1) frequently arise in many applied areas including machine learning, compressed sensing, image processing, etc. Since the regularizer g is generally extended-real-valued, the unconstrained format (5.1) encompasses problems of *constrained optimization*. If, in particular, g is the indicator function of a closed and convex set, then (5.1) becomes a constrained optimization problems studied, e.g., in the book [72] with numerous applications.

One of the most well-recognized and applied algorithms to solve problems (5.1) is the forward–backward splitting (FBS), or proximal splitting, method [14, 52]. Since this method is of first order, its rate of convergence is at most linear. Another approach to solve (5.1) is to use second-order methods such as proximal Newton methods, proximal quasi-Newton methods, etc.; see, e.g., [5, 49, 69]. Although the latter approach has several benefits over first-order methods (as fast convergence and high accuracy), a severe limitation of these methods is the cost of solving subproblems.

To develop here new globally convergent Newton methods to solve convex composite optimization problems of type (5.1), we first recall the classical notions of convex

and variational analysis; see, e.g., [88]. Given an extended-real-valued, proper, l.s.c. function $\varphi: \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ and a number $\gamma > 0$, the *Moreau envelope* $e_\gamma \varphi$ and the *proximal mapping* $\text{Prox}_{\gamma\varphi}$ are defined by, respectively,

$$e_\gamma \varphi(x) := \inf_{y \in \mathbb{R}^n} \left\{ \varphi(y) + \frac{1}{2\gamma} \|y - x\|^2 \right\}, \quad (5.2)$$

$$\text{Prox}_{\gamma\varphi}(x) := \operatorname{argmin}_{y \in \mathbb{R}^n} \left\{ \varphi(y) + \frac{1}{2\gamma} \|y - x\|^2 \right\}. \quad (5.3)$$

If $\lambda = 1$, we use the notation $e_\varphi(x)$ and $\text{Prox}_\varphi(x)$ in (5.2) and (5.3), respectively. These notions have been well investigated in variational analysis and optimization as efficient tools of regularization and approximation of nonsmooth functions. Given a closed set $\emptyset \neq C \subset \mathbb{R}^n$, the *orthogonal projection mapping* $P_C: \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ is

$$P_C(x) := \operatorname{argmin}_{y \in C} \|y - x\| \quad \text{for all } x \in \mathbb{R}^n.$$

It is clear that if $g: \mathbb{R}^n \rightarrow (-\infty, \infty]$ be defined by $g(x) := \delta_C(x)$, we have

$$\text{Prox}_{\gamma g}(x) = P_C(x) \quad \text{for all } x \in \mathbb{R}^n, \gamma > 0.$$

More recently, the following extended notion, known now as the *forward-backward envelope*, has been introduced by Patrino and Bemporad [76] for problems of convex composite optimization.

Definition 4 (forward-backward envelope) Let $\varphi = f + g$ be as in (5.1), and let $\gamma > 0$. The FORWARD- BACKWARD ENVELOPE (FBE) of φ with parameter γ is

$$\varphi_\gamma(x) := \inf_{y \in \mathbb{R}^n} \left\{ f(x) + \langle \nabla f(x), y - x \rangle + g(y) + \frac{1}{2\gamma} \|y - x\|^2 \right\}, \quad (5.4)$$

which, by construction (5.2) of the Moreau envelope, is represented by

$$\varphi_\gamma(x) = f(x) - \frac{\gamma}{2} \|\nabla f(x)\|^2 + e_\gamma g(x - \gamma \nabla f(x)). \quad (5.5)$$

The FBE has already been used for developing some efficient algorithms to solve non-smooth optimization problems; see, e.g., [76, 90, 91] with further references therein. The following results taken from [76, 90] list those properties of the forward-backward envelope for convex composite extended-real-valued functions that are needed to derive the main results of this section.

Proposition 3 (basic properties of FBE) Let $\varphi = f + g$ be as in (5.1), and let $\gamma > 0$. Suppose that f is $\mathcal{C}^2$ -smooth on $\mathbb{R}^n$, and that ∇f is Lipschitz continuous on $\mathbb{R}^n$ with modulus $\ell > 0$. Then we have:

(i) The FBE φ_γ of φ is $\mathcal{C}^1$ -smooth on $\mathbb{R}^n$ with the gradient

$$\nabla \varphi_\gamma(x) = \gamma^{-1} (I - \gamma \nabla^2 f(x)) (x - \text{Prox}_{\gamma g}(x - \gamma \nabla f(x))), \quad x \in \mathbb{R}^n. \quad (5.6)$$

Moreover, the set of optimal solutions to (5.1) agrees with the stationary points of φ_γ by

$$\operatorname{argmin} \varphi := \operatorname{zer} \nabla \varphi_\gamma = \{x \in \mathbb{R}^n \mid \nabla \varphi_\gamma(x) = 0\} \text{ for all } \gamma \in (0, 1/\ell).$$

(ii) Let $f(x) := \frac{1}{2}\langle Ax, x \rangle + \langle b, x \rangle + \alpha$, where $A \in \mathbb{R}^{n \times n}$ is a positive-semidefinite symmetric matrix, $b \in \mathbb{R}^n$, and $\alpha \in \mathbb{R}$. Define the numbers

$$L := 2(1 - \gamma \lambda_{\min}(A)) / \gamma \text{ and} \\ K := \min\{(1 - \gamma \lambda_{\min}(A))\lambda_{\min}(A), (1 - \gamma \lambda_{\max}(A))\lambda_{\max}(A)\}.$$

Then for all $\gamma \in (0, 1/\ell)$, the FBE φ_γ is convex and its gradient $\nabla \varphi_\gamma$ is globally Lipschitzian on $\mathbb{R}^n$ with modulus L . If A is positive-definite, then φ_γ is strongly convex with modulus K .

It follows from Proposition 3 that using the forward-backward envelope (5.4) makes it possible to pass from the nonsmooth composite optimization problem (5.1) to the unconstrained one:

$$\text{minimize } \varphi_\gamma(x) \text{ subject to } x \in \mathbb{R}^n \quad (5.7)$$

with a smooth cost function. Thanks to the explicit calculations of φ_γ in (5.5) and its gradient (5.6), we can extend Algorithm 1 and Algorithm 2 to cover problem (5.1) via passing to (5.7). The implementation of this procedure requires revealing appropriate assumptions on φ in (5.1), which ensure the fulfillment of those for φ_γ and thus allow us to apply the results of Sects. 3, 4 to the class of nondifferentiable convex problems (5.7).

Note that (5.7) is not generally a problem of $\mathcal{C}^{1,1}$ optimization, since Proposition 3(i) does not ensure the Lipschitz continuity of $\nabla \varphi_\gamma$. The latter property is guaranteed by Proposition 3(ii) when f is a quadratic function and thus problem (5.1) is written as

$$\text{minimize } \varphi(x) := \frac{1}{2}\langle Ax, x \rangle + \langle b, x \rangle + \alpha + g(x), \quad x \in \mathbb{R}^n, \quad (5.8)$$

where $A \in \mathbb{R}^{n \times n}$ is a positive-semidefinite symmetric matrix, $b \in \mathbb{R}^n$, and $\alpha \in \mathbb{R}$. From now on, this is our standing framework for the rest of the section.

Let us highlight that problems of type (5.8) are important for their own sake, while they also arise frequently as *subproblems* for various efficient numerical algorithms including *sequential quadratic programming methods* (SQP) [7, 44], *augmented Lagrangian methods* [38, 51, 81], *proximal Newton methods* [49, 69], etc. Observe furthermore that optimization problems of this type often appear in practical models related, e.g., to machine learning and statistics. In particular, *Lasso problems* considered in Sect. 6 can be written in form (5.8). Moreover, there are some other important classes of problems that are modeled as (5.8). They include problems in *support vector machine* [43], *convex clustering* [78, 89], *constrained quadratic optimization* [72], etc.

Now we start the procedure of designing and justifying globally convergent generalized Newton algorithms to solve the convex composite problem (5.8) by applying the corresponding results for the $\mathcal{C}^{1,1}$ optimization problem (5.7) obtained in Sects. 3 and 4. The first step is to express the generalized Hessian of the FBE φ_γ from (5.7) in terms of the given data of (5.8).

Proposition 4 (calculating the generalized Hessian of FBE) *Let $\varphi = f + g$ be as in (5.8), and let $\gamma > 0$ be such that $B := I - \gamma A$ is positive-definite. Then we have the calculation formula*

$$\bar{z} \in \partial^2 \varphi_\gamma(\bar{x})(w) \iff B^{-1}\bar{z} - Aw \in \partial^2 g \left(\text{Prox}_{\gamma g}(\bar{u}), \frac{1}{\gamma}(\bar{u} - \text{Prox}_{\gamma g}(\bar{u})) \right) (w - \gamma B^{-1}\bar{z}) \quad (5.9)$$

for any $\bar{x} \in \mathbb{R}^n$, $w \in \mathbb{R}^n$, and $\bar{u} := \bar{x} - \gamma(A\bar{x} + b)$.

Proof Fix x , w , and $\bar{u}$ as above and define the function $h : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ by

$$h(x) := e_\gamma g(x - \gamma(Ax + b)) \quad \text{for all } x \in \mathbb{R}^n.$$

It is clear that h is continuously differentiable with

$$\nabla h(\bar{x}) = (I - \gamma A)^* \nabla e_\gamma g(\bar{u}) = B \nabla e_\gamma g(\bar{u}).$$

It follows from the definition of FBE (5.5) and the second-order sum rule from [59, Proposition 1.121] that

$$\partial^2 \varphi_\gamma(\bar{x})(w) = (A - \gamma A^* A)w + \partial^2 h(\bar{x})(w) = BAw + \partial^2 h(\bar{x})(w). \quad (5.10)$$

By using the second-order chain rule from [59, Theorem 1.127], we have

$$\partial^2 h(\bar{x})(w) = B \partial^2 e_\gamma g(\bar{u})(Bw). \quad (5.11)$$

Combining (5.10) and (5.11) gives us the relationship

$$\partial^2 \varphi_\gamma(\bar{x})(w) = BAw + B \partial^2 e_\gamma g(\bar{u})(Bw),$$

which in turn yields the equivalencies

$$\bar{z} \in \partial^2 \varphi_\gamma(\bar{x})(w) \iff \bar{z} - BAw \in B \partial^2 e_\gamma g(\bar{u})(Bw) \iff B^{-1}\bar{z} - Aw \in \partial^2 e_\gamma g(\bar{u})(Bw).$$

Employing finally [46, Lemma 6.2], we arrive at the inclusion

$$B^{-1}\bar{z} - Aw \in \partial^2 g \left(\text{Prox}_{\gamma g}(\bar{u}), \frac{1}{\gamma}(\bar{u} - \text{Prox}_{\gamma g}(\bar{u})) \right) (Bw - \gamma B^{-1}\bar{z} + \gamma Aw),$$

which completes the proof of the proposition due to $Bw + \gamma Aw = w$. $\square$

The next proposition shows that the metric regularity and tilt stability of the original objective φ in (5.8) is equivalent to the corresponding properties of its FBE φ_γ in (5.7). Moreover, we get a useful estimate of the inverse mapping of $\partial^2\varphi_\gamma$ in terms of the given data of (5.8).

Proposition 5 (metric regularity and tilt-stability of FBE) *Let $\varphi = f + g$ be as in (5.8), and let $\gamma > 0$ be such that $B := I - \gamma A$ is positive-definite. Then for any $\bar{x} \in \mathbb{R}^n$ satisfying $0 \in \partial\varphi(\bar{x})$ we have:*

- (i) $\|\partial^2\varphi_\gamma(\bar{x})^{-1}\| \leq \|\partial^2\varphi(\bar{x}, 0)^{-1}\| + \gamma\|B^{-1}\|$.
- (ii) $\partial\varphi$ is metrically regular around $(\bar{x}, 0)$ if and only if $\nabla\varphi_\gamma$ is metrically regular around $\bar{x}$.
- (iii) $\bar{x}$ is a tilt-stable local minimizer of φ if and only if $\bar{x}$ is a tilt-stable local minimizer of φ_γ .

