

A penalty barrier framework for nonconvex constrained optimization

De Marchi, Alberto

Department of Aerospace Engineering, Institute of Applied Mathematics and Scientific Computing, University of the Bundeswehr Munich

Themelis, Andreas

Faculty of Information Science and Electrical Engineering (ISEE), Kyushu University

<https://hdl.handle.net/2324/7378076>

出版情報 : Journal of Nonsmooth Analysis and Optimization. 5, pp.14585-, 2025-08-19. Centre pour la Communication Scientifique Directe (CCSD)

バージョン :

権利関係 : © the authors



A PENALTY BARRIER FRAMEWORK FOR NONCONVEX CONSTRAINED OPTIMIZATION

Alberto De Marchi*

Andreas Themelis†

Abstract We consider minimization problems with structured objective function and smooth constraints, and present a flexible framework that combines the beneficial regularization effects of (exact) penalty and interior-point methods. In the fully nonconvex setting, a pure barrier approach requires careful steps when approaching the infeasible set, thus hindering convergence. We show how a tight integration with a penalty scheme mitigates this issue and enables the construction of subproblems whose domain is independent of the explicit constraints. This decoupling allows us to leverage efficient solvers designed for unconstrained or suitably structured optimization tasks. The key behind all this is a marginalization step: closely related to a conjugacy operation, this step effectively merges (exact) penalty and barrier into a smooth, full domain functional object. When the penalty exactness takes effect, the generated subproblems do not suffer the ill-conditioning typical of barrier methods, nor do they exhibit the nonsmoothness of exact penalty terms. We provide a theoretical characterization of the algorithm and its asymptotic properties, deriving convergence results for fully nonconvex problems. Stronger conclusions are available for the convex setting, where optimality can be guaranteed. Illustrative examples and numerical simulations demonstrate the wide range of problems our theory and algorithm are able to cover.

Keywords Nonsmooth nonconvex optimization, exact penalty methods, interior point methods, proximal algorithms

AMS subject classifications [49J52](#), [49J53](#), [65K05](#), [90Co6](#), [90C30](#)

1 INTRODUCTION

We are interested in developing numerical methods for constrained optimization problems of the form

$$(P) \quad \underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \quad q(\mathbf{x}) \quad \text{subject to} \quad \mathbf{c}(\mathbf{x}) \leq \mathbf{0}, \quad \mathbf{c}_{\text{eq}}(\mathbf{x}) = \mathbf{0},$$

where $\mathbf{x}$ is the decision variable and q , $\mathbf{c}$ and $\mathbf{c}_{\text{eq}}$ are problem functions. (Throughout, we stick to the convention of bold-facing vector variables and vector-valued functions, so that $\mathbf{0}$ indicates the zero vector of suitable size and similarly $\mathbf{1}$ is the vector with all entries equal to one.) Henceforth we consider (P) under the following standing assumptions.

Work supported by the Japan Society for the Promotion of Science (JSPS) KAKENHI grants JP21K17710 and JP24K20737. The authors also gratefully acknowledge the constructive feedback provided by the anonymous reviewers, whose careful reading and thoughtful remarks have been instrumental in refining this manuscript.

*Department of Aerospace Engineering, Institute of Applied Mathematics and Scientific Computing, University of the Bundeswehr Munich, Werner-Heisenberg-Weg 39, 85577 Neubiberg, Germany. alberto.demarchi@unibw.de, ORCID: [0000-0002-3545-6898](https://orcid.org/0000-0002-3545-6898)

†Faculty of Information Science and Electrical Engineering (ISEE), Kyushu University, 744 Motooka, Nishi-ku 819-0395, Fukuoka, Japan. andreas.themelis@ees.kyushu-u.ac.jp, ORCID: [0000-0002-6044-0169](https://orcid.org/0000-0002-6044-0169)

Assumption 1. The following hold in problem (P):

- A1 $q : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ is proper and lower semicontinuous (lsc).
- A2 $c : \mathbb{R}^n \rightarrow \mathbb{R}^m$ and $c_{\text{eq}} : \mathbb{R}^n \rightarrow \mathbb{R}^{m_{\text{eq}}}$ are continuously differentiable.
- A3 The problem is well posed: $q_{\star} := \inf_{\{x | c(x) \leq 0, c_{\text{eq}}(x) = 0\}} q(x)$ is finite.

Notice that no differentiability requirements are imposed on the cost q , nor convexity on any term in the formulation (hence the connotation of *fully* nonconvex). This modeling flexibility allows one to include simple constraints directly in q , forcing all generated iterates to honor them, as an alternative to explicit constraints in the format $c(x) \leq 0$, $c_{\text{eq}}(x) = 0$, which may be violated along the iterates. We remark that our framework allows for (and is robust to) equality constraints $c_{\text{eq}}(x) = 0$ to be expressed as two-sided inequalities $c_{\text{eq}}(x) \leq 0$ and $-c_{\text{eq}}(x) \leq 0$, despite the lack of constraint qualifications and in contrast to purely interior-point schemes, though a dedicated treatment of equalities yields an advantage in terms of algorithmic performance.

The primary objective of this paper is to devise an abstract algorithmic framework in the generality of this setting. The methodology requires an oracle for solving, up to approximate local optimality, minimization instances of the sum of q with a differentiable term. In practice, some structure is required in the cost function q to efficiently address these subproblems. The general setting of (P) under [Assumption 1](#) is considered without significantly weakening the convergence guarantees with respect to, say, assuming q smooth. Nevertheless, some stronger results are established under additional assumptions, such as locally Lipschitz continuity of the cost q or convexity of (P). In our numerical experiments we will invoke off-the-shelf routines based on proximal gradient iterations, thereby restricting our attention to problem instances in which q is structured as $q = f + g$ for a differentiable function f and a function g that enjoys an easily computable proximal map. Most nonsmooth functions widely used in practice comply with all these requirements. For instance, g can include indicators of any nonempty and closed set, and thus enforce arbitrary closed constraints that are easy to project onto. This modeling flexibility extends beyond the handling of constraints. While nonsmooth functions commonly encountered in practice can often be reformulated into smooth equivalents using slack variables and additional constraints, such reformulations typically come at the cost of introducing auxiliary variables and complicating the problem structure, penalizing algorithmic efficiency. By allowing nonsmooth terms to appear directly in the objective, our framework eliminates the need for such artificial constructs.

A particularly illustrative example is the so-called L^0 -(pseudo)norm penalty (number of nonzero entries) $\|x\|_0$ for $x \in \mathbb{R}^n$. As shown in [3, Lem. 3.1], this function can be represented as the *linear program*

$$(1.1) \quad \|x\|_0 = \min_{u \in \mathbb{R}^n} \|u\|_1 \quad \text{subject to} \quad -1 \leq u \leq 1, \langle u, x \rangle = \|x\|_1,$$

(more generally, matrix rank can also be cast in a similar fashion). In turn, nonsmoothness of each of the L^1 -norms can be resolved via the introduction of n slack variables and $2n$ inequality constraints, leading to a substantial increase in both problem size and constraint count. In contrast, our approach accommodates nonsmooth terms such as the L^0 -norm directly in the objective, avoiding any such inflation. The computational advantages of this modeling flexibility are also evident from the numerical experiments presented in [Section 5.3](#).

Motivations and related work The class of problems (P) with structured cost q has been recently studied in [5] and [13], respectively, for the fully convex and nonconvex setting, developing methods that bear strong convergence guarantees under some restrictive assumptions. Above all, building on a

pure barrier approach, these methods demand a feasible set with nonempty interior, thus excluding problems with equality constraints. Although restricted to simple bounds, a similar interior-point technique is investigated in [18] and manifests analogous pros and cons. In contrast to these works, we intend to address equality constraints as well. An augmented Lagrangian scheme for constrained structured problems was developed in [11], which also allows the specification of constraints in a function-in-set format.

Constrained structured programs (P) are also closely related to the template of structured *composite* optimization

$$\underset{x \in \mathbb{R}^n}{\text{minimize}} \quad q(x) + h(c(x))$$

with $h : \mathbb{R}^m \rightarrow \overline{\mathbb{R}}$. By introducing additional variables, composite problems can be rewritten in (equality) constrained form recovering the class of problems (P), with a one-to-one relationship between (local and global) solutions and stationary points [11, Lem. 3.1]. The recent literature on structured composite optimization includes [24], only for convex h , and [16, 10] for fully nonconvex problems, and concentrates almost exclusively on the augmented Lagrangian framework. Relying essentially on a penalty approach, in contrast to a barrier, the algorithmic characterization in [11] involved weaker assumptions and yet retrieved standard convergence results in constrained nonconvex optimization. However, the dependency on dual estimates makes methods of this family sensitive to the initialization of Lagrange multipliers. Moreover, they require some safeguards to ensure convergence from arbitrary starting points [4, 6]. In contrast, thanks to their ‘primal’ nature and inherent regularizing effect, penalty-barrier techniques can conveniently cope with degenerate problems.

The idea of adopting and merging penalty and barrier approaches, in a variety of possible flavors and combinations, is certainly not new, tracing back at least to [15]. Among several recent concretizations of this avenue, we refer to Curtis’ work [8] for a comprehensive discussion and further references. Our motivation for developing this technique for constrained structured problems stems from insights gained during the design of the interior point scheme IPprox [13]. The key observation therein is that, with a pure barrier approach, the arising subproblems have a smooth term *without* full domain. This nonstandard situation, together with a nonconvex and possibly extended-real-valued cost q and nonlinear constraints $c(x) \leq 0$, significantly restricts the range of subsolvers that can be employed. As a result, one cannot fully exploit more efficient optimization routines that would otherwise be suitable in an unconstrained or more structured setting.

In the broad setting of (P) under [Assumption 1](#), a blind application of penalty-barrier strategies in the spirit of [8] would bear no advantages, since the inconvenience in IPprox of a restricted domain would persist, hindering again the practical performance. In this paper we propose and investigate in detail a simple technique to overcome this limitation. The crucial step consists in the *marginalization* of auxiliary variables: after applying some penalty and barrier modifications, the auxiliary variables are optimized *pointwise*, for any given decision variable x .¹ Before proceeding with the technical content, we emphasize that the marginalization step not only reduces the subproblems’ size (recovering that of the original decision variable x only), but it also—and especially—results in a smooth penalty term for the subproblems that has always *full* domain. The emergence of this penalty-barrier envelope enables the adoption of generic (efficient) subsolvers, as well as tailored routines that exploit the problem’s original structure. This claim will be substantiated in [Section 3.2](#), where we show that properties such as convexity and Lipschitz differentiability, whenever present, are preserved in the transformed problems.

¹This approach can be interpreted as a drastic version of the so-called *magical steps* [7, 4], or *slack reset* in [8], and was inspired by the *proximal* approaches in [14, 10].

2 PRELIMINARIES

In this section we comment on useful notation and preliminary results before discussing optimality notions to characterize solutions of (P).

2.1 NOTATION AND KNOWN FACTS

With $\mathbb{R}$ and $\overline{\mathbb{R}} := \mathbb{R} \cup \{\pm\infty\}$ we denote the real and extended-real line, respectively, and with $\mathbb{R}_+ := [0, \infty)$ and $\mathbb{R}_- := (-\infty, 0]$ the set of nonnegative and nonpositive real numbers, respectively. The positive and negative parts of a number $r \in \mathbb{R}$ are respectively denoted as $[r]_+ := \max\{0, r\}$ and $[r]_- := \max\{0, -r\}$, so that $r = [r]_+ - [r]_-$ and $|r| = [r]_+ + [r]_-$. We stick to the convention of bold-facing vector variables and vector-valued functions, and use $\mathbf{0}$ to denote the zero vector of suitable size and similarly $\mathbf{1}$ for the vector with all entries equal to one. When applying unary operators to a vector $\mathbf{r}$, such as $|\mathbf{r}|$ or $[\mathbf{r}]_+$, the operation is meant elementwise.

The notation $T : \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ indicates a set-valued mapping T that maps any $\mathbf{x} \in \mathbb{R}^n$ to a (possibly empty) subset $T(\mathbf{x})$ of $\mathbb{R}^m$. Its (effective) domain and graph are the sets $\text{dom } T := \{\mathbf{x} \in \mathbb{R}^n \mid T(\mathbf{x}) \neq \emptyset\}$ and $\text{gph } T := \{(\mathbf{x}, \mathbf{y}) \in \mathbb{R}^n \times \mathbb{R}^m \mid \mathbf{y} \in T(\mathbf{x})\}$. T is said to be *outer semicontinuous* (osc) if its graph is a closed subset of $\mathbb{R}^n \times \mathbb{R}^m$. Algebraic operations with or among set-valued mappings are meant in a componentwise sense; for instance, the sum of $T_1, T_2 : \mathbb{R}^n \rightrightarrows \mathbb{R}^m$ is defined as $(T_1 + T_2)(\mathbf{x}) := \{\mathbf{y}^1 + \mathbf{y}^2 \mid (\mathbf{y}^1, \mathbf{y}^2) \in T_1(\mathbf{x}) \times T_2(\mathbf{x})\}$ for all $\mathbf{x} \in \mathbb{R}^n$.

The distance from a nonempty set $E \subseteq \mathbb{R}^n$ $\text{dist}_E : \mathbb{R}^n \rightarrow [0, \infty)$ is $\text{dist}_E(\mathbf{x}) := \inf_{\mathbf{y} \in E} \|\mathbf{y} - \mathbf{x}\|$. With $\delta_E : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ we denote the indicator function of E , namely such that $\delta_E(\mathbf{x}) = 0$ if $\mathbf{x} \in E$ and ∞ otherwise. For an extended-real-valued function $h : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$, the (effective) domain, graph, and epigraph are given by $\text{dom } h := \{\mathbf{x} \in \mathbb{R}^n \mid h(\mathbf{x}) < \infty\}$, $\text{gph } h := \{(\mathbf{x}, h(\mathbf{x})) \mid \mathbf{x} \in \text{dom } h\}$, and $\text{epi } h := \{(\mathbf{x}, \alpha) \in \mathbb{R}^n \times \mathbb{R} \mid \alpha \geq h(\mathbf{x})\}$. We say that h is *proper* if $\text{dom } h \neq \emptyset$ and $h > -\infty$, and *lower semicontinuous* (lsc) if $h(\bar{\mathbf{x}}) \leq \liminf_{\mathbf{x} \rightarrow \bar{\mathbf{x}}} h(\mathbf{x})$ for all $\bar{\mathbf{x}} \in \mathbb{R}^n$ or, equivalently, if $\text{epi } h$ is a closed subset of $\mathbb{R}^{n+1}$. Following [25, Def. 8.3], we denote by $\hat{\partial}h : \mathbb{R}^n \rightrightarrows \mathbb{R}^n$ the *regular subdifferential* of h , where

$$\mathbf{v} \in \hat{\partial}h(\bar{\mathbf{x}}) \iff \liminf_{\substack{\mathbf{x} \rightarrow \bar{\mathbf{x}} \\ \mathbf{x} \neq \bar{\mathbf{x}}}} \frac{h(\mathbf{x}) - h(\bar{\mathbf{x}}) - \langle \mathbf{v}, \mathbf{x} - \bar{\mathbf{x}} \rangle}{\|\mathbf{x} - \bar{\mathbf{x}}\|} \geq 0.$$

The (limiting, or Mordukhovich) subdifferential of h is $\partial h : \mathbb{R}^n \rightrightarrows \mathbb{R}^n$, where $\bar{\mathbf{v}} \in \partial h(\bar{\mathbf{x}})$ if and only if $\bar{\mathbf{x}} \in \text{dom } h$ and there exists a sequence $(\mathbf{x}^k, \mathbf{v}^k)_{k \in \mathbb{N}}$ in $\text{gph } \hat{\partial}h$ such that $(\mathbf{x}^k, \mathbf{v}^k, h(\mathbf{x}^k)) \rightarrow (\bar{\mathbf{x}}, \bar{\mathbf{v}}, h(\bar{\mathbf{x}}))$. In particular, $\hat{\partial}h(\mathbf{x}) \subseteq \partial h(\mathbf{x})$ holds at any $\mathbf{x} \in \mathbb{R}^n$; moreover, $\mathbf{0} \in \hat{\partial}h(\mathbf{x})$ is a necessary condition for local minimality of h at $\mathbf{x}$ [25, Thm. 10.1]. The subdifferential of h at $\bar{\mathbf{x}}$ satisfies $\partial(h + h_0)(\bar{\mathbf{x}}) = \partial h(\bar{\mathbf{x}}) + \nabla h_0(\bar{\mathbf{x}})$ for any $h_0 : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ continuously differentiable around $\bar{\mathbf{x}}$ [25, Ex. 8.8]. If h is convex, then $\hat{\partial}h = \partial h$ coincide with the *convex subdifferential*

$$\mathbb{R}^n \ni \bar{\mathbf{x}} \mapsto \{\mathbf{v} \in \mathbb{R}^n \mid h(\mathbf{x}) - h(\bar{\mathbf{x}}) - \langle \mathbf{v}, \mathbf{x} - \bar{\mathbf{x}} \rangle \geq 0 \ \forall \mathbf{x} \in \mathbb{R}^n\}.$$

For a convex set $C \subseteq \mathbb{R}^m$ and a point $\mathbf{x} \in C$ one has that $\partial \delta_C(\mathbf{x}) = N_C(\mathbf{x})$, where

$$N_C(\mathbf{x}) := \{\mathbf{v} \in \mathbb{R}^n \mid \langle \mathbf{v}, \mathbf{x}' - \mathbf{x} \rangle \leq 0 \ \forall \mathbf{x}' \in C\}$$

denotes the *normal cone* of C at $\mathbf{x}$, while $N_C(\mathbf{x}) = \emptyset$ for $\mathbf{x} \notin C$.

We use the symbol $JF : \mathbb{R}^n \rightarrow \mathbb{R}^{m \times n}$ to indicate the Jacobian of a differentiable mapping $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$, namely $JF(\bar{\mathbf{x}})_{i,j} = \frac{\partial F_i}{\partial x_j}(\bar{\mathbf{x}})$ for all $\bar{\mathbf{x}} \in \mathbb{R}^n$. For a real-valued function h , we instead use the gradient notation $\nabla h := Jh^\top$ to indicate the column vector of its partial derivatives. Finally, we remind that the *convex conjugate* of a proper lsc convex function $b : \mathbb{R} \rightarrow \overline{\mathbb{R}}$ is the proper lsc convex function $b^* : \mathbb{R} \rightarrow \overline{\mathbb{R}}$ defined as $b^*(\tau) := \sup_{t \in \mathbb{R}} \{\tau t - b(t)\}$, and that one then has $\tau \in \partial b(t)$ if and only if $t \in \partial b^*(\tau)$.

2.2 STATIONARITY CONCEPTS

This subsection summarizes well-known standard local optimality measures which were adopted in the proximal interior point framework of [13], and which will be further developed in the following Section 3 into conditions tailored to the setting of this paper. The interested reader is referred to [13, §2] for a verbose introduction and to [4, §3] for a detailed treatise. We start with the usual notion of (approximate) stationarity for general minimization problems of an extended-real-valued function.

Definition 2.1 (stationarity). Relative to the problem $\text{minimize}_{\mathbf{x} \in \mathbb{R}^n} \varphi(\mathbf{x})$ for a function $\varphi : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$, a point $\bar{\mathbf{x}} \in \mathbb{R}^n$ is called

- (i) *stationary* if it satisfies $\mathbf{0} \in \partial\varphi(\bar{\mathbf{x}})$;
- (ii) ε -*stationary* (with $\varepsilon > 0$) if it satisfies $\text{dist}_{\partial\varphi(\bar{\mathbf{x}})}(\mathbf{0}) \leq \varepsilon$.

A standard optimality notion that reflects the constrained structure of (P) is given by the Karush-Kuhn-Tucker (KKT) conditions.

Definition 2.2 (KKT optimality). Relative to problem (P), we say that $\bar{\mathbf{x}} \in \mathbb{R}^n$ is *KKT-optimal* if there exist $\bar{\mathbf{y}} \in \mathbb{R}^m$ and $\bar{\mathbf{y}}_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}$ such that

$$(KKT) \quad \begin{cases} -\mathbf{J}\mathbf{c}(\bar{\mathbf{x}})^\top \bar{\mathbf{y}} - \mathbf{J}\mathbf{c}_{\text{eq}}(\bar{\mathbf{x}})^\top \bar{\mathbf{y}}_{\text{eq}} \in \partial q(\bar{\mathbf{x}}) \\ \mathbf{c}(\bar{\mathbf{x}}) \leq \mathbf{0} \text{ and } \mathbf{c}_{\text{eq}}(\bar{\mathbf{x}}) = \mathbf{0} \\ \bar{\mathbf{y}} \geq \mathbf{0} \\ \bar{y}_i c_i(\bar{\mathbf{x}}) = 0 \quad i = 1, \dots, m. \end{cases}$$

In such case, we say that $(\bar{\mathbf{x}}, \bar{\mathbf{y}}, \bar{\mathbf{y}}_{\text{eq}}) \in \mathbb{R}^n \times \mathbb{R}^m \times \mathbb{R}^{m_{\text{eq}}}$ is a *KKT-optimal triplet* for (P).

Even for convex problems, unless suitable constraint and epigraphical qualifications are met, local minimizers may fail to be KKT-optimal. Necessary conditions in the generality of problem (P) are provided by the following asymptotic counterpart.²

Definition 2.3 (A-KKT optimality). Relative to problem (P), we say that $\bar{\mathbf{x}} \in \mathbb{R}^n$ is *asymptotically KKT-optimal* if $\bar{\mathbf{x}} \in \text{dom } q$ and there exist sequences $(\mathbf{x}^k)_{k \in \mathbb{N}} \rightarrow \bar{\mathbf{x}}$, $(\mathbf{y}^k)_{k \in \mathbb{N}} \subset \mathbb{R}^m$ and $(\mathbf{y}_{\text{eq}}^k)_{k \in \mathbb{N}} \subset \mathbb{R}^{m_{\text{eq}}}$ such that

$$(A-KKT) \quad \begin{cases} \text{dist}_{\partial q(\mathbf{x}^k)}(-\mathbf{J}\mathbf{c}(\mathbf{x}^k)^\top \mathbf{y}^k - \mathbf{J}\mathbf{c}_{\text{eq}}(\mathbf{x}^k)^\top \mathbf{y}_{\text{eq}}^k) \rightarrow 0 \\ [\mathbf{c}(\mathbf{x}^k)]_+ \rightarrow \mathbf{0} \text{ and } \mathbf{c}_{\text{eq}}(\mathbf{x}^k) \rightarrow \mathbf{0} \\ \mathbf{y}^k \geq \mathbf{0} \\ \mathbf{y}_i^k c_i(\bar{\mathbf{x}}) = 0 \quad i = 1, \dots, m. \end{cases}$$

The requirement $\bar{\mathbf{x}} \in \text{dom } q$, while superfluous in the original [4, Def. 3.1], is a necessary technicality to cope with possible nonclosedness of $\text{dom } q$ in the generality of Assumption 1. Taking the unconstrained minimization of $q(x) = \frac{1}{|x|} + \sin \frac{1}{x}$ as an example, this requirements prevents $\bar{x} = 0 \notin \text{dom } q$ to be considered A-KKT-optimal despite the fact that $\mathbf{x}^k = \frac{1}{(2k+1)\pi} \rightarrow \bar{\mathbf{x}}$ constitutes a valid sequence in the definition (having $\text{dist}_{\partial q(\mathbf{x}^k)}(\mathbf{0}) = 0$ for all k).

Proposition 2.4 ([4, Thm. 3.1], [11, Prop. 2.5]). *Any local minimizer for (P) is A-KKT-optimal.*

For the sake of designing suitable algorithmic stopping criteria, we also define an approximate variant which provides a further weaker notion of optimality.

²This definition is inspired by [4, Def. 3.1], where the ‘A’ in A-KKT is short for ‘approximate’. We however find ‘asymptotic’ more fit to emphasize its dependency on sequences, and reserve the ‘approximate’ label to characterize points satisfying KKT optimality up to some tolerance as in Definition 2.5.

Definition 2.5 (ϵ -KKT optimality). Relative to problem (P), for $\epsilon = (\epsilon_p, \epsilon_d) > (0, 0)$ we say that $\bar{x}$ is an (approximate) ϵ -KKT point if there exist $\bar{y} \in \mathbb{R}^m$ and $\bar{y}_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}$ such that

$$(\epsilon\text{-KKT}) \quad \begin{cases} \text{dist}_{\partial q(\bar{x})}(-\text{Jc}(\bar{x})^\top \bar{y} - \text{Jc}_{\text{eq}}(\bar{x})^\top \bar{y}_{\text{eq}}) \leq \epsilon_d \\ \|[c(\bar{x})]_+\|_\infty \leq \epsilon_p \text{ and } \|\text{c}_{\text{eq}}(\bar{x})\|_\infty \leq \epsilon_p \\ \bar{y} \geq 0 \\ \min\{\bar{y}_i, [c_i(\bar{x})]_-\} \leq \epsilon_p \quad i = 1, \dots, m. \end{cases}$$

It is handy to name points satisfying (approximate) feasibility as in Definition 2.5 and Definition 2.2 in order to soften symbolic clutter in the sequel.

