



Research article

A minimal residual method for large-scale trust-region subproblem

Yusen Zhang and Bo Feng*

School of Mathematics, Yangzhou University, Yangzhou 225000, China

* **Correspondence:** Email: bofeng@yzu.edu.cn.

Abstract: This paper was devoted to solving the large-scale trust-region subproblem (TRS), a core subproblem in nonlinear optimization with extensive applications in scientific computing fields such as the Lorentz eigenvalue problem, Tikhonov regularization, nonlinear least squares, and so on. The generalized Lanczos trust-region (GLTR) method is one of the most popular iterative algorithms for large-scale TRS. In this method, the approximate solution converged more slowly than the approximate Lagrange multiplier. Inspired by this characteristic and the generalized minimal residual method (GMRES) and minimal residual method (MINRES), we proposed a minimal residual method based on the GLTR framework (MRTR). The MRTR method fixed the approximate Lagrange multiplier from GLTR and sought the approximate solution that minimized the residual norm within the Krylov subspace. We established some theoretical results to characterize the relationship between the approximate solutions of MRTR and GLTR. The experimental results showed that the MRTR method outperformed the GLTR method in terms of iteration numbers and CPU running time, and its residual norm convergence curve was smoother with better numerical stability, while achieving approximate optimal values with the same high precision as GLTR.

Keywords: trust-region subproblem; generalized Lanczos trust-region (GLTR) method; Krylov subspace; minimal residual method

1. Introduction

In this paper, we focus on the following large-scale *trust-region subproblem* (TRS) [1, 2]:

$$\min_{\|\mathbf{x}\|_2 \leq \Delta} \left\{ f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T A \mathbf{x} + \mathbf{x}^T \mathbf{g} \right\}, \quad (1.1)$$

where $A \in \mathbb{R}^{n \times n}$ is a symmetric matrix, $\mathbf{g} \in \mathbb{R}^n$ is a nonzero vector, $\|\cdot\|_2$ denotes the 2-norm of a vector (or matrix), and $\Delta > 0$ is the trust-region radius. Trust-region methods [1, 2] are a widely used technique for solving nonlinear optimization problems that aim to minimize objective function. Each

iteration of such methods entails finding a solution to the TRS (1.1). Within the trust-region framework, the objective function of TRS (1.1) serves as a quadratic approximation of objective function around the current iterate $\tilde{x}$, where A denotes the Hessian matrix and g stands for the gradient of objective function evaluated at $\tilde{x}$. In addition, the TRS (1.1) finds wide applications in various fields, including the Lorentz eigenvalue problem [3], graph partitioning problems [4], the Levenberg–Marquardt algorithm for nonlinear least squares [2], and the Tikhonov regularization [5, 6].

Numerous algorithms have been developed to address large-scale TRS (1.1). Gould et al. proposed a modified version of the Moré–Sorensen method using Taylor series expansion [7], while Steihaug and Toint [8, 9] introduced a truncated conjugate gradient (tCG) approach. In [10], Gander et al. addressed the TRS (1.1) from the perspective of an eigenproblem. Building on this approach, Adachi, Iwata, Nakatsukasa, and Takeda extended the strategy to weighted-norm TRS in [11], where the problem can be transformed into a generalized eigenvalue problem of size $2n$. The generalized Lanczos trust-region (GLTR) method [12] leverages projection techniques to reduce the problem size. Other notable methods include a nested Lanczos method [13] (regarded as an accelerated variant of GLTR), a matrix-free algorithm [5, 6], and a method based on semi-definite programming (SDP) [14]. For more details, the reader is referred to [15–20].

The GLTR method [12] is among the most widely used techniques for solving large-scale TRS (1.1) in the *easy case* (for the definition of the *easy case*, please see Section 2). The principle of GLTR consists of the Lanczos process combined with the Rayleigh–Ritz procedure. Analogous to the classical Lanczos method for eigenvalue problems (see, e.g., [21]), for solving large-scale TRS, the original problem (1.1) is first projected onto the Krylov subspace to yield a reduced-size TRS, which is subsequently solved by an efficient and accurate method. The convergence of GLTR has been investigated in [22–27]. From these convergence results, the convergence rate of the approximate Lagrange multiplier is significantly higher than that of the approximate solution in the GLTR method. Inspired by this fact and the generalized minimal residual (GMRES)/minimal residual (MINRES) methods [28], we propose a minimal residual method for solving large-scale TRS problems based on the GLTR method. This method fixes the approximate Lagrange multiplier obtained from the GLTR method and seeks the approximate solution that minimizes the residual norm within the Krylov subspace, replacing the approximate solution yielded by the GLTR to approximate the global solution. We show through several numerical examples that the proposed method yields a certain improvement over the GLTR method. Furthermore, it can be seen from the convergence curves of the residual norm that the convergence curve of the new method is smoother than that of the GLTR method (see Section 4).

The paper is organized as follows. In Section 2, we first review the theoretical characterization of the global solution to the TRS (1.1) and classify the *easy case* and *hard case* of the problem, then give a detailed introduction to the GLTR method, including its basic principle, two-phase iteration process, and existing convergence results. In Section 3, inspired by the GMRES/MINRES methods and the faster convergence rate of the Lagrange multiplier in the GLTR method, we propose a minimal residual method for large-scale TRS (MRTR), derive the explicit form of the approximate solution of the new method, and establish key theoretical results to analyze the relationship between the approximate solutions of MRTR and GLTR. Section 4 presents some numerical experiments, where we test the performance of the MRTR method on both constructed diagonal matrices and real sparse symmetric matrices from the University of Florida Sparse Matrix Collection, and compare the number of iterations, CPU running time, residual norm convergence curves, and approximate optimal values with those of

the GLTR method to verify the numerical efficiency and stability of the proposed method. Finally, concluding remarks are drawn in Section 5.

Throughout this paper, the transpose of a matrix or vector is denoted by $(\cdot)^T$; $\mathcal{N}(A)$ stands for the null space of matrix A ; $\|\cdot\|$ represents the Euclidean norm of a matrix or vector. We use $\mathbf{0}$ and I to denote the zero vector (matrix) and identity matrix, respectively, with their dimensions being self-evident from the context. Let $\mathbf{e}_i$ denote the i -th column of I . Additionally, $A \succcurlyeq \mathbf{0}$ (resp., $A \succ \mathbf{0}$) indicates that A is a symmetric positive semi-definite (resp., positive definite) matrix.

2. Preliminaries

In this section, we briefly introduce the solutions to the TRS (1.1), and the GLTR method [12].

2.1. Solutions of the TRS (1.1)

A global solution to the TRS (1.1) is characterized by the following theorem.

Theorem 2.1. [1, 17, 18] *The vector $\mathbf{x}_{\text{opt}}$ is a global optimal solution of the trust-region problem (1.1), if, and only if, $\|\mathbf{x}_{\text{opt}}\| \leq \Delta$ and there exists Lagrange multiplier $\lambda_{\text{opt}} \geq 0$ such that*

$$(A + \lambda_{\text{opt}}I)\mathbf{x}_{\text{opt}} = -\mathbf{g}, \quad \lambda_{\text{opt}}(\Delta - \|\mathbf{x}_{\text{opt}}\|) = 0 \quad \text{and} \quad A + \lambda_{\text{opt}}I \succcurlyeq \mathbf{0}. \quad (2.1)$$

Let $s_{\min}$ be the smallest eigenvalue of A . In particular, the TRS (1.1) admits two scenarios [2, 17, 29]:

Easy case: $\lambda_{\text{opt}} \geq \max\{0, -s_{\min}\}$. In this case, $A + \lambda_{\text{opt}}I \succ \mathbf{0}$, and the TRS (1.1) has a *unique* solution given by:

$$\mathbf{x}_{\text{opt}} = -(A + \lambda_{\text{opt}}I)^{-1}\mathbf{g}.$$

Hard case: $\lambda_{\text{opt}} = -s_{\min}$. In this case, there exist multiple solutions that can be expressed as follows:

$$\mathbf{x}_{\text{opt}} = -(A + \lambda_{\text{opt}}I)^{\dagger}\mathbf{g} + \sqrt{\Delta^2 - \|(A + \lambda_{\text{opt}}I)^{\dagger}\mathbf{g}\|^2} \cdot \mathbf{u},$$

where the superscript $\dagger$ denotes the Moore-Penrose pseudoinverse of a matrix and $\mathbf{u} \in \mathcal{N}(A - s_{\min}I)$ with $\|\mathbf{u}\|_2 = 1$.