Proof It follows from Proposition 4 that

$$z \in \partial^2\varphi_\gamma(\bar{x})(w) \iff 0 \in Aw + \partial^2g\left(\text{Prox}_{\gamma g}(\bar{u}), \frac{1}{\gamma}(\bar{u} - \text{Prox}_{\gamma g}(\bar{u}))\right)(w - \gamma B^{-1}z) \quad (5.12)$$

with $\bar{u}$ defined therein. The convexity of φ ensures that $\bar{x}$ is an optimal solution to (5.8), and thus $\bar{x} - \text{Prox}_{\gamma g}(\bar{u}) = 0$ by [1, Theorem 27.2]. Therefore, (5.12) is equivalent to

$$\begin{aligned} z \in Aw + \partial^2g(\bar{x}, -A\bar{x} - b)(w - \gamma B^{-1}z) \\ = A(w - \gamma B^{-1}z) + \partial^2g(\bar{x}, -A\bar{x} - b)(w - \gamma B^{-1}z) + \gamma AB^{-1}z. \end{aligned} \quad (5.13)$$

The second-order subdifferential sum rule from [59, Proposition 1.121] yields

$$\begin{aligned} A(w - \gamma B^{-1}z) + \partial^2g(\bar{x}, -A\bar{x} - b)(w - \gamma B^{-1}z) &= \partial^2(f + g)(\bar{x}, 0)(w - \gamma B^{-1}z) \\ &= \partial^2\varphi(\bar{x}, 0)(w - \gamma B^{-1}z). \end{aligned} \quad (5.14)$$

Combining (5.12), (5.13), and (5.14) gives us the equivalence

$$z \in \partial^2\varphi_\gamma(\bar{x})(w) \iff z \in \partial^2\varphi(\bar{x}, 0)(w - \gamma B^{-1}z), \quad (5.15)$$

which verifies (i). It follows from the coderivative criterion (2.4) and the equivalence (5.15) that $\partial\varphi$ is metrically regular around $(\bar{x}, 0)$ if and only if $\nabla\varphi_\gamma$ is metrically regular around $\bar{x}$, which justifies assertion (ii). Finally, [22, Proposition 4.5] tells us that a stationary point $\bar{x}$ is a tilt-stable local minimizer of an l.s.c. convex function if and only if its subgradient mapping is metrically regular around $(\bar{x}, 0)$. Using this observation together with (ii), we obtain (iii) and complete the proof of the proposition. $\square$

Now we recall some other notions of variational analysis that are used to establish the superlinear convergence of both algorithms developed in this section to solve problem (5.8). These notions, introduced by Rockafellar, are taken from the book [88]. A set-valued mapping $S : \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ is *proto-differentiable* at $(\bar{x}, \bar{y}) \in \text{gph } S$ if for any $\bar{w} \in \mathbb{R}^n$, $\bar{z} \in \text{Lim}, \sup_{\substack{t \downarrow 0 \\ w \rightarrow \bar{w}}} (S(\bar{x} + tw) - \bar{y})/t$, and $t_k \downarrow 0$ there exist $w_k \rightarrow \bar{w}$ and $z_k \rightarrow \bar{z}$ such that $z_k \in (S(\bar{x} + t_k w_k) - \bar{y})/t_k$ whenever $k \in \mathbb{N}$. Given $\varphi : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ with $\bar{x} \in \text{dom } \varphi$, consider the family of second-order finite differences

$$\Delta_\tau^2 \varphi(\bar{x}, v)(u) := \frac{\varphi(\bar{x} + \tau u) - \varphi(\bar{x}) - \tau \langle v, u \rangle}{\frac{1}{2} \tau^2}$$

and define the *second subderivative* of φ at $\bar{x}$ for $v \in \mathbb{R}^n$ and $w \in \mathbb{R}^n$ by

$$d^2 \varphi(\bar{x}, v)(w) := \liminf_{\tau \downarrow 0, u \rightarrow w} \Delta_\tau^2 \varphi(\bar{x}, v)(u).$$

Then φ is said to be *twice epi-differentiable* at $\bar{x}$ for v if for every $w \in \mathbb{R}^n$ and every choice of $\tau_k \downarrow 0$ there exists a sequence $w^k \rightarrow w$ such that

$$\frac{\varphi(\bar{x} + \tau_k w^k) - \varphi(\bar{x}) - \tau_k \langle v, w^k \rangle}{\frac{1}{2} \tau_k^2} \rightarrow d^2 \varphi(\bar{x}, v)(w) \text{ as } k \rightarrow \infty.$$

Twice epi-differentiability has been recognized as an important concept of second-order variational analysis with numerous applications to optimization; see the aforementioned monograph by Rockafellar and Wets and the recent papers [54–56] developing a systematic approach to verify epi-differentiability via *parabolic regularity*, which is a major second-order property of sets and functions.

The next proposition expresses the properties of the FBE φ_γ in (5.7), which are needed for the superlinear convergence of our algorithms, in terms of the given data of (5.8). Recall that the sign ‘>’ before a matrix indicates the matrix positive-definiteness.

Proposition 6 (semismoothness* and directional differentiability of FBE derivatives) *Let $\varphi = f + g$ be as in (5.8), and let $\gamma > 0$ be such that $B := I - \gamma A$ is positive-definite. Then for any $\bar{x} \in \mathbb{R}^n$ satisfying the stationary condition $0 \in \partial \varphi(\bar{x})$ the following assertions hold:*

- (i) $\nabla \varphi_\gamma$ is semismooth* at $\bar{x}$ if ∂g is semismooth* at $(\bar{x}, \bar{v})$, where $\bar{v} := -A\bar{x} - b$.
- (ii) $\nabla \varphi_\gamma$ is directionally differentiable at $\bar{x}$ if g is twice epi-differentiable at $\bar{x}$ for $\bar{v}$.

Proof Denote $h_\gamma(x) := \text{Prox}_{\gamma g}(x - \gamma(Ax + b))$ for all $x \in \mathbb{R}^n$ and get by Proposition 3 that

$$\nabla \varphi_\gamma(x) = \gamma^{-1}(I - \gamma A)(x - h_\gamma(x)) = \gamma^{-1}Bx - \gamma B h_\gamma^{-1}(x), \quad x \in \mathbb{R}^n. \quad (5.16)$$

Since the stationary point $\bar{x}$ is an optimal solution to (5.8) due the convexity of φ , we have that $\bar{x} = h_\gamma(\bar{x})$ by [1, Theorem 27.2]. Since ∂g is semismooth* at $(\bar{x}, \bar{v})$, we

have that $\text{Prox}_{\gamma g}$ is semismooth* at $\bar{x} - \gamma(A\bar{x} + b)$ by using [27, Proposition 6]. It follows from Lemma 5 in the Appendix that h_γ is semismooth* at $\bar{x}$. Employing now (5.16) and [32, Proposition 3.6] tells us that $\nabla\varphi_\gamma$ is semismooth* at $\bar{x}$. This verifies assertion (i).

To proceed with the proof of (ii), observe by [88, Theorem 13.40] that the twice epi-differentiability of g at $\bar{x}$ for $\bar{v}$ amounts to saying that the subgradient mapping ∂g is proto-differentiable at $(\bar{x}, \bar{v})$. Using [27, Corollary 8], we conclude that $\text{Prox}_{\gamma g}$ is directionally differentiable at $\bar{x} - \gamma(A\bar{x} + b)$, which yields in turn the directional differentiability of h_γ at $\bar{x}$. Thus the mapping $\nabla\varphi_\gamma$ is directionally differentiable at $\bar{x}$ due to (5.16). This verifies (ii) and completes the proof of the proposition. $\square$

Now we are ready to describe and then justify the proposed globally coderivative-based damped Newton method for solving the convex composite optimization problem (5.8).

Algorithm 3 Coderivative-based damped Newton algorithm for convex composite optimization

Input: $x^0 \in \mathbb{R}^n$, $\gamma > 0$ such that $B := I - \gamma A > 0$, $\sigma \in (0, \frac{1}{2})$, $\beta \in (0, 1)$, and φ_γ as in (5.4)

```

1: for  $k = 0, 1, \dots$  do
2:   If  $\nabla\varphi_\gamma(x^k) = 0$ , stop. Otherwise set  $u^k := x^k - \gamma(Ax^k + b)$ ,  $v^k := \text{Prox}_{\gamma g}(u^k)$ 
3:   Find  $d^k \in \mathbb{R}^n$  st  $-\frac{1}{\gamma}(x^k - v^k) - Ad^k \in \partial^2 g\left(v^k, \frac{1}{\gamma}(u^k - v^k)\right)(x^k - v^k + d^k)$ 
4:   Set  $\tau_k = 1$ 
5:   while  $\varphi_\gamma(x^k + \tau_k d^k) > \varphi_\gamma(x^k) + \sigma \tau_k \langle \nabla\varphi_\gamma(x^k), d^k \rangle$  do
6:     set  $\tau_k := \beta \tau_k$ 
7:   end while
8:   Set  $x^{k+1} := x^k + \tau_k d^k$ 
9: end for
```

Explicit expressions for the sequences $\{v^k\}$ and $\{d^k\}$ in Algorithm 3 depend on given structures of the regularizers g , which are efficiently specified in applied models of machine learning and statistics; see, e.g., Sect. 6. The next theorem provides explicit sufficient conditions to run Algorithm 3 for solving the class of convex composite optimization problems (5.8).

Theorem 5 (global convergence of coderivative-based damped Newton algorithm in convex composite optimization) *Consider problem (5.8), where the matrix A is positive-definite. Then Algorithm 3 either stops after finitely many iterations, or produces a sequence $\{x^k\}$ such that it globally R -linearly converges to $\bar{x}$, which is the unique solution to (5.8) being a tilt-stable local minimizer of φ with modulus $\kappa := 1/\lambda_{\min}(A)$. Furthermore, the convergence rate of $\{x^k\}$ is at least Q -superlinear if the subgradient mapping ∂g is semismooth* at $(\bar{x}, \bar{v})$, where $\bar{v} := -A\bar{x} - b$, and if either one of two following conditions is satisfied:*

- (i) $\sigma \in (0, 1/(2LK))$, where $L := 2(1 - \gamma\lambda_{\min}(A))/\gamma$ and $K := \kappa + \gamma\|B^{-1}\|$.
- (ii) g is twice epi-differentiable at $\bar{x}$ for $\bar{v}$.

Proof We deduce from Propositions 3(ii) and 4 that solving the convex composite optimization problem (5.8) by Algorithm 3 reduces to solving the $\mathcal{C}^{1,1}$ optimization problem (5.7) by using Algorithm 1. The rest of the proof is split into the following two claims.

Claim 1 Algorithm 3 either stops after finitely many iterations, or produces a sequence $\{x^k\}$ that globally R -linearly converges to the unique solution $\bar{x}$ of (5.8), which is a tilt-stable local minimizer of φ . Indeed, we get from Proposition 3(ii) that φ_γ is a strongly convex function, and its gradient is globally Lipschitz continuous with modulus L . It follows from [10, Theorem 5.1] that $\partial^2 \varphi_\gamma(x)$ is positive-definite for all $x \in \mathbb{R}^n$. Then Theorem 1 tells us that Algorithm 3 either stops after finitely many iterations, or produces a sequence $\{x^k\}$ that globally R -linearly converges to $\bar{x}$, which is a stationary point of φ_γ . Employing again Proposition 3(ii) confirms that $\bar{x}$ is an optimal solution to (5.8). Furthermore, the strong convexity of φ with modulus $\lambda_{\min}(A)$ and Lemma 6 from the Appendix yield the uniqueness and tilt stability conclusions for $\bar{x}$.

