



Simple stopping criteria for the LSQR method applied to discrete ill-posed problems

Lothar Reichel¹ · Hassane Sadok² · Wei-Hong Zhang^{1,3}

Dedicated to Gérard Meurant on the occasion of his 70th birthday.

Received: 25 August 2019 / Accepted: 8 November 2019 / Published online: 22 January 2020
© Springer Science+Business Media, LLC, part of Springer Nature 2020

Abstract

The LSQR iterative method is one of the most popular numerical schemes for computing an approximate solution of large linear discrete ill-posed problems with an error-contaminated right-hand side, which represents available data. It is important to terminate the iterations after a suitable number of steps, because too many steps yield an approximate solution that suffers from a large propagated error due to the error in the data, and too few iterations give an approximate solution that may lack many details that can be of interest. When the error in the right-hand side is white Gaussian and a tight bound on its variance is known, the discrepancy principle typically furnishes a suitable termination criterion for the LSQR iterations. However, in many applications in science and engineering that give rise to large linear discrete ill-posed problems, the variance of the error is not known. This has spurred the development of a variety of stopping rules for assessing when to terminate the iterations in this situation. The present paper proposes new simple stopping rules that are based on comparing the residual errors associated with iterates generated by the LSQR and Craig iterative methods.

Keywords Discrete ill-posed problem · Golub–Kahan bidiagonalization · LSQR · Craig’s method · Stopping criterion

Mathematics Subject Classification (2010) 65F10 · 65F22

✉ Lothar Reichel
reichel@math.kent.edu

1 Introduction

We are concerned with the solution of large linear systems of equations

$$\mathbf{A}\mathbf{x} = \mathbf{b}, \quad \mathbf{A} \in \mathbb{R}^{n \times n}, \quad \mathbf{x}, \mathbf{b} \in \mathbb{R}^n, \quad (1)$$

with a matrix $\mathbf{A}$, whose singular values “cluster” at the origin. In particular, $\mathbf{A}$ is severely ill-conditioned and may be rank-deficient; the system (1) may not have a solution. Matrices of this kind arise, for instance, from the discretization of linear ill-posed problems such as Fredholm integral equations of the first kind. Therefore, linear systems of (1) with this type of matrices are commonly referred to as linear discrete ill-posed problems. The right-hand side $\mathbf{b}$ in discrete ill-posed problems that arise in science and engineering represents available data and typically is contaminated by a measurement error $\mathbf{e} \in \mathbb{R}^n$. For insightful discussions on linear discrete ill-posed problems, we refer to [12, 14].

Let $\mathbf{b}_{\text{exact}}$ denote the unknown error-free vector associated with the available right-hand side $\mathbf{b}$, i.e.,

$$\mathbf{b} = \mathbf{b}_{\text{exact}} + \mathbf{e}. \quad (2)$$

We would like to compute the solution $\mathbf{x}_{\text{exact}}$ of minimal Euclidean norm of (1) with $\mathbf{b}$ replaced by $\mathbf{b}_{\text{exact}}$. It can be expressed as

$$\mathbf{x}_{\text{exact}} = \mathbf{A}^\dagger \mathbf{b}_{\text{exact}}, \quad (3)$$

where $\mathbf{A}^\dagger$ denotes the Moore–Penrose pseudoinverse of $\mathbf{A}$. Note that the solution of minimal Euclidean norm of (1), which is given by $\mathbf{A}^\dagger \mathbf{b}$, generally is not a useful approximation of $\mathbf{x}_{\text{exact}}$ due to severe propagation of the error $\mathbf{e}$ into this solution. We have

$$\mathbf{A}^\dagger \mathbf{b} = \mathbf{x}_{\text{exact}} + \mathbf{A}^\dagger \mathbf{e},$$

and, typically, $\|\mathbf{x}_{\text{exact}}\| \ll \|\mathbf{A}^\dagger \mathbf{e}\|$. Throughout this paper, $\|\cdot\|$ denotes the Euclidean vector norm or the associated induced matrix norm.

The above discussions indicate that even when the linear system of (1) has a solution, one typically should not determine it. Instead, one should compute a suitable approximate solution of (1) that furnishes an accurate approximation of $\mathbf{x}_{\text{exact}}$.

It is natural to try to solve large-scale linear discrete ill-posed problems (1) by an iterative method. Discussions and analyses of a variety of iterative methods are provided by Meurant [17, 18]. One of the most popular iterative methods for computing an approximate solution of large-scale linear discrete ill-posed problems (1) is the LSQR iterative method; see, e.g., [12, 21] for descriptions of this method. When LSQR is used to solve linear discrete ill-posed problems, the initial iterate often is chosen to be $\mathbf{x}_0 = \mathbf{0}$. We will use this choice in the present paper. Then the k th step of the LSQR method generates an approximate solution $\mathbf{x}_k \in \mathbb{R}^n$ of (1) in the Krylov subspace

$$\mathcal{K}_k(\mathbf{A}^T \mathbf{A}, \mathbf{A}^T \mathbf{b}) = \text{span} \left\{ \mathbf{A}^T \mathbf{b}, (\mathbf{A}^T \mathbf{A}) \mathbf{A}^T \mathbf{b}, \dots, (\mathbf{A}^T \mathbf{A})^{k-1} \mathbf{A}^T \mathbf{b} \right\}, \quad (4)$$

where we assume $k \geq 1$ to be small enough so that $\dim(\mathcal{K}_k(A^T A, A^T \mathbf{b})) = k$. We will comment on this restriction below. The iterate $\mathbf{x}_k$ is characterized by

$$\mathbf{x}_k \in \mathcal{K}_k(A^T A, A^T \mathbf{b}) \quad \text{and} \quad \|\mathbf{A}\mathbf{x}_k - \mathbf{b}\| = \min_{\mathbf{x} \in \mathcal{K}_k(A^T A, A^T \mathbf{b})} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|; \quad (5)$$

see, e.g., [12, 21] for details on LSQR. It follows from (5) that LSQR is a minimal residual method in the sense that at step k , LSQR determines the element $\mathbf{x}_k \in \mathcal{K}_k(A^T A, A^T \mathbf{b})$ that minimizes the residual error $\mathbf{r}_k = \mathbf{b} - \mathbf{A}\mathbf{x}_k$. Clearly,

$$\|\mathbf{r}_{k-1}\| \geq \|\mathbf{r}_k\|$$

and, generally, this inequality is strict.

