

# PROJECTED NEWTON METHOD FOR LARGE-SCALE BAYESIAN LINEAR INVERSE PROBLEMS \*

HAIBO LI †

**Abstract.** Computing the regularized solution of Bayesian linear inverse problems as well as the corresponding regularization parameter is highly desirable in many applications. This paper proposes a novel iterative method, termed the *Projected Newton method* (PNT), that can simultaneously update the regularization parameter and solution step by step without requiring any expensive matrix inversions or decompositions. By reformulating the Tikhonov regularization as a constrained minimization problem and leveraging its Lagrangian function, a Newton-type method coupled with a Krylov subspace method is designed for the unconstrained Lagrangian function. The resulting PNT algorithm only needs solving a small-scale linear system to get a descent direction of a merit function at each iteration, thus significantly reducing computational overhead. Rigorous convergence results are proved, showing that PNT always converges to the unique regularized solution and the corresponding Lagrangian multiplier. Experimental results on both small and large-scale Bayesian inverse problems demonstrate its excellent convergence property, robustness and efficiency. Given that the most demanding computational tasks in PNT are primarily matrix-vector products, it is particularly well-suited for large-scale problems.

**Key words.** Bayesian inverse problem, Tikhonov regularization, constrained optimization, Newton method, generalized Golub-Kahan bidiagonalization, projected Newton direction

**MSC codes.** 65J22, 65J20, 65K10, 90C06

**1. Introduction.** Inverse problems arise in various scientific and engineering fields, where the aim is to recover unknown parameters or functions from noisy observed data. Applications include image reconstruction, computed tomography, medical imaging, geoscience, data assimilation and so on [6, 25, 28, 30, 49]. A linear inverse problem of the discrete form can be written as

$$(1.1) \quad \mathbf{b} = \mathbf{A}\mathbf{x} + \boldsymbol{\epsilon},$$

where  $\mathbf{x} \in \mathbb{R}^n$  is the underlying quantity to reconstruct,  $\mathbf{A} \in \mathbb{R}^{m \times n}$  is the discretized forward model matrix,  $\mathbf{b} \in \mathbb{R}^m$  is the vector of observation with noise  $\boldsymbol{\epsilon}$ . We assume that the distribution of  $\boldsymbol{\epsilon}$  is known, which follows a zero mean Gaussian distribution with positive definite covariance matrix  $\mathbf{M}$ , i.e.,  $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{M})$ . A big challenge for reconstructing a good solution is the ill-posedness of inverse problems, which means that there may be multiple solutions that fit the observation equally well, or the solution is very sensitive with respect to observation perturbation.

To overcome the ill-posedness, regularization is a commonly used technique. From a Bayesian perspective [28, 53], this corresponds to adding a prior distribution of the desired solution to constrain the set of possible solutions to improve stability and uniqueness. By treating  $\mathbf{x}$  and  $\mathbf{b}$  as random variables, the observation vector  $\mathbf{b}$  has a conditional probability density function (pdf) of the form  $p(\mathbf{b}|\mathbf{x}) \propto \exp(-\frac{1}{2}\|\mathbf{A}\mathbf{x} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2)$ . To get a regularized solution, this paper considers a Gaussian prior about the desired solution with the form  $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mu^{-1}\mathbf{N})$ , where  $\mathbf{N}$  is a positive definite covariance matrix. Then the Bayes' formula leads to

$$p(\mathbf{x}|\mathbf{b}, \lambda) \propto p(\mathbf{x}|\lambda)p(\mathbf{b}|\mathbf{x}) \propto \exp\left(-\frac{1}{2}\|\mathbf{A}\mathbf{x} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \frac{\mu}{2}\|\mathbf{x}\|_{\mathbf{N}^{-1}}^2\right),$$

---

\*Submitted to the editors DATE.

†School of Mathematics and Statistics, The University of Melbourne, Parkville, VIC 3010, Australia. (haibo.li@unimelb.edu.au).

where  $\|\mathbf{x}\|_{\mathbf{B}} := (\mathbf{x}^\top \mathbf{B} \mathbf{x})^{1/2}$  is the  $\mathbf{B}$ -norm of  $\mathbf{x}$  for a positive definite matrix  $\mathbf{B}$ . Maximize the posterior pdf  $p(\mathbf{x}|\mathbf{b}, \lambda)$  leads to the Tikhonov regularization problem

$$(1.2) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \{ \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 + \mu \|\mathbf{x}\|_{\mathbf{N}^{-1}}^2 \},$$

where the regularization term  $\mu \|\mathbf{x}\|_{\mathbf{N}^{-1}}^2$  enforces extra structure on the solution that comes from the prior distribution of  $\mathbf{x}$ .

The parameter  $\mu$  in the Gaussian prior  $\mathcal{N}(\mathbf{0}, \mu^{-1} \mathbf{N})$  is crucial for obtaining a good regularized solution, which controls the trade-off between the data-fit term and regularization term. There is tremendous effort in determining a proper value of  $\mu$ . For the standard 2-norm problem, i.e.  $\mathbf{M} = \mathbf{I}$  and  $\mathbf{N} = \mathbf{I}$ , the classical parameter-selection methods include the L-curve criterion [23], generalized cross-validation [22], unbiased predictive risk estimation [47] and discrepancy principle [40]. There are also some iterative methods based on solving a nonlinear equation of  $\mu$ ; see e.g. [2, 20, 37, 46]. However, the aforementioned methods can not be directly applied to (1.2). A common procedure needs to first transform (1.2) into standard 2-norm form

$$(1.3) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \{ \|\mathbf{L}_M(\mathbf{A}\mathbf{x} - \mathbf{b})\|_2^2 + \mu \|\mathbf{L}_N \mathbf{x}\|_2^2 \},$$

where  $\mathbf{M}^{-1} = \mathbf{L}_M^\top \mathbf{L}_M$  and  $\mathbf{N}^{-1} = \mathbf{L}_N^\top \mathbf{L}_N$  are the Cholesky factorizations, and then apply the parameter-selection methods. This procedure needs the matrix inversions of  $\mathbf{M}$  and  $\mathbf{N}$  as well as the Cholesky factorizations of  $\mathbf{M}^{-1}$  and  $\mathbf{N}^{-1}$ . For large-scale matrices, these two types of computations are almost impossible or extremely expensive.

For large-scale problems, there exist some iterative regularization methods that can avoid choosing  $\mu$  in advance. A class of commonly used iterative methods is based on Krylov subspace [35], where the original linear system is projected onto lower-dimensional subspaces to become a series of small-scale problems [19, 27, 33, 42]. For dealing with the general-form Tikhonov regularization term  $\|\mathbf{L}_N \mathbf{x}\|_2^2$ , some recent Krylov iterative methods include [26, 29, 34, 39, 45] and so on. When the Cholesky factor  $\mathbf{L}_N$  is not accessible, a key difficulty is dealing with the prior covariance  $\mathbf{N}$ , which means that the subspaces should be constructed elaborately such that the prior information of  $\mathbf{x}$  can be effectively incorporated into these subspaces [8, 32]. Such methods have been proposed in [7–9], where a statistically inspired priorconditioning technique is used to whiten the noise and the desired solution. However, these methods still require large-scale matrix inversions and Cholesky factorizations, which prohibits their applications to large-scale problems.

Recently, there are several Krylov methods for directly solving (1.1) without choosing  $\mu$  in advance and can avoid the matrix inversions and Cholesky factorizations [12, 32]. These methods use the generalized Golub-Kahan bidiagonalization (gen-GKB), which can iteratively reduce the original large-scale problem to small-scale ones and generate Krylov subspaces that effectively incorporate the prior information of  $\mathbf{x}$  encoded by  $\mathbf{N}$ . In [32], the regularization effect of the proposed method comes from early stopping the iteration, where the iteration number plays the role of the regularization parameter, while in [12], the authors proposed a hybrid regularization method that simultaneously computes the regularized parameter and solution step by step. Although these two methods are very efficient for large-scale problems, there may be some issues in certain situations. The method in [32] only computes a good regularized solution but not a good  $\mu$ . However, in some applications, we need an accurate estimate of  $\mu$  to get the posterior distribution of  $\mathbf{x}$  for sampling and uncertainty quantification [17, 52, 53]. For the hybrid method in [12], the convergence

property does not have a solid theoretical foundation, and it has been numerically found that the method sometimes does not converge to a good solution, which is a common potential flaw for hybrid methods [11, 48].

Many optimization methods have been proposed for inverse problems, particularly those stemming from image processing that leads to total variation regularization or  $\ell_p$  regularization. These methods include the Bregman iteration [21, 43, 56], iterative shrinkage thresholding [3, 15], and many others [1, 36, 54]. However, these methods either need a good parameter  $\mu$  in advance or can not well deal with  $\mathbf{M}^{-1}$  and  $\mathbf{N}^{-1}$ . In [31] the author proposed a modification of the Newton method that can iteratively compute a good  $\mu$  and regularized solution simultaneously. However, this method needs to solve a large-scale linear system at each iteration, which is very costly for large-scale problems. This method was improved in [13, 14], where the Newton method is successfully combined with a Krylov subspace method to get a so-called projected Newton method. Compared with the original method, the projected Newton method only needs to solve a small-scale linear system at each iteration, thereby very efficient for large-scale Tikhonov regularization (1.3). However, for solving (1.2), this method needs to compute  $\nabla(\frac{1}{2}\|\mathbf{x}\|_{\mathbf{N}^{-1}}^2) = \mathbf{N}^{-1}\mathbf{x}$  to construct subspaces, which is also very costly. Besides, their methods lack rigorous proof of convergence.

In this paper, we develop a new efficient iterative method for (1.2) that simultaneously updates the regularization parameter and solution step by step, and it does not require any expensive matrix inversions or Cholesky factorizations. This method follows the Newton-type approach for noise constrained Tikhonov regularization proposed in [13], where the **gen-GKB** process is integrated to compute a projected Newton direction by solving a small-scale linear system at each iteration, thereby it is also named the *projected Newton method* (PNT). The main contributions of this paper are listed as follows:

- We reformulate the regularization of the original Bayesian linear inverse problem as a noise constrained minimization problem and prove the existence, uniqueness and positivity of its Lagrangian multiplier  $\lambda$  under a very reasonable assumption. The correspondence between the constrained minimization problem and Tikhonov regularization (1.2) is connected by  $\mu = 1/\lambda$ .
- We propose a **gen-GKB** based Newton-type method to compute the regularized solution by optimizing its Lagrangian function and obtaining the corresponding Lagrangian multiplier. A series of Krylov subspaces is generated by **gen-GKB**, avoiding the need for costly matrix inversions or Cholesky factorizations. Using the subspace projection technique, we only need to solve a small-scale linear system to compute the descent direction at each iteration.
- A rigorous proof of convergence for the proposed method is provided. With a very practical initialization  $(\mathbf{x}_0, \lambda_0)$ , we prove that PNT always converges to the unique solution of the constrained minimization problem and the corresponding Lagrangian multiplier.

We use both small-scale and large-scale inverse problems to test the proposed method and compare it with other state-of-the-art methods. The experimental results demonstrate excellent convergence properties of PNT, and it is very robust and efficient for regularizing Bayesian linear inverse problems. Since the most computationally intensive operations in PNT primarily involve matrix-vector products, it is especially appropriate for large-scale problems.

This paper is organized as follows. In Section 2, we formulate the noise constrained minimization problem for regularizing (1.1) and study its properties. In Section 3, we propose the PNT method. In Section 4, we prove the convergence of PNT. Numerical

results are presented in [Section 5](#) and conclusions are provided in [Section 6](#).

**2. Noise constrained minimization for Bayesian inverse problems.** In order to get a good estimate of  $\mu$  in [\(1.2\)](#), the discrepancy principle (DP) criterion is commonly used, which depends on the variance of the noise. Based on DP, we can rewrite [\(1.2\)](#) as an equivalent form of noise constrained minimization problem.

**2.1. Noise constrained minimization.** If  $\epsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$  is a white Gaussian noise, the DP criterion states that the 2-norm discrepancy between the data and predicted output  $\|\mathbf{Ax}(\mu) - \mathbf{b}\|_2$  should be of the order of  $\|\epsilon\|_2 \approx \sqrt{m}\sigma$ , where  $\mathbf{x}(\mu)$  is the solution to [\(1.2\)](#); see [\[28, §5.6\]](#). If  $\epsilon$  is a general Gaussian noise, notice that [\(1.1\)](#) leads to  $\mathbf{L}_M \mathbf{b} = \mathbf{L}_M \mathbf{Ax} + \mathbf{L}_M \epsilon$ , and  $\mathbf{L}_M \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ , thereby this transformation whitens the noise. Since  $\bar{\epsilon} := \mathbf{L}_M \epsilon$  is a white Gaussian noise with zero mean and covariance  $\mathbf{I}$ , it follows that

$$\mathbb{E} [\|\bar{\epsilon}\|_2^2] = \mathbb{E} [\text{trace}(\bar{\epsilon}^\top \bar{\epsilon})] = \mathbb{E} [\text{trace}(\bar{\epsilon} \bar{\epsilon}^\top)] = \text{trace}(\mathbb{E} [\bar{\epsilon} \bar{\epsilon}^\top]) = \text{trace}(\mathbf{I}) = m.$$

Therefore, the DP for [\(1.1\)](#) can be written as

$$(2.1) \quad \|\mathbf{Ax}(\mu) - \mathbf{b}\|_{M^{-1}}^2 = \|\mathbf{L}_M \mathbf{Ax}(\mu) - \mathbf{L}_M \mathbf{b}\|_2^2 = \tau m,$$

where  $\tau$  is chosen to be marginally greater than 1, such as  $\tau = 1.01$ .

Using this expression of DP, we rewrite the regularization of [\(1.1\)](#) as the noise constrained minimization problem

$$(2.2) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \frac{1}{2} \|\mathbf{x}\|_{N^{-1}}^2 \quad \text{s.t.} \quad \frac{1}{2} \|\mathbf{Ax} - \mathbf{b}\|_{M^{-1}}^2 \leq \frac{\tau m}{2},$$

where its Lagrangian is

$$(2.3) \quad \mathcal{L}(\mathbf{x}, \lambda) = \frac{1}{2} \|\mathbf{x}\|_{N^{-1}}^2 + \frac{\lambda}{2} (\|\mathbf{Ax} - \mathbf{b}\|_{M^{-1}}^2 - \tau m)$$

with  $\lambda \geq 0$  the Lagrangian multiplier. To further investigate [\(2.2\)](#) and [\(2.3\)](#), we first state the following basic assumption, which is used throughout the paper.

ASSUMPTION 1. For all  $\mathbf{x} \in \{\mathbf{x} \in \mathbb{R}^n : \|\mathbf{Ax} - \mathbf{b}\|_{M^{-1}} = \min\}$ , it holds

$$(2.4) \quad \|\mathbf{Ax} - \mathbf{b}\|_{M^{-1}}^2 < \tau m < \|\mathbf{b}\|_{M^{-1}}^2.$$

The first inequality means that the naive solutions to [\(1.1\)](#) fit the observation very well, and it ensures the feasible set of [\(2.2\)](#) is nonempty. The second inequality comes from the condition  $\|\mathbf{L}_M \epsilon\|_2 < \|\mathbf{L}_M \mathbf{b}\|_2$ , meaning that the noise does not dominate the observation, which ensures the effectiveness of the regularization. Under this assumption, the following result describes the solution to [\(2.2\)](#).

**THEOREM 2.1.** *The noise constrained minimization [\(2.2\)](#) has a unique solution  $\mathbf{x}^*$  satisfying  $\|\mathbf{Ax}^* - \mathbf{b}\|_{M^{-1}}^2 = \tau m$ . Furthermore, there is a unique  $\lambda^* > 0$ , which is the Lagrangian multiplier corresponding to  $\mathbf{x}^*$  in [\(2.3\)](#).*

*Proof.* Let  $\varphi(\mathbf{x}) := \frac{1}{2}(\|\mathbf{Ax} - \mathbf{b}\|_{M^{-1}}^2 - \tau m)$ , which is a convex function. In [\(2.2\)](#) we seek solutions to  $\min \frac{1}{2} \|\mathbf{x}\|_{N^{-1}}^2$  in the feasible set  $S := \{\mathbf{x} \in \mathbb{R}^n : \varphi(\mathbf{x}) \leq 0\}$ , which is the 0-lower level set of  $\varphi(\mathbf{x})$ . Note that  $S$  is a compact and convex set and  $\frac{1}{2} \|\mathbf{x}\|_{N^{-1}}^2$  is continuous and strictly convex. Thus, there is a unique solution  $\mathbf{x}^*$  to

(2.2). Suppose  $\lambda^*$  is a Lagrangian multiplier corresponding to  $\mathbf{x}^*$ . By the Karush-Kuhn-Tucker (KKT) condition [41, §12.3], the solution  $(\mathbf{x}^*, \lambda^*)$  satisfies

$$\begin{cases} \mathbf{N}^{-1}\mathbf{x}^* + \lambda^*\nabla\varphi(\mathbf{x}^*) = \mathbf{0}, \\ \lambda^*\varphi(\mathbf{x}^*) = 0, \\ \lambda^* \geq 0. \end{cases}$$

If  $\lambda^* = 0$ , then  $\mathbf{N}^{-1}\mathbf{x}^*$ , leading to  $\mathbf{x}^* = \mathbf{0}$ . This means  $\mathbf{0} \in S$ , i.e.  $\|\mathbf{b}\|_{\mathbf{M}^{-1}}^2 \leq \tau m$ , a contradiction. Consequently, it must hold  $\lambda^* > 0$ . From the relation  $\lambda^*\varphi(\mathbf{x}^*) = 0$  we have  $\varphi(\mathbf{x}^*) = 0$ , i.e.  $\|\mathbf{Ax}^* - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \tau m$ .

For the uniqueness of  $\lambda^*$ , here we give two proofs. In the first proof, we note that  $\nabla\varphi(\mathbf{x}^*) = \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{Ax}^* - \mathbf{b}) \neq \mathbf{0}$  since  $\mathbf{x}^* \notin \{\mathbf{x} \in \mathbb{R}^n : \|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}} = \min\}$  by Assumption 1. Therefore, the linear independence constraint qualification (LICQ) holds at  $\mathbf{x}^*$ , which leads to the uniqueness of  $\lambda^*$ ; see [41, §12.3].