Claim 2 The Q -superlinear convergence of $x^k \rightarrow \bar{x}$ holds under the assumptions of the theorem. The imposed semismooth* property of ∂g at $(\bar{x}, \bar{v})$ ensures the fulfillment of this property for $\nabla \varphi_\gamma$ at $\bar{x}$ by Lemma 6. Assume now that condition (i) of the theorem is satisfied. Then we get from Claim 1 that L is a Lipschitz constant of $\nabla \varphi_\gamma$ around $\bar{x}$. As follows from [22, Proposition 4.5], the modulus of tilt-stability of the l.s.c. convex function φ under consideration at $\bar{x}$ is the same as the modulus of metric regularity of φ around this point. Combining the latter with the statement of Proposition 5(i) and the precise calculation in (2.5) of the exact bound of metric regularity, we conclude that $\bar{x}$ is tilt-stable local minimizer of φ_γ with modulus K . Thus the claimed assertion on the superlinear convergence in this case follows directly from Theorem 2. Assuming finally by (ii) that g is twice epi-differentiable at $\bar{x}$ for $\bar{v}$, we deduce from Proposition 6(ii) that the mapping $\nabla \varphi_\gamma$ is directionally differentiable at $\bar{x}$. Thus the claimed superlinear convergence of $\{x^k\}$ follows in this case from the corresponding statement of Theorem 2. This completes the proof of the theorem. $\square$

Theorem 5 and the results of numerical experiments in Sect. 6 show that Algorithm 3, designed in terms of the computable data of (5.8), exhibits an excellent performance when A is positive-definite, i.e., in the strongly convex setting of (5.8). Otherwise, this algorithm is not even well-defined. To relax this positive-definiteness/strong convex assumption, we now propose and justify a new coderivative-based algorithm of the regularized Newton type, which is well-defined and globally convergent to solutions of (5.8) for merely *positive-semidefinite* matrices A with linear and superlinear convergence rates under some additional assumptions that include the *metric regularity* of the subgradient mapping $\partial \varphi$. Observe that the latter assumption is *weaker* than the *strong convexity assumption* on f in problems of composite optimization (5.1), including those with quadratic functions f as in (5.8). A simple class of functions φ illustrating this observation is given by $\varphi = f + |x|$. In particular, for $f \equiv 0$ we clearly have that

$$\partial^2 \varphi(0, 0)(v) = \{w \in \mathbb{R} \mid (w, -v) \in \mathbb{R} \times \{0\}\},$$

and thus $0 \in \partial^2 \varphi(0, 0)(v) \implies v = 0$, which tells us by (2.4) that $\partial \varphi$ is metrically regular around $(0, 0)$.

Here is the aforementioned algorithm, which is more complicated than Algorithm 3, while being applied for problems (5.8) with positive-semidefinite matrices A . Note that the new algorithm *does not* require performing operations like computing inverse matrices that are expensive in large dimensions.

Algorithm 4 Coderivative-based regularized Newton algorithm for convex composite optimization

Input: $x^0 \in \mathbb{R}^n$, $\gamma > 0$ such that $B := I - \gamma A > 0$, $\lambda > 0$, $\sigma \in (0, \frac{1}{2})$, $\beta \in (0, 1)$, and φ_γ as in (5.4)

```

1: for  $k = 0, 1, \dots$  do
2:   If  $\nabla \varphi_\gamma(x^k) = 0$ , stop. Otherwise set  $u^k := x^k - \gamma(Ax^k + b)$ ,  $v^k := \text{Prox}_{\gamma g}(u^k)$ ,  $\mu_k := \lambda \|\nabla \varphi_\gamma(x^k)\|$ 
3:   Find  $z^k \in \mathbb{R}^n$  st  $-\frac{1}{\gamma}(x^k - v^k) - (\mu_k I + AB)z^k \in \partial^2 g\left(v^k, \frac{1}{\gamma}(u^k - v^k)\right)(x^k - v^k + (B + \gamma \mu_k I)z^k)$ 
4:   Set  $d^k = Bz^k$ 
5:   Set  $\tau_k = 1$ 
6:   while  $\varphi_\gamma(x^k + \tau_k d^k) > \varphi_\gamma(x^k) + \sigma \tau_k \langle \nabla \varphi_\gamma(x^k), d^k \rangle$  do
7:     set  $\tau_k := \beta \tau_k$ 
8:   end while
9:   Set  $x^{k+1} := x^k + \tau_k d^k$ 
10: end for
    
```

The next theorem fully describes the well-posedness and performance of Algorithm 4.

Theorem 6 (global convergence of coderivative-based regularized Newton algorithm in convex composite optimization) *Consider problem (5.8) of convex composite optimization, where the matrix A is positive-semidefinite. Then we have the assertions:*

- (i) *Algorithm 4 either stops after finitely many iterations, or produces a sequence $\{x^k\}$ for which all the accumulation points of this sequence are optimal solutions to (5.8).*
- (ii) *If in addition the subgradient mapping $\partial \varphi$ is metrically regular around $(\bar{x}, 0)$ with modulus $\kappa > 0$, where $\bar{x}$ is an accumulation point of $\{x^k\}$, then the sequence $\{x^k\}$ globally R -linearly converges to $\bar{x}$, and $\bar{x}$ is a tilt-stable local minimizer of φ with modulus κ .*
- (iii) *The rate of convergence of $\{x^k\}$ is at least Q -superlinear if the subgradient mapping ∂g is semismooth* at $(\bar{x}, \bar{v})$, where $\bar{v} := -A\bar{x} - b$, and if one of two following conditions holds:*

- (a) $\sigma \in (0, 1/(2LK))$, where $L := 2(1 - \gamma \lambda_{\min(A)})/\gamma$ and $K := \kappa + \gamma \|B^{-1}\|$.
- (b) g is twice epi-differentiable at $\bar{x}$ for $\bar{v}$.

Proof Using Propositions 3(ii) and 4, we can reduce Algorithm 4 for solving the convex composite optimization problem (5.8) to Algorithm 2 for solving the $\mathcal{C}^{1,1}$ optimization problem (5.7). Let us now proceed with the justification of this procedure by verifying each claim of the theorem.

Claim 1 *Assertion (i) holds.* Indeed, we have from Proposition 3(ii) that the gradient mapping $\nabla \varphi_\gamma$ is globally Lipschitz continuous with modulus L and φ_γ is a convex

function; thus the generalized Hessian $\partial^2\varphi_\gamma(x)$ is positive-semidefinite for all $x \in \mathbb{R}^n$ by [10, Theorem 3.2]. Therefore, it follows from Theorem 3 that Algorithm 4 either stops after finitely many iterations, or produces a sequence $\{x^k\}$ whose accumulation points are solutions to (5.7). This verifies assertion (i).

Claim 2 *Assertion (ii) holds.* Under the assumptions made in (ii), it follows from the established relationships between problems (5.8) and (5.7) and the application of Theorem 4 to the latter that the sequence $\{x^k\}$ globally R -linearly converges to $\bar{x}$ as $k \rightarrow \infty$. Then we get by [22, Proposition 4.5] that the tilt-stability of φ at $\bar{x}$ with modulus κ follows from the metric regularity of $\partial\varphi$ and the convexity of φ , which therefore verifies (ii).

Claim 3 *Assertion (iii) holds.* We deduce from Proposition 6(i) that the semismoothness* of ∂g at $(\bar{x}, \bar{v})$ yields this property for $\nabla\varphi_\gamma$ at $\bar{x}$. Assuming first that condition (a) is satisfied, we deduce from Proposition 3(ii) that L is a Lipschitz constant of $\nabla\varphi_\gamma$ around $\bar{x}$. It follows from [22, Proposition 4.5] that the modulus of tilt-stability of the l.s.c. convex function φ at $\bar{x}$ is equal to the modulus of metric regularity of φ around this point. Combining the latter with Proposition 5(i) and the calculation formula for the exact regularity bound in (2.5) tells us that $\bar{x}$ is a tilt-stable local minimizer of φ_γ with modulus K . Hence assertion (iii) in case (a) follows from Theorem 4(a). Assuming now (b) implies by Proposition 6(ii) that $\nabla\varphi_\gamma$ is directionally differentiable at $\bar{x}$. Applying finally Theorem 4(b) to problem (5.7) ensures the Q -superlinear convergence of sequence $x^k \rightarrow \bar{x}$ as $k \rightarrow \infty$ and therefore completes the proof of the theorem. $\square$

Remark 6 (second-order subdifferential computations for regularizers) One of the crucial steps in implementing Algorithms 3 and 4 is deriving explicit second-order subdifferential computations for the regularizers $g : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ in the quadratic composite optimization problem (5.8). This has been accomplished in many publications, some of which we mention below for the reader's convenience and further applications. When g is the indicator function of general *polyhedral sets*, explicit formulas of different types for $\partial^2 g$ are derived in [19, 35, 36, 95]. In the case of *moving polyhedra*, such computations are provided in [13, 70, 84] with applications to optimal control of sweeping processes in [13] among other publications. When g is described by *non-linear inequality systems*, $\partial^2 g$ is efficiently evaluated in [37], and for various types of *maximum functions* the computations $\partial^2 g$ are achieved in [23, 67]. Furthermore, in [67], the reader can find explicit computation formulas for $\partial^2 g$ addressing the general class of extended-real-valued *convex piecewise linear* functions, while papers [65, 66] contain second-order subdifferential computations for some different subclasses of piecewise linear-quadratic functions. Complete computations of $\partial^2 g$ for the indicator function of the *Lorentz cone* in second-order cone programming are provided in [64, 73]. Finally in this list, we mention the most involved second-order subdifferential computations accomplished in [17] for general classes of the general class of nonpolyhedral systems including *semidefinite cone complementarity constraints*.

6 Applications and numerical experiments

This section provides numerical experiments for our methods and their comparison with the classical semismooth Newton methods in $\mathcal{C}^{1,1}$ optimization. Furthermore, we also provide applications of the developed generalized Newton algorithms to problems of nonsmooth convex composite optimization and their comparison with some well-recognized methods of nonsmooth optimization in some classes of practical models.

6.1 Testing $\mathcal{C}^{1,1}$ optimization problems

The first subsection here is devoted to comparing our approach with the well-recognized semismooth Newton method (SNM) to solve the *testing optimization problem*:

$$\text{minimize } \varphi(x) := e_f(x) + \frac{1}{2}\|x\|^2 \quad \text{subject to } x \in \mathbb{R}^n, \quad (6.1)$$

where $e_f(x)$ is the Moreau envelope with the parameter $\lambda = 1$ of the maximum function given by

$$f(x) := \max \{x_1, \dots, x_n\} \quad \text{for all } x = (x_1, \dots, x_n) \in \mathbb{R}^n. \quad (6.2)$$

Due to [1, Proposition 12.30], $e_f(x)$ is continuously differentiable with the Lipschitzian gradient, and (6.1) belongs to the class of $\mathcal{C}^{1,1}$ optimization problems (5.8).

The following lemma is needed for implementing our generalized damped Newton algorithm (GDNM) to solve the formulated optimization problem (6.1).

