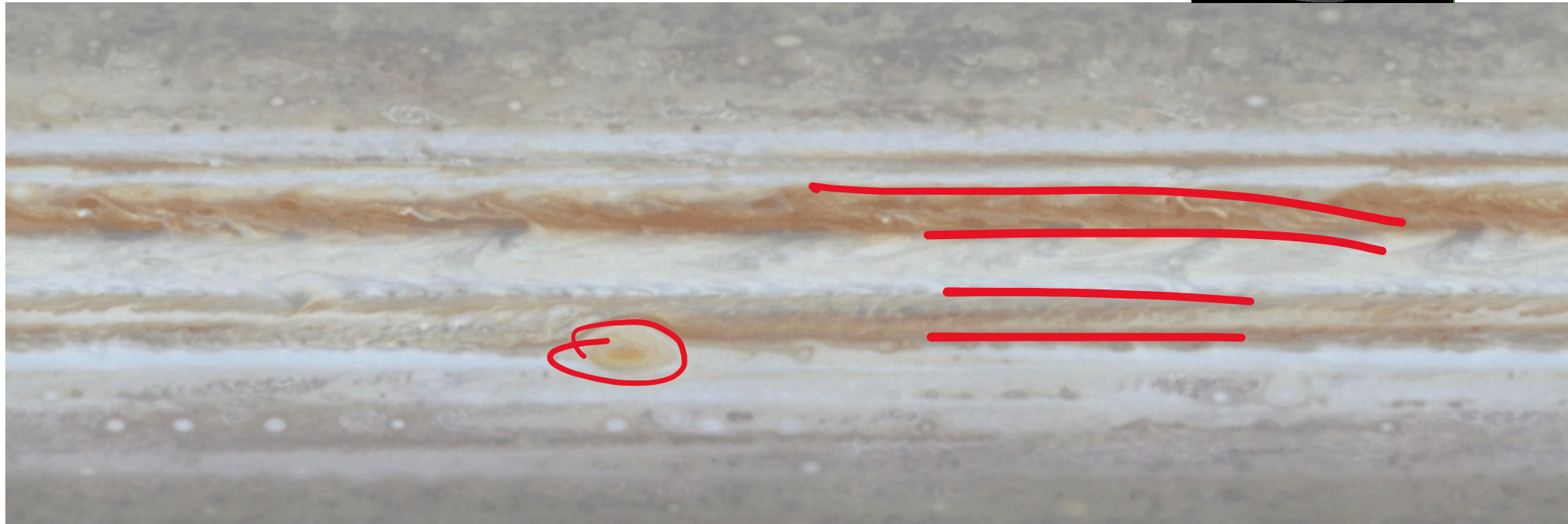
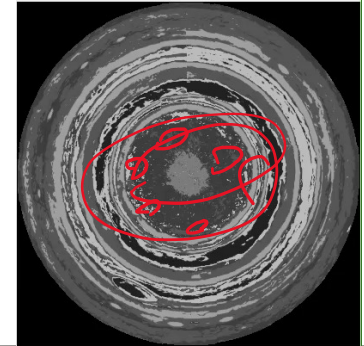
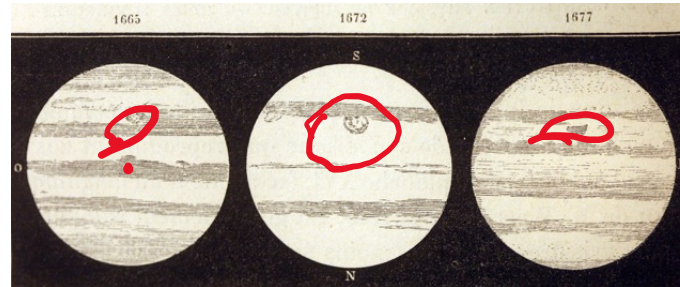


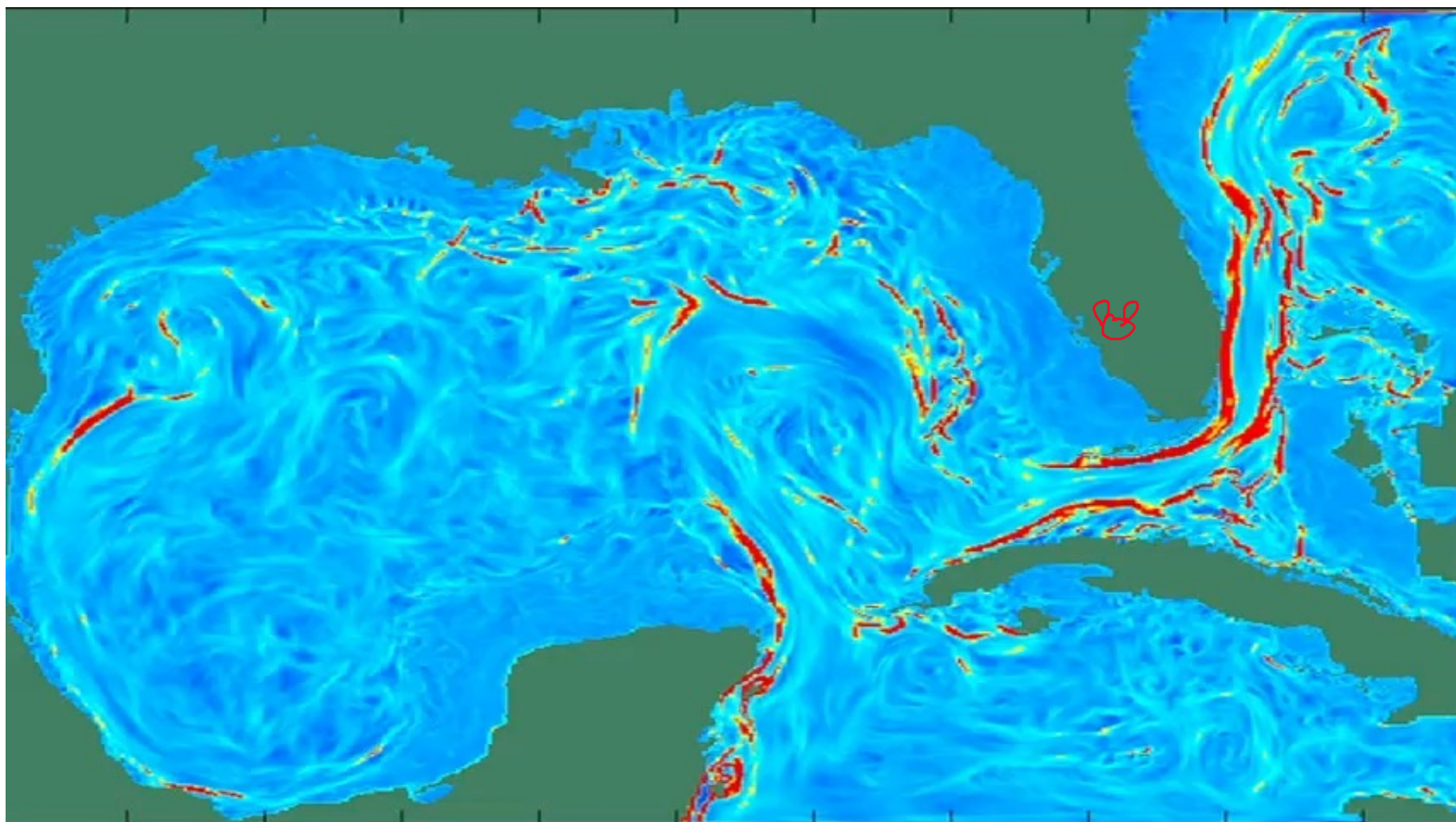
EE520 Data Driven Analysis of Complex Systems

