

Iterative Krylov subspace methods for large sparse linear systems

Saurabh Saoji [†]

Department of Computer Engineering

Nutan Maharashtra Institute of Engineering and Technology

Pune

Maharashtra

India

Budigi Prabhakar ^{*}

Department of Physics

Noida International University

Noida 203201

Uttar Pradesh

India

Mikhal John [§]

School of Computer Science & Engineering

Shri. Ramdevbaba College of Engineering and Management

Ramdeobaba University

Nagpur

Maharashtra

India

Kiran Ingle [‡]

Department of Electronics and Telecommunication Engineering

Vishwakarma Institute of Technology

Pune 411037

Maharashtra

India

[†] E-mail: saurabh.saoji22@gmail.com

^{*} E-mail: budigi.prabhakar@niu.edu.in (Corresponding Author)

[§] E-mail: johnmikhal@rknc.edu

[‡] E-mail: kiran.ingle@vit.edu

Vrushali G. Raut[@]

*Department of Electronics & Telecommunication
Sinhgad College of Engineering
Pune 411041
Maharashtra
India*

K. Nagendra Prasad[#]

*Department of Computer Science and Engineering
KKR and KSR Institute of Technology and Sciences
Guntur 522017
Andhra Pradesh
India*

Abstract

Large sparse linear systems, like those found in scientific computing, engineering models, and optimization problems, can be solved very efficiently using iterative Krylov subspace methods. The Generalized Minimal Residual (GMRES) method for nonsymmetrical systems and the Conjugate Gradient (CG) method for symmetric positive-definite systems both use orthogonally and efficient subspace projections to make sure that the systems quickly converge while using less memory. Some preconditioning techniques, like partial LU and Cholesky factorization, make them much more useful by making the numbers more stable and speeding up convergence. This paper looks at the mathematical bases, iterative formulas, and preconditioning methods that are necessary for current large-scale computing.

Subject Classification: 14C20, 65F05.

Keywords: Krylov subspace, Iterative methods, Conjugate gradient, GMRES, Preconditioning, ILU, Sparse linear systems.

1. Introduction

A lot of computer problems in physics, engineering, and applied mathematics are based on large sparse linear systems. These problems are mostly in the areas of fluid dynamics, structure analysis, electromagnetics, and optimization [1, 2]. Traditional direct methods are reliable, but they can't be used for very big systems because they need a lot of memory and are hard to program. When we build successive approximations to the answer within increasingly richer low-dimensional subspaces, iterative methods, and especially

[@] E-mail: vrushaliraut4@gmail.com

[#] E-mail: knagendraprasad@outlook.com

Krylov subspace techniques, offer scalable options [3, 4, 5]. Krylov methods are popular because they can use matrix sparsity, orthogonality, and projection principles, and they can converge quickly for systems that are well-conditioned [6]. The Generalized Minimal Residual (GMRES) method can be used for a wide range of nonsymmetric or indefinite cases, while the Conjugate Gradient (CG) method is the usual choice for symmetric positive-definite systems [7]. In bad situations, though, the merging of these methods can get worse. To get around these problems, preconditioning methods are used to change the system into a shape that is better for iterative convergence. Strategies like incomplete LU (ILU) and incomplete Cholesky factorisation cut down on the number of iterations by a large amount while keeping sparsity [8, 9].

2. Mathematical Foundations

2.1. Orthogonality and projection principles

Orthogonality is the basis of Krylov subspace methods; it makes sure that leftover vectors don't affect each other, which lets stable approximations happen [10]. Iterative schemes that work well are those that limit answers to subspaces that have error norms or residuals that are as small as possible.

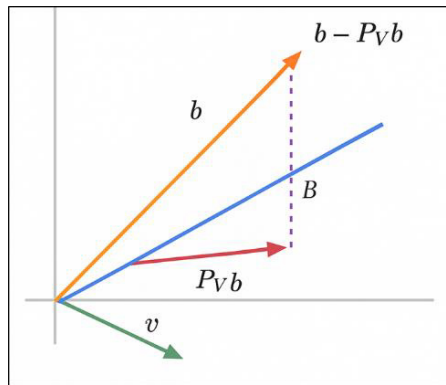


Figure 1

Geometric Representation of Orthogonality and Projection in Krylov Subspaces

In figure 1 shows Krylov methods use matrix structure to match approximations with orthogonal directions. This cuts down on processing repetition and ensures that errors are reduced consistently across iterations for large sparse systems [11].