Lemma 2 (proximal mapping and second-order subdifferentials of maximum functions) *Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be taken from (6.2). Then the proximal mapping of f is calculated by*

$$\text{Prox}_f(x) = (\min\{x_1, s\}, \dots, \min\{x_n, s\}), \quad (6.3)$$

where the number $s \in \mathbb{R}$ is such that

$$\sum_{i=1}^n \max \{0, x_n - s\} = 1. \quad (6.4)$$

For each $x \in \mathbb{R}^n$ and $y \in \partial f(x)$, define the index sets

$$J(x) := \{i \in \{1, \dots, n\} \mid x_i = f(x)\} \quad \text{and} \quad L(x) := \{i \in J(x) \mid y_i > 0\}.$$

The calculation formula for the generalized Hessian of f is

$$w \in \partial^2 f(x, y)(v) \iff v_i = c \quad \forall i \in L(x),$$

$$\sum_{i=1}^n w_i = 0, \quad w_i \geq 0 \quad \forall i \in J_{>}(x), \quad w_i = 0 \quad \forall i \in J^c(x) \cup J_{<}(x), \quad (6.5)$$

where $c \in \mathbb{R}$ is an arbitrary constant, and where

$$J_{>}(x) := \{i \in J(x) \mid v_i > c\}, \quad J_{<}(x) := \{i \in J(x) \mid v_i < c\}.$$

Proof The proof for the proximal mapping formula (6.3) can be found in [3, Example 6.49]. The calculation formula for the second-order subdifferential of the maximum function (6.5) is obtained in [23, Theorem 3.1]. $\square$

We start implementing Algorithms 1 to solve (6.1) with the explicit computation of their ingredients (gradient and second-order subdifferential of φ) given entirely via the problem data. More specifically, we need to explicitly determine the gradient $\nabla\varphi(x)$ and the coderivative-based Newton direction $d \in \mathbb{R}^n$ generated by the inclusion

$$-\nabla\varphi(x) \in \partial^2\varphi(x)(d). \quad (6.6)$$

The expressions for the gradient follows directly from formula (6.3) telling us that

$$\nabla\varphi(x) = \nabla e_f(x) + x = 2x - \text{Prox}_f(x). \quad (6.7)$$

Now we show how to construct $d \in \mathbb{R}^n$ satisfying inclusion (6.6). For each $x \in \mathbb{R}^n$, define the vector d by

$$d_i := \begin{cases} \frac{1}{2}(c_x - (\nabla\varphi(x))_i) & \text{if } i \in L(x), \\ (-\nabla\varphi(x))_i & \text{otherwise,} \end{cases} \quad (6.8)$$

where the index set $L(x)$ and the number c_x are given by

$$L(x) := \{i \in \{1, \dots, n\} \mid (\text{Prox}_f(x))_i = f(\text{Prox}_f(x)), (x - \text{Prox}_f(x))_i > 0\},$$

$$c_x := -\frac{1}{|L(x)|} \sum_{i \in L(x)} (\nabla\varphi(x))_i.$$

The next proposition verifies that the vector d constructed in (6.8) indeed satisfies inclusion (6.6).

Proposition 7 Let $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}$ be given in (6.1). For each $x \in \mathbb{R}^n$, consider the vector $d \in \mathbb{R}^n$ defined in (6.8). Then inclusion (6.6) holds for this vector d .

Proof Applying the second-order subdifferential sum rule from [59, Proposition 1.121] to φ in (6.1) gives us

$$\partial^2\varphi(x)(w) = \partial^2 e_f(x)(w) + w \quad \text{for all } w \in \mathbb{R}^n.$$

This tells us that the inclusion $-\nabla\varphi(x) \in \partial^2\varphi(x)(w)$ is equivalent to

$$-\nabla\varphi(x) - w \in \partial^2 e_f(x)(w).$$

It follows from [46, Lemma 6.4] that the latter inclusion can be equivalently rewritten as

$$-\nabla\varphi(x) - w \in \partial^2 f(\text{Prox}_f(x), x - \text{Prox}_f(x))(\nabla\varphi(x) + 2w). \quad (6.9)$$

We only need to show that $w := d$, for d defined in (6.8), solves (6.9). Indeed, fixing any indices $i \in L(x)$ and $j \in \{1, \dots, n\} \setminus L(x)$ gives us the equalities

$$(\nabla\varphi(x) + 2d)_i = c_x, \quad (-\nabla\varphi(x) - d)_j = 0, \quad \text{and} \quad (6.10)$$

$$\begin{aligned} \sum_{i=1}^n (-\nabla\varphi(x) - d)_i &= \sum_{i \in L(x)} (-\nabla\varphi(x) - d)_i + \sum_{i \in \{1, \dots, n\} \setminus L(x)} (-\nabla\varphi(x) - d)_i \\ &= -\frac{1}{2} \sum_{i \in L(x)} (c_x + (\nabla\varphi(x))_i) = 0. \end{aligned} \quad (6.11)$$

Combining (6.10) and (6.11), we deduce from Lemma 2 that d solves (6.9) and thus complete the proof. $\square$

Next we present specifications of Algorithm 1 and Theorems 5, 6 on its performance for the case of the testing optimization problem (6.1).

Theorem 7 (solving the testing $\mathcal{C}^{1,1}$ optimization problem by GDNM) *Consider the testing optimization problem (6.1). Then Algorithm 1, with all its ingredients calculated in (6.7) and (6.8), either stops after finitely many iterations, or produces a sequence $\{x^k\}$ such that it globally Q -superlinearly converges to $\bar{x}$, which is the unique solution to problem (6.1), while being also a tilt-stable local minimizer of the function φ from (6.1). Furthermore, the sequences $\{\varphi(x^k)\}$ and $\{\nabla\varphi(x^k)\}$ Q -superlinearly converge to $\min \varphi$ and 0, respectively.*

Proof It is clear that the cost function φ in (6.1) is strongly convex, and thus the generalized Hessian $\partial^2\varphi(x)$ is positive-definite for all $x \in \mathbb{R}^n$. Furthermore, it follows from (6.7) and [26, Proposition 7.4.6] that the mapping $\nabla\varphi$ is semismooth on $\mathbb{R}^n$ due to its piecewise linearity. Applying now Theorems 1, 2 and Proposition 7 to this setting, we arrive at all the conclusions of the theorem. $\square$

Due to Remark 3(i), the underlying difference between our approach and the conventional semismooth Newton method is in finding the *generalized Newton directions*. To be more specific, the semismooth Newton method to solve (6.1) is based on using generalized Jacobian of $\nabla\varphi$ to determine the Newton directions d as solutions to the system of linear equations

$$-\nabla\varphi(x) = Ad, \quad \text{where } A \in \partial_C \nabla\varphi(x). \quad (6.12)$$

We have therefore in the setting of (6.1) that

$$\partial_C \nabla \varphi(x) = 2I - \partial_C \text{Prox}_f(x)$$

Consider further the index sets

$$J(x) := \{i \in \{1, \dots, n\} \mid \max\{0, x_i - s\} = 0\}, \quad \text{and} \quad J^c(x) := \{1, \dots, n\} \setminus J(x),$$

where $s \in \mathbb{R}$ is such that (6.4) holds. It follows from [77, Example 5.4] that an element of $\partial_C \text{Prox}_f(x)$ is the matrix P with the entries

$$P_{ij} := \begin{cases} 1 + 1/(n - |J|) & \text{if } i \neq j, i, j \in J^c(x), \\ 1/(n - |J|) & \text{if } i = j, i, j \in J^c(x), \\ 1 & \text{if } i, j \in J(x). \end{cases}$$

This tells us that finding a Newton direction d in SNM requires solving the following system of linear equation:

$$-\nabla \varphi(x) = (2I - P)d. \quad (6.13)$$

Now we are ready to conduct numerical experiments to solve the testing problem (6.1) by using our algorithm GDNM and the semismooth Newton method SNM. All the numerical experiments are conducted on a desktop with 10th Gen Intel(R) Core(TM) i5-10,400 processor (6-Core, 12 M Cache, 2.9 GHz to 4.3 GHz) and 16 GB memory. All the codes are written in MATLAB 2016a. We test with different dimensions n ranging from 200 to 2000. The stopping criterion $\|\nabla \varphi(x^k)\| \leq 10^{-6}$ is used in all the tests. The initial points for GDNM and SNM are the same being chosen as a randomly generated vector in $\mathbb{R}^n$ with i.i.d. (identically and independent distributed) standard Gaussian entries. The results are shown in Fig. 1. As we can see from the results presented in this figure, our algorithm GDNM is highly efficient in solving problem (6.1).

A clear explanation for this is that in the setting under consideration, we can find a precise formula for generalized Newton directions from the second-order subdifferential inclusion (6.6) while SNM requires solving the system of linear equation (6.13). Typically, although solving the inclusion (6.6) might be more difficult than solving the system of linear equation (6.12), observe that calculus rules for the second-order subdifferential $\partial^2 \varphi$ are much more developed than in the generalized Jacobian case $\partial_C \nabla \varphi$ of the semismooth Newton method (6.12).

6.2 Solving Lasso problems

This subsection is devoted to specifying both Algorithms 3 and 4 for the case of the basic Lasso problem stated below, and then to conducting numerical experiments for this problem and comparing them with the performances of some major first-order and second-order algorithms. The *basic Lasso problem*, known also as the ℓ^1 -regularized

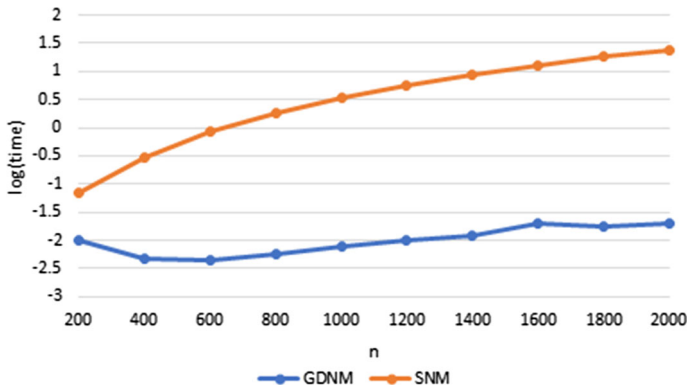


Fig. 1 CPU time for GDNM and SNM in the $C^{1,1}$ optimization problem (6.1)

least square optimization problem, was introduced by Tibshirani [92], and since then it has been largely investigated and applied to various issues in statistics, machine learning, image processing, etc. This problem is formulated as follows:

$$\text{minimize } \varphi(x) := \frac{1}{2} \|Ax - b\|_2^2 + \mu \|x\|_1 \quad \text{subject to } x \in \mathbb{R}^n, \quad (6.14)$$

where A is an $m \times n$ matrix, $\mu > 0$, and $b \in \mathbb{R}^m$, and where $\|\cdot\|_1$ and $\|\cdot\|_2$ stand for the standard p -norms on $\mathbb{R}^n$. It is easy to see that the Lasso problem (6.14) belongs to the class of convex composite optimization problems (5.8). Indeed, we can represent (6.14) as minimizing the nonsmooth convex function $\varphi(x) := f(x) + g(x)$, where

$$f(x) := \frac{1}{2} \langle \tilde{A}x, x \rangle + \langle \tilde{b}, x \rangle + \tilde{\alpha} \quad \text{and} \quad g(x) := \mu \|x\|_1 \quad (6.15)$$

with $\tilde{A} := A^*A$, $\tilde{b} := -A^*b$, and $\tilde{\alpha} := \frac{1}{2} \|b\|^2$, and where the matrix $\tilde{A} = A^*A$ is symmetric and positive-semidefinite. Observe that the Lasso problem (6.14) always admits an optimal solution [92].

We start implementing Algorithms 3 and 4 to solve (6.14) with the explicit computation of their ingredients (proximal and subgradient mappings, generalized Hessian of g) given entirely via the problem data.