Erik Bollt

~

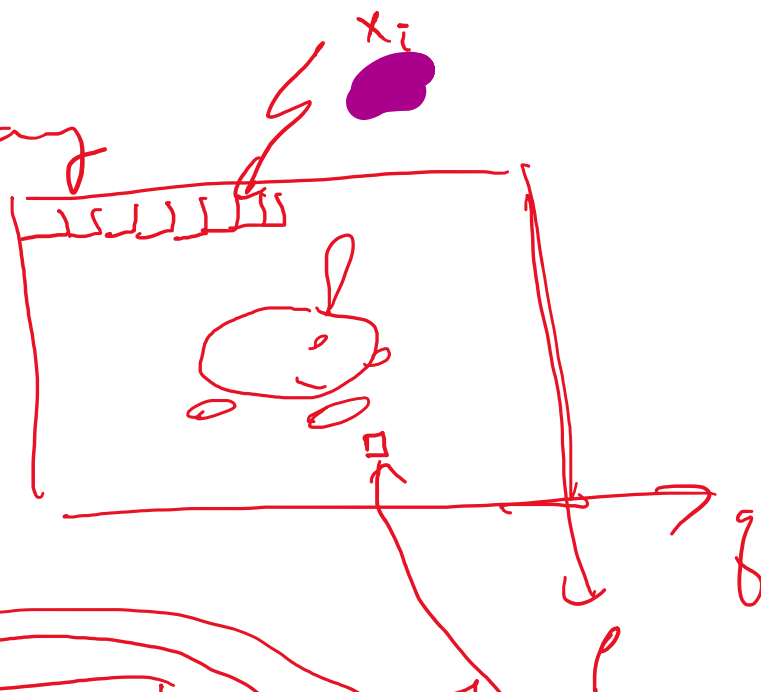
~





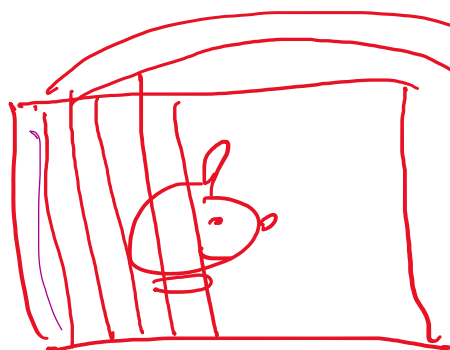
Data as
an array

$X =$

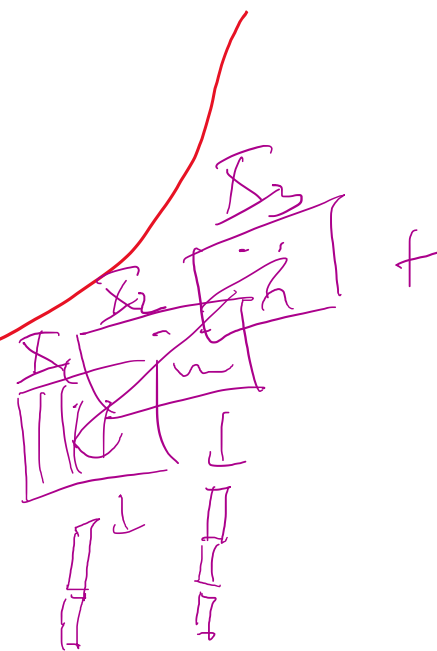


$X_{p \times q} \in \mathbb{R}^{p \times q}$
matrix

$[X]_{l,m} = \text{one pixel}$



slice & stack



Data as an array

On Matrix Multiplication

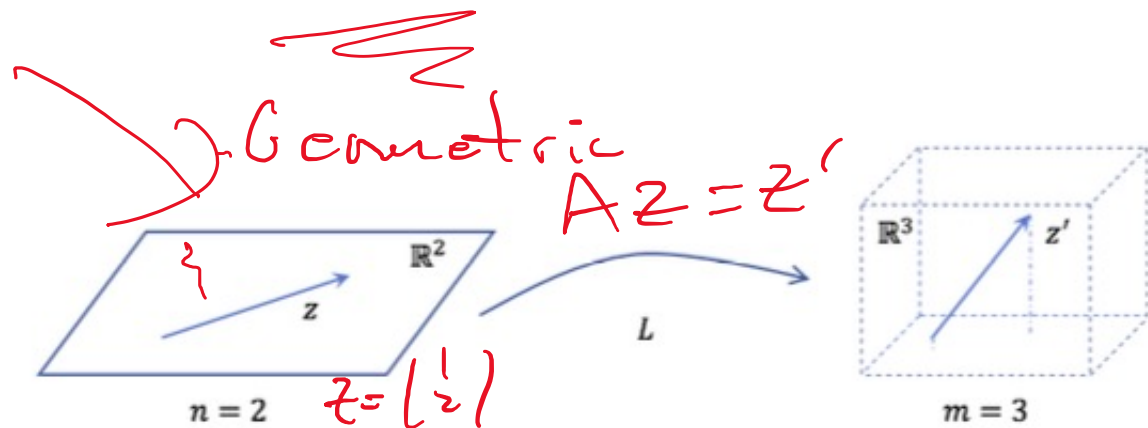
$$L(\mathbf{z}) : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

$$\mathbf{z} \mapsto \mathbf{z}' = A\mathbf{z},$$

in terms of the usual matrix times vector multiplication,

$$[\mathbf{z}']_i = \sum_{j=1}^n A_{i,j} [\mathbf{z}]_j, \text{ for each } i = 1, \dots, m,$$

a vector has length & direction.

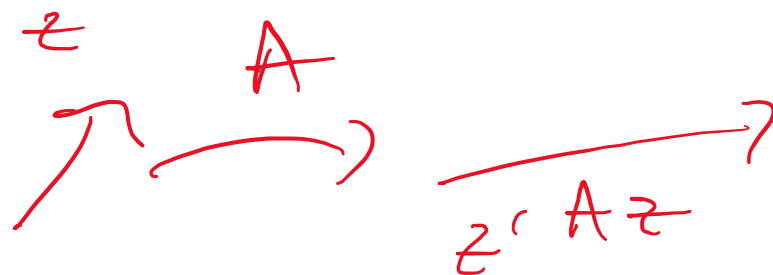


But as linear algebra

$$A_{2 \times 2} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}; \text{ matrices } \times \text{ vectors}$$

$$A \begin{pmatrix} 3 \\ 4 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \begin{pmatrix} 3 \\ 4 \end{pmatrix} = \begin{pmatrix} 1 \cdot 3 + 2 \cdot 4 \\ 3 \cdot 3 + 4 \cdot 4 \end{pmatrix} = \begin{pmatrix} 11 \\ 25 \end{pmatrix}$$

$$z' = Az$$



new direction, new length.

Eig for square -

$$AV = \lambda V$$

2×2



Characterize matrices by knowing just these
 • Eig. special directions

Linear

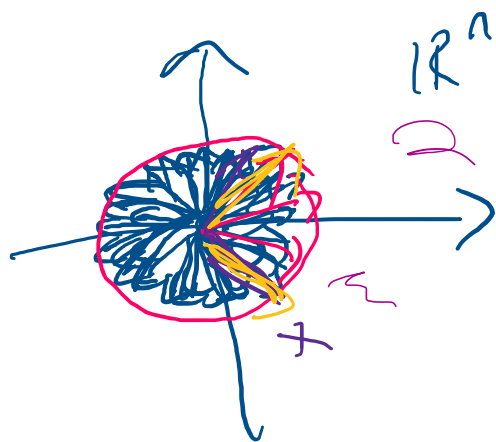
$$Au = A(a_1 v_1 + a_2 v_2) = a_1 A v_1 + a_2 A v_2 = a_1 \lambda_1 v_1 + a_2 \lambda_2 v_2$$

$$\det(A - \lambda I) = 0$$

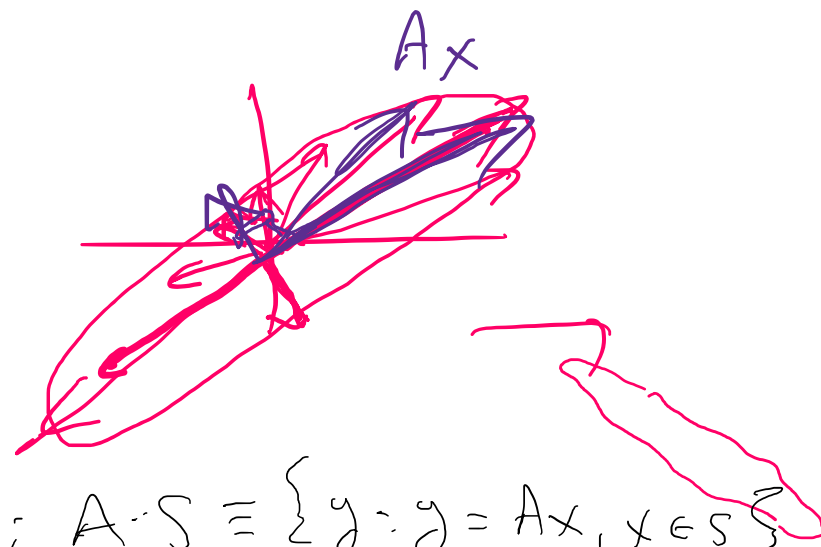
$$(A - \lambda I)v = 0$$

Matrix time circle =
 all vectors of length 1.

? Matrix $\times$ circle ? But matrix
 times vector. $\nearrow = \nearrow$



$$\begin{matrix} A \\ \hline Ax \end{matrix}$$



$$S = \{x \mid \|x\|_2 = 1, x \in E \cong \mathbb{R}^2\}; \quad A \cdot S = \{y \mid y = Ax, x \in S\}$$

Theorem 2.1.1 — Singular Value Decomposition. Let A be an $m \times n$ matrix whose entries come from the field $\mathcal{K}$, which is either the field of real numbers or the field of complex numbers. Then the singular value decomposition of A exists, and it takes the form of a product of matrices:

$$A_{m \times n} = U_{m \times m} \Sigma_{m \times n} V_{n \times n}^* \quad (2.5)$$

where

- U is an $m \times m$ unitary matrix.
- Σ is a diagonal $m \times n$ matrix with non-negative real numbers on the diagonal.
- V is an $n \times n$ unitary matrix, and V^* is the conjugate transpose of V .

The singular values are the nonnegative values: $\sigma_i \geq 0, i = 1, \dots, n$,

The left singular vectors: u_i are the columns of $U = [u_1, u_2, \dots, u_m]$.

The right singular vectors: v_i are the columns of $V = [v_1, v_2, \dots, v_n]$.

Definition 2.1.1 — Singular values and singular vectors. The singular values of A are the scalar values, σ_i , and the columns of U and V have columns that are the corresponding i^{th} left and right singular vectors, u_i and v_i :

The singular values are the nonnegative values: $\sigma_i \geq 0, i = 1, \dots, n$,

The left singular vectors: u_i are the columns of $U = [u_1, u_2, \dots, u_m]$.

The right singular vectors: v_i are the columns of $V = [v_1, v_2, \dots, v_n]$.

Since V is orthogonal, then right multiplying Eq. (2.5) by V ,

$$AV = U\Sigma V^*V = U\Sigma, \quad (2.8)$$

$$\Sigma := \text{diag}(\sigma_1, \sigma_2, \dots, \sigma_p), p = \min(m, n),$$

$$\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_p \geq 0.$$

$$\Sigma = \begin{pmatrix} \sigma_1 & & 0 \\ & \sigma_2 & \\ 0 & & \ddots \\ & & & \sigma_p \\ & & & & 0 \end{pmatrix}$$

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

$$u_i^T u_i = u_i \cdot u_i$$

$$u_i^T u_j = \begin{pmatrix} u_1^T & u_2^T & \dots & u_m^T \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \dots \\ u_m \end{pmatrix} = \begin{pmatrix} \sigma_1 & 0 & \dots & 0 \\ 0 & \sigma_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_p \end{pmatrix} = \delta_{ij} = \begin{cases} 1 & i=j \\ 0 & i \neq j \end{cases}$$

$$U^T U = \begin{pmatrix} u_1^T & u_2^T & \dots & u_m^T \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \dots \\ u_m \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix} = I$$

$$U \text{ unitary} \Leftrightarrow U^* U = U U^* = I$$

$$U^T U = \begin{pmatrix} u_1^T & u_2^T & \dots & u_m^T \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \dots \\ u_m \end{pmatrix} = I = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

$$A x = b$$

$$\begin{pmatrix} 1 & 2 & 3 \\ 2 & 4 & 6 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} b_1 \\ b_2 \end{pmatrix}$$

if A is square

$$A \underline{V} = \underline{U} \Sigma$$

$$[A] [v_1 v_2 \cdots v_n] = [u_1 u_2 \cdots u_n] \text{diag}(\sigma_1, \sigma_2, \cdots, \sigma_n).$$

■ **Example 2.1** Let $A = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 4 & 6 \end{pmatrix}_{2 \times 3}$. By SVD of the matrix A we have:

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

$$= \begin{pmatrix} \frac{1}{\sqrt{5}} & \frac{2}{\sqrt{5}} \\ \frac{-2}{\sqrt{5}} & \frac{1}{\sqrt{5}} \end{pmatrix} \begin{pmatrix} \sqrt{70} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} \frac{1}{\sqrt{14}} & \sqrt{\frac{2}{7}} & \frac{3}{\sqrt{14}} \\ \frac{-3}{\sqrt{10}} & 0 & \frac{1}{\sqrt{10}} \\ \frac{-1}{\sqrt{35}} & \sqrt{\frac{5}{7}} & \frac{-3}{\sqrt{35}} \end{pmatrix}. \quad (2.28)$$

We see that the second singular value, $\sigma_2 = 2$, meaning that number of non-zero singular values $r < \min\{m, n\}$. Such matrix is called rank deficient matrix. If we take the economy version (with $r = 1$) of the SVD we will have:

$$u_1 \sigma_1 v_1^T = \begin{pmatrix} \frac{1}{\sqrt{5}} \\ \frac{-2}{\sqrt{5}} \end{pmatrix} (\sqrt{70}) \begin{pmatrix} \frac{1}{\sqrt{14}} & \sqrt{\frac{2}{7}} & \frac{3}{\sqrt{14}} \end{pmatrix}$$

$$\approx \begin{pmatrix} 1 & 2 & 3 \\ 2 & 4 & 6 \end{pmatrix} \quad (2.29)$$

$$[A] [v_1 v_2 \cdots v_n] = [u_1 u_2 \cdots u_n] \text{diag}(\sigma_1, \sigma_2, \cdots, \sigma_n).$$

$$\underline{A^T A} = \underline{V} \underline{\Sigma^T \Sigma} \underline{V^T}$$

$$\underline{A^T A} \underline{V} = \underline{\Sigma^T \Sigma} \underline{V}$$

$$\underline{V} = \begin{pmatrix} v_1 & v_2 & \cdots & v_n \end{pmatrix}$$

$$AA^T = U \Sigma V^T (V \Sigma^T U^T)$$

$$= U \Sigma \Sigma^T U^T$$

$$(AA^T) \underline{U} = U(\Sigma \Sigma^T) = (\Sigma \Sigma^T) \underline{U}$$

$$\underline{U} = \begin{pmatrix} \underset{|}{u_1} & \underset{|}{u_2} & \dots & \underset{|}{u_n} \end{pmatrix}$$

Full *

The Economy SVD, and Reduced Rank SVD

The general SVD, Eq. (2.5) may be written in terms of submatrices.

Definition 2.1.3 — The Economy SVD. For any matrix $A \in \mathbb{R}^{m \times n}$, the general SVD Eq. (2.5) can be written in terms of smaller matrices,

$$A_{m \times n} = \hat{U}_{m \times n} \hat{\Sigma}_{n \times n} V_{n \times n}^* \quad (2.21)$$

and $U = [\hat{U}_{m \times n} | \hat{U}_{(n-m) \times n}]$, written in terms of an orthogonal "buffer" matrix

Definition 2.1.4 — Rank Deficient SVD. For a matrix $A \in \mathbb{R}^{m \times n}$ such that the SVD results in singular values

$$\sigma_r > \sigma_{r+1} = 0, \text{ for some } r < n. \quad (2.22)$$

then the SVD can be written in terms of an economy form as smaller matrices,

$$A_{m \times n} = \hat{U}_{m \times r} \hat{\Sigma}_{n \times n} V_{n \times r}^* \quad (2.23)$$

and related to the general SVD Eq. (2.5) by $U = [\hat{U}_{m \times r} | \hat{U}_{(n-r) \times n}]$, but $r < n$.

$$\sigma_1 > \sigma_2 > \dots > \sigma_r > \sigma_{r+1} = 0 \\ = 0 \dots$$

Rank ill conditioned

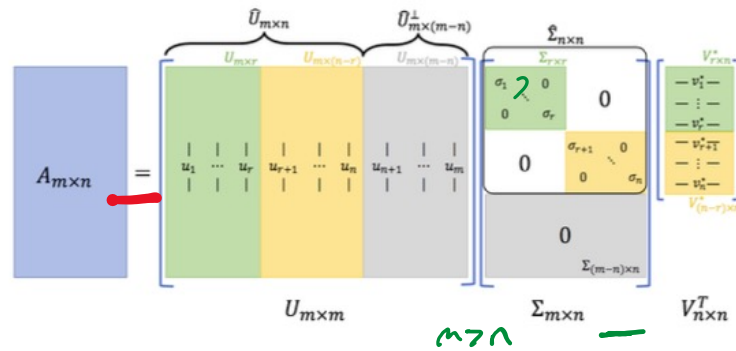


Figure 2.3: $m > n$ tall skinny

Recall that,

$$A_{m \times n} = \hat{U}_{m \times n} \hat{\Sigma}_{n \times n} \hat{V}_{n \times n}^T$$

$$= \begin{bmatrix} | & | & | & | \\ u_1 & u_2 & \dots & u_n \\ | & | & | & | \end{bmatrix} \begin{bmatrix} \sigma_1 & 0 & \dots & 0 \\ 0 & \sigma_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_n \end{bmatrix} \begin{bmatrix} - & v_1^T & - \\ - & v_2^T & - \\ - & \vdots & - \\ - & v_n^T & - \end{bmatrix} \quad (2.24)$$

but $V^T V = I$, orthogonality allows:

$$A_{m \times n} \hat{V}_{n \times n} = \hat{U}_{m \times n} \hat{\Sigma}_{n \times n} \quad (2.25)$$

so,

$$A_{m \times n} \begin{bmatrix} | & | & | & | \\ v_1 & v_2 & \dots & v_n \\ | & | & | & | \end{bmatrix} = \begin{bmatrix} | & | & | & | \\ u_1 & u_2 & \dots & u_n \\ | & | & | & | \end{bmatrix} \begin{bmatrix} \sigma_1 & 0 & \dots & 0 \\ 0 & \sigma_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_n \end{bmatrix} \quad (2.26)$$

but this just states n -matrix times vector statements:

$$\begin{aligned} Av_1 &= \sigma_1 u_1 \\ Av_2 &= \sigma_2 u_2 \\ &\vdots \\ Av_n &= \sigma_n u_n \end{aligned} \quad (2.27)$$

Full,
Economy,
Truncated
SVD

Handwritten notes:

$$b_{r+1} = 0$$

$$b_r = 10^{-6}$$

$$b_{r+1} \approx 10^{-13} < \epsilon$$

Full,
Economy,
Truncated
SVD

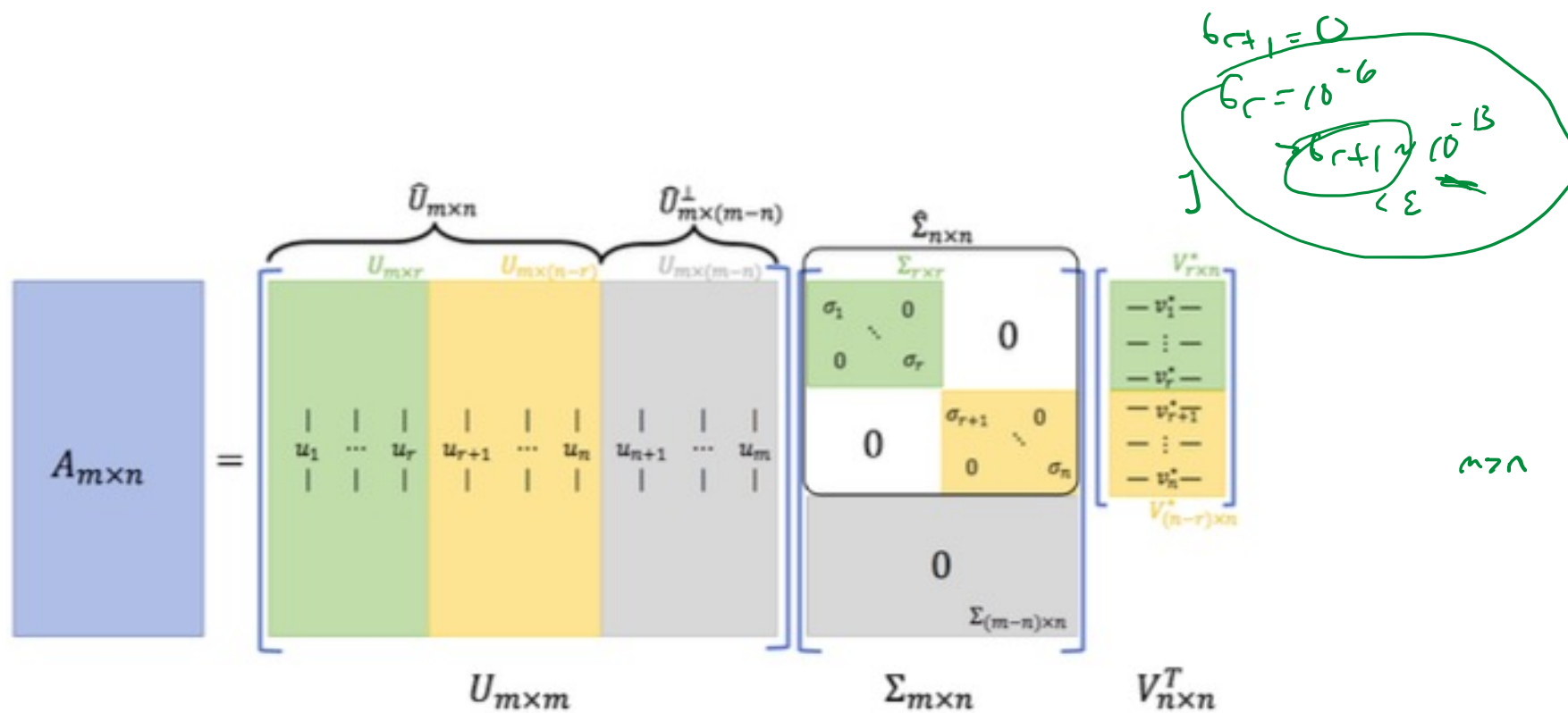


Figure 2.3: $m > n$ tall skinny

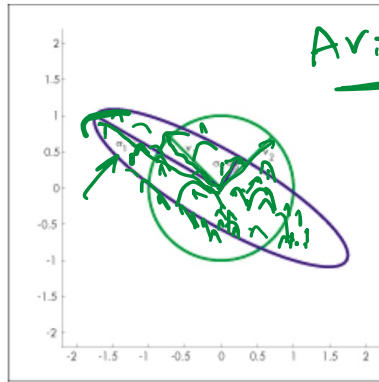
$$A = \underline{U} \underline{\Sigma} \underline{V}^T$$

$$A \underline{V} = \underline{U} \underline{\Sigma} \underline{e}$$

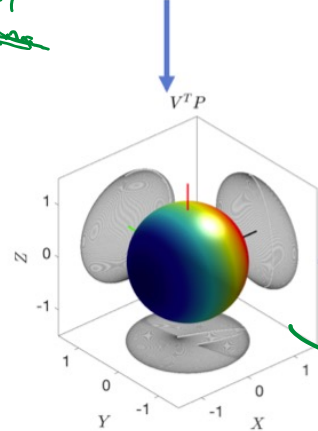
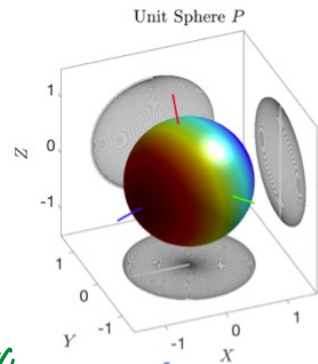
$$A \underline{v}_i = \underline{\sigma}_i \underline{u}_i$$

Geometry:.

