

Eigshow

- Eigen values of 2×2 matrix represent transformations in the plane
- Ideas of symmetry

The Power Method

- Label the eigenvalues in order of decreasing absolute value so $|\lambda_1| > |\lambda_2| > \dots > |\lambda_n|$.
- Consider the iteration formula:

$$\mathbf{y}_{k+1} = A\mathbf{y}_k$$

where we start with some initial $\mathbf{y}_0$, so that:

$$\mathbf{y}_k = A^k \mathbf{y}_0$$

- Then $\mathbf{y}_k$ converges to the eigenvector $\mathbf{x}_1$ corresponding the eigenvalue λ_1 .

Eigen Decomposition

- Recall from previous class the Eigen decomposition
- Let $\lambda_1, \lambda_2, \dots, \lambda_n$ be the eigenvalues of the $n \times n$ matrix A and $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ the corresponding eigenvectors.
- Let Λ be the diagonal matrix with $\lambda_1, \lambda_2, \dots, \lambda_n$ on the main diagonal.
- Let X be the $n \times n$ matrix whose j th column is $\mathbf{x}_j$.
- Then $AX = X\Lambda$, and so we have the *eigen decomposition* of A :

$$A = X\Lambda X^{-1}$$

- This requires X to be invertible, thus the eigenvectors of A must be linearly independent.

Powers of Matrices

- Also recall
- If $A = X\Lambda X^{-1}$ then:

$$A^2 = (X\Lambda X^{-1})(X\Lambda X^{-1}) = X\Lambda (X^{-1}X)\Lambda X^{-1} = X\Lambda^2 X^{-1}$$
 Hence we have:

$$A^p = X\Lambda^p X^{-1}$$
- Thus, A^p has the same eigenvectors as A , and its eigenvalues are $\lambda_1^p, \lambda_2^p, \dots, \lambda_n^p$.
- We can use these results as the basis of an iterative algorithm for finding the eigenvalues of a matrix.

Proof

- We know that $A^k = X \Lambda^k X^{-1}$, so:

$$\mathbf{y}_k = A^k \mathbf{y}_0 = X \Lambda^k X^{-1} \mathbf{y}_0$$

- Now we have:

$$\Lambda^k = \begin{pmatrix} \lambda_1^k & & \\ & \lambda_2^k & \\ & & \ddots \\ & & & \lambda_n^k \end{pmatrix} = \lambda_1^k \begin{pmatrix} 1 & & \\ & \lambda_2^k / \lambda_1^k & \\ & & \ddots \\ & & & \lambda_n^k / \lambda_1^k \end{pmatrix}$$

- The terms on the diagonal get smaller in absolute value as k increases, since λ_1 is the dominant eigenvalue.

Proof (continued)

- So we have

$$\mathbf{y}_k = \lambda_1^k \begin{pmatrix} \vdots & & \vdots \\ x_1 & \cdots & x_n \\ \vdots & & \vdots \end{pmatrix} \begin{pmatrix} 1 & & \\ & 0 & \\ & & \ddots \\ & & & 0 \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \\ \vdots \\ c_n \end{pmatrix} = \lambda_1^k c_1 \mathbf{x}_1$$

- Since $\lambda_1^k c_1 \mathbf{x}_1$ is just a constant times $\mathbf{x}_1$ then we have the required result.

Example

- Let $A = \begin{bmatrix} 2 & -12 \\ 1 & -5 \end{bmatrix}$ and $\mathbf{y}_0 = \begin{bmatrix} 1 & 1 \end{bmatrix}'$
- $\mathbf{y}_1 = -4 \begin{bmatrix} 2.50 & 1.00 \end{bmatrix}'$
- $\mathbf{y}_2 = 10 \begin{bmatrix} 2.80 & 1.00 \end{bmatrix}'$
- $\mathbf{y}_3 = -22 \begin{bmatrix} 2.91 & 1.00 \end{bmatrix}'$
- $\mathbf{y}_4 = 46 \begin{bmatrix} 2.96 & 1.00 \end{bmatrix}'$
- $\mathbf{y}_5 = -94 \begin{bmatrix} 2.98 & 1.00 \end{bmatrix}'$
- $\mathbf{y}_6 = -190 \begin{bmatrix} 2.99 & 1.00 \end{bmatrix}'$
- The iteration is converging on a scalar multiple of $\begin{bmatrix} 3 & 1 \end{bmatrix}'$, which is the correct dominant eigenvector.