Proposition 8 (explicit computations for the Lasso problem) *Let $g(x) = \mu \|x\|_1$ be the regularizer in the Lasso problem (6.14). Then we have the calculation formulas:*

$$(\text{Prox}_{\gamma g}(x))_i = \begin{cases} x_i - \mu\gamma & \text{if } x_i > \mu\gamma, \\ 0 & \text{if } -\mu\gamma \leq x_i \leq \mu\gamma, \\ x_i + \mu\gamma & \text{if } x_i < -\mu\gamma. \end{cases} \quad (6.16)$$

$$\partial g(x) = \left\{ v \in \mathbb{R}^n \mid \begin{array}{l} v_i = \mu \cdot \text{sgn}(x_i), \ x_i \neq 0, \\ v_i \in [-\mu, \mu], \ x_i = 0 \end{array} \right\} \quad \text{whenever } x \in \mathbb{R}^n. \quad (6.17)$$

The generalized Hessian of g is calculated by

$$\partial^2 g(x, y)(v) = \left\{ w \in \mathbb{R}^n \mid \left(\frac{1}{\mu} w_i, -v_i \right) \in G\left(x_i, \frac{1}{\mu} y_i\right), i = 1, \dots, n \right\} \quad (6.18)$$

for each $(x, y) \in \text{gph } \partial g$ and $v = (v_1, \dots, v_n) \in \mathbb{R}^n$, where the mapping $G: \mathbb{R}^2 \rightrightarrows \mathbb{R}^2$ is defined by

$$G(t, p) := \begin{cases} \{0\} \times \mathbb{R} & \text{if } t \neq 0, p \in \{-1, 1\}, \\ \mathbb{R} \times \{0\} & \text{if } t = 0, p \in (-1, 1), \\ (\mathbb{R}_+ \times \mathbb{R}_-) \cup (\{0\} \times \mathbb{R}) \cup (\mathbb{R} \times \{0\}) & \text{if } t = 0, p = -1, \\ (\mathbb{R}_- \times \mathbb{R}_+) \cup (\{0\} \times \mathbb{R}) \cup (\mathbb{R} \times \{0\}) & \text{if } t = 0, p = 1, \\ \emptyset & \text{otherwise.} \end{cases} \quad (6.19)$$

Proof The formula for the proximal mapping (6.16) follows from definition (5.3) and the form of $g(\cdot) = \|\cdot\|_1$. The calculations of ∂g and $\partial^2 g$ are taken from [46, Propositions 7.1 and 7.2], respectively. $\square$

Let us present specifications of Algorithms 3 and 4 as well as Theorems 5 and 6 on their performances, respectively, for the Lasso problem (6.14).

Theorem 8 (solving Lasso) *Considering the Lasso problem (6.14), we have the following:*

- (i) *Algorithm 3, with all its ingredients calculated in Proposition 8, either stops after finitely many iterations, or produces a sequence $\{x^k\}$ such that it globally Q -superlinearly converges to $\bar{x}$, which is the unique solution to (6.14) and a tilt-stable local minimizer of φ with modulus $\kappa := 1/\lambda_{\min}(A^*A)$, provided that the matrix A^*A is positive-definite.*
- (ii) *Algorithm 4, with the positive-semidefinite matrix A^*A and the ingredients calculated in Proposition 8, either stops after finitely many steps, or produces a sequence $\{x^k\}$ such that any accumulation point $\bar{x}$ of it is a solution to (6.14). If in addition $\partial\varphi$ is metrically regular around $(\bar{x}, 0)$ with modulus $\kappa > 0$, then the sequence $\{x^k\}$ globally Q -superlinearly converges to $\bar{x}$, which is a tilt-stable local minimizer of φ with the same modulus.*

Proof Observe by (6.17) that the graph of ∂g is the union of finitely many closed convex sets, and hence ∂g is semismooth* on its graph; see [32]. Furthermore, g is a proper, convex, and piecewise linear function on $\mathbb{R}^n$. Then it follows from [88, Proposition 13.9] that g is twice epi-differentiable on $\mathbb{R}^n$. Applying Theorems 5 and 6, we arrive at all the conclusions in (i) and (ii) of this theorem, respectively. $\square$

To run Algorithms 3 and 4, we need to explicitly determine the sequences $\{v^k\}$ and $\{d^k\}$ generated by these algorithms. The expressions for v^k follows directly from formula (6.16) telling us that

$$\left(v^k\right)_i = \begin{cases} (u^k)_i - \mu\gamma & \text{if } (u^k)_i > \mu\gamma, \\ 0 & \text{if } -\mu\gamma \leq (u^k)_i \leq \mu\gamma, \\ (u^k)_i + \mu\gamma & \text{if } (u^k)_i < -\mu\gamma, \end{cases} \quad \text{where } u^k = x^k - \gamma(A^*Ax^k + b).$$

Using further the formulas in (6.17)–(6.19), we express d^k in Algorithm 3 via the conditions

$$\begin{cases} \left(-\frac{1}{\gamma}(x^k - v^k) - A^*Ad^k\right)_i = 0 & \text{if } (v^k)_i \neq 0, \\ (x^k - v^k + d^k)_i = 0 & \text{if } (v^k)_i = 0. \end{cases}$$

Thus d^k can be computed for each $k \in \mathbb{N}$ by solving the linear equation $X^k d = v^k - x^k$, where

$$(X^k)_i := \begin{cases} \gamma(A^*A)_i & \text{if } (v^k)_i \neq 0, \\ I_i & \text{if } (v^k)_i = 0. \end{cases} \quad (6.20)$$

Similarly, by employing the calculations of Proposition 8 in the framework of Algorithm 4 and by performing elementary transformations, we get the linear equation

$$(BX^k + \gamma\mu_k I)d^k = B(v^k - x^k)$$

to find the direction d^k , where X^k is computed in (6.20), and $B := I - \gamma A^*A$.

Now we are ready to conduct numerical experiments for solving the Lasso problem (6.14) by using our globally convergent coderivative-based Generalized Damped Newton Method (GDNM) via Algorithm 3 and globally convergent coderivative-based Generalized Regularized Newton Method (GRNM) via Algorithm 4. The obtained calculations are compared with those obtained by implementing the following highly recognized first-order and second-order algorithms:

- (i) The Alternating Direction Methods of Multipliers¹ (ADMM); see [8, 28, 29].
- (ii) The Fast Iterative Shrinkage-Thresholding Algorithm² (FISTA) with the code presented in [4].
- (iii) The Semismooth Newton Augmented Lagrangian Method³ (SSNAL) recently developed in [51].

All the numerical experiments are conducted in the same desktop and software described in Sect. 6.1. All the codes are written in MATLAB 2016a. In our numerical experiment, A is generated randomly with i.i.d. (independent and identically distributed) standard Gaussian entries, where b is generated randomly with values of components are from 0 to 1. In some particular tests, we normalize each column of A so that A^*A is close to singular and mark them with symbol * for identification. In summary, A^*A is nonsingular in Tests 3, 4, 7, 8, and it is singular or close to singular in all the other tests. Table 1 contains 2 tests where $n > m$ and the matrix A^*A is nonsingular, 2 tests where $n > m$ and the matrix A^*A is singular, 2 tests where $n = m$ and the matrix A^*A is nonsingular, 2 tests where $n = m$ and the matrix A^*A is singular, and 2 tests when $m > n$. To simplify the numerical implementations for solving the Lasso problem (6.14), we set $\mu := 10^{-3}$ as the tuning parameter for all the tests. If

¹ <https://web.stanford.edu/boyd/papers/admm/lasso/lasso.html>.

² <https://github.com/he9180/FISTA-lasso>.

³ <https://www.polyu.edu.hk/ama/profile/dfsun/>.

an algorithm cannot start iterating in the first step, it is marked by ‘Error’ word; this concerns only some cases of GDNM when the matrix A^*A is not positive-definite. In our numerical experiments, $x^0 := 0$ is the starting point for each algorithm, and the following *relative KKT residual* η_k in (6.21) suggested in [51] is used to measure the accuracy of an approximate optimal solution x^k for (6.14):

$$\eta_k := \frac{\|x^k - \text{Prox}_{\mu\|\cdot\|_1}(x^k - A^*(Ax^k - b))\|}{1 + \|x^k\| + \|Ax^k - b\|}. \quad (6.21)$$

We stop the algorithms when either the condition $\eta_k < 10^{-6}$ is satisfied, or the maximum computation time of 10,000 s is reached. The results of computations for this part are displayed in Table 1. There ‘TN’ stands for the test number, ‘iter’ indicates the number of performed iterations, and ‘CPU time’ stands for the time needed to achieve the prescribed accuracy of approximate solutions (the smaller the better).

As we can see from the results presented in Table 1, our algorithms GDNM and GRNM are highly efficient when A^*A is *nonsingular*, where the Q -superlinear convergence is guaranteed by Theorem 8. They may behave even better than the other compared algorithms when $m \geq n$, which is the setting of various practically important models; see, e.g., [24] for Lasso applications to diabetes studies where m is much larger than n , and [4] for $m = n$ with applications to image processing.

When the matrix A^*A is *singular* (or close to be singular), our theoretical results do not guarantee the fast convergence of GDNM and GRNM, while the conducted numerical experiments show that GRNM performs better than GDNM and better than FISTA and ADMM in Table 1, while usually worse than SSNAL. A partial explanation for this is that SSNAL is actually a *hybrid* algorithm, which combines the first-order augmented Lagrangian method to solve *dual* subproblems, which are strongly convex and of lower dimensions, with the subsequent applications of the second-order semismooth Newton method. Such a combination exhibits a high efficiency in solving Lasso problems in the singular case.

Table 1 Solving (6.14) on random instances

Problem size and TN			Iter		CPU time							
TN	m	n	SSNAL	FISTA	ADMM	GRNM	GDNM	SSNAL	FISTA	ADMM	GRNM	GDNM
1	400	800	25	37,742	22,873	1813	Error	0.45	145.52	10.89	45.62	Error
2	4000	8000	153	19,173	19,173	2499	Error	847.87	10,000	2359.36	10,000	Error
3	2000	2000	43	239,701	12,785	59	12	78.38	8138.94	158.12	11.07	2.24
4	4000	4000	246	73,374	5970	59	218	1253.45	10,000	320.81	48.16	178.91
5*	2000	2000	22	3619	90,501	394	292	18.11	123.38	1141.64	65.60	58.80
6*	4000	4000	24	3629	103,868	520	555	231.40	462.53	5166.16	369.27	474.74
7	800	400	4	430	10	6	3	0.14	0.86	0.02	0.11	0.08
8	8000	4000	13	487	11	7	3	18.80	117.92	3.67	8.46	4.39
9*	800	400	11	245	426	31	7	0.18	0.53	0.12	0.23	0.11
10*	8000	4000	11	238	411	72	9	8.37	59.18	32.17	56.37	8.88

6.3 Box constrained quadratic programming

This subsection is devoted to specifying Algorithms 3 and 4 for quadratic programming problems of the form

$$\begin{aligned} \text{minimize} \quad & \varphi(x) := \frac{1}{2} \langle Ax, x \rangle + \langle b, x \rangle \\ \text{subject to} \quad & l_i \leq x_i \leq L_i \text{ for all } i = 1, \dots, n \end{aligned} \quad (6.22)$$

where A is an $n \times n$ positive-semidefinite matrix, and where $b, l, L \in \mathbb{R}^n$ are such that $l_i \leq L_i$ as $i = 1, \dots, n$.

