

Lecture 20: Unsupervised Learning: Principal Component Analysis (PCA)
 Derivations of Maximum Likelihood Estimation, minimizing the variance and minimizing the sum of squared projections. Eigenfaces for face recognition.

Unsupervised Learning:

we have sample points but no labels i.e., we don't have any classes or y-values against which we can regress to attain predictions.

The goal of unsupervised learning though is to discover structure in the data.

Some examples of unsupervised learning:

- ① Clustering: partition data into groups of similar/neighboring points
- ② Dimensionality reduction: data often lies near a low dimensional subspace (or manifold) in feature space. **MARKET** have low rank approximations.
- ③ Density Estimation:

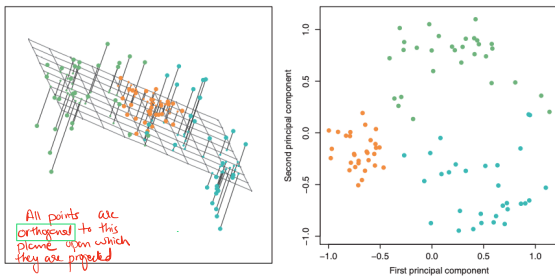
Goal is to fit a continuous probability distribution to discrete data (MLE is an archetypal example of Density Estimation) but, more complex density estimation problems also exist.

PRINCIPAL COMPONENT ANALYSIS (PCA)

(KARL PEARSON, 1901)

Goal: Given sample points in $\mathbb{R}^d$, find k -directions in our d dimensions that capture most of the information / variance / variations between the different points.

for example:



[3dpca.pdf] [Example of 3D points projected to 2D by PCA.]

Goal: how do I choose a plane so that we don't lose more information than necessary after dimensionality reduction.

An example of higher dimension reduction to a 2-d dimensional space:

5 4 1 7 7 6 6 1 1
 4 5 7 7 6 6 1 1
 2 1 7 7 1 2 4 5
 4 1 1 0 1 8 8 4
 1 1 1 1 4 1 0 0
 7 8 9 2 6 3 1 7
 2 2 2 2 3 4 1 0
 8 5 8 0 7 8 8 7
 0 1 6 6 4 6 0 2 0 3
 7 2 8 1 6 1 8 5 1



[pcadigits.pdf] [The (high-dimensional) MNIST digits projected to 2D (from 784D). Two dimensions aren't enough to fully separate the digits, but observe that the digits 0 (red) and 1 (orange) are well on their way to being separated.]

maybe would do classification well if we reduced dimensions to 4 dimensions then 2, but obviously one can't visualize that.

But, why do we need to do this?

- ① Reducing # of dimensions makes some computation cheaper e.g. regression $\mathbb{R}^{1000} \rightarrow \mathbb{R}^{100}$ so regression code will run 10 times faster
- ② Remove irrelevant dimensions to reduce overfitting in learning algorithms (fewer parameters to fit $\rightarrow$ less overfitting though you won't be left with features per se rather linear combinations of those features)
 very much like subset selection however the features that we lose keeping versus the features we are throwing away are linear combinations of those input features $\rightarrow$ LOSS OF INTERPRETABILITY (we don't get our original features back)
- ③ Find a small basis for representing variations in complex things e.g. faces, genes etc.
- ④ Just for the sake of visualization (PCA plots)

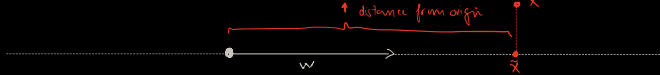
Let X be $n \times d$ design matrix AND center the matrix.

From now on, assume X is centered i.e., $\mu_{X_i} = \text{mean}_i X_i = 0$

To start with the math while using or computing principle components, let's start with assuming that we have found one of the directions that we think is important, one of the directions that captures a lot of the variance in our dataset $\rightarrow$ we'll call this direction a principle component or a principle direction $\rightarrow$ Let's call it w

so, assume w is a unit vector.

Now, The ORTHOGONAL PROJECTION of some sample point x onto vector w is $\tilde{x} = (x \cdot w)w$



so, $(x \cdot w)w$ is basically the distance from the origin to the orthogonal projection of a point x onto the direction w .

If w is not a unit vector then

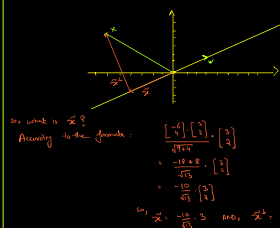
$$\tilde{x} = \frac{(x \cdot w)w}{\|w\|^2}$$

And so, the idea here is that we are going to pick the best direction w so that we can analyze then project all the data down onto w so that we can analyze it in a one-dimensional space. Otherwise we are losing a lot of information when we are projecting down from d -dimensions to just one. So, suppose we pick several directions. Those directions span a subspace and we want to project points orthogonally onto the subspace. This is easy if the directions are orthogonal to each other.

An example to understand the formula

suppose line $\tilde{w}$ in 2-D is given by the vector $\tilde{w} = \begin{bmatrix} 3 \\ 4 \end{bmatrix}$ and we have a point $\tilde{x} = \begin{bmatrix} -6 \\ 4 \end{bmatrix}$ and we are asked to find $\tilde{x}$: Projection of x onto the line spanned by $\tilde{w}$

Graphically drawing this $\Rightarrow$

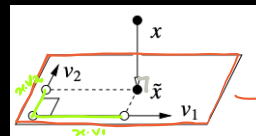


so, basically, we can say that $\tilde{x} = c \cdot w$
 Now, $x - \tilde{x} = x - c \cdot w$ which would be orthogonal
 so, $(x - c \cdot w) \cdot w = 0$
 $x \cdot w = (c \cdot w) \cdot w$
 $c = \frac{x \cdot w}{\|w\|^2}$
 And so, $\tilde{x} = c \cdot w = \frac{x \cdot w}{\|w\|^2} \cdot w$

Given orthonormal directions mutually orthogonal to each other and unit vectors.

$$v_1, \dots, v_k = \sum_{i=1}^k (x \cdot v_i) \cdot v_i$$

coordinates in some special coordinate system that lives inside the actual dimension space of x .



$$x \in \mathbb{R}^n$$

$$v \in \mathbb{R}^n \begin{bmatrix} v_1 \\ v_2 \end{bmatrix}$$

often, we just want the k -principle coordinates $x \cdot v_i$ in principle component space (eigenvectors) and don't actually want the projected point in a lower dimensional space $\mathbb{R}^d$.

