

An ADMM-LAP method for total variation myopic deconvolution of adaptive optics retinal images

Xiaotong Chen¹, James L Herring², James G Nagy^{3,*} ,
Yuanzhe Xi³  and Bo Yu¹

¹ School of Mathematical Sciences, Dalian University of Technology, Dalian, Liaoning 116024, People's Republic of China

² Department of Mathematics, University of Houston, Houston, TX 77204, United States of America

³ Department of Mathematics, Emory University, Atlanta, GA 30322, United States of America

E-mail: chenxiaotong@mail.dlut.edu.cn, jnagy@emory.edu, herrinj@gmail.com, yxi26@emory.edu and yubo@dlut.edu.cn

Received 23 December 2019, revised 30 August 2020

Accepted for publication 4 September 2020

Published 3 December 2020



CrossMark

Abstract

Adaptive optics corrected flood imaging of the retina is a popular technique for studying the retinal structure and function in the living eye. However, the raw retinal images are usually of poor contrast and the interpretation of such images requires image deconvolution. Different from standard deconvolution problems where the point spread function (PSF) is completely known, the PSF in these retinal imaging problems is only partially known which leads to the more complicated myopic (mildly blind) deconvolution problem. In this paper, we propose an efficient numerical scheme for solving this myopic deconvolution problem with total variational (TV) regularization. First, we apply the alternating direction method of multipliers (ADMM) to tackle the TV regularizer. Specifically, we reformulate the TV problem as an equivalent equality constrained problem where the objective function is separable, and then minimize the augmented Lagrangian function by alternating between two (separated) blocks of unknowns to obtain the solution. Due to the structure of the retinal images, the subproblems with respect to the fidelity term appearing within each ADMM iteration are tightly coupled and a variation of the linearize and project method is designed to solve these subproblems efficiently. The proposed method is called the ADMM-LAP method. Theoretically, we establish the subsequence convergence of the ADMM-LAP method to a stationary point.

*Author to whom any correspondence should be addressed.

Both the theoretical complexity analysis and numerical results are provided to demonstrate the efficiency of the ADMM-LAP method.

Keywords: total variation, retinal imaging, image restoration, myopic deconvolution

(Some figures may appear in colour only in the online journal)

1. Introduction

Image restoration is an important topic in image processing which is widely used in many areas, such as astronomical imaging, medical imaging and restoring aging and deteriorated films. The goal of image restoration is to reconstruct the best possible approximation of the clean, original image from an observed, blurred and noisy image. The basic image restoration problem can be described as a linear inverse problem

$$\mathbf{d} = \mathbf{A}\mathbf{x} + \mathbf{e}, \quad (1)$$

where $\mathbf{A} \in \mathbb{R}^{n^2 \times n^2}$ is an ill-conditioned blurring matrix defined by the point spread function (PSF) [9], $\mathbf{d} \in \mathbb{R}^{n^2}$ represents the observed, blurred and noisy image, $\mathbf{e} \in \mathbb{R}^{n^2}$ denotes the additive noise and $\mathbf{x} \in \mathbb{R}^{n^2}$ denotes the unknown true image to be restored.

The PSF is generally assumed to be perfectly known in standard image restoration techniques. However, this is not always the case. In many applications, the true PSF (and therefore the blurring matrix $\mathbf{A}$) is unknown or partially known. This results in blind deconvolution problems where image restoration also requires recovering or approximating the PSF. If the PSF can be parameterized by a small number of unknown parameters, the problem can be considered mildly/partially blind, or *myopic*. Adaptive optics (AO) corrected flood imaging of the retina is one such myopic deconvolution problem that has received much attention. In [2, 4], the authors present an imaging model that can transfer the 3D model to a 2D model with the global PSF being an unknown linear combination of a few PSFs. Thus, problem (1) requires estimating both the combination coefficients of $\mathbf{A}$ and the true image $\mathbf{x}$ in this model.

In order to guarantee the fidelity of the recovery, it is necessary to add a regularizer. There are two well-known types of regularizers for problem (1): one is Tikhonov and the other is total variation (TV). Tikhonov regularization was first proposed in [20] with a quadratic penalty added to the objective function. Due to its quadratic property, it is inexpensive to minimize the objective function. Thus, Tikhonov regularization is computationally efficient and has been widely used. However, one disadvantage of the Tikhonov approach is that it tends to over smooth the image and fails to preserve important image details such as sharp edges. The TV regularizer, first proposed by Rudin, Osher and Fatemi in [16], has also been widely adopted in image reconstruction problems [11, 17, 19, 21]. Since the TV approach uses the summation of the variation of the image $\mathbf{x}$ at all pixels to control the norm (or semi-norm) and the smoothness of the solution, it has been shown both experimentally and theoretically that the TV approach can effectively preserve sharp edges and keep the important features of the restored image.

Following the AO retinal imaging model in [2, 4], in this paper, we consider the myopic deconvolution model with TV regularization as follows

$$\begin{aligned} \min_{\mathbf{x} \in \mathcal{C}_x, \mathbf{w} \in \mathcal{C}_w} \quad & \Phi(\mathbf{x}, \mathbf{w}) = \frac{\mu}{2} \|\mathbf{A}(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + \text{TV}(\mathbf{x}) \\ \text{s.t.} \quad & \sum_{j=1}^p w_j = 1, \end{aligned} \quad (2)$$

where $\mathcal{C}_x = \{x | x_i \geq 0 \text{ for } i = 1, \dots, n^2\}$, $\mathcal{C}_w = \{w | w_j \geq 0 \text{ for } j = 1, \dots, p\}$, μ is a positive parameter that is used to balance the two terms in the objective function, TV denotes the TV regularization term, $x \in \mathbb{R}^{n^2}$ is a vectorized version of the unknown $n \times n$ image to be recovered, $w \in \mathbb{R}^p$ denotes unknown weights parameterizing the blurring matrix $A(w) \in \mathbb{R}^{n^2 \times n^2}$, and $d \in \mathbb{R}^{n^2}$ is the observed, blurred noisy image (data). For the imaging model we use, the blurring matrix $A(w)$ is a weighted sum of p known blurring matrices A_j with the form

$$A(w) = \sum_{j=1}^p w_j A_j = w_1 A_1 + w_2 A_2 + \dots + w_p A_p.$$

Note that the fidelity term in the objective function in (2) is nonconvex and the TV regularization term is nondifferentiable and nonlinear. This poses some computational challenges in the optimization problem. In particular, due to the existence of the parameters w and the nonconvexity of the objective function, approaches such as the fast iterative shrinkage-thresholding algorithm [1, 6] cannot be applied directly to solve the problem. However, alternative approaches have been proposed to solve the optimization problem with TV regularization in the literature. One particularly efficient scheme is the alternating direction method of multipliers (ADMM) [19, 21, 28]. ADMM was first developed to solve convex optimization problems by breaking them into subproblems, each of which are then easier to handle. Recent work has shown that ADMM can also perform well for a variety of applications involving nonconvex objective functions or nonconvex sets [5, 13, 22, 26, 27]. Inspired by the success of ADMM on nonconvex problems, we consider using it as the optimization method to tackle problem (2).

The main contribution of this paper can be summarized as follows:

- We implement an efficient algorithm called ADMM-LAP for myopic deconvolution problems with TV regularization arising from the AO retinal image restoration. Specifically, we first apply ADMM as an outer optimization method to tackle the TV regularizer and then apply the linearize and project (LAP) method as an inner optimization method to solve the tightly coupled subproblems arising within each ADMM iteration. We establish the subsequent convergence of ADMM-LAP to a stationary point and conduct a complexity analysis of ADMM-LAP to demonstrate its computational efficiency for myopic deconvolution problems with TV regularization.
- We present extensive numerical experiments to illustrate the effectiveness of the ADMM-LAP method. In addition, we compare the performance of ADMM-LAP with a benchmark method called ADMM-BCD where ADMM is applied to tackle the TV regularization while block coordinate descent (BCD) is applied to solve the coupled subproblems. Compared to ADMM-BCD, ADMM-LAP converges faster and can reach smaller relative errors for both the restored image and the obtained parameters.

The paper is organized as follows. In section 2, we introduce a general formulation of the AO retinal imaging problem. In section 3, we first briefly review the iteration format of ADMM and introduce a variation of LAP for our problem, then propose the ADMM-LAP method and present the convergence results. A benchmark method ADMM-BCD is also discussed. The computational complexity of both methods are analyzed in this section. Numerical results are given in section 4, and concluding remarks are drawn in section 5.

2. Retinal imaging problem

2.1. AO retinal imaging model

AO is a well-known optoelectronic technique that compensates for the time-varying aberrations of the eye [3]. However, AO flood imaging suffers from an intrinsic limitation that leads to a loss in resolution because the object is three-dimensional, the image contains information from both the in-focus plane and the out-of-focus planes of the object. Hence, interpretation of such images requires an appropriate post-processing, including image deconvolution. In this paper, we focus on the myopic image deconvolution problem that arises from the AO retinal image restoration and propose the efficient ADMM-LAP method to solve the problem.

First, we describe the structure of the AO retinal imaging model. In the continuous setting, retinal imaging is typically modeled as a three-dimensional (3D) convolution [2, 4]:

$$\mathbf{d}_{3D} = \mathbf{h}_{3D} *_{3D} \mathbf{x}_{3D} + \mathbf{e}, \quad (3)$$

where $\mathbf{d}_{3D}$ is the observed image, $\mathbf{h}_{3D}$ is the PSF, $*_{3D}$ is the three-dimensional convolution operator, $\mathbf{x}_{3D}$ is the true object, and $\mathbf{e}$ is the additive noise. If the true object is assumed to be shift invariant along the optical axis, i.e., the z -axis, then $\mathbf{x}_{3D}$ becomes separable with

$$\mathbf{x}_{3D}(x, y, z) = \mathbf{x}_{2D}(x, y) \mathbf{w}(z),$$

where $\mathbf{w}(z)$ is the normalized flux emitted by the plane at depth z such that $\int \mathbf{w}(z) dz = 1$. For reasonable optical setups, this shift invariance along the optical axis can be guaranteed to a sufficient degree to make this separability assumption meaningful [2].

In practice, retinal flood imaging systems typically image along a single plane of interest, a departure from the 3D model in (3). This results in a 2D data image taken at a single depth. For depth $z = 0$, this gives the observed image

$$\mathbf{d}_{2D}(x, y) = \mathbf{d}_{3D}(x, y, 0).$$

The shift invariance assumption for the true image $\mathbf{x}_{3D}$ then implies that the two-dimensional PSF for the observed image $\mathbf{d}_{2D}(x, y)$ can be expressed as

$$\mathbf{h}_{2D}(x, y) = \int \mathbf{w}(-z) \mathbf{h}_{3D}(x, y, z) dz.$$

Discretizing the above integral using a quadrature rule (a simple rectangle rule is often sufficient) we obtain the PSF for the observed retinal flood images as

$$\mathbf{h}_{2D}(x, y) \approx \sum_{j=1}^p w_j \mathbf{h}_j(x, y),$$

where $\mathbf{h}_j(x, y) = \mathbf{h}_{3D}(x, y, z_j)$ is the 2D PSF taken at depth z_j and $w_j = \mathbf{w}(z_j) \Delta z_j$ are weights with Δz_j the thickness of the j th slice of the 3D image in the quadrature sum. Thus, in our myopic deconvolution model (2), the problem becomes determining the unknown weights w_j that parameterize the PSF with the constraints that $\sum_{j=1}^p w_j = 1$ and $w_j \geq 0$ for all j . We point out that it can be seen from numerical experiments that the global PSF in (2) is normalized (i.e., all entries sum to approximately 1). Therefore, the myopic deconvolution model satisfies the fact that the PSF is energy preserving.

2.2. Myopic deconvolution model with TV regularization

For TV regularization in (2), we follow the notations used in [19, 21]: the discrete form of TV for a grayscale image $\mathbf{x} \in \mathbb{R}^{n^2}$ is defined as

$$\text{TV}(\mathbf{x}) = \sum_{i=1}^{n^2} \|D_i \mathbf{x}\|_2, \quad (4)$$

where for each i , $D_i \mathbf{x} \in \mathbb{R}^2$ represents the first-order finite difference of $\mathbf{x}$ at pixel i in both horizontal and vertical directions. We note that the 2-norm in (4) can be replaced by the 1-norm. If the 2-norm is used, then we obtain the isotropic version of TV, and if the 1-norm is used, we obtain the anisotropic version. In this paper, we will only treat the isotropic case for simplicity, but the treatment for the anisotropic case is completely analogous.

Notations: the two first-order global finite difference operators in horizontal and vertical directions are, respectively, denoted by $D^{(1)}, D^{(2)} \in \mathbb{R}^{n^2 \times n^2}$. $D_i \in \mathbb{R}^{2 \times n^2}$ is a two-row matrix formed by stacking the i th row of $D^{(1)}$ on top of the i th row of $D^{(2)}$ and $D := (D^{(1)}, D^{(2)}) \in \mathbb{R}^{2n^2 \times n^2}$ is the global first-order finite difference operator.

In order to deal with the restriction on the sum of the weights, we introduce an appropriate regularizer on $\mathbf{w}$:

$$\min_{\mathbf{x} \in \mathcal{C}_x, \mathbf{w} \in \mathcal{C}_w} \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + \sum_{i=1}^{n^2} \|D_i \mathbf{x}\|_2 + S(\mathbf{w}) \quad (5)$$

where $\mathcal{C}_x = \{\mathbf{x} | x_i \geq 0 \text{ for } i = 1, \dots, n^2\}$, $\mathcal{C}_w = \{\mathbf{w} | w_j \geq 0 \text{ for } j = 1, \dots, p\}$ and $A(\mathbf{w}) = \sum_{j=1}^p w_j A_j$. The regularizer $S(\mathbf{w})$ is defined as

$$S(\mathbf{w}) = \frac{\xi}{2} (\mathbf{e}^\top \mathbf{w} - 1)^2,$$

where $\mathbf{e} \in \mathbb{R}^p$ is the vector all ones and $\xi > 0$ is a weighting parameter. This regularizer penalizes solutions where the weights w_k do not sum to 1, so for appropriately large ξ , it can effectively enforce the summation constraint. We choose this option because it fits conveniently within the LAP framework in section 3.2. Another possible alternative is a Lagrangian multiplier approach, which would strictly enforce the summation [15].