It is evident that when $A \succ \mathbf{0}$, the global solution $\mathbf{x}_{\text{opt}}$ must take one of two forms: either $\mathbf{x}_{\text{opt}} = -A^{-1}\mathbf{g}$ (corresponding to $\lambda_{\text{opt}} = 0$) or $\mathbf{x}_{\text{opt}} = -(A + \lambda_{\text{opt}}I)^{-1}\mathbf{g}$ (corresponding to $\lambda_{\text{opt}} > 0$) lying on the boundary. Consequently, the *hard case* arises if, and only if, $s_{\min} \leq 0$. Throughout this paper, we assume that the TRS is the *easy case*.

2.2. The GLTR method

The GLTR method proposed by Gould et al. [12] is one of the most commonly used approaches for solving large-scale TRS (1.1) in the *easy case*. Specifically, the GLTR method yields an approximate solution to the original TRS by solving the following problem:

$$\min_{\|\mathbf{x}\| \leq \Delta, \mathbf{x} \in \mathcal{K}_k} \left\{ f(\mathbf{x}) = \frac{1}{2}\mathbf{x}^T A \mathbf{x} + \mathbf{x}^T \mathbf{g} \right\}, \quad (2.2)$$

where $\mathcal{K}_k$ denotes the Krylov subspace

$$\mathcal{K}_k := \mathcal{K}_k(A, \mathbf{g}) = \text{span}\{\mathbf{g}, A\mathbf{g}, \dots, A^{k-1}\mathbf{g}\} = \left\{ \sum_{i=0}^{k-1} \alpha_i A^i \mathbf{g} \mid \alpha_0, \alpha_1, \dots, \alpha_{k-1} \in \mathbb{R} \right\}.$$

We can obtain an orthonormal basis $Q_k = (\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k)$ for $\mathcal{K}_k$ by the Lanczos process, which can be written in matrix form [28]

$$AQ_k = Q_k T_k + \beta_k \mathbf{q}_{k+1} \mathbf{e}_k^T \quad \text{with } \mathbf{g} = \beta_0 \mathbf{q}_1, \text{ and } \beta_0 = \|\mathbf{g}\|, \quad (2.3)$$

where $\mathbf{e}_i$ is the i -th column of the identity matrix, and

$$T_k = Q_k^T A Q_k = \begin{pmatrix} \alpha_1 & \beta_1 & & & \\ \beta_1 & \alpha_2 & \beta_2 & & \\ & \ddots & \ddots & \ddots & \\ & & \beta_{k-2} & \alpha_{k-1} & \beta_{k-1} \\ & & & \beta_{k-1} & \alpha_k \end{pmatrix} \in \mathbb{R}^{k \times k}.$$

Let

$$\mathbf{y}_k = \arg \min_{\|\mathbf{y}\| \leq \Delta} \left\{ \frac{1}{2} \mathbf{y}^T T_k \mathbf{y} + \beta_0 \mathbf{y}^T \mathbf{e}_1 \right\}. \quad (2.4)$$

It can be easily verified that

$$\mathbf{x}_k := Q_k \mathbf{y}_k = \arg \min_{\|\mathbf{x}\| \leq \Delta, \mathbf{x} \in \mathcal{K}_k} f(\mathbf{x}).$$

Therefore, we obtain the solution $\mathbf{x}_k$ to Problem (2.2), which can be used as an approximation to $\mathbf{x}_{\text{opt}}$.

Since T_k is a tridiagonal and irreducible matrix, the TRS (2.4) is always in the *easy case* [12, Theorem 5.3], i.e., the solution $\mathbf{y}_k$ is unique. By Theorem 2.1, there exists a Lagrange multiplier $\lambda_k \geq 0$, such that

$$\mathbf{y}_k = -\beta_0 (T_k + \lambda_k I)^{-1} \mathbf{e}_1 \quad \text{and} \quad \mathbf{x}_k = -\beta_0 Q_k (T_k + \lambda_k I)^{-1} \mathbf{e}_1. \quad (2.5)$$

So we have from (2.3) that

$$(A + \lambda_k I) \mathbf{x}_k + \mathbf{g} = (\beta_k \mathbf{e}_k^T \mathbf{y}_k) \mathbf{q}_{k+1}. \quad (2.6)$$

The convergence of $\mathbf{x}_k$, λ_k , and $\|(A + \lambda_k I) \mathbf{x}_k + \mathbf{g}\|$ has been thoroughly investigated [24, 27]:

Theorem 2.2. Suppose that TRS (1.1) is in the *easy case* and $\|\mathbf{x}_{\text{opt}}\| = \|\mathbf{x}_k\| = \Delta$. Then

$$\begin{aligned} \frac{\|\mathbf{x}_k - \mathbf{x}_{\text{opt}}\|}{\Delta} &\leq 4 \sqrt{\kappa} \left(\frac{\sqrt{\kappa} - 1}{\sqrt{\kappa} + 1} \right)^k, \\ \|(A + \lambda_k I) \mathbf{x}_k + \mathbf{g}\| &\leq 2 \|A_{\text{opt}}\| \sqrt{\Delta^2 + \|(I - Q_k Q_k^T) \mathbf{x}_{\text{opt}}\|^2} \cdot \left(\frac{\sqrt{\kappa} - 1}{\sqrt{\kappa} + 1} \right)^k, \\ 0 \leq \lambda_{\text{opt}} - \lambda_k &\leq 2 s_k \cdot \left(1 + \frac{2k \sqrt{\kappa}}{\|A_{\text{opt}}\|} \left(\frac{\|\mathbf{x}_{\text{opt}}\|_{A_{\text{opt}}}}{\Delta} \right)^2 \right) \left(\frac{\sqrt{\kappa} - 1}{\sqrt{\kappa} + 1} \right)^{2k}, \end{aligned}$$

where κ denotes the spectral condition number of $A + \lambda_{\text{opt}} I$ and

$$s_k = \frac{\Delta^2}{\|\mathbf{g}\|^2 \cdot \mathbf{e}_1^T (T_k + \lambda_{\text{opt}} I)^{-3} \mathbf{e}_1}.$$

The GLTR can be divided into two phases. In this first phases, $\mathbf{x}_k$ can recurs from the conjugate-gradient (CG) iteration [12]:

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \rho_k \mathbf{p}_k, \quad k \geq 1$$

where

$$\begin{aligned} \mathbf{g}_1 = -\mathbf{p}_1 = \mathbf{g}, \quad \mathbf{x}_1 = \mathbf{0}, \quad \rho_k &= \frac{\|\mathbf{g}_k\|^2}{\mathbf{p}_k^T A \mathbf{p}_k}, \quad \mathbf{g}_{k+1} = \mathbf{g}_k + \rho_k A \mathbf{p}_k, \\ \tau_k &= \frac{\|\mathbf{g}_{k+1}\|^2}{\|\mathbf{g}_k\|^2}, \quad \mathbf{p}_{k+1} = -\mathbf{g}_{k+1} + \tau_k \mathbf{p}_k. \end{aligned}$$

Moreover, $\mathbf{q}_k = \frac{\sigma_k \mathbf{g}_k}{\|\mathbf{g}_k\|}$ where $\sigma_k = -\text{sign}(\rho_{k-1})\sigma_{k-1}$, $\sigma_1 = 1$, and

$$T_k = \begin{pmatrix} \frac{1}{\rho_1} & \frac{\sqrt{\tau_1}}{|\rho_1|} & & & \\ \frac{\sqrt{\tau_1}}{|\rho_1|} & \frac{1}{\rho_2} + \frac{\tau_1}{\rho_1} & \frac{\sqrt{\tau_2}}{|\rho_2|} & & \\ & \ddots & \ddots & \ddots & \\ & & \frac{\sqrt{\tau_{k-2}}}{|\rho_{k-2}|} & \frac{1}{\rho_{k-1}} + \frac{\tau_{k-2}}{\rho_{k-2}} & \frac{\sqrt{\tau_{k-1}}}{|\rho_{k-1}|} \\ & & & \frac{\sqrt{\tau_{k-1}}}{|\rho_{k-1}|} & \frac{1}{\rho_k} + \frac{\tau_{k-1}}{\rho_{k-1}} \end{pmatrix}. \quad (2.7)$$