Now how do we get those principle components i.e., the directions that define our subspace?

we get them from the matrix $X^T X$

- 1) square matrix
- 2) symmetric matrix
- 3) Positive Semi-Definite

why is $X^T X$ a positive semi definite matrix?

Going back to HW-2,

2. Let $A \in \mathbb{R}^{n \times n}$ be a symmetric matrix. Prove equivalence between these three different definitions of positive semidefiniteness (PSD). Note that when we talk about PSD matrices, they are defined to be symmetric matrices. There are nonsymmetric matrices that exhibit PSD properties, like the first definition below, but not all three.

- (a) For all $x \in \mathbb{R}^n$, $x^T A x \geq 0$.
- (b) All the eigenvalues of A are nonnegative.
- (c) There exists a matrix $U \in \mathbb{R}^{n \times n}$ such that $A = U U^T$.

Positive semidefiniteness will be denoted as $A \succeq 0$.

Solution: Let λ be an eigen value of A with corresponding eigen vector x . Then,

(a) to prove (b) $ATAx = \lambda TTx = \lambda TTx$

From part (a), we know that $ATx = \lambda x$

$ATAx \geq 0$ so, $ATTx \geq 0$

so, $\lambda TTx \geq 0$

hence $\lambda \geq 0$ i.e. all eigenvalues of A are non-negative.

using (b) to prove (c)

consider the eigen decomposition of A i.e.

$A = V\Lambda V^T$ where Λ is the diagonal matrix with entries equal to eigenvalues of A i.e. $\lambda_1, \dots, \lambda_n$

If $U = A\sqrt{\Lambda}$

Then $U^T = (A\sqrt{\Lambda})^T = \sqrt{\Lambda}^T A^T$

Now we compare $A = V\Lambda V^T$ with U and observe that:

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

$A = V\Lambda V^T = \sqrt{\Lambda} \sqrt{\Lambda}^T V^T V^T$

$= V\Lambda V^T$

This definition of $U = A\sqrt{\Lambda}$ is only possible for non-negative eigenvalues which is why we use using (b) to prove (c).

Understanding this requires a quick recap of lecture 6 & context!

When you scale an eigenvalue, it changes the eigen vector. You need only the direction vector.

So $Ax = \lambda x$

① $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

② $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

③ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

④ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑤ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑥ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑦ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑧ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑨ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑩ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑪ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑫ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑬ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑭ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑮ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑯ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑰ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑱ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑲ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

⑳ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉑ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉒ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉓ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉔ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉕ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉖ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉗ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉘ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉙ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉚ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉛ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉜ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉝ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉞ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㉟ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊱ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊲ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊳ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊴ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊵ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊶ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊷ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊸ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊹ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊺ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊻ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊼ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

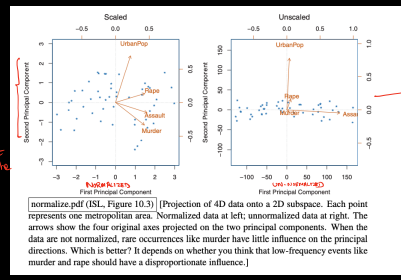
㊽ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊾ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

㊿ $A(1/\lambda)x = 1/\lambda Ax = 1/\lambda \lambda x = x$

PCA Algorithm:-

- Center X
- Optional: Normalise X . (if units of measurements are different)
scale of different features really matters (as linear combinations)
- Compute unit eigen vectors and corresponding eigen values of XTX
- Choose k (optional): based on eigenvalue size
- For the best k -dimensional subspace, pick eigen vectors $v_1, \dots, v_k$
- compute the k -principle coordinate each training point. [when we do this projection, we have 2 choices: we can project the original, uncentered training data OR we can project the centered training data. But if we do the latter, we have to translate the test data by the same vector we used to translate the training data when we centered it.]



Could be useful to quantify change why? because murders (despite being lowest in numbers) pose the greatest danger! (As ones deserve equal consideration)

↓

Means this is application specific!

Application Specific → if we want to analyse the time spent by car, force is contributing crimes. Because it is the most frequent.

using (c) to prove (a)

$XTAx = \lambda TTx = \lambda TTx$

So a square symmetric matrix A is P.S.D since it satisfies the following 3 conditions i.e.

- $XTAx \geq 0$
- all λ 's are non-negative
- There exists a matrix U such that $A = UU^T$

So,

XTX is a square & symmetric, positive semi-definite 2d matrix [As its symmetric, its eigenvalues are real] → spectral theorem

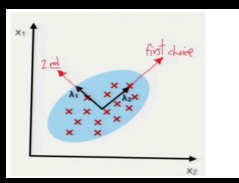
Let, $0 \leq \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_n$ be its eigenvalues [ascending order]

Let $v_1, v_2, \dots, v_n$ be corresponding orthogonal unit vectors. These are called the principal components.

And, the most imp. principle components (eigen vectors) are the one with the greatest eigen value!

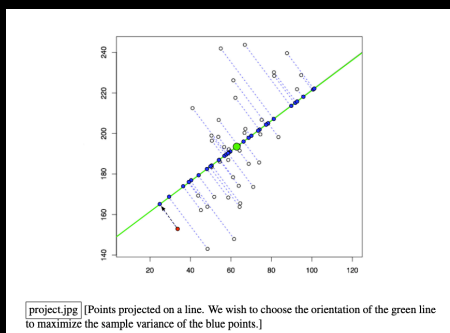
showing this in 3-different motivating derivations:-

- PCA Derivation 1: Fit a Gaussian to data with maximum likelihood
- Choose k Gaussian axes of greatest variance-



Recall that MLE estimates a covariance matrix $\hat{\Sigma} = \frac{1}{n} X^T X$ [Presuming X is centered]

PCA Derivation 2: Find direction w that maximizes the sample variance of projected data.



Writing this as a formal optimization problem:

Find w that maximizes $V_{\text{var}}(\{x_1, x_2, \dots, x_n\}) = \frac{1}{n} \sum_{i=1}^n (x_i \cdot w)^2 = \frac{1}{n} \frac{\|Xw\|^2}{\|w\|^2}$

standard variance formula, when mean is 0

$= \frac{1}{n} \frac{w^T X^T X w}{w^T w}$

when you see Rayleigh Quotient of $X^T X$ and w

we can simplify eigen vectors easily!