3. Optimization schemes

In this section, we first briefly review the ADMM, discuss how to adapt ADMM to solve (5) and discuss how to adapt a variation of the LAP method to solve the related coupled subproblems. We then propose the ADMM-LAP method for solving (5) and give the convergence analysis. A benchmark method called ADMM-BCD is also introduced in this section. We compare their computational complexity and show that ADMM-LAP is more efficient when solving (5).

3.1. ADMM splitting

The classical ADMM is designed to solve the following two-block optimization problem with linear constraints

$$\begin{aligned} \min_{\mathbf{u}, \mathbf{v} \in \mathbb{R}^{n^2}} & f(\mathbf{u}) + g(\mathbf{v}) \\ \text{s.t.} & \mathbf{B}\mathbf{u} + \mathbf{C}\mathbf{v} = \mathbf{z}, \end{aligned} \quad (6)$$

where $f(\mathbf{u}) : \mathbb{R}^{n^2} \rightarrow (-\infty, +\infty]$, $g(\mathbf{v}) : \mathbb{R}^{n^2} \rightarrow (-\infty, +\infty]$ are convex functions, and $B, C \in \mathbb{R}^{n^2 \times n^2}$ and $\mathbf{z} \in \mathbb{R}^{n^2}$ are given.

The augmented Lagrangian function for (6) is given by

$$\mathcal{L}_\beta(\mathbf{u}, \mathbf{v}, \boldsymbol{\lambda}) = f(\mathbf{u}) + g(\mathbf{v}) + (\boldsymbol{\lambda}, B\mathbf{u} + C\mathbf{v} - \mathbf{z}) + \frac{\beta}{2} \|B\mathbf{u} + C\mathbf{v} - \mathbf{z}\|_2^2,$$

where $\boldsymbol{\lambda} \in \mathbb{R}^{n^2}$ denotes the Lagrange multiplier, $\beta > 0$ is the penalty parameter.

Given $(\mathbf{u}^0, \mathbf{v}^0, \boldsymbol{\lambda}^0) \in \mathbb{R}^{n^2} \times \mathbb{R}^{n^2} \times \mathbb{R}^{n^2}$, the penalty parameter $\beta > 0$, ADMM iterates as follows:

$$\begin{cases} \mathbf{u}^{k+1} = \underset{\mathbf{u}}{\operatorname{argmin}} \mathcal{L}_\beta(\mathbf{u}, \mathbf{v}^k, \boldsymbol{\lambda}^k), \\ \mathbf{v}^{k+1} = \underset{\mathbf{v}}{\operatorname{argmin}} \mathcal{L}_\beta(\mathbf{u}^{k+1}, \mathbf{v}, \boldsymbol{\lambda}^k), \\ \boldsymbol{\lambda}^{k+1} = \boldsymbol{\lambda}^k - \beta(B\mathbf{u}^{k+1} + C\mathbf{v}^{k+1} - \mathbf{z}). \end{cases}$$

In order to apply ADMM to solve (5), we introduce artificial vectors $\mathbf{y}_i \in \mathbb{R}^2, i = 1, \dots, n^2$, then we can rewrite (5) in an equivalent form:

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{w}, \mathbf{y}} \quad & \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}) + \sum_{i=1}^{n^2} \|\mathbf{y}_i\|_2 + \delta_{\mathcal{C}_x}(\mathbf{x}) + \delta_{\mathcal{C}_w}(\mathbf{w}), \\ \text{s.t.} \quad & \mathbf{y}_i = D_i \mathbf{x}, \quad i = 1, \dots, n^2, \end{aligned} \quad (7)$$

where $\mathcal{C}_x = \{\mathbf{x} | x_i \geq 0 \text{ for } i = 1, \dots, n^2\}$, $\mathcal{C}_w = \{\mathbf{w} | w_j \geq 0 \text{ for } j = 1, \dots, p\}$, $A(\mathbf{w}) = \sum_{j=1}^p w_j A_j$ and $\delta_{\mathcal{C}}(\cdot)$ denotes the indicator function of $\mathcal{C}$, i.e.,

$$\delta_{\mathcal{C}}(s) = \begin{cases} 0, & s \in \mathcal{C}, \\ \infty, & s \notin \mathcal{C}. \end{cases}$$

For convenience, let⁴ $\mathbf{y} = [\mathbf{y}_1; \mathbf{y}_2] \in \mathbb{R}^{2n^2}$, where $\mathbf{y}_1, \mathbf{y}_2$ are vectors of length n^2 and $[(\mathbf{y}_1)_i; (\mathbf{y}_2)_i] = \mathbf{y}_i \in \mathbb{R}^2$ for $i = 1, \dots, n^2$. The augmented Lagrangian function for (7) is then given by

$$\begin{aligned} \mathcal{L}_\beta(\mathbf{x}, \mathbf{w}, \mathbf{y}, \boldsymbol{\lambda}) := & \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}) + \delta_{\mathcal{C}_x}(\mathbf{x}) + \delta_{\mathcal{C}_w}(\mathbf{w}) \\ & + \sum_{i=1}^{n^2} \left(\|\mathbf{y}_i\|_2 - \boldsymbol{\lambda}_i^T (\mathbf{y}_i - D_i \mathbf{x}) + \frac{\beta}{2} \|\mathbf{y}_i - D_i \mathbf{x}\|_2^2 \right), \end{aligned}$$

where each $\boldsymbol{\lambda}_i \in \mathbb{R}^2$ and $\boldsymbol{\lambda} \in \mathbb{R}^{2n^2}$ is a reordering of $\boldsymbol{\lambda}_i$ similar to $\mathbf{y}$.

Consider $(\mathbf{x}, \mathbf{w})$ as one block of variables and $\mathbf{y}$ as the other. We can now apply ADMM as the outer optimization method to tackle the TV regularization term in (7). Given the initial point

⁴ Borrowing MATLAB notation, we use the semicolon in a vector to denote concatenation of terms.

$(\mathbf{y}^0, \mathbf{x}^0, \mathbf{w}^0, \boldsymbol{\lambda}^0)$, the ADMM algorithm iteratively solves the following three subproblems:

$$\begin{cases} \mathbf{y}^{k+1} = \underset{\mathbf{y}}{\operatorname{argmin}} \mathcal{L}_\beta(\mathbf{x}^k, \mathbf{w}^k, \mathbf{y}, \boldsymbol{\lambda}^k), \\ (\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) = \underset{\mathbf{x}, \mathbf{w}}{\operatorname{argmin}} \mathcal{L}_\beta(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k), \\ \boldsymbol{\lambda}^{k+1} = \boldsymbol{\lambda}^k - \beta(\mathbf{y}^{k+1} - D\mathbf{x}^{k+1}). \end{cases}$$

First, for the $\mathbf{y}$ -subproblem, notice that minimizing $\mathcal{L}_\beta(\mathbf{x}^k, \mathbf{w}^k, \mathbf{y}, \boldsymbol{\lambda}^k)$ with respect to $\mathbf{y}$ is equivalent to minimizing n^2 two-dimensional problems of the form

$$\min_{\mathbf{y}_i \in \mathbb{R}^2} \|\mathbf{y}_i\|_2 + \frac{\beta}{2} \|\mathbf{y}_i - \left(D_i \mathbf{x}^k + \frac{1}{\beta}(\boldsymbol{\lambda}^k)_i\right)\|_2^2, \quad i = 1, 2, \dots, n^2. \quad (8)$$

The solution to (8) is given explicitly by the two-dimensional shrinkage [19]

$$\mathbf{y}_i^{k+1} = \max \left\{ \left\| D_i \mathbf{x}^k + \frac{1}{\beta}(\boldsymbol{\lambda}^k)_i \right\|_2 - \frac{1}{\beta}, 0 \right\} \frac{D_i \mathbf{x}^k + \frac{1}{\beta}(\boldsymbol{\lambda}^k)_i}{\|D_i \mathbf{x}^k + \frac{1}{\beta}(\boldsymbol{\lambda}^k)_i\|_2}, \quad i = 1, 2, \dots, n^2. \quad (9)$$

Second, denote

$$R(\mathbf{x}, \mathbf{y}, \boldsymbol{\lambda}) := \sum_{i=1}^{n^2} \left(\|\mathbf{y}_i\|_2 - \boldsymbol{\lambda}_i^\top (\mathbf{y}_i - D_i \mathbf{x}) + \frac{\beta}{2} \|\mathbf{y}_i - D_i \mathbf{x}\|_2^2 \right),$$

and define

$$\hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}, \boldsymbol{\lambda}) := \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}) + R(\mathbf{x}, \mathbf{y}, \boldsymbol{\lambda}).$$

Then it can be shown that the $(\mathbf{x}, \mathbf{w})$ -subproblem is equivalent to the following problem:

$$(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) = \underset{\mathbf{x} \in \mathcal{C}_x, \mathbf{w} \in \mathcal{C}_w}{\operatorname{argmin}} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k). \quad (10)$$

Note that this is a tightly coupled optimization problem with element-wise bound constraints.

Finally, update the multiplier,

$$\boldsymbol{\lambda}^{k+1} = \boldsymbol{\lambda}^k - \beta(\mathbf{y}^{k+1} - D\mathbf{x}^{k+1}). \quad (11)$$

If the termination criterion is met, stop; else, set $k := k + 1$ and go to the $\mathbf{y}$ -subproblem.

3.2. LAP method

To solve the tightly coupled $(\mathbf{x}, \mathbf{w})$ -subproblem (10), we develop a variation of the LAP method proposed by Herring *et al* [10]. The LAP method is efficient for inverse problems with multiple, tightly coupled blocks of variables such as the problem under consideration. Its strengths include the option to impose element-wise bound constraints on all blocks of variables.

In this paper, we present the LAP method based on the normal equation approach instead of the least squares approach presented in the original paper [10]. First, we consider the unconstrained problem where $\mathcal{C}_x = \mathbb{R}^{n^2}$, $\mathcal{C}_w = \mathbb{R}^p$. Denote the residual as $\mathbf{r}(\mathbf{x}, \mathbf{w}) := A(\mathbf{w})\mathbf{x} - \mathbf{d}$. Then at the iterate $(\mathbf{x}, \mathbf{w})$, the Jacobian with respect to the $\mathbf{x}$ block of variables is

$$\mathbf{J}_x = \nabla_{\mathbf{x}} \mathbf{r}(\mathbf{x}, \mathbf{w})^\top = A(\mathbf{w})^\top, \quad (12)$$

and the Jacobian with respect to the w block of variables is

$$\mathbf{J}_w = \begin{bmatrix} (A_1 \mathbf{x})^\top \\ (A_2 \mathbf{x})^\top \\ \vdots \\ (A_p \mathbf{x})^\top \end{bmatrix}. \quad (13)$$

Computing the update step around the current iterate $(\mathbf{x}, w)$ requires the gradient and Hessian of the objective function $\hat{\Phi}(\mathbf{x}, w, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$. These are given by

$$\nabla_{\mathbf{x}, w} \hat{\Phi}(\mathbf{x}, w, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) = \mu \begin{bmatrix} \mathbf{J}_x^\top \mathbf{r} \\ \mathbf{J}_w^\top \mathbf{r} \end{bmatrix} + \begin{bmatrix} \nabla_x R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ \nabla_w S(w) \end{bmatrix}, \quad (14)$$

$$\nabla_{\mathbf{x}, w}^2 \hat{\Phi}(\mathbf{x}, w, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \approx \mu \begin{bmatrix} \mathbf{J}_x^\top \mathbf{J}_x & \mathbf{J}_x^\top \mathbf{J}_w \\ \mathbf{J}_w^\top \mathbf{J}_x & \mathbf{J}_w^\top \mathbf{J}_w \end{bmatrix} + \begin{bmatrix} \nabla_x^2 R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) & \mathbf{0} \\ \mathbf{0} & \nabla_w^2 S(w) \end{bmatrix}, \quad (15)$$

where $\mathbf{r} := \mathbf{r}(\mathbf{x}, w)$. Here, the Hessian of the regularizer $R(\mathbf{x})$ can be exact or a linearized approximation. Let $\delta \mathbf{x}$ and δw denote the update step for the image and the parameters, respectively. Then the update steps $\delta \mathbf{x}$ and δw are given by the solution of the following block linear system

$$\begin{bmatrix} \mathbf{J}_x^\top \mathbf{J}_x + \frac{1}{\mu} \nabla_x^2 R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) & \mathbf{J}_x^\top \mathbf{J}_w \\ \mathbf{J}_w^\top \mathbf{J}_x & \mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(w) \end{bmatrix} \begin{bmatrix} \delta \mathbf{x} \\ \delta w \end{bmatrix} = - \begin{bmatrix} \mathbf{J}_x^\top \mathbf{r} + \frac{1}{\mu} \nabla_x R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ \mathbf{J}_w^\top \mathbf{r} + \nabla_w S(w) \end{bmatrix}. \quad (16)$$

Note that omitting the regularizer terms, (16) is the normal equation corresponding to the least-squares problem

$$\frac{\mu}{2} \left\| \begin{bmatrix} \mathbf{J}_x & \mathbf{J}_w \end{bmatrix} \begin{bmatrix} \delta \mathbf{x} \\ \delta w \end{bmatrix} + \mathbf{r} \right\|_2^2$$

that can be obtained by the *Linearize* step of LAP in [10].