Definition 2.6 (ϵ -feasibility). Given $\epsilon \geq 0$, a point $\bar{x} \in \mathbb{R}^n$ is said to be ϵ -feasible if $\|[c(\bar{x})]_+\|_\infty \leq \epsilon$ and $\|\text{c}_{\text{eq}}(\bar{x})\|_\infty \leq \epsilon$. When $\epsilon = 0$, we simply say that $\bar{x}$ is feasible.

As discussed in the commentary after [4, Thm. 3.1], A-KKT optimality of $\bar{x} \in \text{dom } q$ is tantamount to the existence of a sequence $x^k \rightarrow \bar{x}$ of ϵ^k -KKT points for some $\epsilon^k \rightarrow (0, 0)$. More generally, any KKT point is both A-KKT and ϵ -KKT for any $\epsilon \geq (0, 0)$. We conclude by listing the observations in [13, Lem. 8 and Rem. 9] that will be useful in the sequel.

Remark 2.7. Relative to the conditions A-KKT in Definition 2.3:

- (i) Up to possibly perturbing the sequence of multipliers, the complementarity slackness $y_i^k c_i(x^k) = 0$ can be equivalently expressed as $y_i^k c_i(x^k) \rightarrow 0$.
- (ii) Suppose that ∂q is osc on $\text{dom } q$ (as is the case when q is continuous relative to its domain). If the sequence $(y^k, y_{\text{eq}}^k)_{k \in \mathbb{N}}$ contains a bounded subsequence, then $\bar{x}$ is a KKT-optimal point, not merely asymptotically.

We note that [13] imposes a standing assumption that q be continuous at every point $\bar{x} \in \text{dom } q$. This is used to ensure that $\bar{v} \in \partial q(\bar{x})$ whenever $(x^k, v^k) \rightarrow (\bar{x}, \bar{v})$ with $(x^k, v^k) \in \text{gph } \partial q$, that is, that the condition $q(x^k) \rightarrow q(\bar{x})$ is superfluous in the definition of limiting subdifferential. In Remark 2.7(ii) we have relaxed this requirement by directly assuming this limiting property, that is, that ∂q is osc on $\text{dom } q$. Many functions of practical interest that are not continuous, such as the L^0 -norm, still comply with this requirement.

3 SUBPROBLEMS GENERATION

In this section we operate a two-step modification of problem (P), whose conceptual roadmap is as follows. We begin with a relaxed reformulation (P_α) in which violation of the constraints $c(x) \leq 0$ and $c_{\text{eq}}(x) = 0$ is penalized with an L^1 -norm in the cost function. An equivalent reformulation (Q_α) with slack variables $z \in \mathbb{R}^m$ and $z_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}$ simplifies this formulation by promoting separability. Next, a new problem $(Q_{\alpha, \mu})$ is created by adding a barrier term to enforce strict satisfaction of the inequality constraints in the L^1 -penalized reformulation (Q_α) . The pointwise minimization with respect to the slack variables z and z_{eq} can be carried out explicitly, with negligible computational overhead, resulting in a new problem $(P_{\alpha, \mu})$ in which the original constraints $c(x) \leq 0$ and $c_{\text{eq}}(x) = 0$ are softened with a smooth penalty. Increasing the L^1 -penalty and decreasing the barrier coefficients gives rise to a homotopic transition between smooth reformulations $(P_{\alpha, \mu})$ and the original nonsmooth problem (P).

Compared to envelope-type smoothings such as in [28] that preserve one-to-one correspondence of minimizers for any parameter, the smoothened subproblems here are equivalent to the original (P) only in the limit. In this sense, our method is more closely related to the approach in [26], but without the practical restrictions. Unlike the latter, which is tailored to problems where nonsmooth terms admit an easily computable ‘double envelope’, our technique applies more broadly with minimal limitations.

We adopt the convention of using the label ‘P’ in problems (P_α) and $(P_{\alpha,\mu})$ that share the minimization variable x with that in the original problem (P) . Label ‘Q’ is instead used in problems (Q_α) and $(Q_{\alpha,\mu})$ which come with additional slack variables z and z_{eq} . Each ‘P-problem’ amounts to the corresponding ‘Q-problem’ after marginal minimization with respect to the slack variables.

3.1 L^1 -PENALIZATION

Given $\alpha > 0$, we consider the following L^1 relaxation of (P) :

$$(P_\alpha) \quad \underset{x \in \mathbb{R}^n}{\text{minimize}} \quad q(x) + \alpha \| [c(x)]_+ \|_1 + \alpha \| c_{\text{eq}}(x) \|_1.$$

By introducing slack variables $z \in \mathbb{R}^m$ and $z_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}$, (P_α) can equivalently be cast as

$$(Q_\alpha) \quad \begin{aligned} & \underset{\substack{x \in \mathbb{R}^n, z \in \mathbb{R}^m \\ z_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}}}{\text{minimize}} \quad q(x) + \alpha \langle \mathbf{1}, z \rangle + \delta_{\mathbb{R}_+^m}(z) + \alpha \langle \mathbf{1}, z_{\text{eq}} \rangle \\ & \text{subject to} \quad c(x) \leq z \quad \text{and} \quad -z_{\text{eq}} \leq c_{\text{eq}}(x) \leq z_{\text{eq}}, \end{aligned}$$

as one can easily verify that

$$[c(x)]_+ = \arg \min_{z \in \mathbb{R}^m} \{ \alpha \langle \mathbf{1}, z \rangle + \delta_{\mathbb{R}_+^m}(z) \mid c(x) \leq z \}$$

and

$$|c_{\text{eq}}(x)| = \arg \min_{z_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}} \{ \alpha \langle \mathbf{1}, z_{\text{eq}} \rangle \mid -z_{\text{eq}} \leq c_{\text{eq}}(x) \leq z_{\text{eq}} \}$$

hold for any $x \in \mathbb{R}^n$ and $\alpha > 0$. In other words, (P_α) amounts to (Q_α) after a marginal minimization with respect to the slack variables z and z_{eq} . The KKT conditions associated to (Q_α) are of particular interest to us. As it can be deduced from the following lemma, they correspond to the stationarity condition $0 \in \partial q + \alpha \partial [\| [c(\cdot)]_+ \|_1] + \alpha \partial [\| [c_{\text{eq}}(\cdot)]_1]$ for (P_α) . The result is textbook, but its proof is nevertheless detailed in [Appendix A](#) for completeness.

Lemma 3.1. *Let [Assumption 1](#) hold. Then, a point $(x, z, z_{\text{eq}}) \in \mathbb{R}^n \times \mathbb{R}^m \times \mathbb{R}^{m_{\text{eq}}}$ is KKT-optimal for (Q_α) if and only if $z = [c(x)]_+$, $z_{\text{eq}} = |c_{\text{eq}}(x)|$, and there exist $y \in \mathbb{R}^m$ and $y_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}$ such that*

$$(KKT_\alpha) \quad \begin{cases} -Jc(x)^\top y - Jc_{\text{eq}}(x)^\top y_{\text{eq}} \in \partial q(x) \\ 0 \leq y \leq \alpha \mathbf{1} \\ |y_{\text{eq}}| \leq \alpha \mathbf{1} \\ y_i [c_i(x)]_- = 0 = (\alpha - y_i) [c_i(x)]_+, \quad i = 1, \dots, m \\ (\alpha - y_{\text{eq},j}) [c_{\text{eq},j}(x)]_+ = 0 = (\alpha + y_{\text{eq},j}) [c_{\text{eq},j}(x)]_-, \quad j = 1, \dots, m_{\text{eq}}. \end{cases}$$

Proof. See [Appendix A](#). □

Lemma 3.1 suggests the following relaxed optimality notion for problem (P) , which in light of the connection with KKT-optimality for (Q_α) we shall refer to as KKT_α -optimality.

Definition 3.2 (KKT_α optimality). Given $\alpha > 0$, we say that a point $\bar{x}^\alpha \in \mathbb{R}^n$ is KKT_α -optimal for (P) if there exist $\bar{y}^\alpha \in \mathbb{R}^m$ and $\bar{y}_{\text{eq}}^\alpha \in \mathbb{R}^{m_{\text{eq}}}$ such that $(\bar{x}^\alpha, \bar{y}^\alpha, \bar{y}_{\text{eq}}^\alpha)$ satisfy (KKT_α) , and call $(\bar{x}^\alpha, \bar{y}^\alpha, \bar{y}_{\text{eq}}^\alpha) \in \mathbb{R}^n \times \mathbb{R}^m \times \mathbb{R}^{m_{\text{eq}}}$ a KKT_α -optimal triplet for (P) .

Similarly to what done in $(\epsilon\text{-KKT})$ with respect to (KKT) , we may introduce an approximate KKT_α -optimality condition in which stationarity and complementarity slackness are satisfied up to some tolerance parameters. When said tolerance is zero, the nonapproximate KKT_α notion is recovered.

Definition 3.3 (ϵ -KKT $_{\alpha}$ optimality). Given $\alpha > 0$ and $\epsilon = (\epsilon_p, \epsilon_d) \geq (0, 0)$, we say that a point $\bar{x}^{\alpha} \in \mathbb{R}^n$ is ϵ -KKT $_{\alpha}$ -optimal for (P) if there exist $\bar{y}^{\alpha} \in \mathbb{R}^m$ and $\bar{y}_{\text{eq}}^{\alpha} \in \mathbb{R}^{m_{\text{eq}}}$ such that

$$(\epsilon\text{-KKT}_{\alpha}) \quad \begin{cases} \text{dist}_{\partial q(\bar{x}^{\alpha})}(-\text{Jc}(\bar{x}^{\alpha})^{\top} \bar{y}^{\alpha} - \text{Jc}_{\text{eq}}(\bar{x}^{\alpha})^{\top} \bar{y}_{\text{eq}}^{\alpha}) \leq \epsilon_d \\ \mathbf{0} \leq \bar{y}^{\alpha} \leq \alpha \mathbf{1} \\ -\alpha \mathbf{1} \leq \bar{y}_{\text{eq}}^{\alpha} \leq \alpha \mathbf{1} \\ s(\bar{x}^{\alpha}, \bar{y}^{\alpha}, \bar{y}_{\text{eq}}^{\alpha}) \leq \epsilon_p, \end{cases}$$

where

$$(3.1) \quad s(\bar{x}^{\alpha}, \bar{y}^{\alpha}, \bar{y}_{\text{eq}}^{\alpha}) := \left\| \min \left\{ \begin{pmatrix} \bar{y}^{\alpha} \\ \alpha \mathbf{1} - \bar{y}^{\alpha} \\ \alpha \mathbf{1} + \bar{y}_{\text{eq}}^{\alpha} \\ \alpha \mathbf{1} - \bar{y}_{\text{eq}}^{\alpha} \end{pmatrix}, \begin{pmatrix} [\text{c}(\bar{x}^{\alpha})]_{-} \\ [\text{c}(\bar{x}^{\alpha})]_{+} \\ [\text{c}_{\text{eq}}(\bar{x}^{\alpha})]_{-} \\ [\text{c}_{\text{eq}}(\bar{x}^{\alpha})]_{+} \end{pmatrix} \right\} \right\|_{\infty},$$

and we say that $(\bar{x}^{\alpha}, \bar{y}^{\alpha}, \bar{y}_{\text{eq}}^{\alpha}) \in \mathbb{R}^n \times \mathbb{R}^m \times \mathbb{R}^{m_{\text{eq}}}$ is an ϵ -KKT $_{\alpha}$ -optimal triplet for (P).

As a next step, we clarify how ϵ -KKT- and ϵ -KKT $_{\alpha}$ -optimality for problem (P) are interrelated.

Lemma 3.4. For any $\epsilon = (\epsilon_p, \epsilon_d) \geq (0, 0)$ the following hold:

- (i) An ϵ -KKT $_{\alpha}$ -optimal triplet $(\bar{x}^{\alpha}, \bar{y}^{\alpha}, \bar{y}_{\text{eq}}^{\alpha})$ with $\bar{x}^{\alpha}$ ϵ_p -feasible is also ϵ -KKT-optimal.
- (ii) An ϵ -KKT-optimal triplet $(\bar{x}, \bar{y}, \bar{y}_{\text{eq}})$ is also ϵ -KKT $_{\alpha}$ -optimal for any $\alpha \geq \max \{ \|\bar{y}\|_{\infty}, \|\bar{y}_{\text{eq}}\|_{\infty} \}$.

Once again the result is standard, and the proof is obvious by comparing the ϵ -KKT and ϵ -KKT $_{\alpha}$ optimality conditions, as schematically summarized below:

$$\begin{array}{l} \epsilon\text{-KKT} \left\{ \begin{array}{l} \text{dist}_{\partial q(\bar{x})}(-\text{Jc}(\bar{x})^{\top} \bar{y} - \text{Jc}_{\text{eq}}(\bar{x})^{\top} \bar{y}_{\text{eq}}) \leq \epsilon_d \\ \bar{y} \geq \mathbf{0} \\ \|\text{c}_{\text{eq}}(\bar{x})\|_{\infty} \leq \epsilon_p \\ \|[\text{c}(\bar{x})]_{+}\|_{\infty} \leq \epsilon_p \\ \min \{ \bar{y}_i, [c_i(\bar{x})]_{-} \} \leq \epsilon_p \end{array} \right. \end{array} \quad \begin{array}{l} \epsilon\text{-KKT}_{\alpha} \left\{ \begin{array}{l} \text{dist}_{\partial q(\bar{x})}(-\text{Jc}(\bar{x})^{\top} \bar{y} - \text{Jc}_{\text{eq}}(\bar{x})^{\top} \bar{y}_{\text{eq}}) \leq \epsilon_d \\ \mathbf{0} \leq \bar{y} \leq \alpha \mathbf{1}, -\alpha \mathbf{1} \leq \bar{y}_{\text{eq}} \leq \alpha \mathbf{1} \\ \min \{ \alpha - \bar{y}_{\text{eq},j} \text{sgn}(\text{c}_{\text{eq},j}(\bar{x})), |\text{c}_{\text{eq},j}(\bar{x})| \} \leq \epsilon_p \\ \min \{ \alpha - \bar{y}_i, [c_i(\bar{x})]_{+} \} \leq \epsilon_p \\ \min \{ \bar{y}_i, [c_i(\bar{x})]_{-} \} \leq \epsilon_p \end{array} \right. \end{array}$$

with $i = 1, \dots, m$ and $j = 1, \dots, m_{\text{eq}}$.

3.2 IP-TYPE BARRIER REFORMULATION

To carry on with the second modification of the problem, in what follows we fix a barrier b satisfying the following requirements.

Assumption 2. The barrier function $b : \mathbb{R} \rightarrow \overline{\mathbb{R}}$ is proper, lsc, and twice continuously differentiable on its domain $\text{dom } b = (-\infty, 0)$ with $b' > 0$ and $b'' > 0$.

For reasons that will be elaborated on later, convenient choices of barriers are $b(t) = -\frac{1}{t}$, $b(t) = \ln(1 - \frac{1}{t})$, and the classical logarithmic barrier $b(t) = -\ln(-t)$ (all extended as ∞ on $\mathbb{R}_{+}$), see Table 2 in Section 4.2. Once such b is fixed, in the spirit of interior point methods we enforce strict satisfaction

of the constraint in (Q_α) by considering the following barrier version

$$(Q_{\alpha,\mu}) \quad \begin{aligned} & \underset{\substack{\mathbf{x} \in \mathbb{R}^n, \mathbf{z} \in \mathbb{R}^m \\ \mathbf{z}_{\text{eq}} \in \mathbb{R}^{m_{\text{eq}}}}}{\text{minimize}} \quad q(\mathbf{x}) + \alpha \langle \mathbf{1}, \mathbf{z} \rangle + \delta_{\mathbb{R}_+^m}(\mathbf{z}) + \mu \sum_{i=1}^m b(c_i(\mathbf{x}) - z_i) \\ & \quad + \alpha \langle \mathbf{1}, \mathbf{z}_{\text{eq}} \rangle + \mu \sum_{j=1}^{m_{\text{eq}}} [b(c_{\text{eq},j}(\mathbf{x}) - z_{\text{eq},j}) + b(-c_{\text{eq},j}(\mathbf{x}) - z_{\text{eq},j})] \end{aligned}$$

for some given parameter $\mu > 0$. Differently from the IP frameworks of [5, 13], we here enforce a barrier in the relaxed version (Q_α) , and *not* on the original problem (P) . As such, it is only triplets $(\mathbf{x}, \mathbf{z}, \mathbf{z}_{\text{eq}})$ that need to lie in the interior of the constraints, but $\mathbf{x}$ is otherwise ‘unconstrained’: for any $\mathbf{x} \in \mathbb{R}^n$, any $\mathbf{z} > \mathbf{c}(\mathbf{x})$ and $\mathbf{z}_{\text{eq}} > |\mathbf{c}_{\text{eq}}(\mathbf{x})|$ (elementwise) yield a triplet $(\mathbf{x}, \mathbf{z}, \mathbf{z}_{\text{eq}})$ that satisfies the strict constraints $\mathbf{c}(\mathbf{x}) - \mathbf{z} < \mathbf{0}$ and $-\mathbf{z}_{\text{eq}} < \mathbf{c}_{\text{eq}}(\mathbf{x}) < \mathbf{z}_{\text{eq}}$. Furthermore, notice that the positivity constraint $\mathbf{z} \geq \mathbf{0}$ remains untouched, formally imposed by an indicator and not by the barrier. In fact, observing that the cost in $(Q_{\alpha,\mu})$ is separable, we may explicitly minimize with respect to the slack variables $\mathbf{z}$ and $\mathbf{z}_{\text{eq}}$. Plugging their optimal values into $(Q_{\alpha,\mu})$ results in an unconstrained reformulation of the form

$$(P_{\alpha,\mu}) \quad \underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \quad q(\mathbf{x}) + \mu \Psi_{\alpha/\mu}(\mathbf{c}(\mathbf{x})) + \mu \Psi_{\alpha/\mu}^{\text{eq}}(\mathbf{c}_{\text{eq}}(\mathbf{x})),$$

where, for any $\rho^* > 0$,³

$$(3.2) \quad \Psi_{\rho^*}(\mathbf{y}) := \sum_{i=1}^m \psi_{\rho^*}(y_i) \quad \text{and} \quad \Psi_{\rho^*}^{\text{eq}}(\mathbf{y}_{\text{eq}}) := \sum_{j=1}^{m_{\text{eq}}} \psi_{\rho^*}^{\text{eq}}(y_{\text{eq},j})$$

are separable functions with

$$(3.3) \quad \psi_{\rho^*}(t) := \min_{z \in \mathbb{R}_+} \{\rho^* z + b(t - z)\}$$

and

$$(3.4) \quad \psi_{\rho^*}^{\text{eq}}(t) := \min_{z \in \mathbb{R}} \{\rho^* z + b(t - z) + b(-t - z)\}.$$

These functions satisfy appealing properties summarized in the following theorems.

Theorem 3.5. *Suppose that [Assumption 2](#) holds. Then, for any $\rho^* > 0$ one has that*

$$(3.5) \quad \psi_{\rho^*}(t) = \begin{cases} b(t) & \text{if } b'(t) \leq \rho^* \\ \rho^* t - b^*(\rho^*) & \text{otherwise} \end{cases}$$

is convex, Lipschitz differentiable, and ρ^ -Lipschitz continuous with derivative*

$$(3.6) \quad \psi'_{\rho^*}(t) = \min \{b'(t), \rho^*\}.$$

Moreover, for any $c : \mathbb{R}^n \rightarrow \mathbb{R}$ convex, the composition $\psi_{\rho^} \circ c$ is also convex.*

Proof. See [Appendix A](#). □

³The choice of the starred symbol ρ^* stems from the fact that, as shown in [Theorem 3.5](#) (see also [Figure 2a](#)), this quantity represents a ‘slope’ of b , that is, a value of its derivative, and we thus treat it as a ‘dual’ object.

$b(t)$ (for $t < 0$)	$b^*(\tau)$ (for $\tau \geq 0$)	$z_{\rho^*}(t)$ (for $t \in \mathbb{R}$)
$-\frac{1}{t}$	$-2\sqrt{\tau}$	$\sqrt{t^2 + \frac{1}{\rho^*}} + \sqrt{\frac{4}{\rho^*}t^2 + \frac{1}{(\rho^*)^2}}$
$\ln(1 - \frac{1}{t})$	$-2\left(\frac{\sqrt{\tau}}{\sqrt{\tau} + \sqrt{\tau+4}} + \ln\left(\frac{\sqrt{\tau} + \sqrt{\tau+4}}{2}\right)\right)$	$\sqrt{t^2 + \frac{1}{4} + \frac{1}{\rho^*}} + \sqrt{t^2 + \frac{1}{(\rho^*)^2} + \frac{4t^2}{\rho^*}} - \frac{1}{2}$
$-\ln(-t)$	$-1 - \ln(\tau)$	$\frac{1}{\rho^*} + \sqrt{t^2 + \frac{1}{(\rho^*)^2}}$

Table 1: Examples of barriers with their conjugates and analytic expressions for $z_{\rho^*}(t)$, needed to compute the equality penalty $\psi_{\rho^*}^{\text{eq}}(t)$ and its derivative $(\psi_{\rho^*}^{\text{eq}})'(t)$ as in (3.7) and (3.8).

As is apparent from (3.5), ψ_{ρ^*} coincides with the barrier b up to when its slope is ρ^* , and after that point it reduces to its tangent line. As such, ψ_{ρ^*} coincides with a McShane Lipschitz (and globally Lipschitz differentiable) extension [20] of a portion of the barrier b , as depicted in Figure 1a and 2a. This feature is also evident by viewing ψ_{ρ^*} as the ρ^* -Pasch-Hausdorff envelope of b , as detailed in the proof.

Similar properties are true for $\psi_{\rho^*}^{\text{eq}}$, though a corresponding closed-form expression for generic barriers b is more cumbersome and not particularly helpful. An analytic expression is nevertheless available for specific choices of barriers b , see Table 1, or their value at any point can more generally be retrieved at negligible cost by solving a one-dimensional smooth monotone equation.

Theorem 3.6. Suppose that Assumption 2 holds. Then, for any $\rho^* > 0$ one has that

$$(3.7) \quad \psi_{\rho^*}^{\text{eq}}(t) = \rho^* z_{\rho^*}(t) + b(t - z_{\rho^*}(t)) + b(-t - z_{\rho^*}(t))$$

is convex, Lipschitz differentiable, and ρ^* -Lipschitz continuous with derivative

$$(3.8) \quad (\psi_{\rho^*}^{\text{eq}})'(t) = \rho^* - 2b'(-t - z_{\rho^*}(t)) \in (-\rho^*, \rho^*),$$

where, denoting $\rho := (b^*)'(\rho^*) < 0$, $z_{\rho^*}(t) > |t| - \rho$ is the unique solution $z \in \mathbb{R}$ to the smooth monotone equation

$$(3.9) \quad b'(t - z) + b'(-t - z) = \rho^*.$$

Moreover, for any $t \neq 0$ one has that $|(\psi_{\rho^*}^{\text{eq}})'(t)| \geq \rho^* - 2b'(-|t|)$.

Proof. See Appendix A. □

The specific barriers included in Table 1 are visualized for comparison in Figure 1, along with their corresponding envelopes ψ_{ρ^*} and $\psi_{\rho^*}^{\text{eq}}$. Notably, the log-like barrier offers an intermediate between the inverse and the logarithmic barriers, bringing together the positive valuedness of the former with the behavior of the latter near $t = 0$. On the one hand, positive valuedness guarantees via Theorems 3.5 and 3.6 that this property is inherited by ψ_{ρ^*} and $\psi_{\rho^*}^{\text{eq}}$, resulting in $(P_{\alpha, \mu})$ being more likely well posed. On the other hand, we will see below in Section 4.2 that the logarithmic barrier behavior is optimal near $t = 0$, in a certain sense. For these reasons, the log-like barrier function is a practical substitute for the classical logarithmic barrier and will be our default choice in the numerical validations, which support these claims.