However, while the norm of the residual errors associated with iterates generated by LSQR decreases monotonically as k increases, the distance $\|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ generally does not. Typically, this distance decreases during the first few iterations, but increases during subsequent iterations. This behavior of the iterates $\mathbf{x}_k$ is commonly referred to as semi-convergence; see [10] for illustrations. We would like to terminate the iterations with LSQR when $\|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ is minimal.

When a sharp bound $\|\mathbf{e}\| \leq \delta$ is known and $\mathbf{b}_{\text{exact}}$ is in the range of A , the discrepancy principle can be used to determine how many iterations to carry out with LSQR. The discrepancy principle prescribes that the iterations with LSQR be terminates as soon as

$$\|\mathbf{A}\mathbf{x}_k - \mathbf{b}\| \leq \tau\delta,$$

where $\tau > 1$ is a user-chosen constant that is independent of δ . Note that $\|\mathbf{A}\mathbf{x}_{\text{exact}} - \mathbf{b}\| \leq \delta$.

We are interested in determining an iterate $\mathbf{x}_k$ that furnishes an as accurate approximation as possible of $\mathbf{x}_{\text{exact}}$ when no information about the norm of $\mathbf{e}$ is available. The determination of such an iterate is an important problem, because for many linear discrete ill-posed problems (1) that arise in applications, an accurate bound for $\|\mathbf{e}\|$ is not available. Methods for determining a suitable iterate $\mathbf{x}_k$ without using a bound for $\|\mathbf{e}\|$ are commonly referred to as “heuristic,” because they may fail for certain problems; see Kindermann [16] for an insightful discussion. Many heuristic methods have been proposed in the literature for determining a suitable iterate, or for the related problem of choosing an appropriate value of the regularization parameter in Tikhonov regularization. The methods include the L-curve criterion, generalized cross-validation, and extrapolation-based approaches; see [2–8, 11, 12, 22, 24–26] and references therein.

It is the purpose of the present paper to describe new heuristic stopping rules that are attractive to use with LSQR. Our interest in these rules stems from the fact that they are inexpensive to use and simple to implement. These rules belong to a new class of rules that seek to determine a suitable LSQR iterate $\mathbf{x}_k$ by comparing properties of iterates generated by two methods. Specifically, we will compare the residual errors associated with LSQR iterates with residual errors associated with iterates determined by Craig’s method. The k th iterate generated by Craig’s method lives in the Krylov subspace (4) and satisfies

$$\|\mathbf{x}_k - A^\dagger \mathbf{b}\| = \min_{\mathbf{x} \in \mathcal{K}_k(A^T A, A^T \mathbf{b})} \|\mathbf{x} - A^\dagger \mathbf{b}\|. \quad (6)$$

An implementation of Craig's method that is fairly similar to the standard implementation of LSQR is described by Saunders [27]; see also Paige [20]. We remark that methods that determine a suitable LSQR iterate, $\mathbf{x}_k$, by comparing it to approximate solutions computed by Tikhonov regularization, are discussed in [15, 22].

Craig's method is usually not used to compute approximate solutions of linear discrete ill-posed problems (1), because it is not a minimal residual method. However, as we will illustrate, the residual errors associated with iterates determined by Craig's method are helpful for deciding when to terminate the iterations with LSQR.

Section 2 discusses LSQR, describes Craig's method, and shows some properties that shed light on how these methods perform when applied to the solution of linear discrete ill-posed problems. An algorithm that implements both LSQR and Craig's method is described. The computational effort and storage requirement of this algorithm are essentially the same as for LSQR. A post-processing step that refines the choice of iterate by considering the distance between consecutive iterates also is introduced. Section 3 presents a few computed examples and Section 4 contains concluding remarks.

2 Novel stopping criteria

Both the LSQR and Craig iterative methods are based on Golub–Kahan bidiagonalization; see [20, 21, 27]. Application of k steps of Golub–Kahan bidiagonalization to the matrix A with initial vector $\mathbf{b}$ gives the decompositions

$$\begin{aligned} AV_k &= U_{k+1}B_k = U_kL_k + \beta_{k+1}\mathbf{u}_{k+1}\mathbf{e}_k^T, \\ A^T U_k &= V_kL_k^T, \end{aligned} \quad (7)$$

where the matrices $U_{k+1} = [\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{k+1}] \in \mathbb{R}^{n \times (k+1)}$ and $V_k = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k] \in \mathbb{R}^{n \times k}$ have orthonormal columns, the first column of U_{k+1} is $\mathbf{b}/\|\mathbf{b}\|$, and the matrix U_k is made up of the first k columns of U_{k+1} . The vector $\mathbf{e}_k$ is the k th column of an identity matrix of appropriate order and the superscript T denotes transposition. The columns of V_k form an orthonormal basis for the Krylov subspace (4). Finally, the matrix $B_k \in \mathbb{R}^{(k+1) \times k}$ is lower bidiagonal and the matrix $L_k \in \mathbb{R}^{k \times k}$ is made up of the first k rows and columns of B_k . Thus, L_k is lower bidiagonal and lower triangular. The computation of the decompositions (7) requires k matrix-vector product evaluations with A and k matrix-vector product evaluations with A^T . This is the dominating computational effort for computing the decompositions. We will assume that k is small enough so that the decompositions (7) exist. This is the generic situation. The computations simplify in the rare situation when the decompositions (7) only exist for some small value of k . We omit the details.

The decompositions (7) form the basis for the LSQR iterative method. The k th iterate determined by LSQR lives in the Krylov subspace (4) and can be represented

as $\mathbf{x}_k = V_k \mathbf{y}_k$ for some vector $\mathbf{y}_k \in \mathbb{R}^k$ that is determined by solving the small minimization problem on the right-hand side of

$$\|\mathbf{A}\mathbf{x}_k - \mathbf{b}\| = \min_{\mathbf{y} \in \mathbb{R}^k} \|A V_k \mathbf{y} - \mathbf{b}\| = \min_{\mathbf{y} \in \mathbb{R}^k} \|B_k \mathbf{y} - \mathbf{e}_1\| \|\mathbf{b}\|. \quad (8)$$

LSQR is a clever implementation of the solution of these minimization problems for increasing values of k . The method does not require storage of the whole matrices U_{k+1} and V_k ; only a few of the most recently generated columns of the matrices U_{k+1} and V_k in (7) have to be stored simultaneously.