LAP solves (16) by projecting the original problem onto a reduced space. In this paper, we choose to project the problem onto the image space, i.e., we eliminate the block of variables corresponding to w . When projecting the problem onto the image space, δw can be computed by

$$\delta w = -(\mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(w))^{-1} (\mathbf{J}_w^\top \mathbf{J}_x \delta \mathbf{x} + \mathbf{J}_w^\top \mathbf{r} + \nabla_w S(w)). \quad (17)$$

Plug (17) into (16) and get

$$\begin{aligned} & \left(\mathbf{J}_x^\top \mathbf{J}_x + \frac{1}{\mu} \nabla_x^2 R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \right) \delta \mathbf{x} - \mathbf{J}_x^\top \mathbf{J}_w (\mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(w))^{-1} \\ & \times (\mathbf{J}_w^\top \mathbf{J}_x \delta \mathbf{x} + \mathbf{J}_w^\top \mathbf{r} + \nabla_w S(w)) = - \left(\mathbf{J}_x^\top \mathbf{r} + \frac{1}{\mu} \nabla_x R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \right), \end{aligned}$$

which can be simplified as

$$\mathbf{M}\delta\mathbf{x} = \mathbf{b}, \quad (18)$$

the operator and the right-hand side are given by

$$\begin{aligned} \mathbf{M} &:= \mathbf{J}_x^\top (\mathbf{I} - \mathbf{J}_w (\mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(\mathbf{w}))^{-1} \mathbf{J}_w^\top) \mathbf{J}_x + \frac{1}{\mu} \nabla_x^2 R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k), \\ \mathbf{b} &:= -\mathbf{J}_x^\top (\mathbf{I} - \mathbf{J}_w (\mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(\mathbf{w}))^{-1} \mathbf{J}_w^\top) \mathbf{r} - \frac{1}{\mu} \nabla_x R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ &\quad + \mathbf{J}_x^\top \mathbf{J}_w (\mathbf{J}_w^\top \mathbf{J}_w + \nabla_w^2 S(\mathbf{w}))^{-1} \nabla_w S(\mathbf{w}). \end{aligned}$$

Moreover, it is easy to see that the gradient and Hessian of $R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ satisfy

$$\begin{aligned} \nabla_x R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) &= \sum_{i=1}^{n^2} (D_i^\top (\boldsymbol{\lambda}^k)_i - \beta D_i^\top (\mathbf{y}_i^{k+1} - D_i \mathbf{x})) \\ &= D^\top \boldsymbol{\lambda}^k - \beta D^\top \mathbf{y}^{k+1} + \beta D^\top D \mathbf{x}, \end{aligned} \quad (19)$$

and

$$\nabla_x^2 R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) = \sum_{i=1}^{n^2} \beta D_i^\top D_i = \beta D^\top D. \quad (20)$$

Following the above procedures, we are able to compute the unconstrained update steps $\delta\mathbf{x}$ and $\delta\mathbf{w}$.

Next, we consider modifying the above procedures in order to handle the element-wise bound constraints on $\mathbf{x}$ and $\mathbf{w}$. We know that a simple extension to the Gauss–Newton method does not work [12], however, a simple correction, i.e. the projected Gauss–Newton method [8] can be made convergent and it works well for many inverse problems. To this end, the variables are divided into active set variables and inactive set variables and the step $\delta\mathbf{x}$ and $\delta\mathbf{w}$ are computed through these two separate sets. Let the feasible index set be defined as

$$\mathcal{N} := \left\{ q \in \mathbb{N} \mid \begin{bmatrix} \mathbf{x}_q \\ \mathbf{w}_q \end{bmatrix} \geq \boldsymbol{\theta} \right\}.$$

Define the active and inactive sets as

$$\mathcal{A} := \left\{ q \in \mathbb{N} \mid \begin{bmatrix} \mathbf{x}_q \\ \mathbf{w}_q \end{bmatrix} = \boldsymbol{\theta} \right\},$$

$$\mathcal{I} := \mathcal{N} \setminus \mathcal{A}.$$

Then, we can divide the variables into the active set variables $\begin{bmatrix} \mathbf{x}_{\mathcal{A}} \\ \mathbf{w}_{\mathcal{A}} \end{bmatrix}$ and the inactive set variables $\begin{bmatrix} \mathbf{x}_{\mathcal{I}} \\ \mathbf{w}_{\mathcal{I}} \end{bmatrix}$. Denote the steps taken on $\begin{bmatrix} \mathbf{x}_{\mathcal{A}} \\ \mathbf{w}_{\mathcal{A}} \end{bmatrix}$ by $\begin{bmatrix} \delta\mathbf{x}_{\mathcal{A}} \\ \delta\mathbf{w}_{\mathcal{A}} \end{bmatrix}$ and the steps taken on $\begin{bmatrix} \mathbf{x}_{\mathcal{I}} \\ \mathbf{w}_{\mathcal{I}} \end{bmatrix}$ by $\begin{bmatrix} \delta\mathbf{x}_{\mathcal{I}} \\ \delta\mathbf{w}_{\mathcal{I}} \end{bmatrix}$, respectively.

$\delta \mathbf{x}_I$ and $\delta \mathbf{w}_I$ can be computed in a similar way as the unconstrained case except that the variables need to be projected onto the inactive set. That is, $\delta \mathbf{x}_I$ at the current iterate $(\mathbf{x}, \mathbf{w})$ is computed as:

$$\hat{\mathbf{M}} \delta \mathbf{x}_I = \hat{\mathbf{b}}, \quad (21)$$

the operator and right-hand side are given by

$$\begin{aligned} \hat{\mathbf{M}} &:= \hat{\mathbf{J}}_x^\top (\mathbf{I} - \hat{\mathbf{J}}_w (\hat{\mathbf{J}}_x^\top \hat{\mathbf{J}}_w + \nabla_w^2 \hat{S}(w))^{-1} \hat{\mathbf{J}}_w^\top) \hat{\mathbf{J}}_x + \frac{1}{\mu} \nabla_x^2 \hat{R}(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k), \\ \hat{\mathbf{b}} &:= -\hat{\mathbf{J}}_x^\top (\mathbf{I} - \hat{\mathbf{J}}_w (\hat{\mathbf{J}}_x^\top \hat{\mathbf{J}}_w + \nabla_w^2 \hat{S}(w))^{-1} \hat{\mathbf{J}}_w^\top) \mathbf{r} - \frac{1}{\mu} \nabla_x \hat{R}(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ &\quad + \hat{\mathbf{J}}_x^\top \hat{\mathbf{J}}_w (\hat{\mathbf{J}}_w^\top \hat{\mathbf{J}}_w + \nabla_w^2 \hat{S}(w))^{-1} \nabla_w \hat{S}(w), \end{aligned}$$

where $\hat{\mathbf{J}}_x$, $\hat{\mathbf{J}}_w$, $\nabla_x \hat{R}$, $\nabla_x^2 \hat{R}$, $\nabla_w \hat{S}$ and $\nabla_w^2 \hat{S}$ represent $\mathbf{J}_x$, $\mathbf{J}_w$, $\nabla_x R$, $\nabla_x^2 R$, $\nabla_w S$ and $\nabla_w^2 S$ restricted to the inactive set via projection, respectively. The reduced problem (21) does not need to be solved to a high accuracy. For example in [10], a stopping tolerance of 10^{-1} is used to solve the reduced problem iteratively. After solving for $\delta \mathbf{x}_I$, $\delta \mathbf{w}_I$ can be computed by

$$\delta \mathbf{w}_I = -\left(\hat{\mathbf{J}}_w^\top \hat{\mathbf{J}}_w + \nabla_w^2 \hat{S}(w)\right)^{-1} \left(\hat{\mathbf{J}}_w^\top \hat{\mathbf{J}}_x \delta \mathbf{x}_I + \hat{\mathbf{J}}_w^\top \mathbf{r} + \nabla_w \hat{S}(w)\right). \quad (22)$$

For the active set, $\delta \mathbf{x}_A$ and $\delta \mathbf{w}_A$ are given by a scaled projected gradient descent step

$$\begin{bmatrix} \delta \mathbf{x}_A \\ \delta \mathbf{w}_A \end{bmatrix} = -\mu \begin{bmatrix} \tilde{\mathbf{J}}_x^\top \mathbf{r} \\ \tilde{\mathbf{J}}_w^\top \mathbf{r} \end{bmatrix} - \begin{bmatrix} \nabla_x \tilde{R}(\mathbf{x} + \delta \mathbf{x}_A, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ \nabla_w \tilde{S}(w) \end{bmatrix}, \quad (23)$$

where $\tilde{\mathbf{J}}_x$, $\tilde{\mathbf{J}}_w$, $\nabla_x \tilde{R}$ and $\nabla_w \tilde{S}$ represent the projection of $\mathbf{J}_x$, $\mathbf{J}_w$, $\nabla_x R$ and $\nabla_w S$ onto the active set, respectively.

Then $\delta \mathbf{x}$ and $\delta \mathbf{w}$ can be calculated as a scaled combination of $\delta \mathbf{x}_A$, $\delta \mathbf{w}_A$, $\delta \mathbf{x}_I$ and $\delta \mathbf{w}_I$ by

$$\begin{bmatrix} \delta \mathbf{x} \\ \delta \mathbf{w} \end{bmatrix} = \begin{bmatrix} \delta \mathbf{x}_I \\ \delta \mathbf{w}_I \end{bmatrix} + \gamma \begin{bmatrix} \delta \mathbf{x}_A \\ \delta \mathbf{w}_A \end{bmatrix}, \quad (24)$$

and the parameter γ is selected based on the recommendation in [8]

$$\gamma = \frac{\max(\|\delta \mathbf{x}_I\|_\infty, \|\delta \mathbf{w}_I\|_\infty)}{\max(\|\delta \mathbf{x}_A\|_\infty, \|\delta \mathbf{w}_A\|_\infty)}.$$

Finally, the projected Armijo line search is applied to find the solution. Here the modified Armijo condition is given by

$$\begin{aligned} &\hat{\Phi}(\mathbf{P}_{C_x}(\mathbf{x} + \eta \delta \mathbf{x}), \mathbf{P}_{C_w}(\mathbf{w} + \eta \delta \mathbf{w}), \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \\ &\leq \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) + c\eta \mathbf{Q}(\nabla_{x,w} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k))^\top \begin{bmatrix} \delta \mathbf{x} \\ \delta \mathbf{w} \end{bmatrix}, \end{aligned} \quad (25)$$

where $\mathbf{P}_{C_x}$ and $\mathbf{P}_{C_w}$ denotes the projections onto the feasible set for the image and parameters, respectively, $\mathbf{Q}(\nabla_{x,w} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k))$ denotes the projected gradient, $0 < \eta \leq 1$ denotes the step size by backtracking, and we set $c = 10^{-4}$ as suggested in [15]. We point out that it is necessary to update the inactive set and the active set in each iteration because the projection

Algorithm 1. ADMM-LAP method for (7).

Input: $(\mathbf{y}^0, \mathbf{x}^0, \mathbf{w}^0, \boldsymbol{\lambda}^0) \in \mathbb{R}^{2n^2} \times \mathbb{R}^{n^2} \times \mathbb{R}^p \times \mathbb{R}^{2n^2}$, the penalty parameter $\beta > 0$.

Let $\{\epsilon_{k+1}\}_{k=0}^\infty$ be a sequence satisfying $\{\epsilon_{k+1}\}_{k=0}^\infty \subseteq [0, +\infty)$ and $\sum_{k=0}^\infty \epsilon_{k+1} < \infty$. Set $k = 0$.

Output: $\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k$.

Step 1 Compute $\mathbf{y}^{k+1}$ using (9).

Step 2 Find a minimizer of

$$\min_{\mathbf{x} \in \mathcal{C}_x, \mathbf{w} \in \mathcal{C}_w} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k).$$

Specifically,

1. Compute the step on the inactive set by the LAP method using (21) and (22).
2. Compute the step on the active set by (23) using the projected gradient descent method.
3. Combine the steps using (24).
4. Perform the projected Armijo line search satisfying (25) to update $\mathbf{x}^{k+1}, \mathbf{w}^{k+1}$.
5. Update active and inactive sets.

Step 3 if the residual $\begin{bmatrix} \eta_x^{k+1} \\ \eta_w^{k+1} \end{bmatrix} := \nabla_{(\mathbf{x}, \mathbf{w})} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ satisfies $\left\| \begin{bmatrix} \eta_x^{k+1} \\ \eta_w^{k+1} \end{bmatrix} \right\|_2 \leq \epsilon_{k+1}$, stop and go to step 4; else, repeat step 2.

Step 4 Compute $\boldsymbol{\lambda}^{k+1}$ using (11).

Step 5 If a termination criteria is met, stop; else, set $k := k + 1$ and go to step 1.

does not prevent variables from leaving the active set and joining the inactive set. Hence, it makes sense that in algorithm 1, if the termination criterion is not met in step 3, we go back to step 2 and get a new solution.

3.3. ADMM-LAP method

The proposed ADMM-LAP method for solving myopic deconvolution problems with TV regularization is summarized in algorithm 1.

Notice that ADMM-LAP is an inexact ADMM, where the $\mathbf{y}$ -subproblems are exactly solved while the $(\mathbf{x}, \mathbf{w})$ -subproblems are inexactly solved. Inspired by the recent results in [14, 22], we prove that the ADMM-LAP algorithm converges subsequently to a stationary point.

Theorem 1. Let $(\mathbf{y}^0, \mathbf{x}^0, \mathbf{w}^0, \boldsymbol{\lambda}^0)$ be any initial point and $\{(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)\}$ be the sequence of iterates generated by algorithm 1. Then if $\beta > \frac{1}{a_1} \lambda_{\max}(A(\mathbf{w}^k)^T A(\mathbf{w}^k))$ and β satisfies $\beta^2 a_1 - L\beta - 4a_2 C_0 > 0$, where C_0, a_1, a_2, L are constants specified in lemmas 1 and 2, algorithm 1 converges subsequently, i.e., it generates a sequence that has a convergent subsequence, whose limit $(\mathbf{y}^*, \mathbf{x}^*, \mathbf{w}^*, \boldsymbol{\lambda}^*)$ is a stationary point of $\mathcal{L}_\beta$. That is, $0 \in \partial \mathcal{L}_\beta(\mathbf{y}^*, \mathbf{x}^*, \mathbf{w}^*, \boldsymbol{\lambda}^*)$.

To prove theorem 1, we define the following functions:

$$\begin{aligned} F(\mathbf{y}) &= \sum_{i=1}^{n^2} \|\mathbf{y}_i\|_2, \\ G(\mathbf{x}, \mathbf{w}) &= \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}) + \delta_{\mathcal{C}_x}(\mathbf{x}) + \delta_{\mathcal{C}_w}(\mathbf{w}), \\ g(\mathbf{x}, \mathbf{w}) &= \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}). \end{aligned}$$

Lemma 1. The iterates of algorithm 1 satisfy:

$$(a) \quad \nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) = -\sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^{k+1} + \eta_x^{k+1}.$$

(b) $\|\sum_{i=1}^{n^2} D_i^T \lambda_i^{k+1} - \sum_{i=1}^{n^2} D_i^T \lambda_i^k\|_2 \leq C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \|\eta_x^{k+1} - \eta_x^k\|_2$, where C_0 is a constant.