The appeal of $\psi_{\alpha/\mu}$ and $\psi_{\alpha/\mu}^{\text{eq}}$ for the sake of addressing problem (P) lies in their behavior when μ is driven to 0, as they respectively approximate the inequality and equality sharp L^1 penalization $\alpha[\cdot]_+$

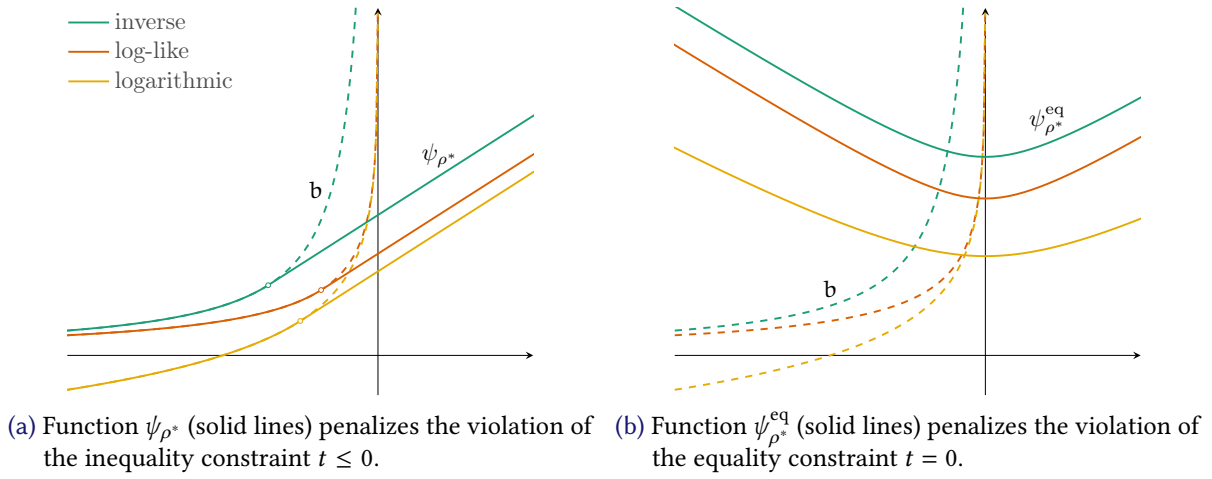


Figure 1: Graph of ψ_{ρ^*} (left) and $\psi_{\rho^*}^{\text{eq}}$ (right) for different barriers b (dashed lines). The log-like barrier behaves as an intermediate between the classical inverse and logarithmic barriers.

and $\alpha|\cdot|$. In this respect, **Figure 2b** demonstrates the advantage of our tailored treatment of equality constraints $\mathbf{c}_{\text{eq}}(\mathbf{x}) = \mathbf{0}$ (solid lines) as opposed to a naive use of double inequalities $\pm \mathbf{c}_{\text{eq}}(\mathbf{x}) \leq \mathbf{0}$ (dotted lines). Indeed, the latter approach results in the sum

$$(3.10) \quad t \mapsto \psi_{\alpha/\mu}^{\pm}(t) := \psi_{\alpha/\mu}(t) + \psi_{\alpha/\mu}(-t)$$

appearing in the formulation $(Q_{\alpha,\mu})$, and because of their opposite slopes around the origin a flat region appears that hinders algorithmic efficiency. In contrast, the combined marginalization in the definition (3.7) results in an envelope function $\psi_{\alpha/\mu}^{\text{eq}}$ that better approximates the sharp L^1 penalty, see also **Figure 3b**.

Theorem 3.7. Suppose that *Assumption 2* holds. Then, $\psi_{\rho^*}/\rho^* \rightarrow [\cdot]_+$ and $\psi_{\rho^*}^{\text{eq}}/\rho^* \rightarrow |\cdot|$ pointwise as $\rho^* \nearrow \infty$. Under the assumption that $b > 0$, both sequences are pointwise decreasing.

Proof. See **Appendix A**. □

Problem $(P_{\alpha,\mu})$ is ‘unconstrained’, in the sense that no explicit ambient constraints are provided, yet stationarity notions relative to it bear a close resemblance with KKT_{α} -optimality.

Lemma 3.8. Suppose that *Assumptions 1* and *2* hold. Then, for any $\alpha, \mu > 0$ and $\mathbf{x} \in \mathbb{R}^n$ one has

$$\begin{aligned} \partial \left[q + \mu \Psi_{\alpha/\mu} \circ \mathbf{c} + \mu \Psi_{\alpha/\mu}^{\text{eq}} \circ \mathbf{c}_{\text{eq}} \right] (\mathbf{x}) &= \partial q(\mathbf{x}) + \mu \sum_{i=1}^m \psi'_{\alpha/\mu}(c_i(\mathbf{x})) \nabla c_i(\mathbf{x}) \\ &\quad + \mu \sum_{j=1}^{m_{\text{eq}}} (\psi_{\alpha/\mu}^{\text{eq}})'(c_{\text{eq},j}(\mathbf{x})) \nabla c_{\text{eq},j}(\mathbf{x}). \end{aligned}$$

In particular, for any $\varepsilon \geq 0$ a point $\bar{\mathbf{x}}^{\alpha,\mu} \in \mathbb{R}^n$ is ε -stationary for $(P_{\alpha,\mu})$ if the pair $(\bar{\mathbf{y}}^{\alpha,\mu}, \bar{\mathbf{y}}_{\text{eq}}^{\alpha,\mu}) \in \mathbb{R}^m \times \mathbb{R}^{m_{\text{eq}}}$ given by

$$\bar{y}_i^{\alpha,\mu} := \mu \psi'_{\alpha/\mu}(c_i(\bar{\mathbf{x}}^{\alpha,\mu})) \in (0, \alpha] \quad \text{and} \quad \bar{y}_{\text{eq},j}^{\alpha,\mu} := \mu (\psi_{\alpha/\mu}^{\text{eq}})'(c_{\text{eq},j}(\bar{\mathbf{x}}^{\alpha,\mu})) \in (-\alpha, \alpha),$$

$i = 1, \dots, m, j = 1, \dots, m_{\text{eq}}$, satisfies

$$\text{dist}_{\partial q(\bar{\mathbf{x}}^{\alpha,\mu})} (-\mathbf{Jc}(\bar{\mathbf{x}}^{\alpha,\mu})^\top \bar{\mathbf{y}}^{\alpha,\mu} - \mathbf{Jc}_{\text{eq}}(\bar{\mathbf{x}}^{\alpha,\mu})^\top \bar{\mathbf{y}}_{\text{eq}}^{\alpha,\mu}) \leq \varepsilon.$$

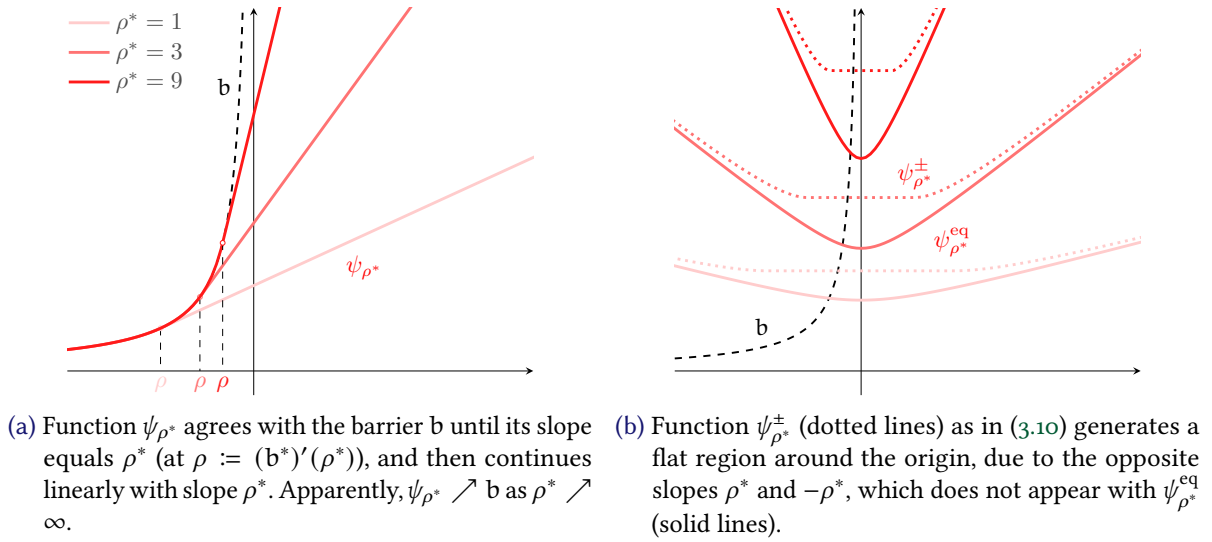


Figure 2: Graph of ψ_{ρ^*} (left) and $\psi_{\rho^*}^{\text{eq}}$ (right) for different values of ρ^* . These examples employ the inverse barrier $b(t) = -\frac{1}{t} + \delta_{(-\infty, 0)}$.

Proof. The expression of the subdifferential follows from the continuous differentiability of c and c_{eq} (Assumption 1), that of $\psi_{\alpha/\mu}$ and $\psi_{\alpha/\mu}^{\text{eq}}$ (Theorems 3.5 and 3.6), and their separable structure as in (3.2). The inclusion of each element of $\bar{y}^{\alpha, \mu}$ in the appropriate interval $(0, \alpha]$ follows from (3.6), and similarly the claimed bounds on the components of $\bar{y}_{\text{eq}}^{\alpha, \mu}$ follow from (3.8). $\square$

The information we gain from Lemma 3.8 is that whenever $\bar{x}$ is ε -stationary for $(P_{\alpha, \mu})$, then the triplet $(\bar{x}, \bar{y}, \bar{y}_{\text{eq}})$ with $\bar{y} = \mu \Psi'_{\alpha/\mu}(c(\bar{x}))$ and $\bar{y}_{\text{eq}} = \mu (\Psi^{\text{eq}}_{\alpha/\mu})'(c_{\text{eq}}(\bar{x}))$ satisfies all the conditions of $(\varepsilon_p, \varepsilon)$ -KKT $_{\alpha}$ optimality, possibly with the exception of the complementarity slackness involving $s(\bar{x}, \bar{y}, \bar{y}_{\text{eq}})$.

In conclusion of this section we emphasize that favorable features of the original problem (P) are likely to be preserved in the formulation (Q_{α}) . As expectable, and as explicitly mentioned in Theorems 3.5 and 3.6, convexity is one such property. More notably, under minimal additional assumptions, whenever c and c_{eq} are Lipschitz differentiable they remain so after the composition with $\psi_{\alpha/\mu}$ and $\psi_{\alpha/\mu}^{\text{eq}}$. This constitutes a significant departure from, say, augmented Lagrangian or penalty methods where such property is lost in the composition with quadratic functions. We exemplify this fact with the following lemma; we note that the statement holds under more general conditions, but a detailed exploration of these broader cases lies beyond the scope of this paper.

Lemma 3.9. Suppose that Assumption 2 holds, and let $c : \mathbb{R}^n \rightarrow \mathbb{R}$ be a Lipschitz-differentiable function. Suppose that c is Lipschitz continuous on the sublevel set $\{x \in \mathbb{R}^n \mid c(x) \leq 0\}$ (as is the case when it is lower bounded [17, Lem. 2.3]). Then, $\psi_{\rho^*} \circ c$ is Lipschitz differentiable for any $\rho^* > 0$.

Proof. Let L_c and ℓ denote the Lipschitz constant of ∇c (on $\mathbb{R}^n$) and that of c on $\{x \in \mathbb{R}^n \mid c(x) \leq 0\}$, respectively. According to Theorem 3.5, ψ_{ρ^*} is ρ^* -Lipschitz continuous, coincides with b on $(-\infty, \rho]$, and is then linear with slope ρ^* on (ρ, ∞) , where $\rho := (b^*)'(\rho^*) < 0$. Fix $x, y \in \mathbb{R}^n$, and without loss of generality assume that $c(x) \leq c(y)$. We have

$$\begin{aligned} \|\nabla(\psi_{\rho^*} \circ c)(x) - \nabla(\psi_{\rho^*} \circ c)(y)\| &= \|\psi'_{\rho^*}(c(x))\nabla c(x) - \psi'_{\rho^*}(c(y))\nabla c(y)\| \\ &\leq \psi'_{\rho^*}(c(x))\|\nabla c(x) - \nabla c(y)\| + \|\nabla c(y)\|\|\psi'_{\rho^*}(c(x)) - \psi'_{\rho^*}(c(y))\| \\ &\leq \rho^* L_c \|x - y\| + \|\nabla c(y)\|(\psi'_{\rho^*}(c(y)) - \psi'_{\rho^*}(c(x))). \end{aligned}$$

It remains to account for the second term in the last sum. If $c(x) \leq c(y) \leq \rho$, then ψ_{ρ^*} coincides with b in all occurrences, and the term can be upper bounded as $B\ell^2\|x - y\|$, where $B := \max_{(-\infty, \rho]} b''$ is a

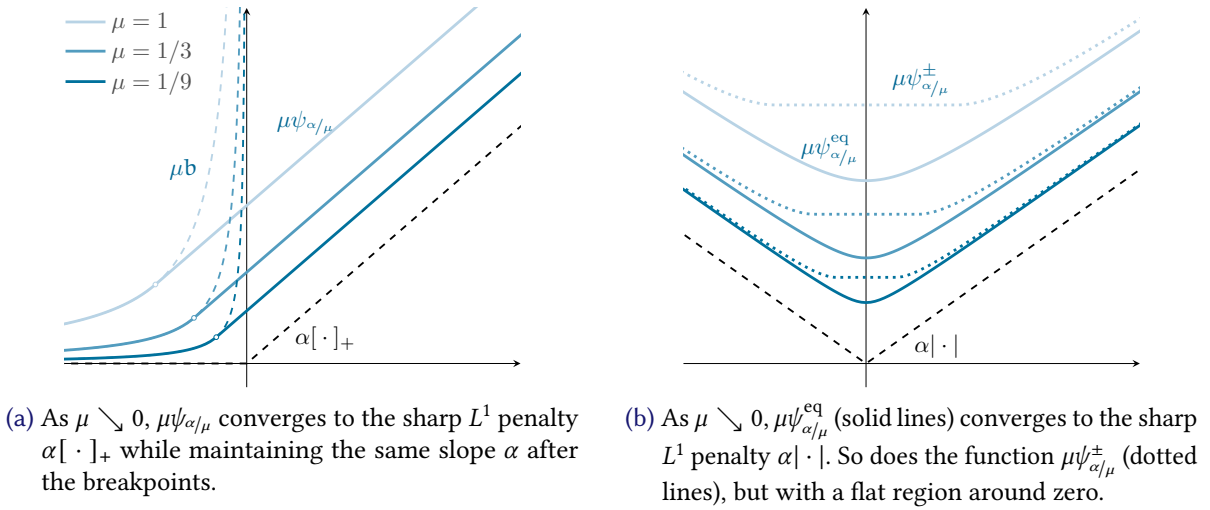


Figure 3: Limiting behavior of $\mu\psi_{\alpha/\mu}$ (left) and $\mu\psi_{\alpha/\mu}^{\text{eq}}$ (right) with constant α as $\mu \searrow 0$. These examples employ the inverse barrier $b(t) = -\frac{1}{t} + \delta_{(-\infty, 0)}$.

Lipschitz modulus for b' on $(-\infty, \rho]$. If $c(\mathbf{x}) \leq \rho < c(\mathbf{y})$, then $\psi'_{\rho^*}(c(\mathbf{y})) = \psi'_{\rho^*}(\rho)$ and, by continuity, there exists $t \in [0, 1]$ such that $c(\mathbf{x} + t(\mathbf{y} - \mathbf{x})) = \rho$, so that $\psi'_{\rho^*}(c(\mathbf{x} + t(\mathbf{y} - \mathbf{x}))) = \psi'_{\rho^*}(\rho)$, resulting in the same bound $Bt\ell^2\|\mathbf{x} - \mathbf{y}\| \leq B\ell^2\|\mathbf{x} - \mathbf{y}\|$. Lastly, if $\rho \leq c(\mathbf{x}) \leq c(\mathbf{y})$ then the last term is zero. In all cases we conclude that

$$\|\nabla(\psi_{\rho^*} \circ c)(\mathbf{x}) - \nabla(\psi_{\rho^*} \circ c)(\mathbf{y})\| \leq \left(\frac{\alpha}{\mu}L + B\ell^2\right)\|\mathbf{x} - \mathbf{y}\| \quad \forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^n,$$

proving the claim. $\square$

4 ALGORITHMIC FRAMEWORK

The main ingredient for the proposed numerical scheme is the penalty-barrier problem $(P_{\alpha, \mu})$. As shown in the previous section, the cost function in $(P_{\alpha, \mu})$ converges pointwise to the original hard-constrained cost $q + \delta_{\mathbb{R}^m} \circ c + \delta_{\{0\}} \circ c_{\text{eq}}$ of (P) as $\mu \searrow 0$ and $\alpha \nearrow \infty$. Following a homotopic rationale, this motivates solving (up to approximate local optimality) instances of $(P_{\alpha, \mu})$ for progressively small values of μ and larger values of α . This is the leading idea of the algorithmic framework of [Algorithm 1](#) presented in this section, whose name ‘*Marge*’ evokes the key underlying feature of *marginalization* discussed in [Section 3.2](#). The update rules for the coefficients are carefully designed so as to ensure that the output satisfies suitable optimality conditions for the original problem (P) , as well as to prevent the L^1 penalization parameter α in $(P_{\alpha, \mu})$ from divergent behaviors under favorable conditions on the problem. This is the reason behind the involvement of the conjugate b^* in the update criterion at [Step 1.8](#), as will be revealed in [Sections 4.2](#) and [4.3](#) through a systematic study of the properties of the barrier b in the generality of [Assumption 1](#) as well as when specialized to the convex case.

[Algorithm 1](#) is not tied to any particular solver for addressing each instance of $(P_{\alpha, \mu})$ at [Step 1.1](#). Whenever q amounts to the sum of a differentiable and a prox-friendly function (in the sense that its proximal mapping is easily computable), such structure is retained by the cost function in $(P_{\alpha, \mu})$, indicating that proximal-gradient based methods are suitable candidates. This was also the case in the purely interior-point based IPprox of [\[13\]](#), which considers a plain proximal gradient with a backtracking routine for selecting the stepsizes. Differently from the subproblems of IPprox in which the differentiable term is extended-real valued, the differentiable term in $(P_{\alpha, \mu})$ is smooth on the whole $\mathbb{R}^n$. This enables the employment of more sophisticated proximal-gradient-type algorithms such as

Algorithm 1 Marge: A combined penalty and barrier framework for constrained optimization

REQUIRE tolerances $\epsilon_p, \epsilon_d \geq 0$; parameters $\alpha_0, \mu_0 > 0$, $\epsilon_0 \geq \epsilon_d$, $\delta_\alpha > 1$ and $\delta_\epsilon, \delta_\mu \in (0, 1)$

REPEAT FOR $k = 0, 1, \dots$

- 1.1: Find an ϵ_k -stationary point $\mathbf{x}^k$ for $(P_{\alpha, \mu})$ with $(\alpha, \mu) = (\alpha_k, \mu_k)$
- 1.2: set $\mathbf{y}_i^k = \mu_k \psi'_{\alpha_k/\mu_k}(c_i(\mathbf{x}^k))$ as in (3.6), $i = 1, \dots, m$
- 1.3: set $\mathbf{y}_{\text{eq},j}^k = \mu_k (\psi'_{\alpha_k/\mu_k})'(c_{\text{eq},j}(\mathbf{x}^k))$ as in (3.8), $j = 1, \dots, m_{\text{eq}}$
- 1.4: $p_k = \max \{ \|[c(\mathbf{x}^k)]_+\|_\infty, \|\mathbf{c}_{\text{eq}}(\mathbf{x}^k)\|_\infty \}$ % constraints violation
- 1.5: $s_k = s(\mathbf{x}^k, \mathbf{y}^k, \mathbf{y}_{\text{eq}}^k)$ as in (3.1) % complementarity violation
- 1.6: IF $(\epsilon_k, p_k, s_k) \leq (\epsilon_d, \epsilon_p, \epsilon_p)$ THEN
 RETURN $(\mathbf{x}^k, \mathbf{y}^k)$ (ϵ_p, ϵ_d) -KKT pair for (P)
- 1.7: $\epsilon_{k+1} = \max \{ \delta_\epsilon \epsilon_k, \epsilon_d \}$
- 1.8: IF $p_k > \max \left\{ \epsilon_p, 2(m + m_{\text{eq}}) \frac{-b^*(\alpha_k/\mu_k)}{\alpha_k/\mu_k} \right\}$ THEN $\alpha_{k+1} = \delta_\alpha \alpha_k$, ELSE $\alpha_{k+1} = \alpha_k$
- 1.9: IF $s_k > \epsilon_p$ OR $\alpha_{k+1} = \alpha_k$ THEN $\mu_{k+1} = \delta_\mu \mu_k$, ELSE $\mu_{k+1} = \mu_k$

PANOC⁺ [29, 12] that make use of higher-order information to considerably enhance convergence speed. This claim will be substantiated with numerical evidence in Section 5; in this section, we instead focus on properties of the *outer* Algorithm 1 that are independent of the *inner* solver.

Remark 4.1. Throughout our convergence analysis, it is assumed that Algorithm 1 is well-defined, thus requiring that each subproblem $(P_{\alpha, \mu})$ at Step 1.1 admits an approximate stationary point. Moreover, some of the following statements assume the existence of an accumulation point $\bar{\mathbf{x}}$ for a sequence $(\mathbf{x}^k)_{k \in \mathbb{N}}$ generated by Algorithm 1. In general, these preconditions can be verified with coercivity or (level) boundedness arguments; for instance, the objective function of $(P_{\alpha, \mu})$ is bounded from below whenever $\text{dom } q$ is a compact set.

Remark 4.2 (parameter update variant). According to Steps 1.8 and 1.9, at each iteration either α_k or μ_k (possibly both) is updated. With the aim of slowing down the reduction of μ_k to attenuate unnecessary ill-conditioning, another viable option is to restrict the update condition as in

1.9': IF $s_k > \epsilon_p$ OR $(\alpha_{k+1}, \epsilon_{k+1}) = (\alpha_k, \epsilon_k)$ THEN $\mu_{k+1} = \delta_\mu \mu_k$, ELSE $\mu_{k+1} = \mu_k$

allowing also the possibility of neither parameter being updated. Such circumstance takes place only if $\epsilon_k > \epsilon_d$, owing to Step 1.7. In this event, then, it is only the stationarity tolerance ϵ_k for Step 1.1 that is decreased, and the next iteration reduces to solving the same subproblem with higher accuracy. To avoid unnecessary notational complexity in the proofs, we adhere to the steps outlined in Algorithm 1, while noting that the entire theoretical framework remains valid for this variant, with only minor changes required to the iteration indexing. For our numerical tests in Section 5, however, Step 1.9 of Algorithm 1 is replaced by the variant above.

4.1 CONVERGENCE ANALYSIS

Lemma 4.3 (properties of the iterates). Suppose that Assumptions 1 and 2 hold, and consider the iterates generated by Algorithm 1. At every iteration k the following hold:

- (i) $0 < \mathbf{y}^k \leq \alpha_k \mathbf{1}$ and $|\mathbf{y}_{\text{eq}}^k| < \alpha_k \mathbf{1}$.
- (ii) $(\mathbf{x}^k \in \text{dom } q \text{ and}) \text{dist}_{\partial q(\mathbf{x}^k)}(-\mathbf{J}c(\mathbf{x}^k)^\top \mathbf{y}^k - \mathbf{J}c_{\text{eq}}(\mathbf{x}^k)^\top \mathbf{y}_{\text{eq}}^k) \leq \epsilon_k$.

- (iii) If $(\epsilon_p > 0 \text{ and}) \mu_k \leq \frac{\epsilon_p}{2b'(-\epsilon_p)}$, then $s_k \leq \epsilon_p$.
- (iv) For $k \geq 1$, either $\alpha_k = \delta_\alpha \alpha_{k-1}$ or $\mu_k = \delta_\mu \mu_{k-1}$ (possibly both); in particular, letting $\rho_k^* := \alpha_k / \mu_k$ and $\delta_{\rho^*} := \min \left\{ \delta_\alpha, \delta_\mu^{-1} \right\}$ it holds that $\rho_k^* \geq \delta_{\rho^*} \rho_{k-1}^*$.

Proof. Assertions 4.3(i) and 4.3(ii) follow from Lemma 3.8. Assertion 4.3(iv) is obvious by observing that whenever $\alpha_{k+1} = \alpha_k$ the update $\mu_{k+1} = \delta_\mu \mu_k$ is enforced.

