



## Log-Sigmoid Multipliers Method in Constrained Optimization \*

ROMAN A. POLYAK \*\*

*Systems Engineering and Operations Research Department and Mathematical Sciences Department,  
George Mason University, Fairfax, VA 22030, USA*

**Abstract.** In this paper we introduced and analyzed the Log-Sigmoid (LS) multipliers method for constrained optimization. The LS method is to the recently developed smoothing technique as augmented Lagrangian to the penalty method or modified barrier to classical barrier methods. At the same time the LS method has some specific properties, which make it substantially different from other nonquadratic augmented Lagrangian techniques.

We established convergence of the LS type penalty method under very mild assumptions on the input data and estimated the rate of convergence of the LS multipliers method under the standard second order optimality condition for both exact and nonexact minimization.

Some important properties of the dual function and the dual problem, which are based on the LS Lagrangian, were discovered and the primal–dual LS method was introduced.

**Keywords:** log-sigmoid, multipliers method, duality, smoothing technique

### 1. Introduction

Recently Chen and Mangasarian used the integral of the scaled sigmoid function  $S(t, k) = (1 + \exp(-kt))^{-1}$  as an approximation for  $x_+ = \max\{0, x\}$  to develop the smoothing technique for solving convex system of inequalities and linear complementarity problems [6].

Later Auslender et al. analyzed the smoothing technique for constrained optimization [1].

The smoothing method for constrained optimization employs a smooth approximation of  $x_+$  to transform a constrained optimization problem into a sequence of unconstrained optimization problems. The convergence of the correspondent sequence of the unconstrained minimizers to the primal solution is due to the unbounded increase of the scaling parameter. So the smoothing technique is in fact a penalty type method with a smooth penalty function and can be considered as a particular case of SUMT [7].

There are few well known difficulties associated with the penalty type approach: rather slow convergence, the Hessian of the penalty function became ill conditioned and the area where Newton method is “well” defined shrinks to a point when the scaling parameter  $k \rightarrow \infty$ .

\* This paper is dedicated to Professor Anthony V. Fiacco on the occasion of his 70th birthday.

\*\* Partially supported by NSF under Grant DMS-9705672.

It motivates an alternative approach. We use the Log-Sigmoid (LS) function

$$\psi(t) = 2 \ln 2S(t, 1)$$

to transform the constraints of a given constrained optimization problem into an equivalent one. The transformation  $\psi(t)$  is parametrized by a positive scaling parameter. Simultaneously we transform the objective function with log-sigmoid type transformation. The classical Lagrangian for the equivalent problem – the Log-Sigmoid Lagrangian (LSL) is our basic instrument.

There are three basic reasons for using the LS transformation and the corresponding Lagrangian:

- (1)  $\psi \in C^\infty$  on  $(-\infty, \infty)$ ;
- (2) the LSL is as smooth as the initial functions in the entire primal space;
- (3)  $\psi'$  and  $\psi''$  are bounded on  $(-\infty, \infty)$ .

Sequential unconstrained minimization of the LSL in primal space followed by explicit formula for the Lagrange multipliers update forms the LS multipliers method. Our first contribution is the convergence proof of the LS multipliers method. It is proven that for inequality constrained optimization problem, which satisfies the standard second order optimality conditions the LS method converges with  $Q$ -linear rate for any fixed but large enough scaling parameter. If one changes the scaling parameter from step to step as it takes place in the smoothing methods then the rate of convergence is  $Q$ -superlinear. It is worth to mention that such substantial improvement of the rate of convergence is possible to achieve without increase computational efforts per step as compare with the smoothing technique.

Our second contribution is the proof that a particular modification of the LS method retains the up to  $Q$ -superlinear rate of convergence if instead of the exact primal minimizer one uses its approximation. It makes the LS multipliers method practical and together with the properties (1)–(3) of the transformation  $\psi$  increases the efficiency of the Newton method for constrained optimization.

We also discovered that the dual function and the dual problem, which are based on LSL have some extra important properties on the top of those which are typical for the classical dual function and the corresponded dual problem.

The new properties of the dual function allow to use Newton type methods for solving the dual problem, which leads to the second order multipliers methods with up to quadratic rate of convergence.

Finally we introduced the primal–dual LS method, which has been tested numerically on a number of LP and NLP problems. The numerical results obtained clearly indicate that the primal–dual LS method can be very efficient in the final phase of the computational process.

The paper is organized as follows. The problem formulation and the basic assumptions are given in the next section. In section 3, we consider the LS transformation and its properties. In section 4 we consider the equivalent problem and correspondent

Lagrangian. The LS multiplier method is introduced in section 5. In section 6 we establish the convergence and estimate the rate of convergence of the LS multipliers method. In section 7 we consider the modification of the LS method and show that the rate of convergence of LS methods can be retained for inexact minimization. The primal–dual LS method is introduced in section 8. Duality issues related to the LSL are considered in section 9. We conclude the paper with some remarks related to the future research.

## 2. Statement of the problem and basic assumptions

Let  $f: \mathfrak{R}^n \rightarrow \mathfrak{R}^1$  be convex and all  $c_i: \mathfrak{R}^n \rightarrow \mathfrak{R}^1, i = 1, \dots, p$ , be concave and smooth functions. We consider a convex set  $\Omega = \{x \in \mathfrak{R}_+^n: c_i(x) \geq 0, i = 1, \dots, p\}$  and the following convex optimization problem.

$$x^* \in X^* = \arg \min \{f(x) \mid x \in \Omega\}. \quad (\text{P})$$

We will assume that:

- (A) The optimal set  $X^*$  is not empty and bounded.
- (B) The Slater's condition holds, i.e., there exists  $\hat{x} \in \mathfrak{R}_{++}^n: c_i(\hat{x}) > 0, i = 1, \dots, p$ .

To simplify consideration we will include the nonnegativity constraints  $x_i \geq 0, i = 1, \dots, n$ , into the set  $c_i(x) \geq 0$ , i.e.,

$$\begin{aligned} \Omega &= \{x \in \mathfrak{R}^n: c_i(x) \geq 0, i = 1, \dots, p, c_{p+1}(x) = x_1 \geq 0, \dots, c_{p+n}(x) = x_n \geq 0\} \\ &= \{x \in \mathfrak{R}^n: c_i(x) \geq 0, i = 1, \dots, m\}, \quad m = p + n. \end{aligned}$$

If (B) holds and  $f(x), c_i(x), i = 1, \dots, m$ , are smooth, then the Karush–Kuhn–Tucker's (KKT's) conditions hold true, i.e., there exists a nonnegative vector  $\lambda^* = (\lambda_1^*, \dots, \lambda_m^*)$  such that

$$\nabla_x L(x^*, \lambda^*) = \nabla f(x^*) - \sum_{i=1}^m \lambda_i^* \nabla c_i(x^*) = 0, \quad (2.1)$$

$$\lambda_i^* c_i(x^*) = 0, \quad i = 1, \dots, m, \quad (2.2)$$

where  $L(x, \lambda) = f(x) - \sum_{i=1}^m \lambda_i c_i(x)$  is the Lagrangian for the primal problem (P).

Also due to (B), the optimal dual set

$$L^* = \left\{ \lambda \in \mathfrak{R}_+^m: \nabla f(x^*) - \sum_{i=1}^m \lambda_i \nabla c_i(x^*) = 0, x^* \in X^* \right\} \quad (2.3)$$

is bounded.

Along with the primal problem (P) we consider the dual problem

$$\lambda^* \in L^* = \text{Arg max} \{d(\lambda) \mid \lambda \in \mathfrak{R}_+^m\}, \quad (\text{D})$$

where  $d(\lambda) = \inf_x L(x, \lambda)$  is the dual function.

Later we will use the standard second order optimality condition. Let us assume that the active constraint set at  $x^*$  is  $I^* = \{i: c_i(x^*) = 0\} = \{1, \dots, r\}$ . We consider the vector-functions  $c^T(x) = (c_1(x), \dots, c_m(x))$ ,  $c_{(r)}^T(x) = (c_1(x), \dots, c_r(x))$  and their Jacobians

$$\nabla c(x) = J(c(x)) = \begin{bmatrix} \nabla c_1(x) \\ \vdots \\ \nabla c_m(x) \end{bmatrix}, \quad \nabla c_{(r)}(x) = J(c_{(r)}(x)) = \begin{bmatrix} \nabla c_1(x) \\ \vdots \\ \nabla c_r(x) \end{bmatrix}.$$

The sufficient regularity conditions

$$\text{rank } \nabla c_{(r)}(x^*) = r, \quad \lambda_i^* > 0, \quad i \in I^*, \quad (2.4)$$

together with the sufficient condition for the minimum  $x^*$  to be isolated

$$(\nabla_{xx}^2 L(x^*, \lambda^*)y, y) \geq \rho(y, y), \quad \rho > 0, \quad \forall y \neq 0: \nabla c_{(r)}(x^*)y = 0 \quad (2.5)$$

comprise the standard second order optimality sufficient conditions.

We conclude the section with an assertion, which will be used later. The following assertion is a slight modification of Debreu theorem (see, for example, [11]).

**Assertion 2.1.** Let  $A$  be a symmetric  $n \times n$  matrix, let  $B$  an  $r \times n$  matrix,  $\Lambda = \text{diag}(\lambda_i)_{i=1}^r$  and  $\lambda_i > 0$ .

$$(Ay, y) \geq \rho(y, y), \quad \rho > 0, \quad \forall y: By = 0 \quad (2.6)$$

then there exists  $k_0 > 0$  large enough such that for any  $0 < \mu < \rho$  we have

$$((A + kB^T \Lambda B)x, x) \geq \mu(x, x), \quad \forall x \in \mathfrak{R}^n, \quad (2.7)$$

whenever  $k \geq k_0$ .

### 3. Log-sigmoid transformation

The Log-Sigmoid Transformation (LST)  $\psi: \mathfrak{R} \rightarrow (-\infty, 2 \ln 2)$  we define by the formula

$$\psi(t) = 2 \ln 2S(t, 1) = 2 \ln 2(1 + e^{-t})^{-1} = 2(\ln 2 + t - \ln(1 + e^t)). \quad (3.1)$$

For the scaled log-sigmoid transformation we have

$$k^{-1}\psi(kt) = 2k^{-1} \ln 2S(t, k) = 2k^{-1}(\ln 2 - \ln(1 + e^{-kt})), \quad k > 0.$$

Let us consider the following function:

$$v(t, k) = \begin{cases} 2t + 2k^{-1} \ln 2, & t \leq 0, \\ 2k^{-1} \ln 2, & t \geq 0. \end{cases}$$

It is easy to see that

$$v(t, k) - k^{-1}\psi(kt) = \begin{cases} 2k^{-1} \ln(1 + e^{kt}), & t \leq 0, \\ 2k^{-1} \ln(1 + e^{-kt}), & t \geq 0. \end{cases}$$

Therefore the following estimation is taking place

$$0 \leq v(t, k) - k^{-1}\psi(kt) \leq 2k^{-1} \ln 2, \quad -\infty < t < \infty.$$

The assertion below states the basic LST properties.

**Assertion 3.1.** The LST  $\psi$  has the following properties:

- (A1)  $\psi(0) = 0$ ;
- (A2)  $\psi'(t) = 2(1 + e^t)^{-1} > 0, \forall t \in (-\infty, +\infty)$  and  $\psi'(0) = 1$ ;
- (A3)  $\psi''(t) = -2e^t(1 + e^t)^{-2} < 0, \forall t \in (-\infty, +\infty)$  and  $\psi''(0) = -1/2$ ;
- (A4)  $\lim_{t \rightarrow \infty} \psi'(t) = \lim_{t \rightarrow \infty} 2(1 + e^t)^{-1} = 0$ ;
- (A5) (a)  $0 < \psi'(t) < 2$ ;
- (b)  $-0.5 \leq \psi''(t) < 0, -\infty < t < \infty$ .

One can check properties (A1)–(A5) directly. The substantial difference between  $\psi(t)$  and the shifted log-barrier function, which leads to the MBF theory and methods [11], is that  $\psi(t)$  is defined on  $(-\infty, +\infty)$  together with its derivatives of any order.

The properties (A5) distinguish  $\psi(t)$  not only from shifted barrier and exponential transformation (see [9,17]), but also from classes of nonquadratic augmented Lagrangians  $P_l$  and  $\hat{P}_l$  (see [4, p. 309]) as well as transformations which have been considered lately (see [3,8,13,16]).

The properties (A5) have substantial impact on both global and local behavior of the LS multiplier as well as on its dual equivalents – interior prox method with entropy like  $\varphi$ -divergence distance.

Entropy like  $\varphi$ -divergence distance function and correspondent interior prox method for the dual problem have been considered in [15].

The LS transformation and the correspondent LS multipliers method, which we consider in section 5 is equivalent to a prox method with entropy like  $\varphi$ -divergence distance for the dual problem. The  $\varphi$ -divergence distance is based on Fermi–Dirac kernel  $\varphi = -\psi^*$ , because the Fenchel conjugate of LS

$$\psi^*(s) = \inf\{st - \psi(t) \mid t \in \mathfrak{R}\} = (s - 2) \ln(2 - s) - s \ln s$$

is in fact the Fermi–Dirac entropy type function.

The issues related to LS multipliers method and its dual equivalent we are going to consider in the upcoming paper.

#### 4. Equivalent problem and log-sigmoid Lagrangian

We use  $\ln(1 + e^t)$  to transform the objective function and  $\psi(t)$  to transform the constraints. For the objective function we obtain

$$f(x) := \ln(1 + e^{f(x)}) > 0. \quad (4.1)$$

The constraints transformation is scaled by the parameter  $k > 0$ , i.e.,

$$c_i(x) \geq 0 \iff 2k^{-1} \ln 2 (1 + e^{-kc_i(x)})^{-1} \geq 0, \quad i = 1, 2, \dots, m.$$

Therefore for any given  $k > 0$  the problem

$$x^* \in X^* = \arg \min \{ f(x) \mid 2k^{-1} (\ln 2 - \ln(1 + e^{-kc_i(x)})) \geq 0, \quad i = 1, \dots, m \} \quad (4.2)$$

is equivalent to the original problem (P).

The boundness of  $f(x)$  from below is important for our further considerations.

The Lagrangian for the equivalent problem (4.2) – log-sigmoid Lagrangian is the main tool in our analysis

$$\mathcal{L}(x, \lambda, k) = f(x) + 2k^{-1} \sum_{i=1}^m \lambda_i \ln(1 + e^{-kc_i(x)}) - 2k^{-1} \left( \sum_{i=1}^m \lambda_i \right) \ln 2. \quad (4.3)$$

The LSL can be rewritten as follows

$$\begin{aligned} \mathcal{L}(x, \lambda, k) &= f(x) - 2 \sum_{i=1}^m \lambda_i c_i(k) + 2k^{-1} \sum_{i=1}^m \lambda_i \ln(1 + e^{kc_i(x)}) - 2k^{-1} \left( \sum_{i=1}^m \lambda_i \right) \ln 2 \\ &= f(x) - \sum_{i=1}^m \lambda_i c_i(x) + 2k^{-1} \sum_{i=1}^m \lambda_i \ln e^{-kc_i(x)/2} + 2k^{-1} \sum_{i=1}^m \lambda_i \ln \frac{(1 + e^{kc_i(x)})}{2} \\ &= L(x, \lambda) + 2k^{-1} \sum_{i=1}^m \lambda_i \ln \frac{e^{kc_i(x)/2} + e^{-kc_i(x)/2}}{2} \\ &= L(x, \lambda) + 2k^{-1} \sum_{i=1}^m \lambda_i \ln \operatorname{ch} \left( \frac{kc_i(x)}{2} \right). \end{aligned}$$

The following lemma establishes the basic LSL properties at any KKT's pair  $(x^*, \lambda^*)$ .