In the CG iterations, $\|\mathbf{x}_k\|$ increases monotonically. As a result, the CG iteration terminates if $\|\mathbf{x}_{\text{opt}}\| < \Delta$, or $\|\mathbf{x}_k\| > \Delta$, or $\mathbf{p}_k^T A \mathbf{p}_k < 0$. If $\mathbf{x}_{\text{opt}}$ is not obtained after the CG iterations terminate, GLTR proceeds to the second phase, where it needs to solve a smaller-scale TRS (2.4) to obtain $\mathbf{y}_k$. In short, GLTR is divided into the following two phases:

- (i) $k = 1, 2, \dots, k_{\text{CG}}$: $\mathbf{x}_{k+1} = \mathbf{x}_k + \rho_k \mathbf{p}_k$ (CG iteration);
- (ii) $k = k_{\text{CG}} + 1, k_{\text{CG}} + 2, \dots$ until convergence: $\mathbf{x}_k = Q_k \mathbf{y}_k$.

When $A > \mathbf{0}$ and $\|A^{-1}\mathbf{g}\| \leq \Delta$ ($\lambda_{\text{opt}} = 0$), the GLTR method reduces to the standard CG method and $\mathbf{x}_{\text{opt}} = -A^{-1}\mathbf{g}$. In the following, we assume that the CG iterations do not solve TRS (1.1), meaning that the second phase is activated, i.e., $\lambda_{\text{opt}} > 0$ and $\|\mathbf{x}_{\text{opt}}\| = \Delta$. The GLTR algorithm is listed in Algorithm 1 (see [12] in detail).

3. A minimal residual method for the TRS

Inspired by the GMRES/MINRES methods [21, 28], we seek $(\widehat{\lambda}, \widehat{\mathbf{x}})$ to minimize the residual norm $\|(A + \widehat{\lambda}I)\widehat{\mathbf{x}} + \mathbf{g}\|$ in the *second phase*, where $\widehat{\mathbf{x}} \in \mathcal{K}_k$. Note that in the GLTR method, the convergence rate of λ_k is faster than that of $\mathbf{x}_k$ (see Theorem 2.2). Therefore, in our method, we fix λ_k and seek $\widetilde{\mathbf{x}}_k \in \mathcal{K}_k$ such that

$$\widetilde{\mathbf{x}}_k = \arg \min_{\mathbf{x} \in \mathcal{K}_k} \|(A + \lambda_k I)\mathbf{x} + \mathbf{g}\|.$$

Algorithm 1 The GLTR algorithm.

Require: A , $\mathbf{g}$, Δ , tol , and $It_{\max}$.

Ensure: $\mathbf{x}_k$ and $f(\mathbf{x}_k)$

```

1:  $\beta_0 = \|\mathbf{g}\|$ ,  $\mathbf{g}_1 = -\mathbf{p}_1 = \mathbf{g}$ ,  $\mathbf{x}_1 = \mathbf{0}$ ,  $k = 1$ 
2: while  $\frac{\|\mathbf{g}_k\|}{\|\mathbf{g}\|} \geq tol$  do
3:    $\rho_k = \frac{\|\mathbf{g}_k\|^2}{\mathbf{p}_k^T A \mathbf{p}_k}$ 
4:   Obtain  $T_k$  from  $T_{k-1}$  using (2.7)
5:   if  $\rho_k \leq 0$  or  $\|\mathbf{x}_k + \rho_k \mathbf{p}_k\| > \Delta$  then
6:      $k_{CG} = k$ ; Go to Step 15
7:   break
8:   end if
9:    $\mathbf{x}_{k+1} = \mathbf{x}_k + \rho_k \mathbf{p}_k$ 
10:   $\mathbf{g}_{k+1} = \mathbf{g}_k + \rho_k A \mathbf{p}_k$ 
11:   $\tau_k = \frac{\|\mathbf{g}_{k+1}\|^2}{\|\mathbf{g}_k\|^2}$ 
12:   $\mathbf{p}_{k+1} = -\mathbf{g}_{k+1} + \tau_k \mathbf{p}_k$ 
13:   $k = k + 1$ 
14: end while
15: for  $k = k_{CG} + 1, \dots, It_{\max}$  do
16:    $\rho_k = \frac{\|\mathbf{g}_k\|^2}{\mathbf{p}_k^T A \mathbf{p}_k}$ 
17:   Solve TRS (2.4) to obtain  $\mathbf{y}_k$ 
18:   if  $\frac{\beta_k |\mathbf{e}_k^T \mathbf{y}_k|}{\|\mathbf{g}\|} \leq tol$  then
19:     Output  $\mathbf{x}_k = Q_k \mathbf{y}_k$ ; return
20:   end if
21:    $\mathbf{g}_{k+1} = \mathbf{g}_k + \rho_k A \mathbf{p}_k$ 
22:    $\tau_k = \frac{\|\mathbf{g}_{k+1}\|^2}{\|\mathbf{g}_k\|^2}$ 
23:    $\mathbf{p}_{k+1} = -\mathbf{g}_{k+1} + \tau_k \mathbf{p}_k$ 
24: end for

```

We use $\tilde{\mathbf{x}}_k$ instead of $\mathbf{x}_k$ as an approximation for $\mathbf{x}_{\text{opt}}$. Let

$$\tilde{T}_k = \begin{pmatrix} \alpha_1 & \beta_1 & & & \\ \beta_1 & \alpha_2 & \beta_2 & & \\ & \ddots & \ddots & \ddots & \\ & & \beta_{k-2} & \alpha_{k-1} & \beta_{k-1} \\ & & & \beta_{k-1} & \alpha_k \\ & & & & \beta_k \end{pmatrix} = \begin{pmatrix} T_k \\ \beta_k \mathbf{e}_k^T \end{pmatrix}, \quad \tilde{I}_k = \begin{pmatrix} I_k \\ \mathbf{0} \end{pmatrix} \in \mathbb{R}^{(k+1) \times k}.$$

Hence,

$$\begin{aligned} \min_{\mathbf{x} \in \mathcal{K}_k} \|(A + \lambda_k I)\mathbf{x} + \mathbf{g}\| &= \min_{\mathbf{y}} \|(A + \lambda_k I)Q_k \mathbf{y} + \beta_0 Q_{k+1} \mathbf{e}_1\| \\ &\stackrel{(2,3)}{=} \min_{\mathbf{y}} \|Q_{k+1}[(\tilde{T}_k + \lambda_k \tilde{I}_k)\mathbf{y} + \beta_0 \mathbf{e}_1]\| \\ &= \min_{\mathbf{y}} \|(\tilde{T}_k + \lambda_k \tilde{I}_k)\mathbf{y} + \beta_0 \mathbf{e}_1\|. \end{aligned} \quad (3.1)$$

Let

$$\tilde{\mathbf{y}}_k = -\beta_0(\tilde{T}_k + \lambda_k \tilde{I}_k)^\dagger \mathbf{e}_1 = \arg \min_{\mathbf{y}} \|(\tilde{T}_k + \lambda_k \tilde{I}_k)\mathbf{y} + \beta_0 \mathbf{e}_1\|.$$

Thus,

$$\tilde{\mathbf{x}}_k = Q_k \tilde{\mathbf{y}}_k = \arg \min_{\mathbf{x} \in \mathcal{K}_k} \|(A + \lambda_k I)\mathbf{x} + \mathbf{g}\|,$$

is taken as an approximation to $\mathbf{x}_{\text{opt}}$. The proposed algorithm (MRTR) is given in Algorithm 2. The difference between the MRTR method and the GLTR method lies in Lines 17–20.

3.1. Main results

In this section, we provide some theoretical analysis for the MRTR algorithm. To this end, we need the following lemma.