We finally turn to assertion 4.3(iii), and suppose that $\mu_k \leq \frac{\epsilon_p}{2b'(-\epsilon_p)}$. Then, for all i such that $[c_i(\mathbf{x}^k)]_- > \epsilon_p$ (or, equivalently, $c_i(\mathbf{x}^k) < -\epsilon_p$), one has

$$y_i^k = \mu_k \psi'_{\alpha_k/\mu_k}(c_i(\mathbf{x}^k)) \leq \mu_k b'(c_i(\mathbf{x}^k)) \leq \mu_k b'(-\epsilon_p) \leq \frac{1}{2} \epsilon_p \leq \epsilon_p,$$

where the first inequality follows from the definition of $\mathbf{y}^k$ together with (3.6), and the second one owes to monotonicity of b' . On the other hand, for all i such that $c_i(\mathbf{x}^k) > \epsilon_p$ (in fact, more generally when $c_i(\mathbf{x}^k) > 0$), one has that $y_i^k = \mu_k \psi'_{\alpha_k/\mu_k}(c_i(\mathbf{x}^k)) = \alpha_k$, cf. (3.6). Next, if j is such that $c_{\text{eq},j}(\mathbf{x}^k) > \epsilon_p$, one has that

$$y_{\text{eq},j}^k = \mu_k (\psi_{\alpha_k/\mu_k}^{\text{eq}})'(c_{\text{eq},j}(\mathbf{x}^k)) \geq \mu_k (\alpha_k/\mu_k - 2b'(-\epsilon_p)) \geq \alpha - \epsilon_p,$$

where the first inequality follows from Theorem 3.6. By the same arguments, we deduce that $y_{\text{eq},j}^k \leq -\alpha + \epsilon_p$ holds for all j such that $c_{\text{eq},j}(\mathbf{x}^k) < -\epsilon_p$. In summary, for all $i = 1, \dots, m$ at least one among y_i^k and $[c_i(\mathbf{x}^k)]_-$ is not larger than ϵ_p , and at least one among $\alpha - y_i^k$ and $[c_i(\mathbf{x}^k)]_+$ is zero. Similarly, for all $j = 1, \dots, m_{\text{eq}}$ at least one among $\alpha - y_{\text{eq},j}^k \text{sgn}(c_{\text{eq},j}(\mathbf{x}^k))$ and $|c_{\text{eq},j}(\mathbf{x}^k)|$ is not larger than ϵ_p , overall proving that $s_k \leq \epsilon_p$. $\square$

Corollary 4.4 (stationarity of feasible limit points). *Let Assumptions 1 and 2 hold, and consider the iterates generated by Algorithm 1. If the algorithm runs indefinitely, then $-\frac{\mu_k}{\alpha_k} b^*(\alpha_k/\mu_k) \searrow 0$ as $k \rightarrow \infty$. Moreover, any feasible accumulation point of $(\mathbf{x}^k)_{k \in \mathbb{N}}$ that belongs to $\text{dom } q$ is A-KKT-optimal for (P).*

Proof. The monotonic vanishing of $-\frac{\mu_k}{\alpha_k} b^*(\alpha_k/\mu_k)$ follows from Lemmas 4.3(iv) and A.1(iv). Suppose that $(\mathbf{x}^k)_{k \in K} \rightarrow \bar{\mathbf{x}} \in \text{dom } q$ with $\mathbf{c}(\bar{\mathbf{x}}) \leq \mathbf{0}$ and $\mathbf{c}_{\text{eq}}(\bar{\mathbf{x}}) = \mathbf{0}$, and let i be such that $c_i(\bar{\mathbf{x}}) < 0$ (if such an i does not exist, then there is nothing to show). According to Definition 2.3 and Remark 2.7(i), it suffices to show that $(y_i^k)_{k \in K} \rightarrow 0$; in turn, by definition of $\mathbf{y}^k$ and continuity of $\mathbf{c}$ it suffices to show that $\mu_k \searrow 0$. If $\epsilon_p > 0$, then continuity of $\mathbf{c}$ implies that $\alpha_{k+1} = \alpha_k$ for all $k \in K$ large enough, hence, by virtue of Lemma 4.3(iv), $\mu_{k+1} = \delta_\mu \mu_k$ for all such k . Since $(\mu_k)_{k \in \mathbb{N}}$ is monotone, in this case $\mu_k \searrow 0$ as $k \rightarrow \infty$. Suppose instead that $\epsilon_p = 0$, and, to arrive to a contradiction, that μ_k is asymptotically constant. This implies that $s_k \leq \epsilon_p = 0$ eventually always holds, which is a contradiction since

$$s_k \geq \min \left\{ y_i^k, [c_i(\mathbf{x}^k)]_- \right\} \geq \min \left\{ y_i^k, -\frac{1}{2} c_i(\bar{\mathbf{x}}) \right\} > 0 \quad \forall k \in K \text{ large},$$

where the first inequality follows by definition of s_k , cf. Step 1.5, the second one for $k \in K$ large since $c_i(\mathbf{x}^k) \rightarrow c_i(\bar{\mathbf{x}}) < 0$ as $K \ni k \rightarrow \infty$, and the last one because $\mathbf{y}^k > \mathbf{0}$. $\square$

The update rule for the penalty parameter does not demand (approximate) feasibility, but it depends on a relaxed condition at Step 1.8. By (A.1) in the appendix, the second term vanishes as $\alpha/\mu \rightarrow \infty$, so the penalty parameter is eventually increased as needed to achieve ϵ_p -feasibility. The relaxation of this condition using a quantity involving the conjugate b^* mitigates the growth of α . Simultaneously, under suitable choices of the barrier b , it ensures that this parameter remains unchanged only if the constraints violation stays within a controlled range, as will be ultimately demonstrated in Corollary 4.12.

Theorem 4.5. *Suppose that Assumptions 1 and 2 hold, and consider the iterates generated by Algorithm 1 with $\epsilon_p, \epsilon_d > 0$. Then, $\inf_{k \in \mathbb{N}} \mu_k > 0$ and exactly one of the following scenarios occurs:*

- (A) either the algorithm terminates returning an (ϵ_p, ϵ_d) -KKT stationary point for (P),
 (B) or it runs indefinitely with $s_k \leq \epsilon_p < p_k$ for all k large enough, and $(\alpha_k)_{k \in \mathbb{N}} \nearrow \infty$.

In the latter case, if $\text{dom } q$ is closed and q is continuous relative to it, then for any accumulation point $\bar{x}$ of $(x^k)_{k \in \mathbb{N}}$ one has that $(\bar{x}, q(\bar{x}))$ is KKT-stationary for the feasibility problem

$$(4.1) \quad \underset{(x,t) \in \text{epi } q}{\text{minimize}} \quad \| [c(x)]_+ \|_1 + \| c_{\text{eq}}(x) \|_1,$$

in the sense that $(0, 0) \in \partial [\| [c(\bar{x})]_+ \|_1 + \| c_{\text{eq}}(\bar{x}) \|_1] \times \{0\} + N_{\text{epi } q}(\bar{x}, q(\bar{x}))$.

Proof. Since $\mu_{k+1} \leq \mu_k$ for all k , and μ_k is linearly reduced whenever $s_k > \epsilon_p$, we conclude that (either the algorithm terminates or) $s_k \leq \epsilon_p$ eventually always holds.

If the algorithm returns $(x^k, y^k, y_{\text{eq}}^k)$, then the compliance with the termination criteria combined with the fact that $y^k > 0$ for all k , see Lemma 4.3(i), ensures that such triplet meets all conditions in Definition 2.5, and hence it is (ϵ_p, ϵ_d) -KKT-stationary for (P).

Suppose instead that the algorithm does not terminate. Clearly, $\epsilon_k = \epsilon_d$ holds for k large enough, so that the only unmet termination criterion is eventually $p_k \leq \epsilon_d$. Therefore, $p_k > \epsilon_p$ holds for every k large enough. It follows from Lemmas 4.3(iv) and A.1(iv) that $-2(m + m_{\text{eq}})b^*(\rho_k^*)/\rho_k^*$ eventually drops below ϵ_p , implying that the condition for increasing α_{k+1} at Step 1.8 reduces to $p_k > \epsilon_p$. Having shown that this is eventually always the case, $\alpha_{k+1} = \delta_\alpha \alpha_k$ always holds for k large, $\alpha_k \nearrow \infty$, and μ_k is eventually never updated, cf. Step 1.8.

To conclude, suppose that $\text{dom } q$ is closed and that q is continuous relative to this set. By Lemma 3.8, for every k we have that there exists $\eta^k \in \mathbb{R}^n$ with $\|\eta^k\| \leq \epsilon_d$ such that

$$\eta^k - Jc(x^k)^\top y^k - Jc_{\text{eq}}(x^k)^\top y_{\text{eq}}^k \in \partial q(x^k).$$

Let $\bar{x}$ be the limit of a subsequence $(x^k)_{k \in K}$ and, up to extracting, let $\bar{\lambda}$ and $\bar{\lambda}_{\text{eq}}$ be the limits of $(\frac{1}{\alpha_k} y^k)_{k \in K}$ and $(\frac{1}{\alpha_k} y_{\text{eq}}^k)_{k \in K}$, respectively. The definition of y^k and y_{eq}^k together with the continuity of c and c_{eq} yields that

$$\bar{\lambda}_i = \begin{cases} 0 & \text{if } c_i(\bar{x}) < 0 \\ 1 & \text{if } c_i(\bar{x}) > 0 \\ \in [0, 1] & \text{if } c_i(\bar{x}) = 0 \end{cases} \quad \text{and} \quad \bar{\lambda}_{\text{eq},j} = \begin{cases} -1 & \text{if } c_{\text{eq},j}(\bar{x}) < 0 \\ 1 & \text{if } c_{\text{eq},j}(\bar{x}) > 0 \\ \in [-1, 1] & \text{if } c_{\text{eq},j}(\bar{x}) = 0 \end{cases}$$

or, equivalently,

$$(4.2) \quad \bar{\lambda}_i \in \partial [\cdot]_+(c_i(\bar{x})) \quad \text{and} \quad \bar{\lambda}_{\text{eq},j} \in \partial \|\cdot\|_1(c_{\text{eq},j}(\bar{x}))$$

for every $i = 1, \dots, m$ and $j = 1, \dots, m_{\text{eq}}$. Since $\text{dom } q$ is closed and $x^k \in \text{dom } q$ for all k , one has that $q(\bar{x}) < \infty$. Moreover, it follows from the continuity assumption that $q(x^k) \rightarrow q(\bar{x})$ as $K \ni k \rightarrow \infty$, hence that $-Jc(\bar{x})^\top \bar{\lambda} - Jc_{\text{eq}}(\bar{x})^\top \bar{\lambda}_{\text{eq}} \in \partial^\infty q(\bar{x})$, where $\partial^\infty q(\bar{x})$ denotes the horizon subdifferential of q at $\bar{x}$. Appealing to [25, Thm. 8.9 and Ex. 8.14], this means that

$$(4.3) \quad (-Jc(\bar{x})^\top \bar{\lambda} - Jc_{\text{eq}}(\bar{x})^\top \bar{\lambda}_{\text{eq}}, 0) \in N_{\text{epi } q}(\bar{x}, q(\bar{x})) = \partial \delta_{\text{epi } q}(\bar{x}, q(\bar{x})).$$

Since

$$\partial \| [c(x)]_+ \|_1 = \sum_{i=1}^m \partial [c_i(x)]_+ = \sum_{i=1}^m \partial [\cdot]_+(c_i(x)) \nabla c_i(x)$$

and similarly

$$\partial \| c_{\text{eq}}(x) \|_1 = \sum_{j=1}^{m_{\text{eq}}} \partial |c_{\text{eq},j}(x)| = \sum_{j=1}^{m_{\text{eq}}} \partial |\cdot|(c_{\text{eq},j}(x)) \nabla c_{\text{eq},j}(x),$$

see [25, Ex. 10.26], it follows from (4.2) and continuity of $\|[\mathbf{c}(\mathbf{x})]_+\|_1$ and $\|\mathbf{c}_{\text{eq}}(\mathbf{x})\|_1$ that $\bar{\lambda} \in \partial\|[\mathbf{c}(\bar{\mathbf{x}})]_+\|_1$ and $\bar{\lambda}_{\text{eq}} \in \partial\|\mathbf{c}_{\text{eq}}(\bar{\mathbf{x}})\|_1$. Combining with (4.3) concludes the proof. $\square$

Remark 4.6 (relaxing continuity of q). Similarly to the commentary after Remark 2.7(ii) pertaining the *limiting* subdifferential ∂q , note that continuity of q on its domain is used to guarantee that equality $\partial^\infty q(\bar{\mathbf{x}}) = \tilde{\partial}^\infty q(\bar{\mathbf{x}})$ (as opposed to inclusion $\partial^\infty q(\bar{\mathbf{x}}) \subseteq \tilde{\partial}^\infty q(\bar{\mathbf{x}})$) holds for the *horizon* subdifferential $\partial^\infty q(\bar{\mathbf{x}})$ of q at any point $\bar{\mathbf{x}} \in \text{dom } q$, where for reference we denote

$$\tilde{\partial}^\infty q(\bar{\mathbf{x}}) := \left\{ \bar{\mathbf{v}} \in \mathbb{R}^n \mid \exists \lambda_k \rightarrow 0, (\mathbf{x}^k, \mathbf{v}^k) \in \text{gph } \hat{\partial} q \text{ such that } (\mathbf{x}^k, \lambda_k \mathbf{v}^k) \rightarrow (\bar{\mathbf{x}}, \bar{\mathbf{v}}) \right\}.$$

Here, the definition of $\tilde{\partial}^\infty q(\bar{\mathbf{x}})$ matches that of $\partial^\infty q(\bar{\mathbf{x}})$ except that the “ q -attentive” constraint $q(\mathbf{x}^k) \rightarrow q(\bar{\mathbf{x}})$ is not imposed for the sequence $(\mathbf{x}^k)_{k \in \mathbb{N}}$. The validity of Theorem 4.5 is thus unaffected if continuity of q on its domain is replaced by $\partial^\infty q$ being as above at any point $\bar{\mathbf{x}} \in \text{dom } q$, and once again this generalization covers important examples such as functions involving L^0 -norm terms.

On the other hand, note that mere outer semicontinuity of ∂q as in Remark 2.7(ii) is not enough to ensure this identity. To see this, consider the function $h : \mathbb{R} \rightarrow \mathbb{R}$ defined as $h(x) = 0$ for $x \leq 0$ and $h(x) = 1 - \sqrt{x}$ for $x > 0$. One can easily see that $\hat{\partial} h = \partial h$ with $\partial h(x) = \{\nabla h(x)\}$ for $x \neq 0$ and $\partial h(0) = [0, \infty)$ is everywhere osc, yet $\partial^\infty q(0) = [0, \infty) \neq \mathbb{R} = \tilde{\partial}^\infty q(0)$.

The abuse of terminology to express KKT-stationarity in terms of subdifferentials passes through the same construct relating $(\mathbf{P}_\alpha)$ and $(\mathbf{Q}_\alpha)$, in which a slack variable is tacitly introduced to reformulate the L^1 norm; see the discussion in Section 3.1. More importantly, the involvement in (4.1) of the epigraph of q , as opposed to its domain, is a necessary technicality that cannot be avoided in the generality of Assumption 1, as we illustrate next.

Remark 4.7 (epi q vs dom q). Stationarity for (4.1) is, in general, weaker than that for the more natural minimal infeasibility violation problem

$$(4.4) \quad \underset{\mathbf{x} \in \text{dom } q}{\text{minimize}} \quad \|[\mathbf{c}(\mathbf{x})]_+\|_1 + \|\mathbf{c}_{\text{eq}}(\mathbf{x})\|_1.$$

To see how this notion may be violated, consider $q(x) = \sqrt{|x|}$ and $c(x) = x + 1$, so that (P) reads

$$\underset{x \in \mathbb{R}}{\text{minimize}} \quad \sqrt{|x|} \quad \text{subject to } x \leq -1.$$

The point $x^k = 0$ is stationary for any subproblem $(\mathbf{P}_{\alpha, \mu})$ with arbitrary $\alpha, \mu > 0$, and therefore constitutes a feasible choice in Algorithm 1. However, the limit $\bar{x} = 0$ of the corresponding constant sequence is not stationary for the minimization of $[x + 1]_+$ over $\text{dom } q = \mathbb{R}$. Nevertheless,

$$\partial[x + 1]_+(0) \times \{0\} + \text{N}_{\text{epi } q}(0, 0) = \{(1, 0)\} + (\mathbb{R} \times \{0\}) \ni (0, 0),$$

confirming that $(0, 0)$ is stationary for the epigraphical problem (4.1).

We next formally illustrate why stationarity for (4.4) always implies that for (4.1), and identify the culprit of a possible discrepancy in uncontrolled growths around $\bar{\mathbf{x}}$ from within $\text{dom } q$. To this end, we remind that a proper function $h : \mathbb{R}^n \rightarrow \bar{\mathbb{R}}$ is said to be *calm* at a point $\bar{\mathbf{x}} \in \text{dom } h$ relative to a set $X \ni \bar{\mathbf{x}}$ if

$$\liminf_{\substack{X \ni \mathbf{x} \rightarrow \bar{\mathbf{x}} \\ \mathbf{x} \neq \bar{\mathbf{x}}}} \frac{|h(\mathbf{x}) - h(\bar{\mathbf{x}})|}{\|\mathbf{x} - \bar{\mathbf{x}}\|} < \infty,$$

and that this condition is weaker than strict continuity.

Lemma 4.8. *Let $h : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ be proper and lsc. Then, for any $\bar{x} \in \text{dom } h$ one has*

$$\widehat{N}_{\text{dom } h}(\bar{x}) \subseteq \left\{ \bar{v} \in \mathbb{R}^n \mid (\bar{v}, 0) \in \widehat{N}_{\text{epi } h}(\bar{x}, h(\bar{x})) \right\}$$

and

$$(4.5) \quad N_{\text{dom } h}(\bar{x}) \subseteq \left\{ \bar{v} \in \mathbb{R}^n \mid (\bar{v}, 0) \in N_{\text{epi } h}(\bar{x}, h(\bar{x})) \right\} = \partial^\infty h(\bar{x}).$$

When h is convex, both inclusions hold as equality. More generally, when h is calm (in particular, if it is strictly continuous) at $\bar{x}$ relative to $\text{dom } h$, then the first inclusion holds as equality, and so does the second one when such property holds not only at $\bar{x}$, but also at all points in $\text{dom } h$ close to it.

Proof. The relations in the convex case are shown in [25, Thm. 8.9 and Prop. 8.12]; in what follows, we consider an arbitrary proper and lsc function h . Let $\bar{v} \in \widehat{N}_{\text{dom } h}(\bar{x})$ and let $\text{epi } h \ni (x^k, t^k) \rightarrow (\bar{x}, h(\bar{x}))$. Then, there exists $\varepsilon_k \rightarrow 0$ such that $\langle \bar{v}, x^k - \bar{x} \rangle \leq \varepsilon_k \|x^k - \bar{x}\|$ holds for every k , hence

$$\varepsilon_k \left\| \begin{pmatrix} x^k - \bar{x} \\ t^k - h(\bar{x}) \end{pmatrix} \right\| \geq \varepsilon_k \|x^k - \bar{x}\| \geq \langle \bar{v}, x^k - \bar{x} \rangle = \left\langle \begin{pmatrix} \bar{v} \\ 0 \end{pmatrix}, \begin{pmatrix} x^k - \bar{x} \\ t^k - h(\bar{x}) \end{pmatrix} \right\rangle.$$

By the arbitrariness of the sequence, we conclude that $(\bar{v}, 0) \in \widehat{N}_{\text{epi } h}(\bar{x}, h(\bar{x}))$. The same inclusion must then hold for the limiting normal cones, leading to (4.5), where the identity follows from [25, Thm. 8.9].

Suppose now that there exists $\kappa > 0$ such that $|h(x) - h(\bar{x})| \leq \kappa \|x - \bar{x}\|$ for $x \in \text{dom } h$ close to $\bar{x}$, and suppose that $(\bar{v}, 0) \in \widehat{N}_{\text{epi } h}(\bar{x}, h(\bar{x}))$. Let $\text{dom } h \ni x^k \rightarrow \bar{x}$, and note that $\text{epi } h \ni (x^k, h(x^k)) \rightarrow (\bar{x}, h(\bar{x}))$. Then, there exists $\varepsilon_k \rightarrow 0$ such that

$$\begin{aligned} \langle \bar{v}, x^k - \bar{x} \rangle &= \left\langle \begin{pmatrix} \bar{v} \\ 0 \end{pmatrix}, \begin{pmatrix} x^k - \bar{x} \\ t^k - h(\bar{x}) \end{pmatrix} \right\rangle \leq \varepsilon_k \left\| \begin{pmatrix} x^k - \bar{x} \\ h(x^k) - h(\bar{x}) \end{pmatrix} \right\| \\ &\leq \varepsilon_k \left\| \begin{pmatrix} x^k - \bar{x} \\ \kappa \|x^k - \bar{x}\| \end{pmatrix} \right\| = \varepsilon_k \sqrt{1 + \kappa^2} \|x^k - \bar{x}\|, \end{aligned}$$

where the second inequality holds for k large enough. Arguing again by the arbitrariness of the sequence, we conclude that $\bar{v} \in \widehat{N}_{\text{dom } h}(\bar{x})$. Finally, when h is calm relative to its domain at all points $x \in \text{dom } h$ close to $\bar{x}$, then the identity $\widehat{N}_{\text{dom } h}(x) \times \{0\} = \widehat{N}_{\text{epi } h}(x, h(x))$ holds for all such points, and a limiting argument then yields that $N_{\text{dom } h}(\bar{x}) \times \{0\} = N_{\text{epi } h}(\bar{x}, h(\bar{x}))$ holds for the limiting normal cones. Therefore, the inclusion in (4.5) holds as equality, which concludes the proof. $\square$

When q is locally Lipschitz relative to its domain, the inclusion $(-Jc(\bar{x})^\top \bar{\lambda} - Jc_{\text{eq}}(\bar{x})^\top \bar{\lambda}_{\text{eq}}, 0) \in \partial \delta_{\text{epi } q}(\bar{x}, q(\bar{x}))$ derived in (4.3) thus simplifies as $-Jc(\bar{x})^\top \bar{\lambda} - Jc_{\text{eq}}(\bar{x})^\top \bar{\lambda}_{\text{eq}} \in \partial \delta_{\text{dom } q}(\bar{x})$, and the conclusions of Theorem 4.5 can be strengthened as follows.

Corollary 4.9. *Additionally to Assumptions 1 and 2, suppose that q is locally Lipschitz continuous on its domain, assumed to be closed. Consider the iterates generated by Algorithm 1 with $\epsilon_p, \epsilon_d > 0$. Then, $\inf_{k \in \mathbb{N}} \mu_k > 0$ and exactly one of the following scenarios occurs:*

- (A) *either the algorithm terminates returning an (ϵ_p, ϵ_d) -KKT stationary point for (P),*
- (B) *or it runs indefinitely with $s_k \leq \epsilon_p < p_k$ for all k large enough, and $(\alpha_k)_{k \in \mathbb{N}} \nearrow \infty$. Moreover, any accumulation point $\bar{x}$ of $(x^k)_{k \in \mathbb{N}}$ is KKT-stationary for the feasibility problem*

$$\text{minimize } \|[c(x)]_+\|_1 + \|c_{\text{eq}}(x)\|_1,$$

$$x \in \text{dom } q$$

in the sense that $0 \in \partial[\|[c(\bar{x})]_+\|_1 + \|c_{\text{eq}}(\bar{x})\|_1] + N_{\text{dom } q}(\bar{x})$.

Even under smoothness assumptions on q , the outcomes presented in [Corollary 4.9](#) are usual for nonlinear programming solvers in the sense that, in general, one cannot exclude the possibility that the iterates remain trapped in an infeasible region. If [Algorithm 1](#) achieves approximate feasibility, then outcome (A) guarantees that it returns an approximate KKT stationary point for (P). Otherwise, outcome (B) specifies that the constraint violation is increasingly emphasized and [Algorithm 1](#) steers toward the solution of a feasibility problem. Several analogous results appear in the literature, such as [6, Thm. 4.4], [11, Thm. 3.3] and [4, Thm. 6.2], among others. Stronger convergence guarantees can be established based on more restrictive assumptions, which are not of interest here beyond this remark.

In practice, well-established solvers such as Ipopt [32] rely on auxiliary procedures whose purpose is to compute a new iterate that yields sufficient improvement. When in difficulty, Ipopt switches to a *feasibility restoration phase* that attempts to decrease the constraint violation. Equipped with this mechanism, all the limit points it generates are feasible under the (tautological) assumption that “the feasibility restoration phase . . . always terminates successfully” [31, Ass. G and Thm. 1].

Another strategy is that of [8, §3.2], whose update rules of penalty (and barrier) parameters are designed to result in rapid convergence to optimal points of the original problem (P) or, at least, of the feasibility problem (4.1)—in the context of nonlinear programming. Again, this technique promotes feasibility but additional properties are needed to obtain convergence guarantees. Our update rules at [Steps 1.8](#) and [1.9](#) draw inspiration from the Fiacco–McCormick approach [15] and are comparable to the conservative strategy in [8].

Finally, also the question of whether the sequence of penalty parameters $(\alpha_k)_{k \in \mathbb{N}}$ remains bounded or not has practical interest, but theoretical findings in this direction are limited and circumscribed, even in the nonlinear programming setting. Boundedness of the sequence of multipliers, strict complementarity, and constraint qualifications are some of the assumptions adopted to establish the boundedness of penalty parameters; cf. [4, §7] and [16, Ass. 3 and Thm. 1].