We turn to Craig's method. The k th iterate determined by this method also can be expressed with the aid of the decompositions (7). Denote this iterate by $\widehat{\mathbf{x}}_k$. Then $\widehat{\mathbf{x}}_k = V_k \widehat{\mathbf{y}}_k$, where $\widehat{\mathbf{y}}_k \in \mathbb{R}^k$ satisfies

$$L_k \widehat{\mathbf{y}}_k = \|\mathbf{b}\| \mathbf{e}_1. \quad (9)$$

Algorithm 1, which is presented below, implements LSQR and the part of Craig's method that is of interest to us. Knowing the LSQR iterates $\mathbf{x}_k$ and the norms of the residual errors associated with the LSQR and Craig iterates, i.e., the norms of

$$\mathbf{r}_k = \mathbf{b} - \mathbf{A}\mathbf{x}_k, \quad \widehat{\mathbf{r}}_k = \mathbf{b} - \mathbf{A}\widehat{\mathbf{x}}_k, \quad (10)$$

allow us to determine when to stop the iterations with LSQR. Specifically, we terminate the iterations with Algorithm 1 as soon as

$$\frac{\|\widehat{\mathbf{r}}_k\|}{\|\mathbf{r}_k\|} \geq \delta \quad (11)$$

for a user-supplied parameter $\delta > 1$. This stopping criterion can be justified as follows: LSQR is a minimal residual method. The ratio (11) therefore is bounded below by one. The ratio typically is quite close to unity for k small, but increases with k because Craig's method is not a minimal residual method. Due to the property (6) of the iterates $\widehat{\mathbf{x}}_k$ generated by Craig's method, they will converge faster towards $A^\dagger \mathbf{b}$ than the iterates $\mathbf{x}_k$ computed by LSQR. Illustrations of this behavior can be found in [10]. This suggests that when the ratio (11) is significantly larger than unity, the iterates determined by Craig's method are contaminated substantially more by propagated error stemming from the error $\mathbf{e}$ in $\mathbf{b}$ than the corresponding iterates determined by LSQR. Numerous computed examples suggest that $\delta = 1.88$ in (11) to be a suitable choice. We will use this value in the computed examples reported in Section 3. Of course, other values of the constant δ also can be used.

Algorithm 1 computes the LSQR iterates $\mathbf{x}_k$, $k = 1, 2, \dots$, and the norm of the residual vectors (10). Each iteration of the algorithm requires the evaluation of one matrix-vector product with the matrix A and one matrix-vector product with the matrix A^T , just like the LSQR algorithm. Algorithm 1 requires essentially the same computational effort in each iteration as LSQR and furnishes the norms of the residual errors (10) that are used in the stopping criterion (11).

Algorithm 1 LSQR and Craig's methods run simultaneously.**Inputs:** $A \in \mathbb{R}^{n \times n}$, $\mathbf{b} \in \mathbb{R}^n$, $\mathbf{x}_0 \in \mathbb{R}^n$,**Output:** Approximate solution $\mathbf{x}_k$ of ill-posed problems.

- 1: Initialize: $\mathbf{x}_0 = \mathbf{0}$, $\beta_1 \mathbf{u}_1 = \mathbf{b}$, $\alpha_1 \mathbf{v}_1 = A^T \mathbf{u}_1$, $\mathbf{w}_1 = \mathbf{v}_1$, $\bar{\phi}_1 = \beta_1$, $\bar{\rho}_1 = \alpha_1$, $\tau_0 = 1$, $\varpi_0 = 0$, $\zeta_0 = -1$, $\widehat{\mathbf{v}}_0 = \mathbf{0}$, $\widehat{\mathbf{w}}_0 = \mathbf{0}$.
- 2: **for** $k = 1, 2, \dots$ until termination criterion satisfied **do**
- 3: Golub–Kahan bidiagonalization.
 - (1) $\beta_{k+1} \mathbf{u}_{k+1} = A \mathbf{v}_k - \alpha_k \mathbf{u}_k$
 - (2) $\alpha_{k+1} \mathbf{v}_{k+1} = A^T \mathbf{u}_{k+1} - \alpha_k \mathbf{v}_k$
- 4: LSQR method
 - (1) $\rho_k = (\bar{\rho}_k^2 + \beta_{k+1}^2)^{1/2}$
 - (2) $c_k = \bar{\rho}_k / \rho_k$
 - (3) $s_k = \beta_{k+1} / \rho_k$
 - (4) $\theta_{k+1} = s_k \alpha_{k+1}$
 - (5) $\bar{\rho}_{k+1} = -c_k \alpha_{k+1}$
 - (6) $\bar{\phi}_k = c_k \bar{\phi}_k$
 - (7) $\bar{\phi}_{k+1} = s_k \bar{\phi}_k$
 - (8) $\mathbf{x}_k = \mathbf{x}_{k-1} + (\bar{\phi}_k / \rho_k) \mathbf{w}_k$, $r_k = \bar{\phi}_{k+1}$
 - (9) $\mathbf{w}_{k+1} = \mathbf{v}_{k+1} - (\theta_{k+1} / \rho_k) \mathbf{w}_k$
- 5: Craig's method
 - (1) $\zeta_k = -\zeta_{k-1} \beta_k / \alpha_k$, $\widehat{\mathbf{v}}_k = \widehat{\mathbf{v}}_{k-1} + \zeta_k \mathbf{v}_k$
 - (2) $\varpi_k = (\tau_k - \beta_k \varpi_{k-1}) / \alpha_k$, $\widehat{\mathbf{w}}_k = \widehat{\mathbf{w}}_{k-1} + \varpi_k \mathbf{v}_k$
 - (3) **if** $(\beta_{k+1} = 0)$, **then**
 - (4) $\widehat{\mathbf{x}}_k = \widehat{\mathbf{v}}_k$, $\widehat{\mathbf{r}}_k = \mathbf{b} - A \widehat{\mathbf{x}}_k$;
 - (5) **else**
 - (6) $\tau_k = -\tau_{k-1} \alpha_k / \beta_{k+1}$
 - (7) **if** $(\alpha_{k+1} = 0)$, **then**
 - (8) $\gamma = \beta_{k+1} \zeta_k / (\beta_{k+1} \varpi_k - \tau_k)$
 - (9) $\widehat{\mathbf{x}}_k = \widehat{\mathbf{v}}_k - \gamma \cdot \widehat{\mathbf{w}}_k$,
 - (10) $\widehat{\mathbf{r}}_k = \mathbf{b} - A \widehat{\mathbf{x}}_k$;
 - (11) **end if**
- 6: Stopping criterion (11)
 - (1) **if** $\|\widehat{\mathbf{r}}_k / r_k\| \geq \delta$, **then**
 - (2) **return** $k = k$, $\mathbf{x}_{\widetilde{k}}$ **stop**
- 7: **end for**