Lemma 3.1. *For $k > 1$, we have*

$$|\mathbf{e}_1^T (T_k + \lambda_k I_k)^{-3} \mathbf{e}_k| \leq C q^{k-1},$$

where C is a constant and $0 < q < 1$ depends only on function $f(x) = x^{-3}$.

Proof. Note that $T_k + \lambda_k I_k$ is a symmetric positive definite matrix with bandwidth 1; thus, the conclusion follows directly from [30, p. 317, Theorem 14.1]. $\square$

The following result establishes the relationship between $\mathbf{x}_k$ and $\tilde{\mathbf{x}}_k$.

Theorem 3.1. *For $k \geq 1$, we have*

$$\|\mathbf{x}_k - \tilde{\mathbf{x}}_k\| \leq 2\xi_k \Delta \left(\frac{\sqrt{\kappa_k} - 1}{\sqrt{\kappa_k} + 1} \right)^{k-1}, \quad (3.2a)$$

$$\left| \|\mathbf{x}_k\| - \|\tilde{\mathbf{x}}_k\| \right| \leq \frac{4(\xi_k^2 \Delta^2 + \zeta_k)}{\|\mathbf{x}_k\| + \|\tilde{\mathbf{x}}_k\|} \left(\frac{\sqrt{\kappa_k} - 1}{\sqrt{\kappa_k} + 1} \right)^{k-1} \cdot q^{k-1}, \quad (3.2b)$$

Algorithm 2 A minimal residual method for the TRS (1.1) (MRTR).

```

1: Input:  $A, g, \Delta, tol$ , and  $It_{\max}$ .
2: Output:  $\tilde{x}_k$  and  $f(\tilde{x}_k)$ 
3:  $\beta_0 = \|g\|, g_1 = -p_1 = g, x_1 = 0, k = 1$ 
4: while  $\frac{\|g_k\|}{\|g\|} \geq tol$  do
5:    $\rho_k = \frac{\|g_k\|^2}{p_k^T A p_k}$ 
6:   Obtain  $T_k$  from  $T_{k-1}$  using (2.7)
7:   if  $\rho_k \leq 0$  or  $\|x_k + \rho_k p_k\| > \Delta$  then
8:      $k_{CG} = k$ ; Go to Step 15
9:   break
10:  end if
11:   $x_{k+1} = x_k + \rho_k p_k$ 
12:   $g_{k+1} = g_k + \rho_k A p_k$ 
13:   $\tau_k = \frac{\|g_{k+1}\|^2}{\|g_k\|^2}$ 
14:   $p_{k+1} = -g_{k+1} + \tau_k p_k$ 
15:   $k = k + 1$ 
16: end while
17: for  $k = k_{CG} + 1, \dots, It_{\max}$  do
18:    $\rho_k = \frac{\|g_k\|^2}{p_k^T A p_k}$ 
19:   Solve TRS (2.4) to obtain  $\lambda_k$  and compute  $\tilde{y}_k = -\beta_0(\tilde{T}_k + \lambda_k \tilde{I}_k)^\dagger e_1$ 
20:   if  $\frac{\|(\tilde{T}_k + \lambda_k \tilde{I}_k)\tilde{y}_k + \beta_0 e_1\|}{\|g\|} \leq tol$  then
21:     Output  $\tilde{x}_k = Q_k \tilde{y}_k$ ; return
22:   end if
23:    $g_{k+1} = g_k + \rho_k A p_k$ 
24:    $\tau_k = \frac{\|g_{k+1}\|^2}{\|g_k\|^2}$ 
25:    $p_{k+1} = -g_{k+1} + \tau_k p_k$ 
26: end for

```

where

$$\xi_k = \frac{\beta_k^2 \|(T_k + \lambda_k I_k)^{-2} \mathbf{e}_k\|}{1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k} \text{ and } \zeta_k = \frac{C\beta_k^2\beta_0\Delta}{1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k}.$$

Proof. We first prove (3.2a). Recall that TRS (2.4) is in the *easy case*. Consequently, $T_k + \lambda_k I_k > \mathbf{0}$. This immediately implies that

$$(\widetilde{T}_k + \lambda_k \widetilde{I}_k)^T (\widetilde{T}_k + \lambda_k \widetilde{I}_k) = (T_k + \lambda_k I_k)^2 + \beta_k^2 \mathbf{e}_k \mathbf{e}_k^T > \mathbf{0}. \quad (3.3)$$

We further have the identity

$$\begin{aligned} (\widetilde{T}_k + \lambda_k \widetilde{I}_k)^\dagger \mathbf{e}_1 &= [(\widetilde{T}_k + \lambda_k \widetilde{I}_k)^T (\widetilde{T}_k + \lambda_k \widetilde{I}_k)]^{-1} (\widetilde{T}_k + \lambda_k \widetilde{I}_k)^T \mathbf{e}_1 \\ &= [(T_k + \lambda_k I_k)^2 + \beta_k^2 \mathbf{e}_k \mathbf{e}_k^T]^{-1} (T_k + \lambda_k I_k) \mathbf{e}_1 \\ &= (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1 - \frac{\beta_k^2 (\mathbf{e}_k^T (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1)}{1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k} (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k, \end{aligned} \quad (3.4)$$

where the last equality is derived from the Sherman-Morrison-Woodbury formula [21, Section 2.1.4]. It thus follows that

$$\widetilde{\mathbf{x}}_k - \mathbf{x}_k = Q_k(\widetilde{\mathbf{y}}_k - \mathbf{y}_k) = \frac{\beta_0 \beta_k^2 (\mathbf{e}_k^T (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1)}{1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k} \cdot Q_k (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k. \quad (3.5)$$

By (2.5) and [24, Eqs (3.9) and (3.10)],

$$|\mathbf{e}_k^T (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1| = \frac{|\mathbf{e}_k^T \mathbf{y}_k|}{\beta_0} \leq \frac{2\Delta}{\beta_0} \left(\frac{\sqrt{\kappa_k} - 1}{\sqrt{\kappa_k} + 1} \right)^{k-1}. \quad (3.6)$$

The desired bound (3.2a) then follows directly from (3.5) and (3.6).

Next we prove (3.2b). Recall that $\|\mathbf{x}_k\| = \|\mathbf{y}_k\| = \beta_0 \|(T_k + \lambda_k I_k)^{-1} \mathbf{e}_1\|$ and $\|\widetilde{\mathbf{x}}_k\| = \|\widetilde{\mathbf{y}}_k\| = \beta_0 \|(\widetilde{T}_k + \lambda_k \widetilde{I}_k)^\dagger \mathbf{e}_1\|$. It follows from (3.4) that

$$\begin{aligned} \frac{\|\widetilde{\mathbf{x}}_k\|^2 - \|\mathbf{x}_k\|^2}{\beta_0^2} &= \frac{\beta_k^4 (\mathbf{e}_k^T (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1)^2}{(1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k)^2} \|(T_k + \lambda_k I_k)^{-2} \mathbf{e}_k\|^2 \\ &\quad - \frac{2\beta_k^2 (\mathbf{e}_k^T (T_k + \lambda_k I_k)^{-1} \mathbf{e}_1)}{1 + \beta_k^2 \mathbf{e}_k^T (T_k + \lambda_k I_k)^{-2} \mathbf{e}_k} \cdot \mathbf{e}_1^T (T_k + \lambda_k I_k)^{-3} \mathbf{e}_k. \end{aligned} \quad (3.7)$$

The desired bound (3.2b) then follows directly from (3.6) and Lemma 3.1. $\square$

Remark 1. Since $\lambda_k \rightarrow \lambda_{\text{opt}}$, we have $\kappa_k \rightarrow \kappa$ and

$$\left(\frac{\sqrt{\kappa_k} - 1}{\sqrt{\kappa_k} + 1} \right)^{k-1} \rightarrow 0, \text{ as } k \text{ increases.}$$

Thus, the bound in (3.2a) implies that $\mathbf{x}_k$ and $\widetilde{\mathbf{x}}_k$ approach each other as k increase. Furthermore, since $\mathbf{x}_k \rightarrow \mathbf{x}_{\text{opt}}$, it follows that $\widetilde{\mathbf{x}}_k \rightarrow \mathbf{x}_{\text{opt}}$ as k increase.