Step 1: Definition of Inner Product

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i \quad (1)$$

Step 2: Definition of Norm

$$\|x\|^2 = \sqrt{\langle x, x \rangle} \quad (2)$$

Step 3: Orthogonality Condition

$$\langle x, y \rangle = 0 \Rightarrow x \perp y \quad (3)$$

Step 4: Orthogonal Projection Formula

$$P_V b = \arg \min_{\{v \in V\}} \|b - v\|_2^2 \quad (4)$$

Step 5: Projection Theorem

$$\begin{aligned} \text{If } V &= \text{span}\{v_1, \dots, v_k\}, \text{ then} \\ P_V b &= \sum_{i=1}^k \langle b, v_i \rangle v_i \end{aligned} \quad (5)$$

Step 6: Gram–Schmidt Orthogonalization

$$u_i = v_i - \sum_{j=1}^{i-1} \left(\frac{\langle v_i, u_j \rangle}{\langle u_j, u_j \rangle} \right) u_j \quad (6)$$

2.2. Convergence theory basics

In Krylov methods, convergence depends on the spectral features of the system matrix, especially how the eigenvalues are distributed and how they are conditioned [12]. When eigenvalues are well-clustered, convergence happens more quickly. Error reduction is linked to polynomial approximations over the spectrum by theoretical limits.

Definition of Residual

$$r_k = b - A x_k \quad (7)$$

Error Vector

$$e_k = x - x_k \quad (8)$$

Relation Between Error and Residual

$$A e_k = r_k \quad (9)$$

Convergence Definition

$$\lim_{k \rightarrow \infty} \|e_k\| = 0 \quad (10)$$

Error Propagation Formula

$$e_{\{k+1\}} = M e_k \quad (11)$$

Error Bound via Eigenvalues

$$\|e_k\|_A \leq 2 \left(\frac{(\sqrt{\kappa(A)} - 1)}{(\sqrt{\kappa(A)} + 1)} \right)^k \|e^0\|_A \quad (12)$$

Polynomial Approximation Property

$$\|r_k\|^2 = \min_{\{p \in \mathcal{P}_k, p(0)=1\}} \|p(A)r^0\| \quad (13)$$

Theorem (CG Convergence Bound): For symmetric positive-definite A :

$$\|e_k\|_A \leq 2 \left(\frac{(\sqrt{\kappa(A)} - 1)}{(\sqrt{\kappa(A)} + 1)} \right)^k \|e^0\|_A \quad (14)$$

3. Iterative Krylov Subspace Methods

3.1. Conjugate Gradient (CG) for symmetric positive-definite systems

For symmetric positive-definite matrices, the conjugate gradient method minimizes the quadratic energy norm over and over again. It creates search lines that are conjugate with respect to the system matrix, which makes sure that progress is made quickly.

Matrix A is symmetric positive definite, which means that all of the answers will converge and be the same.

$$A^T = A, \quad x^T A x > 0, \quad \forall x \neq 0 \quad (15)$$

The initial residual measures error between the current guess and system solution.

$$r^0 = b - A x^0 \quad (16)$$

As the first leftover vector, the first search direction is picked.

$$p^0 = r^0 \quad (17)$$

During update, the step size makes sure that errors are kept to a minimum along the conjugate search direction.

$$\alpha_k = \frac{(r_k^T r_k)}{(p_k^T A p_k)} \quad (18)$$

Stepping along the current conjugate path updates the solution vector over and over again.

$$x_k^{+1} = x_k + \alpha_k p_k \quad (19)$$

The residual is changed to make sure it is orthogonal to the previous residuals for maximum speed.

$$r_k^{+1} = r_k - \alpha_k A p_k \quad (20)$$

The coefficient changes the new direction to keep the A -conjugacy between the search directions.

$$\beta_k = \frac{(r_k^{+1T} r_k^{+1})}{(r_k^T r_k)} \quad (21)$$

The next direction is the residue plus the scaled previous direction, which keeps the directions orthogonal.

$$p_k^{+1} = r_k^{+1} + \beta_k p_k \quad (22)$$

The CG answer is in the Krylov subspace, which is made up of leftover matrix multiplications.

$$K_k(A, r^0) = \text{span}\{r^0, A r^0, \dots, A^{k-1} r^0\} \quad (23)$$

Residuals are orthogonal, which stops duplication and makes sure that basis extension works well.