The choice of the parameter k can be refined by a technique described by Morigi et al. [19], who proposed its application to the modification of the termination index obtained by the L-curve criterion. The sequence of iterates $\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2, \dots$ may be thought of as a semi-convergent sequence, whose limit we seek to determine. A

good approximation of the limit of a semi-convergent often can be found at a local minimum of the function

$$k \rightarrow \|\mathbf{x}_{k+1} - \mathbf{x}_k\|. \quad (12)$$

This lead Morigi et al. [19] to suggest that the iterate $\mathbf{x}_j$ determined by the L-curve criterion should be replaced by $\mathbf{x}_{\check{k}}$, where $\check{k}$ is an index close to j and the function (12) achieves a local minimum at $\check{k}$; see [19] for further details and many computed examples. We note that the approach used by Morigi et al. [19] is closely related to the quasi-optimality criterion, which is discussed and analyzed in, e.g., [2]. The latter criterion determines the index $k_{\min}$ that minimizes (12) and uses $\mathbf{x}_{k_{\min}}$ as an approximate solution of (1). The quasi-optimality criterion does not perform well in our experience, while the “localized quasi-optimality criterion” proposed in [19] performs better.

We will apply this technique to replace the iterate determined by the criterion (11). Thus, let the iterate $\mathbf{x}_{\check{k}}$ be determined by (11). Then we determine the index $\check{k}$ closest to $\tilde{k}$ such that the function (12) achieves a local minimum at $k = \check{k}$, and we use the LSQR iterate $\mathbf{x}_{\check{k}}$ as an approximation of $\mathbf{x}_{\text{exact}}$. If the index $\check{k}$ is not uniquely defined, then we let $\check{k}$ be the smallest of possible indices.

We conclude this section with a discussion on the conditioning of the matrices B_k and L_k in (7) that are used to compute the LSQR iterate $\mathbf{x}_k = V_k \mathbf{y}_k$ and the Craig iterate $\hat{\mathbf{x}}_k = V_k \hat{\mathbf{y}}_k$, respectively; cf. (8) and (9). We first describe how the singular values of the matrices B_k and L_k relate.

Proposition 1 *Let $\sigma_1^{(B_k)} \geq \sigma_2^{(B_k)} \geq \dots \geq \sigma_k^{(B_k)} \geq 0$ denote the singular values of the matrix B_k in decreasing order, and let $\sigma_1^{(L_k)} \geq \sigma_2^{(L_k)} \geq \dots \geq \sigma_k^{(L_k)} \geq 0$ denote the similarly ordered singular values of L_k . Then $\sigma_j^{(B_k)} \geq \sigma_j^{(L_k)}$ for $j = 1, 2, \dots, k$. Generically, the inequalities are strict.*

Proof Denote the $(k+1, k)$ -entry of B_k by β_{k+1} . Then

$$B_k^T B_k = L_k^T L_k + \beta_{k+1}^2 \mathbf{e}_k \mathbf{e}_k^T.$$

The eigenvalues of $B_k^T B_k$ are $(\sigma_j^{(B_k)})^2$, $1 \leq j \leq k$, and the eigenvalues of $L_k^T L_k$ are $(\sigma_j^{(L_k)})^2$, $1 \leq j \leq k$. The positive semidefinite matrix $B_k^T B_k$ is a rank-one modification of the positive semidefinite matrix $L_k^T L_k$. It follows that the eigenvalues of $B_k^T B_k$ interlace those of $L_k^T L_k$, i.e., $(\sigma_k^{(B_k)})^2 \geq (\sigma_k^{(L_k)})^2$ for $k = 1, 2, \dots$. Generically, the inequalities are strict; see, e.g., [29, pp. 94–97] for a proof. $\square$

It is interesting to investigate whether the matrix B_k is better conditioned than L_k . Define the condition numbers

$$\kappa(B_k) = \frac{\sigma_1^{(B_k)}}{\sigma_k^{(B_k)}}, \quad \kappa(L_k) = \frac{\sigma_1^{(L_k)}}{\sigma_k^{(L_k)}}. \quad (13)$$

We would like to determine whether

$$\kappa(B_k) \leq \kappa(L_k). \quad (14)$$

Let $\sigma_1^{(B_k)} = \sigma_1^{(L_k)} + \epsilon_1$ and $\sigma_k^{(B_k)} = \sigma_k^{(L_k)} + \epsilon_k$. Assume that $\sigma_k^{(B_k)} > 0$. Then (14) is equivalent to

$$\kappa(L_k) \geq \frac{\epsilon_1}{\epsilon_k}.$$

In applications of Golub–Kahan bidiagonalization to the matrix A of a linear discrete ill-posed problem, the condition numbers $\kappa(L_k)$ may be large. In order for the inequality (14) to hold, $\epsilon_k > 0$ cannot be arbitrarily much smaller than ϵ_1 . Proposition 1 does not furnish information about the relative size of ϵ_1 and ϵ_k . Example 5 of Section 3 illustrates that for several linear discrete ill-posed problems (1) and bidiagonalization steps k , the inequality (14) holds. This property is in agreement with the observation that the error $\mathbf{e}$ in $\mathbf{b}$ is propagated less severely into the computed approximate solution when using LSQR than when applying the same number of steps with Craig’s method.