Proof. First, notice that for $(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})$ generated by ADMM, we have $\mathbf{x}^{k+1} = [x_1^{k+1}; x_2^{k+1}; \dots; x_{n^2}^{k+1}] \in \mathcal{C}_x$, $\mathbf{w}^{k+1} = [w_1^{k+1}; w_2^{k+1}; \dots; w_p^{k+1}] \in \mathcal{C}_w$. From the definition of the general subgradient, for all $\mathbf{v}_x \in \partial \delta_{\mathcal{C}_x}(\mathbf{x}^{k+1})$,

$$\delta_{\mathcal{C}_x}(\mathbf{x}) - \delta_{\mathcal{C}_x}(\mathbf{x}^{k+1}) \geq \mathbf{v}_x \cdot (\mathbf{x} - \mathbf{x}^{k+1}), \quad \forall \mathbf{x} \in \mathbb{R}^{n^2}.$$

For any $\mathbf{x} = [x_1; x_2; \dots; x_{n^2}]$, where $x_i \geq x_i^{k+1} \geq 0, i = 1, 2, \dots, n^2$, we have $\mathbf{0} \geq \mathbf{v}_x \cdot (\mathbf{x} - \mathbf{x}^{k+1})$. For any $\mathbf{x} = [x_1; x_2; \dots; x_{n^2}]$, where $0 \leq x_i \leq x_i^{k+1}, i = 1, 2, \dots, n^2$, we have $\mathbf{0} \leq \mathbf{v}_x \cdot (\mathbf{x} - \mathbf{x}^{k+1})$. Thus, we have $\mathbf{v}_x = \mathbf{0}$. Similarly, $\mathbf{v}_w = \mathbf{0}$. By the first-order optimality condition at $(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})$,

$$\begin{aligned} & \begin{bmatrix} \nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) \\ \nabla_w g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) \end{bmatrix} + \begin{bmatrix} \sum_{i=1}^{n^2} D_i^T \lambda_i^k \\ \mathbf{0} \end{bmatrix} - \begin{bmatrix} \sum_{i=1}^{n^2} \beta D_i^T (y_i^{k+1} - D_i \mathbf{x}) \\ \mathbf{0} \end{bmatrix} \\ & - \begin{bmatrix} \eta_x^{k+1} \\ \eta_w^{k+1} \end{bmatrix} \in \partial \delta_{\mathcal{C}_x \times \mathcal{C}_w}(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}). \end{aligned}$$

Hence,

$$\nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) + \sum_{i=1}^{n^2} (D_i^T \lambda_i^k - \beta D_i^T (y_i^{k+1} - D_i \mathbf{x})) - \eta_x^{k+1} = \mathbf{0}.$$

Then by the update of the multiplier $\lambda^{k+1} = \lambda^k - \beta (y^{k+1} - D\mathbf{x}^{k+1})$, we obtain

$$\nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) = - \sum_{i=1}^{n^2} D_i^T \lambda_i^{k+1} + \eta_x^{k+1}. \quad (26)$$

We make use of the decomposition

$$\begin{aligned} \|\nabla_x g(\mathbf{x}^k, \mathbf{w}^k) - \nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})\|_2 & \leq \|\nabla_x g(\mathbf{x}^k, \mathbf{w}^k) - \nabla_x g(\mathbf{x}^k, \mathbf{w}^{k+1})\|_2 \\ & \quad + \|\nabla_x g(\mathbf{x}^k, \mathbf{w}^{k+1}) - \nabla_x g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})\|_2. \end{aligned} \quad (27)$$

For the term $\|\nabla_x g(\mathbf{x}^k, \mathbf{w}^k) - \nabla_x g(\mathbf{x}^k, \mathbf{w}^{k+1})\|_2$, we have the estimate

$$\begin{aligned} & \|\nabla_x g(\mathbf{x}^k, \mathbf{w}^k) - \nabla_x g(\mathbf{x}^k, \mathbf{w}^{k+1})\|_2 \\ & = \|A(\mathbf{w}^k)^T A(\mathbf{w}^k) \mathbf{x}^k - A(\mathbf{w}^k)^T \mathbf{d} - A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1}) \mathbf{x}^k + A(\mathbf{w}^{k+1})^T \mathbf{d}\|_2 \\ & \leq \|A(\mathbf{w}^k)^T A(\mathbf{w}^k) \mathbf{x}^k - A(\mathbf{w}^{k+1})^T A(\mathbf{w}^k) \mathbf{x}^k\|_2 \\ & \quad + \|A(\mathbf{w}^{k+1})^T A(\mathbf{w}^k) \mathbf{x}^k - A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1}) \mathbf{x}^k\|_2 \\ & \quad + \|A(\mathbf{w}^k)^T \mathbf{d} - A(\mathbf{w}^{k+1})^T \mathbf{d}\|_2. \end{aligned} \quad (28)$$

Then, we estimate the terms in (28) separately with the two-weight case, i.e., $\mathbf{w} = [w; 1 - w]$. Cases with $p \geq 3$ can be considered similarly, here we omit it. First,

$$\begin{aligned}
& \|A(\mathbf{w}^k)^T A(\mathbf{w}^k) \mathbf{x}^k - A(\mathbf{w}^{k+1})^T A(\mathbf{w}^k) \mathbf{x}^k\|_2 \\
& \leq \|A(\mathbf{w}^k)^T - A(\mathbf{w}^{k+1})^T\|_2 \cdot \|A(\mathbf{w}^k) \mathbf{x}^k\|_2 \\
& \leq \|w^k A_1^T + (1 - w^k) A_2^T - w^{k+1} A_1^T - (1 - w^{k+1}) A_2^T\|_2 \cdot \|A(\mathbf{w}^k) \mathbf{x}^k\|_2 \\
& \leq \|A_1^T - A_2^T\|_2 \cdot \|A(\mathbf{w}^k) \mathbf{x}^k\|_2 \cdot \|w^k - w^{k+1}\|_2, \\
& = \lambda_{\max}(A_1^T - A_2^T) \cdot \|A(\mathbf{w}^k) \mathbf{x}^k\|_2 \cdot \|w^k - w^{k+1}\|_2, \\
& = c_1 \|w^k - w^{k+1}\|_2,
\end{aligned} \tag{29}$$

where $c_1 := \lambda_{\max}(A_1^T - A_2^T) \cdot \|A(\mathbf{w}^k) \mathbf{x}^k\|_2$ is a constant. Similarly,

$$\begin{aligned}
& \|A(\mathbf{w}^{k+1})^T A(\mathbf{w}^k) \mathbf{x}^k - A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1}) \mathbf{x}^k\|_2 \\
& \leq \lambda_{\max}(A_1 - A_2) \cdot \|A(\mathbf{w}^{k+1})^T\|_2 \cdot \|\mathbf{x}^k\|_2 \cdot \|w^k - w^{k+1}\|_2, \\
& = c_2 \|w^k - w^{k+1}\|_2,
\end{aligned} \tag{30}$$

where $c_2 := \lambda_{\max}(A_1 - A_2) \cdot \|A(\mathbf{w}^{k+1})^T\|_2 \cdot \|\mathbf{x}^k\|_2$ is a constant.

Moreover,

$$\begin{aligned}
& \|A(\mathbf{w}^k)^T \mathbf{d} - A(\mathbf{w}^{k+1})^T \mathbf{d}\|_2 \leq \|(A_1 - A_2)^T \mathbf{d}\|_2 \cdot \|w^k - w^{k+1}\|_2 \\
& = c_3 \|w^k - w^{k+1}\|_2,
\end{aligned} \tag{31}$$

where $c_3 := \|(A_1 - A_2)^T \mathbf{d}\|_2$ is a constant.

For the term $\|\nabla_{\mathbf{x}} g(\mathbf{x}^k, \mathbf{w}^{k+1}) - \nabla_{\mathbf{x}} g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})\|_2$, we have

$$\begin{aligned}
& \|\nabla_{\mathbf{x}} g(\mathbf{x}^k, \mathbf{w}^{k+1}) - \nabla_{\mathbf{x}} g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1})\|_2 = \mu \|A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1})(\mathbf{x}^k - \mathbf{x}^{k+1})\|_2 \\
& \leq \mu \lambda_{\max}(A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1})) \cdot \|\mathbf{x}^k - \mathbf{x}^{k+1}\|_2,
\end{aligned} \tag{32}$$

Then we know from (26)–(32) that there exists a constant C_0 such that

$$\begin{aligned}
& \left\| \sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^{k+1} - \sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^k \right\|_2 - \|\eta_x^k - \eta_x^{k+1}\|_2 \\
& \leq \left\| -\sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^{k+1} + \eta_x^{k+1} + \sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^k - \eta_x^k \right\|_2 \\
& = \|\nabla_{\mathbf{x}} g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) - \nabla_{\mathbf{x}} g(\mathbf{x}^k, \mathbf{w}^k)\|_2 \\
& \leq C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2,
\end{aligned}$$

where $C_0 := \sqrt{2} \max\{c_1 + c_2 + c_3, \mu \lambda_{\max}(A(\mathbf{w}^{k+1})^T A(\mathbf{w}^{k+1}))\}$. Hence, we have

$$\left\| \sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^{k+1} - \sum_{i=1}^{n^2} D_i^T \boldsymbol{\lambda}_i^k \right\|_2 \leq C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 + \|\eta_x^{k+1} - \eta_x^k\|_2.$$

□

Lemma 2. Let $\{(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)\}$ be the sequence of iterates generated by algorithm 1. If $\beta > \frac{1}{a_1} \lambda_{\max}(A(\mathbf{w}^k)^T A(\mathbf{w}^k))$ and β satisfies $\beta^2 a_1 - L\beta - 4a_2 C_0 > 0$, where a_1, a_2 are constants depending on D , L denotes the Lipschitz constant of $-g(\mathbf{x}^{k+1}, \mathbf{w}^k)$, then $\{(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)\}$ satisfies:

(a) $\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)$ is lower bounded and there is a constant $C_1 > 0$ such that for all sufficiently large k , we have

$$\begin{aligned} & \mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k) - \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) \\ & \geq C_1 \left(\sum_{i=1}^{n^2} \|\mathbf{y}_i^k - \mathbf{y}_i^{k+1}\|_2^2 + \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 \right) - \frac{2}{\beta} \|\eta_{\mathbf{x}}^{k+1} - \eta_{\mathbf{x}}^k\|_2^2. \end{aligned}$$

(b) $\{(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)\}$ is bounded.

Proof. According to the optimality condition of the $\mathbf{y}$ -subproblem, we define

$$d_i^{k+1} := (\boldsymbol{\lambda}^k)_i + \beta(\mathbf{y}_i^{k+1} - D_i \mathbf{x}^k) \in \partial \|\mathbf{y}_i^{k+1}\|_2.$$

We know from the definition of $\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)$ that

$$\begin{aligned} & \mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k) - \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k) \\ & = F(\mathbf{y}^k) - F(\mathbf{y}^{k+1}) - \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^k)_i^T (\mathbf{y}_i^k - \mathbf{y}_i^{k+1}) + \frac{\beta}{2} \sum_{i=1}^{n^2} \|\mathbf{y}_i^k - D_i \mathbf{x}^k\|_2^2 - \frac{\beta}{2} \sum_{i=1}^{n^2} \|\mathbf{y}_i^{k+1} - D_i \mathbf{x}^k\|_2^2 \\ & = F(\mathbf{y}^k) - F(\mathbf{y}^{k+1}) - \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^k)_i^T (\mathbf{y}_i^k - \mathbf{y}_i^{k+1}) + \frac{\beta}{2} \sum_{i=1}^{n^2} (\|\mathbf{y}_i^k - \mathbf{y}_i^{k+1}\|_2^2 + 2\langle \mathbf{y}_i^{k+1} - D_i \mathbf{x}^k, \mathbf{y}_i^k - \mathbf{y}_i^{k+1} \rangle) \\ & = \sum_{i=1}^{n^2} (\|\mathbf{y}_i^k\|_2 - \|\mathbf{y}_i^{k+1}\|_2 - \langle d_i^{k+1}, \mathbf{y}_i^k - \mathbf{y}_i^{k+1} \rangle) + \frac{\beta}{2} \sum_{i=1}^{n^2} \|\mathbf{y}_i^k - \mathbf{y}_i^{k+1}\|_2^2 \\ & \geq \frac{\beta}{2} \sum_{i=1}^{n^2} \|\mathbf{y}_i^k - \mathbf{y}_i^{k+1}\|_2^2, \end{aligned}$$

where the second equality follows from the cosine rule: $\|b + c\|^2 - \|a + c\|^2 = \|b - a\|^2 + 2\langle a + c, b - a \rangle$ and the last inequality follows from the convexity of $\|\mathbf{y}\|_2$.