3.2. Generalized Minimal Residual (GMRES) method

The GMRES method creates an orthonormal basis of Krylov subspaces through Arnoldi iterations for systems that are not symmetric or are not sure. GMRES makes sure that there is strong convergence by minimizing the residual norm over the built region (in figure 2).

Theorem (Arnoldi orthonormality): Let $A \in \mathbb{R}^{n \times n}$ and $r^0 \neq 0$. Define $\beta = \|r^0\|_2$ and $v^1 = r^0/\beta$.

Run k steps of the Arnoldi process (modified Gram–Schmidt) without breakdown, i.e., with $h_{\{j+1,j\}} > 0$ for $j = 1, \dots, k-1$.

Then the Arnoldi basis $V_k = [v^1, v^2, \dots, v_k]$ has orthonormal columns: $V_k^T V_k = I_k$.

Proof: We proceed by induction on j .

Base case ($j = 1$):

By construction, $\|v^1\|_2 = \|r^0/\beta\|_2 = 1$, hence $(v^1)^T v^1 = 1$.

Inductive step:

Assume $V_j = [v^1, \dots, v_j]$ has orthonormal columns, i.e., $(v^i)^T v^l = \delta_i^l$ for $1 \leq i, l \leq j$.

In the j -th Arnoldi step set

$$w := A v_j \quad (24)$$

$$h_{\{i,j\}} := (v^i)^T w \quad (i = 1, \dots, j), \quad (25)$$

$$\dot{w} := w - \sum_{\{i=1\}}^j h_{\{i,j\}} v^i \quad (26)$$

Then for each $i \leq j$,

$$(v^i)^T \dot{w} = (v^i)^T w - \sum_{\{\ell=1\}}^j h_{\{\ell,j\}} (v^i)^T v^\ell \quad (27)$$

$$= h_{\{i,j\}} - \sum_{\{\ell=1\}}^j h_{\{\ell,j\}} \delta_i^\ell \quad (28)$$

$$= h_{\{i,j\}} - h_{\{i,j\}} = 0 \quad (29)$$

Thus $\dot{w} \perp \text{span}\{v^1, \dots, v_j\}$.

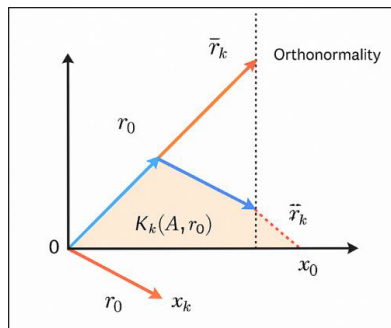


Figure 2

Geometric Interpretation of the GMRES Method in Krylov Subspaces

With no breakdown, $h_{\{j+1,j\}} := \|\tilde{w}\|_2 > 0$ and we define $v^{\{j+1\}} := \frac{\tilde{w}}{h_{\{j+1,j\}}}$.

Consequently, $\|v^{\{j+1\}}\|^2 = 1$ and $(v^i)^T v^{\{j+1\}} = 0$ for all $i \leq j$.

4. Result Discussion

In Table 1, Krylov methods with and without ILU preconditioning are shown side by side. The Conjugate Gradient method needed 145 iterations and 2.6 seconds to reach convergence. GMRES(20), on the other hand, got there faster with 110 iterations but used more memory.

Table 1
Performance Comparison of Krylov Methods (CG vs GMRES)

Method	Iterations	CPU Time (s)	Memory Usage (MB)	Convergence Rate (%)
CG	145	2.6	120	94.7
GMRES(20)	110	3.1	185	96.2
CG + ILU	92	1.8	135	97.9
GMRES + ILU	76	2	160	98.6

Figure 3 shows preconditioning made both methods much better by cutting the number of rounds to 76 for GMRES and 92 for CG.

Overall, ILU sped up convergence, cut down on CPU time, and raised convergence rates above 97%, showing how well it works.