3 Numerical examples

This section illustrates the stopping rule (11) when applied to the solution of several linear discrete ill-posed problems from the MATLAB program package Regularization Tools by Hansen [13]. All computations are carried out using MATLAB R2016b with unit round-off $\varepsilon_{\text{machine}} \approx 2.22 \cdot 10^{-16}$. A Lenovo laptop computer running Windows 10 with 4.87 GB of RAM was used.

Each code from Regularization Tools [13] provides a matrix $A \in \mathbb{R}^{n \times n}$ that is the discretization of a Fredholm integral equation of the first kind, a vector $\mathbf{b}_{\text{exact}}$ that represents the error-free right-hand side, and the desired solution $\mathbf{x}_{\text{exact}}$ given by (3). The noise vector $\mathbf{e} \in \mathbb{R}^n$ is in all examples chosen to have normally distributed random entries with zero mean. The vector is normalized to correspond to a specific noise level

$$\epsilon = \frac{\|\mathbf{e}\|}{\|\mathbf{b}_{\text{exact}}\|}. \quad (15)$$

In the computed examples, we let $\epsilon \in \{10^{-1}, 10^{-2}, 10^{-3}\}$. Noise levels of these orders of magnitude occur in many real-world problems. The error-contaminated right-hand side in (1) is given by (2).

The computed indices are compared with the index k_{best} , which corresponds to an iterate that is closest to $\mathbf{x}_{\text{exact}}$. Thus,

$$k_{\text{best}} = \arg \min_{1 \leq k \leq \ell} \|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|, \quad (16)$$

where ℓ is chosen large enough to secure that k_{best} is the smallest global minimum. Clearly, k_{best} only can be computed when $\mathbf{x}_{\text{exact}}$ is known and therefore can be determined for certain test problems, only. We mark the points corresponding to the

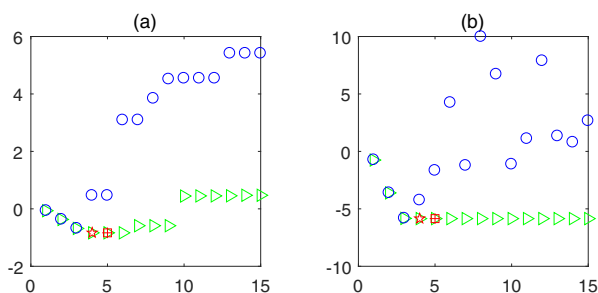


Fig. 1 Example 1. **a** Graph of $\log_{10} \|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ vs. iteration number k for $1 \leq k \leq 15$, and **b** $\log_{10} \|\mathbf{A}\mathbf{x}_k - \mathbf{b}\|$ vs. k for $1 \leq k \leq 15$; $\blacktriangleright$ indicates results for LSQR and $\bigcirc$ stands for the results for Craig's method; $\star$ denotes $\tilde{k}$, $\square$ denotes $\hat{k}$, and $+$ denotes k_{best}

approximate solutions $\mathbf{x}_{\text{best}}$, $\mathbf{x}_{\tilde{k}}$, and $\mathbf{x}_{\hat{k}}$ in Figs. 1, 2, 4, and 5 below. The index $\tilde{k}$ is chosen such that

$$\|\mathbf{x}_{\tilde{k}+1} - \mathbf{x}_{\tilde{k}}\| = \min_{\hat{k} \leq k \leq \tilde{k}(3)} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|, \quad (17)$$

where $\hat{k} = \max(2, \tilde{k} - 3)$ and $\tilde{k}(3)$ denotes the third smallest index $\tilde{k}$ that satisfies (11). Note that $\tilde{k}(3)$ may be larger than $\tilde{k} + 3$. The choice of $\hat{k}$ and $\tilde{k}(3)$ is not critical; the purpose of these parameters is to define an interval around $\tilde{k}$, where to apply the localized quasi-optimality criterion. The matrices A in all examples are of size 500×500 .

Example 1 Consider the Fredholm integral equation of the first kind

$$\int_0^\pi \exp(s \cdot \cos(t))x(t)dt = 2 \frac{\sin(s)}{s}, \quad 0 \leq s \leq \frac{\pi}{2},$$

which is discussed by Baart [1]. It has the solution $x(t) = \sin(t)$. The integral equation is discretized by a Galerkin method with piece-wise constant test and trial functions using the function `baart` from [13].

Figure 1 shows the error and residual error histories for the LSQR iterates determined when the noise level is $\epsilon = 1 \cdot 10^{-3}$. The error norms $\|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ for iterates $\mathbf{x}_k$ determined by LSQR are marked by green triangles ($\blacktriangleright$), and the error norms for iterates computed by Craig's method are marked by blue circles ($\bigcirc$). The red square ($\square$) marks the iterate determined by the criterion (11) and (17), the red star ($\star$) denotes the iterate $\mathbf{x}_{\tilde{k}}$ determined by (11), and the red plus sign ($+$) indicates the iterate $\mathbf{x}_{\text{best}}$, whose index is defined in (16). Thus, we can observe from Fig. 1 that the termination index determined by (11) is $\tilde{k} = 4$ and the closest local minimum of the function (17) is at $\tilde{k} = 5$, which is equal to the index k_{best} . Both LSQR iterate $\mathbf{x}_4$ and $\mathbf{x}_5$ are accurate approximations of $\mathbf{x}_{\text{best}}$. Details are reported in Table 1.

