

Lecture 23

Question:

In $AV_k = V_{k+1}T_{k+1}$ with $V_k^T V_k = I$.

Are there any degrees of freedom?

$V_k e_1 = b / \|b\|$

- 1. Are there cases when $[b, Ab, \dots, A^{k-1}b]$ does have dimension $k+1$ ✓
- For CG does we need to run Arnoldi to completion or do we get a good answer pretty quickly? (2 variables)
- Can we apply Lanczos to non-sym matrix? *Yes but B.C. Lanczos*
- Why does projection onto the Krylov subspace help solve $Ax=b$? *Eigenvalue*
- Why is this a good subspace?
- Why faster than Stochastic?

How do random Error impact orthogonality of Lanczos? → Can we mitigate? ✓

Lanczos → Lanczos tri.

How does Arnoldi build an orthogonal basis for $K_k(A, b)$?

Why does Arnoldi rely on a Hessenberg matrix?

Can we say projection to V_k reveals the Triad structure? That's why it works?

Krylov subspace = new polynomial of E and b . Does this impact convergence?

At choose the "best" poly at each iteration!

$A[V_1 \ V_2 \ V_3] = [V_1 \ V_2 \ V_3 \ V_4] \begin{bmatrix} \alpha_1 & \beta_2 & \gamma_3 \\ & \beta_2 & \alpha_2 & \beta_3 \\ & & \beta_3 & \alpha_3 & \beta_4 \end{bmatrix}$

$V_1 \rightarrow \text{init / fixed}$

$AV_1 = \alpha_1 V_1 + \beta_2 V_2 \quad V_2^T V_1 = 0$

$\alpha_1 = V_1^T A V_1 \rightarrow \text{fixed!}$

$\beta_2 V_2 = A V_1 - \alpha_1 V_1 \quad \text{also } \|V_2\| = 1$

β_2, V_2 are "forced" up to sign.

$AV_2 = \beta_2 V_1 + \alpha_2 V_2 + \beta_3 V_3$

$y = AV_2 = \gamma_2 V_1 + \alpha_2 V_2 + \beta_3 V_3$

$V_1^T y = \gamma_2 V_1^T V_1 = \gamma_2$

$V_1^T A V_2 = V_2^T A V_1$ by symmetry

$V_2^T A V_2 = \alpha_2 \rightarrow \text{fixed?}$

$\beta_3 = \|y - \gamma_2 V_1 - \alpha_2 V_2\|$ up to sign.

What happens if you relax orthogonality?

$AV_k = V_{k+1}T_{k+1}$ but we don't assume orthogonality?

Suppose V_k wasn't orthogonal! (But $\|V_k\|=1$)

Then we choose α_k s.t. $\|AV_k - \beta_k V_{k+1} - \alpha_k V_k\| = \beta_{k+1}$

Is as small as possible.

Then we will get ortho V_k and β_k will be as small as possible.