Rayleigh Quotient

- Note that once we have the eigenvector, the corresponding eigenvalue can be obtained from the *Rayleigh quotient*:

$$\text{dot}(\mathbf{Ax}, \mathbf{x}) / \text{dot}(\mathbf{x}, \mathbf{x})$$

where $\text{dot}(\mathbf{a}, \mathbf{b})$ is the scalar product of vectors $\mathbf{a}$ and $\mathbf{b}$ defined by:

$$\text{dot}(\mathbf{a}, \mathbf{b}) = a_1 b_1 + a_2 b_2 + \dots + a_n b_n$$

- So for our example, $\lambda_1 = -2$.

Scaling

- The λ_1^k can cause problems as it may become very large or small as the iteration progresses.
- To avoid this problem we scale the iteration formula:

$$\mathbf{y}_{k+1} = A(\mathbf{y}_k / r_{k+1})$$

where r_{k+1} is the component of $A\mathbf{y}_k$ with largest absolute value.

Example with Scaling

- Let $A = \begin{bmatrix} 2 & -12 \\ 1 & -5 \end{bmatrix}$ and $\mathbf{y}_0 = [1 \ 1]'$
- $A\mathbf{y}_0 = [-10 \ -4]'$ so $r_1 = -10$ and $\mathbf{y}_1 = [1.00 \ 0.40]'$.
- $A\mathbf{y}_1 = [-2.8 \ -1.0]'$ so $r_2 = -2.8$ and $\mathbf{y}_2 = [1.0 \ 0.3571]'$.
- $A\mathbf{y}_2 = [-2.2857 \ -0.7857]'$ so $r_3 = -2.2857$ and $\mathbf{y}_3 = [1.0 \ 0.3437]'$.
- $A\mathbf{y}_3 = [-2.1250 \ -0.7187]'$ so $r_4 = -2.1250$ and $\mathbf{y}_4 = [1.0 \ 0.3382]'$.
- $A\mathbf{y}_4 = [-2.0588 \ -0.6912]'$ so $r_5 = -2.0588$ and $\mathbf{y}_5 = [1.0 \ 0.3357]'$.
- r is converging to the correct eigenvalue -2.
- At step $k+1$, the scaling factor r_{k+1} is the component with largest absolute value is $A\mathbf{y}_k$.
- When k is sufficiently large $A\mathbf{y}_k \approx \lambda_1 \mathbf{y}_k$.
- The component with largest absolute value in $\lambda_1 \mathbf{y}_k$ is λ_1 (since $\mathbf{y}_k$ was scaled in the previous step to have largest component 1).
- Hence, $r_{k+1} \searrow \lambda_1$ as $k \rightarrow \infty$.

Convergence

- The Power Method relies on us being able to ignore terms of the form $(\lambda_j / \lambda_1)^k$ when k is large enough.
- Thus, the convergence of the Power Method depends on $|\lambda_2 / \lambda_1|$.
- If $|\lambda_2 / \lambda_1| = 1$ the method will not converge.
- If $|\lambda_2 / \lambda_1|$ is close to 1 the method will converge slowly.

The QR Algorithm

- The QR algorithm for finding eigenvalues is based on the QR factorisation we learnt in the least squares part of the course
- Recall the QR factorization represents a matrix A as:

$$A = QR$$
 where Q is a matrix whose columns are orthonormal, and R is an upper triangular matrix.
- Recall that $Q^t Q = I$ and $Q^{-1} = Q^t$.

Similar Matrices

- Two matrices **A** and **B** are said to be similar if it is possible to relate them as

$$\mathbf{B} = \mathbf{T}^{-1} \mathbf{A} \mathbf{T} \qquad \mathbf{T} \mathbf{B} \mathbf{T}^{-1} = \mathbf{A}$$

- Here **T** is any non singular matrix, which is the similarity transform matrix
- Theorem:** Similar matrices have the same eigenvalues and their eigen-vectors are related via the similarity transform.
- Proof.** Let $(\mathbf{x}, \lambda)$ be an eigen-pair for **A**. Then $\mathbf{A}\mathbf{x} = \lambda \mathbf{x}$.

Let $\mathbf{y} = \mathbf{T}^{-1}\mathbf{x}$ and $\mathbf{x} = \mathbf{T}\mathbf{y}$

Then $\mathbf{A}\mathbf{x} = \lambda \mathbf{x}$.