**Lemma 4.1.** For any KKT's pair  $(x^*, \lambda^*)$  the following LSL properties are taking place for any  $k > 0$ .

$$(1^\circ) \quad \mathcal{L}(x^*, \lambda^*, k) = f(x^*);$$

$$(2^\circ) \quad \nabla_x \mathcal{L}(x^*, \lambda^*, k) = \nabla_x L(x^*, \lambda^*) = \nabla f(x^*) - \sum_{i=1}^m \lambda_i^* \nabla c_i(x^*) = 0;$$

$$(3^\circ) \quad \nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k) = \nabla_{xx}^2 L(x^*, \lambda^*) + 0.5k \nabla c(x^*)^T \Lambda^* \nabla c(x^*).$$

*Proof.* In view of the complementarity condition we have

$$\mathcal{L}(x^*, \lambda^*, k) = f(x^*) - 2k^{-1} \sum_{i=1}^m \lambda_i^* (\ln 2 - \ln(1 + e^{-kc_i(x^*)})) = f(x^*)$$

for any  $k > 0$ .

For the LSL gradient in  $x$  we have

$$\nabla_x \mathcal{L}(x, \lambda, k) = \nabla f(x) - \sum_{i=1}^m \frac{2\lambda_i e^{-kc_i(x)}}{1 + e^{-kc_i(x)}} \nabla c_i(x) = \nabla f(x) - \sum_{i=1}^m \frac{2\lambda_i}{1 + e^{kc_i(x)}} \nabla c_i(x).$$

Again due to (2.2) we obtain

$$\nabla_x \mathcal{L}(x^*, \lambda^*, k) = \nabla_x L(x^*, \lambda^*) = \nabla f(x^*) - \sum_{i=1}^m \lambda_i^* \nabla c_i(x^*) = 0.$$

For the LSL Hessian in  $x$  we obtain

$$\nabla_{xx}^2 \mathcal{L}(x, \lambda, k) = \nabla_{xx}^2 L(x, \lambda) + 2k \nabla c(x)^T \Lambda (I + e^{kc(x)})^{-2} \nabla c(x),$$

where  $e^{kc(x)} = \text{diag}(e^{kc_i(x)})_{i=1}^m$ ,  $\Lambda = \text{diag}(\lambda_i)_{i=1}^m$  and  $I$  is the identical matrix in  $\mathfrak{R}^m$ .

Again due to (2.2) for any  $k > 0$  we have

$$\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k) = \nabla_{xx}^2 L(x^*, \lambda^*) + 0.5k \nabla c(x^*)^T \Lambda^* \nabla c(x^*). \quad (4.4)$$

□

The local LSL properties (1°)–(3°) are similar to those of the logarithmic MBF function [11]. Globally the LSL has some extra important features due to the properties (A5). These features effect substantially the behavior of the correspondent multipliers method, which we consider in the next section. The following lemma characterizes the convexity properties of LSL.

**Lemma 4.2.** If  $f(x)$  and all  $c_i(x) \in C^2$  then for any fixed  $\lambda \in \mathfrak{R}_{++}^n$  and  $k > 0$  the LSL Hessian is positive definite for any  $x \in \mathfrak{R}^n$ , i.e.,  $\mathcal{L}(x, \lambda, k)$  is strictly convex in  $\mathfrak{R}^n$  and strongly convex on any bounded set in  $\mathfrak{R}^n$ .

*Proof.* The proof follows directly from the formula of the LSL Hessian

$$\begin{aligned} \nabla_{xx}^2 \mathcal{L}(x, \lambda, k) &= \nabla_{xx}^2 L(x, \lambda) + 2k \nabla c_{(p)}(x)^T \Lambda_{(p)} (I_p + e^{kc_{(p)}(x)})^{-2} \nabla c_{(p)}(x) \\ &\quad + 2k \Lambda_{(n)} (I_n + e^{kx})^{-2} \end{aligned}$$

and the convexity of  $f(x)$  and all  $-c_i(x)$ .

□

The following lemma 4.3 is a consequence of property (3°) and assertion 1.1.

**Lemma 4.3.** If conditions (2.4), (2.5) are satisfied then there exists  $k_0 > 0$  and  $M_0 > \mu_0 > 0$  that the following estimation

$$M_0 k(y, y) \geq (\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k)y, y) \geq \mu_0(y, y), \quad \forall y \in \mathfrak{R}^n, \quad (4.5)$$

takes place for any fixed  $k \geq k_0$ .

*Proof.* We obtain the right inequality (4.5) as a consequence of (2.5), (2.7) and (3°) by taking

$$A = \nabla_{xx}^2 L(x^*, \lambda^*) \quad \text{and} \quad B = \nabla_{c(r)}(x^*).$$

The left inequality follows from the formula (4.4) for  $\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k)$  if  $k \geq k_0$  and  $k_0 > 0$  is large enough.  $\square$

**Corollary.** If  $f(x)$  and all  $c_i(x)$  are twice continuous differentiable and  $\varepsilon > 0$  is small enough then for any fixed  $k \geq k_0$  there exist a pair  $M > \mu > 0$  such that for any primal–dual pair

$$w = (x, \lambda) \in S(w^*, \varepsilon) = \{w: \|w - w^*\| \leq \varepsilon\}$$

the following inequalities

$$\mu(x, y) \leq (\nabla_{xx}^2 \mathcal{L}(x, \lambda, k)y, y) \leq M(y, y), \quad \forall y \in \mathfrak{R}^n, \quad (4.6)$$

hold true.

In other words in the neighborhood of the KKT's pair  $(x^*, \lambda^*)$  the condition number  $\text{cond } \nabla_{xx}^2 \mathcal{L}(x, \lambda, k) \leq \mu M^{-1}$  is stable for any fixed  $k \geq k_0$ .

*Remark 4.1.* Lemma 4.3 is true whether  $f(x)$  and all  $-c_i(x)$ ,  $i = 1, \dots, m$ , are convex or not.

**Lemma 4.4.** If  $X^*$  is bounded, then for any  $\lambda \in \mathfrak{R}_{++}^m$  and  $k > 0$  there exists

$$\hat{x} = \hat{x}(\lambda, k) = \arg \min \{ \mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n \}.$$

*Proof.* If  $X^*$  is bounded then by adding one extra constraint  $c_{m+1}(x) = -f(x) + M \geq 0$ , where  $M > 0$  is large enough we obtain due to corollary 20 (see [7, p. 94]), that the feasible set  $\Omega$  is bounded. Therefore without restriction of generality we assume from the very beginning that  $\Omega$  is bounded. We start by establishing that LSL  $\mathcal{L}(x, \lambda, k)$  has no direction of recession in  $x$ , i.e., for any nontrivial direction  $z \in \mathfrak{R}^n$

$$\lim_{t \rightarrow \infty} \mathcal{L}(x + tz, \lambda, k) = \infty, \quad \forall \lambda \in \mathfrak{R}_{++}^m, \quad k > 0.$$

Let  $x \in \text{int } \Omega$ , i.e.,  $c_i(x) > 0$ . Due to the boundness of  $\Omega$  for any  $z \neq 0$  one can find  $i_0$ :  $c_{i_0}(x + \bar{t}z) = 0$ ,  $\bar{t} > 0$ , in fact, if  $c_i(x + tz) > 0 \forall t > 0$ ,  $i = 1, \dots, m$ , then  $\Omega$  is unbounded.

Let  $\bar{x} = x + \bar{t}z$ , using concavity of  $c_{i_0}(x)$  we obtain

$$c_{i_0}(x) - c_{i_0}(\bar{x}) \leq (\nabla c_{i_0}(\bar{x}), x - \bar{x})$$

or

$$0 < \alpha = c_{i_0}(x) \leq -(\nabla c_{i_0}(\bar{x}), z)\bar{t},$$

i.e.,

$$(\nabla c_{i_0}(x), z) \leq -\alpha\bar{t}^{-1} = \beta < 0. \quad (4.7)$$

Again using the concavity of  $c_{i_0}(x)$  we obtain

$$c_{i_0}(x + tz) \leq c_{i_0}(\bar{x}) + (\nabla c_{i_0}(\bar{x}), z)(t - \bar{t})$$

or

$$-c_{i_0}(x + tz) \geq -\beta(t - \bar{t}), \quad \forall t > 0. \quad (4.8)$$

Hence in view of (4.8) we obtain

$$\begin{aligned} \mathcal{L}(x + tz, \lambda, k) &= f(x + tz) + 2k^{-1} \sum \lambda_i \ln(1 + e^{-kc_{i_0}(x+tz)}) - 2k^{-1} \sum \lambda_i \ln 2 \\ &\geq f(x + tz) - 2k^{-1} \sum \lambda_i + 2k^{-1} \lambda_{i_0} \ln(1 + e^{-kc_{i_0}(x+tz)}) \\ &= f(x + tz) - 2k^{-1} \sum \lambda_i + 2k^{-1} \lambda_{i_0} \ln(1 + e^{-kc_{i_0}(x+tz)}) - 2\lambda_{i_0} c_{i_0}(x + tz) \\ &\geq f(x + tz) - 2k^{-1} \sum \lambda_i - 2\beta \lambda_{i_0}(t - \bar{t}). \end{aligned}$$

Taking into account (4.1) and (4.6) we obtain,

$$\lim_{t \rightarrow \infty} \mathcal{L}(x + tz, \lambda, k) = +\infty, \quad \forall z \in \mathfrak{R}^n,$$

so the set

$$\widehat{X}(\lambda, k) = \left\{ \hat{x} \mid \mathcal{L}(\hat{x}, \lambda, k) = \inf_{x \in \mathfrak{R}^n} \mathcal{L}(x, \lambda, k) \right\}$$

is not empty and bounded (Rockafellar [4, theorem 27.1d]).  $\square$

Moreover, for  $(\lambda, k) \in \mathfrak{R}_{++}^{m+1}$  due to lemma 4.2 the set  $\widehat{X}(\lambda, k)$  contains only one point  $\hat{x}(\lambda, k) = \arg \min\{\mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n\}$ . The uniqueness of  $\hat{x}(\lambda, k)$  means that in contrast to the dual function  $d(\lambda) = \inf\{L(x, \lambda) \mid x \in \mathfrak{R}^n\}$  which is based on the Lagrangian for the initial problem (P), the dual function  $d_k(\lambda) = \min\{\mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n\}$ , which is based on LSL, is as smooth as the initial functions for any  $\lambda \in \mathfrak{R}_{++}^n$ .

*Remark 4.2.* Due to (A5(a)) we have  $\lim_{t \rightarrow -\infty} \psi'(t) = 2 < \infty$ , therefore the existence of  $\hat{x}(\lambda, k)$  does not follow from standard considerations [4, p. 329], see also [1]. In fact let us consider the following LP  $\min\{3x \mid x \geq 0\}$ . We have  $X^* = \{0\}$  and  $L^* = \{3\}$ . For  $\lambda = 1$  and  $k = 1$  the LSL  $L(x, 1, 1) = 3x + 2\ln(1 + e^{-x}) - 2\ln 2$  and  $\inf \mathcal{L}(x, 1, 1) = -\infty$ . The transformation of the objective function is critical for the existence of  $\hat{x}(\lambda, k)$ .

*Remark 4.3.* The convex in  $x \in \mathfrak{R}^n$  LSL  $\mathcal{L}(x, \lambda, k)$  has a bounded level set in  $x$  for any  $\lambda \in \mathfrak{R}_{++}^n$  and  $k > 0$ . It does not follow from lemma 12 (see [7, p. 95]), because  $\psi(t)$  does not satisfy the assumption (a).

## 5. Log-sigmoid multipliers method

We consider the following method. For a chosen  $\lambda^0 \in \mathfrak{R}_{++}^n$  and  $k > 0$  we generate iteratively the sequence  $\{x^s\}$  and  $\{\lambda^s\}$  according to the following formulas:

$$x^{s+1} = \arg \min \{ \mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n \}, \quad (5.1)$$

$$\lambda_i^{s+1} = \lambda_i^s \psi'(k c_i(x^{s+1})) = 2\lambda_i^s (1 + e^{k c_i(x^{s+1})})^{-1}, \quad i = 1, \dots, m. \quad (5.2)$$

Along with multipliers method (5.1), (5.2) we consider a version of this method when the parameter  $k > 0$  is not fixed but one can change it from step to step. For a given positive sequence  $\{k_s\}$ :  $k_{s+1} > k_s$ ,  $\lim_{s \rightarrow \infty} k_s = \infty$ , we find the primal  $\{x^s\}$  and the dual  $\{\lambda^s\}$  sequences by formulas

$$x^{s+1} = \arg \min \{ \mathcal{L}(x, \lambda, k_s) \mid x \in \mathfrak{R}^n \}, \quad (5.3)$$

$$\lambda_i^{s+1} = \lambda_i^s \psi'(k_s c_i(x^{s+1})) = 2\lambda_i^s (1 + e^{k_s c_i(x^{s+1})})^{-1}, \quad i = 1, \dots, m. \quad (5.4)$$

First of all we have to guarantee that the multipliers method (5.1), (5.2) is well defined, i.e., that  $x^{s+1}$  exists for any given  $\lambda^s \in \mathfrak{R}_{++}^m$  and  $k > 0$ .

Due to lemma 4.4 for any  $\lambda^s \in \mathfrak{R}_{++}^m$  there exist  $\hat{x}(\lambda^s, k) = \arg \min \{ \mathcal{L}(x, \lambda^s, k) \mid x \in \mathfrak{R}^n \}$  and due to the formulas (5.2) and (5.4) we have  $\lambda^s \in \mathfrak{R}_{++}^m \Rightarrow \lambda^{s+1} \in \mathfrak{R}_{++}^m$ . Therefore if the starting vector of Lagrange multipliers  $\lambda^0 \in \mathfrak{R}_{++}^m$  then all vectors  $\lambda^s$ ,  $s = 1, 2, \dots$ , will remain positive, so the LS method is executable.

