

Lecture 24: Principal component analysis (PCA)

Tuesday, November 17, 2015 9:30 AM

Singular-value decomposition vs. Spectral decomposition

What is it?

$$A = U S V^{\dagger}$$

U : unitary
 S : diagonal
 V : unitary
 σ_i : sing. values ≥ 0
 $\vec{u}_i$: left singular vectors
 $\vec{v}_i$: right singular vectors
 $\{\vec{u}_i\}, \{\vec{v}_i\}$ orthonormal

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

D : diagonal

if A is normal:

$$A = T D T^{\dagger}$$

T : same unitary!
 λ_i : e-values
 $\vec{t}_i$: e-vectors (orthonormal)

When?

all matrices!

square, diagonalizable matrices

$\mathbb{R}$ or $\mathbb{C}$?

real if A is real

complex even if A is real
 eg. $\begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \rightarrow e^{\pm i\theta}$
 real if A is real & symmetric

Other properties

$$\|A\| = \max \text{ sing. value}(A)$$

$\|A\| \neq \max |\text{eigenvalue}|$ in general
 eg. $\begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$ has e-values 0, 0

$A \rightarrow A + I$ does not shift sing. values
 in general, eg. $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$
 svs 1, 1 svs 2, 0
 evs 1, -1 evs 2, 0

$A \rightarrow A + \alpha I$ shifts e-values by α

They can't be entirely different. What's the connection?

Recall: For any $m \times n$ matrix A ,
 $A^{\dagger} A \geq 0$.

In particular, $A^{\dagger} A$ is Hermitian (and thus normal),
 and its eigenvalues are all real and ≥ 0 .

Proposition: Let A be an $m \times n$ real or complex matrix.

Then,

- The nonzero singular values of A are the square-roots of the eigenvalues of $A^T A$ (same as the square-roots of the e-values of AA^T).
- Eigenvectors of $A^T A$ are right singular vectors of A .
- Eigenvectors of AA^T are left singular vectors of A .

Corollary: AA^T and $A^T A$ have the same nonzero eigenvalues.

Corollary: How to compute the SVD of a matrix?

Answer: Prove this proposition!

Proof:

Let $A = \sum_i \sigma_i \vec{u}_i \vec{v}_i^T$ be the SVD of A .
(We don't know what it is, maybe, but it exists!)

$$\begin{aligned} \Rightarrow A^T A &= \left(\sum_i \sigma_i \vec{v}_i \vec{u}_i^T \right) \left(\sum_j \sigma_j \vec{u}_j \vec{v}_j^T \right) \\ &= \sum_i \sigma_i^2 \vec{v}_i \vec{v}_i^T \quad \text{since } \vec{u}_i^T \vec{u}_j = \begin{cases} 1 & \text{if } i=j \\ 0 & \text{if } i \neq j \end{cases} \end{aligned}$$

$\Rightarrow A^T A \vec{v}_j = \sigma_j^2 \vec{v}_j$
ie., the right singular vector $\vec{v}_j$ is an e-vector of $A^T A$,
with e-value σ_j^2 . ✓

$\Rightarrow$ By diagonalizing $A^T A$, we now know the singular values σ_i and right singular vectors $\vec{v}_i$.

To get $\vec{u}_i$, note $\vec{u}_i = \frac{1}{\sigma_i} A \vec{v}_i$. ✓

□

Corollary: For any A ,

$$\begin{aligned} \|A\| &= \sqrt{\max \text{ e-value of } A^T A} \\ &= \sqrt{\max \text{ e-value of } AA^T}. \end{aligned}$$

For normal matrices, the relationship between singular values and eigenvalues is even simpler.

A normal

$$\Rightarrow A = U D U^\dagger$$

$\uparrow$ unitary $\uparrow$ diag. w/ e-values
 $= \sum_i \lambda_i \vec{u}_i \vec{u}_i^\dagger$
 $\uparrow$ e-vectors (columns of U)

$$= \sum_i |\lambda_i| \cdot \left(\frac{\lambda_i}{|\lambda_i|} \vec{u}_i \right) \cdot \vec{u}_i^\dagger$$

This is an SVD for A !
 $\uparrow$ sing. values $\uparrow$ left/right sing. vectors

Claim: If A is a normal matrix, then the singular values of A are the absolute values of the eigenvalues of A . ✓

Corollary: For any normal matrix A ,
 $\|A\| = \max |\text{eigenvalue of } A|$.

(Again, this is generally false for non-normal matrices.)