4.2 BARRIER’S PROPERTIES

According to its update rule in [Algorithm 1](#), before a desired feasibility violation $p_k \leq \epsilon_p$ has been reached, $\alpha_{k+1} = \alpha_k$ means that $p_k \leq 2(m + m_{\text{eq}}) \frac{-b^*(\rho_k^*)}{\rho_k^*}$, where $\rho_k^* = \alpha_k / \mu_k$. As shown in [Lemma 4.3\(iv\)](#), regardless of whether α_k is updated or not, ρ_k^* grows linearly over the iterations, specifically as $\rho_k^* \geq \rho_0^* \delta_{\rho^*}^k$. Therefore, having $\alpha_{k+1} = \alpha_k$ implies in particular that either the constraint violation p_k is within a desired tolerance ϵ_p , or that it is controlled from above by $2(m + m_{\text{eq}}) \frac{-b^*(\rho_k^*)}{\rho_k^*} \leq 2(m + m_{\text{eq}}) \frac{-b^*(\delta_{\rho^*}^k \rho_0^*)}{\delta_{\rho^*}^k \rho_0^*}$, where the inequality follows from (asymptotic) monotonicity of $-b^*(t^*)/t^*$, cf. [Lemma A.1\(iv\)](#). This means that a desired bound on feasibility violation whenever α_k is not updated can be enforced through suitable choices of the barrier b . This will be particularly significant in the convex case, for it can be shown that α_k does eventually remain constant under reasonable assumptions.

Lemma 4.10. *Let [Assumptions 1](#) and [2](#) hold, and consider the iterates generated by [Algorithm 1](#). Suppose that there exists $\theta \in (0, 1)$ such that the barrier b satisfies $b(\theta t) \leq \theta \delta_{\rho^*} b(t)$ for every $t < 0$ (resp. for every $t < 0$ close enough to 0), where $\delta_{\rho^*} := \min \left\{ \delta_{\mu}^{-1}, \delta_{\alpha} \right\} > 1$. Then,*

$$(4.6) \quad \alpha_{k+1} = \alpha_k \quad \Rightarrow \quad p_k \leq \max \left\{ \epsilon_p, -2(m + m_{\text{eq}}) \frac{\mu_0}{\alpha_0} b^* \left(\frac{\alpha_0}{\mu_0} \right) \theta^k \right\}$$

holds for every k (resp. for every k large enough).

Proof. To simplify the presentation, without loss of generality let us set $\epsilon_p = 0$. We have already argued that $\alpha_{k+1} = \alpha_k$ implies $p_k \leq 2(m + m_{\text{eq}}) \pi_k$, where $\pi_k := \frac{-b^*(\delta_{\rho^*}^k \rho_0^*)}{\delta_{\rho^*}^k \rho_0^*}$ for all $k \in \mathbb{N}$. It thus suffices to

show that $\pi_k \leq -\frac{\mu_0}{\alpha_0} b^*(\frac{\alpha_0}{\mu_0}) \theta^k$. To this end, notice that for every $t^* > 0$ one has

$$\frac{-b^*(\delta_{\rho^*} t^*)}{\delta_{\rho^*} t^*} \leq \theta \frac{-b^*(t^*)}{t^*} \Leftrightarrow b^*(t^*) \leq \frac{b^*(\delta_{\rho^*} t^*)}{\theta \delta_{\rho^*}} = \sup_{\tau} \left\{ t^* \tau - \frac{b(\theta \tau)}{\delta_{\rho^*} \theta} \right\} = \left(\frac{b(\theta \cdot)}{\theta \delta_{\rho^*}} \right)^*(t^*),$$

hence, since $b^*(t^*) = \infty$ for $t^* < 0$,

$$\frac{-b^*(\delta_{\rho^*} t^*)}{\delta_{\rho^*} t^*} \leq \theta \frac{-b^*(t^*)}{t^*} \quad \forall t^* \in \mathbb{R} \quad \Leftrightarrow \quad b(t) \geq \frac{b(\theta t)}{\theta \delta_{\rho^*}} \quad \forall t \in \mathbb{R},$$

which amounts to the condition in the statement. Under such condition, then, $\pi_{k+1} \leq \theta \pi_k$ holds for every k , leading to $\pi_k \leq \pi_0 \theta^k = \frac{-b^*(\rho_0^*)}{\rho_0^*} \theta^k$ as claimed. $\square$

Though it would be tempting to seek barriers for which $(\pi_k)_{k \in \mathbb{N}}$ as in the proof vanishes at any desired rate, it can be easily verified that no choice of b or δ_{ρ^*} can result in $(\pi_k)_{k \in \mathbb{N}}$ converging any faster than linearly. In fact,

$$\pi_{k+1} = \frac{-b^*(\rho_0^* \delta_{\rho^*}^{k+1})}{\rho_0^* \delta_{\rho^*}^{k+1}} > \frac{-b^*(\rho_0^* \delta_{\rho^*}^k)}{\rho_0^* \delta_{\rho^*}^k} = \frac{1}{\delta_{\rho^*}} \pi_k,$$

where the inequality follows from monotonicity of $-b^*$, cf. [Lemma A.1\(ii\)](#). This shows that a linear decrease by a factor $\delta_{\rho^*}^{-1}$ is the fastest worst-case rate this lemma can guarantee, and that this can only happen in the limit. [Lemma 4.10](#) nevertheless identifies a property that allows us to judge the fitness of a barrier b within the framework of [Algorithm 1](#). As we will see in [Section 4.3](#), this will be particularly evident in the convex case, for it can be guaranteed that, under assumptions, α_k eventually does remain constant, so that employing a barrier that complies with this requirement is a guarantee that eventually the infeasibility p_k of the iterates generated by [Algorithm 1](#) vanishes at R-linear rate. This motivates the following definition.

Definition 4.11 (behavior profiles of b). We say that a barrier b complying with [Assumption 2](#) is *asymptotically well behaved* if

$$\forall \theta \in (0, 1) \quad \kappa_b(\theta) := \limsup_{t \rightarrow 0^-} \frac{b(\theta t)}{\theta b(t)} < \infty \quad \text{and} \quad \lim_{\theta \rightarrow 1^-} \kappa_b(\theta) = 1.$$

If this condition can be strengthened to

$$\forall \theta \in (0, 1) \quad \kappa_b^{\max}(\theta) := \sup_{t < 0} \frac{b(\theta t)}{\theta b(t)} < \infty \quad \text{and} \quad \lim_{\theta \rightarrow 1^-} \kappa_b^{\max}(\theta) = 1,$$

then we say that b is *well behaved* (not merely asymptotically). We call the functions $\kappa_b^{\max}, \kappa_b : (0, 1) \rightarrow (1, \infty)$ the *behavior profile* and the *asymptotic behavior profile* of b , respectively.

In penalty-type methods, the update of a penalty parameter is typically decided based on the violation of the corresponding constraints. Under the assumption that the barrier b is (asymptotically) well behaved, [Lemma 4.10](#) demonstrates that in [Algorithm 1](#) (eventually) the condition $\alpha_{k+1} = \alpha_k$ furnishes a guarantee of linear decrease of the infeasibility. Insisting on continuity of κ_b and $\kappa_b^{\max}$ at $\theta = 1$ in [Definition 4.11](#) is a minor technicality ensuring that, regardless of the value of $\delta_{\mu} \in (0, 1)$ and $\delta_{\alpha} > 1$, for any (asymptotically) well behaved barrier there always exists $\theta \in (0, 1)$ such that $b(\theta t) \leq \theta \delta_{\rho^*} b(t)$ holds for every $t < 0$ (close enough to zero) as required in [Lemma 4.10](#). The result can thus be restated as follows.

Corollary 4.12. *Additionally to [Assumptions 1](#) and [2](#), suppose that the barrier b is (asymptotically) well behaved. Then, there exists $\theta \in (0, 1)$ such that the iterates of [Algorithm 1](#) satisfy (4.6) for all $k \in \mathbb{N}$ (large enough).*

$b(t)$ (for $t < 0$)	$\kappa_b(\theta)$	$\kappa_b^{\max}(\theta)$	
$\frac{1}{p}(-t)^{-p}$	$(\frac{1}{\theta})^{1+p}$	$(\frac{1}{\theta})^{1+p}$	$(p > 0)$
$\ln(1 - \frac{1}{t})$	$\frac{1}{\theta}$	$(\frac{1}{\theta})^2$	
$-\ln(-t)$	$\frac{1}{\theta}$	∞	
$\exp(-\frac{1}{t})$	∞	∞	

Table 2: Examples of barriers and their behavior profiles κ_b . A low κ_b is symptomatic of good aptitude of b as barrier within [Algorithm 1](#). Geometrically, it indicates that b well approximates the nonsmooth indicator $\delta_{\mathbb{R}_-}$. Functions like $\exp(-\frac{1}{t})$ growing too fast are unsuited, whereas logarithmic barriers attain an optimal asymptotic behavior profile $\kappa_b(\theta) = \frac{1}{\theta}$. The log-like barrier $b(t) = \ln(1 - \frac{1}{t})$ is well-behaved, whereas the logarithmic barrier $b(t) = -\ln(-t)$ is so only asymptotically.

When it comes to comparing different barriers, lower values of κ_b are clearly preferable. Notice that both $\kappa_b^{\max}$ and κ_b are scaling invariant:

$$\kappa_{\beta b} = \kappa_{b(\beta \cdot)} = \kappa_b \quad \text{and} \quad \kappa_{\beta b}^{\max} = \kappa_{b(\beta \cdot)}^{\max} = \kappa_b^{\max} \quad \forall \beta > 0.$$

Moreover, since

$$\kappa_b(\theta) \geq \frac{1}{\theta} \quad \forall \theta \in (0, 1)$$

(owing to monotonicity of b and the fact that consequently $b(\theta t) \geq b(t)$ for $t < 0$), barriers attaining $\kappa_b(\theta) = \frac{1}{\theta}$ can be considered asymptotically optimal. [Table 2](#) shows that logarithmic barriers can attain such lower bound.

4.3 THE CONVEX CASE

In this section we investigate the behavior of [Algorithm 1](#) when applied to convex problems. In particular, we detail an asymptotic analysis in which the termination tolerances are set to zero, so that the algorithm (may) run indefinitely. We demonstrate that under standard assumptions the iterates subsequentially converge to (global) solutions, and that the L^1 penalty parameter α is eventually never updated.

Theorem 4.13. *Additionally to [Assumptions 1](#) and [2](#), suppose that $\inf b > 0$ and that problem (P) is convex, namely that q and c_i , $i = 1, \dots, m$, are convex functions and that c_{eq} is affine. If there exists an optimal KKT-triplet $(x^*, y^*, y_{\text{eq}}^*)$ for (P), then the following hold for the iterates generated by [Algorithm 1](#) with $\epsilon_p = \epsilon_d = 0$:*

- (i) *Any accumulation point of the sequence $(x^k)_{k \in \mathbb{N}}$ is a solution of (P).*
- (ii) *If, additionally, $(x^k)_{k \in \mathbb{N}}$ remains bounded (as is the case when $\text{dom } q$ is bounded), α_k is eventually never updated.*
- (iii) *Further assuming that the barrier b is asymptotically well behaved, so that there exists $\theta \in (0, 1)$ such that $b(\theta t) \leq \theta \min \{\delta_\mu^{-1}, \delta_\alpha\} b(t)$ for every $t < 0$ close enough to 0, then the feasibility violation eventually vanishes with rate $p_k \leq 2(m + m_{\text{eq}}) \frac{-b^*(\alpha_0/\mu_0)}{\alpha_0/\mu_0} \theta^k$.*

Proof. It follows from [Lemma 3.4\(ii\)](#) that x^* solves (P_α) for $\alpha := \max \{\|y^*\|_\infty, \|y_{\text{eq}}^*\|_\infty\}$. For every k , there exists

$$(4.7) \quad \eta^k \in \partial q_{\alpha_k, \mu_k}(x^k) \quad \text{with} \quad \|\eta^k\| \leq \varepsilon_k,$$

where for $\alpha, \mu > 0$ we let

$$(4.8) \quad q_{\alpha,\mu} := q + \mu \Psi_{\alpha/\mu} \circ \mathbf{c} + \mu \Psi_{\alpha/\mu}^{\text{eq}} \circ \mathbf{c}_{\text{eq}} \geq q + \alpha \|\mathbf{c}(\cdot)\|_1 + \alpha \|\mathbf{c}_{\text{eq}}(\cdot)\|_1.$$

Function $q_{\alpha,\mu}$ is convex, because so are $\Psi_{\alpha/\mu} \circ \mathbf{c}$ and $\Psi_{\alpha/\mu}^{\text{eq}} \circ \mathbf{c}_{\text{eq}}$ by [Theorems 3.5](#) and [3.6](#), and satisfies the inequality as in (4.8) owing to [Theorem 3.7](#), having assumed that $b > 0$. Notice further that

$$\begin{aligned} q_{\alpha,\mu}(\mathbf{x}^\star) &= q(\mathbf{x}^\star) + \mu \Psi_{\alpha/\mu}(\mathbf{c}(\mathbf{x}^\star)) + \mu \Psi_{\alpha/\mu}^{\text{eq}}(0) \\ &\leq q(\mathbf{x}^\star) + \mu \Psi_{\alpha/\mu}(0) + \mu \Psi_{\alpha/\mu}^{\text{eq}}(0) \\ &= q(\mathbf{x}^\star) + m\mu\psi_{\alpha/\mu}(0) + m_{\text{eq}}\mu\psi_{\alpha/\mu}^{\text{eq}}(0) \\ &= q(\mathbf{x}^\star) + m\mu\psi_{\alpha/\mu}(0) + 2m_{\text{eq}}\mu\psi_{\alpha/2\mu}(0) \\ &\stackrel{\text{(by Lemma A.1(iv))}}{\leq} q(\mathbf{x}^\star) + (m + m_{\text{eq}})\mu\psi_{\alpha/\mu}(0) \\ (4.9) \quad &\stackrel{\text{(by (3.7))}}{=} q(\mathbf{x}^\star) + (m + m_{\text{eq}})\alpha \frac{-b^*(\alpha/\mu)}{\alpha/\mu}, \end{aligned}$$

where the first inequality follows from (elementwise) monotonicity of $\Psi_{\alpha/\mu}$, and the third identity uses the fact that $\psi_{\rho^*}^{\text{eq}}(0) = 2\psi_{\rho^*/2}(0)$ which is apparent from (3.3) and (3.4). Next observe that, since $\mathbf{x}^\star$ solves (P_α) and is feasible, one has

$$\begin{aligned} q(\mathbf{x}^\star) &= q(\mathbf{x}^\star) + \alpha \|\mathbf{c}(\mathbf{x}^\star)\|_1 + \alpha \|\mathbf{c}_{\text{eq}}(\mathbf{x}^\star)\|_1 \\ &\leq q(\mathbf{x}^k) + (\alpha_k - (\alpha_k - \alpha)) \underbrace{(\|\mathbf{c}(\mathbf{x}^k)\|_1 + \|\mathbf{c}_{\text{eq}}(\mathbf{x}^k)\|_1)}_{=: \tilde{p}_k} \\ (4.10) \quad &\stackrel{\text{(by (4.8))}}{\leq} q_{\alpha_k,\mu_k}(\mathbf{x}^k) - (\alpha_k - \alpha)\tilde{p}_k \\ &\stackrel{\text{(by (4.7))}}{\leq} q_{\alpha_k,\mu_k}(\mathbf{x}^\star) + \langle \boldsymbol{\eta}^k, \mathbf{x}^\star - \mathbf{x}^k \rangle - (\alpha_k - \alpha)\tilde{p}_k \\ &\stackrel{\text{(by (4.9))}}{\leq} q(\mathbf{x}^\star) + (m + m_{\text{eq}})\alpha_k \frac{-b^*(\alpha_k/\mu_k)}{\alpha_k/\mu_k} + \varepsilon_k \|\mathbf{x}^\star - \mathbf{x}^k\| - (\alpha_k - \alpha)\tilde{p}_k. \end{aligned}$$

We next prove the assertions one by one.

◇ [4.13\(i\)](#) If $(\alpha_k)_{k \in \mathbb{N}}$ is asymptotically constant, then according to the update rule in [Algorithm 1](#) one has that $p_k \leq 2(m + m_{\text{eq}}) \frac{-b^*(\alpha_k/\mu_k)}{\alpha_k/\mu_k}$ eventually always holds, and thus vanishes as $k \rightarrow \infty$, see [Lemma A.1\(iv\)](#). Any limit point $\bar{\mathbf{x}}$ of $(\mathbf{x}^k)_{k \in \mathbb{N}}$ is thus feasible, and furthermore belongs to $\text{dom } q$ since $q(\mathbf{x}^k)$ remains bounded as is evident from the inequalities in (4.10). In this case, we conclude from [Corollary 4.4](#) that $\mathbf{x}^\star$ is A-KKT-optimal, and thus optimal because of convexity.

Suppose instead that $(\alpha_k)_{k \in \mathbb{N}} \nearrow \infty$, and by removing early iterates let us assume without loss of generality that $\alpha_k > \alpha$ holds for any k . Denoting $\rho_k^* := \alpha_k/\mu_k \nearrow \infty$, the inequality in (4.10) yields that

$$(4.11) \quad p_k \leq \tilde{p}_k \leq \frac{\alpha_k}{\alpha_k - \alpha} \left((m + m_{\text{eq}}) \frac{\overbrace{-b^*(\rho_k^*)}^{\rightarrow 0}}{\rho_k^*} + \frac{\overbrace{\varepsilon_k}^{\rightarrow 0}}{\alpha_k} \|\mathbf{x}^\star - \mathbf{x}^k\| \right),$$

where the fact that $\frac{-b^*(\rho_k^*)}{\rho_k^*} \rightarrow 0$ follows from [Lemma A.1\(iv\)](#). Along any convergent subsequence, it is clear that $p_k \rightarrow 0$ and that $q(\mathbf{x}^k)$ remains bounded, and again we conclude that the limit point is optimal for (P).

◇ [4.13\(ii\)](#) To arrive to a contradiction, suppose that $\alpha_k \nearrow \infty$, and again assume without loss of generality that $\alpha_k > \alpha$ holds for all k . It then follows from (4.11) that

$$\frac{p_k}{2(m + m_{\text{eq}}) \frac{-b^*(\rho_k^*)}{\rho_k^*}} \leq \frac{1}{2} \frac{\alpha_k}{\alpha_k - \alpha} \left(1 + \frac{\varepsilon_k}{\alpha_k} \frac{\rho_k^*}{-b^*(\rho_k^*)} R \right) \rightarrow \frac{1}{2} \quad \text{as } k \rightarrow \infty,$$

and in particular p_k is eventually always smaller than $2(m + m_{\text{eq}}) \frac{-b^*(\rho_k^*)}{\rho_k^*}$. According to the update rule of α_k in [Algorithm 1](#), this means that α_k is eventually constant, which is a contradiction.

◇ [4.13\(iii\)](#) Follows from [Lemma 4.10](#). □

The PIPA algorithm of [5], specialized to the convex setting, generates primal-dual sequences that remain bounded, offering a theoretical advantage over [Theorem 4.13](#), where boundedness is assumed. The difficulty in recovering the results of [5] appears to stem from the constraint relaxation occurring in (P_α) , which on the other hand is the key for [Marge](#) to handle equality constraints and infeasible starting points.

5 NUMERICAL EXPERIMENTS

This section is dedicated to experimental results and comparisons with other numerical approaches for constrained structured optimization. The modular structure of the proposed framework allows us to combine [Marge](#) with a variety of penalty-barrier envelopes and inner solvers (insofar as they provide suitable guarantees). The performance and behavior of [Marge](#) is illustrated in different variants, considering three barrier functions, namely $b(t) = -\frac{1}{t}$, $b(t) = \ln(1 - \frac{1}{t})$, and $b(t) = -\ln(-t)$ (all extended as ∞ on $\mathbb{R}_+$) denominated *inverse*, *log-like*, and *log*, respectively. The numerical comparison will comprise two data science tasks and two *ad hoc* illustrative problems. These numerical tests will also highlight the influence of the barrier function on the performance of [Marge](#), supporting the quality assessment of [Section 4.2](#).

The performance of [Marge](#) is compared against those of IPprox [13, Alg. 1], ALPS [10, Alg. 4.1], and the well-known Ipopt [32]. IPprox builds upon a pure interior point scheme and solves the barrier subproblems with a tailored adaptive proximal-gradient algorithm, extending the strategy of PIPA [5] to the nonconvex setting. ALPS belongs to the family of augmented Lagrangian algorithms and does not require a custom subsolver—suitable subsolvers for [Marge](#) can be applied within ALPS and viceversa. The closely related solvers OpEn [27] and Alpaqa [22] also build on the augmented Lagrangian framework and provide easy-to-use high-performance implementations; however, these focus on nonlinear programming problems and cannot support generic prox-friendly cost functions as ALPS does. Finally, Ipopt is a software package for large-scale sparse nonlinear optimization; it implements an interior point line search filter method. Ipopt can address problems of the form (P) where the input functions, particularly q , should be at least continuously differentiable.

[Marge](#)'s minimal assumptions on the subsolver leave us much freedom in its selection. To handle potential (structured) nonsmoothness of the cost q , we chose PANOC⁺ [12, 29], a proximal-gradient-based solver that can exploit directions of quasi-Newton type while ensuring convergence with a backtracking linesearch, see also [30, §5.1]. We also tested the nonmonotone proximal-gradient method with adaptive stepsizes of [9], but in our comparisons we only retained the results with PANOC⁺, set as default solver in our implementation, as it consistently demonstrated better performance.

Patterning the simulations of [13, §5.2], we first examine the nonnegative PCA problem in [Section 5.2](#) to evaluate [Marge](#) in several variants and compare it against the other solvers. Then, [Section 5.3](#) focuses on a low-rank matrix completion task, a fully nonconvex problem with thousands of variables and constraints, contrasting [Marge](#) to ALPS and Ipopt. Finally, the exact penalty behavior and the ability to handle hidden equalities are illustrated and discussed in [Sections 5.4](#) and [5.5](#), respectively. Additional observations that support the analysis of barriers' properties are included in [Appendix B](#).

The source code of our implementation has been archived for reproducibility of the numerical results presented in this paper; it can be found on Zenodo at DOI: [10.5281/zenodo.11098283](https://doi.org/10.5281/zenodo.11098283).

5.1 IMPLEMENTATION DETAILS

We describe here details pertinent to our implementation *Marge* of Algorithm 1, such as the initialization and update of algorithmic parameters. These numerical features tend to improve the practical performances, without compromising the convergence guarantees. IPprox is available from [13] and adopted as is, whereas ALPS is a slight modification of the code from [10] to be comparable with *Marge*, as detailed below.

Our implementation of *Marge* accepts problems formulated in the form

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \quad f(\mathbf{x}) + g(\mathbf{x}) \quad \text{subject to} \quad \mathbf{l} \leq \mathbf{c}(\mathbf{x}) \leq \mathbf{u},$$

with bounds defined by extended-real-valued vectors $\mathbf{l}$ and $\mathbf{u}$. In a preprocessing phase, these vectors are parsed to reformulate the problem data in the format (P) and to instantiate the appropriate penalty-barrier functions to treat inequality and equality constraints as described above.

Default parameters for *Marge* are $\mu_0 = 1$ and $\delta_\epsilon = \delta_\mu = 1/4$ as in IPprox, $\delta_\alpha = 2$ as in ALPS, and $\alpha_0 = 1$. The initial tolerance ϵ_0 for *Marge* (and ALPS) is chosen adaptively, based on the user-provided starting point $\mathbf{x}^0$ and penalty-barrier parameters. Following the mechanism implemented in IPprox, we set $\epsilon_0 = \max \{\epsilon_d, \epsilon_{\min}, \min \{\kappa_\epsilon \eta_0, \epsilon_{\max}\}\}$, where $\kappa_\epsilon \in (0, 1)$ and $\epsilon_{\max} \geq \epsilon_{\min} \geq 0$ are user-specified parameters (default $\kappa_\epsilon = 10^{-2}$, $\epsilon_{\max} = 1$, $\epsilon_{\min} = 10^{-6}$) and η_0 is an estimate of the initial stationarity measure, as evaluated by (executing one iteration of) the inner solver invoked at $(\mathbf{x}^0, \alpha_0, \mu_0)$. The barrier parameter is updated according to the rule presented in Remark 4.2. For simplicity, no infeasibility detection mechanism nor artificial bounds on penalty and barrier parameters have been included.

We run ALPS with the same settings as in [11, 10] apart from the following features to match *Marge*: the initial penalty parameter is fixed ($\alpha_0 = 1$) and not adaptive, the tolerance reduction factor is set to $\delta_\epsilon = 1/4$ instead of $\delta_\epsilon = 1/10$, and the initial inner tolerance is selected adaptively. We always initialize ALPS with dual estimate $\mathbf{y}^0 = \mathbf{0}$. The PANOC⁺ subsolver is considered with its default tuning, namely L-BFGS directions (memory 5) and monotone linesearch strategy as in [12]. We use Ipopt version 3.13.3, called via mexIPOPT [2], with linear solver MUMPS. The only non-default optional values we set are the limited-memory Hessian approximation and the desired termination tolerance `tol`, as specified below. For consistency with Ipopt's termination condition [32, §2.1], *Marge* requires approximate stationarity measured in the L^∞ -norm, deviating from Definition 2.5; this applies also to ALPS and IPprox.