Moreover, we have

$$\begin{aligned} & \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k) - \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) \\ & = g(\mathbf{x}^k, \mathbf{w}^k) - g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) - \sum_{i=1}^{n^2} ((\boldsymbol{\lambda}^k)_i^T - (\boldsymbol{\lambda}^{k+1})_i^T) \mathbf{y}_i^{k+1} + \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^k)_i^T D_i \mathbf{x}^k \\ & \quad - \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^{k+1})_i^T D_i \mathbf{x}^{k+1} + \sum_{i=1}^{n^2} \frac{\beta}{2} \|\mathbf{y}_i^{k+1} - D_i \mathbf{x}^k\|_2^2 + \sum_{i=1}^{n^2} \frac{\beta}{2} \|\mathbf{y}_i^{k+1} - D_i \mathbf{x}^{k+1}\|_2^2 \\ & = g(\mathbf{x}^k, \mathbf{w}^k) - g(\mathbf{x}^k, \mathbf{w}^{k+1}) + g(\mathbf{x}^k, \mathbf{w}^{k+1}) - g(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}) - \sum_{i=1}^{n^2} ((\boldsymbol{\lambda}^k)_i^T - (\boldsymbol{\lambda}^{k+1})_i^T) \mathbf{y}_i^{k+1} \\ & \quad + \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^k)_i^T D_i \mathbf{x}^k - \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^{k+1})_i^T D_i \mathbf{x}^k + \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^{k+1})_i^T D_i \mathbf{x}^k - \sum_{i=1}^{n^2} (\boldsymbol{\lambda}^{k+1})_i^T D_i \mathbf{x}^k \end{aligned}$$

$$\begin{aligned}
& + \frac{\beta}{2} \sum_{i=1}^{n^2} (\| -D_i \mathbf{x}^k + D_i \mathbf{x}^{k+1} \|_2 + 2 \langle -D_i \mathbf{x}^{k+1} + \mathbf{y}_i^{k+1}, -D_i \mathbf{x}^k + D_i \mathbf{x}^{k+1} \rangle) \\
& \geq -\frac{L}{2} \| \mathbf{w}^{k+1} - \mathbf{w}^k \|_2^2 + \frac{\beta}{2} \sum_{i=1}^{n^2} (\| -D_i \mathbf{x}^k + D_i \mathbf{x}^{k+1} \|_2 + 2 \langle -D_i \mathbf{x}^{k+1} + \mathbf{y}_i^{k+1}, -D_i \mathbf{x}^k + D_i \mathbf{x}^{k+1} \rangle) \\
& \quad + \sum_{i=1}^{n^2} \beta \langle \mathbf{y}_i^{k+1} - D_i \mathbf{x}^{k+1}, -\mathbf{y}_i^{k+1} + D_i \mathbf{x}^{k+1} \rangle \\
& = -\frac{L}{2} \| \mathbf{w}^{k+1} - \mathbf{w}^k \|_2^2 - \frac{1}{\beta} \sum_{i=1}^{n^2} \| (\lambda^{k+1})_i - (\lambda^k)_i \|_2^2 + \frac{\beta}{2} \sum_{i=1}^{n^2} \| -D_i \mathbf{x}^k + D_i \mathbf{x}^{k+1} \|_2^2 \\
& \geq -\frac{L}{2} \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 - \frac{2a_2 C_0}{\beta} \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 + \frac{\beta a_1}{2} \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 - \frac{2}{\beta} \| \eta_{\mathbf{x}}^k - \eta_{\mathbf{x}}^{k+1} \|_2^2 \\
& = C \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 - \frac{2}{\beta} \| \eta_{\mathbf{x}}^k - \eta_{\mathbf{x}}^{k+1} \|_2^2,
\end{aligned}$$

where $C := -\frac{L}{2} - \frac{2a_2 C_0}{\beta} + \frac{\beta a_1}{2}$ is a constant, a_1, a_2 are constants such that for $i = 1, 2, \dots, n^2$, $\|D_i(\mathbf{x}^{k+1} - \mathbf{x}^k)\|_2^2 \geq a_1 \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2$ and $\|(\lambda^{k+1})_i - (\lambda^k)_i\|_2^2 \leq a_2 \|D_i^T((\lambda^{k+1})_i - (\lambda^k)_i)\|_2^2$, respectively. The second equality follows from the cosine rule; the third inequality follows from the convexity of $g(\mathbf{x}^k, \mathbf{w}^{k+1})$ with respect to $\mathbf{x}$ and the Lipschitz differentiable property of $-g(\mathbf{x}^{k+1}, \mathbf{w}^k)$ with respect to $\mathbf{w}$ [22], L is the Lipschitz constant; the fourth equality follows from the definition of the multiplier $\lambda^{k+1} = \lambda^k - \beta(\mathbf{y}^{k+1} - D\mathbf{x}^{k+1})$; the fifth inequality follows from lemma 1. In order to ensure $C > 0$, we require β satisfies $\beta^2 a_1 - L\beta - 4a_2 C_0 > 0$.

Then we know from two equalities above that

$$\begin{aligned}
& \mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \lambda^k) - \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \lambda^{k+1}) \\
& \geq C_1 \left(\sum_{i=1}^{n^2} \| \mathbf{y}_i^k - \mathbf{y}_i^{k+1} \|_2^2 + \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 \right) - \frac{2}{\beta} \| \eta_{\mathbf{x}}^k - \eta_{\mathbf{x}}^{k+1} \|_2^2. \quad (33)
\end{aligned}$$

This means for all sufficiently large k , $\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \lambda^k)$ is nonincreasing. As $D_i \in \mathbb{R}^{2 \times n^2}$ has full row rank, then there exists at one $\hat{\mathbf{x}}$ such that $D_i \hat{\mathbf{x}} = \mathbf{y}_i^k, \forall i = 1, 2, \dots, n^2$. Thus, we arrive at

$$\begin{aligned}
\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \lambda^k) &= \frac{\mu}{2} \|A(\mathbf{w}^k) \mathbf{x}^k - \mathbf{d}\|_2^2 + \sum_{i=1}^{n^2} \| \mathbf{y}_i^k \|_2^2 - \sum_{i=1}^{n^2} (\lambda^k)_i^T (\mathbf{y}_i^k - D_i \mathbf{x}^k) + \frac{\beta}{2} \sum_{i=1}^{n^2} \| \mathbf{y}_i^k - D_i \mathbf{x}^k \|_2^2 \\
&= \sum_{i=1}^{n^2} \| \mathbf{y}_i^k \|_2^2 + \frac{\beta}{2} \sum_{i=1}^{n^2} \| \mathbf{y}_i^k - D_i \mathbf{x}^k \|_2^2 + \frac{\mu}{2} \|A(\mathbf{w}^k) \mathbf{x}^k - \mathbf{d}\|_2^2 - \sum_{i=1}^{n^2} \langle D_i^T(\lambda^k)_i, \hat{\mathbf{x}} - \mathbf{x}^k \rangle, \\
&= \sum_{i=1}^{n^2} \| \mathbf{y}_i^k \|_2^2 + \frac{\beta}{2} \sum_{i=1}^{n^2} \| \mathbf{y}_i^k - D_i \mathbf{x}^k \|_2^2 + \frac{\mu}{2} \|A(\mathbf{w}^k) \hat{\mathbf{x}} - \mathbf{d}\|_2^2 \\
&= \sum_{i=1}^{n^2} \| \mathbf{y}_i^k \|_2^2 + \frac{\mu}{2} \|A(\mathbf{w}^k) \hat{\mathbf{x}} - \mathbf{d}\|_2^2 + \frac{\beta}{2} \sum_{i=1}^{n^2} \|D_i(\hat{\mathbf{x}} - \mathbf{x}^k)\|_2^2 \\
&\geq \sum_{i=1}^{n^2} \| \mathbf{y}_i^k \|_2^2 + \frac{\mu}{2} \|A(\mathbf{w}^k) \hat{\mathbf{x}} - \mathbf{d}\|_2^2 + \frac{\beta a_1}{2} \| \hat{\mathbf{x}} - \mathbf{x}^k \|_2^2 > -\infty.
\end{aligned}$$

Since $\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)$ is lower bounded and $\sum_{i=1}^{n^2} \|\mathbf{y}_i^k\|_2 + \frac{\mu}{2} \|A(\mathbf{w}^k)\hat{\mathbf{x}}^k\|_2^2 + (\frac{1}{2}\beta a_1 - \frac{1}{2}\lambda_{\max}(A(\mathbf{w}^k)^T A(\mathbf{w}^k))) \|\hat{\mathbf{x}} - \mathbf{x}^k\|_2^2$ is coercive over the feasible set $\Omega_F := \{(\mathbf{y}, \mathbf{x}, \mathbf{w}) : D_i \mathbf{x} = \mathbf{y}_i, i = 1, 2, \dots, n^2\}$, we conclude that $\{\mathbf{y}^k\}$ and $\{(\mathbf{x}^k, \mathbf{w}^k)\}$ are bounded. Hence, by lemma 1, $\{\boldsymbol{\lambda}^k\}$ is bounded. $\square$

Lemma 3. Let $\partial \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) = (\partial_{\mathbf{y}} \mathcal{L}_\beta, \nabla_{(\mathbf{x}, \mathbf{w})} \mathcal{L}_\beta, \nabla_{\boldsymbol{\lambda}} \mathcal{L}_\beta)$. Then there exists a constant $\tilde{C} > 0$ such that for all $k \geq 1$, for some $\mathbf{p}^{k+1} \in \partial \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1})$, we have

$$\|\mathbf{p}^{k+1}\|_2 \leq \tilde{C} \left(\left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \|\eta_{\mathbf{x}}^{k+1} - \eta_{\mathbf{x}}^k\|_2 \right).$$

Proof. Because

$$\begin{aligned} \nabla_{\lambda_i} L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) &= -(\mathbf{y}_i^{k+1} - D_i \mathbf{x}^{k+1}) \\ &= \frac{1}{\beta} ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i), \end{aligned}$$

and

$$\nabla_{(\mathbf{x}, \mathbf{w})} L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) = \begin{bmatrix} \sum_{i=1}^{n^2} D_i^T ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i) \\ \mathbf{0} \end{bmatrix},$$

we have

$$\begin{aligned} \|\nabla_{\boldsymbol{\lambda}} L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1})\|_2 &= \frac{1}{\beta} \sum_{i=1}^{n^2} \|(\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i\|_2 \\ &\leq \frac{1}{\beta \sqrt{a_1}} \sum_{i=1}^{n^2} \|D_i^T ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i)\|_2 \\ &\leq \frac{1}{\beta \sqrt{a_1}} \left(C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \|\eta_{\mathbf{x}}^{k+1} - \eta_{\mathbf{x}}^k\|_2 \right) \end{aligned}$$

and

$$\|\nabla_{(\mathbf{x}, \mathbf{w})} L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1})\|_2 \leq C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \|\eta_{\mathbf{x}}^{k+1} - \eta_{\mathbf{x}}^k\|_2.$$

By the definition of the subgradient, we make use of the decomposition and obtain

$$\begin{aligned} \partial_{\mathbf{y}} L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) &= \partial \sum_{i=1}^{n^2} (\|\mathbf{y}_i^{k+1}\|_2 - (\boldsymbol{\lambda}^{k+1})_i + \beta(\mathbf{y}_i^{k+1} - D_i \mathbf{x}^{k+1})) \\ &= \partial \sum_{i=1}^{n^2} (\|\mathbf{y}_i^{k+1}\|_2 - (\boldsymbol{\lambda}^k)_i + \beta(\mathbf{y}_i^{k+1} - D_i \mathbf{x}^k) - (\boldsymbol{\lambda}^{k+1})_i \\ &\quad + (\boldsymbol{\lambda}^k)_i + \beta(D_i \mathbf{x}^k - D_i \mathbf{x}^{k+1})). \end{aligned}$$

Thus, according to the optimal condition

$$0 \in \partial \|\mathbf{y}_i^{k+1}\|_2 - (\boldsymbol{\lambda}^k)_i + \beta(D_i \mathbf{x}^k - D_i \mathbf{x}^{k+1}),$$

we have

$$-\sum_{i=1}^{n^2} ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i) + \beta(D_i \mathbf{x}^k - D_i \mathbf{x}^{k+1}) \in \partial_y L_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}).$$

Letting

$$\mathbf{p}^{k+1} := (\mathbf{p}_y^{k+1}, \mathbf{p}_{(\mathbf{x}, \mathbf{w})}^{k+1}, \mathbf{p}_\lambda^{k+1}),$$

where

$$\mathbf{p}_y^{k+1} := -\sum_{i=1}^{n^2} ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i) + \beta(D_i \mathbf{x}^k - D_i \mathbf{x}^{k+1}),$$

$$\mathbf{p}_{(\mathbf{x}, \mathbf{w})}^{k+1} := \sum_{i=1}^{n^2} D_i^T ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i),$$

$$\mathbf{p}_\lambda^{k+1} := \frac{1}{\beta} \sum_{i=1}^{n^2} ((\boldsymbol{\lambda}^{k+1})_i - (\boldsymbol{\lambda}^k)_i).$$

Then we have

$$\begin{aligned} \|\mathbf{p}^{k+1}\|_2 &\leq \frac{C_0}{\beta\sqrt{a_1}} \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \frac{C_0}{\sqrt{a_1}} \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + C_0 \left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 \\ &\quad + \beta \sum_{i=1}^{n^2} \|D_i\|_2 \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2 + \left(\frac{C_0}{\beta\sqrt{a_1}} + \frac{C_0}{\sqrt{a_1}} + C_0 \right) \|\eta_x^{k+1} - \eta_x^k\|_2 \\ &\leq \tilde{C} \left(\left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2 + \|\eta_x^{k+1} - \eta_x^k\|_2 \right), \end{aligned}$$

where $\tilde{C} := C_0(1 + \frac{1}{\beta\sqrt{a_1}} \frac{1}{\sqrt{a_1}}) + \beta \sum_{i=1}^{n^2} \|D_i\|_2$ is a constant. $\square$

Now we give the proof to theorem 1.

Proof of theorem 1. By lemma 2, the sequence $\{(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)\}$ is bounded, so there exists a convergent subsequence $\{(\mathbf{y}^{n_k}, \mathbf{x}^{n_k}, \mathbf{w}^{n_k}, \boldsymbol{\lambda}^{n_k})\}$, i.e., $\{(\mathbf{y}^{n_k}, \mathbf{x}^{n_k}, \mathbf{w}^{n_k}, \boldsymbol{\lambda}^{n_k})\}$ converges to $(\mathbf{y}^*, \mathbf{x}^*, \mathbf{w}^*, \boldsymbol{\lambda}^*)$ as $k \rightarrow \infty$. Notice that the residual is summable, then we know from $\|\eta_x^k - \eta_x^{k+1}\|_2 \leq \|\eta_x^k\|_2 + \|\eta_x^{k+1}\|_2$ that when $k \rightarrow \infty$, $\|\eta_x^k - \eta_x^{k+1}\|_2 \rightarrow 0$. Hence, for sufficiently large k , we have $L_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k)$ is nonincreasing and lower-bounded. We derive that when $k \rightarrow \infty$,

$$\mathcal{L}_\beta(\mathbf{y}^k, \mathbf{x}^k, \mathbf{w}^k, \boldsymbol{\lambda}^k) - \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1}) \rightarrow 0.$$

Then we know from the proof of lemma 2 that when $k \rightarrow \infty$,

$$\sum_{i=1}^{n^2} \|\mathbf{y}_i^k - \mathbf{y}_i^{k+1}\|_2^2 \rightarrow 0,$$

$$\left\| \begin{bmatrix} \mathbf{x}^{k+1} - \mathbf{x}^k \\ \mathbf{w}^{k+1} - \mathbf{w}^k \end{bmatrix} \right\|_2^2 \rightarrow 0.$$

Then we know from lemma 3 that there exists $\mathbf{p}^{k+1} \in \partial \mathcal{L}_\beta(\mathbf{y}^{k+1}, \mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \boldsymbol{\lambda}^{k+1})$ such that when $k \rightarrow \infty$, $\|\mathbf{p}^{k+1}\|_2 \rightarrow 0$. This leads to $\|\mathbf{p}^k\|_2 \rightarrow 0$ as $k \rightarrow \infty$. Based on the definition of the general subgradient in [18], we obtain that $0 \in \partial \mathcal{L}_\beta(\mathbf{y}^*, \mathbf{x}^*, \mathbf{w}^*, \boldsymbol{\lambda}^*)$, i.e., $(\mathbf{y}^*, \mathbf{x}^*, \mathbf{w}^*, \boldsymbol{\lambda}^*)$ is a stationary point. $\square$

3.4. Complexity analysis

Now we consider the computational complexity of algorithm 1 for (7). In step 1, the main computational costs come from computing $D_i \mathbf{x}^k$ and $\frac{1}{\beta}(\boldsymbol{\lambda}^k)_i$. Since D_i denotes the first-order finite difference of $\mathbf{x}$ at the i th pixel, computing $D_i \mathbf{x}^k$, $i = 1, \dots, n^2$ requires $2n^2$ flops. Computing the scalar-vector product $\frac{1}{\beta}(\boldsymbol{\lambda}^k)_i$, $i = 1, \dots, n^2$ requires $2n^2$ flops. Finally, computing the two-norm of the vectors $D_i \mathbf{x}^k + \frac{1}{\beta}(\boldsymbol{\lambda}^k)_i \in \mathbb{R}^2$, $i = 1, \dots, n^2$ needs $6n^2$ flops. Thus, step 1 costs $O(n^2)$.