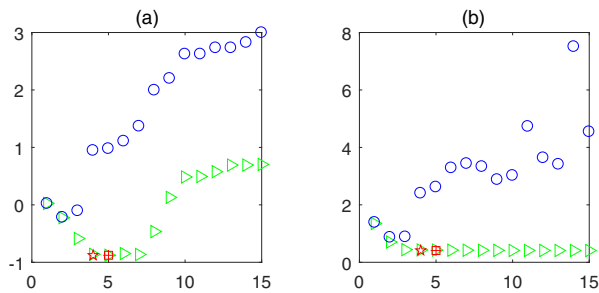


Fig. 2 Example 2. **a** Graph of $\log_{10} \|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ vs. iteration number k for $1 \leq k \leq 15$, and **b** $\log_{10} \|A\mathbf{x}_k - \mathbf{b}\|$ vs. k for $1 \leq k \leq 15$; $\blacktriangleright$ denote results for LSQR and $\circ$ denote results for Craig's method; $\star$ denotes $\tilde{k}$, $\square$ denotes $\tilde{k}$, and $+$ denotes k_{best}

Example 2 Regard the Fredholm integral equation of the first kind discussed by Phillips [23],

$$\int_{-6}^6 K(s, t)x(t)dt = g(s), \quad -6 \leq s \leq 6,$$

whose solution $x(t)$, kernel $K(s, t)$, and right-hand side $g(s)$ are given by

$$x(t) = \begin{cases} 1 + \cos\left(\frac{\pi t}{3}\right), & |t| < 3, \\ 0, & |t| \geq 3, \end{cases}$$

$$K(s, t) = x(s - t),$$

$$g(s) = (6 - |s|) \left(1 + \frac{1}{2} \cos\left(\frac{\pi s}{3}\right)\right) + \frac{9}{2\pi} \sin\left(\frac{\pi |s|}{3}\right).$$

This integral equation is discretized by a Galerkin method using the MATLAB function `phillips` from [13].

The numerical results are listed in Table 1. The index selection criteria based on (11) and (17) can be seen to perform better for the noise level $\epsilon = 1 \cdot 10^{-1}$ than for the noise levels $\epsilon \leq 1 \cdot 10^{-2}$. The error and residual norms versus the number of iteration steps are shown in Fig. 2 for the former noise level.

Figure 3 displays for $\epsilon = 1 \cdot 10^{-1}$ the function (12). The local minimum is achieved at $k = 5$, which agrees with the index determined by (11) and (17).

Example 3 Consider the Fredholm integral equation of the first kind discussed by Shaw [28],

$$\int_{-\frac{\pi}{2}}^{\frac{\pi}{2}} K(s, t)x(t)dt = g(s), \quad -\frac{\pi}{2} \leq s \leq \frac{\pi}{2},$$

with kernel

$$K(s, t) = (\cos(s) + \cos(t))^2 \left(\frac{\sin(\pi(\sin(s) + \sin(t)))}{\pi(\sin(s) + \sin(t))} \right)^2$$

and solution

$$x(t) = 2 \exp\left(-6(t - 0.8)^2\right) + \exp(-2(t + 0.5)^2),$$

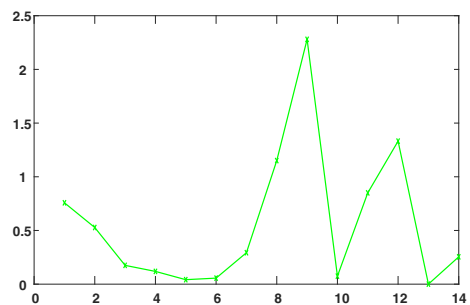
Table 1 Numerical results of tested examples

Problem	Noise level	$\tilde{k}$	$E_{\tilde{k}}$	$\check{k}$	$E_{\check{k}}$	k_{best}	$E_{k_{\text{best}}}$
baart (15)	$1 \cdot 10^{-1}$	3	0.3549	4	1.6059	2	0.3451
baart (15)	$1 \cdot 10^{-2}$	4	0.2325	4	0.2325	3	0.1671
baart (15)	$1 \cdot 10^{-3}$	4	0.1169	5	0.1168	5	0.1168
deriv2 (15)	$1 \cdot 10^{-1}$	2	0.4380	3	0.3526	4	0.3240
deriv2 (15)	$1 \cdot 10^{-2}$	5	0.2664	6	0.2381	9	0.2230
deriv2 (20)	$1 \cdot 10^{-3}$	7	0.2279	6	0.2279	19	0.1490
foxgood (15)	$1 \cdot 10^{-1}$	2	0.0332	4	2.0687	2	0.0332
foxgood (15)	$1 \cdot 10^{-2}$	3	0.0089	3	0.0089	3	0.0089
foxgood (15)	$1 \cdot 10^{-3}$	3	0.0074	3	0.0074	5	0.0061
gravity (15)	$1 \cdot 10^{-1}$	2	0.1370	2	0.1370	5	0.0636
gravity (15)	$1 \cdot 10^{-2}$	6	0.0418	7	0.0352	8	0.0346
gravity (15)	$1 \cdot 10^{-3}$	7	0.0399	7	0.0399	11	0.0202
heat (20)	$1 \cdot 10^{-1}$	5	0.3465	7	0.2439	10	0.1931
heat (30)	$1 \cdot 10^{-2}$	9	0.1855	9	0.1855	17	0.0686
heat (35)	$1 \cdot 10^{-3}$	10	0.1741	16	0.0709	20	0.0227
phillips (15)	$1 \cdot 10^{-1}$	4	0.0458	5	0.0442	5	0.0442
phillips (20)	$1 \cdot 10^{-2}$	5	0.0256	5	0.0256	7	0.0254
phillips (20)	$1 \cdot 10^{-3}$	5	0.0243	4	0.0244	9	0.0087
shaw (15)	$1 \cdot 10^{-1}$	4	0.1720	6	0.6932	4	0.1720
shaw (15)	$1 \cdot 10^{-2}$	5	0.1033	6	0.0536	6	0.0536
shaw (15)	$1 \cdot 10^{-3}$	7	0.0598	9	0.0480	9	0.0480

which define the right-hand side function g . Discretization is carried out by a Nyström method based on the midpoint quadrature rule using the function shaw from [13].

The stopping criteria perform well for this example for all noise levels. Figure 4 displays the numerical results for noise level $\epsilon = 1 \cdot 10^{-2}$. The error in $\mathbf{x}_{\tilde{k}}$ with $\tilde{k} = 6$ is smaller than that in $\mathbf{x}_{\check{k}}$ with $\check{k} = 5$, where the former is the same as in $\mathbf{x}_{k_{\text{best}}}$. Detailed results are shown in Table 1.