Example: If A is orthogonal or unitary, all its singular values are 1.

(Of course, we already knew this; an SVD for A is)
 $A = \sum_i (A e_i) e_i^\dagger$, since the set $\{A e_i\}$ is orthonormal.

THE SINGULAR-VALUE DECOMPOSITION AND DISTANCE TO LOWER-RANK MATRICES

Example:

$$A = \begin{pmatrix} 1 & & & \\ & 1/2 & & \\ & & 1/3 & \\ & & & 1/4 \end{pmatrix} \quad \text{rank}(A) = 4$$

Intuitively:
 What is the rank-one matrix closest to A ?

$$B = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad \text{rank}(B) = 1$$

$$\|B - A\| = \left\| \begin{pmatrix} 0 & 1/2 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1/3 & 0 \\ 0 & 0 & 0 & 1/4 \end{pmatrix} \right\| = \frac{1}{2}$$

What is the rank-two

$$C = \begin{pmatrix} 1 & & & \\ & 1/2 & & \\ & & 0 & \end{pmatrix} \quad \text{rank}(C) = 2$$

What is the rank-2 matrix closest to A ?

$$C = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{2} & 0 \\ 0 & 0 & 0 \end{pmatrix} \quad \text{rank}(C) = 2$$

$$\|C - A\| = \frac{1}{3}$$

What is the rank-3 matrix closest to A ?

$$D = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{2} & 0 \\ 0 & \frac{1}{3} & 0 \end{pmatrix}, \quad \|D - A\| = \frac{1}{4}$$

Theorem: Let A have SVD

$$A = \sum_{i=1}^r \lambda_i \vec{u}_i \vec{v}_i^T,$$

with sorted singular values

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_r > 0.$$

Then for any $k \leq r$, the closest rank k matrix to A is

$$B = \sum_{i=1}^k \lambda_i \vec{u}_i \vec{v}_i^T.$$

Observe: B truncates the smaller singular values of A .

$$\|A - B\| = \lambda_{k+1}$$

Example: $\begin{pmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{1000} & 0 \\ 0 & 0 & \frac{1}{1000} \end{pmatrix}$

has full rank, but is $1/1000$ -close to a rank-one (singular) matrix.

Proof: Let C have rank k .

Goal: Find $\vec{x}$ with $\|\vec{x}\| = 1$,

$$\|(C - A)\vec{x}\| \geq \lambda_{k+1}$$

$$\Rightarrow \|C - A\| \geq \lambda_{k+1}.$$

Intuition: All vectors x in $\text{Span}\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_{k+1}\}$ are stretched by at least λ_{k+1} , and there's no way C cancel out all these $k+1$ dimensions!

find $x \in N(C) \cap \text{Span}\{\vec{v}_1, \dots, \vec{v}_{k+1}\}$ ← new goal!

If these are all $m \times n$ matrices, then $\dim N(C) = n - k$.

Since $\dim \text{Span}\{\vec{v}_1, \dots, \vec{v}_{k+1}\} = k+1$, they must intersect.

For this x , $\|(C - A)x\| = \|Ax\| \geq \lambda_{k+1} \|x\|$. ✓

□

Remark:

Dimension reduction is a major theme in computational

linear algebra and most approaches for data analysis.

- Linear regression, as we have presented it, effectively reduces the dimension of the data by 1, e.g., fitting a 1D line to a 2D cloud of data points.

The Johnson-Lindenstrauss Lemma projects n data points to $\log n$ (random) dimensions, approximately preserving angles & lengths (after rescaling).

- Matrix rank reduction via the SVD is another example. Moreover, it is used in **Principal Component Analysis** for reducing the dimensionality of data sets. We'll see this next.

PRINCIPAL COMPONENT ANALYSIS (PCA)

Example: Campaign contributions to US senators

1. Get the data

MapLight REVEALING MONEY'S INFLUENCE ON POLITICS

U.S. Congress ▶ Change

Guide Bills Legislators Interest Groups Contributions Companies

Find Contributions

Monetary and non-monetary contributions to candidate campaign committees of legislators serving in the 109th, 110th, 111th

Office/Party	Office or party affiliation	☒ Senate ☐ House ☐ Democrat ☐ Republican ☐ Independent
Legislator	Recipient (elected official)	All
Contributor	Name, occupation, or employer *	
Interest Group	Category of contributor	advanced any any any
Date	Date contribution was made	to
Election Cycle	Election in following year	☐ 2002 ☐ 2004 ☐ 2006 ☒ 2008 ☒ 2010 ☒ 2012
Location	Location of individual contributors	, State ,
Source	Individual or organization	☒ All ☐ Non-PAC ☐ PAC