The critical part of any multipliers method is the formula for the Lagrange multipliers update. It follows from (5.2) and (5.4) that  $\lambda_i^{s+1} > \lambda_i^s$  if  $c(x^{s+1}) < 0$  and  $\lambda_i^{s+1} < \lambda_i^s$  if  $c(x^{s+1}) > 0$ . In this respect the LS method is similar to other multipliers method, however due to (A5(a)) the LS method has some very specific properties. In particular the Lagrange multipliers cannot be increased more than twice independent on the constraint violation and the value of the scaling parameter  $k > 0$ . It means, for instance, if  $\lambda^0 = e = (1, \dots, 1) \in \mathfrak{R}^m$  is the starting Lagrange multipliers vector then for any  $k > 0$  large enough and any constraint violation the new Lagrange multipliers cannot be more than two. Therefore in contrast to the exponential [17] or MBF methods [11] it is impossible to find approximation close enough to  $\lambda^*$  by using

$$\mathcal{L}(x, e, k) = f(x) + 2k^{-1} \sum_{i=1}^m \ln 0.5(1 + e^{-k c_i(x)})$$

no matter how large  $k > 0$  we are ready to use. Therefore to guarantee convergence when  $k \rightarrow \infty$  we modified  $\mathcal{L}(x, e, k)$ . The convergence and the rate of convergence of the LS multipliers methods we consider in the next section.

## 6. Convergence and rate of convergence

We start with a modification of  $\mathcal{L}(x, e, k)$ . For a chosen  $0 < \alpha < 1$  we define the penalty LS function  $P : \mathfrak{R}^n \times \mathfrak{R}_{++} \rightarrow \mathfrak{R}_{++}$  by formula

$$P(x, k) = f(x) + 2k^{-1+\alpha} \sum_{i=1}^m \ln 0.5(1 + e^{-kc_i(x)}). \quad (6.1)$$

The existence and uniqueness of the minimizer

$$x(\cdot) = x(k) = \arg \min \{P(x, k) \mid x \in \mathfrak{R}^n\}$$

follows from lemmas 4.2 and 4.4. For the minimizer  $x(\cdot)$  we have

$$\nabla_x P(x(\cdot), \cdot) = \nabla f(x(\cdot)) - \sum_{i=1}^m 2k^\alpha (1 + e^{kc_i(x)})^{-1} \nabla c_i(x(\cdot)) = 0. \quad (6.2)$$

By introducing

$$\lambda_i(\cdot) \equiv \lambda_i(k) = 2k^\alpha (1 + e^{kc_i(x)})^{-1}, \quad i = 1, \dots, m, \quad (6.3)$$

we obtain

$$\nabla_x P(x(\cdot), \cdot) = \nabla f(x(\cdot)) - \sum_{i=1}^m \lambda_i(\cdot) \nabla c_i(x(\cdot)) = \nabla_x L(x(\cdot), \lambda(\cdot)) = 0. \quad (6.4)$$

The following theorem establishes the convergence of  $\{x(k)\}_{k>0}^\infty$  and  $\{\lambda(k)\}_{k>0}^\infty$  to  $X^*$  and  $L^*$ .

### Theorem 6.1.

- (1) If conditions (A) and (B) hold then the primal and dual trajectories  $\{x(k)\}_{k>0}^\infty$  and  $\{\lambda(k)\}_{k>0}^\infty$  are bounded and their limit points belong to  $X^*$  and  $L^*$ .
- (2) If the standard second order optimality conditions (2.4), (2.5) are satisfied then  $\lim_{k \rightarrow \infty} x(k) = x^*$  and  $\lim_{k \rightarrow \infty} \lambda(k) = \lambda^*$ . If, in addition,  $f(x)$  and all  $c_i(x) \in C^2$ , then the following bound

$$0 \leq f(x^*) - f(x(k)) \leq (0.5k^{-\alpha} f(x^*) + k^{-1}m \ln 2) \left( \sum_{i=1}^m \lambda^* \right) \quad (6.5)$$

holds true for any  $k \geq k_0$ , where  $k_0 > 0$  is large enough.

*Proof.* (1) As we mentioned in the proof of lemma 4.4 if  $X^*$  is bounded then by adding an extra constraint we can assume that the feasible set  $\Omega$  is bounded. Also due to lemma 4.2 the minimizer  $x(k)$  is unique, therefore  $\lambda(k)$  is uniquely defined by (6.3).

For the vector  $x(k)$  we define two sets of indexes  $I_+ = I_+(k) = \{i: c_i(x(k)) \geq 0\}$  and  $I_- = I_-(k) = \{i: c_i(x(k)) < 0\}$ . Then

$$\begin{aligned} P(x(k), k) &= P(\cdot, k) = f(\cdot) + 2k^{-1+\alpha} \left[ \sum_{i=1}^m \ln(1 + e^{-kc_i(\cdot)}) - m \ln 2 \right] \\ &= f(\cdot) + 2k^{-1+\alpha} \left[ \sum_{i \in I_+} \ln(1 + e^{-kc_i(\cdot)}) + \sum_{i \in I_-} \ln(1 + e^{kc_i(\cdot)}) - k \sum_{i \in I_-} c_i(\cdot) - m \ln 2 \right]. \end{aligned}$$

Therefore

$$P(\cdot, k) \geq f(\cdot) - 2k^\alpha \sum_{i \in I_-} c_i(\cdot) - 2k^{-1+\alpha} m \ln 2.$$

On the other hand

$$P(\cdot, k) \leq P(x^*, k) = f(x^*) + 2k^{-1+\alpha} \sum_{i=1}^m \ln 0.5(1 + e^{-kc_i(x^*)}).$$

In view of  $0.5(1 + e^{-kc_i(x^*)}) \leq 1$ ,  $i = 1, \dots, m$ , we have  $P(\cdot, k) \leq f(x^*)$ . Therefore keeping in mind  $f(x(k)) > 0$  we obtain

$$-2k^\alpha \sum_{i \in I_-} c_i(\cdot) \leq f(x^*) + 2k^{-1+\alpha} m \ln 2$$

or

$$\sum_{i \in I_-} |c_i(\cdot)| \leq 0.5k^{-\alpha} f(x^*) + k^{-1} m \ln 2.$$

In other words, for the maximum constraint violation at  $x(k)$  we have

$$\max_{i \in I_-} |c_i(x(k))| \leq 0.5k^{-\alpha} f(x^*) + k^{-1} m \ln 2 = v(k). \quad (6.6)$$

Therefore due to corollary 20 (see [7, p. 94]), the boundness  $\{x(k)\}_{k=0}^\infty$  follows from the boundness of  $\Omega$ .

The boundness of the dual trajectory  $\{\lambda(k)\}_{k=0}^\infty$  follows from Slater's condition (B), (6.4) and the boundness of the primal trajectory  $\{x(k)\}_{k=0}^\infty$ .

Let  $\{x(k_s)\}_{s=1}^\infty$  and  $\{\lambda(k_s)\}_{s=1}^\infty$  be the primal and dual converging subsequences and  $\bar{x} = \lim_{s \rightarrow \infty} x(k_s)$  and  $\bar{\lambda} = \lim_{s \rightarrow \infty} \lambda(k_s)$ , then by passing to the limit in (6.4) we obtain

$$\nabla_x L(\bar{x}, \bar{\lambda}) = \nabla f(\bar{x}) - \sum_{i=1}^m \bar{\lambda}_i \nabla c_i(\bar{x}) = 0$$

and from (6.3) and (6.6) we have  $\bar{\lambda} \in \mathfrak{R}_+^m$  and

$$c_i(\bar{x}) \geq 0, \quad i = 1, \dots, m, \quad \bar{\lambda}_i = 0, \quad i \notin I(\bar{x}) = \{i: c_i(\bar{x}) = 0\}$$

hence  $(\bar{x}, \bar{\lambda})$  is a KKT's pair, i.e.,  $\bar{x} \in X^*$ ,  $\bar{\lambda} \in \Lambda^*$ .

(2) If the standard second order optimality conditions (2.4), (2.5) are satisfied, then the pair  $(x^*, \lambda^*)$  is unique, therefore  $x^* = \lim_{k \rightarrow \infty} x(k)$  and  $\lambda^* = \lim_{k \rightarrow \infty} \lambda(k)$ .

To obtain the bound (6.5) we consider the enlarged feasible set  $\Omega(k) = \{x: c_i(x) \geq -v(k), i = 1, \dots, m\}$ , which is bounded because  $\Omega$  is bounded. Therefore  $f_k^* = \arg \min\{f(x) \mid x \in \Omega(k)\}$  exists and  $f_k^* \leq f(x(k))$ , hence  $f(x^*) - f(x(k)) \leq f(x^*) - f_k^*$ . Due to the conditions (2.4), (2.5) and keeping in mind that  $f(x), c_i(x) \in C^2$  we can use theorem 6 (see [7, p. 34]), to estimate  $f(x^*) - f_k^*$  for any  $k \geq k_0$  and  $k_0 > 0$  large enough we obtain

$$f(x^*) - f(x(k)) \leq f(x^*) - f_k^* \leq v(k) \sum_{i=1}^m \lambda_i^*.$$

Using (6.6) we obtain the bound (6.5).  $\square$

Convergence results for general classes of smoothing methods have been considered in [1].

Before we establish the rate of convergence for the LS multipliers method we would like to discuss one intrinsic property of the smoothing methods.

Let us consider the penalty LS function's Hessian. We have

$$\begin{aligned} H(x(\cdot), \cdot) &= \nabla_{xx}^2 P(x(\cdot), \cdot) \\ &= \nabla_{xx}^2 L(x(\cdot), \lambda(\cdot)) + 2k \nabla c(x(\cdot))^T \Lambda(\cdot) e^{kc(x(\cdot))} (I + e^{kc(x(\cdot))})^{-2} \nabla c(x(\cdot)), \end{aligned}$$

where  $\Lambda(\cdot) = \text{diag}(\lambda(\cdot))_{i=1}^m$  and  $e^{kc(x(\cdot))} = \text{diag}(e^{kc_i(x(\cdot))})_{i=1}^m$ . In view of (6.5) for  $k > 0$  large enough the pair  $(x(\cdot), \lambda(\cdot))$  is close to  $(x^*, \lambda^*)$ , therefore

$$H(x(\cdot), \cdot) \approx \nabla_{xx}^2 L(x^*, \lambda^*) + 0.5k \nabla c(x^*)^T \Lambda^* \nabla c(x^*).$$

Due to assertion 2.1 for  $k > 0$  large enough the min eigval  $H(x(\cdot), \cdot) = \mu > 0$ , while the max eigval  $H(x(\cdot), \cdot) = Mk, M > 0$ . Therefore

$$\text{cond } H(x(\cdot), \cdot) = \mu(Mk)^{-1} = O(k^{-1}).$$

Hence the  $\text{cond } H(x(\cdot), \cdot)$  converges to zero faster than  $f(x(k))$  converges to  $f(x^*)$ . The infinite increase of the scaling parameter  $k > 0$  is the only way to insure the convergence of the smoothing method. Therefore from some point on the smooth unconstrained minimization methods and in particular the Newton method might loose its efficiency.

The method (5.1), (5.2) allows to speed up the rate of convergence substantially and at the same time keeps stable the condition number of the LS Hessian.

Now we will prove that under the standard second order optimality conditions the primal-dual sequence  $\{x^s, \lambda^s\}$  generated by the LS multipliers method (5.1), (5.2) converges to the primal-dual solution with  $Q$ -linear rate under a fixed but large enough scaling parameter  $k > 0$ .

In our analysis we follow the scheme [11], in which the quadratic augmented Lagrangian proof (see [4, p. 109]) for equality constraints has been generalized for non-quadratic augmented Lagrangians applied to inequality constrained optimization.

In the course of our analysis we will estimate the threshold  $k_0 > 0$  for the scaling parameter when the  $Q$ -linear rates occurs. First, we specify the extended dual feasible domain in  $\mathfrak{R}_+^m \times [k_0, \infty)$ , where the  $Q$ -linear convergence takes place.

Let  $\|x\| = \|x\|_\infty = \max_{1 \leq i \leq n} |x_i|$ , we choose a small enough

$$0 < \delta < \min_{1 \leq i \leq r} \lambda_i^*.$$

We will split the dual optimal vector  $\lambda^*$  on active  $\lambda_{(r)}^* = (\lambda_1^*, \dots, \lambda_r^*) \in \mathfrak{R}_{++}^r$  and passive  $\lambda_{(m-r)}^* = (\lambda_{r+1}^*, \dots, \lambda_m^*) = 0$  parts. The neighborhood of  $\lambda^*$  we define as follows:

$$\begin{aligned} D(\cdot) &\equiv D(\lambda^*, k_0, \delta) \\ &= \{(\lambda, k) \in \mathfrak{R}_{++}^{m+1}: \lambda_i \geq \delta > 0, |\lambda_i - \lambda_i^*| \leq \delta k, i = 1, \dots, r, \\ &\quad 0 \leq \lambda_i \leq \delta k, k \geq k_0, i = r+1, \dots, m\}. \end{aligned}$$

By introducing the vector  $t = (t_i, i = 1, \dots, m) = (t_{(r)}, t_{(m-r)})$  with  $t_i = (\lambda_i - \lambda_i^*)k^{-1}$ ,  $k \geq k_0$ ,  $i = 1, \dots, m$ , we transform the dual set  $D(\cdot)$  into the neighborhood of the origin in the extended dual space

$$\begin{aligned} S(\delta, k_0) &= \{(t, k): (t_1, \dots, t_m; k): t_i \geq (\delta - \lambda_i^*)k^{-1}, i = 1, \dots, r, \\ &\quad t_i \geq 0, i = r+1, \dots, m, \|t\| \leq \delta, k \geq k_0\}. \end{aligned}$$

Then for LSL we obtain

$$\mathcal{L}(x, t, k) = f(x) + 2k^{-1} \sum_{i=1}^m (kt_i + \lambda_i^*) \ln 0.5(1 + e^{-kc_i(x)}).$$

For each  $\lambda \in D(\cdot)$  and  $k \geq k_0$  we can find the correspondent  $(t, k) \in S(\delta, k_0)$ , the minimizer

$$\hat{x} = \hat{x}(t, k) = \arg \min \{\mathcal{L}(x, t, k) \mid x \in \mathfrak{R}^n\}$$

and the new vector of the Lagrange multipliers

$$\hat{\lambda} = \hat{\lambda}(t, k) = (\hat{\lambda}_i(t, k) = 2(kt_i + \lambda_i^*)(1 + e^{kc_i(\hat{x})})^{-1}, i = 1, \dots, m).$$

Let us split the vector  $\hat{\lambda} = (\hat{\lambda}_{(r)}, \hat{\lambda}_{(m-r)})$  on the active  $\hat{\lambda}_{(r)} = (\hat{\lambda}_i, i = 1, \dots, r)$  and the passive

$$\hat{\lambda}_{(m-r)} \equiv \hat{\lambda}_{(m-r)}(\hat{x}, t, k) = (\lambda_i(\hat{x}, t, k) = 2kt_i(1 + e^{kc_i(\hat{x})})^{-1}, i = r+1, \dots, m)$$

parts. We consider the vector-function

$$h(x, t, k) = \sum_{i=r+1}^m \hat{\lambda}_i(x, t, k) \nabla c_i(x),$$

which correspond to the passive set of constraints.