1. V^* rotates to a standard configuration.
2. Σ stretches each orthogonal axis to the major covariance axis of the corresponding ellipsoid, and
3. U rotates results back to the configuration that associates with A .



$\sigma_1 \sim 1.5$
 $\sigma_2 \sim 0.5$
 $\sigma_3 = 2$
 $\sigma_4 = 0.0003$



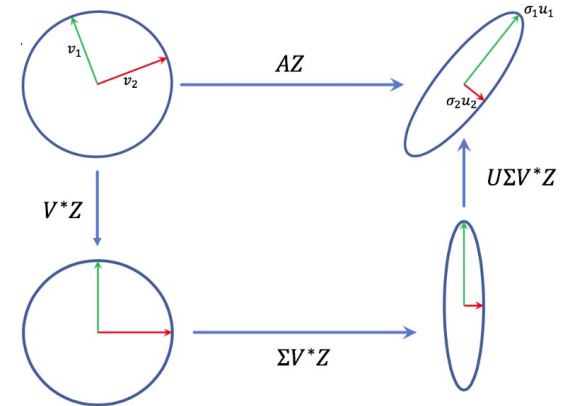
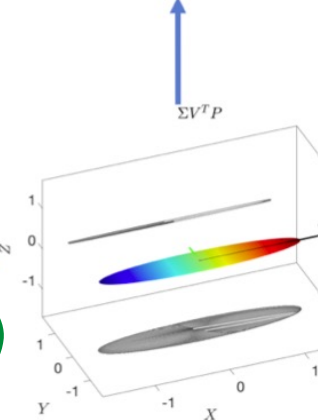
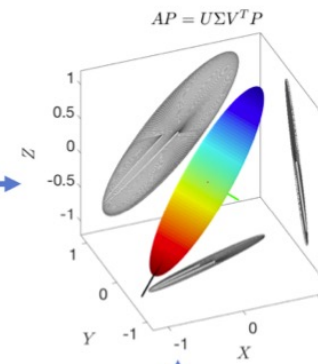
`disp(A)`
 0.8929 0.0417 0.1000
 0.3320 0.1077 0.0000
 0.8212 0.5951 0.0000

`[U,S,V] = svd(A)`

`U = 3x3`
 -0.6330 0.7469 0.2035
 -0.2590 0.0435 -0.9649
 -0.7296 -0.6635 0.1659

`S = 3x3`
 1.3438 0 0
 0 0.3912 0
 0 0 0.0208

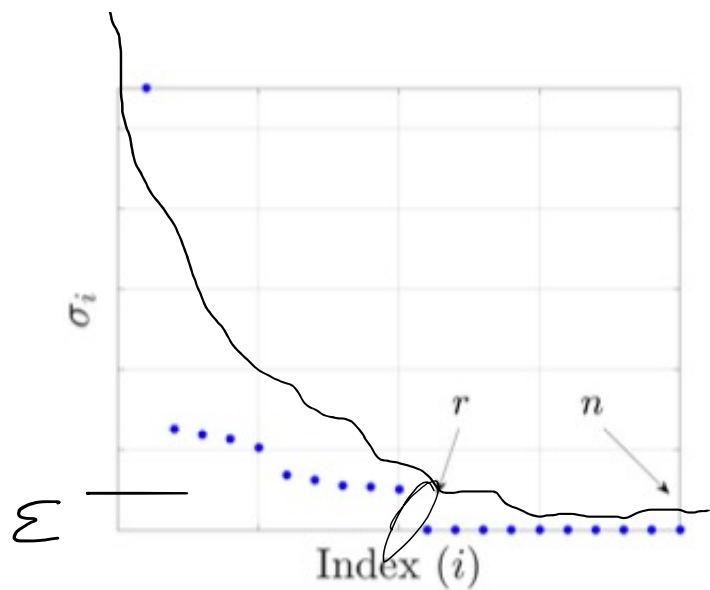
`V = 3x3`
 -0.9304 0.3488 -0.1126
 -0.3635 -0.9176 0.1612
 -0.0471 0.1909 0.9805



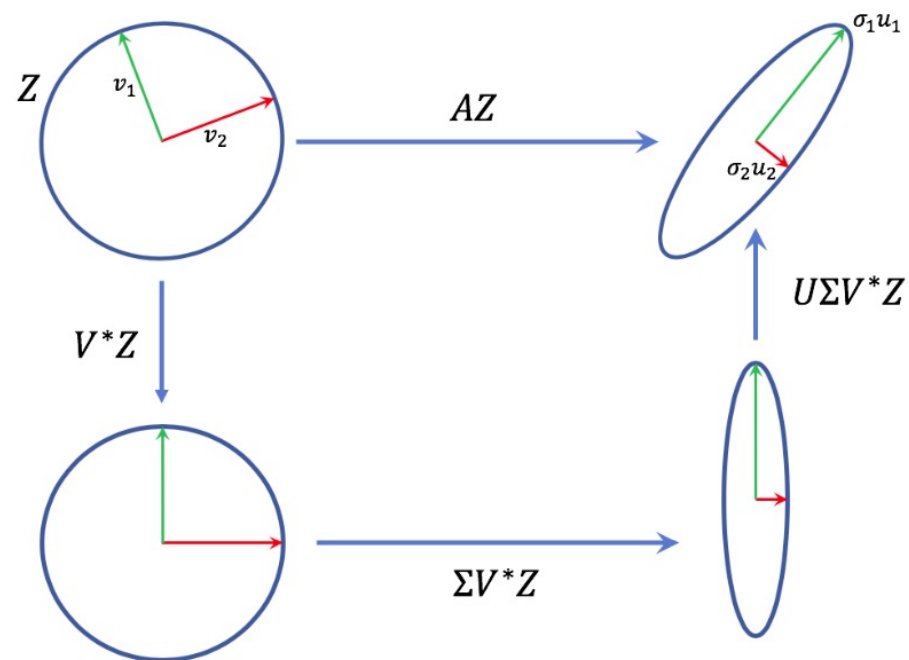
$$A \underline{w} = \underline{U} \underline{\Sigma} \underline{V}^T \underline{w}$$

$$\underline{v}_i^T \underline{w} = \underline{r}_i \cdot \underline{w}$$

And ROM



$$G \subset \mathcal{E} \supset G_{\text{cut}} \\ \sim 0$$



Bunny Compression

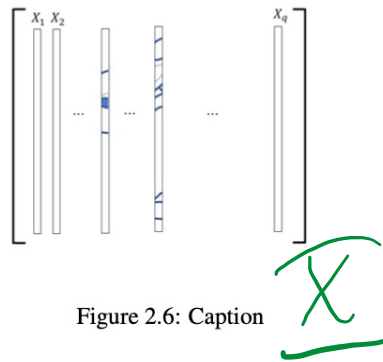
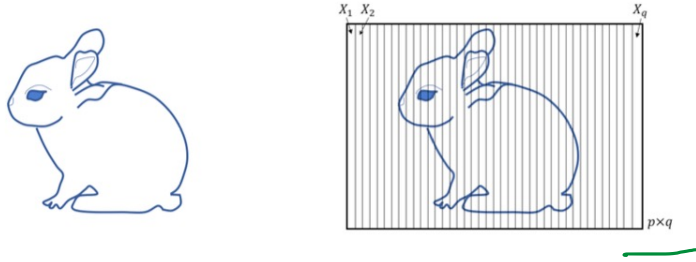
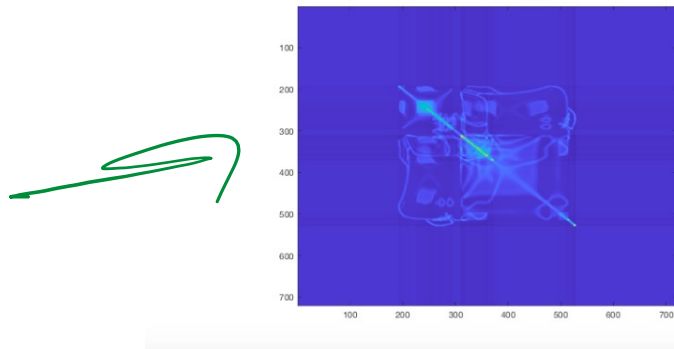