In the second proof, we note that for any  $\lambda \geq 0$ , there is a unique  $\mathbf{x}_\lambda$  that solves the first equality of the KKT condition:

$$(2.5) \quad \mathbf{N}^{-1}\mathbf{x} + \lambda\nabla\varphi(\mathbf{x}) = \mathbf{0} \Leftrightarrow (\mathbf{N}^{-1} + \lambda\mathbf{A}^\top \mathbf{M}^{-1}\mathbf{A})\mathbf{x} = \lambda\mathbf{A}^\top \mathbf{M}^{-1}\mathbf{b},$$

since  $\mathbf{N}^{-1} + \lambda\mathbf{A}^\top \mathbf{M}^{-1}\mathbf{A}$  is positive definite. Here we prove a stronger property: *there exist a unique  $\lambda \geq 0$  such that  $\|\mathbf{Ax}_\lambda - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \tau m$* . The existence of such a  $\lambda$  has been proved, since  $\mathbf{x}^* = \mathbf{x}_{\lambda^*}$ . For the uniqueness, define two functions

$$K(\lambda) := \frac{1}{2}\|\mathbf{x}_\lambda\|_{\mathbf{N}^{-1}}^2, \quad H(\lambda) := \frac{1}{2}(\|\mathbf{Ax}_\lambda - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \tau m).$$

Note that  $\mathcal{L}(\mathbf{x}, \lambda)$  is strictly convex for a fixed  $\lambda > 0$ , which has the unique minimizer  $\mathbf{x}_\lambda$ . Thus, for any two positive  $\lambda_1 \neq \lambda_2$ , we have  $\mathcal{L}(\mathbf{x}_{\lambda_1}, \lambda_1) < \mathcal{L}(\mathbf{x}_{\lambda_2}, \lambda_1) \Leftrightarrow K(\lambda_1) + \lambda_1 H(\lambda_1) < K(\lambda_2) + \lambda_1 H(\lambda_2)$ , since  $\mathbf{x}_{\lambda_1} \neq \mathbf{x}_{\lambda_2}$ ; see the following Lemma 2.2. Similarly, we have  $K(\lambda_2) + \lambda_2 H(\lambda_2) < K(\lambda_1) + \lambda_2 H(\lambda_1)$ . Adding the above two inequalities leads to  $(\lambda_1 - \lambda_2)(H(\lambda_1) - H(\lambda_2)) < 0$ , meaning that  $H(\lambda)$  is a strictly monotonic decreasing function. Therefore, there is a unique  $\lambda$  such that  $H(\lambda) = 0$ .  $\square$

We emphasize that Assumption 1 is essential for ensuring the validity of Theorem 2.1 and plays a key role in the regularization of (1.1). If the left inequality of Assumption 1 is violated, then either the feasible set of (2.2) is empty when  $\|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 > \tau m$  or it becomes the equivalent least squares problem

$$\min \frac{1}{2}\|\mathbf{x}\|_{\mathbf{N}^{-1}}^2 \quad \text{s.t.} \quad \|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \min$$

when  $\|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \tau m$ . The latter case means that no regularization is required, which makes the problem much easier to handle. The right inequality of Assumption 1 ensures the positivity of the Lagrangian multiplier  $\lambda^*$ , which is a necessary condition. To see it, let us assume there exist a  $\lambda^* > 0$  and follow the second proof for the uniqueness of  $\lambda^*$ . From the KKT condition it must hold that  $\varphi(\mathbf{x}_{\lambda^*}) = H(\lambda^*) = 0$ . Now  $H(0) = \frac{1}{2}(\|\mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \tau m) \leq 0$  and  $H(+\infty) = \frac{1}{2}(\|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \tau m) < 0$  for any  $\mathbf{x} \in \arg\min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_{\mathbf{M}^{-1}}$ . Since  $H(\lambda)$  is strictly monotonically decreasing, the only possible zero root of  $H(\lambda)$  is  $\lambda = 0$ , which leads to a contradiction. In this case, it implies that the noise in  $\mathbf{b}$  is too large, resulting in  $\lambda^* = 0$  and a very poor regularized solution  $\mathbf{x}^* = \mathbf{0}$ .

LEMMA 2.2. *For each  $\lambda \geq 0$ , the regularization problem*

$$(2.6) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \{ \lambda \| \mathbf{A}\mathbf{x} - \mathbf{b} \|_{\mathbf{M}^{-1}}^2 + \| \mathbf{x} \|_{\mathbf{N}^{-1}}^2 \}$$

*has the unique solution  $\mathbf{x}_\lambda$ . If  $\lambda_1 \neq \lambda_2$ , then  $\mathbf{x}_{\lambda_1} \neq \mathbf{x}_{\lambda_2}$ .*

*Proof.* Note that the normal equation of (2.6) is equivalent to (2.5). Thus,  $\mathbf{x}_\lambda$  is the unique solution to (2.6). Using the Cholesky factors of  $\mathbf{M}^{-1}$  and  $\mathbf{N}^{-1}$ , and noticing that  $\mathbf{L}_N$  is invertible, we can write the generalized singular value decomposition (GSVD) [55] of  $\{\mathbf{L}_M \mathbf{A}, \mathbf{L}_N\}$  as  $\mathbf{L}_M \mathbf{A} = \mathbf{U}_A \mathbf{\Sigma}_A \mathbf{Z}^{-1}$ ,  $\mathbf{L}_N = \mathbf{U}_N \mathbf{\Sigma}_N \mathbf{Z}^{-1}$  with

$$\mathbf{\Sigma}_A = \begin{pmatrix} \mathbf{D}_A & & \\ & \mathbf{0} & \\ & & \end{pmatrix} \begin{matrix} r \\ m-r \\ r \end{matrix}, \quad \mathbf{\Sigma}_N = \begin{pmatrix} \mathbf{D}_N & & \\ & \mathbf{I} & \\ & & \end{pmatrix} \begin{matrix} r \\ n-r \\ r \end{matrix},$$

where  $\mathbf{U}_A \in \mathbb{R}^{m \times m}$  and  $\mathbf{U}_N \in \mathbb{R}^{n \times n}$  are orthogonal,  $\mathbf{Z} = (\mathbf{z}_1, \dots, \mathbf{z}_n)$  is nonsingular,  $r = \text{rank}(\mathbf{A})$ , and  $\mathbf{D}_A = \text{diag}(\sigma_1, \dots, \sigma_r)$  with  $1 > \sigma_1 \geq \dots \geq \sigma_r > 0$  and  $\mathbf{D}_N = \text{diag}(\rho_1, \dots, \rho_r)$  with  $0 < \rho_1 \leq \dots \leq \rho_r < 1$ , such that  $\sigma_i^2 + \rho_i^2 = 1$ . Then  $\mathbf{x}_\lambda$  can be expressed as  $\mathbf{x}_\lambda = \sum_{i=1}^r \frac{\lambda \sigma_i}{\lambda \sigma_i^2 + \rho_i^2} (\mathbf{u}_{A,i}^\top \mathbf{L}_M \mathbf{b}) \mathbf{z}_i$  where  $\mathbf{u}_{A,i}$  is the  $i$ -th column of  $\mathbf{U}_A$ . Since  $\{\mathbf{z}_i\}_{i=1}^r$  are linear independent, if  $\mathbf{x}_{\lambda_1} = \mathbf{x}_{\lambda_2}$ , then it must hold

$$\frac{\lambda_1 \sigma_i}{\lambda_1 \sigma_i^2 + \rho_i^2} = \frac{\lambda_2 \sigma_i}{\lambda_2 \sigma_i^2 + \rho_i^2} \Leftrightarrow (\lambda_1 - \lambda_2) \sigma_i \rho_i^2 = 0, \quad i = 1, \dots, r.$$

Since  $\sigma_i \rho_i > 0$  for  $i = 1, \dots, r$ , we obtain  $\lambda_1 = \lambda_2$ .  $\square$

REMARK 1. *From the proof of Theorem 2.1, we find that  $\lambda$  plays the role of  $\mu^{-1}$  in (1.2), meaning that  $\mathbf{x}(\mu) = \mathbf{x}_\lambda$  if  $\lambda = \mu^{-1}$ . In fact, there is a one-to-one correspondence between (1.2) and (2.2). Note that  $\mathbf{x}^* = \mathbf{x}_{\lambda^*}$ . Comparing (2.6) with (1.2), we can use  $(\lambda^*)^{-1}$  as a good estimate of the optimal regularization parameter.*

COROLLARY 2.3. *Let  $\mathbb{R}^+ = [0, \infty)$ . Write the gradient of  $\mathcal{L}(\mathbf{x}, \lambda)$  as*

$$(2.7) \quad F(\mathbf{x}, \lambda) = \begin{pmatrix} \lambda \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A}\mathbf{x} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x} \\ \frac{1}{2} \| \mathbf{A}\mathbf{x} - \mathbf{b} \|_{\mathbf{M}^{-1}}^2 - \frac{\tau m}{2} \end{pmatrix}.$$

*Then  $F(\mathbf{x}, \lambda) = \mathbf{0}$  has a unique solution  $(\mathbf{x}^*, \lambda^*)$  in  $\mathbb{R}^n \times \mathbb{R}^+$ , which is the unique minimizer and corresponding Lagrangian multiplier of (2.2).*

**2.2. Newton method.** A modification of the Newton method was proposed in [31] to solve the nonlinear equation  $F(\mathbf{x}, \lambda) = \mathbf{0}$ , which is referred to as the Lagrange method since it is based on the Lagrangian of (2.2). In this method, the Jacobian matrix of  $F(\mathbf{x}, \lambda)$  is first computed as

$$(2.8) \quad J(\mathbf{x}, \lambda) = \begin{pmatrix} \lambda \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{A} + \mathbf{N}^{-1} & \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A}\mathbf{x} - \mathbf{b}) \\ (\mathbf{A}\mathbf{x} - \mathbf{b})^\top \mathbf{M}^{-1} \mathbf{A} & 0 \end{pmatrix}$$

at the current iterate  $(\mathbf{x}, \lambda)$ , and then it computes the Newton direction  $(\Delta \mathbf{x}^\top, \Delta \lambda)^\top$  by solving inexactly the linear system

$$(2.9) \quad J(\mathbf{x}, \lambda) \begin{pmatrix} \Delta \mathbf{x} \\ \Delta \lambda \end{pmatrix} = -F(\mathbf{x}, \lambda)$$

using the MINRES solver [44]. We remark that this method is essentially a Newton-Krylov method [5] for optimizing the nonlinear and nonconvex Lagrangian function

(2.3). It was shown that the computed  $(\Delta \mathbf{x}, \Delta \lambda)$  is a descent direction for the merit function

$$h_w(\mathbf{x}, \lambda) = \frac{1}{2} (\|\nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}, \lambda)\|_2^2 + w |\nabla_{\lambda} \mathcal{L}(\mathbf{x}, \lambda)|^2)$$

with a  $w > 0$ , which means that  $(\Delta \mathbf{x}^\top, \Delta \lambda) \nabla h_w(\mathbf{x}, \lambda) \leq 0$ . By a backtracking line search strategy to determine a step length  $\gamma > 0$ , the iterate is updated as  $(\mathbf{x}, \lambda) \leftarrow (\mathbf{x}, \lambda) + \gamma(\Delta \mathbf{x}, \Delta \lambda)$ .

An advantage of this method is that it can compute a good regularized solution and its regularization parameter simultaneously. However, for large-scale problems, we need to compute  $\mathbf{M}^{-1}$  and  $\mathbf{N}^{-1}$  to form  $F(\mathbf{x}, \lambda)$  and  $J(\mathbf{x}, \lambda)$ , which is almost impossible. Moreover, at each iteration, an  $(n+1) \times (n+1)$  linear system (2.9) needs to be solved, which is very computationally expensive even if we only compute a less accurate solution by an iterative algorithm.

In [13], the authors proposed a projected Newton method, where at each iteration, the large-scale linear system (2.9) is projected to be a small-scale linear system that can be solved cheaply. However, this method can only deal with the standard  $\ell_2 - \ell_2$  regularization, which means we can only apply this method to (1.3) by the substitution  $\bar{\mathbf{x}} = \mathbf{L}_N \mathbf{x}$ , requiring the expensive Cholesky factorization of  $\mathbf{N}^{-1}$ . A generalization of this method [14] can deal with a general-form regularization term. However, for (2.2), it needs to compute  $\nabla(\frac{1}{2}\|\mathbf{x}\|_{\mathbf{N}^{-1}}^2) = \mathbf{N}^{-1}\mathbf{x}$  to construct subspace for projecting (2.9), also very costly.

**3. Projected Newton method based on generalized Golub-Kahan bidiagonalization.** To reduce expensive computations of the Newton method for large-scale problems, we design a new projected Newton method to solve (2.2). This method uses the generalized Golub-Kahan bidiagonalization (gen-GKB) to construct Krylov subspaces to compute projected Newton directions by only solving small-scale problems, and it does not need any expensive matrix inversions or decompositions. This method is composed by the following three main steps:

- Step 1: *Construct Krylov subspaces.* We adopt gen-GKB to iteratively construct a series of low-dimensional Krylov subspaces; see Algorithm 3.1.
- Step 2: *Compute the projected Newton direction.* At each iteration, we compute the projected Newton direction by solving a small-scale problem; see (3.12).
- Step 3: *Determine the step-length to update solution.* We use the Armijo backtracking line search to determine a step-length and update the solution; see Routine 1.

In the next subsection, we present detailed derivations of the whole algorithm. All the proofs can be found in Subsection 3.2.

**3.1. Derivation of projected Newton method.** This subsection presents detailed derivations for the above three steps.

**Step 1: Construct Krylov subspaces by gen-GKB.** The gen-GKB process has been proposed for solving Bayesian linear inverse problems in [12, 32]. The basic idea is to treat  $\mathbf{A}$  as the compact linear operator

$$\mathbf{A} : (\mathbb{R}^n, \langle \cdot, \cdot \rangle_{\mathbf{N}^{-1}}) \rightarrow (\mathbb{R}^m, \langle \cdot, \cdot \rangle_{\mathbf{M}^{-1}}), \quad \mathbf{x} \mapsto \mathbf{A}\mathbf{x},$$

where  $\mathbf{x}$  and  $\mathbf{A}\mathbf{x}$  are vectors under the canonical bases of  $\mathbb{R}^n$  and  $\mathbb{R}^m$ , respectively. Here the two inner products are defined as  $\langle \mathbf{x}, \mathbf{x}' \rangle_{\mathbf{N}^{-1}} := \mathbf{x}^\top \mathbf{N}^{-1} \mathbf{x}'$  and  $\langle \mathbf{y}, \mathbf{y}' \rangle_{\mathbf{M}^{-1}} := \mathbf{y}^\top \mathbf{M}^{-1} \mathbf{y}'$ . Therefore, we can define

$$\mathbf{A}^* : (\mathbb{R}^m, \langle \cdot, \cdot \rangle_{\mathbf{M}^{-1}}) \rightarrow (\mathbb{R}^n, \langle \cdot, \cdot \rangle_{\mathbf{N}^{-1}}), \quad \mathbf{y} \mapsto \mathbf{A}^* \mathbf{y},$$

which is the adjoint operator of  $\mathbf{A}$ , by the relation  $\langle \mathbf{A}\mathbf{x}, \mathbf{y} \rangle_{\mathbf{M}^{-1}} = \langle \mathbf{x}, \mathbf{A}^* \mathbf{y} \rangle_{\mathbf{N}^{-1}}$ . Note that  $(\mathbf{A}\mathbf{x})^\top \mathbf{M}^{-1} \mathbf{y} = \mathbf{x}^\top \mathbf{N}^{-1} \mathbf{A}^* \mathbf{y}$  for any  $\mathbf{x} \in \mathbb{R}^n$  and  $\mathbf{y} \in \mathbb{R}^m$ . Thus, the matrix-form expression of  $\mathbf{A}^*$  is  $\mathbf{A}^* = \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1}$ .

Applying the standard Golub-Kahan bidiagonalization (GKB) to the compact operator  $\mathbf{A}$  with starting vector  $\mathbf{b}$  between the two Hilbert spaces  $(\mathbb{R}^n, \langle \cdot, \cdot \rangle_{\mathbf{N}^{-1}})$  and  $(\mathbb{R}^m, \langle \cdot, \cdot \rangle_{\mathbf{M}^{-1}})$ , we can obtain the gen-GKB process; see [10] for GKB for compact operators. The basic recursive relations of gen-GKB are as follows:

$$\begin{aligned} (3.1a) \quad & \beta_1 \mathbf{u}_1 = \mathbf{b}, \\ (3.1b) \quad & \alpha_i \mathbf{v}_i = \mathbf{A}^* \mathbf{u}_i - \beta_i \mathbf{v}_{i-1}, \\ (3.1c) \quad & \beta_{i+1} \mathbf{u}_{i+1} = \mathbf{A} \mathbf{v}_i - \alpha_i \mathbf{u}_i, \end{aligned}$$

where  $\alpha_i$  and  $\beta_i$  are computed such that  $\|\mathbf{u}_i\|_{\mathbf{M}^{-1}} = \|\mathbf{v}_i\|_{\mathbf{N}^{-1}} = 1$ , and  $\mathbf{v}_0 := \mathbf{0}$ . The whole iterative process is summarized in Algorithm 3.1. For more details of the derivation, please see [32].

We remark that computing with  $\mathbf{M}^{-1}$  can not be avoided, but for the most commonly encountered cases that  $\epsilon$  is a Gaussian noise with uncorrelated components,  $\mathbf{M}$  is diagonal and thereby  $\mathbf{M}^{-1}$  can be directly obtained. For applications that  $\epsilon$  is a colored Gaussian noise such that  $\mathbf{M}$  is not diagonal, computing  $\mathbf{M}$  is the most expensive operation. In these cases, the proposed PNT method based on gen-GKB may not be the optimal choice.

---

**Algorithm 3.1** Generalized Golub-Kahan bidiagonalization (gen-GKB)

---

**Input:**  $\mathbf{A} \in \mathbb{R}^{m \times n}$ ,  $\mathbf{b} \in \mathbb{R}^m$ ,  $\mathbf{M} \in \mathbb{R}^{m \times m}$ ,  $\mathbf{N} \in \mathbb{R}^{n \times n}$   
1:  $\bar{\mathbf{s}} = \mathbf{M}^{-1} \mathbf{b}$ ,  $\beta_1 = \bar{\mathbf{s}}^\top \mathbf{b}$ ,  $\mathbf{u}_1 = \mathbf{b} / \beta_1$ ,  $\bar{\mathbf{u}}_1 = \bar{\mathbf{s}} / \beta_1$   
2:  $\bar{\mathbf{r}} = \mathbf{A}^\top \bar{\mathbf{u}}_1$ ,  $\mathbf{r} = \mathbf{N} \bar{\mathbf{r}}$   
3:  $\alpha_1 = (\mathbf{r}^\top \bar{\mathbf{r}})^{1/2}$ ,  $\bar{\mathbf{v}}_1 = \bar{\mathbf{r}} / \alpha_1$ ,  $\mathbf{v}_1 = \mathbf{r} / \alpha_1$   
4: **for**  $i = 1, 2, \dots, k$  **do**  
5:    $\mathbf{s} = \mathbf{A} \mathbf{v}_i - \alpha_i \mathbf{u}_i$ ,  $\bar{\mathbf{s}} = \mathbf{M}^{-1} \mathbf{s}$   
6:    $\beta_{i+1} = (\bar{\mathbf{s}}^\top \bar{\mathbf{s}})^{1/2}$ ,  $\mathbf{u}_{i+1} = \bar{\mathbf{s}} / \beta_{i+1}$ ,  $\bar{\mathbf{u}}_{i+1} = \bar{\mathbf{s}} / \beta_{i+1}$   
7:    $\bar{\mathbf{r}} = \mathbf{A}^\top \bar{\mathbf{u}}_{i+1} - \beta_{i+1} \bar{\mathbf{v}}_i$ ,  $\mathbf{r} = \mathbf{N} \bar{\mathbf{r}}$   
8:    $\alpha_{i+1} = (\mathbf{r}^\top \bar{\mathbf{r}})^{1/2}$ ,  $\bar{\mathbf{v}}_{i+1} = \bar{\mathbf{r}} / \alpha_{i+1}$ ,  $\mathbf{v}_{i+1} = \mathbf{r} / \alpha_{i+1}$   
9: **end for**  
**Output:**  $\{\alpha_i, \beta_i\}_{i=1}^{k+1}$ ,  $\{\mathbf{u}_i, \mathbf{v}_i\}_{i=1}^{k+1}$

---

The following result gives the basic property of gen-GKB; see [32] for the proof.