Let

$$a \in \mathfrak{R}^n, \quad b \in \mathfrak{R}^n, \quad \theta(\tau): \mathfrak{R} \rightarrow \mathfrak{R}, \quad \theta(b) = (\theta(b_1), \dots, \theta(b_n)), \\ a\theta(b) = (a_1\theta(b_1), \dots, a_n\theta(b_n)), \quad a + \theta(b) = (a_1 + \theta(b_1), \dots, a_n + \theta(b_n)).$$

Now we are ready for the basic statement in this section. We would like to emphasize that results of the following theorem remain true if neither  $f(x)$  nor  $-c_i(x)$ ,  $i = 1, \dots, m$ , are convex.

**Theorem 6.2.** If  $f(x)$  and all  $c_i(x) \in C^2$  and the conditions (2.4) and (2.5) hold, then there exists such a small  $\delta > 0$  and large  $k_0 > 0$  that for any  $\lambda \in D(\cdot)$  and  $k \geq k_0$ :

- (1) there exist  $\hat{x} = \hat{x}(\lambda, k) = \arg \min\{\mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n\}$ :  $\nabla_x \mathcal{L}(\hat{x}, \lambda, k) = 0$  and  $\hat{\lambda} = \hat{\lambda}(\lambda, k) = 2\lambda(1 + e^{kc(\hat{x})})^{-1}$ ;
- (2) for the pair  $(\hat{x}, \hat{\lambda})$  the estimate

$$\max\{\|\hat{x} - x^*\|, \|\hat{\lambda} - \lambda^*\|\} \leq ck^{-1}\|\lambda - \lambda^*\| \quad (6.7)$$

holds and  $c > 0$  is independent on  $k \geq k_0$ ;

- (3) the LS function  $\mathcal{L}(x, \lambda, k)$  is strongly convex in a neighborhood of  $\hat{x}$ .

*Proof.* For any  $x \in \mathfrak{R}^n$ , any  $k > 0$  and  $t \in \mathfrak{R}^m$  the vector-function  $h(x, t, k)$  is smooth in  $x$ , also

$$h(x^*, 0, k) = 0 \in \mathfrak{R}^n, \quad \nabla_x h(x^*, 0, k) = 0^{n,n}, \quad \nabla_{\lambda_{(r)}} h(x^*, 0, k) = 0^{n,r},$$

where  $0^{p,q}$  is  $p \times q$  matrix with zero elements.

Let  $\sigma = \min\{c_i(x^*) \mid i = r+1, \dots, m\} > 0$ .

We consider the following map  $\Phi: \mathfrak{R}^{n+r+m+1} \rightarrow \mathfrak{R}^{n+r}$ , which is defined by

$$\Phi(x, \hat{\lambda}_{(r)}, t, k) = \begin{bmatrix} \nabla f(x) - \sum_{i=1}^r \hat{\lambda}_i \nabla c_i(x) - h(x, t, k) \\ 2k^{-1}(kt_{(r)} + \lambda_{(r)}^*)(1 + e^{kc_{(r)}(x)})^{-1} - k^{-1}\hat{\lambda}_{(r)} \end{bmatrix}. \quad (6.8)$$

Taking into account (2.1), (2.2) we obtain

$$\Phi(x^*, \lambda_{(r)}^*, 0, k) = 0^{n+r}, \quad \forall k > 0. \quad (6.9)$$

Let  $\nabla_{x\hat{\lambda}_{(r)}} \Phi = \nabla_{x\hat{\lambda}_{(r)}} \Phi(x^*, \lambda_{(r)}^*, 0, k)$ ,  $I^r$  – identical matrix in  $\mathfrak{R}^r$ ,

$$\Lambda_{(r)}^* = \text{diag}(\lambda_i^*)_{i=1}^r, \quad \nabla_{xx} L = \nabla_{xx} L(x^*, \lambda^*), \quad \nabla c = \nabla c(x^*), \quad \nabla c_{(r)} = \nabla c_{(r)}(x^*).$$

In view of  $\nabla_x h(x^*, 0, k) = 0^{n,n}$  and  $\nabla_{\hat{\lambda}_{(r)}} h(x^*, 0, k) = 0^{n,r}$  we obtain

$$\Phi_k \equiv \nabla_{x\hat{\lambda}_{(r)}} \Phi = \begin{bmatrix} \nabla_{xx} L & -\nabla c_{(r)}^T \\ -\frac{1}{2}\Lambda_{(r)}^* \nabla c_{(r)} & -k^{-1}I^r \end{bmatrix}.$$

Using reasoning similar to those in [11] we obtain that  $\Phi_k^{-1}$  exists and there is a number  $\varkappa > 0$ , which is independent on  $k \geq k_0$  that

$$\|\Phi_k^{-1}\| \leq \varkappa. \quad (6.10)$$

By applying the second implicit function theorem (see [3, p. 12]) to the map (6.8) we find that on the set

$$S(\delta, k_0, k_1) = \{(t, k): t_i \geq (\delta - \lambda_i^*)k^{-1}, i = 1, \dots, r, t_i \geq 0, i = r+1, \dots, m, \\ \|t\| \leq \delta k^{-1}, k_0 \leq k \leq k_1\}$$

there exists two vector-functions

$$x(\cdot) = x(t, k) = (x_1(t, k), \dots, x_m(t, k)) \quad \text{and} \\ \hat{\lambda}_{(r)}(\cdot) = \hat{\lambda}_{(r)}(t, k) = (\hat{\lambda}_1(t, k), \dots, \hat{\lambda}_r(t, k))$$

such that

$$\Phi(x(t, k), \hat{\lambda}_r(t, k), t, k) \equiv \Phi(x(\cdot), \hat{\lambda}(\cdot), \cdot) \equiv 0. \quad (6.11)$$

One can rewrite system (6.11) as follows:

$$\nabla f(x(\cdot)) - \sum_{i=1}^r \hat{\lambda}_i(\cdot) \nabla c_i(x(\cdot)) - h(x(\cdot), \cdot) = 0, \quad (6.12)$$

$$\hat{\lambda}_i(\cdot) = 2(kt_i + \lambda_i^*)(1 + e^{kc_i(x(\cdot))})^{-1}, \quad i = 1, \dots, r. \quad (6.13)$$

We also have  $\hat{\lambda}_i(\cdot) = 2\lambda_i(1 + e^{kc_i(\hat{x}(\cdot))})^{-1}$ ,  $i = r+1, \dots, m$ . Recalling that  $\lambda_{(m-r)}^* = (\lambda_{r+1}^*, \dots, \lambda_m^*) = 0 \in \mathfrak{R}^{m-r}$  we first estimate  $\|\hat{\lambda}_{(m-r)} - \lambda_{(m-r)}^*\|$ .

For any small  $\varepsilon > 0$  we can find  $\delta > 0$  small enough that  $\|x(t, k) - x^*(0, k)\| = \|x(\cdot) - x^*\| \leq \varepsilon$  for any  $t \in S(\delta, k_0)$ . Taking into account  $c_i(x^*) \geq \sigma > 0$  we obtain

$$c_i(x(t, k)) \geq \frac{\sigma}{2}, \quad i = r+1, \dots, m \text{ for } t \in S(\delta, k_0).$$

Therefore in view of  $e^x \geq x + 1$  we obtain

$$0 < \hat{\lambda}_i(t, k) \leq \frac{2\lambda_i}{1 + e^{0.5k\sigma}} \leq \frac{2\lambda_i}{2 + 0.5k\sigma} \leq \frac{4}{\sigma} \frac{\lambda_i}{k}. \quad (6.14)$$

Therefore

$$\|\hat{\lambda}_{(m-r)} - \lambda_{(m-r)}^*\| \leq \frac{4}{\sigma} k^{-1} \|\lambda_{(m-r)} - \lambda_{(m-r)}^*\|. \quad (6.15)$$

Now we will consider the vector-functions  $x(t, k) = x(\cdot)$  and  $\hat{\lambda}_{(r)}(t, k) = \hat{\lambda}_{(r)}(\cdot)$ . By differentiating (6.12) and (6.13) in  $t$  we find the Jacobians  $\nabla_t x(\cdot) \equiv \nabla_t x(t, k)$  and  $\nabla_t \hat{\lambda}_{(r)}(\cdot) = \nabla_t \hat{\lambda}_{(r)}(t, k)$  from the following system:

$$\begin{bmatrix} \nabla_t x(\cdot) \\ \nabla_t \hat{\lambda}_{(r)}(\cdot) \end{bmatrix} = (\nabla_{x\hat{\lambda}_{(r)}} \Phi(\cdot))^{-1} \begin{bmatrix} \nabla_t h(x(\cdot), \cdot) \\ -2 \operatorname{diag}((1 + e^{kc_i(x(\cdot))})^{-1})_{i=1}^r; 0^{r, m-r} \end{bmatrix}. \quad (6.16)$$

Considering the system (6.16) for  $t = 0 \in \mathfrak{R}^m$ , we obtain

$$\begin{bmatrix} \nabla_t x(0, k) \\ \nabla_t \hat{\lambda}_{(r)}(0, k) \end{bmatrix} = (\nabla_{x \hat{\lambda}_{(r)}} \Phi)^{-1} \begin{bmatrix} \nabla_t h(x(0, k), 0, k) \\ -I^r; \ 0^{r, m-r} \end{bmatrix} = (\Phi_k)^{-1} \begin{bmatrix} \nabla_t h(x(0, k), 0, k) \\ -I^r; \ 0^{r, m-r} \end{bmatrix}.$$

In view of (6.10) and estimation

$$\|\nabla_t h(x^*, 0, k)\| \leq 2k(1 + e^{k\sigma/2})^{-1} \|\nabla_{C(m-r)}(x^*)\| \leq 4\sigma^{-1} \|\nabla_{C(m-r)}(x^*)\|$$

which holds true for any  $k \geq k_0$ , we obtain

$$\begin{aligned} \max\{\|\nabla_t x(0, k)\|, \|\nabla_t \hat{\lambda}_{(r)}(0, k)\|\} &\leq \kappa(\|\nabla_{C(m-r)}(x^*)\| + \|I^r\|) \\ &= \kappa(4\sigma^{-1} \|\nabla_{C(m-r)}(x^*)\| + 1) = c_0. \end{aligned} \quad (6.17)$$

Therefore

$$\max\{\|\nabla_t x(t, k)\|, \|\nabla_t \hat{\lambda}_{(r)}(t, k)\|\} \leq 2c_0 \quad (6.18)$$

for any  $(t, k) \in S(\delta, k_0)$  and  $\delta > 0$  small enough.

Keeping in mind that  $x(0, k) = x^*$  and  $\hat{\lambda}_{(r)}(0, k) = \hat{\lambda}_{(r)}^*$  and using arguments similar to those in [11], we obtain

$$\max\{\|x(t, k) - x^*\|, \|\hat{\lambda}_{(r)}(t, k) - \hat{\lambda}_{(r)}^*\|\} \leq 2c_0 k^{-1} \|\lambda - \lambda^*\|. \quad (6.19)$$

Let

$$\hat{x}(\lambda, k) = x\left(\frac{\lambda - \lambda^*}{k}, k\right), \quad \hat{\lambda}(\lambda, k) = \left(\hat{\lambda}_{(r)}\left(\frac{\lambda - \lambda^*}{k}, k\right), \hat{\lambda}_{(m-r)}\left(\frac{\lambda - \lambda^*}{k}, k\right)\right)$$

then taking  $c = \max\{2c_0, 4/\sigma\}$  from (6.15) and (6.19) we obtain (6.7).

To prove the final part of the theorem we consider the LS Hessian  $\nabla_{xx} \mathcal{L}(x, \lambda, k)$  at the point  $\hat{x} = \hat{x}(\lambda, k)$ . We have

$$\nabla_x \mathcal{L}(x, \lambda, k) = \nabla f(x) - \sum_{i=1}^m \frac{2\lambda_i}{1 + e^{kc_i(x)}} \nabla c_i(x)$$

and for the Hessian  $\nabla_{xx}^2 \mathcal{L}(x, \lambda, k)$  we obtain

$$\begin{aligned} \nabla_{xx}^2 \mathcal{L}(x, \lambda, k) &= \nabla_{xx}^2 f(x) - \sum_{i=1}^m \frac{2\lambda_i}{1 + e^{kc_i(x)}} \nabla_{xx}^2 c_i(x) + 2k(\nabla c(x))^T (I + e^{kc(x)})^{-2} \Lambda \nabla c(x), \end{aligned}$$

where  $I$  – identical matrix in  $\mathfrak{R}^m$  and  $e^{kc(x)} = \text{diag}(e^{kc_i(x)})_{i=1}^m$ ,  $\Lambda = \text{diag}(\lambda_i)_{i=1}^m$ .

Therefore

$$\nabla_{xx}^2 \mathcal{L}(\hat{x}, \lambda, k) = \nabla_{xx}^2 L(\hat{x}, \hat{\lambda}) + k(\nabla c(\hat{x}))^T (I + e^{kc(\hat{x})})^{-1} \hat{\Lambda} \nabla c(\hat{x}).$$

Using the estimation (6.1) for any  $(\lambda, k) \in D(\cdot)$  we obtain

$$\begin{aligned}\nabla_{xx}^2 \mathcal{L}(\hat{x}, \lambda, k) &\approx \nabla_{xx}^2 L(x^*, \lambda^*) + k(\nabla c(x^*))^T (I + e^{kc(x^*)})^{-1} \Lambda^* \nabla c(x^*) \\ &= \nabla_{xx}^2 L(x^*, \lambda^*) + \frac{1}{2} k(\nabla c(x^*))^T \Lambda^* \nabla c(x^*).\end{aligned}$$

The strong convexity of LS  $\mathcal{L}(x, \lambda, k)$  in  $x$  in the neighborhood of  $\hat{x}$  follows from continuity of  $\nabla_{xx}^2 \mathcal{L}(x, \lambda, k)$  in  $x$  and assertion 2.1. The proof of theorem is completed.  $\square$

**Corollary 6.1.** The  $Q$ -linear rate of convergence for the method (5.1), (5.2) and  $Q$ -superlinear convergence for (5.3), (5.4) follows directly from the estimation (6.7) because  $c > 0$  is independent on  $k > k_0$ .

## 7. Modification of the LS method

The LS method (5.1), (5.2) requires solving unconstrained optimization problem at each step. To make the method practical we have to replace the unconstrained minimizer by an approximation that retains the convergence and the rate of convergence of LS method.

In this section we establish the conditions for the approximation and prove that such an approximation allows to retain for the modified LS method the rate of convergence (6.7).

For a given positive Lagrange multipliers vector  $\lambda \in \mathfrak{R}_{++}^m$ , a large enough penalty parameter  $k > 0$  and a positive scalar  $\tau > 0$  we find an approximation  $\tilde{x}$  for the primal minimizer  $\hat{x}$  from the inequality

$$\tilde{x} \in \mathfrak{R}^n: \|\nabla_x \mathcal{L}(\tilde{x}, \lambda, k)\| \leq \tau k^{-1} \|2(I + e^{kc(\tilde{x})})^{-1} \lambda - \lambda\| \quad (7.1)$$

and the approximation for the Lagrange multipliers by formula

$$\tilde{\lambda} = 2(I + e^{kc(\tilde{x})})^{-1} \lambda. \quad (7.2)$$

It leads to the following modification of the LS multipliers method (5.1), (5.2).