So, how to maximize this variance?

(i.e.)

Find w that maximizes $\frac{1}{n} \frac{w^T X^T X w}{|w|^2}$

If w is an eigenvector v_i of $X^T X$,
then the Rayleigh Quotient will be $\Rightarrow \lambda_i$

$$\frac{|w^T X^T X w|}{|w|^2} \Rightarrow \frac{w^T \lambda_i w}{w^T w} \Rightarrow \lambda_i \frac{w^T w}{w^T w} = \lambda_i$$

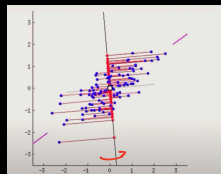
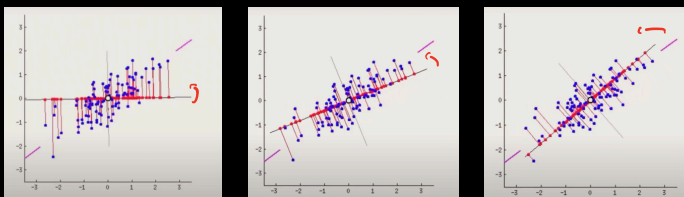
of all eigenvectors, v_1 achieves maximum variance λ_1/n
since v_1 has the greatest eigenvalue.

What if w is not restricted to be an eigenvector? $\Rightarrow$ no other
vector can beat v_1 in capturing the maximum variance.
(Not a difficult proof $\Rightarrow$ on Rayleigh Quotient page on Wikipedia)

So, eigenvector associated with the largest eigenvalue gives you
a principle component that gives you the maximum variance however,
we typically want k -directions. After we've picked one
direction v_1 , we want to pick a second direction that
that is orthogonal to the best direction. But subject to
that constraint, we again pick the direction that maximizes
the sample variance.

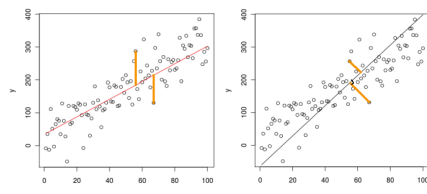
When we constrain w to be orthogonal to
 v_1 , then v_2 is optimal! then
once v_1, v_2 are constrained, next
direction will be v_3 and so on...

PCA Derivation 3: Find direction w that minimizes the
mean squared projection distance.
(sum of squares of v_i of each point
onto a projection axis)



minimizing sum of squares
sad edges.

[You can think of this as a sort of least-squares linear regression, with one subtle but important change. In-
stead of measuring the error in a fixed vertical direction, we're measuring the error in a direction orthogonal
to the principal component direction we choose.]



[projlsq.png, projpca.png] [Least-squares linear regression vs. PCA. In linear regression, the projection direction is always vertical; whereas in PCA, the projection direction is orthogonal to the projection hyperplane. In both methods, however, we minimize the sum of the squares of the projection distances.]

Find w that minimizes $\sum_{i=1}^n \|X_i - \hat{X}_i\|^2 = \sum_{i=1}^n \|X_i - \frac{X_i \cdot w}{|w|} w\|^2 = \sum_{i=1}^n \left(\|X_i\|^2 - \left(X_i \cdot \frac{w}{|w|} \right)^2 \right)$
 $= \text{constant} - n (\text{variance from derivation 2}).$

$$\begin{aligned} \sum_{i=1}^n \|X_i - \hat{X}_i\|^2 &= \sum_{i=1}^n \left\| X_i - \frac{X_i \cdot w}{|w|} w \right\|^2 \\ &= \sum_{i=1}^n \left(\|X_i\|^2 - \left(X_i \cdot \frac{w}{|w|} \right)^2 \right) \\ &= \sum_{i=1}^n \|X_i\|^2 - \sum_{i=1}^n \frac{X_i \cdot w}{|w|} w \\ &= \sum_{i=1}^n \|X_i\|^2 - \frac{2}{|w|^2} \sum_{i=1}^n X_i \cdot w \\ &= \sum_{i=1}^n \|X_i\|^2 - \frac{2}{|w|^2} \left(\sum_{i=1}^n X_i \right) \cdot w \\ &= \sum_{i=1}^n \|X_i\|^2 - \frac{2}{|w|^2} \left(\sum_{i=1}^n X_i \right) \cdot w \end{aligned}$$

EIGEN - FACES

Eigenfaces

X contains n images of faces, d pixels each.

[If we have a 200×200 image of a face, we represent it as a vector of length 40,000, the same way we represent the MNIST digit data.]

Face recognition: Given a query face, compare it to all training faces; find nearest neighbor in $\mathbb{R}^d$.
[This works best if you have several training photos of each person you want to recognize, with different lighting and different facial expressions.]

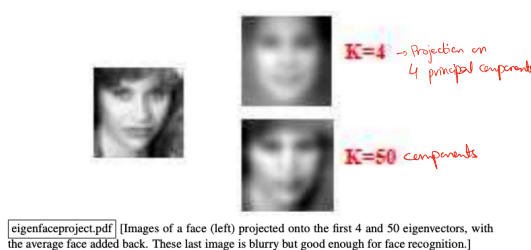
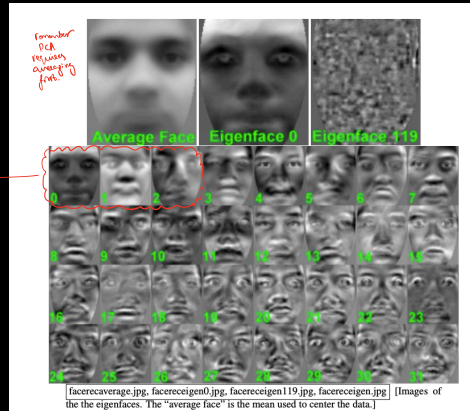
Problem: Each query takes $\Theta(nd)$ time.

Solution: Run PCA on faces. Reduce to much smaller dimension d' .

Now nearest neighbors takes $\Theta(nd')$ time.

[Possibly even less. We'll talk about speeding up nearest-neighbor search at the end of the semester. If the dimension is small enough, you can sometimes do better than linear time.]