**PROPOSITION 3.1.** *The group of vectors  $\{\mathbf{u}_i\}_{i=1}^k$  is an  $\mathbf{M}^{-1}$ -orthonormal basis of the Krylov subspace*

$$(3.2) \quad \mathcal{K}_k(\mathbf{A} \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1}, \mathbf{b}) = \text{span}\{(\mathbf{A} \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1})^i \mathbf{b}\}_{i=0}^{k-1},$$

and  $\{\mathbf{v}_i\}_{i=1}^k$  is an  $\mathbf{N}^{-1}$ -orthonormal basis of the Krylov subspace

$$(3.3) \quad \mathcal{K}_k(\mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{A}, \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{b}) = \text{span}\{(\mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{A})^i \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{b}\}_{i=0}^{k-1}.$$

Define  $\mathbf{U}_{k+1} = (\mathbf{u}_1, \dots, \mathbf{u}_{k+1})$  and  $\mathbf{V}_{k+1} = (\mathbf{v}_1, \dots, \mathbf{v}_{k+1})$ . Then Proposition 3.1 indicates that  $\mathbf{U}_{k+1}^\top \mathbf{M}^{-1} \mathbf{U}_{k+1} = \mathbf{I}$  and  $\mathbf{V}_{k+1}^\top \mathbf{N}^{-1} \mathbf{V}_{k+1} = \mathbf{I}$ . We remark that gen-GKB will eventually terminate in at most  $\min\{m, n\}$  steps, since the column rank of  $\mathbf{U}_k$  or  $\mathbf{V}_k$  can not exceed  $\min\{m, n\}$ . If we define the *termination step* as

$$(3.4) \quad k_t := \max\{k : \alpha_k \beta_k > 0\},$$

then  $\mathbf{V}_k$  will eventually expand to be  $\mathbf{V}_{k_t}$  with  $k_t \leq \min\{m, n\}$ .

By (3.1a)–(3.1c), we can write the  $k$ -step ( $k \leq k_t$ ) gen-GKB in the matrix-form

$$(3.5a) \quad \beta_1 \mathbf{U}_{k+1} \mathbf{e}_1 = \mathbf{b},$$

$$(3.5b) \quad \mathbf{A} \mathbf{V}_k = \mathbf{U}_{k+1} \mathbf{B}_k,$$

$$(3.5c) \quad \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{U}_{k+1} = \mathbf{V}_k \mathbf{B}_k^\top + \alpha_{k+1} \mathbf{v}_{k+1} \mathbf{e}_{k+1}^\top,$$

where  $\mathbf{e}_1$  and  $\mathbf{e}_{k+1}$  are the first and  $(k+1)$ -th columns of the identity matrix of order  $k+1$ , respectively, and

$$(3.6) \quad \mathbf{B}_k = \begin{pmatrix} \alpha_1 & & & & \\ \beta_2 & \alpha_2 & & & \\ & \beta_3 & \ddots & & \\ & & \ddots & \alpha_k & \\ & & & \beta_{k+1} & \end{pmatrix} \in \mathbb{R}^{(k+1) \times k}.$$

Note that  $\mathbf{B}_k$  has full column rank if  $k \leq k_t$ . At the  $k_t$ -th iteration, it is possible that either  $\beta_{k_t+1} = 0$  occurs first or  $\alpha_{k_t+1} = 0$  occurs first. For the former case, the relations (3.5) are replaced by

$$(3.7a) \quad \beta_1 \mathbf{U}_{k_t} \mathbf{e}_1 = \mathbf{b},$$

$$(3.7b) \quad \mathbf{A} \mathbf{V}_{k_t} = \mathbf{U}_{k_t} \underline{\mathbf{B}}_{k_t},$$

$$(3.7c) \quad \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} \mathbf{U}_{k_t} = \mathbf{V}_{k_t} \underline{\mathbf{B}}_{k_t}^\top,$$

where  $\underline{\mathbf{B}}_{k_t}$  is the first  $k \times k$  part of  $\mathbf{B}_{k_t}$  by discarding  $\beta_{k_t+1}$ .

**Step 2: Compute the projected Newton direction.** At the  $k$ -th iteration, we update  $\mathbf{x}_k \in \text{span}\{\mathbf{V}_k\}$  and  $\lambda_k$  from the previous ones. For any  $\mathbf{x} \in \text{span}\{\mathbf{V}_k\}$  of the form  $\mathbf{x} = \mathbf{V}_k \mathbf{y}$  with  $\mathbf{y} \in \mathbb{R}^k$ , define the projected gradient of  $\mathcal{L}(\mathbf{x}, \lambda)$  as

$$(3.8) \quad F^{(k)}(\mathbf{y}, \lambda) = \begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} F(\mathbf{x}, \lambda)$$

and the projected Jacobian of  $F(\mathbf{x}, \lambda)$  as

$$(3.9) \quad J^{(k)}(\mathbf{y}, \lambda) = \begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} J(\mathbf{x}, \lambda) \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix}.$$

**REMARK 2.** Since gen-GKB must terminate at the  $k_t$ -th iteration and  $\mathbf{V}_k$  eventually expands to be  $\mathbf{V}_{k_t}$ , we need to discuss the two different cases that  $k \leq k_t$  and  $k > k_t$ . For notational simplicity, in the rest part of the paper, we use  $\mathbf{V}_k$  and  $\mathbf{B}_k$  by default unless stated otherwise to denote

$$\mathbf{V}_k = \begin{cases} \mathbf{V}_k, & k \leq k_t \\ \mathbf{V}_{k_t}, & k > k_t \end{cases}, \quad \mathbf{B}_k = \begin{cases} \mathbf{B}_k, & k \leq k_t \\ \mathbf{B}_{k_t}, & k > k_t \end{cases}, \quad \mathbf{x}_k = \begin{cases} \mathbf{V}_k \mathbf{y}_k, & k \leq k_t \\ \mathbf{V}_{k_t} \mathbf{y}_k, & k > k_t \end{cases},$$

where  $\mathbf{y} \in \mathbb{R}^k$  for  $k \leq k_t$  and  $\mathbf{y} \in \mathbb{R}^{k_t}$  for  $k > k_t$ . Moreover, for the case  $\beta_{k_t+1} = 0$ , the relations (3.5) are replaced by (3.7) and  $\mathbf{B}_{k_t}$  is replaced by  $\underline{\mathbf{B}}_{k_t}$ . In the subsequent discussions, we employ the unified notations as presented in (3.5), but the readers can readily differentiate between the two cases.

Notice that  $\mathbf{y}$  is uniquely determined by  $\mathbf{x} = \mathbf{V}_k \mathbf{y}$  since  $\mathbf{V}_k$  has full-column rank. Thus,  $F^{(k)}(\mathbf{y}, \lambda)$  and  $J^{(k)}(\mathbf{y}, \lambda)$  are well-defined. The next result shows how we can obtain  $F^{(k)}(\mathbf{y}, \lambda)$  and  $J^{(k)}(\mathbf{y}, \lambda)$  from  $\mathbf{B}_k$  without any additional computations.

LEMMA 3.2. *For any  $\mathbf{x} \in \text{span}\{\mathbf{V}_k\}$  with the form  $\mathbf{x} = \mathbf{V}_k \mathbf{y}$ , the projected gradient of  $\mathcal{L}(\mathbf{x}, \lambda)$  has the expression*

$$(3.10) \quad F^{(k)}(\mathbf{y}, \lambda) = \left( \lambda \mathbf{B}_k^\top (\mathbf{B}_k \mathbf{y} - \beta_1 \mathbf{e}_1) + \mathbf{y} \right) / \left( \frac{1}{2} \|\mathbf{B}_k \mathbf{y} - \beta_1 \mathbf{e}_1\|_2^2 - \frac{\tau m}{2} \right),$$

and the projected Jacobian of  $F(\mathbf{x}, \lambda)$  has the expression

$$(3.11) \quad J^{(k)}(\mathbf{y}, \lambda) = \begin{pmatrix} \lambda \mathbf{B}_k^\top \mathbf{B}_k + \mathbf{I} & \mathbf{B}_k^\top (\mathbf{B}_k \mathbf{y} - \beta_1 \mathbf{e}_1) \\ (\mathbf{B}_k \mathbf{y} - \beta_1 \mathbf{e}_1)^\top \mathbf{B}_k & 0 \end{pmatrix}.$$

Now we can compute the projected Newton direction for updating the solution. Starting from an initial solution  $(\mathbf{x}_0, \lambda_0)$ , consider the following two cases.

**Case 1: Update  $(\mathbf{x}_k, \lambda_k)$  from  $(\mathbf{x}_{k-1}, \lambda_{k-1})$  for  $k \leq k_t$ .** Suppose at the  $(k-1)$ -th iteration, we have  $\mathbf{x}_{k-1} = \mathbf{V}_{k-1} \mathbf{y}_{k-1}$ , where  $\mathbf{x}_0 := \mathbf{0}$  and  $\mathbf{y}_0 := ()$  is an empty vector. Let  $\bar{\mathbf{y}}_{k-1} = (\mathbf{y}_{k-1}^\top, 0)^\top \in \mathbb{R}^k$ . If  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  is nonsingular, we compute the Newton direction for the projected function  $F^{(k)}(\mathbf{y}, \lambda)$  at  $(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$ :

$$(3.12a) \quad \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix} = -J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})^{-1} F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}).$$

Then we update  $(\bar{\mathbf{y}}_k, \lambda_k)$  by

$$(3.12b) \quad \mathbf{y}_k = \bar{\mathbf{y}}_{k-1} + \gamma_k \Delta \mathbf{y}_k, \quad \lambda_k = \lambda_{k-1} + \gamma_k \Delta \lambda_k$$

with a suitably chosen step-length  $\gamma_k > 0$ , and let  $\mathbf{x}_k = \mathbf{V}_k \mathbf{y}_k$ .

**Case 2: Update  $(\mathbf{x}_k, \lambda_k)$  from  $(\mathbf{x}_{k-1}, \lambda_{k-1})$  for  $k > k_t$ .** At each iteration, we seek a solution of the form  $\mathbf{x}_k = \mathbf{V}_{k_t} \mathbf{y}_k$  with  $\mathbf{y}_k \in \mathbb{R}^{k_t}$ . We compute the Newton direction

$$(3.12c) \quad \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix} = -J^{(k_t)}(\mathbf{y}_{k-1}, \lambda_{k-1})^{-1} F^{(k_t)}(\mathbf{y}_k, \lambda_{k-1}),$$

and then compute

$$(3.12d) \quad \mathbf{y}_k = \mathbf{y}_{k-1} + \gamma_k \Delta \mathbf{y}_k, \quad \lambda_k = \lambda_{k-1} + \gamma_k \Delta \lambda_k$$

to get  $\mathbf{x}_k = \mathbf{V}_{k_t} \mathbf{y}_k$ .

For both the two cases, we call  $(\Delta \mathbf{y}_k, \Delta \lambda_k)$  the *projected Newton direction*, since it is the Newton direction of a projected problem. The corresponding update formula for  $\mathbf{x}_k$  is

$$\mathbf{x}_k = \mathbf{x}_{k-1} + \gamma_k \Delta \mathbf{x}_k, \quad \Delta \mathbf{x}_k := \mathbf{V}_k \Delta \mathbf{y}_k,$$

which is easy to be verified. For notational simplicity, in the subsequent part we always use the unified notation

$$(3.13) \quad \bar{\mathbf{y}}_{k-1} = \begin{cases} (\mathbf{y}_{k-1}^\top, 0)^\top, & k \leq k_t \\ \mathbf{y}_{k-1}, & k > k_t \end{cases}$$

for  $\bar{\mathbf{y}}_{k-1}$ . Following the notations stated in Remark 2 and (3.13), we can use (3.12a) and (3.12b) to describe the update procedure for both the two cases.

It is vital to make sure that the projected Jacobian matrix  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  is always nonsingular. This desired property is given in the following result. The proof appears as a part of the proof of Lemma 4.6.

PROPOSITION 3.3. *If we choose  $\mathbf{x}_0 = \mathbf{0}$  and  $\bar{\mathbf{y}}_0 = 0$ , then at each iteration  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  is nonsingular as long as  $\lambda_{k-1} \geq 0$ .*

In order to investigate the convergence behavior of the method, define the following merit function:

$$(3.14) \quad h(\mathbf{x}, \lambda) = \frac{1}{2} [\|\lambda \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x} - \mathbf{b}) + \mathbf{N}^{-1}\mathbf{x}\|_N^2 + (\frac{1}{2}\|\mathbf{A}\mathbf{x} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \frac{\tau m}{2})^2].$$

Notice from Corollary 2.3 that  $(\mathbf{x}^*, \lambda^*)$  is the unique minimizer of  $h(\mathbf{x}, \lambda)$  and that  $h(\mathbf{x}^*, \lambda^*) = 0$ . The following result shows that  $(\Delta \mathbf{x}_k^\top, \Delta \lambda_k)^\top$  is indeed a descent direction for  $h(\mathbf{x}, \lambda)$ .

THEOREM 3.4. *Let  $\Delta \mathbf{x}_k = \mathbf{V}_k \Delta \mathbf{y}_k$ . Then it holds*

$$(3.15) \quad \nabla h(\mathbf{x}_{k-1}, \lambda_{k-1})^\top \begin{pmatrix} \Delta \mathbf{x}_k \\ \Delta \lambda_k \end{pmatrix} = -2h(\mathbf{x}_{k-1}, \lambda_{k-1}) \leq 0.$$

Theorem 3.4 is a desired property for a gradient descent type algorithm. At the  $(k-1)$ -th iteration, if  $h(\mathbf{x}_{k-1}, \lambda_{k-1}) = 0$ , then we have  $(\mathbf{x}_{k-1}, \lambda_{k-1}) = (\mathbf{x}^*, \lambda^*)$ , meaning we have obtained the unique solution to (2.2). Otherwise,  $(\Delta \mathbf{x}_k^\top, \Delta \lambda_k)^\top$  is a descent direction of  $h(\mathbf{x}, \lambda)$  at  $(\mathbf{x}_{k-1}, \lambda_{k-1})$ , thereby we can continue updating the solution by a backtracking line search strategy.

**Step 3: Determine step-length by backtracking line search.** For the case that  $h(\mathbf{x}_{k-1}, \lambda_{k-1}) \neq 0$ , we need to determine a step-length  $\gamma_k$  such that  $h(\mathbf{x}_k, \lambda_k)$  decreases strictly. To this end, we use the backtracking line search procedure to ensure that the Armijo condition [41, §3.1] is satisfied:

$$(3.16) \quad h(\mathbf{x}_k, \lambda_k) \leq h(\mathbf{x}_{k-1}, \lambda_{k-1}) + c\gamma_k(\Delta \mathbf{x}_{k-1}^\top, \Delta \lambda_k)^\top \nabla h(\mathbf{x}_{k-1}, \lambda_{k-1}),$$

where  $(\mathbf{x}_k, \lambda_k) = (\mathbf{x}_{k-1}, \lambda_{k-1}) + \gamma_k(\Delta \mathbf{x}_k, \Delta \lambda_k)$ , and  $c \in (0, 1)$  is a fixed constant. At each iteration, we can quickly compute  $h(\mathbf{x}_k, \lambda_k)$  based on the following result.

LEMMA 3.5. *Let*

$$\bar{F}^{(k)}(\mathbf{y}, \lambda) = \left( \lambda \bar{\mathbf{B}}_k^\top (\bar{\mathbf{B}}_k \mathbf{y} - \beta_1 \mathbf{e}_1) + \mathbf{y} \right), \quad \bar{\mathbf{B}}_k = \begin{pmatrix} \alpha_1 & & & \\ \beta_2 & \alpha_2 & & \\ & \ddots & \ddots & \\ & & \beta_{k+1} & \alpha_{k+1} \end{pmatrix}.$$

Then we have

$$(3.17) \quad h(\mathbf{x}_{k-1}, \lambda_{k-1}) = \frac{1}{2} \|F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})\|_2^2, \quad h(\mathbf{x}_k, \lambda_k) = \frac{1}{2} \|\bar{F}^{(k)}(\bar{\mathbf{y}}_k, \lambda_k)\|_2^2.$$

We remark that in the above expression we have  $\bar{\mathbf{B}}_k = \mathbf{B}_{k_t}$  for  $k \geq k_t$ , and specifically, we have  $\bar{\mathbf{B}}_k = \mathbf{B}_{k_t}^\top$  if  $\beta_{k_t+1} = 0$ . The following theorem shows the existence of a suitable step-length; see e.g. [4, pp. 121, Theorem 2.1] for details.

THEOREM 3.6. *For any continuously differentiable function  $f(\mathbf{s}) : \mathbb{R}^l \rightarrow \mathbb{R}$ , suppose  $\nabla f$  is Lipschitz continuous with constant  $\zeta(\mathbf{s})$  at  $\mathbf{s}$ . If  $\mathbf{p}$  is a descent direction at  $\mathbf{s}$ , i.e.  $\nabla f(\mathbf{s})^\top \mathbf{p} < 0$ , then for a fixed  $c \in (0, 1)$  the Armijo condition  $f(\mathbf{s} + \gamma \mathbf{p}) \leq f(\mathbf{s}) + c\gamma \nabla f(\mathbf{s})^\top \mathbf{p}$  is satisfied for all  $\gamma \in [0, \gamma_{\max}]$  with  $\gamma_{\max} = \frac{2(c-1)\nabla f(\mathbf{s})^\top \mathbf{p}}{\zeta(\mathbf{s})\|\mathbf{p}\|^2}$ .*

With the aid of Lemma 3.5, a suitable step-length  $\gamma_k$  can be determined using the following backtracking line search strategy.

ROUTINE 1. *Armijo backtracking line search:*

1. Given  $\gamma_{\text{init}} > 0$ , let  $\gamma^{(0)} = \gamma_{\text{init}}$  and  $l = 0$ .
2. Until  $\frac{1}{2} \|\bar{F}^{(k)}(\bar{\mathbf{y}}_k, \lambda_k)\|_2^2 \leq (\frac{1}{2} - c\gamma^{(l)}) \|F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})\|_2^2$ ,
  - (i) set  $\gamma^{(l+1)} = \eta\gamma^{(l)}$ , where  $\eta \in (0, 1)$  is a fixed constant;
  - (ii)  $l \leftarrow l + 1$ .
3. Set  $\gamma_k = \gamma^{(l)}$ .

We set  $c = 10^{-4}$ ,  $\gamma_{\text{init}} = 1.0$  and  $\eta = 0.9$  by default. Note that at each iteration we need to ensure  $\lambda_k > 0$ . Suppose at the  $(k-1)$ -th iteration we already have  $\lambda_{k-1} > 0$ . Then at the  $k$ -th iteration, if  $\Delta\lambda_k < 0$ , we only need to enforce  $\gamma_{\text{init}} < -\lambda_{k-1}/\Delta\lambda_k$ .