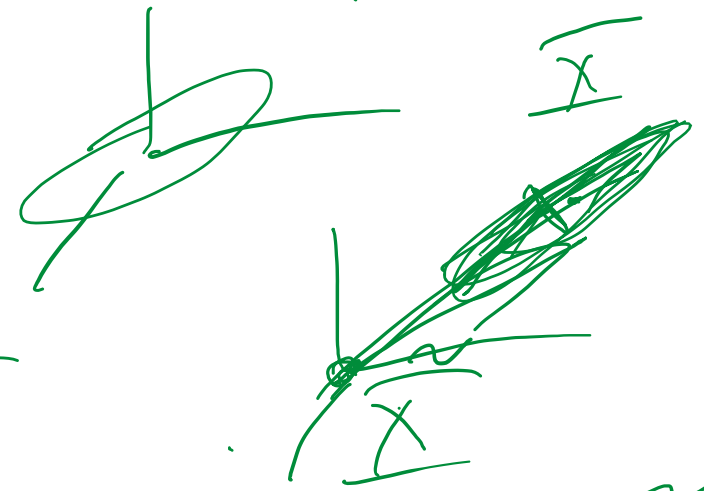
Figure 2.6: Caption



Covariance – notice the demean step

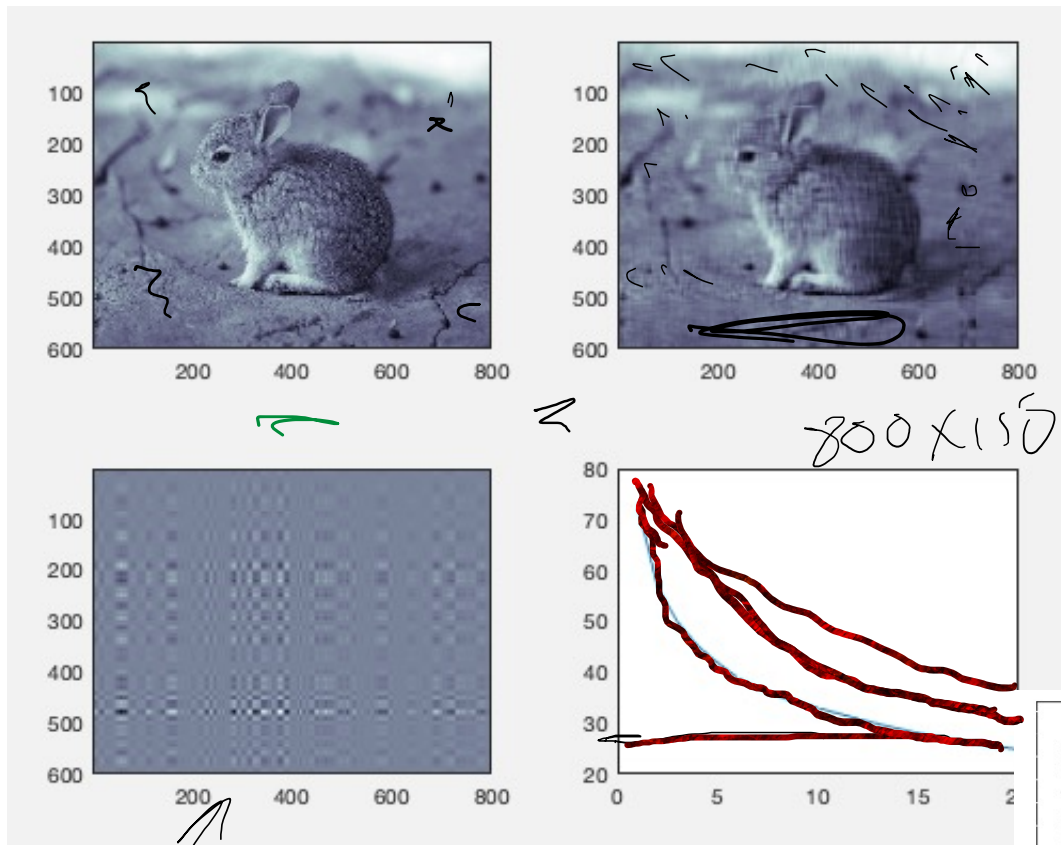
$$C_I = \frac{1}{n-1} (X - \tilde{X})^T (X - \tilde{X})$$

$\nearrow A^T$ A



$$\bar{X}^T \bar{X}$$

$$\bar{X} - \bar{X}$$



$$U(x,t) = \sum_{i=1}^{\infty} a_i(t) \phi_i(x)$$

as fast as possible.

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

Handwritten notes: 800×150 , U , Σ , V^T , U , Σ , V^T

```

1 I = imread('Bunny.jpg');
2
3 figure
4 subplot(1,2,1)
5 imshow(I)
6 xticks({}); yticks({});
7 pbaspect([1 1 1])
8 title('RGB Image')
9
10 I = rgb2gray(I); %Convert the 3D RGB color to 1D grayscale
11 I = im2double(I); %Convert integer value to double (scaled ...
    from 0 to 1)
12
13 subplot(1,2,2)
14 imshow(I)
15 xticks({}); yticks({});
16 pbaspect([1 1 1])
17 title('Grayscale Image')

```

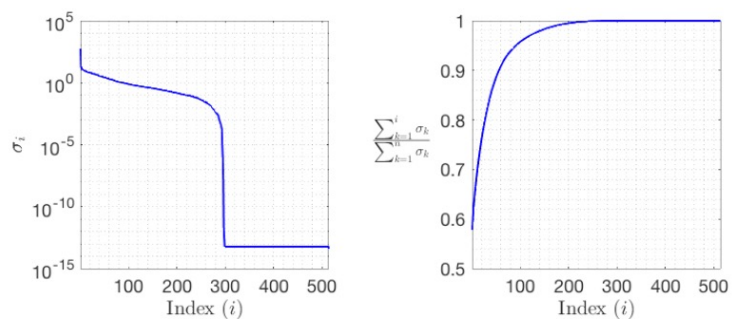
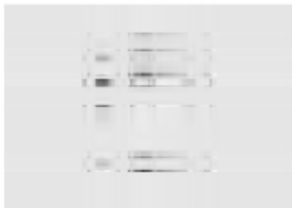


Figure 2.8: (Left) Singular Values. (Right) Energy

Rank $r = 2$



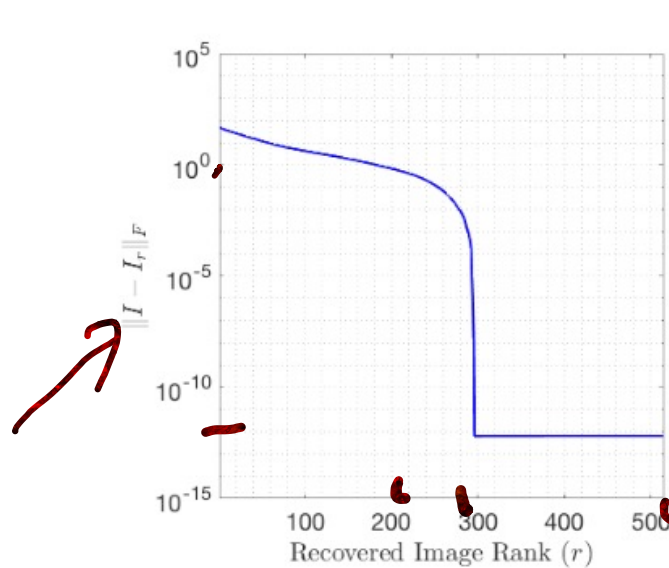
Rank $r = 10$



Rank $r = 20$



Rank $r = 100$



$$e_{ref} = \frac{\|I - I_r\|_F}{\|I\|_F}$$

: Distance $\|I - I_r\|_F$, where I_r is the recovered image using the reduced

Code 2.1: Read, convert, and display images.

```
1 I = imread('Bunny.jpg');
2
3 figure
4 subplot(1,2,1)
5 imshow(I)
6 xticks({}); yticks({});
7 paspect([1 1 1])
8 title('RGB Image')
9
10 I = rgb2gray(I); %Convert the 3D RGB color to 1D grayscale
11 I = im2double(I); %Convert integer value to double (scaled ...
    from 0 to 1)
12
13 subplot(1,2,2)
14 imshow(I)
15 xticks({}); yticks({});
16 paspect([1 1 1])
17 title('Grayscale Image')
```

History



Gene Golub's license plate, photographed by Professor P. M. Kroonenberg of Leiden University. Gene Howard Golub (February 29, 1932 – November 16, 2007), Fletcher Jones Professor of Computer Science at Stanford University. His work made fundamental contributions that have made the singular value decomposition practical as one of the most powerful and widely used tools in modern matrix computation.

Lots of Machine learning

& Data Analysis

is solving an ill-posed

- optimize a cost function.

Definition 2.1.2 — Induced Norm. Suppose a vector norm $\|\cdot\|$ on $\mathcal{K}^m$ is given. Any matrix $A_{m \times n}$ induces a linear operator from $\mathcal{K}^n$ to $\mathcal{K}^m$ with respect to the standard basis, and one defines the corresponding induced norm or operator norm on the space $\mathcal{K}^{m \times n}$ of all $m \times n$ matrices as follows:

$$\|A\|_p = \sup_{x \neq 0} \frac{\|Ax\|_p}{\|x\|_p} \quad (2.14)$$

or, taking a vector x such that $\|x\|_p = 1$, then we have

$$\|A\|_p = \sup_{\|x\|_p=1} \|Ax\|_p \quad (2.15)$$

Some Special (Simple) Matrix Norms

The first 3 of these are induced norms, but the 4th is not.

- For $p = 1$:

$$\|A\|_1 = \max_{1 \leq j \leq n} \sum_{i=1}^m |a_{ij}| \quad (2.16)$$

- For $p = \infty$:

$$\|A\|_\infty = \max_{1 \leq i \leq m} \sum_{j=1}^n |a_{ij}| \quad (2.17)$$

- A special case is the spectral norm when $p = 2$, in which we have:

$$\|A\|_2 = \sqrt{\lambda_{\max}(A^T A)} = \sigma_{\max} \quad (2.18)$$

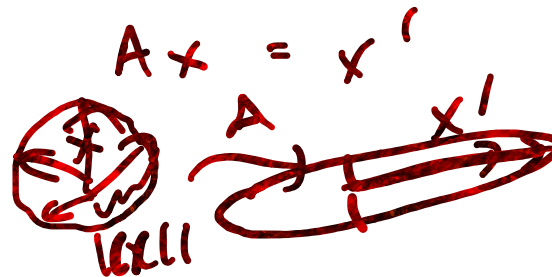
where $\sigma_{\max}$ is the maximum singular value of the matrix A .