For P the set of problems and S the set of solvers, let $t_{s,p}$ denote the user-defined metric for the computational effort required by solver $s \in S$ to solve instance $p \in P$ (lower is better). We will monitor the (total) number of gradient evaluations, so that the computational overhead triggered by backtracking is fairly accounted for, the number of (outer) iterations, and the wall-clock runtime. Then, we display *data profiles* to graphically summarize our numerical results and compare different solvers. A data profile is the graph of the cumulative distribution function $f_s : [0, \infty) \rightarrow [0, 1]$ of the evaluation metric, namely $f_s(t) := |\{p \in P \mid t_{s,p} \leq t\}|/|P|$. As such, each data profile reports the fraction of problems $f_s(t)$ that can be solved (for a given tolerance ϵ) by solver s with a computational budget t , and therefore it is independent of the other solvers [21].

5.2 NONNEGATIVE PCA

Principal component analysis (PCA) aims at estimating the direction of maximal variability of a high-dimensional dataset. Imposing nonnegativity of entries as prior knowledge, we address PCA restricted to the positive orthant:

$$(5.1) \quad \underset{\mathbf{x} \in \mathbb{R}^n}{\text{maximize}} \quad \mathbf{x}^\top \mathbf{Z} \mathbf{x} \quad \text{subject to} \quad \|\mathbf{x}\| = 1, \mathbf{x} \geq \mathbf{0}.$$

This task falls within the scope of (P), with $f(\mathbf{x}) := -\mathbf{x}^\top \mathbf{Z} \mathbf{x}$, $g(\mathbf{x}) := \delta_{\|\cdot\|=1}(\mathbf{x})$, and $\mathbf{c}(\mathbf{x}) = -\mathbf{x}$, and has been considered in [13] for validating IPprox and tuning its hyperparameters. Here, we start by showing that all *Marge*'s variants significantly outperform IPprox on the same instances the latter was fine-tuned with (using the inverse barrier); cf. [13, §5.2]. Then, we investigate how the computational effort of *Marge*'s default setup (that is, with log-like barrier and subsolver PANOC⁺) scales with the problem size and the accuracy requirement.

Setup We generate synthetic problem data as in [13, §5.2]. For a problem size $n \in \mathbb{N}$, let $\mathbf{Z} = \sqrt{\sigma_n} \mathbf{z} \mathbf{z}^\top + \mathbf{N} \in \mathbb{R}^{n \times n}$, where $\mathbf{N} \in \mathbb{R}^{n \times n}$ is a random noise matrix, $\mathbf{z} \in \mathbb{R}^n$ is the true (random) principal direction, and $\sigma_n > 0$ is the signal-to-noise ratio. We consider some dimensions n and, for each dimension, the set of problems parametrized by $\sigma_n \in \{0.05, 0.1, 0.25, 0.5, 1.0\}$ and $\sigma_s \in \{0.1, 0.3, 0.7, 0.9\}$, which control the noise and sparsity level, respectively. There are 5 choices for σ_n , 4 for σ_s , and, for each set of parameters, 2 instances are generated with different problem data $\mathbf{Z}$ and starting point $\mathbf{x}^0$. Overall, each solver-settings pair is invoked on 40 different instances for each dimension n .

A strictly feasible starting point $\mathbf{x}^0$ is generated by sampling a uniform distribution over $[0, 3]^n$ and projecting onto $\text{dom } g = \{\mathbf{x} \in \mathbb{R}^n \mid \|\mathbf{x}\| = 1\}$. This property is necessary for IPprox but not for the other solvers. We will test *Marge*, Ipopt and ALPS also with arbitrary initialization, in which case $\mathbf{x}^0$ is generated by sampling a uniform distribution over $[-3, 3]^n$ and then projecting onto $\text{dom } g$. Furthermore, since IPprox requires a nonnegative-valued barrier function [13, Ass. 2], its variant with log barrier is not included in these tests. Finally, Ipopt tackles (5.1) encoded with the nonlinear equality constraint $\sum_{i=1}^n x_i^2 = 1$ and simple lower bounds on $\mathbf{x}$.

Barriers and subsolvers Algorithm 1 is controlled by, and its performance depends on, several algorithmic hyperparameters, such as the (sequences of) barrier and penalty parameters, the choice of barrier function b , and the subsolver adopted at Step 1.1. We now focus on the effect of the barriers specified in Table 1, for different levels of accuracy requirements, testing all solver variants on *tiny* and *small* instances with problem dimensions $n \in \{10, 15, 20\}$ and $n \in \{100, 150, 200\}$, respectively. Moreover, because of the excessive runtime to perform all simulations for IPprox, we exclude it altogether for the high accuracy tests, for which we consider starting points $\mathbf{x}^0$ that are not necessarily (strictly) feasible.

The results are graphically summarized in Figure 4 with data profiles relative to runtime. All solvers returned successfully within the desired primal-dual tolerances. Across all accuracy levels, *Marge* inverse operates consistently better than the other variants of *Marge*, all of which outperform IPprox. In particular, the overall effort (in terms of runtime) required by the inverse barrier (in both *Marge* and IPprox) is less than with the log-like and log barriers, whose profiles almost overlap. With increasing accuracy it becomes even more efficient to adopt *Marge* over IPprox, which performs gradually more poorly. The slow tail convergence typical of first-order schemes badly affects the scalability of IPprox, whereas the adoption of a quasi-Newton scheme within the subsolver PANOC⁺ of *Marge* and ALPS leads to fast local convergence in practice. Despite the “optimality” of logarithmic barriers ascertained in Section 4.2, the overall computational effort (measured in terms of gradient evaluations or runtime) also depends on the subsolver's efficiency in solving the subproblems.

For tiny instances or low accuracy, the overall runtime of *Marge* tends to be comparable with that of Ipopt. For high accuracy requirements and manageable problem sizes, Ipopt can profit from the more computation-intensive iterations, whereas the cheaper first-order steps of *Marge*'s subsolver pay off with lower accuracies. Adopting the same first-order subsolver, ALPS appears to finish ahead of *Marge* and Ipopt. We can attribute this advantage to the simple structure of the explicit constraints $\mathbf{x} \geq \mathbf{0}$ in (5.1), since the superiority of ALPS vanishes with more difficult constraints, as witnessed by the following Section 5.3.

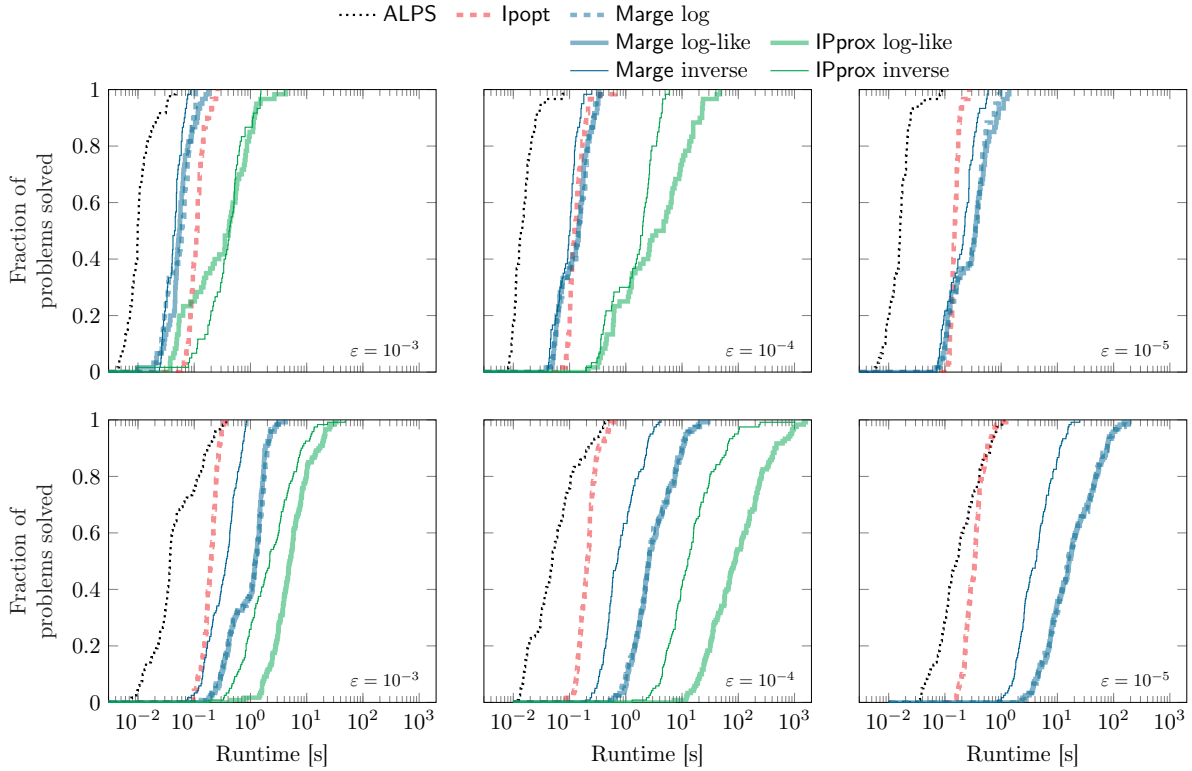


Figure 4: Nonnegative PCA problem (5.1): comparison of solvers on small (bottom) and large instances (top) with low, medium and high accuracy $\epsilon_p = \epsilon_d = \epsilon \in \{10^{-3}, 10^{-4}, 10^{-5}\}$ (left to right) using data profiles relative to runtime. Results for IPprox (green solid lines) are not included for the high accuracy tests due to excessive runtime; in these simulations, the other solvers are initialized with a possibly infeasible guess. In all plots, *Marge*'s profiles for log-like (blue thick solid lines) and log (blue thick dashed lines) barriers almost coincide.

Problem size and accuracy To investigate scalability and influence of accuracy requirements, we consider larger instances of (5.1) with dimensions $n \in \{10, \lceil 10^{1.5} \rceil, 10^2, \lceil 10^{2.5} \rceil, 10^3\}$ and tolerances $\epsilon_p = \epsilon_d = \epsilon \in \{10^{-3}, 10^{-4}, 10^{-5}\}$, and invoke the default solver *Marge* (with PANOC⁺ subsolver and log-like barrier) without time limit. For each of these tolerance parameters, we generate 2 instances (as described above) for each set of parameters, leading to a total of 200 problem instances to be solved for each accuracy level.

All instances are solved up to the desired primal-dual tolerances. The influence of problem size and tolerance is depicted in Figure 5, which displays for each pair (n, ϵ) the number of gradient evaluations and runtime with a jitter plot (for a better visualization of the distribution of numerical values over categories). The empirical cumulative distribution function and the associated median value are also indicated. This chart visualizes how problem size and accuracy requirement affect the solution process, and reveals the stark effect of both n and ϵ .

For low accuracy, *Marge* scales relatively well with the problem size, whereas larger problems become prohibitive for high accuracy. This behavior is typical of first-order methods, due to their slow tail convergence, and we take it as a motivation for investigating the interaction between subproblems and subsolvers in future works. Nevertheless, these experiments (and those forthcoming) demonstrate *Marge*'s capability to handle thousands of variables and constraints in a fully nonconvex optimization landscape. These results witness a tremendous improvement over IPprox, not only in the practical performance but also in the flexibility of use, as *Marge* can be initialized at infeasible points and can

take advantage of general-purpose fast subsolvers, such as PANOC⁺.

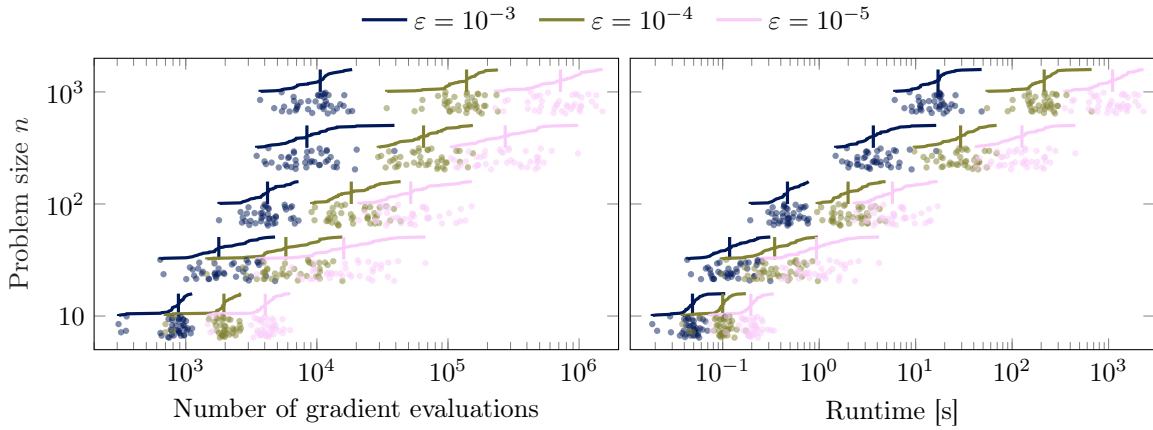


Figure 5: Nonnegative PCA problem (5.1) with *Marge* log-like: comparison for increasing accuracy requirements (decreasing tolerances $\epsilon_p = \epsilon_d = \epsilon$) and problem sizes n , relative to gradient evaluations (left) and runtime (right). Combination of jitter plot (dots) and empirical cumulative distribution function (solid line) with median value (vertical segment).

5.3 LOW-RANK MATRIX COMPLETION

Given an incomplete matrix of (uncertain) ratings Y , a common task is to find a complete ratings matrix X that is a parsimonious representation of Y , in the sense of low-rank, and such that $Y \approx X$ for the entries available [19]. Let $\#_u$ and $\#_m$ denote the number of users and items, respectively, and let the rating $Y_{i,j}$ by the i th user for the j th item range on a scale defined by constants $Y_{\min}$ and $Y_{\max}$. Let Ω represent the index set of observed ratings, and $|\Omega|$ the cardinality of Ω . The ratings matrix Y could be very large and often most of the entries are unobserved, since a given user will only rate a small subset of items. Low-rankness of X can be enforced by construction, with the Ansatz $X \equiv UV^\top$, as in dictionary learning. In practice, for some prescribed embedding dimension $\#_a$, we seek a user embedding matrix $U \in \mathbb{R}^{\#_u \times \#_a}$ and an item embedding matrix $V \in \mathbb{R}^{\#_m \times \#_a}$. Each row $U_{i,:}$ of U is a $\#_a$ -dimensional vector representing user i , while each row $V_{j,:}$ of V is a $\#_a$ -dimensional vector representing item j . We address the joint completion and factorization of the ratings matrix Y , encoded in the following form:

$$\begin{aligned}
 (5.2) \quad & \underset{\substack{U \in \mathbb{R}^{\#_u \times \#_a}, \\ V \in \mathbb{R}^{\#_m \times \#_a}}}{\text{minimize}} && \frac{1}{|\Omega|} \sum_{(i,j) \in \Omega} (\langle U_{i,:}, V_{j,:} \rangle - Y_{i,j})^2 + \frac{\lambda}{\#_m} \sum_{j=1}^{\#_m} \|V_{j,:}\|_0 \\
 & \text{subject to} && \max \{Y_{\min}, Y_{i,j} - 1\} \leq \langle U_{i,:}, V_{j,:} \rangle \leq \min \{Y_{\max}, Y_{i,j} + 1\} && \forall (i,j) \in \Omega, \\
 & && Y_{\min} \leq \langle U_{i,:}, V_{j,:} \rangle \leq Y_{\max} && \forall (i,j) \notin \Omega, \\
 & && \|U_{i,:}\|_2 = 1 && \forall i \in \{1, \dots, \#_u\}.
 \end{aligned}$$

While aiming at $UV^\top \approx Y$, the model in (5.2) sets the rating range $[Y_{\min}, Y_{\max}]$ as a hard constraint for all predictions; a tighter constraint is imposed to observed ratings. Following [30, §6.2], we explicitly constrain the norm of the dictionary atoms $U_{i,:}$, without loss of generality, to reduce the number of equivalent (up to scaling) solutions; this norm specification is included as an indicator in the nonsmooth objective term g . Furthermore, we encourage sparsity of the coefficient representation $V_{j,:}$ with the $\|\cdot\|_0$ penalty, which counts the nonzero elements, scaled with a regularization parameter $\lambda \geq 0$.

Overall, this problem has $n := \#_a(\#_u + \#_m)$ decision variables and $m := 2\#_u\#_m$ inequality constraints. All terms (f , g , and c) are nonconvex, as well as the (unbounded) feasible set.

It appears nontrivial to find a strictly feasible point for (5.2), in the sense of [13, Def. 2], which is required for initializing IPprox, thus highlighting a major advantage of *Marge*. Another feature of *Marge* is the versatility of model (P), which enables an effortless handling of the nonsmooth term $\|\cdot\|_0$, in stark contrast with Ipopt. Although (5.2) can be reformulated as a nonlinear program via (1.1) using $3\#_m\#_a$ auxiliary variables, $4\#_m\#_a$ linear and $\#_m$ nonlinear additional constraints, this extended formulation hinders the scalability to large datasets. For comparison, we will consider instances of (5.2) with and without the sparsity-promoting term $\|\cdot\|_0$ in the objective.

Setup We consider the *MovieLens 100k* dataset,⁴ which contains 1000023 ratings for 3706 unique movies (the dataset contains some repetitions in movie ratings and we have ignored them); these recommendations were made by 6040 users on a discrete rating scale from $Y_{\min} = 1$ to $Y_{\max} = 5$. If we construct a matrix of movie ratings by the users, then it is a sparse unstructured matrix with only 4.47% of the total entries available.

We compare *Marge* to ALPS and Ipopt, testing their scalability with instances of increasing size. The set of problems consists of *small* and *large* instances, as well as instances *with* and *without* the $\|\cdot\|_0$ term in (5.2). We set the regularization parameter $\lambda = 10^{-2}$, randomly generate 4 starting points for each problem instance, and invoke each solver with the primal-dual tolerances $\epsilon_p = \epsilon_d = 10^{-3}$ and without time limit. First, we fix the number of atoms to $\#_a = 5$ and consider the small instances corresponding to subsets of $\#_u \in \{3, 4, \dots, 7\}$ users (always starting from the first one), for a total of 20 calls to each solver variant. For these problem instances the sizes range $n \in [1790, 3450]$ and $m \in [2130, 9562]$. Then, we fix the number of atoms to $\#_a = 10$ and consider the large instances of (5.2) corresponding to subsets of $\#_u \in \{11, 12, \dots, 20\}$ users, for a total of 40 calls to each solver variant. For these problem instances the sizes range $n \in [7630, 9930]$ and $m \in [16544, 38920]$; Ipopt is not tested on these large instances with $\|\cdot\|_0$ because of excessive runtime.

Results A summary of the numerical results is depicted in Figure 6 as data profiles, while median values are reported in Table 3. Except Ipopt on instances with $\|\cdot\|_0$, all solver variants were able to find a solution up to the specified primal-dual tolerance. For small instances, there is not a clear winner among *Marge*'s variants and ALPS, but Ipopt appears to fall behind. Ipopt returned unsuccessfully on 30% of calls for small instances with $\|\cdot\|_0$, requiring more than ten times the runtime of other solvers; for this reason it was not tested on the large instances. Although Ipopt requires significantly fewer gradient evaluations, its overall runtime soars, even on the instances without $\|\cdot\|_0$. Indeed, the metrics reported in Table 3 indicate that, to solve at least half of all large instances without $\|\cdot\|_0$, Ipopt's runtime is six times longer than *Marge*'s with log-like barrier. On the large instances with $\|\cdot\|_0$, *Marge* with log-like and log barriers consistently outperformed the variant with inverse barrier and ALPS, cutting the runtime almost in half. These results show that *Marge* can handle problems with thousands of variables and constraints, with a purely primal approach, and be as effective as well-established general-purpose solvers.

Among all instances of (5.2) considered so far, we observed that the penalty parameter α_k was rarely, if ever, updated by any variants of *Marge*. For illustrative purposes, we solved once again the large instance with $\#_u = 11$ from above, but starting with the much smaller penalty parameter $\alpha_0 = 10^{-4}$. The penalty behavior is displayed in Figure 7 (left panel), tracing the number of updates for α_k along the iterations for each of the 4 starting points. The solution process of each solver is consistent throughout

⁴The entire dataset is available at <https://grouplens.org/datasets/movielens/100k/>.

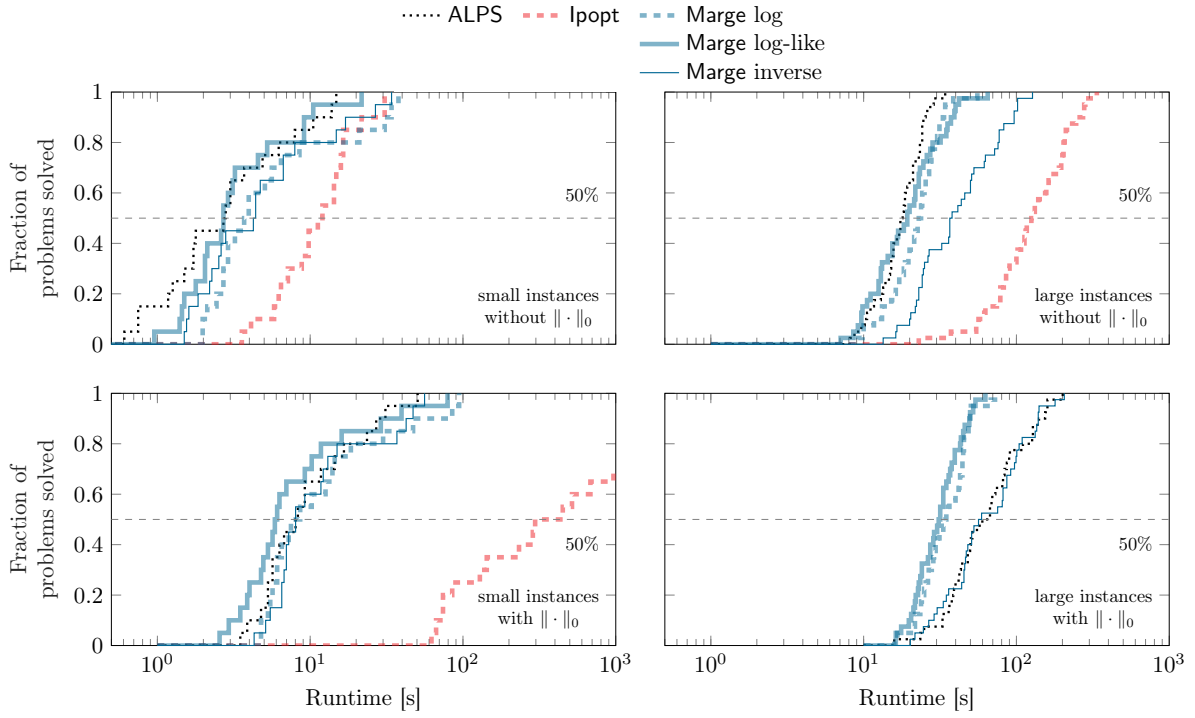


Figure 6: Matrix completion problem (5.2): comparison of different solvers and variants on small (left) and large instances (right), with (bottom) and without the $\|\cdot\|_0$ regularization (top), using data profiles relative to runtime. lpopt (red thick dashed lines) is not included on large instances with $\|\cdot\|_0$ because of excessive runtime.

all initializations, and the total number of penalty updates is also distinctive of *Marge* (except for the log variant) and ALPS; see right panel of Figure 7. The latter takes between 8 and 11 updates, whereas the former only 4 or 5 (7 updates for the log variant). Such updates takes place in ALPS whenever *local* improvement in feasibility is deemed insufficient from one iteration to the next. The same favorable behavior of *Marge* is enabled by the relaxed condition at Step 1.8 of Algorithm 1, which does not require any sufficient improvement at every iteration, but instead monitors *globally* how the constraint violation p_k vanishes. Correspondingly, only the barrier parameter μ_k is decreased in order to reduce the complementarity slackness s_k , see Lemma 4.3(iii). When active, this *exact* penalty quality prevents the barrier to yield too much ill-conditioning.