Find

2. Clean it up

Contributor categories

Senator	Agribusiness	Communic/Electronics	Construction	Defense	Education
{Alexander, Lamar, R}	223 171	208 781	287 750	75 400	145 700
{Ayotte, Kelly, R}	98 000	273 969	124 938	67 150	30 600
{Baldwin, Tammy, D}	85 509	440 506	74 609	4450	287 934
{Barrasso, John, R}	166 150	262 850	186 933	86 500	37 050
{Baucus, Max, D}	627 036	517 595	302 437	101 600	50 538

Senators

Senator	Agribusiness	Communic/Electronics	Construction	Defense	Education
{Alexander, Lamar, R}	223 171	208 781	287 750	75 400	145 700
{Ayotte, Kelly, R}	98 000	273 969	124 938	67 150	30 600
{Baldwin, Tammy, D}	85 509	440 506	74 609	4450	287 934
{Barrasso, John, R}	166 150	262 850	186 933	86 500	37 050
{Baucus, Max, D}	627 036	517 595	302 437	101 600	50 538
{Begich, Mark, D}	16 781	310 583	95 819	30 000	40 175
{Bennet, Michael, D}	214 341	826 273	161 425	53 150	274 355
{Blumenthal, Richard, D}	30 050	309 600	76 400	21 750	155 400
{Blunt, Roy, R}	805 677	728 737	481 335	155 350	70 150
{Boozman, John, R}	172 829	70 296	228 786	3000	26 700
{Boxer, Barbara, D}	207 350	1 634 559	388 955	24 150	338 404
{Brown, Sherrod, D}	180 532	508 863	239 702	102 450	461 577
{Burr, Richard, R}	603 168	264 838	277 904	156 947	55 365
{Cantwell, Maria, D}	90 883	490 380	89 673	40 975	94 278
{Cardin, Benjamin, D}	50 750	350 837	157 822	80 550	121 665
{Carper, Thomas, D}	91 885	276 990	97 810	60 400	25 466
{Casey, Robert, D}	251 250	537 416	244 976	125 075	257 211
{Chambliss, Saxby, R}	1 758 731	380 684	483 351	286 150	44 368
{Coburn, Thomas, R}	100 625	115 630	57 800	36 650	5400
{Cochran, Thad, R}	373 294	98 500	119 250	167 050	21 350
{Collins, Susan, R}	207 838	459 982	269 546	291 801	56 350
{Coons, Chris, D}	32 550	200 450	34 150	19 500	55 077
{Corker, Bob, R}	316 050	436 582	600 050	70 850	81 250
{Cornyn, John, R}	686 106	486 367	559 927	243 076	52 725
{Crapo, Michael, R}	288 499	145 962	132 420	39 000	3930
{Cruz, Ted, R}	247 200	234 178	340 903	38 700	69 500
{Donnelly, Joe, D}	194 546	144 158	181 548	120 283	140 709
{Durbin, Richard, D}	249 014	486 625	191 600	127 350	149 375
{Enzi, Michael, R}	90 444	73 050	50 500	25 500	13 000
{Feinstein, Dianne, D}	522 217	591 836	184 400	198 500	89 386
{Fischer, Deb, R}	291 502	97 400	131 028	10 000	12 500
{Flake, Jeff, R}	399 577	212 700	317 850	26 900	79 000

3. "Standardize" the data

a) Shift the numbers in each column so the mean $\sum_{i=1}^{97} x_j^{(i)} = 0$.

b) Scale each column so the variance $\sigma_j^2 = \sum_{i=1}^{97} (x_j^{(i)})^2 = 1$.

Note: The scaling may be omitted if the different attributes are known to lie on the same scale.

Here it has the effect of making all categories equally significant, even though some give more than others!

```
{sectors, Plus @@ puredata} // Transpose // Sort[#, #1[[2]] > #2[[2]] &] & // MatrixForm
```

```
(Finance/Insur/RealEst 168151725
   Unknown 154842094
   Lawyers & Lobbyists 112167979
   Other 103253137
Ideology/Single-issue 100385417
   Misc Business 92373646
   Health 71844560
Communic/Electronics 45991954
Energy/Nat Resource 40100000
   Agribusiness 30637498
   Construction 28244716
Transportation 23257205
   Labor 20546123
   Education 12200463
   Defense 10484140
   Party Cmte 1690362
Joint Candidate Cmtes 514864)
```

```
{height, width} = Dimensions[puredata]
puredatastandardized = puredata // N;
```

$$\text{means} = \frac{1}{\text{height}} \sum_{j=1}^{\text{height}} \text{puredatastandardized}[[j]];$$

```
For[j = 1, j ≤ height, j++,
  puredatastandardized[[j]] -= means;
];
```

$$\text{variances} = \frac{1}{\text{height}} \sum_{j=1}^{\text{height}} \text{puredatastandardized}[[j]]^2;$$

```
For[j = 1, j ≤ height, j++,
  puredatastandardized[[j]] /= Sqrt[variances];
];
{97, 17}
```