[If you have 500 stored faces with 40,000 pixels each, and you reduce them to 40 principal components, then each query face requires you to read 20,000 stored principal coordinates instead of 20 million pixels.]



For best results, equalize the intensity distributions first.



[facerecequalize.jpg] [Image equalization.]

[Eigenfaces are not perfect. They encode both face shape and lighting. Ideally, we would have some way to factor out lighting and analyze face shape only, but that's harder. Some people say that the first 3 eigenfaces are usually all about lighting, and you sometimes get better facial recognition by dropping the first 3 eigenfaces.]

[Optional: Show Blanz-Vetter face morphing video (morphmod.mpg).]

[Blanz and Vetter use PCA in a more sophisticated way for 3D face modeling. They take 3D scans of people's faces and find correspondences between peoples' faces and an idealized model. For instance, they identify the tip of your nose, the corners of your mouth, and other facial features, which is something the original eigenface work did not do. Instead of feeding an array of pixels into PCA, they feed the 3D locations of various points on your face into PCA. This works more reliably.]

Lecture - 21 ; Singular Value Decomposition

An alternative way of computing eigenvectors / eigenvalues that are associated with PCA.

So, previously, ① Computing eigen decomposition of $X^T X$
takes $\Theta(n^3)$ time.

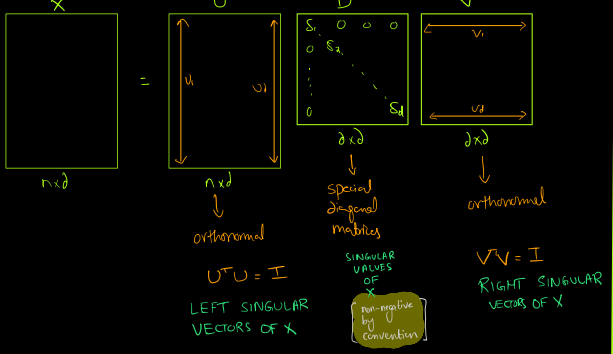
② $X^T X$ is poorly conditioned $\rightarrow$ numerically inaccurate eigenvectors.

Instead of doing an eigen decomposition of $X^T X$, we're going to compute a singular value decomposition of X and we're going to get the same vectors and indirectly the same values.

FACT: Every matrix has an SVD whether it is symmetric or not, square or not

$$X = UDV^T$$

If $n \geq d$, it has the form:



$$\text{So, } X = UDV^T = \sum_{i=1}^d \sigma_i u_i v_i^T$$

outer product
n x d matrix
[rank 1]

→ some of the singular values can be zero!
→ number of singular values is equal to rank of matrix X.

FACT: v_i is an eigenvector of $X^T X$

So, by finding the right singular values of X , we have found eigenvectors of $X^T X$.

$$X^T X = V D^2 U^T U D V^T = V D^2 V^T$$

(as orthogonal)

eigen decomposition of $X^T X$

FACT: We can find the k -greatest singular values & the corresponding right singular vectors in $O(ndk)$ time.

Important: Row i of UD gives the principal coordinates of sample point x_i in k -dim space

(i.e.)

sample point x_i (i.e. $x_i = \sum_{j=1}^k u_j \cdot v_j \cdot \sigma_j$)

Notes on singular value decomposition for Math 54

Recall that if A is a symmetric $n \times n$ matrix, then A has real eigenvalues $\lambda_1, \dots, \lambda_n$ (possibly repeated), and $\mathbb{R}^n$ has an orthonormal basis $v_1, \dots, v_n$, where each vector v_i is an eigenvector of A with eigenvalue λ_i . Then

$$A = PDP^{-1}$$

where P is the matrix whose columns are $v_1, \dots, v_n$, and D is the diagonal matrix whose diagonal entries are $\lambda_1, \dots, \lambda_n$. Since the vectors $v_1, \dots, v_n$ are orthonormal, the matrix P is orthogonal, i.e. $P^T P = I$, so we can alternatively write the above equation as

$$A = PDP^T, \quad (1)$$

A singular value decomposition (SVD) is a generalization of this where A is an $m \times n$ matrix which does not have to be symmetric or even square.

1 Singular values

Let A be an $m \times n$ matrix. Before explaining what a singular value decomposition is, we first need to define the singular values of A .

Consider the matrix $A^T A$. This is a symmetric $n \times n$ matrix, so its eigenvalues are real.

Lemma 1.1. If λ is an eigenvalue of $A^T A$, then $\lambda \geq 0$.

Proof. Let x be an eigenvector of $A^T A$ with eigenvalue λ . We compute that $\|Ax\|^2 = (Ax)^T (Ax) = x^T A^T A x = \lambda x^T x = \lambda \|x\|^2$.

Since $\|Ax\|^2 \geq 0$, it follows from the above equation that $\lambda \|x\|^2 \geq 0$. Since $\|x\|^2 > 0$ (as our convention is that eigenvectors are nonzero), we deduce that $\lambda \geq 0$. $\square$

Let $\lambda_1, \dots, \lambda_n$ denote the eigenvalues of $A^T A$, with repetitions. Order these so that $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq 0$. Let $\sigma_i = \sqrt{\lambda_i}$, so that $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$.

Definition 1.2. The numbers $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$ defined above are called the **singular values** of A .

Proposition 1.3. The number of nonzero singular values of A equals the rank of A .

Proof. The rank of any square matrix equals the number of nonzero eigenvalues (with repetitions), so the number of nonzero singular values of A equals the rank of $A^T A$. By a previous homework problem, $A^T A$ and A have the same kernel. It then follows from the "rank-nullity" theorem that $A^T A$ and A have the same rank. $\square$

Remark 1.4. In particular, if A is an $m \times n$ matrix with $m < n$, then A has at most m nonzero singular values, because $\text{rank}(A) \leq m$.

The singular values of A have the following geometric significance.

Proposition 1.5. Let A be an $m \times n$ matrix. Then the maximum value of $\|Ax\|$, where x ranges over unit vectors in $\mathbb{R}^n$, is the largest singular value σ_1 , and this is achieved when x is an eigenvector of $A^T A$ with eigenvalue σ_1^2 .