If there exist some ℓ such that $\lambda_\ell > 0$, then $\lambda_k > 0$ [31, Theorem 2.6] and $\|\mathbf{x}_k\| = \Delta$ for all $k \geq \ell$. It then follows from (3.2b) that $\|\widetilde{\mathbf{x}}_k\| \rightarrow \Delta$ and that the rate at which $\|\widetilde{\mathbf{x}}_k\|$ approaches Δ is faster than the rate at which $\widetilde{\mathbf{x}}_k$ approaches $\mathbf{x}_{\text{opt}}$.

Combining (3.2a) and Theorem 2.2, we immediately obtain the following corollary.

Corollary 3.1. *For $k \geq 1$, we have*

$$\|\tilde{\mathbf{x}}_k - \mathbf{x}_{\text{opt}}\| \leq 2\xi_k \Delta \left(\frac{\sqrt{\kappa_k} - 1}{\sqrt{\kappa_k} + 1} \right)^{k-1} + 4\Delta \sqrt{\kappa} \left(\frac{\sqrt{\kappa} - 1}{\sqrt{\kappa} + 1} \right)^k.$$

The following theorem shows that, provided the Lanczos process doesn't break down, $\|(A + \lambda_k I)\tilde{\mathbf{x}}_k + \mathbf{g}\|$ is always *strictly* less than $\|(A + \lambda_k I)\mathbf{x}_k + \mathbf{g}\|$.

Theorem 3.2. *If $\beta_k \neq 0$, then*

$$\|(A + \lambda_k I)\tilde{\mathbf{x}}_k + \mathbf{g}\| < \|(A + \lambda_k I)\mathbf{x}_k + \mathbf{g}\|.$$

Proof. By (3.3), $\tilde{T}_k + \lambda_k \tilde{I}_k$ is of full column rank. Therefore, $\tilde{\mathbf{y}}_k$ is the unique solution to the least squares problem (3.1). Since $\beta_k \neq 0$, it follows from (3.4) that $\tilde{\mathbf{y}}_k \neq \mathbf{y}_k$. Hence,

$$\begin{aligned} \|(A + \lambda_k I)\mathbf{x}_k + \mathbf{g}\| &= \|(A + \lambda_k I)Q_k \mathbf{y}_k + \mathbf{g}\| \\ &= \|Q_{k+1} [(\tilde{T}_k + \lambda_k \tilde{I}_k)\mathbf{y}_k + \beta_0 \mathbf{e}_1]\| \\ &= \|(\tilde{T}_k + \lambda_k \tilde{I}_k)\mathbf{y}_k + \beta_0 \mathbf{e}_1\| \\ &> \|(\tilde{T}_k + \lambda_k \tilde{I}_k)\tilde{\mathbf{y}}_k + \beta_0 \mathbf{e}_1\| \\ &= \|(A + \lambda_k I)\tilde{\mathbf{x}}_k + \mathbf{g}\|. \end{aligned}$$

This completes the proof. $\square$

If the Lanczos process terminates at step k , i.e., $\beta_k = 0$, then $\mathbf{x}_k \stackrel{(3.4)}{=} \tilde{\mathbf{x}}_k$ and $\mathbf{x}_k = \mathbf{x}_{\text{opt}}$ [27, Lemma 3.3]. Therefore, both the GLTR and MRTR methods can theoretically obtain the exact solution when the Lanczos process terminates.

The standard GLTR method follows the spirit of the Rayleigh-Ritz procedure [32]: it projects the original TRS onto the Krylov subspace via the Lanczos process, and obtains both the approximate Lagrange multiplier and the approximate solution simultaneously under a Galerkin-type projection condition.

The proposed MRTR method shares the same design logic as the refined Rayleigh-Ritz method [32,33]: it takes the approximate Lagrange multiplier yielded by GLTR as a fixed quantity, and then seeks a new approximate solution within the same Krylov subspace by minimizing the residual norm. This refinement step improves the accuracy of the approximate solution without additional matrix-vector multiplications.

Despite the similar refinement idea, the two classes of methods are essentially different in problem structure. The refined Rayleigh-Ritz method reduces a large-scale eigenvalue problem to a small-scale *singular value problem* to compute refined eigenpairs approximations. In contrast, the MRTR method transforms the large-scale TRS into a small-scale *least-squares problem* after fixing the approximate Lagrange multiplier. Due to this essential difference in the nature of the underlying problems, the mature theoretical analysis framework of the refined Rayleigh-Ritz method seems difficult to be directly borrowed or applied to the MRTR method.

For the current work, we focus on algorithm construction, implementation details, and numerical performance verification. We will carry out in-depth theoretical research along this direction in our future work.

4. Numerical experiments

In this section, we evaluate the numerical performance of the proposed method through some numerical examples. All the numerical experiments were run on an AMD R7 5800H CPU 3.20 GHz with 16 GB RAM under the Windows 11 operating system. The experimental results are obtained from using MATLAB R2022a implementation with machine precision $u_{\text{machine}} \approx 2.22 \times 10^{-16}$.

We compare the MRTR method (Algorithm 2) with the GLTR method. In all the experiments, we let $tol = 10^{-8}$, $It_{\max} = 1000$, and three distinct vectors $\mathbf{g}$:

$$\begin{aligned}\mathbf{g}_1 &= \left(\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}, \dots, \frac{1}{\sqrt{n}} \right)^T \in \mathbb{R}^n, \\ \mathbf{g}_2 &= \frac{\mathbf{c}_1}{\|\mathbf{c}_1\|} \text{ with } \mathbf{c}_1 = (\cos 1, \cos 2, \dots, \cos n)^T \in \mathbb{R}^n, \text{ and} \\ \mathbf{g}_3 &= \frac{\mathbf{c}_2}{\|\mathbf{c}_2\|} \text{ with } \mathbf{c}_2 = \text{randn}(n, 1) \in \mathbb{R}^n\end{aligned}$$

To ensure the orthogonality of the basis $\mathbf{Q}_k$, we employ the full reorthogonalization process such that $\|\mathbf{Q}_k^T \mathbf{Q}_k - \mathbf{I}_k\| = O(u_{\text{machine}})$. It is proven in [11] that the TRS can be equivalently transformed into an eigenvalue problem of a matrix. More precisely, taking TRS (2.4) as an example, λ_k is the rightmost eigenvalue of

$$\begin{pmatrix} -T_k & \frac{\beta_0^2 \mathbf{e}_1 \mathbf{e}_1^T}{\Delta^2} \\ \mathbf{I} & -T_k \end{pmatrix} \in \mathbb{R}^{2k \times 2k}.$$

If $\lambda_k > 0$, then $\mathbf{y}_k = \frac{\Delta \cdot \mathbf{z}_k(1:k)}{\|\mathbf{z}_k(1:k)\|}$, where $\mathbf{z}_k \in \mathbb{R}^{2k}$ is the eigenvector corresponding to λ_k . Therefore, we use the built-in MATLAB function `eig` to compute $(\lambda_k, \mathbf{y}_k)$ in the second phase. For the least squares problem (3.1), we use MATLAB's division command `\` to compute $\tilde{\mathbf{y}}_k$.