Overall, the whole procedure of PNT is presented in [Algorithm 3.2](#). In the PNT algorithm, at each  $k$ -th iteration, computing the projected Newton direction requires solving only the  $(k+1)$ -order linear system [\(3.12a\)](#), which can be done very quickly when  $k \ll n$ . Starting from the termination step  $k_t$ , at each subsequent iteration, a  $(k_t+1)$ -order linear system [\(3.12c\)](#) needs to be solved. We numerically find that the algorithm almost always obtains a satisfied solution before **gen-GKB** terminates. The PNT method is a natural generalization of the projected Newton method in [\[13\]](#). Specifically, when  $\mathbf{M} = \mathbf{I}$  and  $\mathbf{N} = \mathbf{I}$ , it can be confirmed that both methods are identical.

We remark that for very large-scale problems, it may take too many iterations for PNT to converge. In this case, we can update the solution starting from the  $k_0$ -th step of **gen-GKB** to save some computation for solving [\(3.12a\)](#). This means that we first run  $k_0 - 1$  steps **gen-GKB** to construct a  $(k_0 - 1)$ -dimensional subspace and then start to update the solution from the  $k_0$ -th iteration. From the derivation of PNT, it can be easily verified that if we set  $\bar{\mathbf{y}}_{k_0-1} = \mathbf{0} \in \mathbb{R}^{k_0}$ , then  $(\Delta\mathbf{x}_k^\top, \Delta\lambda_k)^\top$  is a descent direction of  $h(\mathbf{x}, \lambda)$  at each iteration  $k \geq k_0$ . We refer to this modified PNT method as PNT-md.

---

**Algorithm 3.2** Projected Newton method (PNT) for [\(2.2\)](#) and [\(2.3\)](#)

---

**Input:**  $\mathbf{A} \in \mathbb{R}^{m \times n}$ ,  $\mathbf{b} \in \mathbb{R}^m$ ,  $\mathbf{M} \in \mathbb{R}^{m \times m}$ ,  $\mathbf{N} \in \mathbb{R}^{n \times n}$ ,  $\tau \gtrsim 1$

- 1: Initialization:  $\lambda_0 > 0$ ,  $\bar{\mathbf{y}}_0 = \mathbf{0}$ ;  $c = 10^{-4}$ ,  $\eta = 0.9$ ;  $\text{tol} > 0$
- 2: Compute  $\beta_1$ ,  $\alpha_1$ ,  $\mathbf{u}_1$ ,  $\mathbf{v}_1$  by [Algorithm 3.1](#)
- 3: **for**  $k = 1, 2, \dots$  **do**
- 4:   Compute  $\beta_{k+1}$ ,  $\alpha_{k+1}$ ,  $\mathbf{u}_{k+1}$ ,  $\mathbf{v}_{k+1}$  by [Algorithm 3.1](#); Form  $\mathbf{B}_{k+1}$  and  $\mathbf{V}_k$
- 5:   (Terminate **gen-GKB** if  $\beta_{k+1}$  or  $\alpha_{k+1}$  is extremely small)
- 6:   Compute  $F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  and  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  by [\(3.10\)](#) and [\(3.11\)](#)
- 7:   Compute  $(\Delta\mathbf{y}_k, \Delta\lambda_k)$  by [\(3.12a\)](#)
- 8:   **if**  $\Delta\lambda_k > 0$  **then**
- 9:      $\gamma_{\text{init}} = 1$
- 10:   **else**
- 11:      $\gamma_{\text{init}} = \min\{1, -\eta\lambda_{k-1}/\Delta\lambda_k\}$  ▷ Ensure the positivity of  $\lambda_k$
- 12:   **end if**
- 13:   Determine the step-length  $\gamma_k$  by [Routine 1](#)
- 14:   Update  $(\mathbf{y}_k, \lambda_k)$  by [\(3.12b\)](#)
- 15:   **if**  $\frac{1}{2} \|\bar{F}^{(k)}(\bar{\mathbf{y}}_k, \lambda_k)\|_2 \leq \text{tol}$  **then**
- 16:     Compute  $\mathbf{x}_k = \mathbf{V}_k \mathbf{y}_k$ ; Stop iteration
- 17:   **end if**
- 18: **end for**

**Output:** Final solution  $(\mathbf{x}_k, \lambda_k)$

---

From the derivations, we find that the success of PNT is attributed to the fact that  $\mathbf{M}$  and  $\mathbf{N}$  can induce inner products. Therefore, the PNT method can not be directly used to handle the total variation (TV) or  $\ell_1$  regularization terms. One possible approach for handling TV or  $\ell_1$  norms is to approximate them with weighted  $\ell_2$  norms at each iterated point [50,51]. Furthermore, for the nonlinear inverse problem  $\mathbf{b} = G(\mathbf{x}) + \boldsymbol{\epsilon}$  with differentiable  $G$ , at each iterated point we can approximate  $\|G(\mathbf{x}) - \mathbf{b}\|_{\mathbf{M}^{-1}}^2$  by a quadratic convex function using the first-order Taylor expansion of  $G$ . The above approaches follow a similar idea to the sequential quadratically constrained quadratic programming (SQCQP) method [16,38]. This allows us to obtain a sequence of optimization problems similar to (2.2), which can be solved efficiently by PNT. Theoretical and computational aspects of this approach will be further studied in the future.

**3.2. Proofs.** We give the proofs of all the results in Subsection 3.1. Remember again that we always follow the notations as stated in Remark 2 and (3.13).

**Proof of Lemma 3.2.** By (3.8) and (3.9) we have

$$F^{(k)}(\mathbf{y}, \lambda) = \left( \lambda(\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b}) + \mathbf{V}_k^\top \mathbf{N}^{-1} \mathbf{V}_k \mathbf{y} \quad \frac{1}{2} \|\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \frac{\tau m}{2} \right),$$

and

$$J^{(k)}(\mathbf{y}, \lambda) = \begin{pmatrix} \lambda(\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{V}_k) + \mathbf{V}_k^\top \mathbf{N}^{-1} \mathbf{V}_k & (\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b}) \\ (\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b})^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{V}_k) & 0 \end{pmatrix}.$$

Using relations (3.5) and Proposition 3.1, we have  $\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b} = \mathbf{U}_{k+1}(\mathbf{B}_k - \beta_1 \mathbf{e}_1)$ , leading to

$$(\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b}) = (\mathbf{U}_{k+1} \mathbf{B}_k)^\top \mathbf{M}^{-1} \mathbf{U}_{k+1}(\mathbf{B}_k - \beta_1 \mathbf{e}_1) = \mathbf{B}_k^\top (\mathbf{B}_k - \beta_1 \mathbf{e}_1),$$

and

$$\|\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \|\mathbf{U}_{k+1}(\mathbf{B}_k - \beta_1 \mathbf{e}_1)\|_{\mathbf{M}^{-1}}^2 = \|\mathbf{B}_k - \beta_1 \mathbf{e}_1\|_2^2.$$

If  $\beta_{k_t+1} = 0$ , then for  $k \geq k_t$ , the relation  $\mathbf{A}\mathbf{V}_k \mathbf{y} - \mathbf{b} = \mathbf{U}_{k+1}(\mathbf{B}_k - \beta_1 \mathbf{e}_1)$  is replaced by  $\mathbf{A}\mathbf{V}_{k_t} \mathbf{y} - \mathbf{b} = \mathbf{U}_{k_t}(\mathbf{B}_{k_t} - \beta_1 \mathbf{e}_1)$ . Therefore, the above identity is also applied to the case  $k \geq k_t$ . Now we have proved (3.10). The expression (3.11) can be proved similarly.  $\square$

In order to prove Lemma 3.5, we first give the following result.

**LEMMA 3.7.** Let  $\widehat{\mathbf{N}} = \begin{pmatrix} \mathbf{N} \\ 1 \end{pmatrix}$ . Then we have the following identity:

$$(3.18) \quad \|F(\mathbf{x}_{k-1}, \lambda_{k-1})\|_{\widehat{\mathbf{N}}} = \|F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})\|_2.$$

*Proof.* First notice that

$$\|F(\mathbf{x}_{k-1}, \lambda_{k-1})\|_{\widehat{\mathbf{N}}}^2 = \|\lambda_{k-1} \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_{k-1}\|_{\widehat{\mathbf{N}}}^2 + \left(\frac{1}{2} \|\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \frac{\tau m}{2}\right)^2. \blacksquare$$

For the first term of the above summation, we have

$$\begin{aligned} & \|\lambda_{k-1} \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_{k-1}\|_{\widehat{\mathbf{N}}}^2 \\ &= (\lambda_{k-1} \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_{k-1})^\top \mathbf{N} (\lambda_{k-1} \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_{k-1}). \blacksquare \end{aligned}$$

and

$$\begin{aligned}
& \mathbf{N}(\lambda_{k-1}\mathbf{A}^\top\mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b})+\mathbf{N}^{-1}\mathbf{x}_{k-1}) \\
&= \lambda_{k-1}\mathbf{N}\mathbf{A}^\top\mathbf{M}^{-1}\mathbf{U}_{k+1}(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)+\mathbf{V}_k\bar{\mathbf{y}}_{k-1} \\
&= \lambda_{k-1}(\mathbf{V}_k\mathbf{B}_k^\top+\alpha_{k+1}\mathbf{v}_{k+1}\mathbf{e}_{k+1}^\top)(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)+\mathbf{V}_k\bar{\mathbf{y}}_{k-1} \\
&= \mathbf{V}_k(\lambda_{k-1}\mathbf{B}_k^\top(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)+\bar{\mathbf{y}}_{k-1}),
\end{aligned}$$

where we have used

$$\alpha_{k+1}\mathbf{v}_{k+1}\mathbf{e}_{k+1}^\top(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)=\alpha_{k+1}\beta_{k+1}\mathbf{v}_{k+1}\mathbf{e}_k^\top\bar{\mathbf{y}}_{k-1}=0,$$

because  $\mathbf{e}_k^\top\bar{\mathbf{y}}_{k-1}=0$  for  $k\leq k_t$  and  $\alpha_{k+1}\beta_{k+1}=0$  for  $k>k_t$ . Similarly, we have

$$\lambda_{k-1}\mathbf{A}^\top\mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b})+\mathbf{N}^{-1}\mathbf{x}_{k-1}=\mathbf{N}^{-1}\mathbf{V}_k(\lambda_{k-1}\mathbf{B}_k^\top(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)+\bar{\mathbf{y}}_{k-1}).$$

We also have

$$\|\lambda_{k-1}\mathbf{A}^\top\mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b})+\mathbf{N}^{-1}\mathbf{x}_{k-1}\|_{\mathbf{N}}=\|\lambda_{k-1}\mathbf{B}_k^\top(\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)+\bar{\mathbf{y}}_{k-1}\|_2$$

and

$$\frac{1}{2}\|\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b}\|_{\mathbf{M}^{-1}}^2-\frac{\tau m}{2}=\frac{1}{2}\|\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1\|_2^2-\frac{\tau m}{2}.$$

The desired result immediately follows by using (3.10).  $\square$

**Proof of Lemma 3.5.** First notice that  $h(\mathbf{x},\lambda)=\frac{1}{2}\|F(\mathbf{x},\lambda)\|_{\hat{\mathbf{N}}}^2$ . Combining the above relation with Lemma 3.7 we obtain the first identity of (3.17). Also, for  $k<k_t$  we have  $h(\mathbf{x}_k,\lambda_k)=\frac{1}{2}\|F^{(k+1)}(\bar{\mathbf{y}}_k,\lambda_k)\|_2^2$  with

$$F^{(k+1)}(\bar{\mathbf{y}}_k,\lambda_k)=\begin{pmatrix} \lambda\mathbf{B}_{k+1}^\top(\mathbf{B}_{k+1}\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1)+\bar{\mathbf{y}}_k \\ \frac{1}{2}\|\mathbf{B}_{k+1}\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1\|_2^2-\frac{\tau m}{2} \end{pmatrix}.$$

Since the last element of  $\beta_1\mathbf{e}_1$  and  $\bar{\mathbf{y}}_k$  is zero, it is easy to verify that

$$\mathbf{B}_{k+1}\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1=\begin{pmatrix} \bar{\mathbf{B}}_k\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1 \\ 0 \end{pmatrix}, \quad \mathbf{B}_{k+1}^\top(\mathbf{B}_{k+1}\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1)=\bar{\mathbf{B}}_k^\top(\bar{\mathbf{B}}_k\bar{\mathbf{y}}_k-\beta_1\mathbf{e}_1).$$

For  $k\geq k_t$ , we have  $F^{(k)}(\mathbf{y},\lambda)=F^{(k+1)}(\mathbf{y},\lambda)=\bar{F}^{(k)}(\mathbf{y},\lambda)$ . Therefore, we prove the second identity of (3.17).  $\square$

In order to prove Theorem 3.4, we need Lemma 3.7 and the following result.

LEMMA 3.8. *For any  $k\geq 1$ , we have the following identity:*

$$\begin{pmatrix} \mathbf{V}_k^\top \\ 1 \end{pmatrix} J(\mathbf{x}_{k-1},\lambda_{k-1})\widehat{\mathbf{N}}F(\mathbf{x}_{k-1},\lambda_{k-1})=J^{(k)}(\bar{\mathbf{y}}_{k-1},\lambda_{k-1})F^{(k)}(\bar{\mathbf{y}}_{k-1},\lambda_{k-1}).$$

*Proof.* First notice from (2.8) that

$$\begin{pmatrix} \mathbf{V}_k^\top \\ 1 \end{pmatrix} J(\mathbf{x}_{k-1},\lambda_{k-1})\widehat{\mathbf{N}}=\begin{pmatrix} \lambda_{k-1}(\mathbf{A}\mathbf{V}_k)^\top\mathbf{M}^{-1}\mathbf{A}\mathbf{N}+\mathbf{V}_k^\top & (\mathbf{A}\mathbf{V}_k)^\top\mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b}) \\ (\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b})^\top\mathbf{M}^{-1}\mathbf{A}\mathbf{N} & 0 \end{pmatrix}. \quad \blacksquare$$

Using (3.5c) and similar derivations to the proof of Lemma 3.7, we get

$$\begin{aligned}
(\mathbf{A}\mathbf{x}_{k-1}-\mathbf{b})^\top\mathbf{M}^{-1}\mathbf{A}\mathbf{N} &= (\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)^\top\mathbf{U}_{k+1}^\top\mathbf{M}^{-1}\mathbf{A}\mathbf{N} \\
&= (\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)^\top(\mathbf{B}_k\mathbf{V}_k^\top+\alpha_{k+1}\mathbf{e}_{k+1}\mathbf{v}_{k+1}^\top) \\
(3.19) \quad &= (\mathbf{B}_k\bar{\mathbf{y}}_{k-1}-\beta_1\mathbf{e}_1)^\top\mathbf{B}_k\mathbf{V}_k^\top.
\end{aligned}$$

Also, we can get

$$\begin{aligned} (\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1} \mathbf{A}\mathbf{N} &= (\mathbf{U}_{k+1} \mathbf{B}_k)^\top \mathbf{M}^{-1} \mathbf{A}\mathbf{N} = \mathbf{B}_k^\top (\mathbf{B}_k \mathbf{V}_k^\top + \alpha_{k+1} e_{k+1} \mathbf{v}_{k+1}^\top) \\ &= \mathbf{B}_k^\top \mathbf{B}_k \mathbf{V}_k^\top + \alpha_{k+1} \beta_{k+1} e_k \mathbf{v}_{k+1}^\top, \end{aligned}$$

and

$$\begin{aligned} (\mathbf{A}\mathbf{V}_k)^\top \mathbf{M}^{-1} (\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) &= (\mathbf{B}_k \mathbf{U}_{k+1})^\top \mathbf{M}^{-1} \mathbf{U}_{k+1} (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1) \\ &= \mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1). \end{aligned}$$

Using (3.11), we get

$$\begin{aligned} &\begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} J(\mathbf{x}_{k-1}, \lambda_{k-1}) \widehat{\mathbf{N}} \\ &= \begin{pmatrix} (\lambda_{k-1} \mathbf{B}_k^\top \mathbf{B}_k + \mathbf{I}) \mathbf{V}_k^\top & \mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1) \\ (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top & \mathbf{B}_k \mathbf{V}_k^\top \\ & 0 \end{pmatrix} + \begin{pmatrix} \alpha_{k+1} \beta_{k+1} e_{k+1} \mathbf{v}_{k+1}^\top & \\ & 0 \end{pmatrix} \\ &= J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}) \begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} + \begin{pmatrix} \alpha_{k+1} \beta_{k+1} e_k \mathbf{v}_{k+1}^\top & \\ & 0 \end{pmatrix}. \end{aligned}$$

Using similar derivations to the proof of Lemma 3.7, we get

$$\begin{aligned} F(\mathbf{x}_{k-1}, \lambda_{k-1}) &= \begin{pmatrix} \mathbf{N}^{-1} & \\ & 1 \end{pmatrix} \begin{pmatrix} \lambda_{k-1} \mathbf{N} \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}) + \mathbf{x}_{k-1} \\ \frac{1}{2} \|\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \frac{\tau m}{2} \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{N}^{-1} & \\ & 1 \end{pmatrix} \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix} \begin{pmatrix} \lambda_{k-1} \mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1) + \bar{\mathbf{y}}_{k-1} \\ \frac{1}{2} \|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2^2 - \frac{\tau m}{2} \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{N}^{-1} & \\ & 1 \end{pmatrix} \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix} F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}). \end{aligned}$$

Using the relations

$$\begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} \begin{pmatrix} \mathbf{N}^{-1} & \\ & 1 \end{pmatrix} \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix} = \mathbf{I}, \quad \begin{pmatrix} \alpha_{k+1} \beta_{k+1} e_k \mathbf{v}_{k+1}^\top & \\ & 0 \end{pmatrix} \begin{pmatrix} \mathbf{N}^{-1} & \\ & 1 \end{pmatrix} \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix} = \mathbf{0},$$

we finally obtain the desired result.  $\square$

**Proof of Theorem 3.4.** Notice  $J(\mathbf{x}, \lambda)$  is the Jacobian of  $F(\mathbf{x}, \lambda)$ . Using  $h(\mathbf{x}, \lambda) = \frac{1}{2} \|F(\mathbf{x}, \lambda)\|_{\widehat{\mathbf{N}}}^2$  we get  $\nabla h(\mathbf{x}, \lambda) = J(\mathbf{x}, \lambda) \widehat{\mathbf{N}} F(\mathbf{x}, \lambda)$ , leading to

$$\begin{aligned} \nabla h(\mathbf{x}_{k-1}, \lambda_{k-1})^\top \begin{pmatrix} \Delta \mathbf{x}_k \\ \Delta \lambda_k \end{pmatrix} &= \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix}^\top \begin{pmatrix} \mathbf{V}_k^\top & \\ & 1 \end{pmatrix} J(\mathbf{x}_{k-1}, \lambda_{k-1}) \widehat{\mathbf{N}} F(\mathbf{x}_{k-1}, \lambda_{k-1}) \\ &= \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix}^\top J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}) F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}) \\ &= -\|F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})\|_2^2 = -\|F(\mathbf{x}_{k-1}, \lambda_{k-1})\|_{\widehat{\mathbf{N}}}^2 \\ &= -2h(\mathbf{x}_{k-1}, \lambda_{k-1}) \leq 0, \end{aligned}$$

where we have used Lemma 3.7 and Lemma 3.8.  $\square$

**4. Convergence analysis.** The objective of this section is to prove the convergence of PNT, which is stated in the following result and [Corollary 4.8](#).