Proof. Let $v_1, \dots, v_n$ be an orthonormal basis for $\mathbb{R}^n$ consisting of eigenvectors of $A^T A$ with eigenvalues σ_i^2 . If $x \in \mathbb{R}^n$, then we can expand x in this basis as

$$x = c_1 v_1 + \dots + c_n v_n \quad (2)$$

for scalars $c_1, \dots, c_n$. Since x is a unit vector, $\|x\|^2 = 1$, which (since the vectors $v_1, \dots, v_n$ are orthonormal) means that

$$c_1^2 + \dots + c_n^2 = 1.$$

On the other hand,

$$\|Ax\|^2 = (Ax)^T (Ax) = x^T A^T A x = x^T (A^T A) x.$$

By (2), since v_i is an eigenvector of $A^T A$ with eigenvalue σ_i^2 , we have

$$A^T A x = c_1 \sigma_1^2 v_1 + \dots + c_n \sigma_n^2 v_n.$$

Taking the dot product with (2), and using the fact that the vectors $v_1, \dots, v_n$ are orthonormal, we get

$$\|Ax\|^2 = x \cdot (A^T A x) = \sigma_1^2 c_1^2 + \dots + \sigma_n^2 c_n^2.$$

Since σ_1 is the largest singular value, we get

$$\|Ax\|^2 \leq \sigma_1^2 (c_1^2 + \dots + c_n^2).$$

Equality holds when $c_1 = 1$ and $c_2 = \dots = c_n = 0$. Thus the maximum value of $\|Ax\|^2$ for a unit vector x is σ_1^2 , which is achieved when $x = v_1$. $\square$

One can similarly show that σ_2 is the maximum of $\|Ax\|$ where x ranges over unit vectors that are orthogonal to v_1 (exercise). Likewise, σ_3 is the maximum of $\|Ax\|$ where x ranges over unit vectors that are orthogonal to v_1 and v_2 , and so forth.

Proof. We compute

$$(Av_i) \cdot (Av_i) = (Av_i)^T (Av_i) = v_i^T A^T A v_i = \sigma_i^2 v_i^T v_i = \sigma_i^2 (v_i \cdot v_i).$$

If $i = j$, then since $\|v_i\| = 1$, this calculation tells us that $\|Av_i\|^2 = \sigma_i^2$, which proves (a). If $i \neq j$, then since $v_i \cdot v_j = 0$, this calculation shows that $(Av_i) \cdot (Av_j) = 0$. $\square$

Theorem 3.2. Let A be an $m \times n$ matrix. Then A has a (not unique) singular value decomposition $A = U\Sigma V^T$, where U and V are as follows:

- The columns of V are orthonormal eigenvectors $v_1, \dots, v_n$ of $A^T A$, where $A^T A v_i = \sigma_i^2 v_i$.
- If $i \leq r$, so that $\sigma_i \neq 0$, then the i th column of U is $\sigma_i^{-1} A v_i$. By Lemma 3.1, these columns are orthonormal, and the remaining columns of U are obtained by arbitrarily extending to an orthonormal basis for $\mathbb{R}^m$.

Proof. We just have to check that if U and V are defined as above, then $A = U\Sigma V^T$. If $x \in \mathbb{R}^n$, then the components of $V^T x$ are the dot products of the rows of V^T with x , so

$$V^T x = \begin{pmatrix} v_1 \cdot x \\ v_2 \cdot x \\ \vdots \\ v_n \cdot x \end{pmatrix}$$

Then

$$\Sigma V^T x = \begin{pmatrix} \sigma_1 v_1 \cdot x \\ \sigma_2 v_2 \cdot x \\ \vdots \\ \sigma_r v_r \cdot x \\ 0 \\ \vdots \\ 0 \end{pmatrix}.$$

When we multiply on the left by U , we get the sum of the columns of U , weighted by the components of the above vector, so that

$$U\Sigma V^T x = (\sigma_1 v_1 \cdot x) \sigma_1^{-1} A v_1 + \dots + (\sigma_r v_r \cdot x) \sigma_r^{-1} A v_r = (v_1 \cdot x) A v_1 + \dots + (v_r \cdot x) A v_r.$$

Step 3. We now find the matrix U . The first column of U is

$$\sigma_1^{-1} A v_1 = \frac{1}{6\sqrt{10}} \begin{pmatrix} 18 \\ 6 \end{pmatrix} = \frac{1}{\sqrt{10}} \begin{pmatrix} 3 \\ 1 \end{pmatrix}.$$

The second column of U is

$$\sigma_2^{-1} A v_2 = \frac{1}{3\sqrt{10}} \begin{pmatrix} 3 \\ 9 \end{pmatrix} = \frac{1}{\sqrt{10}} \begin{pmatrix} 1 \\ 3 \end{pmatrix}.$$

Since U is a 2×2 matrix, we do not need any more columns. (If A had only one nonzero singular value, then we would need to add another column to U to make it an orthogonal matrix.) Thus

$$U = \frac{1}{\sqrt{10}} \begin{pmatrix} 3 & 1 \\ 1 & 3 \end{pmatrix}.$$

To conclude, we have found the singular value decomposition

$$\begin{pmatrix} 4 & 11 & 14 \\ 8 & 7 & -2 \end{pmatrix} = \frac{1}{\sqrt{10}} \begin{pmatrix} 3 & 1 \\ 1 & 3 \end{pmatrix} \begin{pmatrix} 6\sqrt{10} & 0 & 0 \\ 0 & 3\sqrt{10} & 0 \end{pmatrix} \begin{pmatrix} 1/3 & -2/3 & 2/3 \\ 2/3 & -1/3 & -2/3 \\ 2/3 & 2/3 & 1/3 \end{pmatrix}^T$$

4 Applications

Singular values and singular value decompositions are important in analyzing data.

One simple example of this is "rank estimation". Suppose that we have n data points $v_1, \dots, v_n$, all of which live in $\mathbb{R}^m$, where n is much larger than m . Let A be the $m \times n$ matrix with columns $v_1, \dots, v_n$. Suppose the data points satisfy some linear relations, so that $v_1, \dots, v_n$ all lie in an r -dimensional subspace of $\mathbb{R}^m$. Then we would expect the matrix A to have rank r . However if the data points are obtained from measurements with errors, then the matrix A will probably have full rank m . But only r of the singular values of A will be large, and the other singular values will be close to zero. Thus one can compute an "approximate rank" of A by counting the number of singular values which are much larger than the others, and one expects the measured matrix A to be close to a matrix A' such that the rank of A' is the "approximate rank" of A .