Recall that we assume the CG iterations don't solve TRS (1.1), meaning that the second phase is activated, i.e., $\lambda_{\text{opt}} > 0$ and $\|\mathbf{x}_{\text{opt}}\| = \Delta$. Therefore, to ensure that the norm of output approximation $\tilde{\mathbf{x}}_k$ is strictly equal to Δ numerically, we normalize $\tilde{\mathbf{y}}_k$ when the residual norm in the MRTR algorithm approaches the termination criterion. More precisely, we replace lines 18–20 in Algorithm 2 with the following operations:

```

if  $\frac{\|(\tilde{T}_k + \lambda_k \tilde{I}_k) \tilde{\mathbf{y}}_k + \beta_0 \mathbf{e}_1\|}{\|\mathbf{g}\|} \leq tol$ 
   $\tilde{\mathbf{y}}_k = \frac{\Delta \cdot \tilde{\mathbf{y}}_k}{\|\tilde{\mathbf{y}}_k\|}$ ,  $\tilde{\mathbf{x}}_k = \mathbf{Q}_k \tilde{\mathbf{y}}_k$ 
  if  $\frac{\|A \tilde{\mathbf{x}}_k + \lambda_k \tilde{\mathbf{x}}_k + \mathbf{g}\|}{\|\mathbf{g}\|} \leq tol$ 
    Output  $\tilde{\mathbf{x}}_k$ ; return
  end if
end if

```

Define

$$f_{\text{relerr}} := \left| \frac{f(\mathbf{x}_{\text{GLTR}}) - f(\mathbf{x}_{\text{MRTR}})}{\min \{f(\mathbf{x}_{\text{GLTR}}), f(\mathbf{x}_{\text{MRTR}})\}} \right|$$

where $f(\mathbf{x}_{\text{GLTR}})$ and $f(\mathbf{x}_{\text{MRTR}})$ denote the approximate objective function values obtained by the GLTR and MRTR algorithms, respectively. This metric measures the relative discrepancy between the optimal objective values output by the two iterative methods.

In this section, ‘CPU(s)’ denotes the running time of the algorithms, ‘It’ denotes the number of iterations, and $f(\bar{\mathbf{x}})$ denotes the approximate optimal values obtained by the two algorithms.

Example 1. In this example, we consider two diagonal matrices.

- Example 1(a): $A = \text{diag}(t_{jn}^{[a,b]})$ [26, 27], with

$$t_{jn}^{[a,b]} = \left(\frac{b-a}{2}\right)\left(t_{jn} + \left(\frac{a+b}{b-a}\right)\right) \text{ with } t_{jn} = \cos \frac{(2j-1)\pi}{2n}, \quad j = 1, 2, \dots, n, \quad (4.1)$$

where $n = 10,000$, $-a = b = 100$, and $t_{jn}^{[a,b]}$ is the n -th translated Chebyshev zero nodes on $[a, b]$.

- Example 1(b): $A = \text{diag}(\{i - n - 1\}_{i=1}^n)$ [25], with $n = 10,000$.

In this example, we set $\Delta = 1$. The numerical results of Example 1 are presented in Figure 1 and Table 1. Figure 1 shows the convergence curves of the residual norms $\frac{\|(A+\lambda_k I)\bar{\mathbf{x}}+\mathbf{g}\|}{\|\mathbf{g}\|}$ for the two algorithms.

Table 1. Numerical experiments on Example 1.

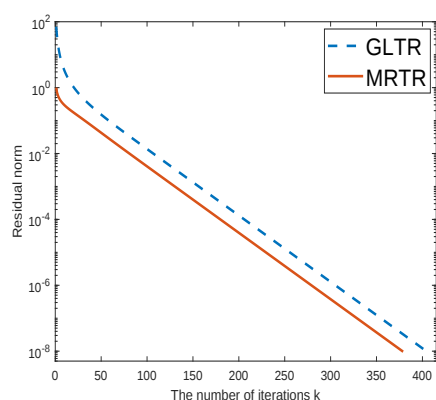
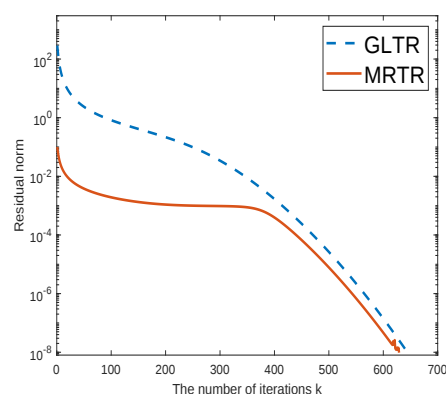
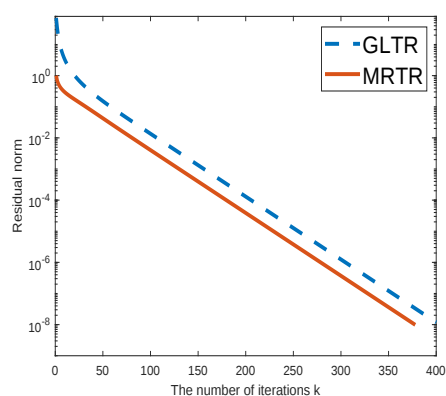
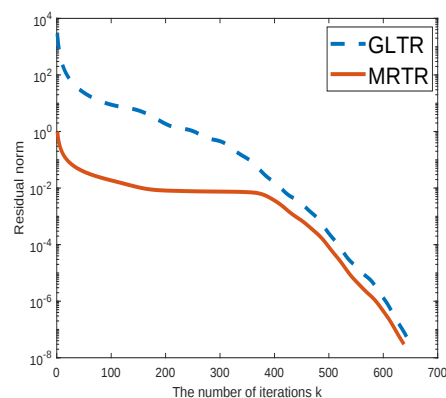
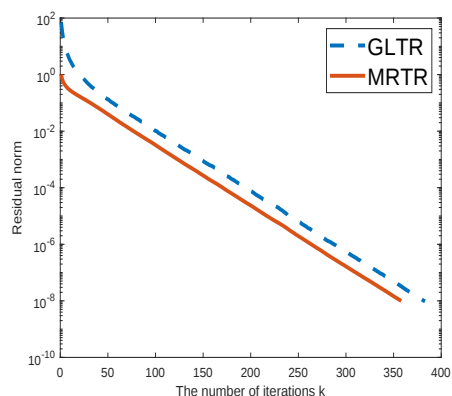
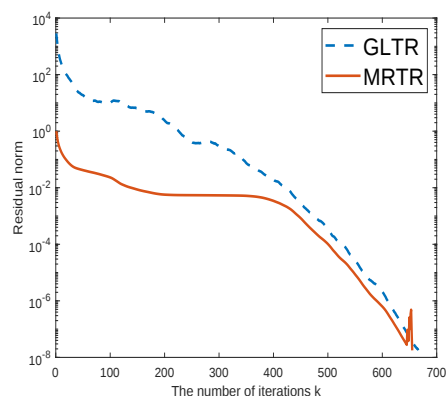
	Matrix	Method	CPU (s)	It	$f(\bar{\mathbf{x}})$	f_{relerr}
$\mathbf{g} = \mathbf{g}_1$	$A = \text{diag}(t_{jn}^{[a,b]})$	GLTR	69.0	405	-50.161553602236729	1.5×10^{-15}
		MRTR	56.8	379	-50.161553602236651	
	$A = \text{diag}(\{i - n - 1\}_{i=1}^n)$	GLTR	282.2	645	-5.000010488558799e+03	9.1×10^{-16}
		MRTR	258.5	629	-5.000010488558803e+03	
$\mathbf{g} = \mathbf{g}_2$	$A = \text{diag}(t_{jn}^{[a,b]})$	GLTR	72.3	404	-50.161827411264966	7.1×10^{-16}
		MRTR	51.7	363	-50.161827411264966	
	$A = \text{diag}(\{i - n - 1\}_{i=1}^n)$	GLTR	286.7	651	-5.000008114793698e+03	2.5×10^{-15}
		MRTR	280.1	650	-5.000008114793711e+03	
$\mathbf{g} = \mathbf{g}_3$	$A = \text{diag}(t_{jn}^{[a,b]})$	GLTR	62.7	383	-50.165962388210183	1.5×10^{-15}
		MRTR	50.2	358	-50.165962388210261	
	$A = \text{diag}(\{i - n - 1\}_{i=1}^n)$	GLTR	385.5	666	-5.000005957484766e+03	1.3×10^{-15}
		MRTR	286.0	655	-5.000005957484773e+03	

Example 2. In this example, we consider three symmetric matrices `nd3k`, `nemeth01`, and `stokes64s` from the University of Florida Sparse Matrix Collection*. The numerical results are presented in Figure 2 and Table 2. Figure 2 shows the convergence curves of the residual norms $\frac{\|(A+\lambda_k I)\bar{\mathbf{x}}+\mathbf{g}\|}{\|\mathbf{g}\|}$ for the two algorithms. In this example, we set $\Delta = 10$.