In step 2, we first consider the computational cost of (21). Suppose we use the conjugate gradient method to solve for $\delta \mathbf{x}_T$ and set the maximum iteration number to a small constant. Then the cost is determined by the cost of multiplying $\tilde{\mathbf{J}}_x^\top (\mathbf{I} - \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \tilde{\mathbf{J}}_w^\top) \tilde{\mathbf{J}}_x + \frac{1}{\mu} \nabla_x^2 \hat{R}(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ on a vector and computing $-\tilde{\mathbf{J}}_x^\top (\mathbf{I} - \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \tilde{\mathbf{J}}_w^\top) \mathbf{r}$ and $\tilde{\mathbf{J}}_x^\top \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \nabla_w \hat{S}(\mathbf{w})$. Since $\mathbf{J}_x = A(\mathbf{w})^\top$ and each PSF matrix A_i allows the use of fast Fourier transforms (FFTs) [23, 24] to multiply it with a vector, the cost of multiplying $\tilde{\mathbf{J}}_x$ on a vector is $O(n^2 \log n)$. Similarly, the cost of multiplying $\tilde{\mathbf{J}}_w$ on a vector is also $O(n^2 \log n)$ by FFTs. So the cost of the matrix-vector product of $\tilde{\mathbf{J}}_x^\top (\mathbf{I} - \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \tilde{\mathbf{J}}_w^\top) \tilde{\mathbf{J}}_x$ and a vector is $O(n^2 \log n)$. Computing the matrix-vector product of $\nabla_x^2 \hat{R}(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ and a vector equals to computing the matrix-vector product of $\beta D^\top D$ and a vector, and it can be done in $O(n^2)$. Moreover, $-\tilde{\mathbf{J}}_x^\top (\mathbf{I} - \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \tilde{\mathbf{J}}_w^\top) \mathbf{r}$ and $\tilde{\mathbf{J}}_x^\top \tilde{\mathbf{J}}_w (\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_w + \nabla_w^2 \hat{S}(\mathbf{w}))^{-1} \nabla_w \hat{S}(\mathbf{w})$ can be performed in $O(n^2 \log n)$ by FFTs. Thus the computational cost of (21) is $O(n^2 \log n)$. Next we consider the cost of (22). The main computational costs in (22) are from computing the matrix-vector product $\tilde{\mathbf{J}}_w^\top \tilde{\mathbf{J}}_x \delta \mathbf{x}_T$ and $\tilde{\mathbf{J}}_w \mathbf{r}$. Since these matrix-vector products can also be accelerated by FFTs, the computational cost of (22) is $O(n^2 \log n)$. For (23), the main computational costs are from computing $\tilde{\mathbf{J}}_x^\top \mathbf{r}$, $\tilde{\mathbf{J}}_w^\top \mathbf{r}$. Similar to the analysis for (22), the costs of computing $\tilde{\mathbf{J}}_x^\top \mathbf{r}$ and $\tilde{\mathbf{J}}_w^\top \mathbf{r}$ are $O(n^2 \log n)$ due to FFTs. The scalar-vector product $-\mu \begin{bmatrix} \tilde{\mathbf{J}}_x^\top \mathbf{r} \\ \tilde{\mathbf{J}}_w^\top \mathbf{r} \end{bmatrix}$

can be done in $n^2 + p$ flops. In applications, we usually choose the image size to be $n \times n = 256 \times 256$ and the number of the unknown weights parameterizing the blurring matrix is usually chosen to be $p = 2, 3$ or some small values much smaller than n^2 , so we always have $p \ll n^2$. Thus, the cost of (23) is $O(n^2 \log n)$. In (24), computing the

scalar-vector multiplication $\gamma \begin{bmatrix} \delta \mathbf{x}_A \\ \delta \mathbf{w}_A \end{bmatrix}$ requires $n^2 + p$ flops and computing the summation of the vectors $\begin{bmatrix} \delta \mathbf{x}_T \\ \delta \mathbf{w}_T \end{bmatrix}$ and $\gamma \begin{bmatrix} \delta \mathbf{x}_A \\ \delta \mathbf{w}_A \end{bmatrix}$ also requires $n^2 + p$ flops. Hence, the cost of (24) is $O(n^2)$. The cost of (25) is mainly from forming $\nabla_{x,w} \hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ and computing

the inner product $\left(\nabla_{x,w}\hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)\right)^\top \begin{bmatrix} \delta\mathbf{x} \\ \delta\mathbf{w} \end{bmatrix}$. We know that $\nabla_{x,w}\hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) = \begin{bmatrix} A(\mathbf{w})^\top(A(\mathbf{w})\mathbf{x} - \mathbf{d}) \\ \hat{J}_w^\top(A(\mathbf{w})\mathbf{x} - \mathbf{d}) \end{bmatrix}$, so the cost of forming $\nabla_{x,w}\hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)$ is $O(n^2 \log n)$ by FFTs, and computing the inner product $\left(\nabla_{x,w}\hat{\Phi}(\mathbf{x}, \mathbf{w}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k)\right)^\top \begin{bmatrix} \delta\mathbf{x} \\ \delta\mathbf{w} \end{bmatrix}$ costs $O(n^2)$. Thus the computational cost of step 2 is $O(n^2 \log n)$.

In step 4, the computational cost is dominated by the cost of computing $D\mathbf{x}^{k+1}$ and forming $\beta\mathbf{y}^{k+1}$ and $\beta D\mathbf{x}^{k+1}$. It is easy to see that these operations cost $O(n^2)$. Therefore, the total computational complexity of algorithm 1 is $O(n^2 \log n)$.

3.5. Comparison method

In this section, we briefly discuss a benchmark method ADMM-BCD of ADMM-LAP used in the numerical comparisons in section 4.

BCD in ADMM-BCD stands for the BCD method [15, 25], which is another popular approach for solving coupled optimization problems. The main idea of BCD is partitioning the optimization variables into a number of blocks, then minimize the objective function cyclically over each block while fixing the remaining blocks at their last updated values until convergence. For AO retinal image problems we consider in this paper, the variables can be naturally separated into two subsets, one for the image variable $\mathbf{x}$ and another for the parameters $\mathbf{w}$. For the tightly coupled $(\mathbf{x}, \mathbf{w})$ -subproblem

$$\min_{\mathbf{x} \in \mathcal{C}_x, \mathbf{w} \in \mathcal{C}_w} \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x} - \mathbf{d}\|_2^2 + S(\mathbf{w}) + R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k),$$

given the initial point $(\mathbf{x}^k, \mathbf{w}^k)$, the iterative format of BCD is as follows

$$\mathbf{x}^{(l+1)} = \operatorname{argmin}_{\mathbf{x} \in \mathcal{C}_x} \frac{\mu}{2} \|A(\mathbf{w}^{(l)})\mathbf{x} - \mathbf{d}\|_2^2 + R(\mathbf{x}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k), \quad (34)$$

$$\mathbf{w}^{(l+1)} = \operatorname{argmin}_{\mathbf{w} \in \mathcal{C}_w} \frac{\mu}{2} \|A(\mathbf{w})\mathbf{x}^{(l+1)} - \mathbf{d}\|_2^2 + S(\mathbf{w}), \quad (35)$$

where $l \in \mathbb{N}$ denotes the l th iteration of BCD.

For the numerical experiments, we inexactly solve (34) and (35) by the projected Gauss–Newton method with bound constraints [8]. The search directions $\delta\mathbf{x}_I$ at $\mathbf{x}^{(l)}$ can be computed as follows:

$$\begin{aligned} & \left(\mu \hat{\mathbf{J}}_x^\top \hat{\mathbf{J}}_x + \nabla_x^2 \hat{R}(\mathbf{x}^{(l)}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \right) \delta\mathbf{x}_I \\ &= - \left(\mu \hat{\mathbf{J}}_x^\top \mathbf{r}(\mathbf{x}^{(l)}, \mathbf{w}^{(l)}) + \nabla_x \hat{R}(\mathbf{x}^{(l)}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^k) \right), \end{aligned} \quad (36)$$

where $\hat{\mathbf{J}}_x$, $\nabla_x \hat{R}$ and $\nabla_x^2 \hat{R}$ represent $\mathbf{J}_x$, $\nabla_x R$ and $\nabla_x^2 R$ restricted to the inactive set via projection, respectively. Moreover, $\delta\mathbf{x}_A$ is given by the projected gradient descent step and $\delta\mathbf{x}$ can be computed by a scaled combination [8]. Finally, the projected Armijo line search is applied to find the solution $\mathbf{x}^{(l+1)}$. Then the search directions $\delta\mathbf{w}_I$ at $\mathbf{w}^{(l)}$ can be computed as follows:

$$\left(\hat{\mathbf{J}}_w^\top \hat{\mathbf{J}}_w + \nabla_w^2 S(\mathbf{w}) \right) \delta\mathbf{w}_I = -\hat{\mathbf{J}}_w^\top \mathbf{r}(\mathbf{x}^{(l+1)}, \mathbf{w}^{(l)}) + \nabla_w S(\mathbf{w}), \quad (37)$$

Algorithm 2. ADMM-BCD method for (7).

Input: $(y^0, x^0, w^0, \lambda^0) \in \mathbb{R}^{2n^2} \times \mathbb{R}^{n^2} \times \mathbb{R}^p \times \mathbb{R}^{2n^2}$, the penalty parameter $\beta > 0$. Set $k = 0$.

Output: y^k, x^k, w^k, λ^k .

Step 1 Compute y^{k+1} using (9).

Step 2 Compute x^{k+1}, w^{k+1} by iteratively solving (34) and (35).

Step 3 If a termination criteria is met, stop and go to step 4; else, repeat step 2.

Step 4 Compute λ^{k+1} using (11).

Step 5 If a termination criteria is met, stop; else, set $k := k + 1$ and go to step 1.

where $\hat{\mathbf{J}}_w$, $\nabla_w \hat{S}(w)$ and $\nabla_w^2 \hat{S}(w)$ represent $\mathbf{J}_w$, $\nabla_w S(w)$ and $\nabla_w^2 S(w)$ restricted to the inactive set via projection, respectively. Similar to the above discussion, we can obtain δw_A by projected gradient descent and obtain δw by a scaled combination. Finally, the solution $w^{(l+1)}$ can be obtained by the projected Armijo line search.

For (36), due to the large dimension, the search direction can be inexactly computed by the conjugate gradients method. For (37), because the dimension is small, the search direction can be computed by a direct solver. Moreover, we note that BCD is fully decoupled while optimizing over one set of variables, and the optimization over another set of variables is neglected. This degrades the convergence for tightly coupled problems [15].

Similar to ADMM-LAP discussed in the previous section, ADMM-BCD uses ADMM as the outer optimization method to tackle the TV regularization and applies BCD to tackle the related (x, w) -subproblems appearing in each ADMM iteration. The main operations of ADMM-BCD are summarized in algorithm 2.

The complexity analysis of algorithm 2 is presented as follows. First notice that step 1 and step 4 in algorithm 2 are identical to those in algorithm 1, thus these two steps cost $O(n^2)$.

For step 2 in algorithm 2, we first analyze the computational cost of (34). To compute δx_I , we use the conjugate gradient method and set the maximum iteration number to a small constant. The main computational cost is from computing the multiplication of $\mu \hat{\mathbf{J}}_x^T \hat{\mathbf{J}}_x + \nabla_x^2 \hat{R}(x, y^{k+1}, \lambda^k)$ and a vector, and computing $\hat{\mathbf{J}}_x^T r(x^{(l)}, w^{(l)})$. Similar to the complexity analysis in algorithm 1, this matrix-vector product can be done in $O(n^2 \log n)$ by FFTs. Then δx_A , δx can be computed in a similar way as discussed in algorithm 1 and a projected Armijo line search is applied to find the solution. The cost of these steps is $O(n^2 \log n)$. Hence, the total computational cost of (34) is $O(n^2 \log n)$. As for the computational cost of (35), the main computational costs of computing δw_I are from computing $\hat{\mathbf{J}}_w^T \hat{\mathbf{J}}_w$ and $\hat{\mathbf{J}}_w^T r(x^{(l+1)}, w^{(l)})$. We know from the definition (13) that $\hat{\mathbf{J}}_w^T$ is a $p \times n^2$ matrix, $p \ll n^2$, hence the cost of computing $\hat{\mathbf{J}}_w^T \hat{\mathbf{J}}_w$ is $O(n^2)$. The matrix-vector multiplication $\hat{\mathbf{J}}_w^T r(x^{(l+1)}, w^{(l)})$ can be done in $O(n^2 \log n)$ by FFTs. Moreover, the cost of computing δw_A , δw and applying the projected Armijo line search is $O(n^2 \log n)$. Therefore, the cost of (35) is $O(n^2 \log n)$. As a result, the total computational cost of algorithm 2 is $O(n^2 \log n)$.

As shown from numerical experiments from section 4, algorithm 2 takes longer time than that of algorithm 1. This is mainly because BCD requires computing the matrix-vector product $\hat{\mathbf{J}}_x^T r(x^{(l)}, w^{(l)})$ and $\hat{\mathbf{J}}_w^T r(x^{(l+1)}, w^{(l)})$ in each iteration, while LAP only requires computing $\hat{\mathbf{J}}_w^T r(x^{(l)}, w^{(l)})$. Hence, BCD takes more computational costs to compute the current residual value r and the matrix-vector multiplication $\hat{\mathbf{J}}_w^T r$, which results in a larger number of FFTs. Moreover, the line search is applied twice in BCD to solve the search directions for both x and w , while in LAP, only one line search is enough to obtain the search directions for both variables. Hence, algorithm 2 costs more. Moreover, it is important to point out that algorithm 1 takes the structure of the tightly coupled problem into consideration and converges faster.