$Ax = b$

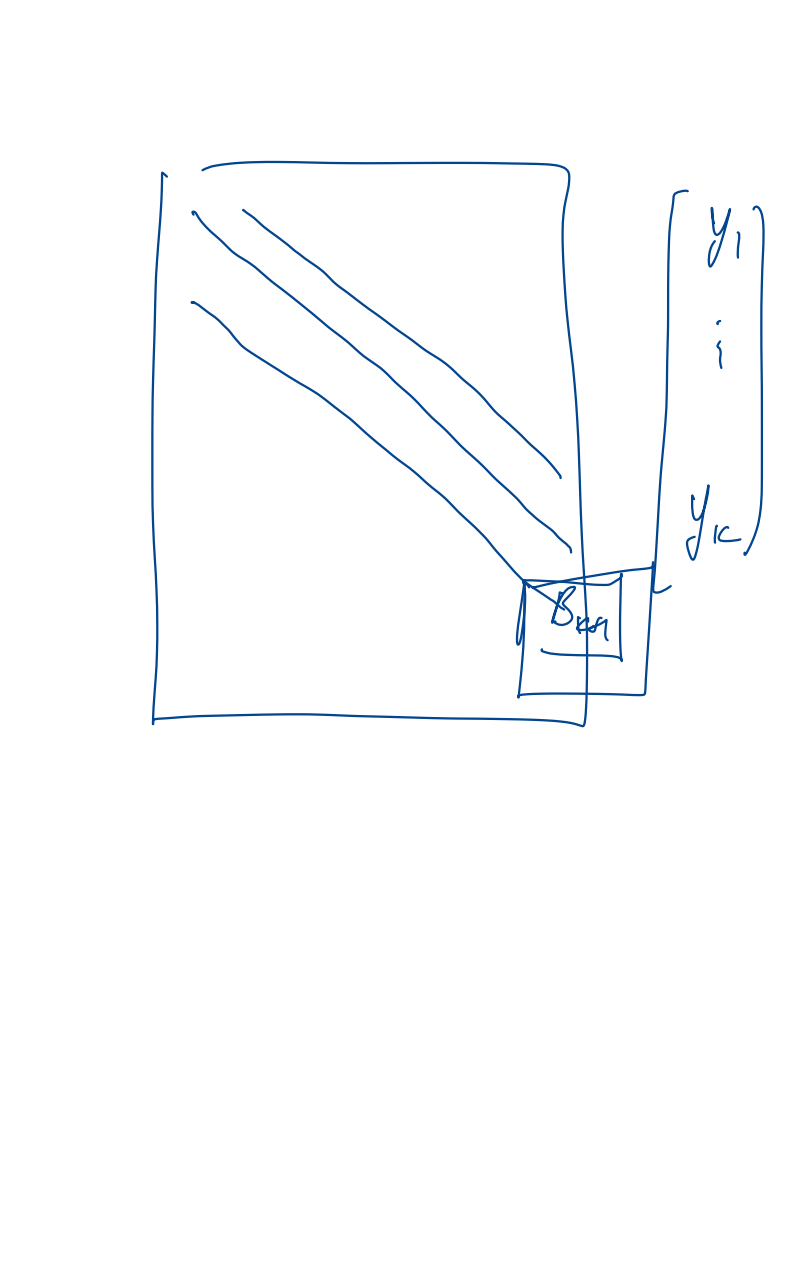
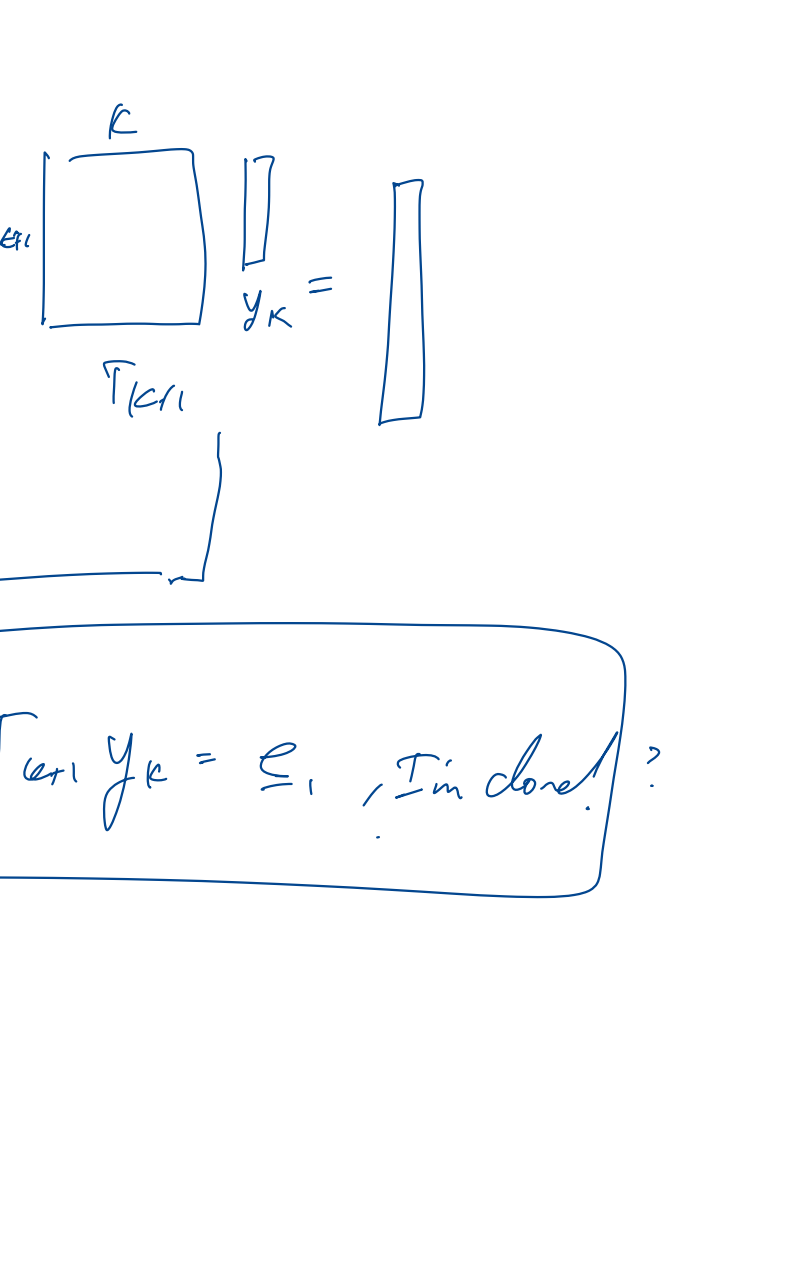
↓

If $x_k \in K_k(A, b)$ then $x_k = V_k y_k$ for some y_k .

$AV_k y_k = b$

$V_{k+1}^T V_{k+1} y_k = V_{k+1}^T b = \|b\|$

To



CG: $\min \frac{1}{2} x^T A x - x^T b$ s.t. $x_k \in K_k(A, b)$

Minkes: $\min \|Ax - b\|_2$ s.t. $x_k \in K_k(A, b)$

If you stop iteration when $\frac{\|Ax - b\|_2}{\|b\|_2} \leq \text{tol}$ then Minkes must use no more iterations than CG.

$[b, Ab, \dots, A^{k-1}b] = QR = M_k$

I want to update to $[b, Ab, \dots, A^{k-1}b, A^k b] = M_{k+1}$

$M_{k+1} = [QR \quad V]$

$Q^T M_{k+1} = [R \quad Q^T V] = \begin{bmatrix} \nabla & | \\ \hline & 1 \end{bmatrix}$

Why faster than Stochastic? SO computes $x_k \in K_k(A, b)$.

Assume Q_k spans $[b, Ab, \dots, A^{k-1}b]$ then AQ_k has vector in $[Ab, \dots, A^k b]$.

$AQ_k = Q_{k+1}H_{k+1}$

CS 515 - Purdue University - Computer Science - David F. Gleich - 2025