As illustrated in Figures 1 and 2, MRTR achieves a lower residual norm at every iteration step compared to GLTR. In particular, it can be seen from Figure 2 that among the three numerical examples in Example 2, the residual convergence curve of the GLTR method exhibits obvious oscillations, whereas that of the MRTR method is smooth, which reflects that the MRTR method has better stability.

It can be seen from Tables 1 and 2 that the MRTR outperforms the GLTR in terms of both running time and the number of iterations. In addition, the relative errors f_{relerr} of the approximate optimal

*<https://sparse.tamu.edu>.

(a) Example 1(a): $A = \text{diag}(t_{jn}^{[a,b]}), \mathbf{g} = \mathbf{g}_1$ (b) Example 1(b): $A = \text{diag}(\{i - n - 1\}_{i=1}^n), \mathbf{g} = \mathbf{g}_1$ (c) Example 1(a): $A = \text{diag}(t_{jn}^{[a,b]}), \mathbf{g} = \mathbf{g}_2$ (d) Example 1(b): $A = \text{diag}(\{i - n - 1\}_{i=1}^n), \mathbf{g} = \mathbf{g}_2$ (e) Example 1(a): $A = \text{diag}(t_{jn}^{[a,b]}), \mathbf{g} = \mathbf{g}_3$ (f) Example 1(b): $A = \text{diag}(\{i - n - 1\}_{i=1}^n), \mathbf{g} = \mathbf{g}_3$ **Figure 1.** Numerical experiments on Example 1.

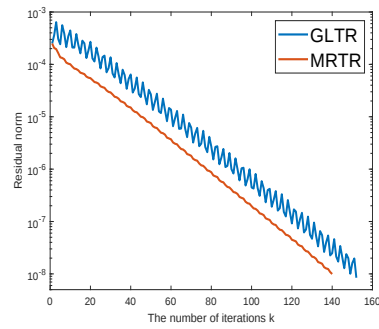
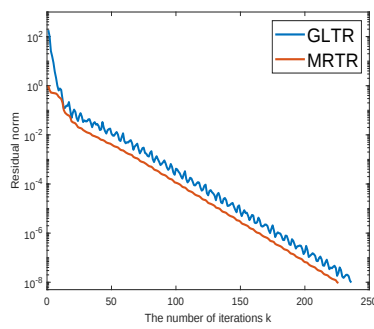
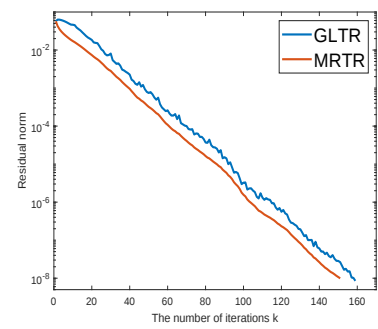
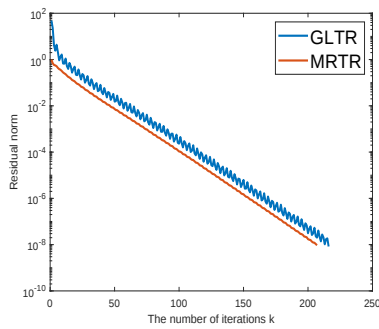
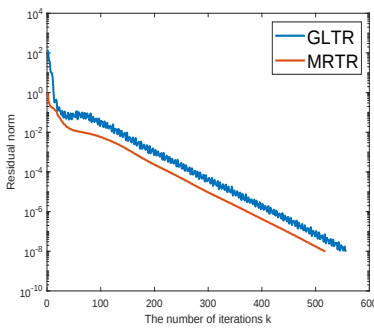
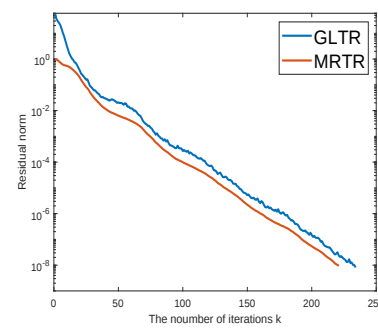
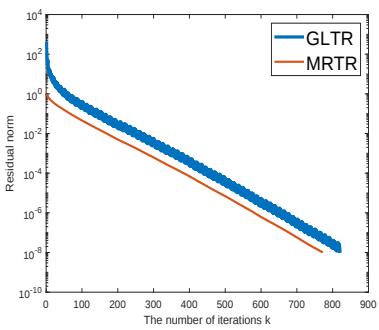
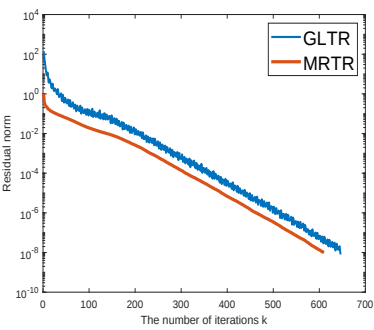
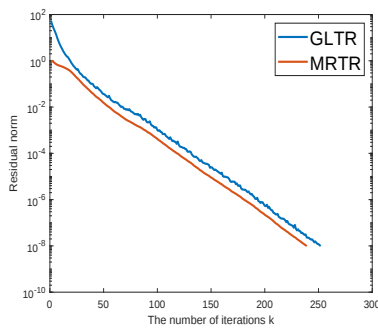
(a) $A = \text{nd3k}, g = g_1$ (b) $A = \text{nemeth01}, g = g_1$ (c) $A = \text{stokes64s}, g = g_1$ (d) $A = \text{nd3k}, g = g_2$ (e) $A = \text{nemeth01}, g = g_2$ (f) $A = \text{stokes64s}, g = g_2$ (g) $A = \text{nd3k}, g = g_3$ (h) $A = \text{nemeth01}, g = g_3$ (i) $A = \text{stokes64s}, g = g_3$ **Figure 2.** Numerical experiments on Example 2.

Table 2. Numerical experiments on Example 2.

	Matrix	Method	CPU (s)	It	$f(\tilde{\mathbf{x}})$	f_{relerr}
$\mathbf{g} = \mathbf{g}_1$	nd3k	GLTR	3.5	152	-9.999539149018121	1.8×10^{-16}
		MRTR	2.4	140	-9.999539149018119	
	nemeth01	GLTR	12.6	236	$-2.158396347063678 \times 10^3$	4.2×10^{-16}
		MRTR	10.2	226	$-2.158396347063679 \times 10^3$	
	stokes64s	GLTR	3.8	159	-10.001772476459125	5.3×10^{-16}
		MRTR	2.5	151	-10.001772476459120	
$\mathbf{g} = \mathbf{g}_2$	nd3k	GLTR	10.6	216	-0.279360075402613	0
		MRTR	8.4	207	-0.279360075402613	
	nemeth01	GLTR	184.9	556	$-2.154654144026191 \times 10^3$	4.2×10^{-16}
		MRTR	148.7	518	$-2.154654144026190 \times 10^3$	
	stokes64s	GLTR	12.1	234	-5.858337377260747	0
		MRTR	9.0	221	-5.858337377260747	
$\mathbf{g} = \mathbf{g}_3$	nd3k	GLTR	640.6	821	-1.793664237863516	5.5×10^{-15}
		MRTR	528.6	772	-1.793664237863506	
	nemeth01	GLTR	299.0	646	$-2.153939233613416 \times 10^3$	2.5×10^{-15}
		MRTR	245.3	610	$-2.153939233613411 \times 10^3$	
	stokes64s	GLTR	16.0	252	-6.000328380904500	1.3×10^{-15}
		MRTR	12.3	239	-6.000328380904508	

values obtained by the two methods reach the order of $O(u_{\text{machine}})$. This implies that the quality of the approximate solutions obtained by the MRTR is comparable to that of the GLTR.

Numerical experiments show that the MRTR method outperforms the GLTR method. This makes MRTR a robust and efficient choice for solving the large-scale TRS problems commonly encountered in scientific computing.

5. Conclusions

In this paper, we propose a minimal residual method for solving the large-scale TRS (MRTR) in the easy case, which is inspired by the GMRES/MINRES and the GLTR method. We first establish the theoretical relationship between $\mathbf{x}_k$ and $\tilde{\mathbf{x}}_k$, proving that the two solutions converge to each other as the iteration proceeds and that the MRTR method yields a strictly smaller residual norm than GLTR when the Lanczos process does not break down, which verifies the theoretical superiority of MRTR in residual reduction.