Premultiply by $\mathbf{T}^{-1}$ to get $\mathbf{T}^{-1} \mathbf{A} \mathbf{T} \mathbf{T}^{-1} \mathbf{x} = \lambda \mathbf{T}^{-1} \mathbf{x}$

So $\mathbf{B}\mathbf{y} = \lambda \mathbf{y}$

EVD

- EVD is a similarity transform that takes **A** to a diagonal matrix using a matrix of eigenvectors.
- Eigenvalue decomposition requires solving of a general polynomial equation.
 - Even if matrix has real entries eigenvalues can be complex
 - So can eigenvectors
- Eigenvectors provide a set of basis vectors in which the matrix becomes diagonal

QR Algorithm without Shifts

$$A_0 = A$$

for $k=1,2,\dots$

$$Q_k R_k = A_k$$

$$A_{k+1} = R_k Q_k$$

end

Since:

$$A_{k+1} = R_k Q_k = Q_k^{-1} A_k Q_k$$

then A_k and A_{k+1} are similar and so have the same eigenvalues.

A_{k+1} tends to an upper triangular matrix with the same eigenvalues as A . These eigenvalues lie along the main diagonal of A_{k+1} .

MATLAB Code for QR Algorithm

- Let A be an $n \times n$ matrix
 $n = \text{size}(A,1);$
 $I = \text{eye}(n,n);$
 $s = A(n,n); [Q,R] = \text{qr}(A-s*I); A = R*Q+s*I$
- Use the up arrow key in MATLAB to iterate or put a loop round the last line.

Deflation

- The eigenvalue at $A(n,n)$ will converge first.
- Then we set $s=A(n-1,n-1)$ and continue the iteration until the eigenvalue at $A(n-1,n-1)$ converges.
- Then set $s=A(n-2,n-2)$ and continue the iteration until the eigenvalue at $A(n-2,n-2)$ converges, and so on.
- This process is called *deflation*.

The SVD

- **Definition:** Every matrix A of dimensions $m \times n$ ($m \geq n$) can be decomposed as

$$A = U \Sigma V^T$$

- where
 - U has dimension $m \times m$ and $U^T U = I$,
 - Σ has dimension $m \times n$,
the only nonzeros are on the main diagonal, and they are nonnegative real numbers $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n$,
 - V has dimension $n \times n$ and $V^T V = I$.

Relation with the Eigenvalue Decomposition

- Let $A = U \Sigma V^t$. Then

$$\begin{aligned} A^t A &= (U \Sigma V^t)^t U \Sigma V^t \\ &= V \Sigma^t U^t U \Sigma V^t = V \Sigma^2 V^t \end{aligned}$$

- This tells us that the singular value decomposition of A is related to the Eigenvalue decomposition of $A^t A$
- Recall eigen value decomposition $A = (X \Lambda X^t)$
 - So V which contains the right singular vectors of A has the right eigenvectors of $A^t A$
 - Σ^2 are the eigenvalues of $A^t A$
 - The **singular values** σ_i of A are the square roots of the **eigenvalues** of $A^t A$.

Relation with the Eigenvalue Decomposition (2)

- Let $A = U \Sigma V^t$. Then

$$\begin{aligned} A A^t &= (U \Sigma V^t) (U \Sigma V^t)^t \\ &= U \Sigma V^t V \Sigma U^t = U \Sigma^2 U^t \end{aligned}$$

- This tells us that the singular value decomposition of A is related to the Eigenvalue decomposition of $A A^t$
- Recall eigen value decomposition $A = (X \Lambda X^t)$
 - So U contains the **left singular vectors** of A , which are also the left eigenvectors of $A A^t$
 - Σ^2 are the eigenvalues of $A A^t$ and the **singular values** σ_i of A are the square roots of the **eigenvalues** of $A A^t$

Economy-sized SVD

$$\boxed{A} = \boxed{U} \boxed{\Sigma} \boxed{V'}$$

$$\boxed{A} = \boxed{U} \boxed{\Sigma} \boxed{V'}$$

Computing the SVD

- The algorithm is a variant on algorithms for computing eigendecompositions.
 - rather complicated, so better to use a high-quality existing code rather than writing your own.
- In Matlab: $[U, S, V] = \text{svd}(A)$
- The cost is $O(mn^2)$ when $m \geq n$. The constant is of order 10.