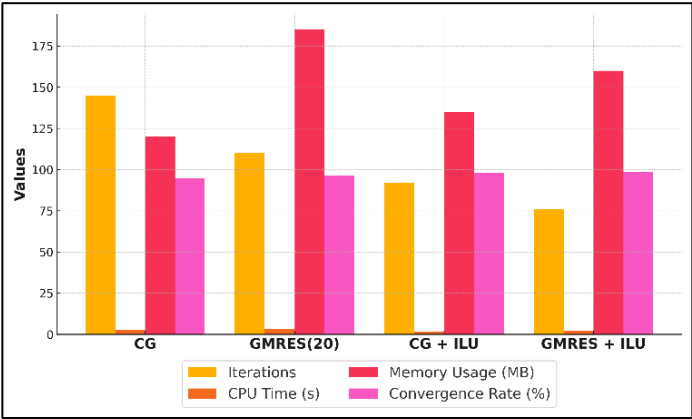


Figure 3
Comparison of Solver Performance Metrics

5. Conclusion

A study that compares Krylov subspace methods shows that preconditioning can speed up convergence and make things run more smoothly.

If we don't precondition CG and GMRES, they need a lot more iterations and CPU time, which makes them less useful for big tasks. Using ILU cuts the number of iterations by a huge amount (from 145 to 92 for CG and from 110 to 76 for GMRES), which speeds up processing while keeping sparsity. GMRES has faster convergence rates than CG, but it needs more memory, which shows that there is a trade-off between speed and storage. This shows that preconditioning is important for improving the scalability, robustness, and performance of iterative solvers, especially for large, sparse, and poorly-conditioned linear systems used in science and engineering.

References

- [1] L. Reichel and M. M. Spalević, "Averaged Gauss quadrature formulas: Properties and applications," *J. Comput. Appl. Math.*, vol. 410, pp. 114232 (2022).
- [2] D. L. Djukić, R. M. Mutavdžić Djukić, L. Reichel, and M. M. Spalević, "Weighted averaged Gaussian quadrature rules for modified Chebyshev measures," *Appl. Numer. Math.*, vol. 200, pp. 195–208 (2024).
- [3] O. Balabanov and L. Grigori, "Randomized Gram-Schmidt process with application to GMRES," *SIAM J. Sci. Comput.*, vol. 44, pp. A1450–A1474 (2022).
- [4] S. Xu and F. Xue, "Inexact rational Krylov subspace methods for approximating the action of functions of matrices," *Electron. Trans. Numer. Anal.*, vol. 58, pp. 538–567 (2023).
- [5] C. S. Liu, C. W. Chang, and C. L. Kuo, "Re-orthogonalized/affine GMRES and orthogonalized maximal projection algorithm for solving linear systems," *Algorithms*, vol. 17, pp. 266 (2024).
- [6] F. Bouyghf, A. Messaoudi, and H. Sadok, "A unified approach to Krylov subspace methods for solving linear systems," *Numer. Algor.*, vol. 96, pp. 305–332 (2024).
- [7] A. B. Gavali, B. P. Gorakh, B. O. Ramchandra, B. T. Tukaram, and E. A. Suresh, "Exploring Assistive Vision Technologies for the Visually Impaired: A Comprehensive Review," *Int. J. Recent Adv. Eng. Technol. (IJRAET)*, vol. 13, no. 2, pp. 28–34 (Mar. 2025).
- [8] C. S. Liu and C. W. Chang, "The SOR and AOR methods with step-wise optimized values of parameters for the iterative solutions of linear systems," *Contemp. Math.*, vol. 5, pp. 4013–4028 (2024).

- [9] V. B. K. Vatti, G. C. Rao, and S. S. Pai, "Reaccelerated over relaxation (ROR) method," *Bull. Int. Math. Virtual Inst.*, vol. 10, pp. 315–324 (2020).
- [10] C. Bedjguelel, H. Gharout, and B. Farhi, "Dynamics analysis of the modified Beverton-Holt model," *J. Interdiscip. Math.*, vol. 27, no. 6, pp. 1257–1271 (2024), doi: 10.47974/JIM-1748.
- [11] S. J. Desai and K. G. Kharade, "Stock Market Price Prediction Using LSTM," *Int. J. Adv. Comput. Theory Eng.*, vol. 14, no. 1, pp. 16–19 (Apr. 2025).
- [12] A. Godbole, R. Dhabliya, V. Deshpande, S. A. Sivakumar, B. M. Shankar, and V. Khetani, "Ethical hacking and penetration testing strengthening cybersecurity posture through offensive security measures," *J. Discrete Math. Sci. Cryptogr.*, vol. 27, no. 4, pp. 1295–1305 (2024).

Received December, 2024