For example, consider the matrix

$$A' = \begin{pmatrix} 1 & 2 & -2 & 3 \\ -4 & 0 & 1 & 2 \\ 3 & -2 & 1 & -5 \end{pmatrix}$$

2 Definition of singular value decomposition

Let A be an $m \times n$ matrix with singular values $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$. Let r denote the number of nonzero singular values of A , or equivalently the rank of A .

Definition 2.1. A singular value decomposition of A is a factorization

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

where:

- U is an $m \times m$ orthogonal matrix.
- V is an $n \times n$ orthogonal matrix.
- Σ is an $m \times n$ matrix whose i th diagonal entry equals the i th singular value σ_i for $i = 1, \dots, r$. All other entries of Σ are zero.

Example 2.2. If $m = n$ and A is symmetric, let $\lambda_1, \dots, \lambda_n$ be the eigenvalues of A , ordered so that $|\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_n|$. The singular values of A are given by $\sigma_i = |\lambda_i|$ (exercise). Let $v_1, \dots, v_n$ be orthonormal eigenvectors of A with $Av_i = \lambda_i v_i$. We can then take V to be the matrix whose columns are $v_1, \dots, v_n$. (This is the matrix P in equation (1).) The matrix Σ is the diagonal matrix with diagonal entries $|\lambda_1|, \dots, |\lambda_n|$. (This is almost the same as the matrix D in equation (1), except for the absolute value signs.) Then U must be the matrix whose columns are $\pm v_1, \dots, \pm v_n$, where the sign next to v_i is $+$ when $\lambda_i \geq 0$, and $-$ when $\lambda_i < 0$. (This is almost the same as P , except we have changed the signs of some of the columns.)

3 How to find a SVD

Let A be an $m \times n$ matrix with singular values $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$, and let r denote the number of nonzero singular values. We now explain how to find a SVD of A .

Let $v_1, \dots, v_n$ be an orthonormal basis of $\mathbb{R}^n$, where v_i is an eigenvector of $A^T A$ with eigenvalue σ_i^2 .

Lemma 3.1. (a) $\|Av_i\| = \sigma_i$.

(b) If $i \neq j$ then Av_i and Av_j are orthogonal.

Since $Av_i = 0$ for $i > r$ by Lemma 3.1(a), we can rewrite the above as

$$\begin{aligned} U\Sigma V^T x &= (v_1 \cdot x) Av_1 + \dots + (v_n \cdot x) Av_n \\ &= Av_1 v_1^T x + \dots + Av_n v_n^T x \\ &= A(v_1 v_1^T + \dots + v_n v_n^T) x \\ &= Ax. \end{aligned}$$

In the last line, we have used the fact that if $\{v_1, \dots, v_n\}$ is an orthonormal basis for $\mathbb{R}^n$, then $v_1 v_1^T + \dots + v_n v_n^T = I$ (exercise). $\square$

Example 3.3. (from Lay's book) Find a singular value decomposition of

$$A = \begin{pmatrix} 4 & 11 & 14 \\ 8 & 7 & -2 \end{pmatrix}.$$

Step 1. We first need to find the eigenvalues of $A^T A$. We compute that

$$A^T A = \begin{pmatrix} 80 & 100 & 40 \\ 100 & 170 & 140 \\ 40 & 140 & 200 \end{pmatrix}$$

We know that at least one of the eigenvalues is 0, because this matrix can have rank at most 2. In fact, we can compute that the eigenvalues are $\lambda_1 = 360$, $\lambda_2 = 90$, and $\lambda_3 = 0$. Thus the singular values of A are $\sigma_1 = \sqrt{360} = 6\sqrt{10}$, $\sigma_2 = \sqrt{90} = 3\sqrt{10}$, and $\sigma_3 = 0$. The matrix Σ in a singular value decomposition of A has to be a 2×3 matrix, so it must be

$$\Sigma = \begin{pmatrix} 6\sqrt{10} & 0 & 0 \\ 0 & 3\sqrt{10} & 0 \end{pmatrix}.$$

Step 2. To find a matrix V that we can use, we need to solve for an orthonormal basis of eigenvectors of $A^T A$. One possibility is

$$v_1 = \begin{pmatrix} 1/3 \\ 2/3 \\ 2/3 \end{pmatrix}, \quad v_2 = \begin{pmatrix} -2/3 \\ -1/3 \\ 2/3 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 2/3 \\ -2/3 \\ 1/3 \end{pmatrix}.$$

(There are seven other possibilities in which some of the above vectors are multiplied by -1 .) Then V is the matrix with v_1, v_2, v_3 as columns, that is

$$V = \begin{pmatrix} 1/3 & -2/3 & 2/3 \\ 2/3 & -1/3 & -2/3 \\ 2/3 & 2/3 & 1/3 \end{pmatrix}.$$

The matrix A' has rank 2, because all of its columns are points in the subspace $x_1 + x_2 + x_3 = 0$ (but the columns do not all lie in a 1-dimensional subspace). Now suppose we perturb A' to the matrix

$$A = \begin{pmatrix} 1.01 & 2.01 & -2 & 2.99 \\ -4.01 & 0.01 & 1.01 & 2.02 \\ 3.01 & -1.99 & 1 & -4.98 \end{pmatrix}$$

This matrix now has rank 3. But the eigenvalues of $A^T A$ are

$$\sigma_1^2 \approx 58.604, \quad \sigma_2^2 \approx 19.3973, \quad \sigma_3^2 \approx 0.00029, \quad \sigma_4^2 = 0.$$

Since two of the singular values are much larger than the others, this suggests that A is close to a rank 2 matrix.

For more discussion of how SVD is used to analyze data, see e.g. Lay's book.