We define the modified primal–dual sequence by the following formulas:

$$\tilde{x}^{s+1} \in \mathfrak{R}^n: \|\nabla_x \mathcal{L}(\tilde{x}^{s+1}, \lambda^s, k)\| \leq \tau k^{-1} \|2(I + e^{kc_i(\tilde{x}^{s+1})})^{-1} \tilde{\lambda}^s - \tilde{\lambda}^s\|, \quad (7.3)$$

$$\tilde{\lambda}^{s+1} = 2(I + e^{kc(\tilde{x}^{s+1})})^{-1} \tilde{\lambda}^s. \quad (7.4)$$

It turns out that the modification (7.3), (7.4) of the LS method (5.1), (5.2) keeps the basic property of the LS method, namely the  $Q$ -linear rate of convergence as soon as the second order optimality conditions hold and the functions  $f(x)$  and  $c_i(x)$ ,  $i = 1, \dots, m$ , are smooth enough.

**Theorem 7.1.** If the standard second order optimality conditions (2.4), (2.5) hold and the Hessians  $\nabla^2 f(x)$  and  $\nabla^2 c_i(x)$ ,  $i = 1, \dots, m$ , satisfy the Lipschitz condition

$$\begin{aligned} \|\nabla^2 f(x_1) - \nabla^2 f(x_2)\| &\leq L_0 \|x_1 - x_2\|, \\ \|\nabla^2 c_i(x_1) - \nabla^2 c_i(x_2)\| &\leq L_i \|x_1 - x_2\| \end{aligned} \quad (7.5)$$

then there is  $k_0 > 0$  that for any  $\lambda \in D(\cdot)$  and  $k \geq k_0$  the following bound holds true:

$$\max\{\|\tilde{x} - x^*\|, \|\tilde{\lambda} - \lambda^*\|\} \leq c(5 + \tau)k^{-1}\|\lambda - \lambda^*\| \quad (7.6)$$

and  $c > 0$  is independent on  $k \geq k_0$ .

*Proof.* Let us assume that  $\varepsilon > 0$  is small enough and

$$\begin{aligned} \tilde{x} &\in S(x^*, \varepsilon) = \{x \in \mathfrak{R}^n: \|x - x^*\| \leq \varepsilon\}, \\ \tilde{\lambda} &= (1 + e^{kc(\tilde{x})})^{-1}\lambda \in S(\lambda^*, \varepsilon) = \{\lambda \in \mathfrak{R}_{++}^m: \|\lambda - \lambda^*\| \leq \varepsilon\}. \end{aligned}$$

We consider vectors

$$\begin{aligned} \Delta x &= \tilde{x} - x^*, \quad \Delta \lambda = \tilde{\lambda} - \lambda^* = (\Delta \lambda_{(r)}, \Delta \lambda_{(m-r)}), \quad \Delta \lambda_{(r)} = \tilde{\lambda}_{(r)} - \lambda_{(r)}^*, \\ \Delta \lambda_{(m-r)} &= \tilde{\lambda}_{(m-r)} - \lambda_{(m-r)}^* = \lambda_{(m-r)}, \quad \Delta y_{(r)} = (\Delta x, \Delta \lambda_{(r)}) \quad \text{and} \quad \Delta y = (\Delta x, \Delta \lambda). \end{aligned}$$

Due to (7.5) we have

$$\nabla f(\tilde{x}) = \nabla f(x^*) + \nabla^2 f(x^*)\Delta x + r_0(\Delta x), \quad (7.7)$$

$$\nabla c_i(\tilde{x}) = \nabla c_i(x^*) + \nabla^2 c_i(x^*)\Delta x + r_i(\Delta x), \quad i = 1, \dots, m, \quad (7.8)$$

and  $r_0(0) = 0$ ,  $r_i(0) = 0$ . Also due to (7.5) we have  $\|\nabla r_0(\Delta x)\| \leq L_0 \|\Delta x\|$ ,  $\|\nabla r_i(\Delta x)\| \leq L_i \|\Delta x\|$ . Then

$$\begin{aligned} \nabla_x \mathcal{L}(\tilde{x}, \lambda, k) &= \nabla f(\tilde{x}) - \sum_{i=1}^m \frac{2\lambda_i}{1 + e^{kc_i(\tilde{x})}} \nabla c_i(\tilde{x}) \\ &= \nabla f(\tilde{x}) - \sum_{i=1}^m \tilde{\lambda}_i \nabla c_i(\tilde{x}) \\ &= \nabla f(\tilde{x}) - \sum_{i=1}^r (\Delta \lambda_i + \lambda_i^*) \nabla c_i(\tilde{x}) + h(\tilde{x}, \lambda_{(m-r)}, k), \end{aligned}$$

where  $h(\tilde{x}, \lambda_{(m-r)}, k) = \sum_{i=r+1}^m 2\lambda_i (1 + e^{kc_i(\tilde{x})})^{-1} \nabla c_i(\tilde{x})$ .

Using (7.7), (7.8) we obtain

$$\begin{aligned} \nabla_x \mathcal{L}(\tilde{x}, \lambda, k) &= \nabla f(x^*) + \nabla^2 f(x^*)\Delta x + r_0(\Delta x) \\ &\quad - \sum_{i=1}^r (\Delta \lambda_i + \lambda_i^*) (\nabla c_i(x^*) + \nabla^2 c_i(x^*)\Delta x + r_i(\Delta x)) + h(\tilde{x}, \lambda_{(m-r)}, k) \end{aligned}$$

$$\begin{aligned}
&= \nabla f(x^*) - \sum_{i=1}^r \lambda_i^* \nabla c_i(x^*) + \left( \nabla^2 f(x^*) - \sum_{i=1}^r \lambda_i^* \nabla^2 c_i(x^*) \right) \Delta x \\
&\quad - \sum_{i=1}^r \Delta \lambda_i \nabla c_i(x^*) + r_0(\Delta x) - \sum_{i=1}^r \Delta \lambda_i \nabla^2 c_i(x^*) \Delta x \\
&\quad - \sum_{i=1}^r (\Delta \lambda_i + \lambda_i^*) r_i(\Delta x) + h(\tilde{x}, \lambda_{(m-r)}, k).
\end{aligned}$$

Let

$$r^{(1)}(\Delta y) = r_0(\Delta x) - \sum_{i=1}^r \Delta \lambda_i \nabla^2 c_i(x^*) \Delta x + \sum_{i=1}^r (\Delta \lambda_i + \lambda_i^*) r_i(\Delta x)$$

then keeping in mind the KKT's condition we can rewrite the expression above as

$$\nabla_x \mathcal{L}(\tilde{x}, \lambda, k) = \nabla_{xx} \mathcal{L}(x^*, \lambda^*) \Delta x - \nabla_{c(r)}(x^*)^T \Delta \lambda_{(r)} + h(\tilde{x}, \lambda_{(m-r)}, k) + r^{(1)}(\Delta y), \quad (7.9)$$

where  $r^{(1)}(0) = 0$  and there is  $L^{(1)} > 0$  that  $\|\nabla_{r_x}(\Delta y)\| \leq L^{(1)} \|\Delta y\|$ .

Then  $\Delta \lambda_i = \tilde{\lambda}_i - \lambda_i^* = \tilde{\lambda}_i - \lambda_i + \lambda_i - \lambda_i^*$ , i.e.,  $(\tilde{\lambda}_i - \lambda_i) - \Delta \lambda_i = \lambda_i^* - \lambda_i$ ,  $i = 1, \dots, r$ , or

$$\Lambda_{(r)} e_{(r)}(\tilde{x}, k) - \Delta \lambda_{(r)} = \lambda_{(r)}^* - \lambda_{(r)}, \quad (7.10)$$

where  $e_{(r)}^T(\tilde{x}, k) = (e_1(\tilde{x}, k), \dots, e_r(\tilde{x}, k))$  and  $e_i(\tilde{x}, k) = (1 - e^{kc_i(x)})(1 + e^{kc_i(x)})^{-1}$ ,  $i = 1, \dots, r$ .

Further,

$$\begin{aligned}
e_i(\tilde{x}, k) &= e_i(x^*, k) + k \nabla e_i^T(x^*, k) \Delta x + r_i^e(\Delta x), \quad i = 1, \dots, r, \quad \text{and} \\
r_i^e(0) &= 0, \quad \|\nabla r_i^e(\Delta x)\| \leq L_i^e \|\Delta x\|.
\end{aligned}$$

In view of  $e_i(x^*, k) = 0$  and

$$\nabla e_i(x^*, k) = -2k e^{2kc_i(x^*)} (1 + e^{kc_i(x^*)})^{-2} \nabla c_i(x^*) = -\frac{k}{2} \nabla c_i(x^*), \quad i = 1, \dots, r,$$

we have

$$e_i(\tilde{x}, k) = -\frac{1}{2} k \nabla c_i(x^*) + r_i^e(\Delta x), \quad i = 1, \dots, r. \quad (7.11)$$

Therefore the system (7.10) can be rewritten as

$$\Lambda_{(r)}^* e_{(r)}(\tilde{x}, k) - \Delta \lambda_{(r)} = (I^r + E_{(r)}(\tilde{x}, k))(\lambda_{(r)}^* - \lambda_{(r)}),$$

where  $E_{(r)}(x, k) = \text{diag}(e_i(x, k))_{i=1}^r$ . Using (7.11) we obtain

$$\begin{aligned}
&-\frac{1}{2} \Lambda_{(r)}^* \nabla c_{(r)}(x^*) \Delta x - k^{-1} \Delta \lambda_{(r)} \\
&= k^{-1} (I^r + E_{(r)}(\tilde{x}, k))(\lambda_{(r)}^* - \lambda_{(r)}) - k^{-1} \Lambda_{(r)}^* r_{(r)}^e(\Delta x)
\end{aligned}$$

or

$$-\Lambda_{(r)}^* \nabla c_{(r)}(x^*) \Delta x - 2k^{-1} \Delta \lambda_{(r)} = 2k^{-1} (I^r + E_{(r)}(\tilde{x}, k)) (\lambda_{(r)}^* - \lambda_{(r)}) - r_{\lambda}(\Delta x), \quad (7.12)$$

where

$$r^{(2)}(\Delta x) = 2k^{-1} \Lambda_{(r)}^* r_{(r)}^e(\Delta x), \quad r_{(r)}^e(\Delta x)^T = (r_1^e(\Delta x), \dots, r_r^e(\Delta x)),$$

$r^{(2)}(0) = 0$ , and there is  $L^{(2)} > 0$  that  $\|\nabla r^{(2)}(\Delta x)\| \leq L^{(2)} \|\Delta x\|$ .

Combining (7.9) and (7.12) we obtain

$$\begin{aligned} & \nabla_{xx}^2 L \Delta x - \nabla c_{(r)} \Delta \lambda_{(r)} \\ &= \nabla_x \mathcal{L}(\tilde{x}, \lambda, k) - h(\tilde{x}, \lambda_{(m-r)}, k) - r_x(\Delta y) - \Lambda_{(r)}^* \nabla c_{(r)} \Delta x - 2k^{-1} \Delta \lambda_{(r)} \\ &= 2k^{-1} (I^r + E_{(r)}(\tilde{x}, k)) (\lambda_{(r)}^* - \lambda_{(r)}) - r^{(2)}(\Delta x) \end{aligned}$$

or

$$\bar{\varphi}_k \Delta y_{(r)} = a(\tilde{x}, \lambda, k) + b(\tilde{x}, \lambda, k) + r(\Delta y_{(r)}), \quad (7.13)$$

where

$$\begin{aligned} \bar{\varphi}_k &= \begin{bmatrix} \nabla_{xx} L & -\nabla c_{(r)} \\ -\Lambda_{(r)}^* \nabla c_{(r)} & -2k^{-1} \end{bmatrix}, \quad a(\tilde{x}, \lambda, k) = \begin{bmatrix} \nabla_x \mathcal{L}(\tilde{x}, \lambda, k) - h(\tilde{x}, \lambda_{(m-r)}, k) \\ 0 \end{bmatrix}, \\ b(\tilde{x}, \lambda_{(r)}, k) &= \begin{bmatrix} 0 \\ 2k^{-1} (I^r + E_{(r)}(\tilde{x}, k)) (\lambda_{(r)} - \lambda_{(r)}^*) \end{bmatrix}, \quad r(\Delta y_{(r)}) = \begin{bmatrix} -r^{(1)}(\Delta y_{(r)}) \\ -r^{(2)}(\Delta x) \end{bmatrix}, \end{aligned}$$

$r(0) = 0$ , and there is  $L > 0$  that  $\|\nabla r(\Delta y_{(r)})\| \leq L \|\Delta y_{(r)}\|$ .

As we know already for  $k_0 > 0$  large enough and any  $k \geq k_0$  the inverse matrix  $\bar{\varphi}_k^{-1}$  exists and there is  $\bar{\kappa} > 0$  independent on  $k \geq k_0$  that  $\|\bar{\varphi}_k^{-1}\| \leq \bar{\kappa}$ . Therefore we can solve the system (7.13) for  $\Delta y_{(r)}$ .

$$\begin{aligned} \Delta y_{(r)} &= \bar{\varphi}_k^{-1} [a(\tilde{x}, \lambda, k) + b(\tilde{x}, \lambda, k) + r(\Delta y)] = \bar{\varphi}_k^{-1} [a(\cdot) + b(\cdot) + r(\Delta y_{(r)})] \\ &= C(\Delta y_{(r)}). \end{aligned} \quad (7.14)$$

Therefore

$$\nabla C(\Delta y_{(r)}) = \bar{\varphi}_k^{-1} \nabla r(\Delta y_{(r)})$$

and

$$\|\nabla C(\Delta y_{(r)})\| \leq \|\bar{\varphi}_k^{-1}\| \|\nabla r(\Delta y_{(r)})\| \leq \bar{\kappa} L \|\Delta y_{(r)}\|.$$

So, for  $\|\Delta y_{(r)}\|$  small enough we have

$$\|\nabla C(\Delta y_{(r)})\| \leq q < 1.$$

In other words the operator  $C(\Delta y_{(r)})$  is a contractive operator for  $\|\Delta y_{(r)}\|$  small enough.