**THEOREM 4.1.** *Suppose the PNT algorithm is initialized with  $\bar{\mathbf{y}}_0 = \mathbf{0}$ ,  $\mathbf{x}_0 = \mathbf{0}$  and  $\lambda_0 > 0$ . Then we either have*

$$(4.1) \quad h(\mathbf{x}_k, \lambda_k) = 0$$

for some  $k < \infty$ , or have

$$(4.2) \quad \lim_{k \rightarrow \infty} h(\mathbf{x}_k, \lambda_k) = 0.$$

Notice that  $(\mathbf{x}^*, \lambda^*)$  is the unique minimizer of  $h(\mathbf{x}, \lambda)$  and  $h(\mathbf{x}^*, \lambda^*) = 0$ . Therefore, (4.1) implies that the algorithm finds the exact solution to (2.2) and (2.3) at the  $k$ -th iteration. In the following part, we prove (4.2) under the assumption that  $h(\mathbf{x}_k, \lambda_k) > 0$  for any  $k \geq 1$ . We need a series of lemmas, which are [Lemma 4.2](#)–[Lemma 4.7](#). All these lemmas follow the same assumption of [Theorem 4.1](#).

**LEMMA 4.2.** *For any matrix  $\mathbf{C} \in \mathbb{R}^{m \times n}$  with full column rank and  $\mathbf{d} \in \mathbb{R}^m$ , if the vector sequence  $\{\mathbf{w}_k\} \in \mathbb{R}^n$  satisfies  $\lim_{k \rightarrow \infty} \|\mathbf{C}^\top(\mathbf{C}\mathbf{w}_k - \mathbf{d})\|_2 = 0$ , then  $\lim_{k \rightarrow \infty} \mathbf{w}_k = \mathbf{w}^* := \operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^n} \|\mathbf{C}\mathbf{w} - \mathbf{d}\|_2$ .*

*Proof.* First notice that  $\mathbf{w}^*$  is well-defined, since  $\operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^n} \|\mathbf{C}\mathbf{w} - \mathbf{d}\|_2$  has a unique solution for the full column rank matrix  $\mathbf{C}$ . For any  $\mathbf{w}_k$ , let  $\mathbf{w}_k = \mathbf{w}^* + \bar{\mathbf{w}}_k$ . Then we have

$$\lim_{k \rightarrow \infty} \|\mathbf{C}^\top(\mathbf{C}\mathbf{w}_k - \mathbf{d})\|_2 = \lim_{k \rightarrow \infty} \|\mathbf{C}^\top(\mathbf{C}\mathbf{w}^* - \mathbf{d}) + \mathbf{C}^\top\mathbf{C}\bar{\mathbf{w}}_k\|_2 = \lim_{k \rightarrow \infty} \|\mathbf{C}^\top\mathbf{C}\bar{\mathbf{w}}_k\|_2 = 0,$$

since  $\mathbf{C}^\top(\mathbf{C}\mathbf{w}^* - \mathbf{d}) = \mathbf{0}$ . Now we have  $\|\bar{\mathbf{w}}_k\|_2 \rightarrow 0$  since  $\mathbf{C}^\top\mathbf{C}$  is positive definite and all norms of  $\mathbb{R}^n$  are equivalent. Therefore, we have  $\|\mathbf{w}_k - \mathbf{w}^*\|_2 \rightarrow 0$  or the equivalent form  $\lim_{k \rightarrow \infty} \mathbf{w}_k = \mathbf{w}^*$ .  $\square$

**LEMMA 4.3.** *If the unique solution to  $\min_{\mathbf{y} \in \mathbb{R}^{k_t}} \|\mathbf{B}_{k_t}\mathbf{y} - \beta_1\mathbf{e}_1\|_2$  is  $\mathbf{y}_{\min}$ , then  $\mathbf{x}_{\min} := \mathbf{V}_{k_t}\mathbf{y}_{\min}$  is the unique solution to*

$$(4.3) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{x}\|_{\mathbf{N}^{-1}} \quad \text{s.t.} \quad \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_{\mathbf{M}^{-1}} = \min.$$

*Proof.* It is easy to verify that both  $\min_{\mathbf{y} \in \mathbb{R}^{k_t}} \|\mathbf{B}_{k_t}\mathbf{y} - \beta_1\mathbf{e}_1\|_2$  and (4.3) have a unique solution. A vector  $\mathbf{x}$  is the unique solution to (4.3) if and only if

$$\mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x} - \mathbf{b}) = \mathbf{0}, \quad \mathbf{x} \perp_{\mathbf{N}^{-1}} \mathcal{N}(\mathbf{A}),$$

where  $\perp_{\mathbf{N}^{-1}}$  means the orthogonality relation under the  $\mathbf{N}^{-1}$ -inner product. Now we verify the above two conditions for  $\mathbf{x}_{\min}$ . For the first condition, using the relations  $\mathbf{A}\mathbf{x}_{\min} = \mathbf{A}\mathbf{V}_{k_t}\mathbf{y}_{\min} = \mathbf{U}_{k_t+1}\mathbf{B}_{k_t}\mathbf{y}_{\min}$  and (3.5c), we get

$$\begin{aligned} \mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{\min} - \mathbf{b}) &= \mathbf{A}^\top \mathbf{M}^{-1}\mathbf{U}_{k_t+1}(\mathbf{B}_{k_t}\mathbf{y}_{\min} - \beta_1\mathbf{e}_1) \\ &= \mathbf{N}^{-1}(\mathbf{V}_{k_t}\mathbf{B}_{k_t}^\top + \alpha_{k_t+1}\mathbf{v}_{k_t+1}\mathbf{e}_{k_t+1}^\top)(\mathbf{B}_{k_t}\mathbf{y}_{\min} - \beta_1\mathbf{e}_1) \\ &= \mathbf{N}^{-1}[\mathbf{V}_{k_t}\mathbf{B}_{k_t}^\top(\mathbf{B}_{k_t}\mathbf{y}_{\min} - \beta_1\mathbf{e}_1) + \alpha_{k_t+1}\beta_{k_t+1}\mathbf{v}_{k_t+1}\mathbf{e}_{k_t+1}^\top\mathbf{y}_{\min}] \\ &= \mathbf{0}, \end{aligned}$$

since  $\mathbf{B}_{k_t}^\top (\mathbf{B}_{k_t} \mathbf{y}_{\min} - \beta_1 e_1) = \mathbf{0}$  and  $\alpha_{k_t+1} \beta_{k_t+1} = 0$ . For the second condition, by [Proposition 3.1](#) we have

$$\begin{aligned} \mathbf{x}_{\min} &\in \text{span}\{\mathbf{V}_{k_t}\} = \text{span}\{(\mathbf{N}\mathbf{A}^\top \mathbf{M}^{-1} \mathbf{A})^i \mathbf{N}\mathbf{A}^\top \mathbf{M}^{-1} \mathbf{b}\}_{i=0}^{k_t-1} \\ &\subseteq \mathcal{R}(\mathbf{N}\mathbf{A}^\top) = \mathbf{N}\mathcal{N}(\mathbf{A})^\perp. \end{aligned}$$

Write  $\mathbf{x}_{\min} = \mathbf{N}\bar{\mathbf{x}}_{\min}$  with  $\bar{\mathbf{x}}_{\min} \in \mathcal{N}(\mathbf{A})^\perp$ . For any  $\mathbf{w} \in \mathcal{N}(\mathbf{A})$ , we have

$$\langle \mathbf{x}_{\min}, \mathbf{w} \rangle_{\mathbf{N}^{-1}} = \langle \mathbf{N}\bar{\mathbf{x}}_{\min}, \mathbf{w} \rangle_{\mathbf{N}^{-1}} = \langle \bar{\mathbf{x}}_{\min}, \mathbf{w} \rangle_2 = 0.$$

Therefore, it holds that  $\mathbf{x}_{\min} \perp_{\mathbf{N}^{-1}} \mathcal{N}(\mathbf{A})$ . □

LEMMA 4.4. *There exist a positive constant  $C_1$  such that for any  $k \geq 1$ ,*

$$(4.4) \quad \|\mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})\|_{\mathbf{N}} = \|\mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)\|_2 \geq C_1 > 0$$

*Proof.* First, we get from (3.19) the first identity:

$$\begin{aligned} &\|\mathbf{A}^\top \mathbf{M}^{-1}(\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})\|_{\mathbf{N}}^2 \\ &= (\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})^\top \mathbf{M}^{-1} \mathbf{A} \mathbf{N} \mathbf{N}^{-1} ((\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})^\top \mathbf{M}^{-1} \mathbf{A} \mathbf{N})^\top \\ &= (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top \mathbf{B}_k \mathbf{V}_k^\top \mathbf{N}^{-1} \mathbf{V}_k \mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1) \\ &= \|\mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)\|_2^2. \end{aligned}$$

Then, we prove

$$(4.5) \quad \|\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}} = \|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2 \geq \sqrt{\tau m}$$

by mathematical induction. For  $k = 1$ , we have  $\|\mathbf{A}\mathbf{x}_0 - \mathbf{b}\|_{\mathbf{M}^{-1}} = \|\mathbf{b}\|_{\mathbf{M}^{-1}} > \sqrt{\tau m}$  since  $\mathbf{x}_0 = \mathbf{0}$ . Suppose  $\|\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}} \geq \sqrt{\tau m}$  for  $k \geq 1$ . We have

$$\begin{aligned} &\|\mathbf{A}\mathbf{x}_k - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 = \|\mathbf{A}\mathbf{V}_k(\bar{\mathbf{y}}_{k-1} + \gamma_k \Delta \mathbf{y}_k) - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 \\ &= \|\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 + \gamma_k^2 \|\mathbf{A}\mathbf{V}_k \Delta \mathbf{y}_k\|_{\mathbf{M}^{-1}}^2 + 2\gamma_k (\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})^\top \mathbf{M}^{-1} \mathbf{A}\mathbf{V}_k \Delta \mathbf{y}_k \\ &= \|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2^2 + \gamma_k^2 \|\mathbf{A}\mathbf{V}_k \Delta \mathbf{y}_k\|_{\mathbf{M}^{-1}}^2 + 2\gamma_k (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top \mathbf{B}_k \Delta \mathbf{y}_k, \end{aligned}$$

since

$$\begin{aligned} (\mathbf{A}\mathbf{x}_{k-1} - \mathbf{b})^\top \mathbf{M}^{-1} \mathbf{A}\mathbf{V}_k &= (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top \mathbf{U}_{k+1}^\top \mathbf{M}^{-1} \mathbf{U}_{k+1} \mathbf{B}_k \\ &= (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top \mathbf{B}_k. \end{aligned}$$

Writing  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}) \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix} = -F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  in the matrix form and using  $\|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2 \geq \sqrt{\tau m}$ , we get from the second equality of the above equation that  $(\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1)^\top \mathbf{B}_k \Delta \mathbf{y}_k = -\frac{1}{2} (\|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2^2 - \tau m) \leq 0$ . Since  $\gamma_k \leq 1$ , we get

$$\begin{aligned} &\|\mathbf{A}\mathbf{x}_k - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 \\ &\geq \|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2^2 + \gamma_k^2 \|\mathbf{A}\mathbf{V}_k \Delta \mathbf{y}_k\|_{\mathbf{M}^{-1}}^2 - (\|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 e_1\|_2^2 - \tau m) \\ &= \tau m + \gamma_k^2 \|\mathbf{A}\mathbf{V}_k \Delta \mathbf{y}_k\|_{\mathbf{M}^{-1}}^2 \geq \tau m. \end{aligned}$$

Therefore, we prove (4.5).

To obtain the lower bound in (4.4), we investigate two cases:  $k < k_t$  and  $k \geq k_t$ .

**Case 1:**  $k < k_t$ . For this case, we have

$$\|\mathbf{B}_k^\top (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1)\|_2 \geq \sigma_{\min}(\mathbf{B}_k) \|\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1\|_2 \geq \sigma_{\min}(\underline{\mathbf{B}}_{k_t}) \sqrt{\tau m} > 0,$$

where  $\sigma_{\min}(\cdot)$  is the smallest singular value of a matrix, and  $\underline{\mathbf{B}}_{k_t}$  is the first  $k_t \times k_t$  part of  $\mathbf{B}_{k_t}$ .

**Case 2:**  $k \geq k_t$ . For this case, we can write  $\mathbf{x}_{k-1}$  as  $\mathbf{x}_{k-1} = \mathbf{V}_{k_t} \bar{\mathbf{y}}_{k-1}$ . Remember that  $\bar{\mathbf{y}}_{k-1} = \mathbf{y}_{k-1}$  if  $k > k_t$ . We first prove  $\|\mathbf{B}_{k_t}^\top (\mathbf{B}_{k_t} \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1)\|_2 \neq 0$ . If it is not true, then  $\bar{\mathbf{y}}_{k-1} = \arg\min_{\mathbf{y}} \|\mathbf{B}_{k_t} \mathbf{y} - \beta_1 \mathbf{e}_1\|_2$ . By Lemma 4.3,  $\mathbf{x}_{k-1}$  is the solution to (4.3). Thus, it must hold that  $\|\mathbf{A} \mathbf{x}_{k-1} - \mathbf{b}\|_{\mathbf{M}^{-1}} < \sqrt{\tau m}$  by Assumption 1, which contradicts (4.5). Now suppose the lower bound in (4.4) is not true. Then there exists a subsequence  $\{\bar{\mathbf{y}}_{k_j-1}\}$  with  $k_j \geq k_t$  such that  $\lim_{j \rightarrow \infty} \|\mathbf{B}_{k_t}^\top (\mathbf{B}_{k_t} \bar{\mathbf{y}}_{k_j-1} - \beta_1 \mathbf{e}_1)\|_2 = 0$ . By Lemma 4.2, we have  $\lim_{j \rightarrow \infty} \bar{\mathbf{y}}_{k_j-1} = \mathbf{y}_{\min} := \arg\min_{\mathbf{y}} \|\mathbf{B}_{k_t} \mathbf{y} - \beta_1 \mathbf{e}_1\|_2$ , leading to  $\lim_{j \rightarrow \infty} \mathbf{x}_{k_j-1} = \lim_{k_j \rightarrow \infty} \mathbf{V}_{k_t} \bar{\mathbf{y}}_{k_j-1} = \mathbf{V}_{k_t} \mathbf{y}_{\min}$ . It follows from Lemma 4.3 that  $\mathbf{V}_{k_t} \mathbf{y}_{\min} = \mathbf{x}_{\min}$ , which is the solution to (4.3). Therefore, it must hold  $\lim_{j \rightarrow \infty} \|\mathbf{A} \mathbf{x}_{k_j-1} - \mathbf{b}\|_{\mathbf{M}^{-1}} = \|\mathbf{A} \mathbf{x}_{\min} - \mathbf{b}\|_{\mathbf{M}^{-1}} < \sqrt{\tau m}$  by Assumption 1, which contradicts (4.5). Summarizing both the two cases, the desired result is proved.  $\square$

LEMMA 4.5. *The points  $\{(\mathbf{x}_k, \lambda_k)\}_{i=0}^\infty$  generated by the PNT algorithm lie in a bounded set of  $\mathbb{R}^n \times \mathbb{R}^+$ .*

*Proof.* First notice that  $h(\mathbf{x}_0, \lambda_0) \geq h(\mathbf{x}_1, \lambda_1) \geq \dots$ . We only need to prove  $\{(\mathbf{x}_k, \lambda_k)\}_{k \geq k_t}$  is bounded above. In this case, notice that  $\mathbf{x}_k = \mathbf{V}_{k_t} \mathbf{y}_k$  and

$$h(\mathbf{x}_k, \lambda_k) = \frac{1}{2} [\|\lambda_k \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A} \mathbf{x}_k - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_k\|_{\mathbf{N}}^2 + (\frac{1}{2} \|\mathbf{B}_{k_t} \mathbf{y}_k - \beta_1 \mathbf{e}_1\|_2^2 - \frac{\tau m}{2})^2].$$

If the points do not lie in a bounded set, there exists a subsequence  $\{(\mathbf{x}_{k_j}, \lambda_{k_j})\}$  with  $k_j \geq k_t$  such that  $(\mathbf{x}_{k_j}, \lambda_{k_j}) \rightarrow \infty$ . If  $\|\mathbf{x}_{k_j}\|_2 \rightarrow \infty$ , then  $\|\mathbf{y}_{k_j}\|_2 \rightarrow \infty$ , since  $\mathbf{x}_{k_j} = \mathbf{V}_{k_t} \mathbf{y}_{k_j}$  and  $\mathbf{V}_{k_t}$  has full column rank. This leads to  $\|\mathbf{B}_{k_t} \mathbf{y}_{k_j}\|_2 \rightarrow \infty$  since  $\mathbf{B}_{k_t}$  has full column rank. It follows that the second term of  $h(\mathbf{x}_{k_j}, \lambda_{k_j})$  tends to infinity and  $h(\mathbf{x}_{k_j}, \lambda_{k_j}) \rightarrow \infty$ , a contradiction. Therefore, it must hold that  $\|\mathbf{x}_{k_j}\|_2$  is bounded above and  $\lambda_{k_j} \rightarrow \infty$ . By Lemma 4.4, we have  $\|\lambda_{k_j} \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A} \mathbf{x}_{k_j} - \mathbf{b})\|_{\mathbf{N}} \geq \lambda_{k_j} C_1$ . Notice that  $\{\mathbf{N}^{-1} \mathbf{x}_{k_j}\}$  lie in a bounded set. It follows that  $\|\lambda_{k_j} \mathbf{A}^\top \mathbf{M}^{-1} (\mathbf{A} \mathbf{x}_{k_j} - \mathbf{b}) + \mathbf{N}^{-1} \mathbf{x}_{k_j}\|_{\mathbf{N}} \rightarrow \infty$  and  $h(\mathbf{x}_{k_j}, \lambda_{k_j}) \rightarrow \infty$ , also a contradiction.  $\square$

LEMMA 4.6. *There exist a positive constant  $C_2 < +\infty$  such that for any  $k \geq 1$ ,*

$$(4.6) \quad \|J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})^{-1}\|_2 \leq C_2.$$

*Proof.* First, we prove that  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  is always nonsingular. Write it as

$$J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1}) = \begin{pmatrix} \lambda_{k-1} \mathbf{B}_{k_t}^\top \mathbf{B}_k + \mathbf{I} & \mathbf{B}_k^\top (\mathbf{B}_{k_t} \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1) \\ (\mathbf{B}_k \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1)^\top \mathbf{B}_k & 0 \end{pmatrix} =: \begin{pmatrix} \mathbf{C}_k & \mathbf{d}_k \\ \mathbf{d}_k^\top & 0 \end{pmatrix}$$

and notice that

$$\begin{pmatrix} \mathbf{C}_k & \mathbf{d}_k \\ \mathbf{d}_k^\top & 0 \end{pmatrix} = \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ -\mathbf{d}_k^\top \mathbf{C}_k^{-1} & 1 \end{pmatrix}^{-1} \begin{pmatrix} \mathbf{C}_k & \mathbf{0} \\ \mathbf{0} & -\mathbf{d}_k^\top \mathbf{C}_k^{-1} \mathbf{d}_k \end{pmatrix} \begin{pmatrix} \mathbf{I} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & 1 \end{pmatrix}^{-1}.$$

It follows that

$$(4.7) \quad \begin{pmatrix} \mathbf{C}_k & \mathbf{d}_k \\ \mathbf{d}_k^\top & 0 \end{pmatrix}^{-1} = \begin{pmatrix} \mathbf{I} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & 1 \end{pmatrix} \begin{pmatrix} \mathbf{C}_k^{-1} & \mathbf{0} \\ \mathbf{0} & -(\mathbf{d}_k^\top \mathbf{C}_k^{-1} \mathbf{d}_k)^{-1} \end{pmatrix} \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ -\mathbf{d}_k^\top \mathbf{C}_k^{-1} & 1 \end{pmatrix},$$

Since  $\mathbf{C}_k$  is positive definite and  $\|\mathbf{d}_k\|_2 \geq C_1 > 0$ .