5 Exercises (some from Lay's book)

- (a) Find a singular value decomposition of the matrix $A = \begin{pmatrix} 2 & -1 \\ 2 & 2 \end{pmatrix}$.
(b) Find a unit vector x for which $\|Ax\|$ is maximized.
- Find a singular value decomposition of $A = \begin{pmatrix} 3 & 2 & 2 \\ 2 & 3 & -2 \end{pmatrix}$.
- (a) Show that if A is an $n \times n$ symmetric matrix, then the singular values of A are the absolute values of the eigenvalues of A .
(b) Give an example to show that if A is a 2×2 matrix which is not symmetric, then the singular values of A might not equal the absolute values of the eigenvalues of A .
- Let A be an $m \times n$ matrix with singular values $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$. Let v_1 be an eigenvector of $A^T A$ with eigenvalue σ_1^2 . Show that σ_2 is the maximum value of $\|Ax\|$ where x ranges over unit vectors in $\mathbb{R}^n$ that are orthogonal to v_1 .
- Show that if $\{v_1, \dots, v_n\}$ is an orthonormal basis for $\mathbb{R}^n$, then $v_1 v_1^T + \dots + v_n v_n^T = I$.
- Let A be an $m \times n$ matrix, and let P be an orthogonal $m \times m$ matrix. Show that PA has the same singular values as A .

FROM EECS 16-B

Relevant Videos Shared

by Prof. Saha...

Topic: Moving into Spectral Theorem & SVD.

If M is an $n \times n$ matrix with real eigen values then we can find an orthonormal basis $U = [\vec{u}_1, \vec{u}_2, \dots, \vec{u}_n]$ such that $U^T M U = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix}$ (Upper triangular matrix)

Changing coordinates doesn't change eigen values!

$M = U \tilde{M} U^T$
what happens if you take transpose on both sides?

$$M^T = U \tilde{M}^T U^T$$

What if M were symmetric to begin with i.e. $M = M^T$ which means that $U^T M U = U^T M^T U$

$$\tilde{M} = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} \quad \tilde{M}^T = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix}$$

If LHS = RHS, then stuff = 0 which means that if M was symmetric $\tilde{M}$ will be a DIAGONAL matrix (and not an upper triangle one)

So, REAL if M was symmetric $\Rightarrow$ has real eigen values & then $M = U \Lambda U^T$ i.e. eigen vectors are all orthogonal to one another.

Remember, symmetric condition can be easily checked BUT, checking for real eigen values is not that straightforward.

SAHA'S PROOF OF SPECTRAL THEOREM.

Do symmetric matrices have truly complex eigenvalues?

Let λ be an eigenvalue and $\vec{v}$ be a corresponding eigenvector such that

$$\lambda \vec{v} = S \vec{v}$$

If S is real, any complex in the LHS, it would be in $\vec{v}$ (comparing with RHS)

Nexts consider $\vec{v}^* = \vec{v}^T$ conjugating LHS & RHS.

$$\vec{v}^* \vec{v} = \sum_{i=1}^n \vec{v}^*(i) \cdot \vec{v}(i) = \sum_{i=1}^n |\vec{v}(i)|^2 \Rightarrow \text{MAGNITUDE}$$

$$\vec{v}^* S \vec{v} = \lambda \vec{v}^* \vec{v} = \lambda \text{ real number}$$

$$\vec{v}^* S \vec{v} = \lambda \vec{v}^* \vec{v} = \lambda \text{ real number}$$

$$\vec{v}^* S \vec{v} = \lambda \vec{v}^* \vec{v} = \lambda \text{ real number}$$

So, ALL REAL SYMMETRIC MATRICES HAVE REAL EIGEN VALUES AND A REAL SYMMETRIC MATRIX IS DIAGONIZABLE BY ORTHONORMAL EIGEN BASIS!

FROM SAHA'S NOTES ON SVDs

https://inst.eecs.berkeley.edu/~ee16b/fa19/notes/notes13.pdf

PLANNING:

$$\vec{x}(i+1) = A \vec{x}(i) + B \vec{u}(i)$$

$\Rightarrow$ Start at $\vec{x}(0) = \vec{x}_0$
 $\Rightarrow$ and our desired goal $\Rightarrow$ to get to $\vec{x}(T) = \vec{x}_T$

So, Planning problem is to pick $\vec{u}(0), \vec{u}(1), \dots, \vec{u}(T-1)$

Controllability tells us that

$$C = [B, AB, \dots, A^{n-1}B]$$

suppose $\vec{0} = U$ (scalar), $B = \vec{b}$

$$\Delta = 100ms$$

$$n = 10$$

$$T = 10 \text{ seconds} \Rightarrow 100 \text{ controls i.e. } \frac{10}{0.1}$$

$$\vec{x}_0 = A^{100} \vec{x}_0 + \underbrace{[B, AB, \dots, A^{99}B]}_{10 \times 100} \begin{bmatrix} u(0) \\ u(1) \\ \vdots \\ u(99) \end{bmatrix}$$

NATURAL $\Rightarrow$ 'LEAST EFFORT'

want

$$\vec{u}^* = \arg \min \|\vec{u}\|$$

$$C_{100} \vec{u} = \vec{x}_0 - A^{100} \vec{x}_0$$

minimizing length

$\Rightarrow$ Geometry tells us that $\Rightarrow$ we want $\vec{u}^*$ (not in the null space of C_{100})

IMPORTANT $\Rightarrow C_{100}$ is a fat matrix so it means it has a Null space.

C_{100} is NOT A SQUARE MATRIX so we can break it down into eigen basis? ABSOLUTELY NOT.

C_{100} is NOT SQUARE so it doesn't have any eigen vectors!

A non-square matrix doesn't have any eigen basis $\Rightarrow$

NOTHING TIMES A NON-SQUARE MATRIX WILL BE A MULTIPLE OF ITSELF

e.g.

$$A \vec{x} = \lambda \vec{x} \quad \text{but,}$$

$A \vec{x} \neq \lambda \vec{x}$ because the result will NOT have the same dimensions $\Rightarrow$ VERY IMP.

So, we want a basis that has the property that the first 10 vectors are orthogonal & not in the Nullspace & 90 vectors are orthogonal & in the Null Space.