Let us estimate the contractibility at the operator  $C(\Delta y_{(r)})$  with more details. First of all we shall estimate  $\|a(\cdot)\|$  and  $\|b(\cdot)\|$ . We have

$$\|a(\cdot)\| \leq \|\nabla_x \mathcal{L}(\tilde{x}, \lambda, k)\| + \|h(\tilde{x}, \lambda_{(m-r)}, k)\|.$$

Note that for  $\varepsilon > 0$  small enough and  $\tilde{x} \in S(x^*, \varepsilon)$  we obtain

$$\tilde{\lambda}_i = 2\lambda_i(1 + e^{kc_i(\tilde{x})})^{-1} \leq 4\lambda_i(1 + e^{kc_i(x^*)})^{-1} \leq 4\lambda_i(1 + e^{k\sigma/2})^{-1}, \quad i = r+1, \dots, m.$$

Hence for  $k_0 > 0$  large enough and  $k \geq k_0$  one can find a small enough  $\eta > 0$  that

$$\tilde{\lambda}_i \leq \eta k^{-1} \lambda_i = \eta k^{-1} (\lambda_i - \lambda_i^*), \quad i = r+1, \dots, m,$$

i.e.,

$$\|\Delta \lambda_{(m-r)}\| = \|\tilde{\lambda}_{(m-r)} - \lambda_{(m-r)}^*\| \leq \eta k^{-1} \|\lambda_{(m-r)} - \lambda_{(m-r)}^*\| \quad (7.15)$$

and

$$\begin{aligned} \|h(\tilde{x}, \lambda_{(m-r)}, k)\| &= \left\| \sum_{i=r+1}^m 2\lambda_i(1 + e^{kc_i(\tilde{x})})^{-1} \nabla c_i(\tilde{x}) \right\| \\ &\leq \sum_{i=r+1}^m 8\lambda_i(1 + e^{k\sigma/2})^{-1} \|\nabla c_i(x^*)\|. \end{aligned}$$

For  $k_0 > 0$  large enough and any  $k \geq k_0$  we have  $4(1 + e^{k\sigma/2})^{-1} \|\nabla c_i(x^*)\| \leq 2k^{-1}$ ,  $i = r+1, \dots, m$ . Therefore

$$\|h(\tilde{x}, \lambda_{(m-r)}, k)\| \leq 2k^{-1} \|\lambda_{(m-r)} - \lambda_{(m-r)}^*\| \leq 2k^{-1} \|\lambda - \lambda^*\|.$$

Further, from (7.11) we obtain

$$\|\nabla_x \mathcal{L}(\tilde{x}, \lambda, k)\| \leq \tau k^{-1} \|\tilde{\lambda} - \lambda\| \leq \tau k^{-1} \|\tilde{\lambda} - \lambda^*\| + \tau k^{-1} \|\lambda - \lambda^*\|. \quad (7.16)$$

Therefore

$$\|a(\cdot)\| \leq \tau k^{-1} \|\tilde{\lambda} - \lambda^*\| + (2 + \tau)k^{-1} \|\lambda - \lambda^*\|. \quad (7.17)$$

Further

$$I^r + E_{(r)}(\tilde{x}, k) = \text{diag}\left(1 + \frac{1 - e^{kc_i(\tilde{x})}}{1 + e^{kc_i(\tilde{x})}}\right) = \text{diag}(2(1 + e^{kc_i(\tilde{x})})^{-1})_{i=1}^r.$$

Hence

$$\|b(\cdot)\| \leq 2k^{-1} \|\lambda_{(r)} - \lambda_{(r)}^*\| \leq 2k^{-1} \|\lambda - \lambda^*\|. \quad (7.18)$$

From (7.14), (7.15) we obtain

$$\begin{aligned} \|\Delta y_{(r)}\| &\leq \|\bar{\varphi}_k^{-1} [\|a(\cdot)\| + \|b(\cdot)\| + \|r(\Delta y_{(r)})\|]\| \\ &\leq \bar{x}(\tau k^{-1} \|\tilde{\lambda}_{(r)} - \lambda_{(r)}^*\| + \tau k^{-1} \|\Delta \lambda_{(m-r)}\| + (4 + \tau)k^{-1} \|\lambda - \lambda^*\| \\ &\quad + \|r(\Delta y_{(r)})\|). \end{aligned}$$

Then in view of

$$\|\tilde{\lambda}_{(r)} - \lambda_{(r)}^*\| \leq \|\Delta y_{(r)}\|, \quad \|\Delta \lambda_{(m-r)}\| \leq k^{-1} \|\lambda_{(m-r)} - \lambda_{(m-r)}^*\| \leq k^{-1} \|\lambda - \lambda^*\|$$

and  $\|r(\Delta y_{(r)})\| \leq \frac{L}{2} \|\Delta y_{(r)}\|^2$  we obtain

$$\|\Delta y_{(r)}\| \leq \bar{\kappa} \left( \tau k^{-1} \|\Delta y_{(r)}\| + (5 + \tau) k^{-1} \|\lambda - \lambda^*\| + \frac{L}{2} \|\Delta y_{(r)}\|^2 \right)$$

or

$$\frac{\bar{\kappa} L}{2} \|\Delta y_{(r)}\|^2 - (1 - \bar{\kappa} \tau k^{-1}) \|\Delta y_{(r)}\| + (5 + \tau) \bar{\kappa} k^{-1} \|\lambda - \lambda^*\| \geq 0$$

and

$$\|\Delta y_{(r)}\| \leq \frac{1}{\bar{\kappa} L} \left[ \left( 1 - \frac{\bar{\kappa} \tau}{k} \right) - \left( \left( 1 - \frac{\bar{\kappa} \tau}{k} \right)^2 - \frac{2L\bar{\kappa}^2}{k} (5 + \tau) \|\lambda - \lambda^*\| \right)^{1/2} \right].$$

If  $k_0 > 0$  is large enough then for any  $k \geq k_0$  we have

$$\left[ \left( 1 - \frac{\bar{\kappa} \tau}{k} \right)^2 - \frac{2L\bar{\kappa}^2(5 + \tau)}{k} \|\lambda - \lambda^*\| \right]^{1/2} \geq \left( 1 - \frac{\bar{\kappa} \tau}{k} \right) - \frac{2L\bar{\kappa}^2(5 + \tau)}{k} \|\lambda - \lambda^*\|.$$

Therefore

$$\|\Delta y_{(r)}\| \leq \frac{2\bar{\kappa}(5 + \tau)}{k} \|\lambda - \lambda^*\|.$$

So in view of (7.14) for  $c = \max\{2\bar{\kappa}, \eta\}$  we have

$$\max\{\|\tilde{x} - x^*\|, \|\tilde{\lambda} - \lambda^*\|\} \leq \frac{c(5 + \tau)}{k} \|\lambda - \lambda^*\|.$$

The proof is completed.  $\square$

*Remark 7.1.* The results of theorem 7.1 remain true whenever  $f(x)$  and all  $-c_i(x)$  are convex or not.

## 8. Primal–dual LS method

The numerical realization of the LS method (5.1), (5.2) leads to finding an approximation  $\tilde{x}$  from (7.1) and updating the Lagrange multipliers by formula (7.2). To find  $\tilde{x}$  one can use Newton method. The Newton LS method has been described in [12]. In this section we consider another approach to numerical realization of the LS multipliers method (5.1), (5.2). Instead of using Newton method to find  $\tilde{x}$  and then to update the Lagrange multipliers we will use Newton method for solving the following primal–dual system

$$\nabla_x L(\hat{x}, \hat{\lambda}) = \nabla f(\hat{x}) - \sum \hat{\lambda}_i \nabla c_i(\hat{x}) = 0, \quad (8.1)$$

$$\hat{\lambda} = \psi'(kc(\hat{x}))\lambda \quad (8.2)$$

for  $\hat{x}$  and  $\hat{\lambda}$  under the fixed  $k > 0$  and  $\lambda \in \mathfrak{R}_{++}^m$ , where  $\psi'(kc(\hat{x})) = \text{diag}(\psi'(kc_i(\hat{x})))_{i=1}^m$ . After finding an approximation  $(\tilde{x}, \tilde{\lambda})$  for the primal–dual pair  $(\hat{x}, \hat{\lambda})$  we replace  $\lambda$  for  $\tilde{\lambda}$  and take  $(\tilde{x}, \tilde{\lambda})$  as a starting point for the new system.

We apply the Newton method for solving (8.1), (8.2) using  $(x, \lambda)$  as a starting point. By linearizing (8.1), (8.2) we obtain the following system for the Newton direction  $(\Delta x, \Delta \lambda)$

$$\nabla f(x) + \nabla^2 f(x) \Delta x - \sum_{i=1}^m (\lambda_i + \Delta \lambda_i) (\nabla c_i(x) + \nabla^2 c_i(x) \Delta x) = 0, \quad (8.3)$$

$$\lambda + \Delta \lambda = \psi'(k(c(x) + \nabla c(x) \Delta x)) \lambda. \quad (8.4)$$

The system (8.3) we can rewrite as follows:

$$\nabla_{xx}^2 L(x, \lambda) \Delta x - \nabla c(x)^T \Delta \lambda + \nabla_x L(x, \lambda) - \sum_{i=1}^m \Delta \lambda_i \nabla^2 c_i(x) \Delta x = 0. \quad (8.5)$$

By ignoring the last term we obtain

$$\nabla_{xx}^2 L(x, \lambda) \Delta x - \nabla c(x)^T \Delta \lambda = -\nabla_x L(x, \lambda). \quad (8.6)$$

By ignoring the second order components in the representation of  $\psi'(kc(x) + k\nabla c(x) \Delta x)$  we obtain

$$\psi'(kc(x) + k\nabla c(x) \Delta x) = \psi'(kc(x)) + k\psi''(kc(x)) \Delta c(x) \Delta x.$$

So we can rewrite the system (8.4) as follows

$$-k\psi''(kc(x)) \Lambda \nabla c(x) \Delta x + \Delta \lambda = \bar{\lambda} - \lambda, \quad (8.7)$$

where  $\psi''(kc(x)) = \text{diag}(\psi''(kc_i(x)))_{i=1}^m$ ,  $\Lambda = \text{diag}(\lambda_i)_{i=1}^m$ ,  $\bar{\lambda} = \psi'(kc(x)) \lambda$ . Combining (8.6), (8.7) we obtain

$$\nabla_{xx}^2 L(x, \lambda) \Delta x - \nabla c(x)^T \Delta \lambda = -\nabla_x L(x, \lambda), \quad (8.8)$$

$$-\nabla c(\cdot) \Delta x + (k\psi''(kc(x)) \Lambda)^{-1} \Delta \lambda = (k\psi''(kc(x)) \Lambda)^{-1} (\bar{\lambda} - \lambda). \quad (8.9)$$

We find  $\Delta \lambda$  from (8.9) and substitute to (8.8). We obtain the following system for  $\Delta x$

$$M(\cdot) \Delta x = -\nabla_x L(x, \bar{\lambda}), \quad (8.10)$$

where

$$M(x, \lambda) = \nabla_{xx}^2 L(x, \lambda) - k\nabla c(x)^T \psi''(kc(x)) \Lambda \nabla c(x)$$

is a symmetric and positive definite matrix. Moreover, the cond  $M(x, \lambda, k)$  is stable in the neighborhood of  $(x^*, \lambda^*)$  for any fixed  $k \geq k_0$ , see lemma 4.3. By solving the system (8.10) for  $\Delta x$  we find the primal direction. The next primal approximation  $\tilde{x}$  for  $\hat{x}$  we obtain as

$$\tilde{x} = x + \Delta x. \quad (8.11)$$

The next dual approximation for  $\hat{\lambda}$  we find by formula

$$\tilde{\lambda} = \lambda + \Delta\lambda = \bar{\lambda} + k\psi''(kc(x))\Lambda\nabla c(x)\Delta x. \quad (8.12)$$

So the method (8.11), (8.12) one can view as dual–primal predictor–corrector. The vector  $\bar{\lambda} = \psi'(kc(\hat{x}))\lambda$ , is the predictor for  $\hat{\lambda}$ . Using this dual predictor we find the primal direction  $\Delta x$  from (8.10) and use it to find the dual corrector  $\Delta\lambda$  by formula

$$\Delta\lambda = k\psi''(kc(x))\Lambda\nabla c(x)\Delta x.$$

The next dual approximation  $\tilde{\lambda}$  for  $\hat{\lambda}$  we find by (8.12).

The primal–dual LS is fast and numerically stable in the neighborhood of  $(x^*, \lambda^*)$ . To make the LS method converge globally one can combine the primal–dual method with Newton LS using the scheme [10]. Such approach produced very encouraging results on a number of LP and NLP problems [12]. We will show some recently obtained results in section 10.

## 9. Log-sigmoid Lagrangian and duality

We have seen already that LS Lagrangian  $\mathcal{L}(x, \lambda, k)$  has some important properties, which the classical Lagrangian  $L(x, \lambda)$  does not possess. Therefore one can expect that the dual function

$$d_k(\lambda) = \inf\{\mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n\} \quad (9.1)$$

and the dual problem

$$\lambda^* = \arg \max\{d_k(\lambda) \mid \lambda \in \mathfrak{R}_+^m\} \quad (9.2)$$

might have some extra properties as compared to the dual function  $d(\lambda)$  and the dual problem (D), which are based on  $L(x, \lambda)$ .

First of all due to the lemmas 4.2 and 4.4 for any  $\lambda \in \mathfrak{R}_{++}^m$  and any  $k > 0$  there exists a unique minimizer

$$\hat{x} = \hat{x}(\lambda, k) = \arg \min\{\mathcal{L}(x, \lambda, k) \mid x \in \mathfrak{R}^n\}$$

for any convex programming problem with a bounded optimal set.

The uniqueness of the minimizer  $\hat{x}(\lambda, k)$  together with smoothness of  $f(x)$  and  $c_i(x)$  provide smoothness for the dual function  $d_k(\lambda)$  which is always concave whether the primal problem (P) is convex or not.

So the dual function  $d_k(\lambda)$  is smooth under reasonable assumption about the primal problem (P). Also the dual problem (9.2) is always convex.

Let us consider the properties of the dual function and the dual problem (9.2) with more details.

Assuming smoothness  $f(x)$  and  $c_i(x)$  and uniqueness  $\hat{x}(\lambda, k)$  we can compute the gradient  $\nabla d_k(\lambda)$  and the Hessian  $\nabla^2 d_k(\lambda)$ . For the gradient  $\nabla d_k(\lambda)$  we obtain

$$\nabla d_k(\lambda) = \nabla_x \mathcal{L}(\hat{x}, \lambda, k) \nabla_\lambda \hat{x}(\lambda, k) + \nabla_\lambda \mathcal{L}(x, \lambda, k),$$

where  $\nabla_\lambda \hat{x}(\lambda, k) = J_\lambda(\hat{x}(\lambda, k))$  is the Jacobian of the vector-function  $\hat{x}(\lambda, k)$ .

In view of  $\nabla_x \mathcal{L}(\hat{x}, \lambda, k) = 0$  we have

$$\begin{aligned} \nabla d_k(\lambda) &= \nabla_\lambda \mathcal{L}(\hat{x}(\lambda, k), \lambda, k) = \nabla_\lambda \mathcal{L}(\hat{x}(\cdot), \cdot) \\ &= 2k^{-1} (\ln 0.5(1 + e^{-kc_1(\hat{x})}), \dots, \ln 0.5(1 + e^{-kc_m(\hat{x})}))^T \\ &= 2k^{-1} \ln 0.5(1 + e^{-kc(\hat{x})}), \end{aligned} \quad (9.3)$$

where  $\ln 0.5(1 + e^{-kc(x)})$  is a column vector with components  $\ln 0.5(1 + e^{-kc_i(x)})$ ,  $i = 1, \dots, m$ .