While the update of the penalty parameter α for log-like and inverse barriers follow similar profiles (though exhibiting a discrepancy coherent with Section 4.2), Figure 7 (left panel) portrays a sharper increase of the penalty parameter when adopting the logarithmic barrier. Most directly, this different behavior emerges because, starting with $(\alpha_0, \mu_0) = (10^{-4}, 1)$, $b^*(\alpha_k/\mu_k) = -1 - \ln(\alpha_k/\mu_k) > 0$ holds for some iterations until the ratio α_k/μ_k becomes large enough. In practice, during these initial iterations, Step 1.8 of Algorithm 1 effectively checks the condition $p_k > \epsilon_p$. Log-like and inverse barriers do not experience this effect since $b^* \leq 0$; see Table 1 and Lemma A.1(ii). Moreover, we do not observe this phenomenon in Figure 6 because, starting with the default values $\alpha_0 = \mu_0 = 1$, it is

$$0 > \frac{b^*(\alpha_0/\mu_0)}{\alpha_0/\mu_0} \leq \frac{b^*(\alpha_k/\mu_k)}{\alpha_k/\mu_k} \nearrow 0$$

for all barrier functions specified in Table 1; see Lemma A.1(iv). More fundamentally, we can attribute the disparity in Figure 7 to the fact that, in contrast to the log-like and inverse barriers, the logarithmic barrier is only asymptotically well-behaved; see Table 2.

Solver variant	Small instances		Large instances		
	# gradient eval.	Runtime [s]	# gradient eval.	Runtime [s]	
Marge log-like	4103	2.7	6874	18.8	without $\ \cdot\ _0$
Marge inverse	7262	4.2	15064	36.6	
Marge log	5941	3.7	8657	23.0	
ALPS	6223	2.7	11959	17.4	
lpopt	161	11.2	215	125.2	
Marge log-like	4248	5.9	7156	31.0	with $\ \cdot\ _0$
Marge inverse	6789	7.6	15239	56.5	
Marge log	5805	7.5	7712	34.4	
ALPS	6906	8.0	15239	59.8	
lpopt	1176	296.7	—	—	

Table 3: Matrix completion problem (5.2): performance comparison of different solvers and variants, reporting the computational effort needed to achieve 50% of problems solved. The reported values arise from the intersection of the data profiles in Figure 6 with the 50% horizontal line; analogously for the number of gradient evaluations. lpopt was not tested on large instances with $\|\cdot\|_0$ because of excessive runtime.

Overall, despite the fully nonconvex setting of problem (5.2), *Marge* is able to solve these instances with only moderate values for the penalty parameter α_k . Although these observations indicate that the assumptions behind Theorem 4.13(ii) could be relaxed, the penalty exactness does not always take effect, as demonstrated in the following section.

5.4 EXACT PENALTY BEHAVIOR

After observing the bounded penalty behavior of *Marge* in Section 5.3, we present now an example problem where *Marge* exhibits $\alpha_k \nearrow \infty$, hence it does *not* boil down to an exact penalty method. For this purpose it suffices to consider the two-dimensional *convex* problem

$$(5.3) \quad \underset{x \in \mathbb{R}^2}{\text{minimize}} \quad x_1 + \delta_{\mathbb{R}_+}(x_2) \quad \text{subject to} \quad x_1^2 + x_2 \leq 0,$$

whose (unique) solution is the only feasible point $x^* = (0, 0)$. Since there exists no suitable multiplier y^* , the minimizer x^* is not KKT-optimal. Hence, there is no contradiction with Theorem 4.13(ii).

We intend to solve problem (5.3) with tolerances $\epsilon_p = \epsilon_d = 10^{-5}$, initializing *Marge* and ALPS from 100 random points generated according to $x_i^0 \sim \mathcal{N}(0, \sigma_x^2)$ with large standard deviation $\sigma_x = 30$.

Results All solver variants find a primal-dual solution to (5.3), up to the specified tolerances, for all starting points. The progression of *Marge* and ALPS in terms of penalty updates are summarized in Figure 8, in analogy with Figure 7. The unbounded behavior of the penalty parameters $(\alpha_k)_{k \in \mathbb{N}}$ appears evident for all solvers. Thus, $\alpha_k \nearrow \infty$ seems necessary to drive the constraint violation p_k to zero, while the barrier parameter $\mu_k \searrow 0$ forces the complementarity slackness s_k .

Considering *Marge*'s variants, the log-like and log barriers are again almost indistinguishable and yield better results than the inverse one. ALPS performs poorly and always returns after more iterations and updates of the penalty parameter. All runs of *Marge* and ALPS terminate with α_k updated 8 and 21 times, respectively, namely increased up to $\alpha_0 \delta_\alpha^8 = 256$ and $\alpha_0 \delta_\alpha^{21} \approx 2.1 \cdot 10^6$. Then, it is clear that the performance of ALPS is badly affected by the lack of regularity in (5.3); without an appropriate dual

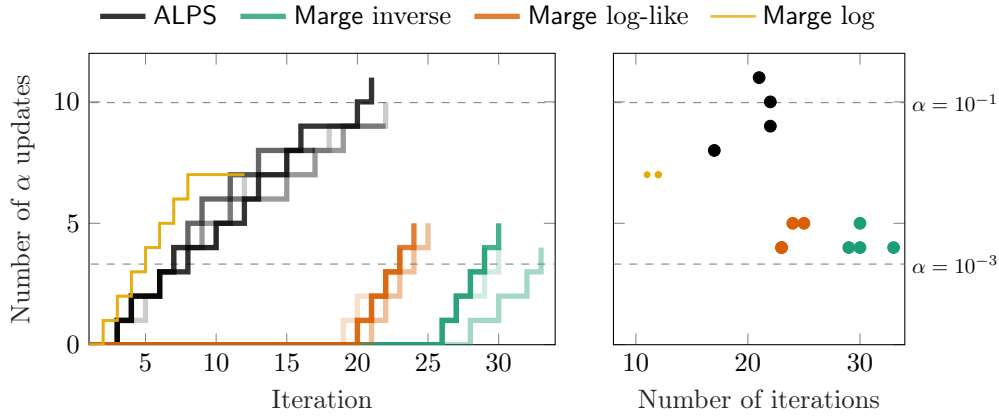


Figure 7: Matrix completion problem (5.2), smallest large instance: comparison of different solvers and variants in terms of updates for the penalty parameter α_k along the iterations (left) and at the solution (right). All solvers start with the penalty value $\alpha_0 = 10^{-4}$ for each of the 4 initializations. Fewer updates are expected to result in better-conditioned subproblems. Notice in the left panel that **Marge** log-like (orange) and inverse (green) increase α_k only after several iterations and fewer times, thanks to the relaxed criterion at **Step 1.8**; **Marge** log (yellow) and ALPS (black) update α_k from the first iterations.

estimate, ALPS essentially falls back to a quadratic penalty scheme. In contrast, all **Marge**'s variants cope well with the lack of penalty exactness and operate consistently better than ALPS in this scenario.

5.5 HANDLING EQUALITIES

Even though equality constraints can be handled explicitly, it is important that **Marge** can cope with hidden equalities too. These may appear as the result of automatic model constructions, and are often difficult to identify by inspection. Here we compare the behavior of **Marge** when the problem specification has explicit equalities against the same problem but whose constraints are described using two inequalities each. The latter approach not only increases the number of constraints, but it has also the drawback that the Mangasarian-Fromovitz constraint qualification (MFCQ) fails to hold at all feasible points. Effectively, the redundancy introduced by splitting into two inequalities undermines the practical relevance of **Lemma 3.4(ii)**; see [8, §4.1.4].

Consider quadratic programming (QP) problems of the form

$$(5.4) \quad \underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \quad \frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \langle \mathbf{q}, \mathbf{x} \rangle \quad \text{subject to} \quad \mathbf{A} \mathbf{x} = \mathbf{b}, \mathbf{x}^{\text{low}} \leq \mathbf{x} \leq \mathbf{x}^{\text{upp}}$$

with matrices $\mathbf{Q} \in \mathbb{R}^{n \times n}$, $\mathbf{A} \in \mathbb{R}^{m \times n}$ and vectors $\mathbf{q}, \mathbf{x}^{\text{low}}, \mathbf{x}^{\text{upp}} \in \mathbb{R}^n$, $\mathbf{b} \in \mathbb{R}^m$ as problem data. Problem (5.4) can be cast as (P) with cost functions $f(\mathbf{x}) := \frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \langle \mathbf{q}, \mathbf{x} \rangle$ and $g(\mathbf{x}) := \delta_{[\mathbf{x}^{\text{low}}, \mathbf{x}^{\text{upp}}]}(\mathbf{x})$, and constraint function $\mathbf{c}_{\text{eq}}(\mathbf{x}) := \mathbf{A} \mathbf{x} - \mathbf{b}$. We are interested in comparing the performance of **Marge** (in different variants) with the two problem formulations described in **Section 3.2** to deal with equalities: either by splitting into two inequalities (leading to the sum $\psi_{\rho^*}^\pm(t) := \psi_{\rho^*}(t) + \psi_{\rho^*}(-t)$ as in (3.10)) or by performing a combined marginalization (resulting in $\psi_{\rho^*}^{\text{eq}}$). Hence, for each solver's variant and problem instance, we contrast these two formulations, symbolized by **Marge** $^\pm$ and **Marge** $^{\text{eq}}$, respectively.

Setup Problem instances are generated as follows: we let either $\mathbf{Q} = \mathbf{M} \mathbf{M}^\top$ or $\mathbf{Q} = \mathbf{M} + \mathbf{M}^\top$, where the elements of $\mathbf{M} \in \mathbb{R}^{n \times n}$ are normally distributed, $M_{i,j} \sim \mathcal{N}(0, 1)$, with only 10% being nonzero. The linear part of the cost $\mathbf{q}$ is also normally distributed, i.e., $q_i \sim \mathcal{N}(0, 1)$. Simple bounds are generated

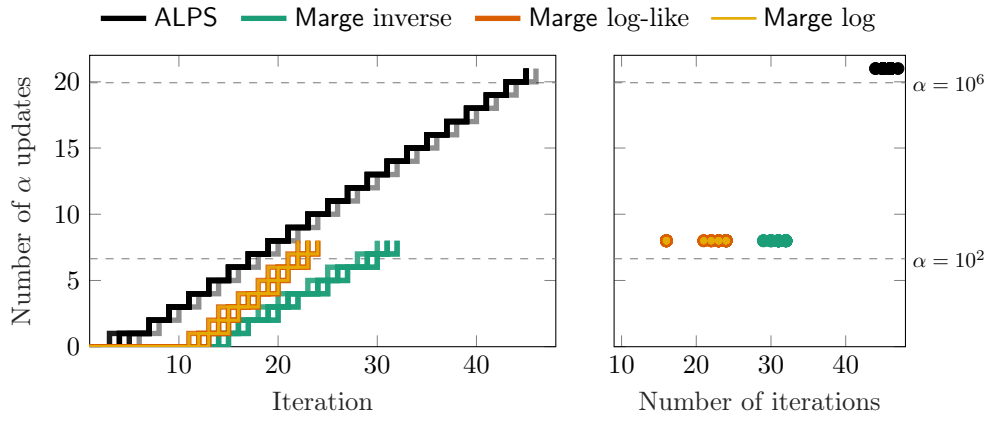


Figure 8: Convex problem without LICQ (5.3), first 10 initializations: comparison of different solvers and variants in terms of updates for the penalty parameter α_k along the iterations (left) and at the solution (right). All solvers start with the (default) penalty value $\alpha_0 = 1$. The left panel depicts the solver trajectories for each initialization, indicating how many times the penalty parameter α_k is updated during the solution process. Fewer updates are expected to result in better-conditioned subproblems. In all cases the sequence of penalty parameters $(\alpha_k)_{k \in \mathbb{N}}$ blows up, but variants of **Marge** terminate sooner and with less penalty updates than ALPS. The results for logarithmic barriers almost overlap and always take less iterations than the inverse barrier, as expected.

according to a uniform distribution, i.e., $x_i^{\text{low}} \sim -\mathcal{U}(0, 1)$ and $x_i^{\text{upp}} \sim \mathcal{U}(0, 1)$. We set the elements of $A \in \mathbb{R}^{m \times n}$ as $A_{i,j} \sim \mathcal{N}(0, 1)$ with only 10% being nonzero. To ensure that the problem is feasible, we draw an element $\hat{x} \in [x^{\text{low}}, x^{\text{upp}}]$ (as $\hat{x}_i = x_i^{\text{low}} + (x_i^{\text{upp}} - x_i^{\text{low}})a_i$ with $a_i \sim \mathcal{U}(0, 1)$) and set $b = A\hat{x}$. An initial guess is randomly generated for each problem instance, as $x_i^0 \sim \mathcal{N}(0, 1)$, and shared across all solvers and formulations.

We consider problems with $m \in \{1, 2, \dots, 20\}$ and $n = 10m$, set the tolerances $\epsilon_p = \epsilon_d = 10^{-5}$, and construct 10 instances for each size, for a total of 200 calls to each solver for each formulation.

Results Numerical results are visualized by means of pairwise (extended) performance profiles. Let $t_{s,p}^{\pm}$ and $t_{s,p}^{\text{eq}}$ denote the evaluation metric of solver $s \in S$ on a certain instance $p \in P$ with the two formulations. Then, for each solver s , the corresponding pairwise performance profile displays the cumulative distribution $\varrho_s : [0, \infty) \rightarrow [0, 1]$ of its performance ratio $\tau_{s,p}$, namely

$$\varrho_s(\tau) := \frac{|\{p \in P \mid \tau_{s,p} \leq \tau\}|}{|P|} \quad \text{where} \quad \tau_{s,p} := \frac{t_{s,p}^{\text{eq}}}{t_{s,p}^{\pm}}.$$

Thus, the profile for solver s indicates the fraction of problems $\varrho_s(\tau)$ for which solver s invoked by **Marge**^{eq} requires at most τ times the computational effort needed by the same solver s when invoked by **Marge**[±].

As depicted in **Figure 9**, all pairwise performance profiles cross the unit ratio with at least 70% problems solved, meaning that, across all variants, **Marge**^{eq} is a more effective formulation than **Marge**[±] for a large majority of problems. Thus, all solver variants benefit from the tailored handling of equality constraints, especially in terms of gradient evaluations (whose profiles are all above 85% at the unit ratio). The flat region of $\psi_{\rho^*}^{\pm}$ displayed in **Figures 2b** and **3b** is arguably responsible for hindering the performance of **Marge** when equalities are split into two inequalities. In the possibly nonconvex case (bottom panels of **Figure 9**), the narrower advantage of **Marge**^{eq} could stem from the fact that solvers

might end up in different local solutions, or even spurious ones due to the reformulation with two inequalities. Moreover, the smaller benefit of Marge^{eq} observed in terms of runtime over gradient evaluations could be ascribed to the slightly more complicated computation of $\psi_{\rho^*}^{\text{eq}}$ over ψ_{ρ^*} . Regardless, all variants exhibited a robust performance with both formulations, confirming that our algorithmic framework can endure redundant, degenerate constraints and hidden equalities.

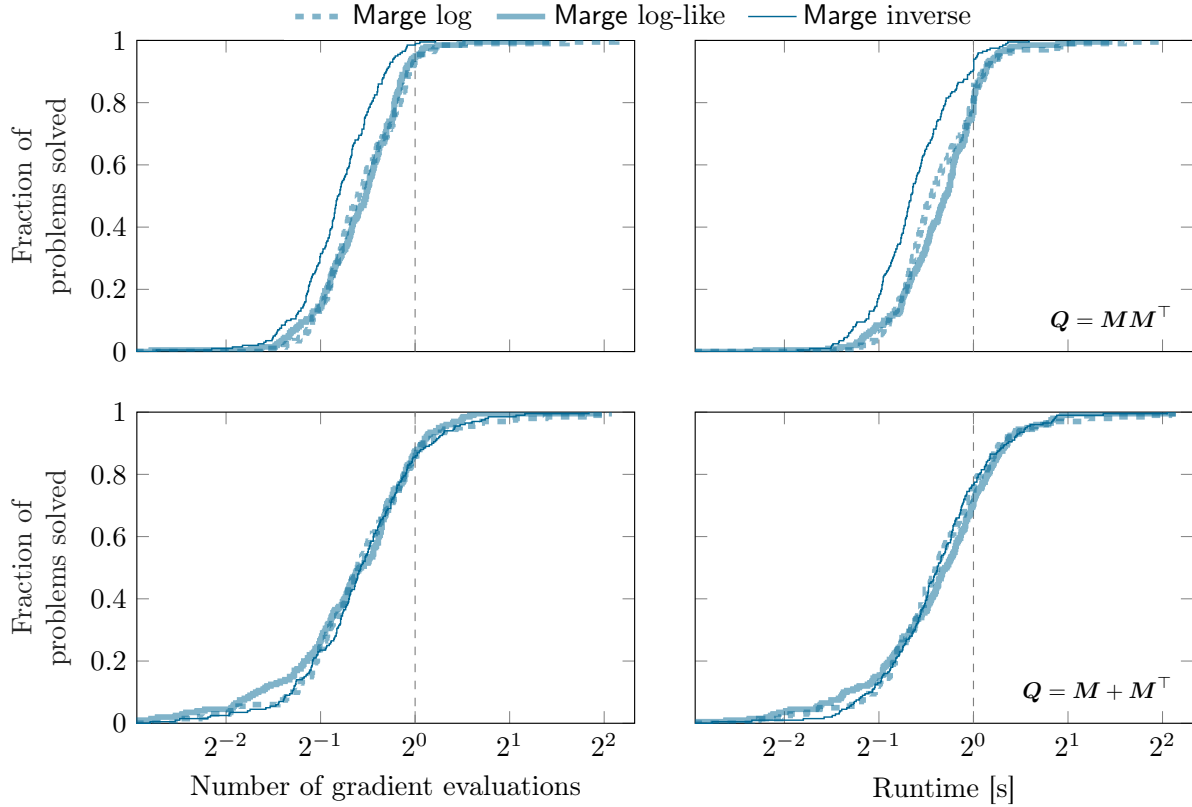


Figure 9: Quadratic programs (5.4): comparison of different solvers and formulations using pairwise performance profiles, relative to gradient evaluations (left) and runtime (right), for Marge^{eq} (explicit equality) over $\text{Marge}^{\pm}$ (split into two inequalities). Profiles located in the top-left indicate that Marge^{eq} tends to outperform $\text{Marge}^{\pm}$. Across all barrier functions, Marge^{eq} is more efficient than $\text{Marge}^{\pm}$ for both convex (top panels) and possibly nonconvex problems (bottom panels).

6 FINAL REMARKS

We proposed Marge , an optimization framework for the numerical solution of constrained structured problems in the fully nonconvex setting. Marge combines (exact) penalty and barrier approaches through a marginalization step, which not only preserves the problem size by avoiding auxiliary variables, but also enables the adoption of generic subsolvers. In particular, by extending the domain of the subproblems' smooth objective term, the proposed methodology overcomes the need for safeguards within the subsolver and the difficulty of accelerating it, a major drawback of IPprox [13]. Under mild assumptions, our theoretical analysis established convergence results on par with those typical for nonconvex continuous optimization. Most notably, all feasible accumulation points are asymptotically KKT optimal. We validated and compared our approach numerically with problems arising in data

science, studying scalability and the effect of accuracy requirements. Furthermore, illustrative examples confirmed the robust behavior of *Marge* on badly formulated problems and degenerate cases.

The methodology in this paper could be applied to a combination of barrier and augmented Lagrangian approaches. By generating a smoother penalty-barrier term, this strategy could benefit from the more effective performance of subsolvers. However, this development comes with the additional challenge of designing suitable updates for the Lagrange multipliers. Future research may also focus on specializing the proposed framework to classical nonlinear programming, taking advantage of the special structure and linear algebra. Finally, mechanisms for rapid infeasibility detection and guaranteed existence of subproblems' solutions should be investigated.

APPENDIX A AUXILIARY RESULTS AND MISSING PROOFS

This appendix contains some auxiliary results and proofs of statements referred to in the main body.

Lemma A.1 (Properties of the barrier b). *Any function b as in Assumption 2 satisfies the following:*

- (i) $\lim_{t \rightarrow -\infty} b(t) = \inf b$ and $\lim_{t \rightarrow 0^-} b(t) = \lim_{t \rightarrow 0^-} b'(t) = \infty$.
- (ii) The conjugate b^* is continuously differentiable on the interior of its domain $\text{int dom } b^* = \mathbb{R}_{++}$ with $(b^*)' < 0$, and satisfies $b^*(0) = -\inf b$ and $\lim_{t^* \rightarrow \infty} b^*(t^*) = -\infty$.
- (iii) $b^*(t^*) = (b^*)'(t^*)t^* - b((b^*)'(t^*))$ for any $t^* > 0$.
- (iv) The function $(0, \infty) \ni t^* \mapsto b^*(t^*)/t^* = t - b(t)/b'(t)$, where $t := (b^*)'(t^*)$, is strictly increasing for t^* large enough with $\lim_{t^* \rightarrow \infty} b^*(t^*)/t^* = 0$. If $\inf b \geq 0$, then it is strictly increasing on the whole $(0, \infty)$.

Proof.

◇ **A.1(i)** Trivial because of strict monotonicity on $(-\infty, 0)$ (since $b' > 0$).

◇ **A.1(ii)** It follows from the definition of Fenchel conjugate that $b^*(0) = -\inf b$, see also [1, Prop. 13.10(i)]. Notice that b is (essentially) strictly convex and *essentially smooth*, in the sense that $b'(t) \rightarrow \infty$ as $t \rightarrow 0^-$, 0 being the only point in the boundary of $\text{dom } b$. As such, the conjugate b^* enjoys the same properties, with $\text{int dom } b^* = \text{range } b' = \mathbb{R}_{++}$ by virtue of [23, Thm. 26.1 and 26.3]. For the same reason, $\text{range}(b^*)' = \text{int dom } b = \mathbb{R}_{--}$, hence $b^{**} < 0$ on $(0, \infty)$. Finally, since $\inf b^* = -b(0) = -\infty$, we conclude that $\lim_{t^* \rightarrow \infty} b^*(t^*) = -\infty$.

◇ **A.1(iii)** This is a standard result of Fenchel conjugacy, see e.g. [1, Prop. 16.10], here specialized to the fact that $\text{range } b' = \mathbb{R}_{++}$.

◇ **A.1(iv)** Observe that $(b^*(t^*)/t^*)' = b(t)/(t^*)^2$ for $t^* > 0$ and $t := (b^*)'(t^*) \rightarrow 0^-$ as $t^* \rightarrow \infty$. Since $b(0^-) = \infty$, for t^* large enough (or for any $t^* > 0$ if $\inf b \geq 0$) this derivative is strictly positive, and as such the function strictly increasing. Lastly,

$$(A.1) \quad \lim_{t^* \rightarrow \infty} \frac{b^*(t^*)}{t^*} = \lim_{t^* \rightarrow \infty} b^{**}(t^*) = \lim_{b'(t) \rightarrow \infty} t = \lim_{t \rightarrow 0^-} t = 0,$$

where the first equality uses L'Hôpital's rule. □

Proof of Lemma 3.1. The Lagrangian associated to (Q_α) reads

$$\begin{aligned} \mathcal{L}(x, z, z_{\text{eq}}; \lambda^+, \lambda^-, y) &= q(x) + \alpha \langle 1, z \rangle + \delta_{\mathbb{R}_+^m}(z) + \alpha \langle 1, z_{\text{eq}} \rangle \\ &\quad + \langle y, c(x) - z \rangle + \langle \lambda^+, c_{\text{eq}}(x) - z_{\text{eq}} \rangle - \langle \lambda^-, c_{\text{eq}}(x) + z_{\text{eq}} \rangle, \end{aligned}$$

so that the corresponding KKT conditions are

$$\begin{cases} \mathbf{0} \in \partial q(\mathbf{x}) + \mathbf{J}\mathbf{c}(\mathbf{x})^\top \mathbf{y} + \mathbf{J}\mathbf{c}_{\text{eq}}(\mathbf{x})^\top (\boldsymbol{\lambda}^+ - \boldsymbol{\lambda}^-) \\ \mathbf{0} \in \boldsymbol{\alpha} - \mathbf{y}_i + \mathbf{N}_{\mathbb{R}_+}(z_i) \\ \mathbf{0} = \boldsymbol{\alpha} - \boldsymbol{\lambda}_j^+ - \boldsymbol{\lambda}_j^- \end{cases} \quad \begin{cases} c_i(\mathbf{x}) \leq z_i \\ |c_{\text{eq},j}(\mathbf{x})| \leq z_{\text{eq},j} \\ y_i, \lambda_j^\pm \geq 0 \end{cases} \quad \begin{cases} 0 = y_i(c_i(\mathbf{x}) - z_i) \\ 0 = \lambda_j^+(z_{\text{eq},j} - c_{\text{eq},j}(\mathbf{x})) \\ 0 = \lambda_j^-(z_{\text{eq},j} + c_{\text{eq},j}(\mathbf{x})) \end{cases}$$

where $i = 1, \dots, m$ and $j = 1, \dots, m_{\text{eq}}$. Here, the first set of conditions corresponds to Lagrangian stationarity (LS), the second one to primal and dual feasibility (PDF), and the last one to complementarity slackness (CS).