To give an upper bound on  $\|J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})^{-1}\|_2$ , we only need to consider  $k \geq k_t$ , where  $\mathbf{B}_k = \mathbf{B}_{k_t}$  in  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$ . Since  $\sigma_{\min}(\mathbf{C}_k) = \sigma_{\min}(\lambda_{k-1} \mathbf{B}_{k_t}^\top \mathbf{B}_{k_t} + \mathbf{I}) \geq 1$ , we have  $\|\mathbf{C}_k^{-1}\|_2 \leq 1$ . By Lemma 4.5, there exist a positive constant  $C_3 < +\infty$  such that  $\lambda_k \leq C_3$ , thereby

$$\sigma_{\max}(\mathbf{C}_k) \leq \sigma_{\max}(\lambda_{k-1} \mathbf{B}_{k_t}^\top \mathbf{B}_{k_t}) + \sigma_{\max}(\mathbf{I}) \leq C_3 \sigma_{\max}(\mathbf{B}_{k_t}^\top \mathbf{B}_{k_t}) + 1 =: \bar{C}_3.$$

By Lemma 4.4 we have  $\|\mathbf{d}_k\|_2 = \|\mathbf{B}_{k_t}^\top (\mathbf{B}_{k_t} \bar{\mathbf{y}}_{k-1} - \beta_1 \mathbf{e}_1)\|_2 \geq C_1$ . On the other hand, by Lemma 4.5 we know that  $\|\mathbf{x}_{k-1}\|_2 = \|\mathbf{V}_{k_t} \bar{\mathbf{y}}_{k-1}\|_2$  is bounded above, thereby  $\|\bar{\mathbf{y}}_{k-1}\|_2$  is bounded above since  $\mathbf{V}_{k_t}$  has full column rank. Thus, there exists a positive constant  $\bar{C}_1$  such that  $\|\mathbf{d}_k\|_2 \leq \bar{C}_1$ , leading to

$$\|\mathbf{C}_k^{-1} \mathbf{d}_k\|_2 \leq \|\mathbf{C}_k^{-1}\|_2 \|\mathbf{d}_k\|_2 \leq \bar{C}_1 \sigma_{\min}(\mathbf{C}_k)^{-1} \leq \bar{C}_1,$$

and

$$\mathbf{d}_k^\top \mathbf{C}_k^{-1} \mathbf{d}_k \geq \sigma_{\min}(\mathbf{C}_k^{-1}) \|\mathbf{d}_k\|_2^2 = \sigma_{\max}(\mathbf{C}_k)^{-1} \|\mathbf{d}_k\|_2^2 \geq C_1^2 / \bar{C}_3 > 0.$$

Therefore, we have

$$\begin{aligned} \left\| \begin{pmatrix} \mathbf{0} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \right\|_2^2 &= \left\| \begin{pmatrix} \mathbf{0} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & \mathbf{0} \end{pmatrix}^\top \begin{pmatrix} \mathbf{0} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \right\|_2 \\ &= \left\| \begin{pmatrix} \mathbf{0} & \|\mathbf{C}_k^{-1} \mathbf{d}_k\|_2^2 \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \right\|_2 \leq \bar{C}_1^2 \end{aligned}$$

and

$$\left\| \begin{pmatrix} \mathbf{I} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & \mathbf{1} \end{pmatrix} \right\|_2 \leq \left\| \begin{pmatrix} \mathbf{I} & \\ & \mathbf{1} \end{pmatrix} \right\|_2 + \left\| \begin{pmatrix} \mathbf{0} & -\mathbf{C}_k^{-1} \mathbf{d}_k \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \right\|_2 \leq 1 + \bar{C}_1.$$

Similarly, we have

$$\left\| \begin{pmatrix} \mathbf{C}_k^{-1} & \mathbf{0} \\ \mathbf{0} & -(\mathbf{d}_k^\top \mathbf{C}_k^{-1} \mathbf{d}_k)^{-1} \end{pmatrix} \right\|_2 \leq \max\{1, \bar{C}_3 / C_1^2\}.$$

Using the expression of  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  in (4.7), we finally obtain the desired result.  $\square$

LEMMA 4.7. *There exists a positive constant  $C_4$  such that for any  $k \geq 1$ ,*

$$(4.8) \quad \gamma_k \geq C_4 > 0.$$

*Proof.* By Theorem 3.6 and Theorem 3.4, at each iteration the Armijo backtracking line search must terminate in finite steps with a  $\gamma_k$  satisfying

$$\gamma_k \geq \min \left\{ 1, \frac{4(1-c)\eta h(\mathbf{x}_{k-1}, \lambda_{k-1})}{\zeta(\mathbf{x}_{k-1}, \lambda_{k-1}) \|(\Delta \mathbf{x}_k^\top, \Delta \lambda_k)\|_2^2} \right\},$$

where  $\zeta(\mathbf{x}_{k-1}, \lambda_{k-1})$  is the Lipschitz constant of  $\nabla h$  at  $(\mathbf{x}_{k-1}, \lambda_{k-1})$ ; see also [4, pp. 122, Corollary 2.1]. Now we prove  $\zeta(\mathbf{x}_{k-1}, \lambda_{k-1})$  are bounded above. Notice that  $\nabla h(\mathbf{x}, \lambda) = J(\mathbf{x}, \lambda) \hat{N}F(\mathbf{x}, \lambda)$ . Thus, all the elements in the Jacobian of  $\nabla h(\mathbf{x}, \lambda)$  are polynomials of  $(\mathbf{x}, \lambda)$  with degrees not bigger than 4. Since  $\{(\mathbf{x}_{k-1}, \lambda_{k-1})\}$  lie in a bounded set, the norms of the Jacobians of  $\nabla h(\mathbf{x}, \lambda)$  at the points  $\{(\mathbf{x}_{k-1}, \lambda_{k-1})\}$  are bounded above. Therefore, the Lipschitz constants  $\zeta(\mathbf{x}_{k-1}, \lambda_{k-1})$  are bounded above.

Let  $\zeta(\mathbf{x}_{k-1}, \lambda_{k-1}) \leq \zeta_0$  with  $0 < \zeta_0 < +\infty$  for any  $k \geq 1$ . Then by Lemma 4.6 and Lemma 3.7, we have

$$\begin{aligned} \left\| \begin{pmatrix} \Delta \mathbf{x}_k \\ \Delta \lambda_k \end{pmatrix} \right\|_2 &\leq \left\| \begin{pmatrix} \mathbf{V}_k & \\ & 1 \end{pmatrix} \right\|_2 \left\| \begin{pmatrix} \Delta \mathbf{y}_k \\ \Delta \lambda_k \end{pmatrix} \right\|_2 \\ &\leq (\|\mathbf{V}_{k_t}\|_2 + 1) \|J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})^{-1}\|_2 \|F^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})\|_2 \\ &\leq C_2 (\|\mathbf{V}_{k_t}\|_2 + 1) (2h(\mathbf{x}_{k-1}, \lambda_{k-1}))^{1/2}. \end{aligned}$$

Then we obtain

$$\gamma_k \geq \min \left\{ 1, \frac{4(1-c)\eta h(\mathbf{x}_{k-1}, \lambda_{k-1})}{\zeta_0 \|(\Delta \mathbf{x}_k^\top, \Delta \lambda_k)^\top\|_2^2} \right\} \geq \min \left\{ 1, \frac{2(1-c)\eta}{\zeta_0 C_2^2 (\|\mathbf{V}_{k_t}\|_2 + 1)^2} \right\} =: C_4.$$

The desired result is obtained.  $\square$

**Proof of Theorem 4.1.** By Lemma 4.5, the sequence  $\{(\mathbf{x}_k, \lambda_k)\}_{k=1}^\infty$  is contained in a bounded set, thereby there exists a convergent subsequence  $\{(\mathbf{x}_{k_j}, \lambda_{k_j})\}_{j=1}^\infty$ . Suppose  $(\mathbf{x}_{k_j}, \lambda_{k_j}) \rightarrow (\hat{\mathbf{x}}, \hat{\lambda})$ . It follows that  $h(\mathbf{x}_{k_j}, \lambda_{k_j}) \rightarrow h(\hat{\mathbf{x}}, \hat{\lambda})$  since  $h(\mathbf{x}, \lambda)$  is continuous. Note that  $h(\mathbf{x}_{k_j}, \lambda_{k_j})$  is nonincreasing, thereby  $h(\hat{\mathbf{x}}, \hat{\lambda}) \leq h(\mathbf{x}_{k_j}, \lambda_{k_j})$  for any  $k_j$ . Thus, for any  $\varepsilon > 0$ , there exist a  $k_\star \in \mathbb{N}$  such that  $h(\mathbf{x}_{k_j}, \lambda_{k_j}) < h(\hat{\mathbf{x}}, \hat{\lambda}) + \varepsilon$ ,  $k_j > k_\star$ . Select one  $k_j$  that satisfies  $k_j > k_\star$ . For any  $k \geq k_j$ , we have  $h(\mathbf{x}_k, \lambda_k) \leq h(\mathbf{x}_{k_j}, \lambda_{k_j}) < h(\hat{\mathbf{x}}, \hat{\lambda}) + \varepsilon$ , which means that  $\lim_{k \rightarrow \infty} h(\mathbf{x}_k, \lambda_k) = h(\hat{\mathbf{x}}, \hat{\lambda})$ . The Armijo condition and Theorem 3.4 lead to  $h(\mathbf{x}_{k+1}, \lambda_{k+1}) - h(\mathbf{x}_k, \lambda_k) \leq c\gamma_k (\Delta \mathbf{x}_k^\top, \Delta \lambda_k) \nabla h(\mathbf{x}_k, \lambda_k) \leq 0$ . Taking the limit on both sides leads to  $\lim_{k \rightarrow \infty} c\gamma_k (\Delta \mathbf{x}_k^\top, \Delta \lambda_k) \nabla h(\mathbf{x}_k, \lambda_k) = 0$ . By Lemma 4.7 we get  $\lim_{k \rightarrow \infty} (\Delta \mathbf{x}_k^\top, \Delta \lambda_k) \nabla h(\mathbf{x}_k, \lambda_k) = 0$ . Noticing by Theorem 3.4 that  $-2h(\mathbf{x}_k, \lambda_k) = (\Delta \mathbf{x}_k^\top, \Delta \lambda_k) \nabla h(\mathbf{x}_k, \lambda_k)$ , we get  $h(\hat{\mathbf{x}}, \hat{\lambda}) = \lim_{k \rightarrow \infty} h(\mathbf{x}_k, \lambda_k) = 0$ . This proves the desired result.  $\square$

Now we can give the convergence result of  $(\mathbf{x}_k, \lambda_k)$ .

**COROLLARY 4.8.** *The sequence  $\{(\mathbf{x}_k, \lambda_k)\}_{k=0}^\infty$  generated by the PNT algorithm eventually converges to  $(\mathbf{x}^*, \lambda^*)$ , i.e. the solution of (2.2) and the corresponding Lagrange multiplier.*

*Proof.* Using the same notations as the proof of Theorem 4.1, we obtain that  $(\hat{\mathbf{x}}, \hat{\lambda}) = (\mathbf{x}^*, \lambda^*)$ , since  $h(\mathbf{x}, \lambda)$  has the unique zero point  $(\mathbf{x}^*, \lambda^*)$ . Therefore, the subsequence  $\{(\mathbf{x}_{k_j}, \lambda_{k_j})\}_{j=1}^\infty$  defined in the proof of Theorem 4.1 converges to  $(\mathbf{x}^*, \lambda^*)$ . Now we need to prove the whole sequence  $\{(\mathbf{x}_k, \lambda_k)\}_{k=1}^\infty$  converges to  $(\mathbf{x}^*, \lambda^*)$ . Assume that there is a subsequence  $\{(\mathbf{x}_{l_j}, \lambda_{l_j})\}_{j=1}^\infty$  that does not converge to  $(\mathbf{x}^*, \lambda^*)$ . We can select a subsequence from  $\{(\mathbf{x}_{l_j}, \lambda_{l_j})\}_{j=1}^\infty$  that converges to a point  $(\bar{\mathbf{x}}, \bar{\lambda}) \neq (\mathbf{x}^*, \lambda^*)$ . Since  $h(\mathbf{x}_{l_j}, \lambda_{l_j})$  is nonincreasing with respect to  $j$ , using the same procedure as the proof of Theorem 4.1, we can obtain again that  $h(\bar{\mathbf{x}}, \bar{\lambda}) = 0$ , leading to  $(\bar{\mathbf{x}}, \bar{\lambda}) = (\mathbf{x}^*, \lambda^*)$ , a contradiction. Therefore, any subsequence of  $\{(\mathbf{x}_k, \lambda_k)\}_{k=0}^\infty$  converges to  $(\mathbf{x}^*, \lambda^*)$ , thereby  $\{(\mathbf{x}_k, \lambda_k)\}_{k=0}^\infty$  converges to  $(\mathbf{x}^*, \lambda^*)$ .  $\square$

**5. Experimental results.** We test the PNT method and compare it with the standard Newton method, which refers to the method in [31] but (2.9) is solved using direct matrix inversions. These two methods use the same initialization and backtracking line search strategy. The setting of hyperparameters follows Algorithm 3.2, and we set  $\tau = 1.001$  and  $\lambda_0 = 0.1$  in all the experiments. We also implement the generalized hybrid iterative method proposed in [12] (denoted by genHyb), which is also based on gen-GKB. The genHyb iteratively computes approximations to  $\mu_{opt}$

and  $\mathbf{x}_{opt} = \mathbf{x}(\mu_{opt})$ , where  $\mu_{opt}$  is the optimal Tikhonov regularization parameter, that is  $\mu_{opt} = \min_{\mu > 0} \|\mathbf{x}(\mu) - \mathbf{x}_{true}\|_2$ ; the  $k$ -th approximate Lagrangian multiplier is  $\lambda_k = 1/\mu_k$ . All the experiments are performed on MATLAB R2023b. The codes are available at <https://github.com/Machealb/InverProb.IterSolver>.

All the inverse problems in the experiments are ill-posed and satisfy [Assumption 1](#). We use three types of ill-posed inverse problems to test the proposed method. The characteristics of these problems are summarized in [Table 5.1](#).

TABLE 5.1  
*Properties of the inverse problems in the experiments.*

| Problem       | $m \times n$         | Ill-posedness | Description           |
|---------------|----------------------|---------------|-----------------------|
| heat          | $2000 \times 2000$   | moderate      | inverse heat equation |
| shaw          | $3000 \times 3000$   | severe        | 1D image restoration  |
| PRblurshake   | $128^2 \times 128^2$ | mild          | 2D image deblurring   |
| PRblurspeckle | $128^2 \times 128^2$ | mild          | 2D image deblurring   |
| PRspherical   | $23168 \times 128^2$ | mild          | computed tomography   |

**5.1. Small-scale problems.** We choose two small-scale 1D inverse problems from [\[24\]](#). The first problem is *heat*, an inverse heat equation described by the Volterra integral equation of the first kind on  $[0, 1]$ . The second problem is *shaw*, a one-dimensional image restoration model described by the Fredholm integral equation of the first kind on  $[-\pi/2, \pi/2]$ . We use the code in [\[24\]](#) to discretize the two problems to generate  $\mathbf{A}$ ,  $\mathbf{x}_{true}$  and  $\mathbf{b}_{true} = \mathbf{A}\mathbf{x}_{true}$ , where  $m = n = 2000$  and  $m = n = 3000$  for *heat* and *shaw*, respectively. We set the noisy observation vector  $\mathbf{b}$  as  $\mathbf{b} = \mathbf{b}_{true} + \boldsymbol{\epsilon}$ , where  $\boldsymbol{\epsilon}$  is a Gaussian noise. For *heat*, we set  $\boldsymbol{\epsilon}$  as a white noise (i.e.  $\mathbf{M}$  is a scalar matrix) with noise level  $\varepsilon := \|\boldsymbol{\epsilon}\|_2 / \|\mathbf{b}_{true}\|_2 = 5 \times 10^{-2}$ ; for *shaw*, we set  $\boldsymbol{\epsilon}$  as a uncorrelated non-white noise (i.e.  $\mathbf{M}$  is a diagonal matrix) with noise level  $\varepsilon = 10^{-2}$ . The true solutions and noisy observed data for these two problems are shown in [Figure 5.1](#).

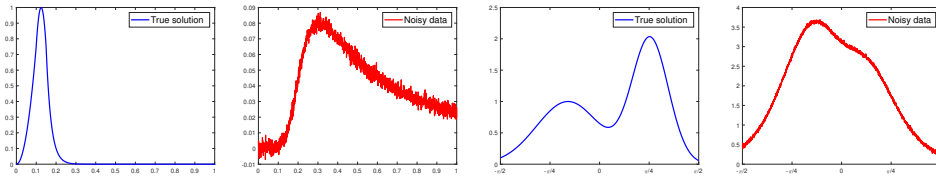


FIG. 5.1. True solution and noisy observed data. The first two: *heat*. The last two: *shaw*.

For *heat*, we assume a Gaussian prior  $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mu^{-1}\mathbf{N})$  with  $\mathbf{N}$  coming from the Gaussian kernel  $\kappa_G$ , i.e. the  $ij$  element of  $\mathbf{N}$  is  $[\mathbf{N}]_{ij} = K_G(r_{ij})$ ,  $K_G(r) := \exp(-r^2/(2l^2))$ , where  $r_{ij} = \|\mathbf{p}_i - \mathbf{p}_j\|_2$  and  $\{\mathbf{p}_i\}_{i=1}^n$  are discretized points of the domain of the true solution; the parameter  $l$  is set as  $l = 0.1$ . For *shaw*, we construct  $\mathbf{N}$  using the exponential kernel  $K_{exp}(r) := \exp(-(r/l)^\nu)$ , where the parameters  $l$  and  $\nu$  are set as  $l = 0.1$  and  $\nu = 1$ . We set  $\tau = 1.001$  for both the two problems. We factorize  $\mathbf{M}^{-1}$  and  $\mathbf{N}^{-1}$  to form (1.3) and solve it directly to find  $\mu_{opt}$  and  $\mathbf{x}_{opt}$ ; the corresponding Lagrangian multiplier is  $\lambda_{opt} = 1/\mu_{opt}$ . We also compute the  $\mu$  of (2.1) and the corresponding regularized solution, which is denoted by  $\mu_{DP}$  and  $\mathbf{x}_{DP}$ , respectively. Therefore, the solution to (2.2) and (2.3) is  $(\mathbf{x}^*, \lambda^*) = (\mathbf{x}_{DP}, 1/\mu_{DP})$ .