Since  $\nabla_{xx}^2 \mathcal{L}(\hat{x}(\lambda, k), \lambda, k)$  is positive definite the system

$$\nabla_x \mathcal{L}(x, \lambda, k) = 0$$

yields a unique vector-function  $\hat{x}(\lambda, k)$  such that  $\hat{x}(\lambda^*, k) = x^*$  and

$$\nabla_x \mathcal{L}(\hat{x}(\lambda, k), \lambda, k) \equiv \nabla_x \mathcal{L}(\hat{x}(\cdot), \cdot) = 0. \quad (9.4)$$

By differentiating (9.4) in  $\lambda$  we obtain

$$\nabla_{xx}^2 \mathcal{L}(\hat{x}(\cdot), \cdot) \nabla_\lambda \hat{x}(\cdot) + \nabla_{x\lambda}^2 \mathcal{L}(\hat{x}(\cdot), \cdot) = 0$$

therefore

$$\nabla_\lambda \hat{x}(\lambda, k) = \nabla_\lambda \hat{x}(\cdot) = -(\nabla_{xx}^2 \mathcal{L}(\hat{x}(\cdot), \cdot))^{-1} \nabla_{x\lambda}^2 \mathcal{L}(\hat{x}(\cdot), \cdot). \quad (9.5)$$

Let us consider the Hessian for the dual function. Using (9.3) and (9.5) we obtain

$$\begin{aligned} \nabla^2 d_k(\lambda) &= \nabla_\lambda (\nabla_\lambda d_k(\lambda)) = 2k^{-1} \nabla_\lambda \ln 0.5(1 + e^{-kc(\hat{x}(\lambda, k))}) = \nabla_{\lambda x}^2 \mathcal{L}(\hat{x}, \lambda, k) \nabla_\lambda \hat{x}(\lambda, k) \\ &= -\nabla_{\lambda x}^2 \mathcal{L}(\hat{x}(\cdot), \cdot) (\nabla_{xx}^2 \mathcal{L}(\hat{x}(\cdot), \cdot))^{-1} \nabla_{x\lambda}^2 \mathcal{L}(\hat{x}(\cdot), \cdot). \end{aligned} \quad (9.6)$$

To compute  $\nabla_{\lambda x} \mathcal{L}(\hat{x}(\cdot), \cdot)$  let us consider

$$\nabla_\lambda \mathcal{L}(x(\cdot), \cdot) = 2k^{-1} \begin{bmatrix} \ln(1 + e^{-kc_1(\hat{x}(\cdot))}) - \ln 2 \\ \vdots \\ \ln(1 + e^{-kc_m(\hat{x}(\cdot))}) - \ln 2 \end{bmatrix}.$$

Then for the Jacobian  $\nabla_x (\nabla_\lambda \mathcal{L}(x(\cdot), \cdot)) = \nabla_{\lambda x}^2 \mathcal{L}(x(\cdot), \cdot)$  we obtain

$$\nabla_{\lambda x}^2 \mathcal{L}(\hat{x}(\cdot), \cdot) = -2 \begin{bmatrix} (1 + e^{kc_1(\hat{x}(\cdot))})^{-1} \nabla c_1(\hat{x}(\cdot)) \\ \vdots \\ (1 + e^{kc_m(\hat{x}(\cdot))})^{-1} \nabla c_m(\hat{x}(\cdot)) \end{bmatrix} = \psi'(kc(\hat{x}(\cdot))) \nabla c(\hat{x}(\cdot)),$$

where  $\psi'(kc(\hat{x}(\cdot))) = -2 \text{diag}[(1 + e^{kc_i(\hat{x}(\cdot))})^{-1}]_{i=1}^m$ . Therefore

$$\nabla_{x\lambda}^2 \mathcal{L}(\hat{x}(\cdot), \cdot) = \nabla_{\lambda x}^T \mathcal{L}(\hat{x}(\cdot), \cdot) = \nabla c(\hat{x}(\cdot))^T \psi'(kc(\hat{x}(\cdot))) \nabla c(\hat{x}(\cdot)).$$

By substituting  $\nabla_{\lambda x}^2 \mathcal{L}(\hat{x}(\cdot), \cdot)$  and  $\nabla_{xx}^2 \mathcal{L}(\hat{x}(\cdot), \cdot)$  into (9.6) we obtain the following formula for the Hessian of the dual function.

$$\nabla^2 d_k(\lambda) = \psi'(kc(\hat{x}(\cdot))) \nabla c(\hat{x}(\cdot)) (\nabla_{xx}^2 \mathcal{L}(\hat{x}(\cdot), \cdot))^{-1} \nabla c(\hat{x}(\cdot))^T \psi'(kc(\hat{x}(\cdot))). \quad (9.7)$$

Note that

$$\nabla^2 d_k(\lambda^*) = \psi'(kc(x^*)) \nabla c(x^*) (\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1} \nabla c(x^*)^T \psi'(kc(x^*)). \quad (9.8)$$

We proved the following lemma.

**Lemma 9.1.** If  $f(x)$  and all  $c_i(x) \in C^2$  then

- (1) if (P) is a convex programming problem and (A) is true then the dual function  $d_k(\lambda) \in C^2$  for any  $\lambda \in \mathfrak{R}_{++}^m$  and any  $k > 0$ ;
- (2) if (P) is a convex programming problem and  $f(x)$  is strongly convex then the dual function  $d_k(\lambda) \in C^2$  for any  $\lambda \in \mathfrak{R}_+^m$  and any  $k > 0$ ;
- (3) if the standard second order optimality conditions (2.4), (2.5) are satisfied then  $d_k(\lambda) \in C^2$  for any pair  $\lambda \in D(\cdot)$  and  $k \geq k_0$  whether the problem (P) is convex or not. The gradient  $\nabla d_k(\lambda)$  is given by (9.3) and the Hessian  $\nabla^2 d_k(\lambda)$  is given by (9.7).

**Theorem 9.1** (Duality).

- (1) If Slater condition (B) holds then the existence of the primal solution implies the existence of the dual solution and

$$f(x^*) = d_k(\lambda^*) \quad (9.9)$$

for any  $k > 0$ .

- (2) If  $f(x)$  is strictly convex,  $f(x)$  and all  $c_i(x)$  are smooth and the dual solution exists, then the primal exists and (9.9) holds for any  $k > 0$ .
- (3) If  $f(x)$  and all  $c_i(x) \in C^2$  and (2.4), (2.5) are satisfied then the second order optimality conditions hold true for the dual problem for any  $k \geq k_0$  if  $k_0$  is large enough.

*Proof.* (1) The primal solution  $x^*$  is at the same time a solution for the equivalent problem (4.1). Therefore keeping in mind the Slater condition (B) we obtain such  $\lambda^* \in \mathfrak{R}_+^m$  such that

$$\lambda_i^* c_i(x^*) = 0, \quad i = 1, \dots, m, \quad \mathcal{L}(x^*, \lambda^*, k) \leq \mathcal{L}(x, \lambda^*, k), \quad \forall x \in \mathfrak{R}^n, \quad k > 0.$$

Therefore

$$\begin{aligned} d_k(\lambda^*) &= \min_{x \in \mathfrak{R}^n} \mathcal{L}(x, \lambda^*, k) = \mathcal{L}(x^*, \lambda^*, k) = f(x^*) \\ &\geq \mathcal{L}(x^*, \lambda, k) \geq \min_{x \in \mathfrak{R}^n} \mathcal{L}(x, \lambda, k) = d_k(\lambda), \quad \forall \lambda \in \mathfrak{R}_+^m. \end{aligned}$$

Hence  $\lambda^* \in \mathfrak{R}_+^m$  is the solution for the dual problem and (9.9) holds.

(2) Let us assume that  $\bar{\lambda} \in \mathfrak{R}_+^m$  is the solution for the dual problem. If  $f(x)$  is strictly convex then the function  $\mathcal{L}(x, \lambda, k)$  is strictly convex too in  $x \in \mathfrak{R}^n$ . Therefore the gradient  $\nabla d_k(\lambda)$  exists. Consider the optimality condition for the dual problem

$$\bar{\lambda}_i = 0 \Rightarrow \nabla_{\lambda_i} d_k(\bar{\lambda}) \leq 0, \quad (9.10)$$

$$\bar{\lambda}_i > 0 \Rightarrow \nabla_{\lambda_i} d_k(\bar{\lambda}) = 0. \quad (9.11)$$

Let  $\bar{x} = \arg \min \{\mathcal{L}(x, \bar{\lambda}, k) \mid x \in \mathfrak{R}^n\}$ , then

$$\nabla_{\lambda_i} d_k(\bar{\lambda}) = 2k^{-1} \ln 0.5(1 + e^{-kc_i(\bar{x})}).$$

From (9.10) we obtain

$$\begin{aligned} \bar{\lambda}_i = 0 &\Rightarrow 2k^{-1} \ln 0.5(1 + e^{-kc_i(\bar{x})}) \leq 0 \Rightarrow 0.5(1 + e^{-kc_i(\bar{x})}) \leq 1 \Rightarrow e^{-kc_i(\bar{x})} \leq 1 \\ &\Rightarrow c_i(\bar{x}) \geq 0. \end{aligned}$$

From (9.11) we have

$$\bar{\lambda}_i > 0 \Rightarrow \ln 0.5(1 + e^{-kc_i(\bar{x})}) = 0 \Rightarrow e^{-kc_i(\bar{x})} = 1 \Rightarrow c_i(\bar{x}) = 0.$$

Therefore  $(\bar{x}, \bar{\lambda})$  is the primal–dual feasible pair for which the complementary conditions hold, i.e.,  $\bar{x} = x^*$ ,  $\bar{\lambda} = \lambda^*$ .

(3) To prove that for the dual problem the standard second order optimality conditions hold true we consider the Lagrangian for the dual problem

$$\lambda^* = \arg \max \{d_k(\lambda) \mid \lambda_i \geq 0, i = 1, \dots, m\}.$$

We have

$$L(\lambda, v, k) = d_k(\lambda) + \sum_{i=1}^m v_i \lambda_i.$$

Then

$$\nabla_{\lambda\lambda}^2 L(\lambda, v, k) = \nabla_{\lambda\lambda}^2 d_k(\lambda). \quad (9.12)$$

The active dual constraints are  $\lambda_i = 0$ ,  $i = r + 1, \dots, m$ , and the vectors  $e_i = (0, \dots, 0, 0, \dots, 1, \dots, 0)$ ,  $i = r + 1, \dots, m$ , are the gradients of the active dual constraints.

Therefore the tangent subspace to the dual active set at the point  $\lambda^*$  is

$$\begin{aligned} Y &= \{y \in \mathfrak{R}^m: (y, e_i) = 0, i = r + 1, \dots, m\} \\ &= \{y \in \mathfrak{R}^m: y = (y_1, \dots, y_r, 0, \dots, 0)\}. \end{aligned}$$

It is clear that the gradients  $e_i$ ,  $i = r + 1, \dots, m$ , of the dual active constraints are linear independent. So, to prove that for the dual problem the second order optimality conditions hold true we have to show

$$(\nabla_{\lambda\lambda}^2 L(\lambda^*, v^*, k)y, y) \leq -\mu \|y\|_2^2, \quad \mu > 0, \quad \forall y \in Y.$$

Using (9.8) and (9.12) we obtain

$$\begin{aligned} (\nabla_{\lambda\lambda}^2 L(\lambda^*, v^*, k)y, y) &= (\nabla_{\lambda\lambda} d_k(\lambda^*)y, y) \\ &= (\psi'(kc(x^*))\nabla c(x^*)(\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1}\nabla c(x^*)^T \psi'(kc(x^*))y, y) \\ &= -(\nabla c(x^*)(\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1}\nabla c(x^*)^T \bar{y}, \bar{y}), \end{aligned}$$

where

$$\begin{aligned} \bar{y} &= \psi'(kc(x^*))y = (\psi'(kc_1(x^*))y_1, \dots, \psi'(kc_r(x^*))y_r, 0, \dots, 0) \\ &= (y_1, \dots, y_r, 0, \dots, 0) = (y_{(r)}, 0) = y. \end{aligned}$$

In other words

$$(\nabla_{\lambda\lambda}^2 L(\lambda^*, v^*, k)y, y) = -(\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1}\nabla c(x^*)^T y, \nabla c(x^*)^T y).$$

Using (4.4) we obtain

$$(M_0 k)^{-1}(y, y) \leq ((\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1}y, y) \leq \mu_0^{-1}(y, y)$$

i.e.,

$$-(M_0 k)^{-1}(y, y) \geq -(\nabla_{xx}^2 \mathcal{L}(x^*, \lambda^*, k))^{-1}y, y \geq -\mu_0^{-1}(y, y).$$

Hence

$$\begin{aligned} (\nabla_{\lambda\lambda}^2 L(\lambda^*, v^*, k)y, y) &\leq -(M_0 k)^{-1}(\nabla c_{(r)}^T(x^*)y_{(r)}\nabla c_{(r)}^T(x^*)y_{(r)}) \\ &= -(M_0 k)^{-1}(\nabla c_{(r)}(x^*)\nabla c_{(r)}^T(x^*)y_{(r)}, y_{(r)}). \end{aligned} \quad (9.13)$$

Due to (2.4) the Gram matrix  $\nabla c_{(r)}\nabla c_{(r)}^T$  is positive definite, therefore there is  $\bar{\mu} > 0$  that  $(\nabla c_{(r)}\nabla c_{(r)}^T y_{(r)}, y_{(r)}) \geq \bar{\mu}\|y_{(r)}\|_2^2$ .

Therefore in view of (9.13) we obtain

$$(\nabla_{\lambda\lambda}^2 L(\lambda^*, v^*, k)y, y) \leq -\mu\|y\|_2^2, \quad \forall y \in Y, \quad (9.14)$$

where  $\mu = (M_0 k)^{-1}\bar{\mu}$ .

So the standard second order optimality condition holds true for the dual problem.  $\square$

**Corollary 9.1.** If (2.4), (2.5) hold and  $k_0 > 0$  is large enough then for any  $k \geq k_0$  the restriction  $\bar{d}_k(\lambda_{(r)}) = d_k(\lambda)|_{\lambda_{r+1}=0, \dots, \lambda_m=0}$  of the dual function to the manifold of the dual active constraints is strongly concave.

The properties of  $\bar{d}_k(\lambda_{(r)})$  allow to use smooth unconstrained optimization technique, in particular Newton method for solving the dual problem. It leads to the second order LS multipliers method.

*Remark 9.1.* The part (3) of theorem 9.1 holds true even for nonconvex optimization problems. It is not true if instead of  $\mathcal{L}(x, \lambda, k)$  one uses the classical Lagrangian  $L(x, \lambda)$ .

## 10. Numerical results

The primal–dual LS method we described in section 8 generally speaking does not converge globally. However locally it converges very fast. Therefore in the first stage of the computation we used the path following type approach with LS penalty function (6.1), i.e., we find an approximation for  $x(k)$  and increase  $k > 0$  from step to step. For the unconstrained minimization  $P(x, k)$  in  $x$  we used Newton method with step-length. When the duality gap becomes reasonably small  $10^{-1} \div 10^{-2}$  we use the primal approximation  $\bar{x}$  for  $x(k)$  to compute approximation  $\bar{\lambda}$  for the Lagrange multipliers  $\lambda(k)$  and then the primal–dual vector  $(\bar{x}, \bar{\lambda})$  is used as a starting point in the primal–dual method (8.11), (8.12).