- The Frobenius norm is given by:

$$\|A\|_F = \sqrt{\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2} = \sqrt{\sum_{i=1}^{\min\{m,n\}} \sigma_i^2} \quad (2.19)$$

Theorem 2.1.2 For a matrix A , the product of the singular values of A , equals the absolute value of its determinant:

$$|\det(A)| = \prod_{i=1}^n \sigma_i \quad (2.20)$$



$$p=1 : \|(x_1, x_2)\|_1 = |x_1| + |x_2|$$

$$p=\infty : \|(x_1, x_2)\|_\infty = \max_i |x_i|$$

$$p=2 : \|(x_1, x_2)\|_2 = \sqrt{x_1^2 + x_2^2}$$

$$x = \begin{pmatrix} 1 \\ 3 \end{pmatrix} ; \|x\|_1 = 1 + 3 = 4$$

$$\|x\|_\infty = 3$$

$$\|x\|_2 = \sqrt{1^2 + 3^2} = \sqrt{10}$$

Definition 2.1.2 — Induced Norm. Suppose a vector norm $\|\cdot\|$ on $\mathcal{K}^m$ is given. Any matrix $A_{m \times n}$ induces a linear operator from $\mathcal{K}^n$ to $\mathcal{K}^m$ with respect to the standard basis, and one defines the corresponding induced norm or operator norm on the space $\mathcal{K}^{m \times n}$ of all $m \times n$ matrices as follows:

$$\|A\|_p = \sup_{x \neq 0} \frac{\|Ax\|_p}{\|x\|_p} \quad (2.14)$$

or, taking a vector x such that $\|x\|_p = 1$, then we have

$$\|A\|_p = \sup_{\|x\|_p=1} \|Ax\|_p \quad (2.15)$$

Some Special (Simple) Matrix Norms

The first 3 of these are induced norms, but the 4th is not.

- For $p = 1$:

$$\|A\|_1 = \max_{1 \leq j \leq n} \sum_{i=1}^m |a_{ij}| \quad (2.16)$$

- For $p = \infty$:

$$\|A\|_\infty = \max_{1 \leq i \leq m} \sum_{j=1}^n |a_{ij}| \quad (2.17)$$

- A special case is the spectral norm when $p = 2$, in which we have:

$$\|A\|_2 = \sqrt{\lambda_{\max}(A^T A)} = \sigma_{\max} \quad (2.18)$$

where $\sigma_{\max}$ is the maximum singular value of the matrix A .

- The Frobenius norm is given by:

$$\|A\|_F = \sqrt{\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2} = \sqrt{\sum_{i=1}^{\min\{m,n\}} \sigma_i^2} \quad (2.19)$$

Theorem 2.1.2 For a matrix A , the product of the singular values of A , equals the absolute value of its determinant:

$$|\det(A)| = \prod_{i=1}^n \sigma_i \quad (2.20)$$

Fun facts about matrix estimation (data estimation)

$$\text{If } A, \quad b_1 \geq \dots \geq b_r > b_{r+1} = 0$$

$$\bullet \text{ range}(A) = \text{span}(u_1, u_2, \dots, u_r)$$

$$\bullet \text{ null}(A) = \text{span}(v_{r+1}, v_{r+2}, \dots, v_n)$$

$$\left. \begin{array}{l} \bullet \text{ range}(A) = \text{span}(u_1, u_2, \dots, u_r) \\ \bullet \text{ null}(A) = \text{span}(v_{r+1}, v_{r+2}, \dots, v_n) \end{array} \right\} \begin{array}{l} \text{A} \\ 0 \rightarrow \end{array}$$

$$\bullet \|A\|_2 = b_1 \quad ; \quad \|A\|_F = \sqrt{b_1^2 + b_2^2 + \dots + b_r^2}$$

$$\bullet A = \sum_{i=1}^r b_i u_i v_i^* = \underbrace{b_1 u_1 v_1^*}_{\text{rank-1}} + \underbrace{b_2 u_2 v_2^*}_{\text{outer products}} + \dots + b_r u_r v_r^*$$

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

$$\begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_m \end{pmatrix} \begin{pmatrix} v_1^T & v_2^T & \dots & v_n^T \end{pmatrix} \begin{pmatrix} b_1 & b_2 & \dots & b_r \end{pmatrix} \begin{pmatrix} -v_1^T \\ -v_2^T \\ \vdots \end{pmatrix}^{1 \times n}$$

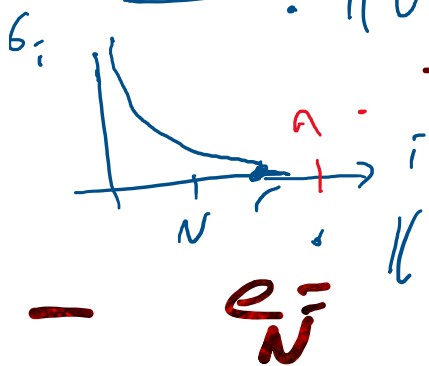
$$\begin{array}{l} \omega_1^T \omega_2 = \omega_1 \cdot \omega_2 \\ \omega_2 = \|\omega_1\| \|\omega_2\| \cos \theta \\ \begin{pmatrix} 1 & 2 \\ 2 & 4 \end{pmatrix} \end{array}$$

Matrix Estimation / Data Estimation. $A_{m \times n}$

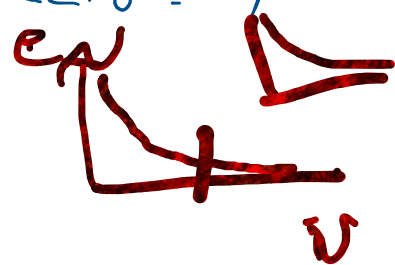
Let $0 \leq N \leq r$ and $A_N = \sum_{i=1}^N \sigma_i u_i v_i^*$
 (so we may be skipping some of them ... $\sum_{i=N+1}^r$)

Then $\|A - A_N\|_2 = \sigma_{N+1}$ (first one skipped)

(what if it zero?)



$$\|A - A_N\|_F = \sqrt{\sigma_{N+1}^2 + \sigma_{N+2}^2 + \dots + \sigma_r^2}$$

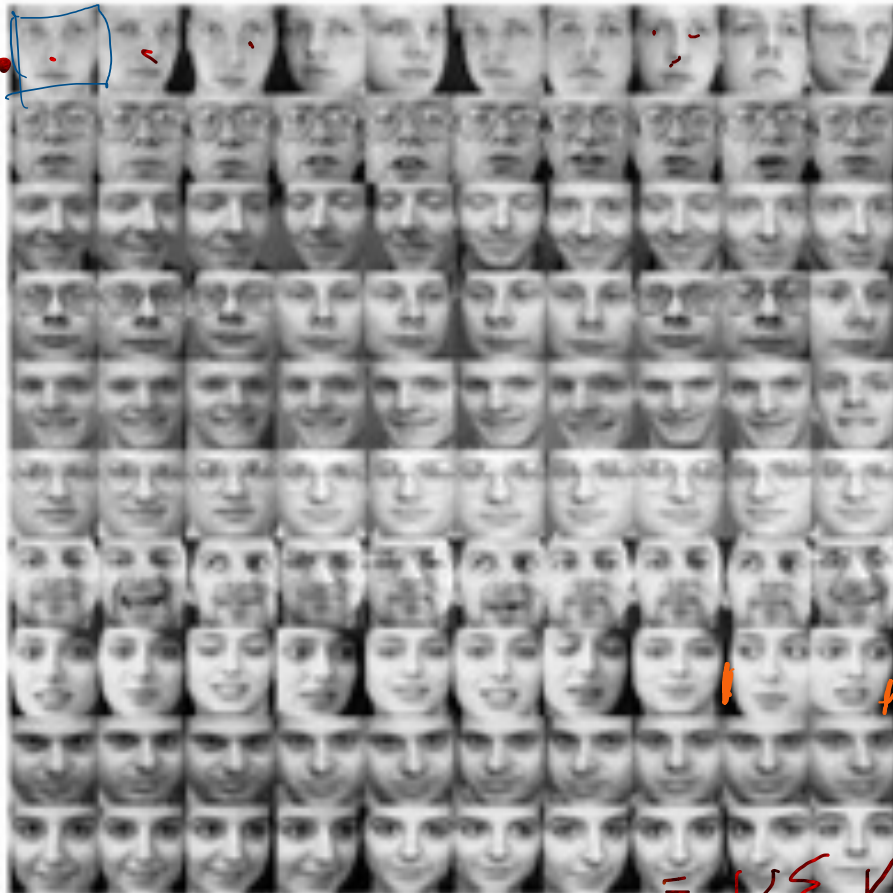


- EigenFace 1st Present

Eigenface



72 picture



i $\begin{bmatrix} \cdot \\ \cdot \\ \cdot \end{bmatrix}$
 $p \times q$
 array

reshape
as vector

$$\begin{bmatrix} 1 \\ \vdots \\ x_i \end{bmatrix}^T = \begin{bmatrix} 1 \\ 0 \\ 0 \\ \vdots \end{bmatrix}$$

$m \times 1$

$$M = P^T Q$$

$$\underline{X} = [x_1, x_2, \dots, x_n]_{m \times n}$$

$$= U \Sigma V^T$$

$$p = 1000, q = 2000$$

08/31/20

remove
variance



Register



$$E_1, E_2, E_3$$

$$x_i \cdot e_i = x_i^T e_i = E$$

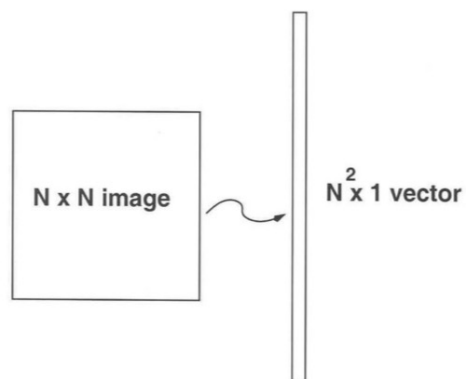
$$R^p \otimes R^q = R^r$$

Eigenfaces for Face Detection/Recognition

(M. Turk and A. Pentland, "Eigenfaces for Recognition", *Journal of Cognitive Neuroscience*, vol. 3, no. 1, pp. 71-86, 1991, hard copy)

• Face Recognition

- The simplest approach is to think of it as a template matching problem:



- Problems arise when performing recognition in a high-dimensional space.
- Significant improvements can be achieved by first mapping the data into a *lower-dimensionality* space.
- How to find this lower-dimensional space?