Fig. 3 Example 2. Graph of $\|\mathbf{x}_{k+1} - \mathbf{x}_k\|$ vs. iteration number k for $1 \leq k \leq 14$



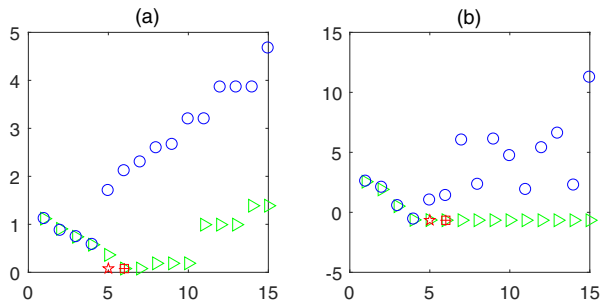


Fig. 4 Example 3. **a** Graph of $\log_{10} \|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ vs. iteration number k for $1 \leq k \leq 15$, and **b** $\log_{10} \|\mathbf{A}\mathbf{x}_k - \mathbf{b}\|$ vs. k for $1 \leq k \leq 15$. The symbol $\blacktriangleright$ denotes results for LSQR and $\circ$ denotes results for Craig's method; $\star$ denotes $\tilde{k}$, $\square$ denotes $\check{k}$, and $+$ denotes k_{best}

Example 4 We consider a discretization of the Fredholm integral equation of the first kind

$$\int_0^1 (s^2 + t^2)^{1/2} x(t) dt = \frac{1}{3}(1 + s^2)^{3/2} - s^3, \quad 0 \leq s \leq 1,$$

with solution $x(t) = t$. This equation is discussed by Fox and Goodwin [9]. We use the function `foxgood` from [13] to determine a discretization by a Nyström method.

Figure 5 illustrates the performance of the stopping rules for the noise level $\epsilon = 1 \cdot 10^{-2}$. In particular, Fig. 5a shows that the criteria (11) and (17) determine the same index for approximate solutions, which is equal to k_{best} . Further details can be found in Table 1.

Example 5 Figure 6 displays the condition numbers (13) for several linear discrete ill-posed problems and different matrix sizes. The figure shows the inequality (14) to hold. This is true for many linear discrete ill-posed problems and various numbers of Golub–Kahan bidiagonalization steps.

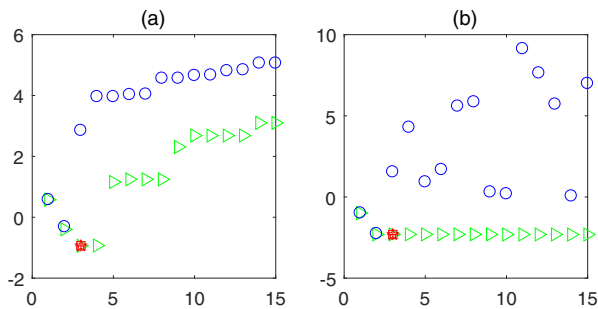


Fig. 5 Example 4. **a** Graph of $\log_{10} \|\mathbf{x}_k - \mathbf{x}_{\text{exact}}\|$ vs. iteration number k for $1 \leq k \leq 15$, and **b** $\log_{10} \|\mathbf{A}\mathbf{x}_k - \mathbf{b}\|$ vs. k for $1 \leq k \leq 15$. The symbol $\blacktriangleright$ denotes results for LSQR and $\circ$ denotes results for Craig's method; $\star$ shows the index $\tilde{k}$, $\square$ the index $\check{k}$, and $+$ the index k_{best}

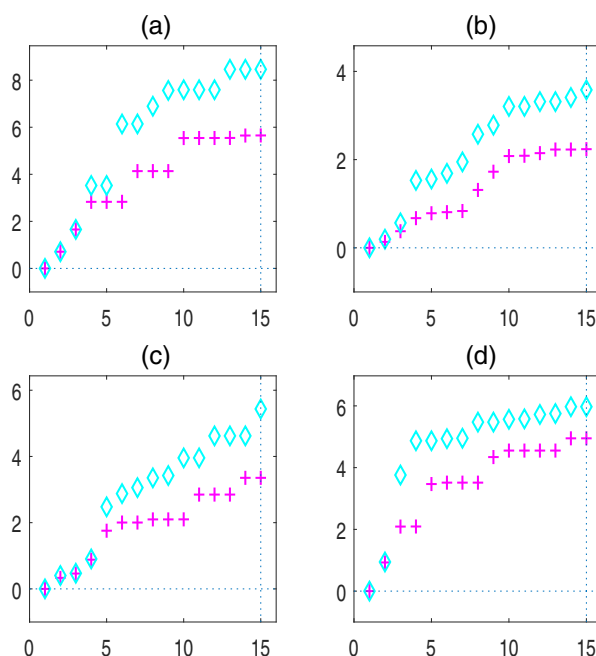


Fig. 6 Graphs of $\log_{10}(\kappa(L_k))$ ($\diamond$) and $\log_{10}(\kappa(B_k))$ ($+$) vs. the number of iteration steps k . **a** Example 1, $\epsilon = 1 \cdot 10^{-3}$. **b** Example 2, $\epsilon = 1 \cdot 10^{-1}$. **c** Example 3, $\epsilon = 1 \cdot 10^{-2}$. **d** Example 4, $\epsilon = 1 \cdot 10^{-2}$

Table 1 summarizes the numerical results obtained with the stopping criteria of the present paper. The column labeled “Problem” displays the names of problems from [13] and contains in parentheses the maximum number of iterations carried out. The column “Noise level” shows the relative amount of error in the right-hand side **b**; see (15). The table also displays the indices $\tilde{k}$, $\hat{k}$, and k_{best} , defined by (11), (17), and (16), respectively. In addition, columns 5, 7, and 9 of Table 1 show the relative errors

$$E_{\tilde{k}} = \frac{\|\mathbf{x}_{\tilde{k}} - \mathbf{x}_{\text{exact}}\|}{\|\mathbf{x}_{\text{exact}}\|}, \quad E_{\hat{k}} = \frac{\|\mathbf{x}_{\hat{k}} - \mathbf{x}_{\text{exact}}\|}{\|\mathbf{x}_{\text{exact}}\|}, \quad E_{k_{\text{best}}} = \frac{\|\mathbf{x}_{k_{\text{best}}} - \mathbf{x}_{\text{exact}}\|}{\|\mathbf{x}_{\text{exact}}\|}.$$

The table shows that for almost all examples, the approximate solution $\mathbf{x}_{\tilde{k}}$ is a good choice.