Proposition 9 *Considering the indicator function $\delta_\Omega : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ of the set*

$$\Omega := \{x \in \mathbb{R}^n \mid \ell \leq x \leq L\}, \quad (6.23)$$

we have the precise calculation formulas

$$\text{Prox}_{\gamma g}(x) = P_\Omega(x) = \left(\min\{\max\{x_i, \ell_i\}, L_i\} \right)_{i=1}^n \text{ for } x \in \mathbb{R}^n, \gamma > 0, \quad (6.24)$$

$$\partial \delta_\Omega(x) = N_\Omega(x) = F_1(x_1) \times \dots \times F_n(x_n) \text{ for } x \in \Omega, \text{ where} \quad (6.25)$$

$$F_i(t) = \begin{cases} (-\infty, 0] & \text{if } t = \ell_i, \\ [0, \infty) & \text{if } t = L_i, \\ \{0\} & \text{if } t \in (\ell_i, L_i). \end{cases}$$

The generalized Hessian of δ_Ω is calculated by

$$\partial^2 \delta_\Omega(x, y)(v) = \{w \in \mathbb{R}^n \mid (w_i, -v_i) \in G_i(x_i, y_i), i = 1, \dots, n\}, \quad (6.26)$$

where the mappings $G_i : \mathbb{R}^2 \rightrightarrows \mathbb{R}^2, i = 1, \dots, n$, are defined by

$$G_i(t, p) := \begin{cases} \mathbb{R} \times \{0\} & \text{if } t = \ell_i, p < 0, \\ (\mathbb{R}_- \times \mathbb{R}_+) \cup (\mathbb{R} \times \{0\}) \cup (\{0\} \times \mathbb{R}) & \text{if } t = \ell_i, p = 0, \\ \{0\} \times \mathbb{R} & \text{if } t \in (\ell_i, L_i), p = 0, \\ (\mathbb{R}_+ \times \mathbb{R}_-) \cup (\{0\} \times \mathbb{R}) \cup (\mathbb{R} \times \{0\}) & \text{if } t = L_i, p = 0, \\ \mathbb{R} \times \{0\} & \text{if } t = L_i, p > 0, \\ \emptyset & \text{otherwise.} \end{cases} \quad (6.27)$$

Proof The formula for the proximal mapping (6.24) follows from [3, Lemma 6.26]. Note that

$$\delta_\Omega(x) = \delta_{\Omega_1}(x_1) + \dots + \delta_{\Omega_n}(x_n) \text{ for all } x = (x_1, \dots, x_n) \in \mathbb{R}^n,$$

where $\Omega_i := [\ell_i, L_i] = \{t \in \mathbb{R} \mid \ell_i \leq t \leq L_i\}$, $i = 1, \dots, n$. Using [88, Exercise 8.14 and Example 6.10], we obtain (6.25). It remains to verify the second-order subdifferential formula (6.26) for δ_Ω at $(x, y) \in \text{gph } \partial\delta_\Omega$. Observe that $N_{\text{gph } \partial\delta_{\Omega_i}} = G_i$ for all $i = 1, \dots, n$. This allows us to deduce from [63, Theorem 4.3] the representation

$$\partial^2\delta_\Omega(x, y)(v) = \{w \in \mathbb{R}^n \mid (w_i, -v_i) \in N_{\text{gph } \partial\delta_{\Omega_i}}(x_i, y_i), i = 1, \dots, n\},$$

which therefore justifies the fulfillment of (6.26) and completes the proof of the proposition. $\square$

Next we obtain specifications of Algorithms 3 and 4 together with Theorem 5 and 6 on their performances, respectively, for the case of box constrained quadratic programming in (6.22).

Theorem 9 (solving box constrained quadratic programs) *Considering the quadratic programming problem (6.22), we have the following assertions:*

- (i) *Assume that the matrix A is positive-definite. Then Algorithm 3, with all its ingredients calculated in Proposition 9, either stops after finitely many iterations, or produces a sequence $\{x^k\}$ such that it globally Q -superlinearly converges to $\bar{x}$, which is the unique solution to (6.22) and a tilt-stable local minimizer of φ with modulus $\kappa := 1/\lambda_{\min}(A^*A)$.*
- (ii) *Assume that the matrix A is positive-semidefinite. Then Algorithm 4, with all its ingredients calculated in Proposition 9, either stops after finitely many iterations, or produces a sequence $\{x^k\}$ such that any accumulation point $\bar{x}$ of it is a solution to (6.22). If in addition $\partial\varphi$ is metrically regular around $(\bar{x}, 0)$ with some modulus $\kappa > 0$, then the sequence $\{x^k\}$ globally Q -superlinearly converges to $\bar{x}$, which is a tilt-stable local minimizer of φ with the same modulus.*

Proof Reduce (6.22) to the equivalent form of convex composite optimization:

$$\begin{aligned} &\text{minimize } f(x) + g(x) \quad \text{subject to } x \in \mathbb{R}^n, \text{ where} \\ &f(x) := \frac{1}{2}\langle Ax, x \rangle + \langle b, x \rangle, \quad g(x) := \delta_\Omega(x), \quad \Omega := \{x \in \mathbb{R}^n \mid \ell \leq x \leq L\}. \end{aligned}$$

Observe by (6.25) that the graph of ∂g is the union of finitely many closed convex sets, and hence ∂g is semismooth* on its graph by [32]. Since Ω is polyhedral, it follows from [88, Example 10.24] that g is *fully amenable* on $\mathbb{R}^n$ in the sense of [88]. Then using [88, Corollary 13.15] tells us that g is twice epi-differentiable on $\mathbb{R}^n$. Applying now Theorems 5 and 6, we verify both assertions of this theorem. $\square$

To run Algorithms 3 and 4, we need to explicitly determine the sequences $\{v^k\}$ and $\{d^k\}$ generated by these algorithms. The expressions for v^k follows directly from (6.24), which tells us that

$$v^k = (\min\{\max\{u_i, \ell_i\}, L_i\})_{i=1}^n.$$

Using further the formulas in (6.25)–(6.27), we express d^k in Algorithm 3 via the conditions

$$\begin{cases} (-\frac{1}{\gamma}(x^k - v^k) - Ad^k)_i = 0 & \text{if } (u^k - v^k)_i = 0, \\ (-x^k - v^k + d^k)_i = 0 & \text{if } (u^k - v^k)_i \neq 0. \end{cases}$$

Thus d^k can be computed for each $k \in \mathbb{N}$ by solving the linear equation $X^k d = v^k - x^k$, where

$$(X^k)_i := \begin{cases} \gamma A_i & \text{if } (u^k - v^k)_i = 0, \\ I_i & \text{if } (u^k - v^k)_i \neq 0. \end{cases} \quad (6.28)$$

Similarly, by employing the calculations of Proposition 9 in the framework of Algorithm 4 and by performing elementary transformations, we get the linear equation

$$(BX^k + \gamma\mu_k I)d^k = B(v^k - x^k)$$

to find the direction d^k , where X^k is computed in (6.28), and $B := I - \gamma A$. Now we are ready to conduct numerical experiments for solving quadratic programming problems with box constraints by using our GDNM via Algorithm 3 and GRNM via Algorithm 4. Our methods are compared with the *trust region reflective algorithm* in MATLAB's quadratic programming solver. All the numerical experiments are conducted in the same desktop and software described in Sect. 6.1.

To get the positive-semidefinite matrix A , we generate a random $n \times n$ matrix C with i.i.d. standard uniform entries and then define $A := C^*C$. In some particular tests, we put $A := C^*C/10^7$ so that A is close to be singular and mark these tests by *. Then the vectors b and l are generated randomly with i.i.d. standard uniform entries. To get the vector $L \in \mathbb{R}^n$ such that $l_i \leq L_i$ for all $i = 1, \dots, n$, entries of L are generated independently with uniform distribution on the interval $(1, 2)$. The initial points are all ones vector for all the tests and all the algorithms. As suggested in the MATLAB built in quadratic programming solver, the stopping criterion used is the function tolerance, i.e.,

$$\frac{|f(x^k) - f(x^{k+1})|}{1 + |f(x^k)|} \leq \varepsilon.$$

The tolerance ε is chosen to be 10^{-9} for all the tests and all the algorithms. The results of numerical experiments in this part are shown in Table 2. In this table, 'TR' refers to the *trust region reflective method* while other information is the same as in Sect. 6.2.

The obtained results show that our algorithms GDNM and GRNM are more efficient when the size of the problem is rather small; see, e.g., Tests 1, 2, 3 and 5, 6. It can also be seen that GRNM is more stable than GDNM when the size of the problems is increasing. Tests 4, 7 and 8 indicate that our methods should be further improved to solve problems in high-dimensional spaces.

Table 2 Solving box constrained quadratic programming on random instances

TN and size		Iter			CPU time		
TN	Size	TR	GDNM	GRNM	TR	GDNM	GRNM
1	200	6	4	6	0.16	0.07	0.02
2	500	7	7	6	0.08	0.05	0.04
3	2000	7	7	6	1.38	1.30	2.40
4	5000	9	8	6	8.03	9.06	18.33
5*	200	7	2	5	0.61	0.08	0.02
6*	500	9	3	5	0.11	0.03	0.04
7*	2000	10	10	8	1.61	2.00	2.97
8*	5000	15	39	6	14.58	54.89	21.54

7 Conclusions and future research

In this paper we propose and develop two globally convergent generalized Newton methods to solve problems of $\mathcal{C}^{1,1}$ optimization and of convex composite optimization with extended-real-valued regularizers, which include nonsmooth problems of constrained optimization. The developed algorithms are far-going extensions of the classical damped Newton method and of the regularized Newton algorithm with the replacement of the standard Hessian by its generalized version applied to nonsmooth (of the second order) functions. The latter construction is coderivative generated, which coins the names of our generalized Newton methods. The obtained results demonstrate the efficiency of both algorithms, their global superlinear convergence under appropriate assumptions, and their applications to the solution of Lasso problems and of box constrained problem of quadratic programming with conducting numerical experiments.

Our future research includes developing hybrid generalized Newton methods, which contain subproblems that can be efficiently solved by using first-order algorithms, and then combining them with the advanced second-order Newton-type techniques. We intend to establish the global superlinear convergence of iterates under relaxed assumptions that do not involve the positive-definiteness of the generalized Hessian in $\mathcal{C}^{1,1}$ optimization as well as the strong convexity requirement for problems of convex composite optimization, which will go beyond those with quadratic smooth parts. The obtained results would allow us to develop new applications to Lasso problems as well as to other important classes of models in machine learning, statistic, and related disciplines.

Acknowledgements The authors are very grateful to three anonymous referees for their helpful remarks and suggestions, which allowed us to significantly improve the original presentation.

Appendix: Some technical lemmas

This section contains four technical lemmas used in the text.

The first lemma is a local version of [44, Lemma 2.20]. The proof of this result is similar to the original one, and thus it is omitted.

Lemma 3 *Let $\Omega \subset \mathbb{R}^n$ be an open set, and let $\varphi : \Omega \rightarrow \mathbb{R}$ be a continuously differentiable function such that $\nabla\varphi$ is Lipschitz continuous with modulus $L > 0$. Then for any $\sigma > 0$, $x \in \Omega$, and $d \in \mathbb{R}^n$ satisfying $\langle \nabla\varphi(x), d \rangle < 0$, the following inequality*

$$\varphi(x + \tau d) \leq \varphi(x) + \sigma \tau \langle \nabla\varphi(x), d \rangle \quad (7.1)$$

holds whenever $\tau \in (0, \bar{\tau}]$ and $x + \tau d \in \Omega$, where

$$\bar{\tau} := \frac{2(\sigma - 1)\langle \nabla\varphi(x), d \rangle}{L\|d\|^2} > 0.$$

The next lemma provides conditions for the R -linear and Q -linear convergence of sequences.