• Main idea behind eigenfaces

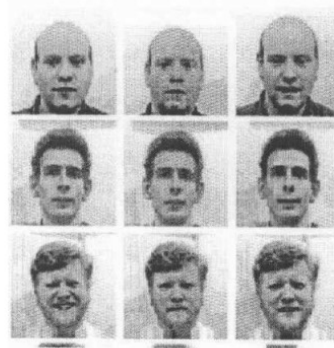
- Suppose Γ is an $N^2 \times 1$ vector, corresponding to an $N \times N$ face image I .
- The idea is to represent Γ ($\Phi = \Gamma$ - mean face) into a low-dimensional space:

$$\hat{\Phi} - \text{mean} = w_1 u_1 + w_2 u_2 + \cdots + w_K u_K \quad (K \ll N^2)$$

Computation of the eigenfaces

Step 1: obtain face images $I_1, I_2, \dots, I_M$ (training faces)

(**very important:** the face images must be *centered* and of the same *size*)



Step 2: represent every image I_i as a vector Γ_i

Step 3: compute the average face vector Ψ :

$$\Psi = \frac{1}{M} \sum_{i=1}^M \Gamma_i$$

Step 4: subtract the mean face:

$$\Phi_i = \Gamma_i - \Psi$$

Step 5: compute the covariance matrix C :

$$C = \frac{1}{M} \sum_{n=1}^M \Phi_n \Phi_n^T = AA^T \quad (N^2 \times N^2 \text{ matrix})$$

$$\text{where } A = [\Phi_1 \ \Phi_2 \ \dots \ \Phi_M] \quad (N^2 \times M \text{ matrix})$$

$$C = \frac{1}{M} X^T X$$

Step 6: compute the eigenvectors u_i of AA^T

The matrix AA^T is very large --> not practical !!

Step 6.1: consider the matrix $A^T A$ ($M \times M$ matrix)

Step 6.2: compute the eigenvectors v_i of $A^T A$

$$A^T A v_i = \mu_i v_i$$

What is the relationship between u_i and v_i ?

$$A^T A v_i = \mu_i v_i \Rightarrow AA^T A v_i = \mu_i A v_i \Rightarrow$$

$$CAv_i = \mu_i Av_i \text{ or } Cu_i = \mu_i u_i \text{ where } u_i = Av_i$$

Thus, AA^T and $A^T A$ have the same eigenvalues and their eigenvectors are related as follows: $u_i = Av_i$!!

Note 1: AA^T can have up to N^2 eigenvalues and eigenvectors.

Note 2: $A^T A$ can have up to M eigenvalues and eigenvectors.

Note 3: The M eigenvalues of $A^T A$ (along with their corresponding eigenvectors) correspond to the M *largest* eigenvalues of AA^T (along with their corresponding eigenvectors).

Step 6.3: compute the M best eigenvectors of AA^T : $u_i = Av_i$

(**important:** normalize u_i such that $\|u_i\| = 1$)

Step 7: keep only K eigenvectors (corresponding to the K largest eigenvalues)

Representing faces onto this basis

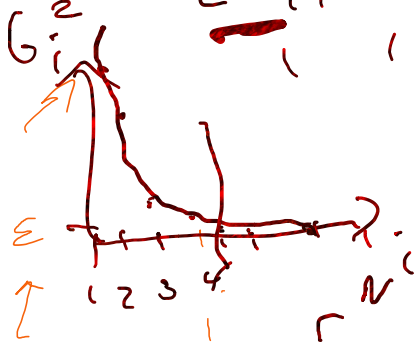
- Each face (minus the mean) Φ_i in the training set can be represented as a linear combination of the best K eigenvectors:

$$\hat{\Phi}_i - \text{mean} = \sum_{j=1}^K w_j u_j, \quad (w_j = u_j^T \Phi_i)$$

$\bullet$ images $cl(g)$
 $\bullet g = \text{reshape}(U(:, i), p, q)$

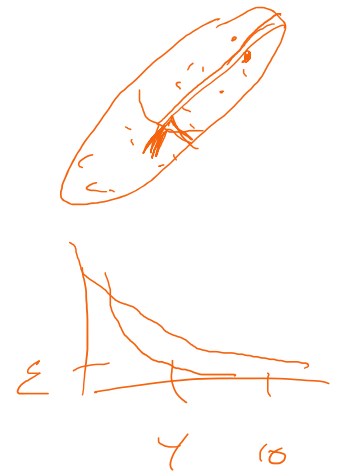
$$\underline{X} = U \Sigma V^T$$

$$U = [u_1 | u_2 | \dots | u_5 | u_{s+1} | \dots | u_r | u_{s+1} | \dots | u_n] \quad (\text{we call the } u_j \text{'s eigenfaces})$$



Each normalized training face Φ_i is represented in this basis by a vector:

$$\Omega_i = \begin{bmatrix} w_1^i \\ w_2^i \\ \dots \\ w_K^i \end{bmatrix}, \quad i = 1, 2, \dots, M$$



$$\| \vec{a} \|_2^2 = \sum_i a_i^2$$

Preserved 'neg.'

$$\| \vec{a} \|_2^2 = \sum_i a_i^2$$

$$\| \vec{a} \|_2^2 = \sum_i a_i^2$$

Fourier

$$f(x) = a_1 \cdot \text{Fourier} + a_2 \cdot \text{Fourier} + a_3 \cdot \text{Fourier} + \dots$$

Donny $\begin{matrix} & & & i \\ & & & | \\ \begin{pmatrix} | & | & | \\ \hline | & | & | \\ \hline | & | & | \end{pmatrix} & \rightarrow & \underline{X} = [x_1^T | x_2^T | \dots | x_N^T] \\ \hline x_1 & x_2 & x_N & n \times n \end{matrix}$

$\leftarrow$ one column

$$\underline{X} = \underline{U} \underline{\Sigma} \underline{V}^T$$



$$x_i = a_1^i u_1 + a_2^i u_2 + \dots + a_N^i u_N$$

$\leftarrow$ bases
full of as possible for the x_i set

$\underline{a}_i = (x_i^T \underline{u}_1) = \text{comp}_{u_1} x_i = \frac{\underline{u}_1 \cdot x_i}{\|\underline{u}_1\|} = \frac{1}{\|\underline{u}_1\|} (x_i \cdot \underline{u}_1)$

- On PCA

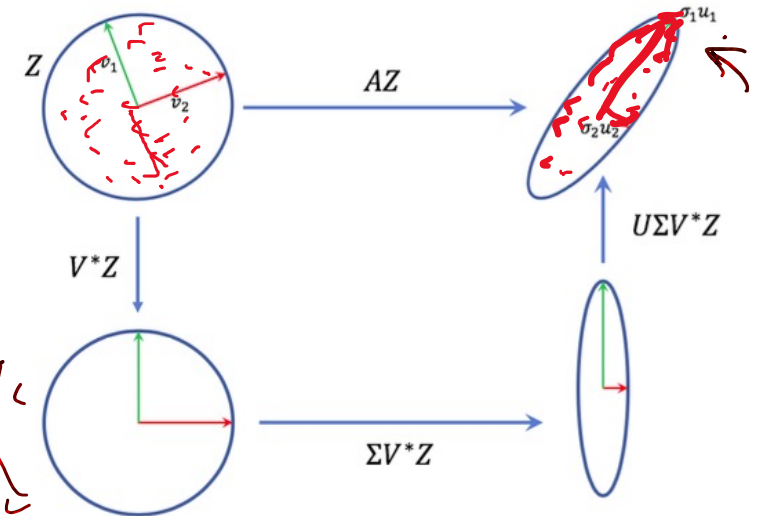
- Optimal

On PCA Principal Component Analysis, Eigenface

-On Raleigh Ritz Quotient

-On Spectral Decomposition Theorem

-On Data Clouds



$$\overline{X} = [x_1 | x_2 | \dots | x_n] \quad \vec{x}_i = a_1^i \vec{v}_1 + a_2^i \vec{v}_2 + \dots + a_r^i \vec{v}_r$$

$\vec{v}_i$ as basis set.

$x_i = \text{PCA gives basis set}$
 where $\langle a_j^i \rangle_i$ as fast as possible vs any other basis set.

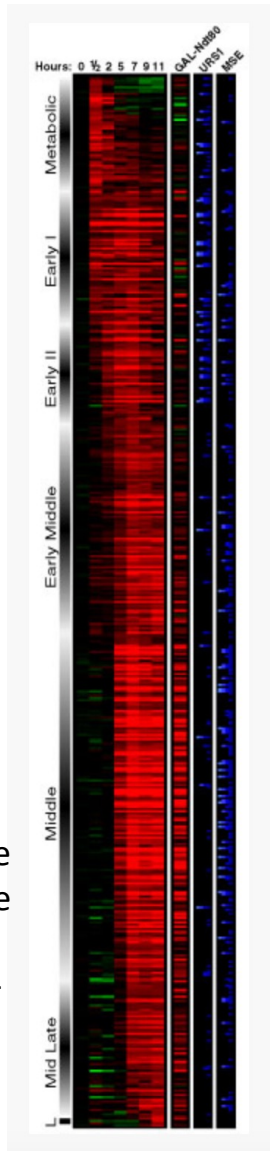
$\langle a_i^i \rangle_i = \frac{1}{n} \sum_{i=1}^n a_i^i$

The diagram shows a 3D plot of a data cloud (red dots) and its projection onto a 2D plane (green axes). The projection is labeled with a_i^i .