Suppose that $z_{\text{eq},j} > |c_{\text{eq},j}(\mathbf{x})|$; then, CS implies that $\lambda_j^\pm = 0$, contradicting the fact that $\lambda_j^+ + \lambda_j^- = \alpha$ in LS. Thus, $z_{\text{eq}} = |c_{\text{eq}}(\mathbf{x})|$ must hold. Suppose instead that $z_i > c_i(\mathbf{x})$; then, $y_i = 0$ by CS, and the second condition in LS then implies that $z_i = 0$ (for otherwise $\mathbf{N}_{\mathbb{R}_+}(z_i) = \{0\}$). Either way, since $y_i - \alpha \in \mathbf{N}_{\mathbb{R}_+}(z_i) \subseteq \mathbb{R}_-$, one has that $y_i - \alpha \leq 0$. These observations show that $\mathbf{z} = [\mathbf{c}(\mathbf{x})]_+$ and that $\mathbf{0} \leq \mathbf{y} \leq \boldsymbol{\alpha}\mathbf{1}$.

Set $\mathbf{y}_{\text{eq}} := \boldsymbol{\lambda}^+ - \boldsymbol{\lambda}^-$, which combined with the last condition in LS yields that $\boldsymbol{\lambda}^+ = \frac{1}{2}(\boldsymbol{\alpha}\mathbf{1} + \mathbf{y}_{\text{eq}})$ and $\boldsymbol{\lambda}^- = \frac{1}{2}(\boldsymbol{\alpha}\mathbf{1} - \mathbf{y}_{\text{eq}})$. Since $\boldsymbol{\lambda}^\pm \geq \mathbf{0}$ by PDF, one has that $|\mathbf{y}_{\text{eq}}| \leq \boldsymbol{\alpha}\mathbf{1}$. With these substitutions, observing that

$$(A.2) \quad z_{\text{eq}} - c_{\text{eq}}(\mathbf{x}) = 2[c_{\text{eq}}(\mathbf{x})]_-, \quad z_{\text{eq}} + c_{\text{eq}}(\mathbf{x}) = 2[c_{\text{eq}}(\mathbf{x})]_+, \quad \text{and} \quad \mathbf{z} - \mathbf{c}(\mathbf{x}) = [\mathbf{c}(\mathbf{x})]_-,$$

the KKT conditions simplify as

$$\begin{cases} \mathbf{0} \in \partial q(\mathbf{x}) + \mathbf{J}\mathbf{c}(\mathbf{x})^\top \mathbf{y} + \mathbf{J}\mathbf{c}_{\text{eq}}(\mathbf{x})^\top \mathbf{y}_{\text{eq}} \\ \mathbf{0} \leq \mathbf{y} \leq \boldsymbol{\alpha}\mathbf{1} \\ |\mathbf{y}_{\text{eq}}| \leq \boldsymbol{\alpha}\mathbf{1} \end{cases} \quad \begin{cases} 0 = y_i[c_i(\mathbf{x})]_-, \quad y_i - \alpha \in \mathbf{N}_{\mathbb{R}_+}([c_i(\mathbf{x})]_+) \\ 0 = (\alpha + y_{\text{eq},j})[c_{\text{eq},j}(\mathbf{x})]_- \\ 0 = (\alpha - y_{\text{eq},j})[c_{\text{eq},j}(\mathbf{x})]_+ \end{cases}$$

where $i = 1, \dots, m$ and $j = 1, \dots, m_{\text{eq}}$. Noticing that $\mathbf{N}_{\mathbb{R}_+}([c_i(\mathbf{x})]_+) = \{0\}$ when $[c_i(\mathbf{x})]_+ > 0$, (KKT _{α}) are obtained. Conversely, by reverting $\mathbf{y}_{\text{eq}} = \boldsymbol{\lambda}^+ - \boldsymbol{\lambda}^-$ and using (A.2) to substitute $[c_{\text{eq}}(\mathbf{x})]_\pm$ and $[\mathbf{c}(\mathbf{x})]_+$ one reobtains the KKT conditions for problem (Q _{α}). $\square$

Proof of Theorem 3.5. We start by observing that

$$\psi_{\rho^*}(t) \stackrel{(\text{def})}{=} \min_{z \geq 0} \{\rho^* z + b(t - z)\} = \min_{z \in \mathbb{R}} \{\rho^* |z| + b(t - z)\},$$

owing to the fact that b is increasing and consequently $\inf_{z \leq 0} b(t - z) = b(t)$ for any t . Hence, ψ_{ρ^*} is the ρ^* -Pasch-Hausdorff envelope of the convex function b as in [1, Def. 12.16], and is itself convex by virtue of [1, Prop. 12.11]. Since $b'' > 0$, and $b'(0^-) = \infty$, there exists $\rho < 0$ such that $b'(\rho) = \rho^*$. In particular, $b'((-\infty, \rho]) = (0, \rho^*]$ and $b'((\rho^*, 0)) = (\rho^*, \infty)$, implying that b is ρ^* -Lipschitz continuous on $(-\infty, \rho]$. The expression (3.7) and ρ^* -Lipschitz continuity then follow from [1, Prop. 12.17(i)], and in turn so does the expression (3.6) of the derivative, which is clearly globally Lipschitz continuous as well. Finally, that $\psi_{\rho^*} \circ c$ is convex whenever c is convex follows from the fact that ψ_{ρ^*} is increasing (additionally to being convex). $\square$

Proof of Theorem 3.6. Function $F(z, t) := \rho^* z + b(t - z) + b(-t - z)$ being minimized on the right-hand side of (3.4) is convex, hence so is its marginalized (wrt z) function $\psi_{\rho^*}^{\text{eq}}$. For every $t \in \mathbb{R}$, $z \mapsto F(z, t)$ is proper, lsc, strictly convex, coercive, and differentiable on its (open) domain, and thus admits a unique minimizer $z_{\rho^*}(t)$, this being the (unique) zero of the derivative, that is, such that (3.9) holds. In

particular, $\psi_{\rho^*}^{\text{eq}}$ is convex and finite valued, and thus everywhere subdifferentiable. Appealing to [25, Thm. 10.13] and denoting $z = z_{\rho^*}(t)$, its (regular, or equivalently, convex) subdifferential satisfies

$$\emptyset \neq \hat{\partial}\psi_{\rho^*}^{\text{eq}}(t) \subseteq \left\{ y \mid \begin{pmatrix} 0 \\ y \end{pmatrix} \in \hat{\partial}F(z, t) \right\} = \left\{ y \mid \begin{pmatrix} 0 \\ y \end{pmatrix} \in \begin{pmatrix} \rho^* - b'(t-z) - b'(-t-z) \\ b'(t-z) - b'(-t-z) \end{pmatrix} \right\} \subseteq \{b'(t-z) - b'(-t-z)\}.$$

This shows that $\psi_{\rho^*}^{\text{eq}}$ is everywhere differentiable with derivative

$$(\psi_{\rho^*}^{\text{eq}})'(t) = b'(t - z_{\rho^*}(t)) - b'(-t - z_{\rho^*}(t)) = \rho^* - 2b'(-t - z_{\rho^*}(t)),$$

as claimed, where the second identity follows from (3.9). Notice that (3.9) also implies that $b'(\pm t - z_{\rho^*}(t)) \leq \rho^*$ (by $b' > 0$); since b' is increasing, one must have that $\pm t - z_{\rho^*}(t) \leq (b^*)'(\rho^*) =: \rho$, yielding the claimed bound $z_{\rho^*}(t) > |t| - \rho$. Notice further that $(\psi_{\rho^*}^{\text{eq}})'(t) < \rho^*$, since $b' > 0$, and consequently $(\psi_{\rho^*}^{\text{eq}})'(t) > -\rho^*$ as well by symmetry. In particular, $\psi_{\rho^*}^{\text{eq}}$ is globally ρ^* -Lipschitz continuous.

We next turn to Lipschitz differentiability. We first demonstrate that the mapping $\mathbb{R}_+ \ni t \mapsto z_{\rho^*}(t)$ is nonexpansive. Fix $t' > t \geq 0$ and let $\varepsilon := z_{\rho^*}(t') - z_{\rho^*}(t) - (t' - t)$; since, apparently, $z_{\rho^*}(t') \geq z_{\rho^*}(t)$, the claim is proven once we show that $\varepsilon \leq 0$. It follows from (3.9) that

$$\begin{aligned} b'(t - z_{\rho^*}(t)) + b'(-t - z_{\rho^*}(t)) &= b'(t' - z_{\rho^*}(t')) + b'(-t' - z_{\rho^*}(t')) \\ &= b'(t - z_{\rho^*}(t) - \varepsilon) + b'(-t - z_{\rho^*}(t) - 2t' - \varepsilon). \end{aligned}$$

For the top left-hand side to equal the bottom right-hand side $\varepsilon < 0$ must hold, for otherwise $b'(t - z_{\rho^*}(t) - \varepsilon) < b'(t - z_{\rho^*}(t))$ and $b'(-t - z_{\rho^*}(t) - 2t' - \varepsilon) < b'(-t - z_{\rho^*}(t))$ (since b' is increasing and $t' > 0$). This shows that $|z_{\rho^*}(t') - z_{\rho^*}(t)| \leq |t' - t|$, as claimed. Next, observe that

$$|(\psi_{\rho^*}^{\text{eq}})'(t') - (\psi_{\rho^*}^{\text{eq}})'(t)| = 2|b'(-t - z_{\rho^*}(t)) - b'(-t' - z_{\rho^*}(t'))|,$$

and both $-t - z_{\rho^*}(t)$ and $-t' - z_{\rho^*}(t')$ are larger than $\rho = (b^*)'(\rho^*) < 0$, as shown above. In particular,

$$|(\psi_{\rho^*}^{\text{eq}})'(t') - (\psi_{\rho^*}^{\text{eq}})'(t)| \leq 2B|t' - t + z_{\rho^*}(t') - z_{\rho^*}(t)| \leq 4B|t' - t|,$$

where $B := \sup_{(-\infty, \rho]} b'' < \infty$ is a finite quantity (by the properties of b in [Assumption 2](#)) that depends only on ρ^* . This shows that $(\psi_{\rho^*}^{\text{eq}})'$ is $4B$ -Lipschitz continuous on $\mathbb{R}_+$, hence on the entire $\mathbb{R}$ by symmetry.

Lastly, take $t > 0$ and observe that, since $z_{\rho^*}(t) > |t| - \rho > 0$ and b' is increasing, one has $(\psi_{\rho^*}^{\text{eq}})'(t) = \rho^* - 2b'(-t - z_{\rho^*}(t)) > \rho^* - 2b'(-t)$. By symmetry, the claimed inequality $|(\psi_{\rho^*}^{\text{eq}})'(t)| > \rho^* - 2b'(-|t|)$ follows. $\square$

Proof of Theorem 3.7. We start by observing that when $b > 0$ the pointwise monotonic decrease of ψ_{ρ^*}/ρ^* and $\psi_{\rho^*}^{\text{eq}}/\rho^*$ as ρ^* grows is apparent from the respective definitions (3.5) and (3.7). Next, the claimed limit of ψ_{ρ^*} follows from the expression (3.5) and [Lemma A.1\(iv\)](#). As to $\psi_{\rho^*}^{\text{eq}}$, notice that it satisfies

$$\psi_{\rho^*}^{\text{eq}}(t) \leq \inf_{z \geq 0} \{ \rho^* z + 2b(|t| - z) \} = 2\psi_{\rho^*/2}(|t|),$$

where the inequality owes to the fact that $b' > 0$ (hence that b is increasing). Next, fix $t \in \mathbb{R}$ and plug $z_{\rho^*}(t) = z_{\rho^*}(|t|)$ into (3.3) to obtain a bound

$$\psi_{\rho^*}(|t|) \stackrel{(\text{def})}{=} \min_{z \geq 0} \{ \rho^* z + b(|t| - z) \} \leq \rho^* z_{\rho^*}(|t|) + b(|t| - z_{\rho^*}(|t|)) = \psi_{\rho^*}^{\text{eq}}(t) - b(-|t| - z_{\rho^*}(t)).$$

Observe that $z_{\rho^*}(t) \searrow |t|$ as $\rho^* \rightarrow \infty$, implying that the term $b(-|t| - z_{\rho^*}(t))$ is positive for ρ^* large enough (or for any ρ^* in case $b > 0$). Thus, for any fixed t and ρ^* large enough one has that $\psi_{\rho^*}^{\text{eq}}(t) \geq \psi_{\rho^*}(|t|)$, which combined with the previous bound results in

$$\psi_{\rho^*}(|t|) \leq \psi_{\rho^*}^{\text{eq}}(t) \leq 2\psi_{\rho^*/2}(|t|)$$

for all ρ^* large enough (depending on t). Dividing by ρ^* and letting $\rho^* \rightarrow \infty$, it follows from the earlier claim on ψ_{ρ^*} that the lower and the upper bounds both converge to $[|t|]_+ = |t|$, demonstrating the claim. $\square$

APPENDIX B BARRIER PROPERTIES

The analysis in Section 4.2 of the barrier's properties is corroborated by examining the performance of solver variants in terms of outer iterations. Computational results relative to the nonnegative PCA example of Section 5.2 are depicted in Figure 10, where (i) the log-like and log barriers have almost indistinguishable profiles and (ii) they invariably demand the solution of fewer subproblems than the inverse barrier. Although observation (i) is supported also by the results in Figure 8, it does not hold in general, as pointed out with the experiments summarized in Figure 7 and the associated discussion. Regardless, both observations are in agreement with the analysis in Section 4.2, in particular with the log-like barrier having a better (in fact, optimal) behavior profile than the inverse barrier. Profiles with an analogous pattern were observed also for the matrix completion task of Section 5.3, but with a less pronounced discrepancy between the inverse and other barrier functions. Interestingly, Figure 10 suggests that this assessment remains valid also for IPprox's profiles, although not being covered by our theory.

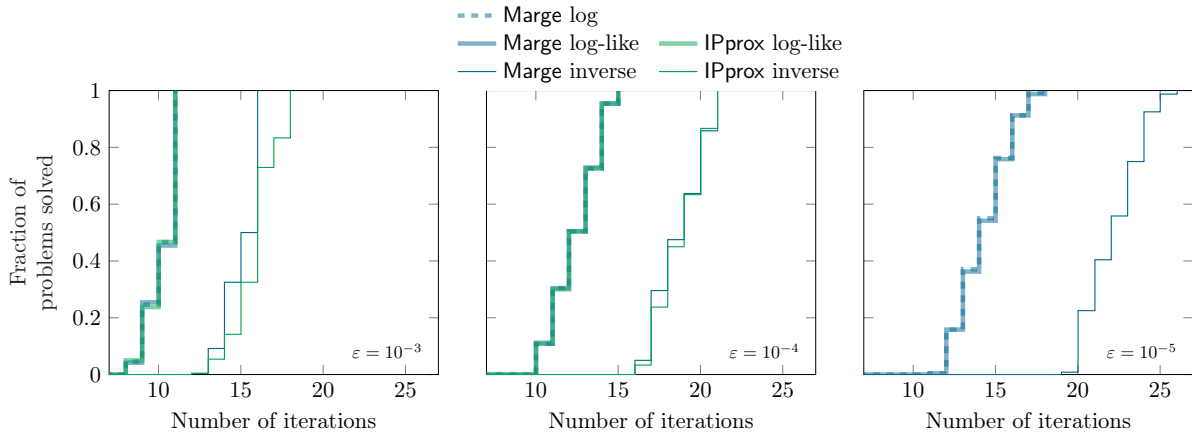


Figure 10: Nonnegative PCA problem (5.1), small and large instances: comparison of solvers with low, medium and high accuracy (left to right) using data profiles relative to the number of outer iterations. With high accuracy, results for IPprox (green solid lines) are not included due to excessive runtime; in this case the other solvers are initialized with a possibly infeasible guess. The comparisons in terms of *outer* iterations for Marge (blue) and IPprox (green) confirms the theoretical appeal of “well-behaved” logarithmic (thick dashed lines) and log-like (thick solid lines) barriers for the former method, with the two profiles almost perfectly overlapping.

REFERENCES

- [1] H. H. Bauschke and P. L. Combettes, *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*, CMS Books in Mathematics, Springer, 2017, [doi:10.1007/978-1-4419-9467-7](https://doi.org/10.1007/978-1-4419-9467-7).
- [2] E. Bertolazzi, mexIPOPT: MATLAB interface for Ipopt (version 1.1.6), 2023, <https://github.com/ebertolazzi/mexIPOPT>.

- [3] S. Bi and S. Pan, Multistage convex relaxation approach to rank regularized minimization problems based on equivalent mathematical program with a generalized complementarity constraint, *SIAM Journal on Control and Optimization* 55 (2017), 2493–2518, [doi:10.1137/15m1037160](https://doi.org/10.1137/15m1037160).
- [4] E. G. Birgin and J. M. Martínez, *Practical Augmented Lagrangian Methods for Constrained Optimization*, Society for Industrial and Applied Mathematics, Philadelphia, PA, 2014, [doi:10.1137/1.9781611973365](https://doi.org/10.1137/1.9781611973365).
- [5] E. Chouzenoux, M. C. Corbineau, and J. C. Pesquet, A proximal interior point algorithm with applications to image processing, *Journal of Mathematical Imaging and Vision* 62 (2020), 919–940, [doi:10.1007/s10851-019-00916-w](https://doi.org/10.1007/s10851-019-00916-w).
- [6] A. R. Conn, N. I. M. Gould, and P. L. Toint, A globally convergent augmented Lagrangian algorithm for optimization with general constraints and simple bounds, *SIAM Journal on Numerical Analysis* 28 (1991), 545–572, [doi:10.1137/0728030](https://doi.org/10.1137/0728030).
- [7] A. R. Conn, N. I. M. Gould, and P. L. Toint, *Trust Region Methods*, Society for Industrial and Applied Mathematics, 2000, [doi:10.1137/1.9780898719857](https://doi.org/10.1137/1.9780898719857).
- [8] F. E. Curtis, A penalty-interior-point algorithm for nonlinear constrained optimization, *Mathematical Programming Computation* 4 (2012), 181–209, [doi:10.1007/s12532-012-0041-4](https://doi.org/10.1007/s12532-012-0041-4).
- [9] A. De Marchi, Proximal gradient methods beyond monotony, *Journal of Nonsmooth Analysis and Optimization* Volume 4 (2023), 10290, [doi:10.46298/jnsao-2023-10290](https://doi.org/10.46298/jnsao-2023-10290).
- [10] A. De Marchi, Implicit augmented Lagrangian and generalized optimization, *Journal of Applied and Numerical Optimization* 6 (2024), 291–320, [doi:10.23952/jano.6.2024.2.08](https://doi.org/10.23952/jano.6.2024.2.08).
- [11] A. De Marchi, X. Jia, C. Kanzow, and P. Mehlitz, Constrained composite optimization and augmented Lagrangian methods, *Mathematical Programming* 201 (2023), 863–896, [doi:10.1007/s10107-022-01922-4](https://doi.org/10.1007/s10107-022-01922-4).
- [12] A. De Marchi and A. Themelis, Proximal gradient algorithms under local Lipschitz gradient continuity, *Journal of Optimization Theory and Applications* 194 (2022), 771–794, [doi:10.1007/s10957-022-02048-5](https://doi.org/10.1007/s10957-022-02048-5).
- [13] A. De Marchi and A. Themelis, An interior proximal gradient method for nonconvex optimization, *Open Journal of Mathematical Optimization* 5 (2024), 1–22, [doi:10.5802/ojmo.30](https://doi.org/10.5802/ojmo.30).
- [14] N. K. Dhingra, S. Z. Khong, and M. R. Jovanović, The proximal augmented Lagrangian method for nonsmooth composite optimization, *IEEE Transactions on Automatic Control* 64 (2019), 2861–2868, [doi:10.1109/tac.2018.2867589](https://doi.org/10.1109/tac.2018.2867589).
- [15] A. V. Fiacco and G. P. McCormick, The sequential unconstrained minimization technique for nonlinear programming, a primal-dual method, *Management Science* 10 (1964), 360–366, [doi:10.1287/mnsc.10.2.360](https://doi.org/10.1287/mnsc.10.2.360).
- [16] N. Hallak and M. Teboulle, An adaptive Lagrangian-based scheme for nonconvex composite optimization, *Mathematics of Operations Research* 48 (2023), 2337–2352, [doi:10.1287/moor.2022.1342](https://doi.org/10.1287/moor.2022.1342).
- [17] B. Hermans, A. Themelis, and P. Patrinos, QPALM: A proximal augmented Lagrangian method for nonconvex quadratic programs, *Mathematical Programming Computation* 14 (2022), 497–541, [doi:10.1007/s12532-022-00218-0](https://doi.org/10.1007/s12532-022-00218-0).

- [18] G. Leconte and D. Orban, An interior-point trust-region method for nonsmooth regularized bound-constrained optimization, *Les Cahiers du GERAD G-2024-17*, GERAD, Montréal, Canada, 2024, <https://www.gerad.ca/en/papers/G-2024-17>.
- [19] J. Mareček, P. Richtárik, and M. Takáč, Matrix completion under interval uncertainty, *European Journal of Operational Research* 256 (2017), 35–43, [doi:10.1016/j.ejor.2016.07.014](https://doi.org/10.1016/j.ejor.2016.07.014).
- [20] E. J. McShane, Extension of range of functions, *Bulletin of the American Mathematical Society* 40 (1934), 837–842, [doi:10.1090/s0002-9904-1934-05978-0](https://doi.org/10.1090/s0002-9904-1934-05978-0).
- [21] J. J. Moré and S. M. Wild, Benchmarking derivative-free optimization algorithms, *SIAM Journal on Optimization* 20 (2009), 172–191, [doi:10.1137/080724083](https://doi.org/10.1137/080724083).
- [22] P. Pas, M. Schuurmans, and P. Patrinos, Alpaqa: A matrix-free solver for nonlinear MPC and large-scale nonconvex optimization, in *2022 European Control Conference (ECC)*, 2022, 417–422, [doi:10.23919/ecc55457.2022.9838172](https://doi.org/10.23919/ecc55457.2022.9838172).
- [23] R. T. Rockafellar, *Convex Analysis*, Princeton University Press, 1970, [doi:10.1515/9781400873173](https://doi.org/10.1515/9781400873173).
- [24] R. T. Rockafellar, Convergence of augmented Lagrangian methods in extensions beyond nonlinear programming, *Mathematical Programming* 199 (2022), 375–420, [doi:10.1007/s10107-022-01832-5](https://doi.org/10.1007/s10107-022-01832-5).
- [25] R. T. Rockafellar and R. J. Wets, *Variational Analysis*, volume 317, Springer, 1998, [doi:10.1007/978-3-642-02431-3](https://doi.org/10.1007/978-3-642-02431-3).
- [26] M. Simões, A. Themelis, and P. Patrinos, Lasry–Lions envelopes and nonconvex optimization: A homotopy approach, in *2021 29th European Signal Processing Conference (EUSIPCO)*, 2021, 2089–2093, [doi:10.23919/eusipco54536.2021.9616167](https://doi.org/10.23919/eusipco54536.2021.9616167).
- [27] P. Sopasakis, E. Fresk, and P. Patrinos, OpEn: Code generation for embedded nonconvex optimization, *IFAC-PapersOnLine* 53 (2020), 6548–6554, [doi:10.1016/j.ifacol.2020.12.071](https://doi.org/10.1016/j.ifacol.2020.12.071). 21st IFAC World Congress.
- [28] L. Stella, A. Themelis, and P. Patrinos, Forward-backward quasi-Newton methods for nonsmooth optimization problems, *Computational Optimization and Applications* 67 (2017), 443–487, [doi:10.1007/s10589-017-9912-y](https://doi.org/10.1007/s10589-017-9912-y).
- [29] L. Stella, A. Themelis, P. Sopasakis, and P. Patrinos, A simple and efficient algorithm for nonlinear model predictive control, in *2017 IEEE 56th Annual Conference on Decision and Control (CDC)*, IEEE, 2017, 1939–1944, [doi:10.1109/cdc.2017.8263933](https://doi.org/10.1109/cdc.2017.8263933).
- [30] A. Themelis, L. Stella, and P. Patrinos, Forward-backward envelope for the sum of two non-convex functions: Further properties and nonmonotone linesearch algorithms, *SIAM Journal on Optimization* 28 (2018), 2274–2303, [doi:10.1137/16m1080240](https://doi.org/10.1137/16m1080240).
- [31] A. Wächter and L. T. Biegler, Line search filter methods for nonlinear programming: Motivation and global convergence, *SIAM Journal on Optimization* 16 (2005), 1–31, [doi:10.1137/s1052623403426556](https://doi.org/10.1137/s1052623403426556).
- [32] A. Wächter and L. T. Biegler, On the implementation of an interior-point filter line-search algorithm for large-scale nonlinear programming, *Mathematical Programming* 106 (2006), 25–57, [doi:10.1007/s10107-004-0559-y](https://doi.org/10.1007/s10107-004-0559-y).