We use the optimal Tikhonov solution and the DP solution as the baseline for the subsequent tests.

For these two small-scale problems, we also implement the projected Newton method in [13] based on the transformation (1.3) as a comparison. This means that we solve

$$\min_{\bar{\mathbf{x}} \in \mathbb{R}^n} \{ \|(\mathbf{L}_M \mathbf{A} \mathbf{L}_N) \bar{\mathbf{x}} - \mathbf{L}_M \mathbf{b}\|_2^2 + \mu \|\bar{\mathbf{x}}\|_2^2 \},$$

using the method in [13] and then compute the regularized solution  $\mathbf{x}_k = \mathbf{L}_N^{-1} \bar{\mathbf{x}}_k$ . This Cholesky factorization based method is abbreviated as Ch-PNT.

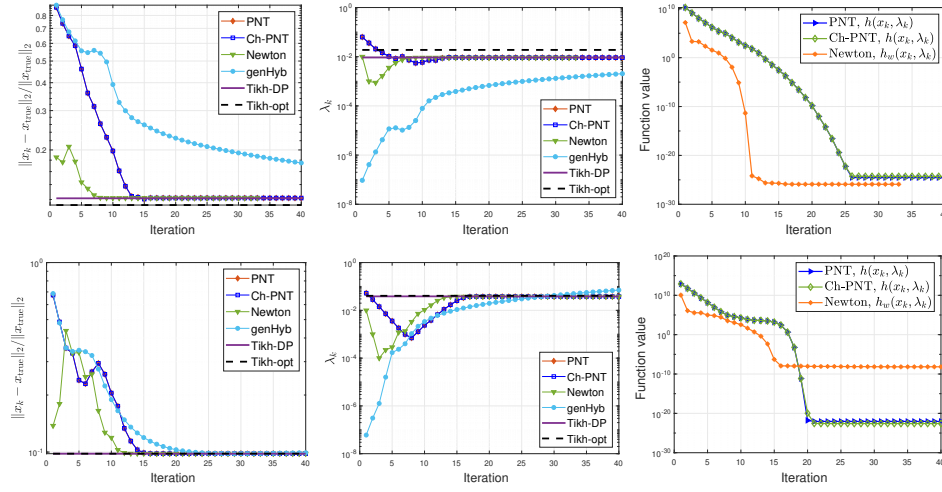


FIG. 5.2. Relative errors of iterative solutions, convergence of  $\lambda_k$ , and convergence of merit functions. Top: heat. Bottom: shaw.

We compare the convergence behaviors of PNT, Ch-PNT, Newton and genHyb methods by plotting the relative error curve of  $\mathbf{x}_k$  with respect to  $\mathbf{x}_{\text{true}}$  and the convergence curves of  $\lambda_k$  and merit functions. The solutions  $(\mathbf{x}_{DP}, 1/\mu_{DP})$  and  $(\mathbf{x}_{opt}, 1/\mu_{opt})$  are used as baselines. From Figure 5.2 we find that both PNT and Newton methods converge very fast to  $\mathbf{x}_{DP}$  and  $\lambda_{DP} := 1/\mu_{DP}$  with very few iterations, and PNT converges only slightly slower than Newton. We also find that the convergence behaviors of PNT and Ch-PNT are almost identical. This is not surprising, as both methods utilize the same subspaces for projecting the large-scale system and employ the same hyperparameters and update procedures. For heat, the error of the DP solution is slightly higher than the optimal Tikhonov solution, because DP slightly under-estimates  $\lambda$ . The merit functions of both PNT and Newton decrease monotonically, and  $h(\mathbf{x}_k, \lambda_k)$  of PNT eventually decreases to an extremely small value for the two problems. We remark that we set  $w = 1$  for  $h_w(\mathbf{x}, \lambda)$  in all the tests. For Newton method for heat, we stop the iterate at  $k = 34$  because the step-length  $\gamma_k$  is too small. In comparison, genHyb converges much slower than the previous two methods, especially for heat.

Figure 5.3 plots the recovered solutions computed by PNT and genHyb methods at the final iterations; the solution by Newton is almost the same as that by PNT, thereby we omit it. We also plot the optimal Tikhonov regularized solution as a comparison, where the DP solution is very similar and omitted. Both PNT and genHyb can recover good regularized solutions, and PNT is slightly better for heat.

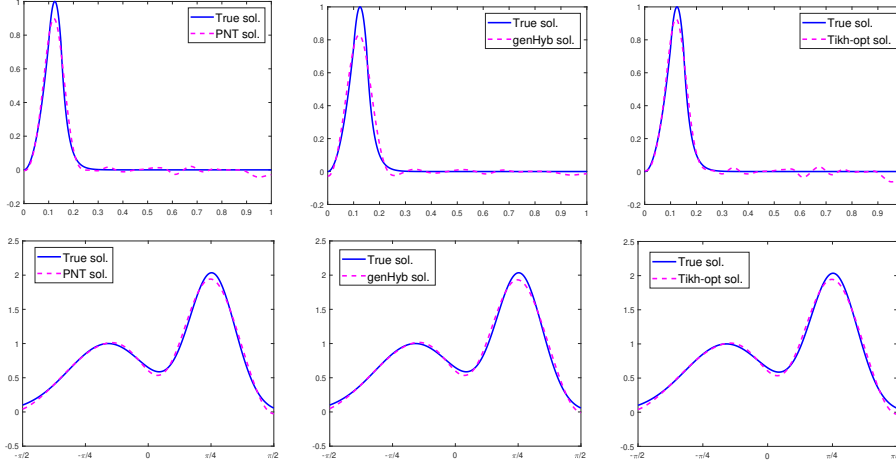


FIG. 5.3. Comparison of reconstructed solutions at the final iterations with the optimal Tikhonov regularized solution. Top: heat. Bottom: shaw.

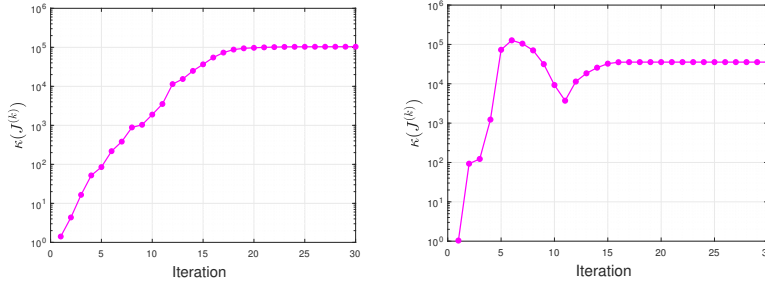


FIG. 5.4. Variation of the condition number of  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  during the iteration of PNT. Left: heat. Right: shaw.

To further demonstrate the performance of PNT, we present the variation of the condition number of  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  during the iteration of PNT in Figure 5.4. This condition number is denoted by  $\kappa(J^{(k)})$  in the two pictures. We observe that the condition number does not increase significantly during the iteration. This ensures that the small-scale linear system (3.12a) can be solved directly via matrix inversions without any issues.

To show the advantage of the computational efficiency of PNT over Ch-PNT and Newton, we gradually increase the scale of the test problems and measure the running time of the three methods, where all of them stop at the first iteration such that  $\|\mathbf{Ax}_k - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \tau m \leq 10^{-8}$ . The time data are listed in Table 5.2. We also compute the ratio of the running time, i.e. the value of Ch-PNT-time/PNT-time and Newton-time/PNT-time. For shaw, we find that all three methods stop with similar iteration numbers, and the computational speed of PNT is much faster than Newton, with the speedup ratio varying from 41 to 157. For heat, we find that Newton stops with only about half iteration numbers of PNT's. However, the total running time of PNT is still much smaller than Newton's, with the speedup ratio varying from 8 to 43. To compare the scalability of PNT, Ch-PNT and Newton more clearly, we use the

TABLE 5.2

Running time (measured in seconds) of *PNT*, *Ch-PNT* and *Newton* methods as the scale of the problems increasing from  $n = 1000$  to  $n = 5000$ . Both the two methods stop at the first  $k$  (in parentheses) such that  $\|\mathbf{Ax}_k - \mathbf{b}\|_{\mathbf{M}^{-1}}^2 - \tau m \leq 10^{-8}$ . The ratio of the running time between *PNT* and *Ch-PNT* is denoted as **ratio-1**, while the ratio of the running time between *PNT* and *Newton* is denoted as **ratio-2**.

| $n$            | 1000        | 2000        | 3000        | 4000        | 5000         |
|----------------|-------------|-------------|-------------|-------------|--------------|
| heat           |             |             |             |             |              |
| PNT            | 0.021 (18)  | 0.114 (21)  | 0.164 (19)  | 0.347 (19)  | 0.492 (19)   |
| Ch-PNT         | 0.032 (18)  | 0.280 (21)  | 0.375 (19)  | 0.764 (19)  | 1.314 (19)   |
| Newton         | 0.249 (10)  | 2.568 (11)  | 6.062 (10)  | 13.914 (10) | 26.127 (11)  |
| <b>ratio-1</b> | <b>1.5</b>  | <b>2.5</b>  | <b>2.3</b>  | <b>2.2</b>  | <b>2.7</b>   |
| <b>ratio-2</b> | <b>11.9</b> | <b>22.5</b> | <b>37.0</b> | <b>40.1</b> | <b>53.1</b>  |
| shaw           |             |             |             |             |              |
| PNT            | 0.014 (17)  | 0.051 (16)  | 0.158 (17)  | 0.293 (18)  | 0.479 (19)   |
| Ch-PNT         | 0.051 (17)  | 0.170 (16)  | 0.438 (19)  | 0.873 (18)  | 1.819 (19)   |
| Newton         | 0.455 (16)  | 3.675 (15)  | 8.904 (14)  | 26.835 (16) | 52.768 (16)  |
| <b>ratio-1</b> | <b>3.6</b>  | <b>3.3</b>  | <b>2.8</b>  | <b>3.0</b>  | <b>3.8</b>   |
| <b>ratio-2</b> | <b>32.5</b> | <b>72.1</b> | <b>56.4</b> | <b>91.6</b> | <b>110.2</b> |

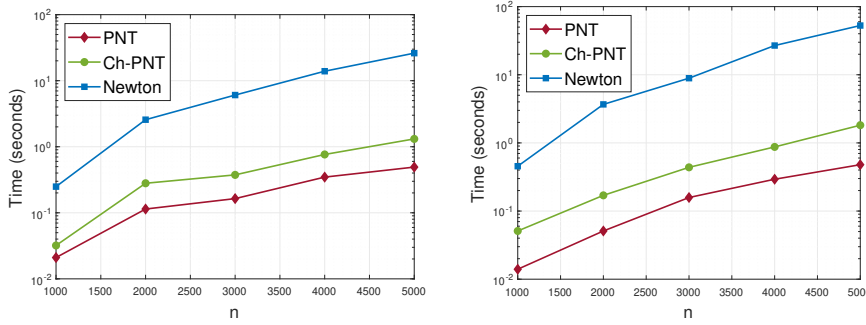


FIG. 5.5. Comparison of scalability of *PNT*, *Ch-PNT* and *Newton* methods as the scale of the problems increasing from  $n = 1000$  to  $n = 5000$ . Left: heat. Right: shaw.

data in Table 5.2 to plot the curve of time growth with respect to  $n$ . Clearly, PNT saves much more time compared to Newton while obtaining solutions with the same accuracy. Although the advantage of PNT over Ch-PNT is not significant for small-scale problems, Ch-PNT is not feasible for large-scale problems due to the prohibitive cost of Cholesky factorization.

**5.2. Large-scale problems.** We choose three 2D image deblurring and computed tomography inverse problems from [18]. The first problem is PRblurshake, which simulates a spatially invariant motion blur caused by the shaking of a camera. The second problem is PRblurspeckle, which simulates a spatially invariant blur caused by atmospheric turbulence. The third problem is PRspherical that models spherical means tomography. The true images and noisy observed data are shown in Figure 5.6, where all the images have  $128 \times 128$  pixels, and  $\epsilon$  are uncorrelated non-white Gaussian

noises with  $\varepsilon = 10^{-3}$ ,  $5 \times 10^{-3}$  and  $10^{-2}$ , respectively. We have  $m = n = 128^2$  for the first two problems, and  $m = 23168, n = 128^2$  for the third problem.

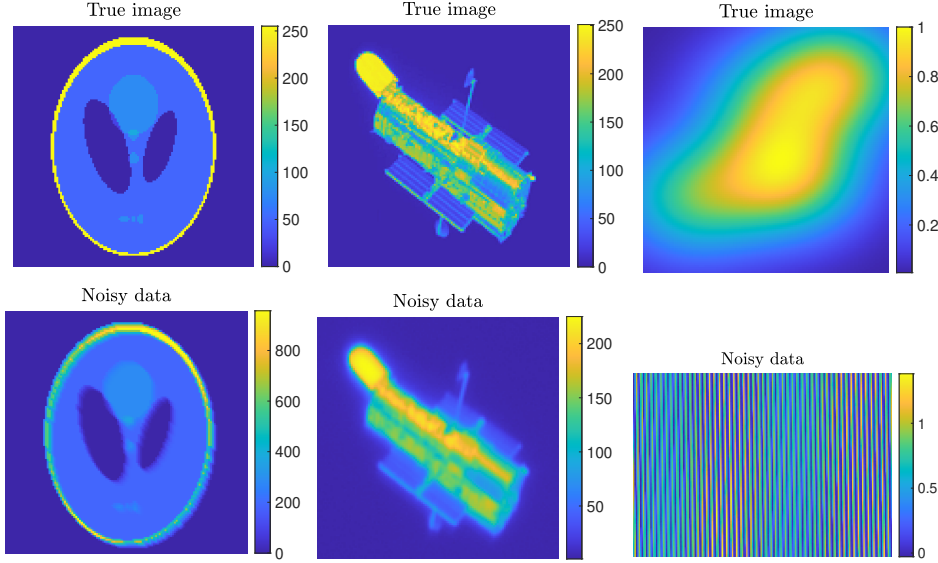


FIG. 5.6. True solution and noisy observed data for deblurring and tomography problems. From the leftmost column to the rightmost column are PRblurshake, PRblurspeckle and PRspherical.

For PRblurshake and PRblurspeckle, we construct  $\mathbf{N}$  using the Gaussian kernel with  $l = 10$  and  $l = 1$ , respectively. For PRspherical, we construct  $\mathbf{N}$  using the Matérn kernel

$$K_M(r) := \frac{2^{1-\nu}}{\Gamma(\nu)} \left( \frac{\sqrt{2\nu}r}{l} \right)^\nu B_\nu \left( \frac{\sqrt{2\nu}r}{l} \right),$$

where  $\Gamma(\cdot)$  is the gamma function,  $B_\nu(\cdot)$  is the modified Bessel function of the second kind, and  $l$  and  $\nu$  are two positive parameters of the covariance; we set  $l = 100$  and  $\nu = 1.5$ . For the three large-scale problems, it is almost impossible to get  $(\mu_{opt}, \mathbf{x}(\mu_{opt}))$  and  $(\mu_{DP}, \mathbf{x}(\mu_{DP}))$  by solving (1.3). The standard **Newton** method and the methods in [13, 14] can not be applied because these methods have to deal with  $\mathbf{N}^{-1}$ . To test the performance of PNT, here we only compare it with **genHyb**. Additionally, we also implement PNT-md to demonstrate that it can save some computation compared to PNT.

The relative error curves of the three methods, the convergence curves of  $\lambda_k$  and  $h(\mathbf{x}_k, \lambda_k)$  are plotted in Figure 5.7. For PNT-md, we set  $k_0 = 150, 80, 50$  for the three problems, respectively. It can be observed that PNT for the last two problems converges very fast: the variations of relative error and  $\lambda_k$  become quickly stabilized after 50 to 150 iterations, although for the second problem  $h(\mathbf{x}_k, \lambda_k)$  are still decreasing significantly after 200 iterations. The **genHyb** method for PRblurshake and PRspherical converges slower, and it obtains two solutions with larger relative errors than that of PNT. This is because **genHyb** under-estimates  $\lambda$  more than PNT. For all three problems, PNT-md converges very quickly from  $k_0$ , achieving solutions with the same accuracy as PNT while requiring nearly the same total number of iterations. The reconstructed images are shown in Figure 5.8, which reveals the effectiveness of PNT and **genHyb**. The variation of the condition number of  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  is shown

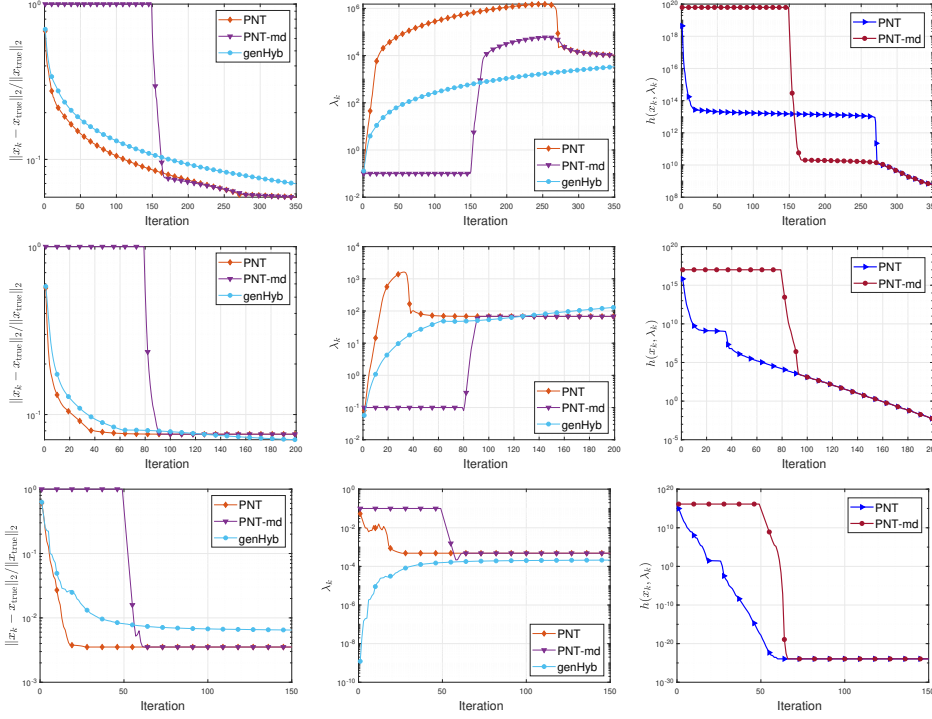


FIG. 5.7. Relative errors of iterative solutions, convergence of  $\lambda_k$ , and convergence of merit functions. Top: PRblurshake. Middle: PRblurspeckle. Bottom: PRspherical.