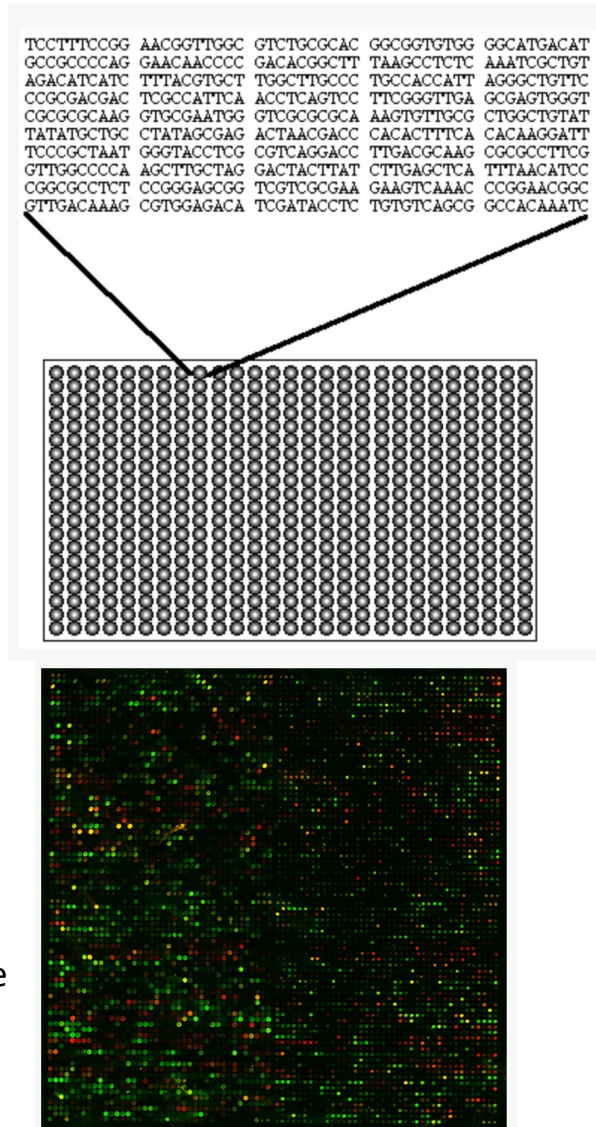
DNA Microarrays

Gene Expression

Microarray results that have been analyzed such that the colors were linked with expression and then similar gene profiles were grouped Together – budding yeast



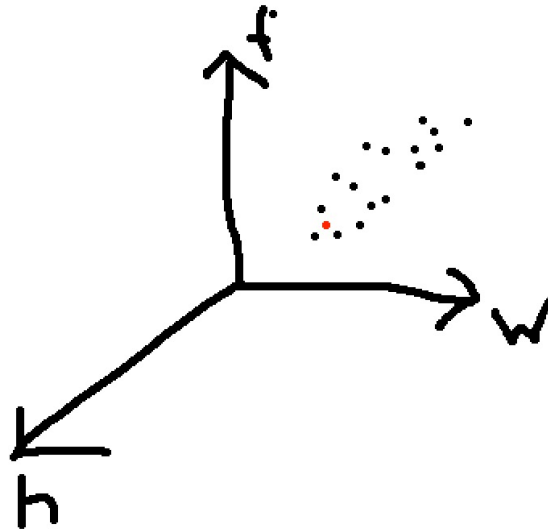
cartoon illustrating an array of DNA snippets on a chip. The top portion depicts a possible nucleotide sequence for the DNA segment immobilized in the position indicated.



DNA Microarray chip containing the entire yeast genome

Morphological

- Height
- Weight
- Footsize
- Belt (waist) size
- Hand size
- Forearm size
- Head circumference
- Femur length



Interpreting as an ellipsoid in the high dimensional space is the simplest geometric interpretation of the data cloud and leads to simplification as major and minor axis, and even Reduced order model (ROM) (meaning a lower dimensional representation).

PCA, SVD, SDT – is optimal

Data for PCA - "Pretend data looks like an ellipsoid"
 Ex. $\underline{X}_i \sim 4000 \times 1$ gene expression table for each i .
 $i = 1 \dots 216$ patients

$$X_i = \begin{bmatrix} x_i^1 \\ x_i^2 \\ \vdots \\ x_i^{4000} \end{bmatrix}$$

$y_i = 0$ or 1 "0" if not cancer "1" if cancer.



$$\mathbb{Z}_2 = \{0, 1\}$$

• Supervised vs. unsupervised.

just x $\xrightarrow{\text{unsupervised}}$ - just input - just structural geometry.
 "Data Cloud" $\sim$ distribution of R.V. $x \sim \underline{X}$
 $\xrightarrow{\text{level sets of } p}$ $p(x): \mathbb{R}^{4000} \rightarrow \mathbb{R}^+$
 • supervised learning is descriptive $f: x \rightarrow y$ $\leftarrow$ labels

THE SPECTRAL DECOMPOSITION

Let A be a $n \times n$ symmetric matrix. From the spectral theorem, we know that there is an orthonormal basis $u_1, \dots, u_n$ of $\mathbb{R}^n$ such that each u_j is an eigenvector of A . Let λ_j be the eigenvalue corresponding to u_j , that is,

$$Au_j = \lambda_j u_j.$$

Then

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

where P is the orthogonal matrix $P = [u_1 \ \dots \ u_n]$ and D is the diagonal matrix with diagonal entries $\lambda_1, \dots, \lambda_n$. The equation $A = PDP^T$ can be rewritten as:

$$A = [u_1 \ \dots \ u_n] \begin{bmatrix} \lambda_1 & & \\ & \dots & \\ & & \lambda_n \end{bmatrix} \begin{bmatrix} u_1^T \\ \vdots \\ u_n^T \end{bmatrix}$$

$$= [\lambda_1 u_1 \ \dots \ \lambda_n u_n] \begin{bmatrix} u_1^T \\ \vdots \\ u_n^T \end{bmatrix}$$

$$= \lambda_1 u_1 u_1^T + \dots + \lambda_n u_n u_n^T.$$

$n \times n \cdot (n \times n)$
 $n \times n$

$$\|u_i\|^2 = u_i \cdot u_i = u_i^T u_i \text{ scalar - inner product}$$

The expression

$$A = \lambda_1 u_1 u_1^T + \dots + \lambda_n u_n u_n^T$$

is called the spectral decomposition of A . Note that each matrix $u_j u_j^T$ has rank 1 and is the matrix of projection onto the one dimensional subspace spanned by u_j . In other words, the linear map P defined by $P(x) = u_j u_j^T x$ is the orthogonal projection onto the subspace spanned by u_j .

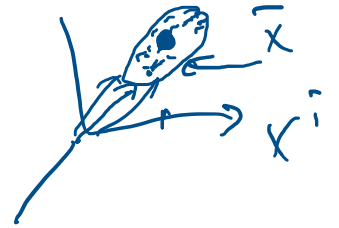
$A = B^T B$ is symmetric
 $\Rightarrow$ spectral decomp. theorem
i.e. also covariance matrices.

A is pos. definite if
 $\lambda_i > 0$ all i .

PCA as algorithm

• Data $\underline{X} = \begin{pmatrix} x_1 & x_2 & \dots & x_n \\ \vdots & \vdots & & \vdots \end{pmatrix}^{ith}$
 $n \times 1$

(say $n = 4000$)
 $n = 216$



• what if $x_i \sim \mathcal{N}(\bar{x}, \Sigma)$ —

covariance matrix.

$$B = \underline{X} - \bar{B}; \quad \bar{B} = \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}_{n \times 1} \bar{x}^T = \underline{1} \bar{x}^T; \quad \bar{x} = \begin{pmatrix} \bar{x}_1 \\ \bar{x}_2 \\ \vdots \\ \bar{x}_n \end{pmatrix}$$

$$\bar{x}_i = \frac{1}{n} \sum_{i=1}^n x_i$$



$$B = U \Sigma V^T; \quad U = [u_1, u_2, \dots, u_n]$$

and u_1 is major axis — most energetic — feature that most explains data.
 u_2 is first minor axis

Side Note:

$\frac{1}{i^2}$ is slowest converging to zero i^2

i.e. $\frac{1}{i}$ is slowest converging i

$$b_i \leq \frac{1}{i}$$

$$\sum_{i=1}^{\infty} \frac{1}{i^p} < \infty \text{ if } p > 1$$

$p < 1$, $p = 1$ harmonic



b_i^2 the power spectrum

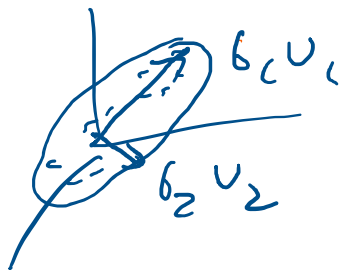
$$= 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots \quad p = 1$$

$$B = \frac{1}{n} \bar{X}^T$$

let $C = \frac{1}{n-1} B^T B$ covariance matrix
is covariance matrix

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

$$U = [u_1, u_2, \dots, u_n]$$



$$\begin{aligned} b_3 &\approx 0 \\ b_2 &\geq 0 \end{aligned}$$

$$u_1 = \underset{\|u\|=1}{\operatorname{argmax}} \frac{u^T B^T B u}{\|B u\|_2^2} = \underset{u}{\operatorname{argmax}} \frac{u^T B^T B u}{\underbrace{u^T u}_{=1}}$$

Rayleigh-Ritz gradient

$$B u \cdot B u = \|B u\|_2^2$$

$$u_2 = \underset{\substack{\|u\|=1 \\ u \perp u_1}}{\operatorname{argmax}} u^T B^T B u$$

• $CV = VD$ eigenvectors of C all stacked.

Ax optimize $x^T Ax$, maximize.
 $(A = B^T B)$
 $x = \begin{pmatrix} x^1 \\ \vdots \\ x^j \\ \vdots \\ x^n \end{pmatrix}$ $r(x) = \frac{x^T Ax}{x^T x}$ ~~r~~
 $\frac{\partial r}{\partial x^j} = \frac{\partial}{\partial x^j} \left(r(x) \right) = \frac{\frac{\partial}{\partial x^j} (x^T Ax)}{(x^T x)^2} - x^T Ax \frac{\partial}{\partial x^j} (\underline{x^T x})$
 $= \frac{2(Ax)_j}{x^T x} - \frac{(x^T Ax) 2x_j}{(x^T x)^2} = \frac{2}{x^T x} (Ax - r(x)x)$
 $\nabla r(x) = \frac{2}{x^T x} (Ax - r(x)x) = \frac{2}{x^T x} (A - r(x)I)x = 0$

$$Ax = \underbrace{r(x)}_{\lambda} x \quad \Rightarrow \quad Ax = \lambda x$$

Conclude

The x that optimizes $r(x) = \frac{x^T A x}{x^T x}$ is an eigenvector and $r(x)$ is its eigenvalue.

• S.D.T. for $A = B^T B = \sum_{i=1}^r \lambda_i \underbrace{U_i U_i^T}_{\substack{\text{cre} \\ \text{the} \\ v_i^T}}$

let $B = U \Sigma V^T$ s.v.d.

$$\underline{B^T B} = V \Sigma^T \cancel{U^T U} \Sigma V^T = V \underbrace{\Sigma^T \Sigma}_{\text{}} V^T$$

$$D = \left(\begin{array}{ccc|c} b_1 & \dots & b_r & \\ \hline & & & 0 \end{array} \right) \quad \left(\begin{array}{ccc|c} b_1 & \dots & b_r & 0 \\ \hline & & & 0 \end{array} \right) = \left(\begin{array}{ccc|c} b_1^2 & & & \\ b_2^2 & \dots & & \\ \hline & & & 0 \end{array} \right)$$

$$= \underline{\underline{V D V^T}} = \sum_{i=1}^r \underline{b_i^2} \underline{v_i} \underline{v_i^T}$$