If I have a U that is orthonormal, then

$$\vec{0} = V^T U$$

$$\text{original } \vec{v} \begin{pmatrix} U \\ \vec{0} \end{pmatrix} V$$

$$\|\vec{0}\|^2 = \vec{0}^T \vec{0}$$

$$= U^T V \cdot V^T U$$

$$= \vec{0}^T \vec{0}$$

$$= \|\vec{0}\|^2$$

length doesn't change $\Rightarrow$ hence orthonormal basis are sought.

optimization problem becomes $\vec{u}^* = \arg \min \|\vec{u}\|^2$

$$\text{s.t. } C_{100} V \vec{u} = \vec{x}_0 - A^{100} \vec{x}_0$$

So, how do we get these basis?

$\Rightarrow$ running Gram Schmidt
 $\Rightarrow$ use the fact that a symmetric real matrix has orthonormal eigenvectors.

IMPORTANT

How can I find an S that is symmetric (100x100)
[null spaces are eigenspaces that correspond to eigen values = 0]

where S has the same nullspace as C_{100}

So, (to get a matrix S that has the same nullspace as C_{100}) $\Rightarrow$ we would make S to assume that C_{100} is a part of that matrix

$$S = \begin{bmatrix} C_{100} \\ 0 \end{bmatrix}$$

Why will S be symmetric?

$$S^T = S = (C_{100}^T C_{100})^T = C_{100}^T (C_{100}^T)^T = C_{100}^T C_{100} = S$$

AND, if $C_{100} \vec{v} = \vec{0}$ then

$$S \vec{v} = \vec{0}$$

if $S \vec{v} = \vec{0}$ then

$$C_{100} C_{100} \vec{v} = \vec{0}$$

$$C_{100} C_{100} \vec{v} = \vec{0}$$

$$(C_{100} \vec{v})^T (C_{100} \vec{v}) = \vec{0}$$

$$\|\vec{0}\|^2 = 0$$

$\Rightarrow$ sum of squares of components can only be zero if all the components are zero!

if something is in the Null space of C_{100} it is also in the Null space of S
AND
if something is in the Null space of S , it is in the Null space of C_{100} because the Null spaces are the same!

So, by the Spectral Theorem for real, symmetric matrices, S has eigenvalues $\vec{v}$ that are orthonormal. i.e. $S \vec{v}_i = \lambda_i \vec{v}_i$ & $\vec{v}_i^T \vec{v}_j = \delta_{ij}$

$$V^T S V = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} \Rightarrow \text{Choose order so } \lambda_1, \lambda_2, \dots, \lambda_n = 0$$

Recall

$$\tilde{C}_{100} = \begin{bmatrix} C_{100} & 0 \\ 0 & 0 \end{bmatrix}$$

I want the first 10 columns of $\tilde{C}_{100}$

$$\text{The } i\text{th column of } \tilde{C}_{100} = C_{100} \vec{u}_i$$

$$\text{Consider } (C_{100} \vec{u}_i)^T (C_{100} \vec{u}_j)$$

$$= \vec{u}_i^T C_{100}^T C_{100} \vec{u}_j = \vec{u}_i^T \lambda_j \vec{u}_j = \lambda_j \vec{u}_i^T \vec{u}_j = \begin{cases} \lambda_j & \text{if } i=j \\ 0 & \text{if } i \neq j \end{cases}$$

So, the columns that we computed U are all ORTHONORMAL to one another!

So,

$$C_{100} \vec{u}_i \text{ has } \|C_{100} \vec{u}_i\| = \sqrt{\lambda_i} \Rightarrow \text{if } i \neq j, \text{ length square can never be } -1!$$

$$\text{Define } \tilde{\vec{u}}_i = \frac{C_{100} \vec{u}_i}{\sqrt{\lambda_i}} \quad \left\{ \begin{array}{l} 10 \text{ long vectors} \end{array} \right\} \text{ ORTHONORMAL}$$

Recall our goal $\Rightarrow \vec{g} = \vec{x}_0 - A^{100} \vec{x}_0$

$$(a) \text{ compute } \alpha_i = \tilde{\vec{u}}_i^T \vec{g}$$

$$(b) \text{ choose } \vec{0}[i] = \frac{\alpha_i}{\sqrt{\lambda_i}}$$

$$(c) \text{ apply control } \vec{u}^* = \sum_{i=1}^{10} \vec{0}[i] \tilde{\vec{u}}_i = \sum_{i=1}^{10} \vec{u}_i \cdot \frac{1}{\sqrt{\lambda_i}} \cdot \tilde{\vec{u}}_i \cdot \sqrt{\lambda_i}$$

Now,

INTERPRETING THE SVD: assemble the single most important way of looking at a matrix! [generalizable to non-square matrices]

$$C_{100} \cdot \vec{0} = \tilde{C}_{100} \vec{0} = \sum_{i=1}^{10} \omega_i \cdot \sqrt{\lambda_i} \cdot \vec{u}_i^T \vec{0} \cdot \vec{u}_i$$

$$= \left(\sum_{i=1}^{10} \omega_i \cdot \sqrt{\lambda_i} \cdot \vec{u}_i^T \right) \vec{0}$$

Outer product form of SVD:

$$C_{100} = \sum_{i=1}^{10} \omega_i \cdot \sqrt{\lambda_i} \cdot \vec{u}_i^T$$

So, in matrix form: FULL MATRIX FORM OF SVD

$$C_{100} \vec{u} = \begin{bmatrix} \omega_1 & & 0 \\ & \ddots & \\ 0 & & \omega_n \end{bmatrix} \begin{bmatrix} \vec{u}_1 \\ \vdots \\ \vec{u}_n \end{bmatrix} \begin{bmatrix} \vec{u}_1^T \\ \vdots \\ \vec{u}_n^T \end{bmatrix}$$

COMPACT FORM OF SVD

$$C_{100} = \begin{bmatrix} \omega & & 0 \\ & \ddots & \\ 0 & & \omega_n \end{bmatrix} \begin{bmatrix} \vec{u}_1 \\ \vdots \\ \vec{u}_n \end{bmatrix} \begin{bmatrix} \vec{u}_1^T \\ \vdots \\ \vec{u}_n^T \end{bmatrix}$$

An outer product CANNOT be of rank greater than 1

Wrapping up SVDs,

Suppose A is an $m \times n$ tall matrix i.e.

$$A = \begin{bmatrix} & & \\ & & \\ & & \end{bmatrix} \quad \text{where } m \geq n, \text{ then,}$$

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