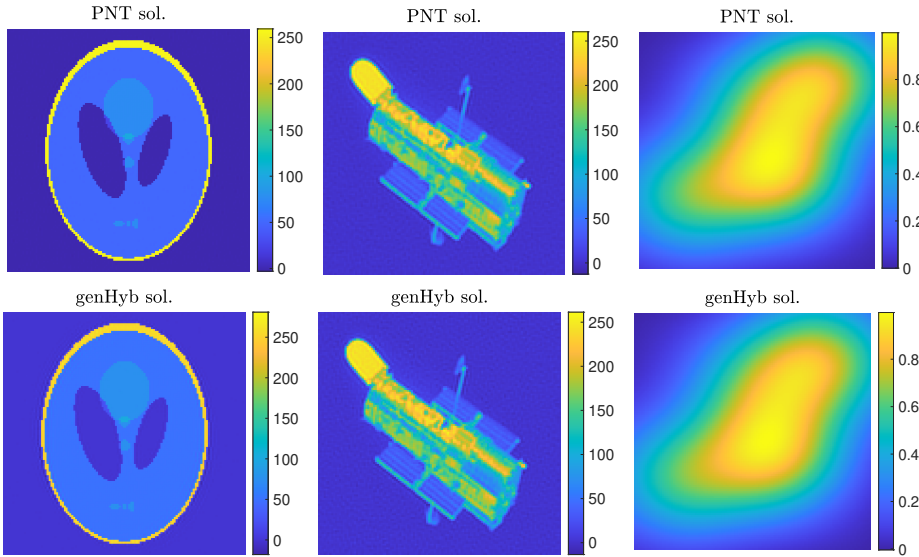


FIG. 5.8. Reconstructed solutions at the final iterations by *PNT* and *genHyb*. From the leftmost column to the rightmost column are PRblurshake, PRblurspeckle and PRspherical.

in Figure 5.9. The condition number does not grow very large during the iteration, allowing the small-scale linear system (3.12a) to be solved directly without issues.

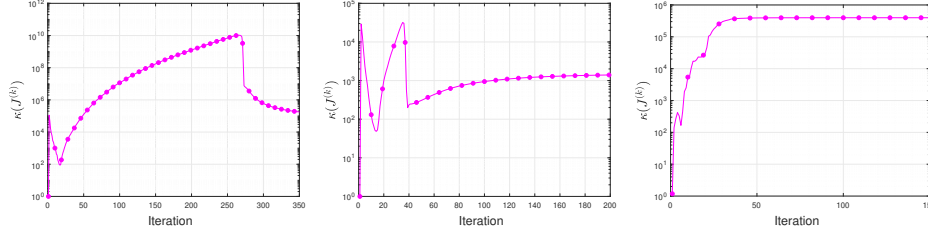


FIG. 5.9. Variation of the condition number of  $J^{(k)}(\bar{\mathbf{y}}_{k-1}, \lambda_{k-1})$  during the iteration of PNT. Left: PRblursake. Middle: PRblurspeckle Right: PRspherical.

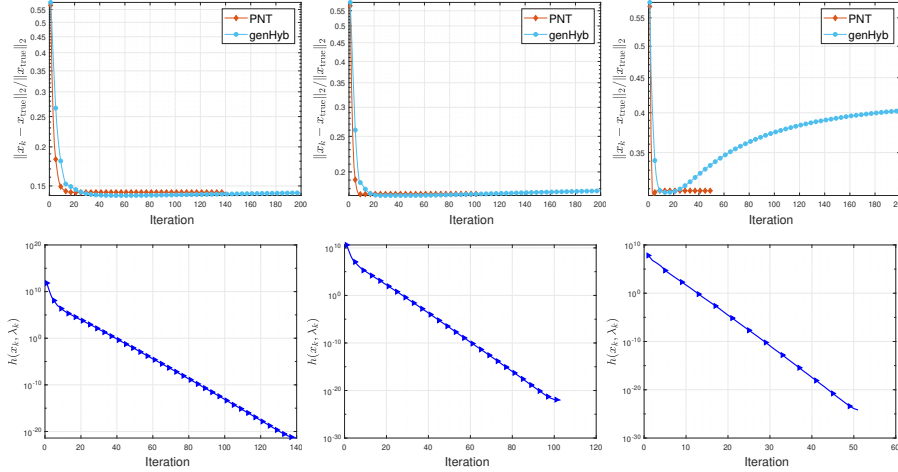


FIG. 5.10. Relative errors of iterative solutions by PNT and *genHyb*, and the decrease of  $h(\mathbf{x}_k, \lambda_k)$ . The test problem is PRblurspeckle. From the leftmost column to the rightmost column, the noise levels are  $\varepsilon = 5 \times 10^{-2}$ ,  $10^{-1}$ ,  $5 \times 10^{-1}$ .

To further test the robustness of PNT and *genHyb* as the noise level gradually increases, we set the noise level of PRblurspeckle to be  $\varepsilon = 5 \times 10^{-2}$ ,  $10^{-1}$ ,  $5 \times 10^{-1}$ . Figure 5.10 shows the corresponding relative error curves and the curves of  $h(\mathbf{x}_k, \lambda_k)$ . We can find that, when the noise is not very big, both PNT and *genHyb* converge stably with almost the same accuracy. However, when the noise gradually increases, the situations are very different. First, we find that as the noise increases, PNT still converges stably, and faster. Second,  $h(\mathbf{x}_k, \lambda_k)$  can always decrease to an extremely small value, which is promised by Theorem 4.1. The iterate of PNT stops when the step-length  $\gamma_k$  becomes too small (less than  $10^{-16}$ ), which happens more early if the noise is bigger. In comparison, the convergence of *genHyb* becomes unstable as the noise increases. For  $\varepsilon = 10^{-1}$ , it can be observed that the relative error for *genHyb* increases slightly after a certain iteration, while for  $\varepsilon = 5 \times 10^{-1}$ , the increase of relative error appears earlier and more clearly. This is a typical potential weakness of hybrid regularization methods, a challenge that the PNT method successfully addresses.

**6. Conclusion.** For large-scale Bayesian linear inverse problems, we have proposed the projected Newton (PNT) method as a novel iterative approach for simultaneously updating both the regularization parameter and solution without any computationally expensive matrix inversions or decompositions. By reformulating the Tikhonov regularization as a corresponding constrained minimization problem and leveraging its Lagrangian function, the regularized solution and the corresponding Lagrangian multiplier can be obtained from the unconstrained Lagrangian function using a Newton-type method. To reduce the computational overhead of the Newton method, the generalized Golub-Kahan bidiagonalization is applied to project the original large-scale problem to become small-scale ones, where the projected Newton direction is obtained by solving the small-scale linear system at each iteration. We have proved that the projected Newton direction is a descent direction of a merit function, and the points generated by PNT eventually converge to the unique minimizer of this merit function, which is just the regularized solution and the corresponding Lagrangian multiplier.

Experimental tests on both small and large-scale Bayesian inverse problems have demonstrated the excellent convergence property, robustness and efficiency of PNT. The most demanding computational tasks in PNT are primarily matrix-vector products, making it particularly well-suited for large-scale problems.

An important remaining question is the convergence rate of PNT, i.e., how fast the three quantities  $\|\mathbf{x}_k - \mathbf{x}_{DP}\|_2$ ,  $|\lambda_k - \lambda_{DP}|$  and  $h(\mathbf{x}_k, \lambda_k)$  converge to zero? The convergence rate may depend on several factors, including the ill-posedness of (1.1), the smoothness of the true solution, the noise level, and the properties of  $\{\mathbf{A}, \mathbf{M}, \mathbf{N}\}$ . We will conduct theoretical investigations into this issue in future work.

**Acknowledgments.** The author thanks Dr. Felipe Atenas for helpful discussions.

#### REFERENCES

- [1] M. V. AFONSO, J. M. BIOUCAS-DIAS, AND M. A. FIGUEIREDO, *An augmented Lagrangian approach to the constrained optimization formulation of imaging inverse problems*, IEEE Trans. Image Process., 20 (2010), pp. 681–695.
- [2] F. S. V. BAZÁN, *Fixed-point iterations in determining the Tikhonov regularization parameter*, Inverse Probl., 24 (2008), p. 035001.
- [3] A. BECK AND M. TEOULLE, *A fast iterative shrinkage-thresholding algorithm for linear inverse problems*, SIAM J. Imaging Sci., 2 (2009), pp. 183–202.
- [4] J. BLOWEY, A. CRAIG, AND T. SHARDLOW, *Frontiers in numerical analysis: Durham 2002*, Springer Science & Business Media, 2003.
- [5] P. N. BROWN AND Y. SAAD, *Convergence theory of nonlinear Newton–Krylov algorithms*, SIAM J. Optim., 4 (1994), pp. 297–330.
- [6] T. BUZUG, *Computed Tomography*, Springer, 2008, <https://doi.org/10.1007/978-3-540-39408-2>.
- [7] D. CALVETTI, F. PITOLLI, J. PREZIOSO, E. SOMERSALO, AND B. VANTAGGI, *Priorconditioned CGLS-based quasi-MAP estimate, statistical stopping rule, and ranking of priors*, SIAM J. Sci. Comput., 39 (2017), pp. S477–S500.
- [8] D. CALVETTI, F. PITOLLI, E. SOMERSALO, AND B. VANTAGGI, *Bayes meets Krylov: Statistically inspired preconditioners for CGLS*, SIAM Rev., 60 (2018), pp. 429–461.
- [9] D. CALVETTI AND E. SOMERSALO, *Priorconditioners for linear systems*, Inverse Probl., 21 (2005), p. 1397.
- [10] N. A. CARUSO AND P. NOVATI, *Convergence analysis of LSQR for compact operator equations*, Linear Algebra Appl., 583 (2019), pp. 146–164.
- [11] J. CHUNG, J. G. NAGY, AND D. P. O’LEARY, *A weighted-GCV method for Lanczos-hybrid regularization*, Electr. Trans. Numer. Anal., 28 (2008), pp. 149–167.
- [12] J. CHUNG AND A. K. SAIBABA, *Generalized hybrid iterative methods for large-scale Bayesian inverse problems*, SIAM J. Sci. Comput., 39 (2017), pp. S24–S46.

- [13] J. CORNELIS, N. SCHENKELS, AND W. VANROOSE, *Projected Newton method for noise constrained Tikhonov regularization*, Inverse Probl., 36 (2020), p. 055002.
- [14] J. CORNELIS AND W. VANROOSE, *Projected Newton method for noise constrained  $\ell_p$  regularization*, Inverse Probl., 36 (2020), p. 125004.
- [15] I. DAUBECHIES, M. DEFRISE, AND C. DE MOL, *An iterative thresholding algorithm for linear inverse problems with a sparsity constraint*, Commun. Pure Appl. Math., 57 (2004), pp. 1413–1457.
- [16] M. FUKUSHIMA, Z.-Q. LUO, AND P. TSENG, *A sequential quadratically constrained quadratic programming method for differentiable convex minimization*, SIAM Journal on Optimization, 13 (2003), pp. 1098–1119.
- [17] A. GARBUNO-INIGO, F. HOFFMANN, W. LI, AND A. M. STUART, *Interacting Langevin diffusions: Gradient structure and ensemble Kalman sampler*, SIAM J. Appl. Dyn. Syst., 19 (2020), pp. 412–441.
- [18] S. GAZZOLA, P. C. HANSEN, AND J. G. NAGY, *IR Tools: A MATLAB package of iterative regularization methods and large-scale test problems*, Numer. Algor., 81 (2019), pp. 773–811.
- [19] S. GAZZOLA, P. NOVATI, AND M. R. RUSSO, *On Krylov projection methods and Tikhonov regularization*, Electr. Trans. Numer. Anal., 44 (2015), pp. 83–123.
- [20] S. GAZZOLA AND M. SABATÉ LANDMAN, *Krylov methods for inverse problems: Surveying classical, and introducing new, algorithmic approaches*, GAMM-Mitteilungen, 43 (2020), p. e202000017.
- [21] T. GOLDSTEIN AND S. OSHER, *The split Bregman method for  $l_1$ -regularized problems*, SIAM J. Imaging Sci., 2 (2009), pp. 323–343.
- [22] G. H. GOLUB AND H. G. WAHBA, *Generalized Cross-Validation as a method for choosing a good ridge parameter*, Technometrics, 21 (1979), pp. 215–223.
- [23] P. C. HANSEN, *Analysis of discrete ill-posed problems by means of the L-curve*, SIAM Rev., 34 (1992), pp. 561–580.
- [24] P. C. HANSEN, *Regularization Tools version 4.0 for Matlab 7.3*, Numer. Algor., 46 (2007), pp. 189–194.
- [25] P. C. HANSEN, J. G. NAGY, AND D. P. O’LEARY, *Deblurring Images: Matrices, Spectra and Filtering*, SIAM, Philadelphia, 2006.
- [26] G. HUANG, L. REICHEL, AND F. YIN, *On the choice of subspace for large-scale Tikhonov regularization problems in general form*, Numer. Algor., 81 (2019), pp. 33–55.
- [27] Z. JIA AND H. LI, *The joint bidiagonalization method for large GSVD computations in finite precision*, SIAM J. Matrix Anal. Appl., 44 (2023), pp. 382–407.
- [28] J. KAIPIO AND E. SOMERSALO, *Statistical and Computational Inverse Problems*, Springer, 2006.
- [29] M. E. KILMER, P. C. HANSEN, AND M. I. ESPANOL, *A projection-based approach to general-form Tikhonov regularization*, SIAM J. Sci. Comput., 29 (2007), pp. 315–330.
- [30] K. Z. KODY LAW, ANDREW STUART, *Data Assimilation: A Mathematical Introduction*, Springer, 2015, <https://doi.org/10.1007/978-3-319-20325-6>.
- [31] G. LANDI, *The Lagrange method for the regularization of discrete ill-posed problems*, Comput. Optim. Appl., 39 (2008), pp. 347–368.
- [32] H. LI, *Subspace projection regularization for large-scale Bayesian linear inverse problems*, arXiv preprint, arXiv:2310.18618, (2023).
- [33] H. LI, *The joint bidiagonalization of a matrix pair with inaccurate inner iterations*, SIAM J. Matrix Anal. Appl., 45 (2024), pp. 232–259, <https://doi.org/10.1137/22M1541083>.
- [34] H. LI, *A preconditioned Krylov subspace method for linear inverse problems with general-form Tikhonov regularization*, SIAM Journal on Scientific Computing, 46 (2024), pp. A2607–A2633, <https://doi.org/10.1137/23M1593802>.
- [35] J. LIESEN AND Z. STRAKOS, *Krylov subspace methods: principles and analysis*, Numerical Mathematics and Scie, 2013.
- [36] N. LUIKEN AND T. VAN LEEUWEN, *Relaxed regularization for linear inverse problems*, SIAM J. Sci. Comput., 43 (2021), pp. S269–S292.
- [37] J. L. MEAD AND R. A. RENAUT, *A Newton root-finding algorithm for estimating the regularization parameter for solving ill-conditioned least squares problems*, Inverse Probl., 25 (2008), p. 025002.
- [38] F. MESSERER, K. BAUMGÄRTNER, AND M. DIEHL, *Survey of sequential convex programming and generalized Gauss-Newton methods*, ESAIM: Proceedings and Surveys, 71 (2021), pp. 64–88.
- [39] S. MORIGI, L. REICHEL, AND F. SGALLARI, *Orthogonal projection regularization operators*, Numer. Algor., 44 (2007), pp. 99–114.
- [40] V. A. MOROZOV, *Regularization of incorrectly posed problems and the choice of regulariza-*

- tion parameter, USSR Computational Mathematics and Mathematical Physics, 6 (1966), pp. 242–251.
- [41] J. NOCEDAL AND S. WRIGHT, *Numerical Optimization*, Springer Series in Operations Research and Financial Engineering, Springer New York, 2006.
  - [42] D. P. O’LEARY AND J. A. SIMMONS, *A bidiagonalization-regularization procedure for large scale discretizations of ill-posed problems*, SIAM J. Sci. Statist. Comput., 2 (1981), pp. 474–489.
  - [43] S. OSHER, M. BURGER, D. GOLDFARB, J. XU, AND W. YIN, *An iterative regularization method for total variation-based image restoration*, Multiscale Model. Simul., 4 (2005), pp. 460–489.
  - [44] C. C. PAIGE AND M. A. SAUNDERS, *Solution of sparse indefinite systems of linear equations*, SIAM J. Numer. Anal., 12 (1975), pp. 617–629.
  - [45] L. REICHEL, F. SGALLARI, AND Q. YE, *Tikhonov regularization based on generalized Krylov subspace methods*, Appl. Numer. Math., 62 (2012), pp. 1215–1228.
  - [46] L. REICHEL AND A. SHYSHKOV, *A new zero-finder for Tikhonov regularization*, BIT Numer. Math., 48 (2008), pp. 627–643.
  - [47] R. A. RENAUT, A. W. HELMSTETTER, AND S. VATANKHAH, *Unbiased predictive risk estimation of the tikhonov regularization parameter: convergence with increasing rank approximations of the singular value decomposition*, BIT Numer. Math., 59 (2019), pp. 1031–1061.
  - [48] R. A. RENAUT, S. VATANKHAH, AND V. E. ARDESTA, *Hybrid and iteratively reweighted regularization by unbiased predictive risk and weighted GCV for projected systems*, SIAM J. Sci. Comput., 39 (2017), pp. B221–B243.
  - [49] M. RICHTER, *Inverse Problems: Basics, Theory and Applications in Geophysics*, Springer, 2016, [https://doi.org/10.1007/978-3-319-48384-9\\_3](https://doi.org/10.1007/978-3-319-48384-9_3).
  - [50] P. RODRIGUEZ AND B. WOHLBERG, *An efficient algorithm for sparse representations with  $\ell_p$  data fidelity term*, in Proceedings of 4th IEEE Andean Technical Conference (ANDESCON), 2008.
  - [51] P. RODRÍGUEZ AND B. WOHLBERG, *Efficient minimization method for a generalized total variation functional*, IEEE Transactions on Image Processing, 18 (2008), pp. 322–332.
  - [52] A. K. SAIBABA, J. CHUNG, AND K. PETROSKE, *Efficient Krylov subspace methods for uncertainty quantification in large Bayesian linear inverse problems*, Numer. Linear Algebra Appl., 27 (2020), p. e2325.
  - [53] A. M. STUART, *Inverse problems: a Bayesian perspective*, Acta Numer., 19 (2010), pp. 451–559.
  - [54] W. TIAN AND X. YUAN, *Linearized primal-dual methods for linear inverse problems with total variation regularization and finite element discretization*, Inverse Probl., 32 (2016), p. 115011.
  - [55] C. F. VAN LOAN, *Generalizing the singular value decomposition*, SIAM J. Numer. Anal., 13 (1976), pp. 76–83.
  - [56] W. YIN, S. OSHER, D. GOLDFARB, AND J. DARBON, *Bregman iterative algorithms for  $\ell_1$ -minimization with applications to compressed sensing*, SIAM J. Imaging Sci., 1 (2008), pp. 143–168.