4. Numerical experiments

In this section, we illustrate the numerical performance of the proposed ADMM-LAP method for the myopic deconvolution problems with TV regularization arising from the AO retinal image restoration. All our computational results are obtained by MATLAB R2018b running on a Macbook Air with Intel Core i5 CPU (1.4 GHz).

The following notations will be used throughout the section:

- BlurLevel: an indicator used to set the severity of the blur to one of the following: ‘mild’, ‘medium’ and ‘severe’;
- NoiseLevel $\triangleq \|e\|_2 / \|\mathbf{d}^*\|_2$: relative level of noise, where e denotes the noise vector of perturbations and $\mathbf{d}^*$ denotes the exact (noise free) data vector;
- #iter: the number of iterations performed by the algorithm;
- Rel. Err. $\mathbf{x}$: the relative error $\|\mathbf{x} - \mathbf{x}^*\|_2 / \|\mathbf{x}^*\|_2$, where $\mathbf{x}^*$ and $\mathbf{x}$ denote the true image and computed image, respectively;
- Rel. Err. $\mathbf{w}$: the relative error $\|\mathbf{w} - \mathbf{w}^*\|_2 / \|\mathbf{w}^*\|_2$, where $\mathbf{w}^*$ and $\mathbf{w}$ denote the true parameters and computed parameters, respectively;
- SNR($\mathbf{x}$) $\triangleq 10 \cdot \log_{10} \frac{\|\bar{\mathbf{x}} - \hat{\mathbf{x}}\|_2^2}{\|\bar{\mathbf{x}} - \mathbf{x}\|_2^2}$: signal-to-noise ratio (SNR), where $\bar{\mathbf{x}}$ is the original image, $\hat{\mathbf{x}}$ is the mean intensity value of $\bar{\mathbf{x}}$ and $\mathbf{x}$ denotes the restored image. SNR($\mathbf{x}$) measures the quality of restoration.

In all test problems, we set the regularization parameter $\mu = 5 \times 10^4$ by trial and errors such that the restored images had reasonable SNR and relative errors. The parameter ξ was set to $\xi = 100$ such that the obtained parameters satisfied $\sum_{j=1}^p w_j = 1$. The parameter β was chosen according to the assumptions in lemma 2. In order to guarantee the sequence $\{\epsilon_{k+1}\}_{k=0}^{\infty} \subseteq [0, +\infty)$ satisfies $\sum_{k=0}^{\infty} \epsilon_{k+1} < \infty$, we choose $\epsilon_{k+1} = \frac{1}{a(k+1)^a}$, where a is a constant. We used one uniform stopping criterion, that is

$$\frac{|\hat{\Phi}(\mathbf{x}^{k+1}, \mathbf{w}^{k+1}, \mathbf{y}^{k+1}, \boldsymbol{\lambda}^{k+1}) - \hat{\Phi}(\mathbf{x}^k, \mathbf{w}^k, \mathbf{y}^k, \boldsymbol{\lambda}^k)|}{|\hat{\Phi}(\mathbf{x}^k, \mathbf{w}^k, \mathbf{y}^k, \boldsymbol{\lambda}^k)|} < \epsilon,$$

where $\hat{\Phi}(\mathbf{x}^k, \mathbf{w}^k, \mathbf{y}^k, \boldsymbol{\lambda}^k)$ is the objective function value in the k th iteration and $\epsilon > 0$ is a given tolerance. In this paper, we set $\epsilon = 10^{-2}$. The maximum number of ADMM iterations was set to 50. We also report the numerical results obtained by running the ADMM-BCD method. We compare the relative error of the image $\mathbf{x}$ and the relative error of the parameter $\mathbf{w}$ estimated by both ADMM-LAP and ADMM-BCD as well as their computational time and SNR of the restored images.

Example 4.1. AO flood illumination retinal imaging is a popular technique which has been in use for more than a decade [2, 4]. AO retinal imaging technique gives us the opportunity to observe and study retinal structures at the cellular level, which is impossible to see directly in the living eye. Due to the poor contrast of the raw AO retinal images, interpretation of such images requires the myopic deconvolution. For this example, we consider the case $p = 2$. The global PSF used in the simulation is the sum of two PSFs with the first one being focused and the second one being defocused. Let A_1, A_2 denote the relative blurring matrices. The test problem is a simulation generated from a 256×256 pixel portion of an actual AO retinal image.

We use the regularization toolbox IR Tools [7] to build up the test problem. *PRblurgauss* and *PRblurdefocus* are functions used to simulate spatially invariant Gaussian blur and spatially

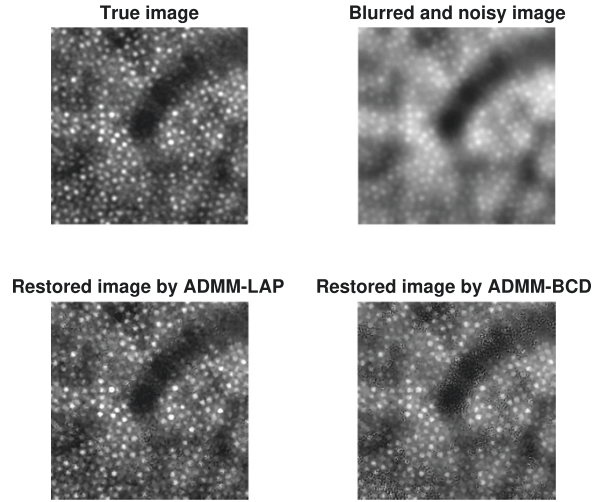


Figure 1. Restored images for the retinal image in example 4.1. The *BlurLevel* of *PRblurdefocus* is set to be ‘medium’. The images from the top left to the bottom right show the true image, the blurred and noisy image, the restored image by ADMM-LAP and the restored image by ADMM-BCD.

invariant out-of-focus blurs, respectively. Specifically, *PRblurgauss* constructs a 256×256 PSF array P with entries

$$p_{ij} = \frac{1}{2\pi\sigma} \exp\left(-\frac{1}{2} \frac{(i-k)^2}{\sigma^2} + \frac{(j-\ell)^2}{\sigma^2}\right)$$

for a given value of σ , and *PRblurdefocus* constructs a 256×256 PSF array P with entries

$$p_{ij} = \begin{cases} 1/(\pi r^2), & \text{if } (i-k)^2 + (j-\ell)^2 \leq r^2, \\ 0, & \text{elsewhere} \end{cases}$$

for a given value of r .

In this example, three cases are considered. We fix one of the PSFs to be built up by *PRblurgauss* with ‘mild’ *BlurLevel* ($\sigma = 2$) and set the other as a combined PSF built up by *PRblurgauss* with ‘mild’ *BlurLevel* and *PRblurdefocus* with three different *BlurLevels* (i.e., ‘mild’ ($r = 7$), ‘medium’ ($r = 15$), ‘severe’ ($r = 31$)). In addition, function *PRnoise* is used to add Gaussian noise with *NoiseLevel* = 0.01. We set the true parameter $w^* = [0.3; 0.7]$, choose the random image with the same size of the true image as the initial guess x^0 and the initial guess $w^0 = [0.5; 0.5]$.

We then test ADMM-LAP and ADMM-BCD for the two-weight case of (2) on this image. For the case with the *BlurLevel* of *PRblurdefocus* being set to be ‘medium’, the restored images by both methods are shown in figure 1. We can clearly see that the restored image obtained by ADMM-LAP is much sharper, hence has a much better contrast than the one obtained by ADMM-BCD. Moreover, the restored image is much closer to the true image.

The plots of the relative errors of the restored image x and the estimated parameters w against iteration can be seen in figure 2. It is easy to see that ADMM-LAP can reach lower

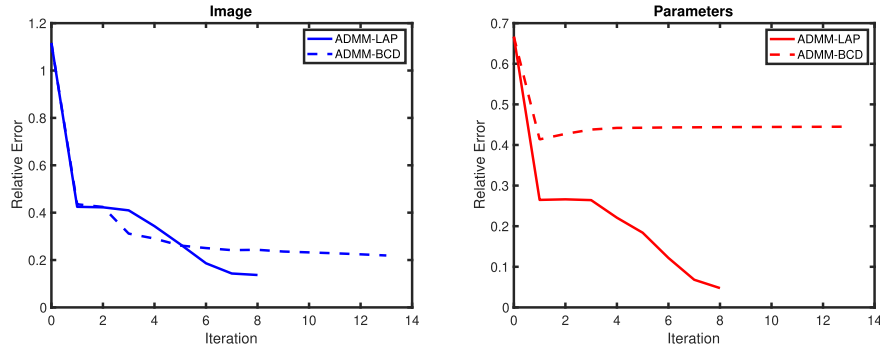


Figure 2. Relative errors of the restored image $\mathbf{x}$ (left) and the parameters $\mathbf{w}$ (right) for the myopic deconvolution problem in example 4.1 with ‘medium’ BlurLevel of *PRblurdefocus*.

Table 1. Numerical results of ADMM-LAP and ADMM-BCD for the AO retinal image in example 4.1. The columns from left to right give the BlurLevel, the method name, the stopping iteration, the relative error of the restored image, the relative error of the solution parameters, the computational time in seconds and SNR of the restored images.

BlurLevel	Method	#iter	Rel. Err. $\mathbf{x}$	Rel. Err. $\mathbf{w}$	Time/s	SNR($\mathbf{x}$)
‘Mild’	ADMM-LAP	6	1.48×10^{-1}	2.63×10^{-2}	50.23	9.57
	ADMM-BCD	11	1.88×10^{-1}	1.11×10^{-1}	154.53	7.48
‘Medium’	ADMM-LAP	8	1.37×10^{-1}	4.76×10^{-2}	56.42	10.26
	ADMM-BCD	13	2.19×10^{-1}	4.45×10^{-1}	145.09	6.16
‘Severe’	ADMM-LAP	12	1.17×10^{-1}	4.39×10^{-2}	69.10	11.60
	ADMM-BCD	12	3.04×10^{-1}	9.13×10^{-1}	122.22	3.32

relative errors than ADMM-BCD for both the restored image and the obtained parameters in this test.

In table 1, we report the relative error of the restored image $\mathbf{x}$, the relative error of the parameters $\mathbf{w}$, the computational time and SNR of the restored images for both methods for all three BlurLevels. As can be seen from the fourth and fifth columns of table 1, ADMM-LAP can obtain more accurate restored images and more accurate parameters than ADMM-BCD on all three cases. The sixth column shows that ADMM-LAP is faster than ADMM-BCD, this is mainly because BCD takes a larger number of FFTs to compute the current residual value $\mathbf{r}$ and the matrix-vector multiplication $\mathbf{J}_w^T \mathbf{r}$, and the line search is applied twice to solve the search directions for $\mathbf{x}$ and $\mathbf{w}$. Moreover, we can see from the seventh column that the quality of restored images obtained by ADMM-LAP is much better than those obtained by ADMM-BCD. We point out that it is very difficult to analyze the difference of the relative errors of $\mathbf{x}$ and $\mathbf{w}$ between the test problems, especially because of the nonlinear relationship between $\mathbf{x}$ and $\mathbf{w}$. We note that even for a linear problem, the quality of the error in $\mathbf{x}$ will depend not only on how ill-conditioned the matrix is, but on the distribution of singular values (e.g., is there a well-defined gap between large and small singular values, or do they decay smoothly). Our main purpose for this table of results is to compare the two methods (ADMM-LAP and

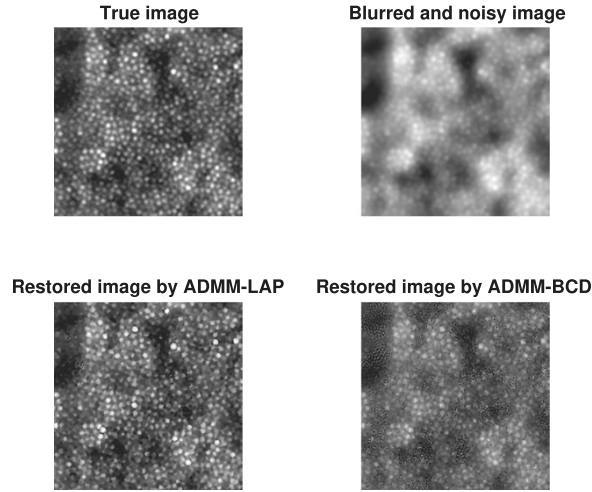


Figure 3. Restored images for the retinal image in example 4.2. The *BlurLevel* of *PRblurdefocus* is set to be ‘medium’. The images from the top left to the bottom right show the true image, the blurred and noisy image, the restored image by ADMM-LAP and the restored image by ADMM-BCD.

ADMM-BCD) for each test problem, and to see that ADMM-LAP consistently outperforms ADMM-BCD even if the test problems change.

Example 4.2. For this example, we consider the test problem from another simulation using a different 256×256 portion of an actual AO retinal image with two blurring matrices. The size of the image is 256×256 . We use *PRblurgauss* with ‘mild’ *BlurLevel*, *PRblurdefocus* with three different *BlurLevels* (i.e., ‘mild’, ‘medium’, ‘severe’) in IR Tools to build up a combined PSF, the other PSF is built up by *PRblurgauss* with ‘mild’ *BlurLevel* only. In addition, function *PRnoise* is used to add Gaussian noise with *NoiseLevel* = 0.01. We set the true parameter $w^* = [0.3; 0.7]$, choose the random image with the same size of the true image as the initial guess x^0 and set the initial guess $w^0 = [w_1^0; w_2^0] \in \mathbb{R}^2$, where w_1^0, w_2^0 are random constants from $[0, 1]$ satisfying $w_1^0 + w_2^0 = 1$.

For the case with the *BlurLevel* of *PRblurdefocus* being set to be ‘medium’, the restored images by both ADMM-LAP and ADMM-BCD method are shown in figure 3. Like example 4.1, it is easy to see that the restored image obtained by ADMM-LAP is much sharper, and is much closer to the true image than the one obtained by ADMM-BCD.