4 Conclusion

This work proposes novel approaches for terminating the iterations with the LSQR iterative method. The stopping criteria described compare the norm of the residual errors associated with iterates determined by the LSQR and Craig’s methods at the same number of steps. A refinement of this strategy also is described. The attractions of the new stopping rules are their ease of implementation, and low computing and storage costs.

Acknowledgments The authors would like to thank the referees for comments that lead to improvements of the presentation. The first author would like to thank the second author for an enjoyable visit to the Université du Littoral, during which this work was initiated. Research by the first author was supported in part by NSF grants DMS-1720259 and DMS-1729509. Research by the third author was supported by the China Scholarship Council (CSC No. 201806180081).

References

1. Baart, M.L.: The use of auto-correlation for pseudorank determination in noisy ill-conditioned linear least-squares problems. *IMA J. Numer. Anal.* **2**, 241–247 (1982)
2. Bauer, F., Kindermann, S.: The quasi-optimality criterion for classical inverse problems. *Inverse Problems* **24**, 035002 (2008). (20pp)
3. Bauer, F., Lukas, M.A.: Comparing parameter choice methods for regularization of ill-posed problems. *Math. Comput. Simulation* **81**, 1795–1841 (2011)
4. Bouhamidi, A., Jbilou, K., Reichel, L., Sadok, H., Wang, Z.: Vector extrapolation applied to truncated singular value decomposition and truncated iteration. *J. Eng. Math.* **93**, 99–112 (2015)
5. Brezinski, C., Redivo Zaglia, M., Rodriguez, G., Seatzu, S.: Extrapolation techniques for ill-conditioned linear systems. *Numer. Math.* **81**, 1–29 (1998)
6. Brezinski, C., Rodriguez, G., Seatzu, S.: Error estimates for linear systems with applications to regularization. *Numer. Algorithms* **49**, 85–104 (2008)
7. Brezinski, C., Rodriguez, G., Seatzu, S.: Error estimates for the regularization of least squares problems. *Numer. Algorithms* **51**, 61–76 (2009)
8. Fenu, C., Reichel, L., Rodriguez, G., Sadok, H.: GCV for Tikhonov regularization by partial SVD. *BIT Numer. Math.* **57**, 1019–1039 (2017)
9. Fox, L., Goodwin, E.T.: The numerical solution of non-singular linear integral equations. *Philos. Trans. Roy. Soc. London. Ser. A.* **245**, 501–534 (1953)
10. Hanke, M.: On Lanczos based methods for the regularization of discrete ill-posed problems. *BIT Numer. Math.* **41**, 1008–1018 (2001)
11. Hansen, P.C.: Analysis of discrete ill-posed problems by means of the L-curve. *SIAM Rev.* **34**, 561–580 (1992)
12. Hansen, P.C.: Rank-deficient and discrete ill-Posed problems: Numerical aspects of linear inversion. SIAM, Philadelphia (1998)
13. Hansen, P.C.: Regularization Tools version 4.0 for Matlab 7.3. *Numer. Algorithms* **46**, 189–194 (2007)
14. Hansen, P.C.: Discrete inverse problems: insight and algorithms. SIAM, Philadelphia (2010)
15. Hochstenbach, M.E., Reichel, L., Rodriguez, G.: Regularization parameter determination for discrete ill-posed problems. *J. Comput. Appl. Math.* **273**, 132–149 (2015)
16. Kindermann, S.: Convergence analysis of minimization-based noise level-free parameter choice rules for linear ill-posed problems. *Electron. Trans. Numer. Anal.* **38**, 233–257 (2011)
17. Meurant, G.: Computer solution of large linear systems. Elsevier, Amsterdam (1999)
18. Meurant, G.: The Lanczos and conjugate gradient algorithms. SIAM, Philadelphia (2006)
19. Morigi, S., Reichel, L., Sgallari, F., Zama, F.: Iterative methods for ill-posed problems and semiconvergent sequences. *J. Comput. Appl. Math.* **193**, 157–167 (2006)
20. Paige, C.C.: Bidiagonalization of matrices and solutions of linear equations. *SIAM J. Numer. Anal.* **11**, 197–209 (1974)
21. Paige, C.C., Saunders, M.A.: LSQR: An algorithm for sparse linear equations and sparse least squares. *ACM Trans. Math. Soft.* **8**, 43–71 (1982)
22. Park, Y., Reichel, L., Rodriguez, G., Yu, X.: Parameter determination for Tikhonov regularization problems in general form. *J. Comput. Appl. Math.* **343**, 12–25 (2018)
23. Phillips, D.L.: A technique for the numerical solution of certain integral equations of the first kind. *J. Assoc. Comput. Mach.* **9**, 84–97 (1962)
24. Regińska, T.: A regularization parameter in discrete ill-posed problems. *SIAM J. Sci. Comput.* **17**, 740–749 (1996)
25. Reichel, L., Rodriguez, G.: Old and new parameter choice rules for discrete ill-posed problems. *Numer. Algorithms* **63**, 65–87 (2013)

26. Reichel, L., Sadok, H.: A new L-curve for ill-posed problems. *J. Comput. Appl. Math.* **219**, 493–508 (2008)
27. Saunders, M.A.: Solution of sparse rectangular systems using LSQR and Craig. *BIT Numer. Math.* **35**, 588–604 (1995)
28. Shaw, C.B. Jr.: Improvement of the resolution of an instrument by numerical solution of an integral equation. *J. Math. Anal. Appl.* **37**, 83–112 (1972)
29. Wilkinson, J.H.: The algebraic eigenvalue problem. Oxford University Press, New York (1965)

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Affiliations

Lothar Reichel¹ · Hassane Sadok² · Wei-Hong Zhang^{1,3}

Hassane Sadok
Hassane.Sadok@lmpa.univ-littoral.fr

Wei-Hong Zhang
zhangwh13@lzu.edu.cn

¹ Department of Mathematical Sciences, Kent State University, Kent, OH, 44242, USA

² Laboratoire de Mathématiques Pures et Appliquées, Université du Littoral, Centre Universitaire de la Mi-Voix, Batiment H. Poincaré, 50 Rue F. Buisson, BP 699, 62228, Calais cedex, France

³ School of Mathematics and Statistics, Lanzhou University, Lanzhou, 730000, People's Republic of China