Uses of the SVD

- Recall to solve least squares problems we could look at the normal equations ($A^*Ax=A^*b$)
 - So, SVD is closely related to solution of least-squares
 - Used for solving ill conditioned least-squares
- Used for creating low-rank approximations
- Both applications are related

SVD and reduced rank approximation

- $Ax=b$ A is $m \times n$, x is $n \times 1$ and b is $m \times 1$.
- $A=USV^t$ where U is $m \times m$, S is $m \times n$ and V is $n \times n$
- $USV^t x=b$. So $SV^t x=U^t b$
- If A has rank r , then r singular values are significant
 $V^t x = \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) U^t b$
 $x = V \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) U^t b$

$$x_r = \sum_{i=1}^r \frac{u_i^t b}{\sigma_i} v_i \quad \sigma_r > \varepsilon, \quad \sigma_{r+1} \leq \varepsilon$$
- We can truncate r at any value and achieve “reduced-rank” approximation to the matrix
- For ordered singular values, this gives the “best reduced rank approximation”

SVD and pseudo inverse

- Pseudoinverse $\mathbf{A}^+ = \mathbf{V} \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) \mathbf{U}^t$
 - $\mathbf{A}^+$ is a $n \times m$ matrix.
 - If $\text{rank}(\mathbf{A}) = n$ then $\mathbf{A}^+ = (\mathbf{A}^t \mathbf{A})^{-1} \mathbf{A}$
 - If $\mathbf{A}$ is square $\mathbf{A}^+ = \mathbf{A}^{-1}$

Well posed problems

- Hadamard postulated that for a problem to be “well posed”
 1. Solution must exist
 2. It must be unique
 3. Small changes to input data should cause small changes to solution
- Many problems in science and computer vision result in “ill-posed” problems.
 - Numerically it is common to have condition 3 violated.
 - Recall from the SVD
$$\mathbf{x} = \sum_{i=1}^n \frac{\mathbf{u}_i^t \mathbf{b}}{\sigma_i} \mathbf{v}_i \quad \sigma_r > \varepsilon, \quad \sigma_{r+1} \leq \varepsilon$$

If the σ are close to zero small changes in the “data” vector $\mathbf{b}$ cause big changes in $\mathbf{x}$.

- Converting ill-posed problem to well-posed one is called *regularization*.

SVD and Regularization

- Pseudoinverse provides one means of regularization
- Another is to solve $(\mathbf{A} + \epsilon \mathbf{I})\mathbf{x} = \mathbf{b}$ $\mathbf{x} = \sum_{i=1}^n \frac{\sigma_i}{\epsilon + \sigma_i^2} (\mathbf{u}_i^t \mathbf{b}) \mathbf{v}_i$
- Solution of the regular problem requires minimizing of $\|\mathbf{Ax} - \mathbf{b}\|^2$
- Solving this modified problem corresponds to minimizing $\|\mathbf{Ax} - \mathbf{b}\|^2 + \epsilon \|\mathbf{x}\|^2$
- Philosophy – pay a “penalty” of $O(\epsilon)$ to ensure solution does not blow up.
- In practice we may know that the data has an uncertainty of a certain magnitude ... so it makes sense to optimize with this constraint.
- Ill-posed problems are also called “ill-conditioned”

SVD and Pseudo-Inverse

- $\mathbf{Ax} = \mathbf{b}$ $\mathbf{A}$ is $m \times n$, $\mathbf{x}$ is $n \times 1$ and $\mathbf{b}$ is $m \times 1$.
- $\mathbf{A} = \mathbf{USV}^t$ where $\mathbf{U}$ is $m \times m$, $\mathbf{S}$ is $m \times n$ and $\mathbf{V}$ is $n \times n$
- $\mathbf{USV}^t \mathbf{x} = \mathbf{b}$. So $\mathbf{SV}^t \mathbf{x} = \mathbf{U}^t \mathbf{b}$
- If $\mathbf{A}$ has rank r , then r singular values are significant
 $\mathbf{V}^t \mathbf{x} = \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) \mathbf{U}^t \mathbf{b}$
 $\mathbf{x} = \mathbf{V} \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) \mathbf{U}^t \mathbf{b}$

$$\mathbf{x}_r = \sum_{i=1}^r \frac{\mathbf{u}_i^t \mathbf{b}}{\sigma_i} \mathbf{v}_i \quad \sigma_r > \epsilon, \quad \sigma_{r+1} \leq \epsilon$$
- Pseudoinverse $\mathbf{A}^+ = \mathbf{V} \text{diag}(\sigma_1^{-1}, \dots, \sigma_r^{-1}, 0, \dots, 0) \mathbf{U}^t$
 - $\mathbf{A}^+$ is a $n \times m$ matrix.
 - If $\text{rank}(\mathbf{A}) = n$ then $\mathbf{A}^+ = (\mathbf{A}^t \mathbf{A})^{-1} \mathbf{A}$
 - If $\mathbf{A}$ is square $\mathbf{A}^+ = \mathbf{A}^{-1}$