Numerical experiments on both constructed diagonal matrices and real sparse symmetric matrices from the University of Florida Sparse Matrix Collection demonstrate that the MRTR method outperforms the GLTR method in terms of iteration numbers and CPU time, while achieving approximate optimal values with the same high precision as GLTR. Future research can focus on extending the MRTR method to TRS in the *hard case* and exploring its applications in more complex optimization models such as cubic regularization and constrained nonlinear optimization.

Use of AI tools declaration

The authors declare they have not used Artificial Intelligence (AI) tools in the creation of this article.

Acknowledgements

This work was supported by the Basic Research Program of Jiangsu Province (Youth Program of Natural Science Foundation of Jiangsu Province, Grant No. BK20250892).

Conflict of interest

The authors declare there is no conflicts of interest.

References

1. A. R. Conn, N. I. M. Gould, P. L. Toint, *Trust-Region Methods*, SIAM, Philadelphia, PA, 2000. <https://doi.org/10.1137/1.9780898719857>
2. J. Nocedal, S. Wright, *Numerical Optimization*, 2nd edition, Springer, New York, 2006.
3. L. Zhang, C. Shen, W. Yang, J. Júdice, A Lanczos method for large-scale extreme Lorentz eigenvalue problems, *SIAM J. Matrix Anal. Appl.*, **39** (2018), 611–631. <https://doi.org/10.1137/17M1111401>
4. W. W. Hager, Y. Krylyuk, Graph partitioning and continuous quadratic programming, *SIAM J. Discrete Math.*, **12** (1999), 500–523. <https://doi.org/10.1137/S0895480199335829>
5. M. Rojas, S. A. Santos, D. C. Sorensen, A new matrix-free algorithm for the large-scale trust-region subproblem, *SIAM J. Optim.*, **11** (2001), 611–646. <https://doi.org/10.1137/S105262349928887X>
6. M. Rojas, S. Santos, D. Sorensen, Algorithm 873: LSTRS: MATLAB software for large-scale trust-region subproblems and regularization, *ACM Trans. Math. Software*, **34** (2008), 1–28. <https://doi.org/10.1145/1326548.1326553>
7. N. M. Gould, D. P. Robinson, H. S. Thorne, On solving trust-region and other regularised subproblems in optimization, *Math. Program. Comput.*, **2** (2010), 21–57. <https://doi.org/10.1007/s12532-010-0011-7>
8. T. Steihaug, The conjugate gradient method and trust regions in large scale optimization, *SIAM J. Numer. Anal.*, **20** (1983), 626–637, <https://doi.org/10.1137/0720042>
9. P. Toint, Towards an efficient sparsity exploiting Newton method for minimization, in *Sparse Matrices and Their Uses*, Academic Press, London, (1981), 57–88.
10. W. Gander, G.H. Golub, U. von Matt, A constrained eigenvalue problem, *Linear Algebra Appl.*, **114–115** (1989), 815–839. [https://doi.org/10.1016/0024-3795\(89\)90494-1](https://doi.org/10.1016/0024-3795(89)90494-1)
11. S. Adachi, S. Iwata, Y. Nakatsukasa, A. Takeda, Solving the trust-region subproblem by a generalized eigenvalue problem, *SIAM J. Optim.*, **27** (2017), 269–291. <https://doi.org/10.1137/16M1058200>
12. N. Gould, S. Lucidi, M. Roma, P. Toint, Solving the trust-region subproblem using the Lanczos method, *SIAM J. Optim.*, **9** (1999), 504–525. <https://doi.org/10.1137/S1052623497322735>

13. L. Zhang, C. Shen, A nested Lanczos method for the trust-region subproblem, *SIAM J. Sci. Comput.*, **40** (2018), A2005–A2032. <https://doi.org/10.1137/17M1145914>
14. F. Rendl, H. Wolkowicz, A semidefinite framework for trust region subproblems with applications to large scale minimization, *Math. Program.*, **77** (1997), 273–299. <https://doi.org/10.1007/BF02614438>
15. J. B. Erway, P. E. Gill, J. D. Griffin, Iterative methods for finding a trust-region step, *SIAM J. Optim.*, **20** (2009), 1110–1131. <https://doi.org/10.1137/070708494>
16. G. H. Golub, U. von Matt, Quadratically constrained least squares and quadratic problems, *Numer. Math.*, **59** (1991), 561–580. <https://doi.org/10.1007/BF01385796>
17. J. J. Moré, D. C. Sorensen, Computing a trust region step, *SIAM J. Sci. Stat. Comput.*, **4** (1983), 553–572. <https://doi.org/10.1137/0904038>
18. D. C. Sorensen, Newton’s method with a model trust region modification, *SIAM J. Numer. Anal.*, **19** (1982), 409–426. <https://doi.org/10.1137/0719026>
19. D. Sorensen, Minimization of a large-scale quadratic function subject to a spherical constraint, *SIAM J. Optim.*, **7** (1997), 141–161. <https://doi.org/10.1137/S1052623494274374>
20. P. Tao, L. An, A DC optimization algorithm for solving the trust-region subproblem, *SIAM J. Optim.*, **8** (1998), 476–505. <https://doi.org/10.1137/S1052623494274313>
21. G. H. Golub, C. F. Van Loan, *Matrix Computations*, 4th edition, Johns Hopkins University Press, Baltimore, MD, 2013.
22. Y. Carmon, J. C. Duchi, First-order method for nonconvex quadratic minimization, *SIAM Rev.*, **62** (2020), 395–436. <https://doi.org/10.1137/20M1321759>
23. Y. Carmon, J. C. Duchi, Analysis of Krylov subspace solutions of regularized nonconvex quadratic problems, in *NIPS’18: Proceedings of the 32nd International Conference on Neural Information Processing Systems*, (2018), 10728–10738.
24. B. Feng, G. Wu, On convergence of the generalized Lanczos trust-region method for trust-region subproblems, *Adv. Comput. Math.*, **51** (2025), 4. <https://doi.org/10.1007/s10444-024-10217-5>
25. N. I. M. Gould, V. Simoncini, Error estimates for iterative algorithms for minimizing regularized quadratic subproblems, *Optim. Methods Software*, **35** (2020), 304–328. <https://doi.org/10.1080/10556788.2019.1670177>
26. Z. Jia, F. Wang, The convergence of the generalized Lanczos trust-region method for the trust-region subproblem, *SIAM J. Optim.*, **31** (2021), 887–914. <https://doi.org/10.1137/19M1279691>
27. L. Zhang, C. Shen, R. Li, On the generalized Lanczos trust-region method, *SIAM J. Optim.*, **27** (2017), 2110–2142. <https://doi.org/10.1137/16M1095056>
28. Y. Saad, *Iterative Methods for Sparse Linear Systems*, 2nd edition, SIAM, Philadelphia, PA, 2003. <https://doi.org/10.1137/1.9780898718003>
29. W. W. Hager, Minimizing a quadratic over a sphere, *SIAM J. Optim.*, **12** (2001), 188–208. <https://doi.org/10.1137/S1052623499356071>
30. N. J. Higham, *Functions of Matrices: Theory and Computation*, SIAM, Philadelphia, 2008. <https://doi.org/10.1137/1.9780898717778>

31. L. Lukšan, C. Matonoha, J. Vlček, On Lagrange multipliers of trust-region subproblems, *BIT Numer. Math.*, **48** (2008), 763–768. <https://doi.org/10.1007/s10543-008-0197-5>
32. G. W. Stewart, *Matrix Algorithms II: Eigensystems*, SIAM, Philadelphia, 2001. <https://doi.org/10.1137/1.9780898718058>
33. Z. Jia, Refined iterative algorithms based on Arnoldi's process for large unsymmetric eigenproblems, *Linear Algebra Appl.*, **259** (1997), 1–23. [https://doi.org/10.1016/S0024-3795\(96\)00238-8](https://doi.org/10.1016/S0024-3795(96)00238-8)



AIMS Press

©2026 the Author(s), licensee AIMS Press. This is an open access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0>)