Lemma 4 (estimates for convergence rates) *Let $\{\alpha_k\}$, $\{\beta_k\}$, and $\{\gamma_k\}$ be sequences of positive numbers. Assume that there exist numbers $c_i > 0$, $i = 1, 2, 3$, and $k_0 \in \mathbb{N}$ such that for all $k \geq k_0$ we have the estimates:*

- (i) $\alpha_k - \alpha_{k+1} \geq c_1 \beta_k^2$.
- (ii) $\beta_k \geq c_2 \gamma_k$.
- (iii) $c_3 \gamma_k^2 \geq \alpha_k$.

Then the sequence $\{\alpha_k\}$ Q -linearly converges to zero, and the sequences $\{\beta_k\}$ and $\{\gamma_k\}$ R -linearly converge to zero as $k \rightarrow \infty$.

Proof Combining (i), (ii), and (iii) yields the inequalities

$$\alpha_k - \alpha_{k+1} \geq c_1 \beta_k^2 \geq c_1 c_2^2 \gamma_k^2 \geq \frac{c_1 c_2^2}{c_3} \alpha_k \quad \text{for all } k \geq k_0,$$

which imply that $\alpha_{k+1} \leq q \alpha_k$, where $q := 1 - (c_1 c_2^2)/c_3 \in (0, 1)$. This verifies that the sequence $\{\alpha_k\}$ Q -linearly converges to zero. Using the latter and the assumed condition (i) ensures that

$$\beta_k^2 \leq \frac{1}{c_1} (\alpha_k - \alpha_{k+1}) \leq \frac{1}{c_1} \alpha_k \leq \frac{1}{c_1} q \alpha_{k-1} \leq \dots \leq \frac{1}{c_1} q^k \alpha_0 \quad \text{for all } k \geq k_0,$$

which tells us that $\beta_k \leq c \mu^k$, where $c := \sqrt{\alpha_0/c_1}$ and $\mu := \sqrt{q}$. This justifies the R -linear convergence of the sequence $\{\beta_k\}$ to zero. Furthermore, it easily follows from (ii) that the sequence $\{\gamma_k\}$ R -linearly converge to zero, and thus we are done with the proof. $\square$

Now we obtain a useful result on the semismooth* property of compositions.

Lemma 5 (semismooth* property of composition mappings) *Let $A \in \mathbb{R}^{n \times n}$ be a symmetric nonsingular matrix, $b \in \mathbb{R}^n$, $\bar{x} \in \mathbb{R}^n$, and $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be continuous and semismooth* at $\bar{y} := A\bar{x} + b$. Then the mapping $g : \mathbb{R}^n \rightarrow \mathbb{R}^n$ defined by $g(x) := f(Ax + b)$ is semismooth* at $\bar{x}$.*

Proof Using the coderivative chain rule from [59, Theorem 1.66], we get

$$D^*g(x)(w) = A^*D^*f(Ax + b)(w) \quad \text{for all } w \in \mathbb{R}^n. \quad (7.2)$$

Denote $\mu := \sqrt{\max\{1, \|A\|^2\} \cdot \max\{1, \|A^{-1}\|^2\}} > 0$. Picking any $\varepsilon > 0$ and employing the semismooth* property of f at $\bar{y}$, we find $\delta > 0$ such that

$$|\langle x^*, y - \bar{y} \rangle - \langle y^*, f(y) - f(\bar{y}) \rangle| \leq \frac{\varepsilon}{\mu} \|(y - \bar{y}, f(y) - f(\bar{y}))\| \cdot \|(x^*, y^*)\| \quad (7.3)$$

for all $y \in \mathbb{B}_\delta(\bar{y})$ and all $(x^*, y^*) \in \text{gph } D^*f(y)$. Denoting $r := \delta/\|A\| > 0$ gives us $y := Ax + b \in \mathbb{B}_\delta(\bar{y})$ whenever $x \in \mathbb{B}_r(\bar{x})$. Picking now $x \in \mathbb{B}_r(\bar{x})$ and $(z^*, w^*) \in \text{gph } D^*g(x)$, we get $(A^{-1}z^*, w^*) \in \text{gph } D^*f(Ax + b)$ due to (7.2). It follows from (7.3) that

$$\begin{aligned} |\langle z^*, x - \bar{x} \rangle - \langle w^*, g(x) - g(\bar{x}) \rangle| &= |\langle A^{-1}z^*, y - \bar{y} \rangle - \langle w^*, f(y) - f(\bar{y}) \rangle| \\ &\leq \frac{\varepsilon}{\mu} \|(y - \bar{y}, f(y) - f(\bar{y}))\| \cdot \|(A^{-1}z^*, w^*)\| \\ &= \frac{\varepsilon}{\mu} \|(Ax - A\bar{x}, g(x) - g(\bar{x}))\| \cdot \|(A^{-1}z^*, w^*)\| \\ &\leq \varepsilon \|(x - \bar{x}, g(x) - g(\bar{x}))\| \cdot \|(z^*, w^*)\|, \end{aligned}$$

which verifies the semismooth* property of g at $\bar{x}$. $\square$

The final lemma establishes tilt stability of strongly convex functions at stationary points.

Lemma 6 (strong convexity and tilt-stability) *Let $\varphi : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ be an l.s.c. and strongly convex function with modulus $\kappa > 0$, and let $\bar{x} \in \text{dom } \varphi$ such that $0 \in \partial\varphi(\bar{x})$. Then $\bar{x}$ is a tilt-stable local minimizer of φ with modulus κ^{-1} .*

Proof By the second-order characterization of strongly convex functions [10, Theorem 5.1], we have

$$\langle z, w \rangle \geq \kappa \|w\|^2 \quad \text{for all } z \in \partial^2\varphi(x, y)(w), \ (x, y) \in \text{gph } \partial\varphi, \text{ and } w \in \mathbb{R}^n.$$

This implies in turn that $\bar{x}$ is a tilt-stable local minimizer with modulus κ^{-1} by second-order characterization of tilt stability taken from [61, Theorem 3.5]. $\square$

References

1. Bauschke, H.H., Combettes, P.L.: *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*, 2nd edn. Springer, New York (2017)
2. Beck, A.: *Introduction to Nonlinear Optimization: Theory, Algorithms, and Applications with MATLAB*. SIAM, Philadelphia, PA (2014)
3. Beck, A.: *First-Order Methods in Optimization*. SIAM, Philadelphia, PA (2017)
4. Beck, A., Teboulle, M.: A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM J. Imaging Sci.* **2**, 183–202 (2009)
5. Becker, S., Fadili, M.J.: A quasi-Newton proximal splitting method. *Adv. Neural Inform. Process. Syst.* **25**, 2618–2626 (2012)
6. Bertsekas, D.P.: *Nonlinear Programming*, 3rd edn. Athena Scientific, Belmont, MA (2016)
7. Bonnans, J.F.: Local analysis of Newton-type methods for variational inequalities and nonlinear programming. *Appl. Math. Optim.* **29**, 161–186 (1994)
8. Boyd, S., Parikh, N., Chu, E., Peleato, B., Eckstein, J.: Distributed optimization and statistical learning via the alternating direction method of multipliers. *Found. Trends Mach. Learning* **3**, 1–122 (2010)
9. Boyd, S., Vandenberghe, L.: *Convex Optimization*. Cambridge University Press, Cambridge (2004)
10. Chieu, N.H., Chuong, T.D., Yao, J.-C., Yen, N.D.: Characterizing convexity of a function by its Fréchet and limiting second-order subdifferentials. *Set-Valued Var. Anal.* **19**, 75–96 (2011)
11. Chieu, N.H., Lee, G.M., Yen, N.D.: Second-order subdifferentials and optimality conditions for C^1 -smooth optimization problems. *Appl. Anal. Optim.* **1**, 461–476 (2017)
12. Chieu, N.M., Hien, L.V., Nghia, T.T.A.: Characterization of tilt stability via subgradient graphical derivative with applications to nonlinear programming. *SIAM J. Optim.* **28**, 2246–2273 (2018)
13. Colombo, G., Henrion, R., Hoang, N.D., Mordukhovich, B.S.: Optimal control of the sweeping process over polyhedral controlled sets. *J. Diff. Eqs.* **260**, 3397–3447 (2016)
14. Combettes, P.L., Pesquet, J.-C.: Proximal splitting methods in signal processing. In: Bauschke, H.H., et al. (eds.) *Fixed-Point Algorithms for Inverse Problems in Science and Engineering*, pp. 185–212. Springer, New York (2011)
15. Dan, H., Yamashita, N., Fukushima, M.: Convergence properties of the inexact Levenberg–Marquardt method under local error bound conditions. *Optim. Meth. Softw.* **17**, 605–626 (2002)
16. Dennis, J.E., Moré, J.J.: Quasi-Newton methods, motivation and theory. *SIAM Rev.* **19**, 46–89 (1977)
17. Ding, C., Sun, D., Ye, J.J.: First-order optimality conditions for mathematical programs with semidefinite cone complementarity constraints. *Math. Program.* **147**, 539–379 (2014)
18. Dias, S., Smirnov, G.: On the Newton method for set-valued maps. *Nonlinear Anal. TMA* **75**, 1219–1230 (2012)
19. Dontchev, A.L., Rockafellar, R.T.: Characterizations of strong regularity for variational inequalities over polyhedral convex sets. *SIAM J. Optim.* **6**, 1087–1105 (1996)
20. Dontchev, A.L., Rockafellar, R.T.: *Implicit Functions and Solution Mappings: A View from Variational Analysis*, 2nd edn. Springer, New York (2014)
21. Drusvyatskiy, D., Lewis, A.S.: Tilt stability, uniform quadratic growth, and strong metric regularity of the subdifferential. *SIAM J. Optim.* **23**, 256–267 (2013)
22. Drusvyatskiy, D., Mordukhovich, B.S., Nghia, T.T.A.: Second-order growth, tilt stability, and metric regularity of the subdifferential. *J. Convex Anal.* **21**, 1165–1192 (2014)
23. Emich, K., Henrion, R.: A simple formula for the second-order subdifferential of maximum functions. *Vietnam J. Math.* **42**, 467–478 (2014)
24. Efron, B., Hastie, T., Johnstone, I., Tibshirani, R.: Least angle regression. *Ann. Statist.* **32**, 407–499 (2004)
25. Facchinei, F.: Minimization of SC1 functions and the Maratos effect. *Oper. Res. Lett.* **17**, 131–137 (1995)
26. Facchinei, F., Pang, J.-C.: *Finite-Dimensional Variational Inequalities and Complementarity Problems*, vol. II. Springer, New York (2003)
27. Friedlander, M.P., Goodwin, A., Hoheisel, T.: From perspective maps to epigraphical projections. *Math. Oper. Res.* (2022). <https://doi.org/10.1287/moor.2022.1317>
28. Gabay, D., Mercier, B.: A dual algorithm for the solution of nonlinear variational problems via finite element approximations. *Comput. Math. Appl.* **2**, 17–40 (1976)