The first stage consumes the most of the computational time, while the primal–dual method (8.11), (8.12) requires only a few steps to reduce the duality gap and the infeasibility from  $10^{-1} \div 10^{-2}$  to  $10^{-8} \div 10^{-10}$ . For all problems which have been solved we observed the “hot” start phenomenon (see [10,11]), when few and from some point on only one Newton step is required for finding an approximation with high accuracy for the solution of the primal–dual system (8.1), (8.2).

In the following tables we show numerical results obtained by using the NR multipliers method for several problems, which we downloaded from Dr. R. Vanderbei web-page.

Name: **catenary** n = 198, m = 298, p = 100;

```
# Objective: linear
# Constraints: convex quadratic
# Feasible set: convex

# This model finds the shape of a hanging chain
# The solution is known to be  $y = \cosh(a \cdot x) + b$ 
# for appropriate a and b.
```

Path following...

| it | f          | g          | gap          | inf          | step |
|----|------------|------------|--------------|--------------|------|
| 1  | -3.642e+03 | 2.0199e-01 | 1.823027e+03 | 3.296342e+00 | 3    |
| 2  | -1.815e+03 | 1.4758e-01 | 9.849452e+02 | 7.518446e-01 | 3    |
| 3  | -8.597e+02 | 7.6270e-02 | 4.644297e+02 | 1.631766e-01 | 3    |
| 4  | -4.032e+02 | 1.5686e-02 | 2.137841e+02 | 3.509197e-02 | 3    |
| 5  | -1.905e+02 | 4.1326e-03 | 9.598980e+01 | 7.457490e-03 | 3    |
| 6  | -9.466e+01 | 3.0647e-03 | 3.848685e+01 | 1.509609e-03 | 3    |
| 7  | -5.659e+01 | 1.6690e-03 | 1.021444e+01 | 2.538015e-04 | 3    |
| 8  | -4.699e+01 | 4.4405e-05 | 1.415691e+00 | 3.058176e-05 | 3    |

Primal–Dual algorithm...

|   |            |            |              |              |   |
|---|------------|------------|--------------|--------------|---|
| 0 | 6.009e+02  | 1.3909e+04 | 1.415691e+00 | 3.058176e-05 | 0 |
| 1 | -4.556e+01 | 1.9293e-05 | 2.154078e-04 | 3.814845e-09 | 6 |

|   |            |            |              |              |   |
|---|------------|------------|--------------|--------------|---|
| 2 | -4.556e+01 | 2.5129e-07 | 2.457058e-08 | 1.175767e-12 | 2 |
| 3 | -4.556e+01 | 2.0630e-10 | 7.838452e-13 | 6.960578e-17 | 1 |

Name: **esfl\_socp** n = 1002, m = 2002, p = 1000;

Pathfollowing...

| it | f          | g          | gap          | inf          | step |
|----|------------|------------|--------------|--------------|------|
| 1  | -6.807e+02 | 4.8419e-03 | 4.378594e+02 | 3.469734e-03 | 3    |
| 2  | -4.902e+02 | 6.0727e-03 | 1.305438e+02 | 2.517168e-03 | 3    |
| 3  | -7.101e+01 | 8.5318e-03 | 2.892838e+01 | 1.865576e-03 | 3    |
| 4  | 1.288e+02  | 8.8694e-03 | 5.795068e+00 | 1.459469e-03 | 3    |
| 5  | 1.738e+02  | 9.9911e-03 | 1.183240e+00 | 1.172302e-03 | 3    |
| 6  | 1.887e+02  | 2.2467e-02 | 2.208563e-01 | 8.762780e-04 | 3    |

Primal-Dual algorithm...

|   |           |            |              |              |   |
|---|-----------|------------|--------------|--------------|---|
| 0 | 1.964e+02 | 6.7414e+03 | 2.208563e-01 | 8.762780e-04 | 0 |
| 1 | 1.939e+02 | 9.5765e-03 | 2.415700e-03 | 5.177679e-06 | 2 |
| 2 | 1.939e+02 | 1.3120e-03 | 2.783116e-04 | 2.500009e-06 | 2 |
| 3 | 1.939e+02 | 1.6533e-04 | 1.514309e-05 | 6.249994e-07 | 1 |
| 4 | 1.939e+02 | 9.5390e-06 | 2.054630e-08 | 3.906232e-08 | 1 |
| 5 | 1.939e+02 | 3.7161e-08 | 2.221649e-11 | 1.525858e-10 | 1 |

Name: **fekete** n = 150, m = 200, p = 50

# Objective: nonconvex nonlinear

# Constraints: convex quadratic

Pathfollowing...

| it | f          | g          | gap          | inf          | step |
|----|------------|------------|--------------|--------------|------|
| 1  | -1.485e+02 | 3.7060e+00 | 1.667138e+02 | 5.547542e+00 | 13   |
| 2  | -4.631e+02 | 5.4162e+00 | 6.017736e+02 | 4.650966e+00 | 13   |
| 3  | -9.168e+02 | 5.0260e+00 | 3.302206e+01 | 2.641380e-01 | 13   |
| 4  | -1.064e+03 | 6.9358e+00 | 9.589543e+01 | 2.098166e-01 | 13   |
| 5  | -1.279e+03 | 6.5917e+00 | 5.954351e+00 | 3.569369e-02 | 13   |
| 6  | -1.427e+03 | 8.6109e+00 | 1.037654e+01 | 2.883360e-02 | 13   |
| 7  | -1.464e+03 | 7.7080e+00 | 3.049917e-03 | 3.698170e-05 | 13   |

Primal-Dual algorithm...

|   |            |            |              |              |   |
|---|------------|------------|--------------|--------------|---|
| 0 | -1.440e+03 | 1.1562e+03 | 3.049917e-03 | 3.698170e-05 | 0 |
| 1 | -1.442e+03 | 7.6695e+00 | 2.296847e-03 | 0.000000e+00 | 3 |
| 2 | -1.442e+03 | 7.6694e+00 | 2.322743e-05 | 0.000000e+00 | 2 |
| 3 | -1.442e+03 | 7.6694e+00 | 2.353873e-07 | 0.000000e+00 | 2 |
| 4 | -1.442e+03 | 7.6694e+00 | 2.386449e-09 | 0.000000e+00 | 2 |

Name: **fir\_socp** n = 12, m = 319, p = 307;

Pathfollowing...

| it | f          | g          | gap          | inf          | step |
|----|------------|------------|--------------|--------------|------|
| 1  | -1.720e+03 | 1.3242e+00 | 1.650751e+02 | 1.234151e-01 | 3    |
| 2  | -3.244e+02 | 5.1078e-03 | 3.754727e+01 | 3.207064e-02 | 3    |
| 3  | -1.380e+01 | 2.3305e-01 | 8.242217e+00 | 1.987763e-02 | 3    |

Primal-Dual algorithm...

|   |           |            |              |              |    |
|---|-----------|------------|--------------|--------------|----|
| 0 | 1.672e+01 | 5.9047e+02 | 8.242217e+00 | 1.987763e-02 | 0  |
| 1 | 9.707e-01 | 1.9810e-03 | 7.764040e-02 | 1.591197e-02 | 12 |
| 2 | 1.044e+00 | 2.6651e-03 | 1.467758e-03 | 4.177845e-04 | 5  |
| 3 | 1.046e+00 | 4.7845e-05 | 9.810459e-06 | 2.525622e-05 | 5  |
| 4 | 1.046e+00 | 6.0862e-04 | 1.095905e-07 | 2.867097e-06 | 4  |
| 5 | 1.046e+00 | 2.3531e-06 | 2.021246e-09 | 1.473551e-07 | 5  |
| 6 | 1.046e+00 | 1.3218e-05 | 1.046812e-08 | 4.219660e-08 | 3  |
| 7 | 1.046e+00 | 2.3915e-05 | 3.806949e-09 | 1.210397e-08 | 4  |
| 8 | 1.046e+00 | 1.3860e-05 | 2.821480e-10 | 8.391990e-10 | 4  |

Name: **hydrothermal** n = 46, m = 55, p = 9;

# Objective: nonconvex nonlinear

# Constraints: nonconvex nonlinear

Pathfollowing...

| it | f          | g          | gap          | inf          | step |
|----|------------|------------|--------------|--------------|------|
| 1  | -1.732e+04 | 1.4540e-01 | 4.277996e+04 | 9.834250e+01 | 13   |
| 2  | 1.788e+04  | 6.7404e-01 | 6.928106e+04 | 5.649293e+01 | 2    |
| 3  | 7.654e+04  | 5.5557e-02 | 2.559329e+04 | 1.527546e+01 | 1    |
| 4  | 9.689e+04  | 2.1795e-02 | 8.946096e+03 | 3.294818e+00 | 1    |
| 5  | 1.015e+05  | 3.6403e-03 | 1.691519e+04 | 2.456099e+00 | 1    |
| 6  | 1.132e+05  | 9.0428e-05 | 5.242193e+04 | 2.019004e+00 | 4    |
| 7  | 1.626e+05  | 2.6742e+01 | 2.548234e+04 | 6.156682e-01 | 25   |
| 8  | 1.818e+05  | 3.8606e+01 | 5.081058e+03 | 1.229012e-01 | 25   |
| 9  | 1.859e+05  | 2.5459e-02 | 1.034586e+03 | 2.480066e-02 | 21   |
| 10 | 1.867e+05  | 3.7396e-01 | 2.076472e+02 | 4.968390e-03 | 1    |

Primal-Dual algorithm...

|   |           |            |              |              |   |
|---|-----------|------------|--------------|--------------|---|
| 0 | 1.910e+05 | 9.7227e+02 | 2.076472e+02 | 4.968390e-03 | 0 |
| 1 | 1.909e+05 | 6.1491e-04 | 9.300815e-02 | 2.583264e-06 | 3 |
| 2 | 1.909e+05 | 2.5799e-03 | 4.112739e-05 | 3.269179e-09 | 1 |
| 3 | 1.909e+05 | 1.5829e-04 | 9.339780e-06 | 1.001445e-08 | 2 |
| 4 | 1.909e+05 | 7.6930e-04 | 1.196375e-08 | 4.407354e-10 | 1 |

|   |           |            |              |              |   |
|---|-----------|------------|--------------|--------------|---|
| 5 | 1.909e+05 | 4.0202e-05 | 5.318109e-08 | 2.343114e-09 | 2 |
| 6 | 1.909e+05 | 2.1153e-04 | 1.352623e-09 | 5.428547e-11 | 1 |

## 11. Concluding remarks

This paper presents a part of our work the purpose of which was to develop an alternative to the smoothing technique approach for constrained optimization.

Different strategies for the scaling parameter update remain an important issue as well as the global convergence of the LS multipliers method. In this respect the equivalence of the LS multipliers method to the interior prox method with  $\varphi$ -divergence Fermi–Dirac distance will play an important role.

The convergence analysis of the primal–dual LS method is another important issue.

One of the most important quality of the LS multipliers method is the opportunity to use the Newton method for primal optimization in the entire primal space without using the extrapolation technique, see [2,3,5,13]. The number of Newton steps required per the Lagrange multipliers update decreases drastically after very few updates. On the other hand in most cases we need only few Lagrange multipliers update to guarantee up to ten digits of accuracy, see [12].

Finding the upper bound for the number of Newton steps required to obtain the primal–dual approximation with a given accuracy remains an important issue.

All these issues as well as a number of questions related to the numerical realization of the LS multipliers method are left for future research.

## Acknowledgements

I am grateful to Professor Marc Teboulle for fruitful discussions at the early stage of this work. I would like to thank my Ph.D. student Igor Griva for his hard work which made possible to conduct extensive numerical experiment and helped us to understand better different aspects of the LS multipliers method.

I also would like to thank the anonymous referees for their comments, which contributed to the improvement of the paper.

## References

- [1] A. Auslender, R. Cominetti and M. Haddou, Asymptotic analysis for penalty and barrier methods in convex and linear programming, *Mathematics of Operations Research* 22 (1) (1997) 43–62.
- [2] A. Ben-Tal, I. Uzefovich and M. Zibulevsky, Penalty/barrier multiplier methods for minmax and constrained smooth convex programs, *Research Report, Optimization Laboratory, Technion, Israel* (1992) pp. 1–16.
- [3] A. Ben-Tal and M. Zibulevsky, Penalty–barrier methods for convex programming problems, *SIAM J. Optimization* 7 (1997) 347–366.
- [4] D. Bertsekas, *Constrained Optimization and Lagrange Multiplier Methods* (Academic Press, New York, 1982).

- [5] M. Breitfeld and D. Shanno, Computational experience with modified log-barrier functions for non-linear programming, *Annals of Operations Research* 62 (1996) 439–464.
- [6] C. Chen and O.L. Mangasarian, Smoothing methods for convex inequalities and linear complementarity problems, *Mathematical Programming* 71 (1995) 51–69.
- [7] A.V. Fiacco and G.P. McCormick, *Nonlinear Programming: Sequential Unconstrained Minimization Techniques*, SIAM Classics in Applied Mathematics (SIAM, Philadelphia, PA, 1990).
- [8] A.N. Iusem, B. Svaiter and M. Teboulle, Entropy-like proximal methods in convex programming, *Mathematics of Operations Research* 19 (1994) 790–814.
- [9] B.W. Kort and D.P. Bertsekas, Multiplier methods for convex programming, in: *Proceedings 1973, IEEE Conference on Decision and Control*, San Diego, CA, pp. 428–432.
- [10] A. Melman and R. Polyak, The Newton modified barrier method for quadratic programming problems, *Annals of Operations Research* 62 (1996) 465–519.
- [11] R. Polyak, Modified barrier functions (theory and methods), *Mathematical Programming* 54 (1992) 177–222.
- [12] R. Polyak, I. Griva and J. Sobieski, The Newton log-sigmoid method in constrained optimization, in: *A Collection of Technical Papers, 7th AIAA/USAF/NASA/ISSMO Symposium on Multidisciplinary Analysis and Optimization* 3 (1998) pp. 2193–2201.
- [13] R. Polyak and M. Teboulle, Nonlinear rescaling and proximal-like methods in convex optimization, *Mathematical Programming* 76 (1997) 965–984.
- [14] R.T. Rockafellar, *Convex Analysis* (Princeton University Press, Princeton, NJ, 1970).
- [15] M. Teboulle, On  $\psi$ -divergence and its applications, in: *Systems and Management Science by Extremal Methods*, eds. F.Y. Philips and J.J. Rousseau (Kluwer Academic, 1992) pp. 255–289.
- [16] M. Teboulle, Entropic proximal mappings with application to nonlinear programming, *Mathematics of Operations Research* 17 (1992) 670–690.
- [17] P. Tseng and O. Bertsekas, On the convergence of the exponential multiplier method for convex programming, *Mathematical Programming* 60 (1993) 1–19.