We plot the relative errors of the restored image x and the estimated parameters w against iteration in figure 4. We can also see that ADMM-LAP reaches lower relative errors of both the restored image and the obtained parameters, hence it tends to recover more accurate restored images and parameters.

Table 2 shows the relative error of the restored image x , the relative error of the parameters w , the computational time and SNR of the restored images associated with three cases for both methods. Similar to example 4.1, ADMM-LAP still outperforms ADMM-BCD in terms of the quality of the restored images and the accuracy of the obtained parameters by large margin. It can be clearly seen from the sixth column of table 2 that ADMM-LAP is faster than ADMM-BCD. This is mainly because ADMM-BCD takes a larger number of FFTs and the line search

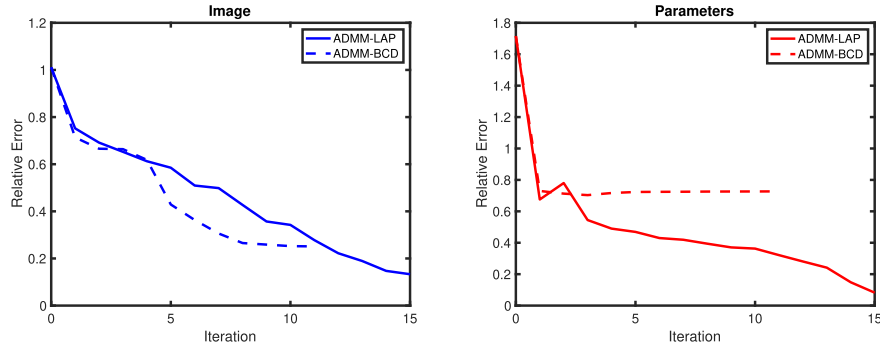


Figure 4. Relative errors for both the restored image x (left) and the parameters w (right) for the myopic deconvolution problem in example 4.2 with ‘medium’ BlurLevel of *PRblurdefocus*.

Table 2. Numerical results of ADMM-LAP and ADMM-BCD for the retinal image in example 4.2. The columns from left to right give the BlurLevel, the method name, the stopping iteration, the relative error of the restored image, the relative error of the solution parameters, the computational time in seconds and SNR of the restored images.

BlurLevel	Method	#iter	Rel. Err. x	Rel. Err. w	Time/s	SNR(x)
‘Mild’	ADMM-LAP	11	1.90×10^{-1}	1.01×10^{-1}	64.76	6.52
	ADMM-BCD	14	2.19×10^{-1}	2.74×10^{-1}	152.13	5.31
‘Medium’	ADMM-LAP	15	1.33×10^{-1}	8.23×10^{-2}	88.12	9.63
	ADMM-BCD	11	2.51×10^{-1}	7.27×10^{-1}	119.20	4.12
‘Severe’	ADMM-LAP	5	1.62×10^{-1}	7.77×10^{-1}	40.99	7.97
	ADMM-BCD	5	2.83×10^{-1}	3.76×10^{-1}	76.84	3.08

is applied twice to compute the search directions, which is its computational bottleneck. The seventh column clearly shows that ADMM-LAP obtains restored images with a much better quality. Above all, we can see that ADMM-LAP is much more efficient and accurate for this test example.

Example 4.3. In this example, we consider the case where three blurring matrices are provided, i.e., $p = 3$. The test problem is generated from the same AO retinal image used in example 4.1. The size of the image is 256×256 . The first PSF is built up by *PRblurgauss* with ‘mild’ BlurLevel. The second PSF is a combined PSF built up by *PRblurgauss* with ‘mild’ BlurLevel and *PRblurdefocus* with ‘medium’ BlurLevel. The third PSF is a combined PSF built up by *PRblurgauss* with ‘mild’ BlurLevel and *PRblurdefocus* with two different BlurLevels (i.e., ‘mild’ and ‘severe’). In addition, function *PRnoise* is used to add Gaussian noise with noiselevel = 0.01. The true parameter vector $w^* = [w_1^*; w_2^*; w_3^*] \in \mathbb{R}^3$ is set to be a random vector that satisfies $w_i^* \geq 0$ and $\sum_{i=1}^3 w_i^* = 1$. We choose the random image with the same size of the true image as the initial guess x^0 and set the initial guess $w^0 = [\frac{1}{3}; \frac{1}{3}; \frac{1}{3}]$.

The restored images by both ADMM-LAP and ADMM-BCD method are shown in figure 5. Like examples 4.1 and 4.2, we can also see that ADMM-LAP reaches a much sharper image,

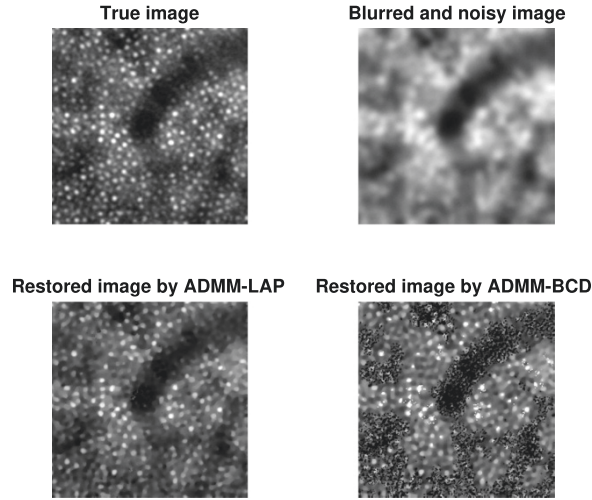


Figure 5. Restored images for the retinal image in example 4.3. The *BlurLevel* of *PRblurdefocus* in the third PSF is set to be ‘mild’. The images from the top left to the bottom right show the true image, the blurred and noisy image, the restored image by ADMM-LAP and the restored image by ADMM-BCD.

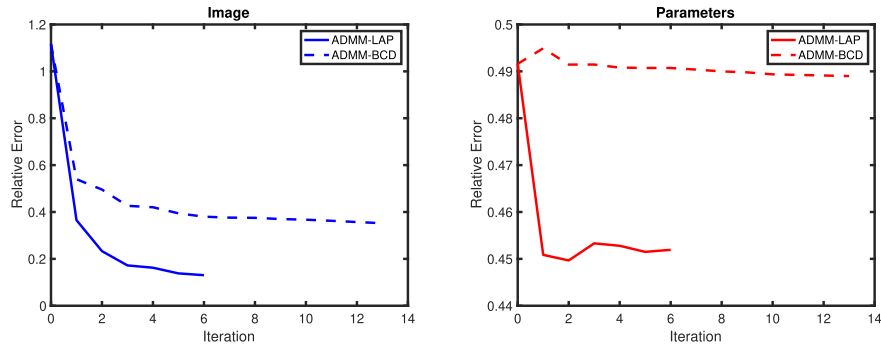


Figure 6. Relative errors of the restored image $\mathbf{x}$ (left) and the parameters $\mathbf{w}$ (right) for the myopic deconvolution problem in example 4.3 with ‘mild’ *BlurLevel* of *PRblurdefocus* in the third PSF.

it tends to recover an image that is much closer to the true image than the one obtained by ADMM-BCD.

We plot the relative errors of the restored image $\mathbf{x}$ and the estimated parameters $\mathbf{w}$ against iteration in figure 6. We can also see that ADMM-LAP achieve lower relative errors of both the restored image and the obtained parameters.

In table 3, we present the relative error of the restored image $\mathbf{x}$, the relative error of the parameters $\mathbf{w}$, the computational time and SNR of the restored images associated with two cases for both methods. We can also see from the sixth column that ADMM-LAP has evident

Table 3. Numerical results of ADMM-LAP and ADMM-BCD for the retinal image in example 4.3. The columns from left to right give the BlurLevel, the method name, the stopping iteration, the relative error of the restored image, the relative error of the solution parameters, the computational time in seconds and SNR of the restored images.

BlurLevel	Method	#iter	Rel. Err. x	Rel. Err. w	Time/s	SNR(x)
‘Mild’	ADMM-LAP	6	1.30×10^{-1}	4.52×10^{-1}	134.03	10.68
	ADMM-BCD	13	3.52×10^{-1}	4.89×10^{-1}	267.21	2.05
‘Severe’	ADMM-LAP	11	1.92×10^{-1}	1.79×10^{-1}	156.05	7.33
	ADMM-BCD	10	2.18×10^{-1}	3.27×10^{-1}	186.31	6.19

advantage over ADMM-BCD in the computational time. Moreover, the seventh column shows that the quality of restored images obtained by ADMM-LAP is much better than those obtained by ADMM-BCD.

5. Conclusion

In this paper, we propose a new efficient ADMM-LAP method for solving large scale ill-posed inverse problems, and more specifically myopic deconvolution problems with TV regularization arising from the AO retinal image restoration. Specifically, ADMM is applied to tackle the nondifferentiable and nonlinear TV regularization term first, then LAP is applied to tackle tightly coupled (x, w) -subproblems appearing within each ADMM iteration. The convergence results of ADMM-LAP are presented. Moreover, the efficiency of the proposed algorithm is demonstrated with a theoretical computational complexity analysis. In our numerical experiments we show that the proposed ADMM-LAP method is superior to ADMM-BCD method in terms of both the accuracy and the efficiency. In our future work, we plan to exploit appropriate preconditioning techniques to further reduce the iteration number and the iteration time.

Acknowledgments

JN acknowledges support from the US National Science Foundation under grant DMS-1819042 and the National Institutes of Health under grant 1R13EB028700-01. The authors also thank Dr John M Nickerson, Prof. of Ophthalmology at Emory University for helpful discussions and for providing data used for numerical experiments in this paper.

ORCID iDs

James G Nagy  <https://orcid.org/0000-0002-0451-3853>

Yuanzhe Xi  <https://orcid.org/0000-0003-0361-0931>

References

- [1] Beck A and Teboulle M 2009 A fast iterative shrinkage-thresholding algorithm for linear inverse problems *SIAM J. Imag. Sci.* **2** 183–202
- [2] Blanco L and Mugnier L M 2011 Marginal blind deconvolution of adaptive optics retinal images *Opt. Express* **19** 23227–39
- [3] Blanco L, Mugnier L M and Glanc M 2011 Myopic deconvolution of adaptive optics retina images *Spie Bios.* (International Society for Optics and Photonics)

- [4] Blanco L, Mugnier L M, Montmerle A M and Paques M 2014 Registration and restoration of adaptive-optics corrected retinal images 2014 *Int. Workshop on Computational Intelligence for Multimedia Understanding (IWCIM) IEEE* 1–5
- [5] Boyd S, Parikh N, Chu E, Peleato B and Eckstein J 2010 Distributed optimization and statistical learning via the alternating direction method of multipliers *Found. Trends Mach. Learn.* **3** 1–122
- [6] Chen Z, Nagy J G, Xi Y and Yu B 2019 Structured FISTA for image restoration *Numer. Lin. Algebra Appl.* **27** e2278
- [7] Gazzola S, Hansen P C and Nagy J G 2019 IR Tools: a MATLAB package of iterative regularization methods and large-scale test problems *Numer. Algorithms* **81** 773–811
- [8] Haber E 2014 *Computational Methods in Geophysical Electromagnetics* (Philadelphia, PA: SIAM)
- [9] Hansen P, Nagy J G and O’Leary D P 2006 *Deblurring Images* (Philadelphia, PA: SIAM)
- [10] Herring J L, Nagy J G and Ruthotto L 2018 LAP: a linearize and project method for solving inverse problems with coupled variables *Sampl. Theory Signal Image Process.* **17** 127–51
- [11] Hu Y, Nagy J G, Zhang J and Andersen M S 2019 Nonlinear optimization for mixed attenuation polyenergetic image reconstruction *Inverse Problems* **35** 064004
- [12] Kelley C T 1999 *Iterative Methods for Optimization* (Philadelphia, PA: SIAM)
- [13] Lai R and Osher S 2014 A splitting method for orthogonality constrained problems *J. Sci. Comput.* **58** 431–49
- [14] Mei J-J, Dong Y, Huang T-Z and Yin W 2018 Cauchy noise removal by nonconvex ADMM with convergence guarantees *J. Sci. Comput.* **74** 743–66
- [15] Nocedal J and Wright S 1999 *Numerical Optimization* (Berlin: Springer)
- [16] Rudin L I, Osher S and Fatemi E 1992 Nonlinear total variation based noise removal algorithms *Phys. D* **60** 259–68
- [17] Rudin L and Osher S 1994 Total variation based image restoration with free local constraints *Proc. of the 1st IEEE Int. Conf. on Image Processing* vol 1 31–5
- [18] Rockafellar R T and Wets R J B 2009 *Variational Analysis* (Berlin: Springer)
- [19] Tao M, Yang J and He B 2009 Alternating direction algorithms for total variation deconvolution in image reconstruction *Tech. Rep. TR0918* (Department of Mathematics, Nanjing University)
- [20] Tikhonov A N and Arsenin V Y 1977 *Solution of Ill-Posed Problems* (Washington, DC: V H Winston and Sons)
- [21] Wang Y, Yang J, Yin W and Zhang Y 2008 A new alternating minimization algorithm for total variation image reconstruction *SIAM J. Imag. Sci.* **1** 248–72
- [22] Wang Y, Yin W and Zeng J 2019 Global convergence of ADMM in nonconvex nonsmooth optimization *J. Sci. Comput.* **78** 29–63
- [23] Xi Y, Xia J, Cauley S and Balakrishnan V 2014 Superfast and stable structured solvers for Toeplitz least squares via randomized sampling *SIAM J. Matrix Anal. Appl.* **35** 44–72
- [24] Xia J, Xi Y and Gu M 2012 A superfast structured solver for Toeplitz linear systems via randomized sampling *SIAM J. Matrix Anal. Appl.* **33** 837–58
- [25] Xu Y and Yin W 2013 A block coordinate descent method for regularized multiconvex optimization with applications to nonnegative tensor factorization and completion *SIAM J. Imag. Sci.* **6** 1758–89
- [26] Xu Y, Yin W, Wen Z and Zhang Y 2012 An alternating direction algorithm for matrix completion with nonnegative factors *Front. Math. China* **7** 365–84
- [27] Yang L, Pong T K and Chen X 2012 Alternating direction method of multipliers for nonconvex background/foreground extraction *SIAM J. Imag. Sci.* **10** 74–110
- [28] Zhang J and Nagy J G 2018 An alternating direction method of multipliers for the solution of matrix equations arising in inverse problems *Numer. Lin. Algebra Appl.* **25** e2123