29. Glowinski, R., Marroco, A.: Sur l'approximation, par éléments finis d'ordre un, et la résolution, par pénalisation-dualité, d'une classe de problèmes de Dirichlet non linéaires. *Revue Francaise d'Automatique, Informatique et Recherche Operationelle* **9**, 41–76 (1975)
30. Gfrerer, H.: On directional metric regularity, subregularity and optimality conditions for nonsmooth mathematical programs. *Set-Valued Var. Anal.* **21**, 151–176 (2013)
31. Gfrerer, H., Mordukhovich, B.S.: Complete characterization of tilt stability in nonlinear programming under weakest qualification conditions. *SIAM J. Optim.* **25**, 2081–2119 (2015)
32. Gfrerer, H., Outrata, J.V.: On a semismooth* Newton method for solving generalized equations. *SIAM J. Optim.* **31**, 489–517 (2021)
33. Ginchev, I., Mordukhovich, B.S.: On directionally dependent subdifferentials. *C. R. Acad. Bulg. Sci.* **64**, 497–508 (2011)
34. Hang, N.T.V., Mordukhovich, B.S., Sarabi, M.E.: Augmented Lagrangian method for second-order conic programs under second-order sufficiency. *J. Glob. Optim.* **82**, 51–81 (2022)
35. Henrion, R., Mordukhovich, B.S., Nam, N.M.: Second-order analysis of polyhedral systems in finite and infinite dimensions with applications to robust stability of variational inequalities. *SIAM J. Optim.* **20**, 2199–2227 (2010)
36. Henrion, R., Römisich, W.: On M -stationary points for a stochastic equilibrium problem under equilibrium constraints in electricity spot market modeling. *Appl. Math.* **52**, 473–494 (2007)
37. Henrion, R., Outrata, J., Surowiec, T.: On the co-derivative of normal cone mappings to inequality systems. *Nonlinear Anal.* **71**, 1213–1226 (2009)
38. Hestenes, M.R.: Multiplier and gradient methods. *J. Optim. Theory Appl.* **4**, 303–320 (1969)
39. Hintermüller, M., Ito, K., Kunisch, K.: The primal-dual active set strategy as a semismooth Newton method. *SIAM J. Optim.* **13**, 865–888 (2002)
40. Hiriart-Urruty, J.-B., Strodio, J.-J., Nguyen, V.H.: Generalized Hessian matrix and second-order optimality conditions for problems with $C^{1,1}$ data. *Appl. Math. Optim.* **11**, 43–56 (1984)
41. Ho, C.H., Lin, C.J.: Large-scale linear support vector regression. *J. Mach. Learn. Res.* **13**, 3323–3348 (2012)
42. Hoheisel, T., Kanzow, C., Mordukhovich, B.S., Phan, H.M.: Generalized Newton's methods for non-smooth equations based on graphical derivatives, *Nonlinear Anal.* **75**, 1324–1340 (2012); Erratum in *Nonlinear Anal.* **86**, 157–158 (2013)
43. Hsieh, C.J., Chang, K.W., Lin, C.J.: A dual coordinate descent method for large-scale linear SVM. *Proceedings 25th International Conference on Machine Learning*, pp. 408–415. Helsinki, Finland (2008)
44. Izmailov, A.F., Solodov, M.V.: *Newton-Type Methods for Optimization and Variational Problems*. Springer, New York (2014)
45. Izmailov, A.F., Solodov, M.V., Uskov, E.T.: Globalizing stabilized sequential quadratic programming method by smooth primal-dual exact penalty function. *J. Optim. Theor. Appl.* **169**, 1–31 (2016)
46. Khanh, P.D., Mordukhovich, B.S., Phat, V.T.: A generalized Newton method for subgradient systems. *Math. Oper. Res.* (2022). <https://doi.org/10.1287/moor.2022.1320>
47. Klatte, D., Kummer, B.: *Nonsmooth Equations in Optimization. Regularity, Calculus, Methods and Applications*. Kluwer Academic Publishers, Dordrecht (2002)
48. Kummer, B.: Newton's method for non-differentiable functions. In: Guddat et al. (eds) *Advances in Mathematical Optimization*, pp. 114–124. Akademie-Verlag, Berlin (1988)
49. Lee, J.D., Sun, Y., Saunders, M.A.: Proximal Newton-type methods for minimizing composite functions. *SIAM J. Optim.* **24**, 1420–1443 (2014)
50. Li, D.H., Fukushima, M., Qi, L., Yamashita, N.: Regularized Newton methods for convex minimization problems with singular solutions. *Comput. Optim. Appl.* **28**, 131–147 (2004)
51. Li, X., Sun, D., Toh, K.-C.: A highly efficient semismooth Newton augmented Lagrangian method for solving Lasso problems. *SIAM J. Optim.* **28**, 433–458 (2018)
52. Lions, P.-L., Mercier, B.: Splitting algorithms for the sum of two nonlinear operators. *SIAM J. Numer. Anal.* **16**, 964–979 (1979)
53. Meng, F., Sun, D., Zhao, Z.: Semismoothness of solutions to generalized equations and the Moreau–Yosida regularization. *Math. Program.* **104**, 561–581 (2005)
54. Mohammadi, A., Mordukhovich, B.S., Sarabi, M.E.: Variational analysis of composite models with applications to continuous optimization. *Math. Oper. Res.* **47**, 397–426 (2022)
55. Mohammadi, A., Mordukhovich, B.S., Sarabi, M.E.: Parabolic regularity in geometric variational analysis. *Trans. Amer. Math. Soc.* **374**, 1711–1763 (2021)

56. Mohammadi, A., Sarabi, M.E.: Twice epi-differentiability of extended-real-valued functions with applications in composite optimization. *SIAM J. Optim.* **30**, 2379–2409 (2020)
57. Mordukhovich, B.S.: Sensitivity analysis in nonsmooth optimization. In: Field, D.A., Komkov, V. (eds) *Theoretical Aspects of Industrial Design*, pp. 32–46. SIAM Proc. Appl. Math. 58. Philadelphia, PA (1992)
58. Mordukhovich, B.S.: Complete characterizations of openness, metric regularity, and Lipschitzian properties of multifunctions. *Trans. Am. Math. Soc.* **340**, 1–35 (1993)
59. Mordukhovich, B.S.: *Variational Analysis and Generalized Differentiation, I: Basic Theory, II: Applications*. Springer, Berlin (2006)
60. Mordukhovich, B.S.: *Variational Analysis and Applications*. Springer, Cham (2018)
61. Mordukhovich, B.S., Nghia, T.T.A.: Second-order characterizations of tilt stability with applications to nonlinear programming. *Math. Program.* **149**, 83–104 (2015)
62. Mordukhovich, B.S., Nghia, T.T.A.: Local monotonicity and full stability of parametric variational systems. *SIAM J. Optim.* **26**, 1032–1059 (2016)
63. Mordukhovich, B.S., Outrata, J.V.: On second-order subdifferentials and their applications. *SIAM J. Optim.* **12**, 139–169 (2001)
64. Mordukhovich, B.S., Outrata, J.V., Sarabi, M.E.: Full stability of local optimal solutions in second-order cone programming. *SIAM J. Optim.* **14**, 1581–1613 (2014)
65. Mordukhovich, B.S., Rockafellar, R.T.: Second-order subdifferential calculus with applications to tilt stability in optimization. *SIAM J. Optim.* **22**, 953–986 (2012)
66. Mordukhovich, B.S., Rockafellar, R.T., Sarabi, M.E.: Characterizations of full stability in constrained optimization. *SIAM J. Optim.* **23**, 1810–1849 (2013)
67. Mordukhovich, B.S., Sarabi, M.E.: Generalized differentiation of piecewise linear functions in second-order variational analysis. *Nonlinear Anal.* **132**, 240–273 (2016)
68. Mordukhovich, B.S., Sarabi, M.E.: Generalized Newton algorithms for tilt-stable minimizers in nonsmooth optimization. *SIAM J. Optim.* **31**, 1184–1214 (2021)
69. Mordukhovich, B.S., Yuan, X., Zheng, S., Zhang, J.: A globally convergent proximal Newton-type method in nonsmooth convex optimization. *Math. Program.* **198**, 899–936 (2023)
70. Nam, N.M.: Coderivatives of normal cone mappings and Lipschitzian stability of parametric variational inequalities. *Nonlinear Anal.* **73**, 2271–2282 (2010)
71. Nesterov, Yu.: *Lectures on Convex Optimization*, 2nd edn. Springer, Cham (2018)
72. Nocedal, J., Wright, S.: *Numerical Optimization*. Springer, New York (2006)
73. Outrata, J.V., Sun, D.: On the coderivative of the projection operator onto the second-order cone. *Set-Valued Anal.* **16**, 999–1014 (2008)
74. Pang, J.S.: Newton's method for B-differentiable equations. *Math. Oper. Res.* **15**, 311–341 (1990)
75. Pang, J.S., Qi, L.: A globally convergent Newton method for convex SC1 minimization problems. *J. Optim. Theory Appl.* **85**, 633–648 (1995)
76. Patrinos, P., Bemporad, A.: Proximal Newton methods for convex composite optimization. In: *IEEE Conference on Decision and Control*, pp. 2358–2363 (2013)
77. Patrinos, P., Stella, L., Bemporad, A.: Forward-backward truncated Newton methods for convex composite optimization. [arXiv:1402.6655](https://arxiv.org/abs/1402.6655) (2014)
78. Pelckmans, K., De Brabanter, J., De Moor, B., Suykens, J.A.K.: Convex clustering shrinkage. In: *PASCAL Workshop on Statistics and Optimization of Clustering*, pp. 1–6. London (2005)
79. Poliquin, R.A., Rockafellar, R.T.: Tilt stability of a local minimum. *SIAM J. Optim.* **8**, 287–299 (1998)
80. Polyak, B.T.: *Introduction to Optimization*. Optimization Software, New York (1987)
81. Powell, M.J.D.: A method for nonlinear constraints in minimization problems. In: Fletcher, R. (ed.) *Optimization*, pp. 283–298. Academic Press, New York (1969)
82. Qi, L.: Convergence analysis of some algorithms for solving nonsmooth equations. *Math. Oper. Res.* **18**, 227–244 (1993)
83. Qi, L., Sun, J.: A nonsmooth version of Newton's method. *Math. Program.* **58**, 353–367 (1993)
84. Qui, N.T.: Generalized differentiation of a class of normal cone operators. *J. Optim. Theory Appl.* **161**, 398–429 (2014)
85. Robinson, S.M.: Newton's method for a class of nonsmooth functions. *Set-Valued Anal.* **2**, 291–305 (1994)
86. Rockafellar, R.T.: Augmented Lagrangian multiplier functions and duality in nonconvex programming. *SIAM J. Control* **12**, 268–285 (1974)

87. Rockafellar, R.T.: Augmented Lagrangians and hidden convexity in sufficient conditions for local optimality. *Math. Program.* **198**, 159–194 (2023)
88. Rockafellar, R.T., Wets, R.J.-B.: *Variational Analysis*. Springer, Berlin (1998)
89. She, Y.: Sparse regression with exact clustering. *Elect. J. Stat.* **4**, 1055–1096 (2010)
90. Stella, L., Themelis, A., Patrinos, P.: Forward-backward quasi-Newton methods for nonsmooth optimization problems. *Comput. Optim. Appl.* **67**, 443–487 (2017)
91. Stella, L., Themelis, A., Patrinos, P.: Forward-backward envelope for the sum of two nonconvex functions: further properties and nonmonotone linesearch algorithms. *SIAM J. Optim.* **28**, 2274–2303 (2018)
92. Tibshirani, R.: Regression shrinkage and selection via the Lasso. *J. R. Stat. Soc.* **58**, 267–288 (1996)
93. Ulbrich, M.: *Semismooth Newton Methods for Variational Inequalities and Constrained Optimization Problems in Function Spaces*. SIAM, Philadelphia, PA (2011)
94. Yamashita, N., Fukushima, M.: On the rate of convergence of the Levenberg–Marquardt method. In: Alefeld, G., Chen, X. (eds.) *Topics in Numerical Analysis*, vol. 15, pp. 239–249. Springer, Vienna (2001)
95. Yao, J.-C., Yen, N.D.: Coderivative calculation related to a parametric affine variational inequality. Part 1: Basic calculation. *Acta Math. Vietnam.* **34**, 157–172 (2009)

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

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.