4. Take the SVD:

```
{U, M, V} = SingularValueDecomposition[puredatastandardized];
Plus @@ Plus @@ Abs /@ (U.M.Transpose[V] - puredatastandardized)
{Dimensions[U], Dimensions[M], Dimensions[V]}
```

5.96535×10^{-12} $\rightarrow U M V^T = \text{data matrix} \checkmark$
 $\text{dim } U \quad \text{dim } M \quad \text{dim } V$
 {{97, 97}, {97, 17}, {17, 17}}

The singular values, sorted from largest on down:

```
M // MatrixForm
```

```
(33.265 0. 0. 0. 0. 0. 0. 0.
 0. 12.4751 0. 0. 0. 0. 0. 0. ....)
```

M // MatrixForm

33.265	0.	0.	0.	0.	0.	0.	0.
0.	12.4751	0.	0.	0.	0.	0.	0.
0.	0.	10.0409	0.	0.	0.	0.	0.
0.	0.	0.	9.3013	0.	0.	0.	0.
0.	0.	0.	0.	7.3885	0.	0.	0.
0.	0.	0.	0.	0.	5.60893	0.	0.
0.	0.	0.	0.	0.	0.	4.79764	0.
0.	0.	0.	0.	0.	0.	0.	4.4487
0.	0.	0.	0.	0.	0.	0.	0.
0.	0.	0.	0.	0.	0.	0.	0.
0.	0.	0.	0.	0.	0.	0.	0.
0.	0.	0.	0.	0.	0.	0.	0.
0.	0.	0.	0.	0.	0.	0.	0.

Note: The columns of U are the left singular vectors, forming a basis for columnspace $R(\text{data})$.
The columns of V are the right singular vectors, forming a basis for row space $R(\text{data}^T)$.

5. Now what?

a) We could approximate the data matrix with a lower-rank matrix, eg., keeping

$$\begin{array}{c}
 \begin{array}{c} \text{first two} \\ \text{columns of } U \end{array} \\
 \left(\begin{array}{cc} 33.265 & 0 \\ 0 & 12.4751 \end{array} \right)^2 \begin{array}{c} \text{first two columns} \\ \text{of } V, \text{ transposed} \end{array}
 \end{array}$$

gives a rank-2 matrix.

Each ² data point (row) now lies in a 2D subspace of $\mathbb{R}^{17}$.

```

"These are the two most important directions:";
principaldirections = V[All, {1, 2}];

principaldirectionslabeled = Append[principaldirections // Transpose, sectors] // Transpose;
"Now let's sort these by the second component (we'll see why in a moment).";
Sort[principaldirectionslabeled,
  #1[[2]] < #2[[2]] &
] // MatrixForm

```

-0.00387268	-0.707777	Labor
-0.238924	-0.329765	Education
-0.261419	-0.268357	Lawyers & Lobbyists
-0.260683	-0.216817	Ideology/Single-issue
-0.266349	-0.211043	Communic/Electronics
-0.282418	-0.00418024	Other
-0.288545	0.00920218	Unknown
-0.286914	0.0197931	Finance/Insur/RealEst
-0.291324	0.0353523	Misc Business
-0.273799	0.0478608	Health
-0.2148	0.05117	Defense
-0.0122403	0.0995367	Joint Candidate Cmtes
-0.285753	0.110642	Construction
-0.273507	0.174331	Transportation
-0.0855998	0.186206	Party Cmte
-0.241713	0.244521	Agribusiness
-0.248974	0.270343	Energy/Nat Resource

b) Or, if you care only about the data points' positions in the 2D subspace — and not about the subspace's position in $\mathbb{R}^{17}$ — then we can use the first two columns of V to project down to $\mathbb{R}^2$:

```

puredatastandardized // Dimensions
principaldirections // Dimensions
dataproyectedmanually = puredatastandardized.principaldirections;
{97, 17}
{17, 2}

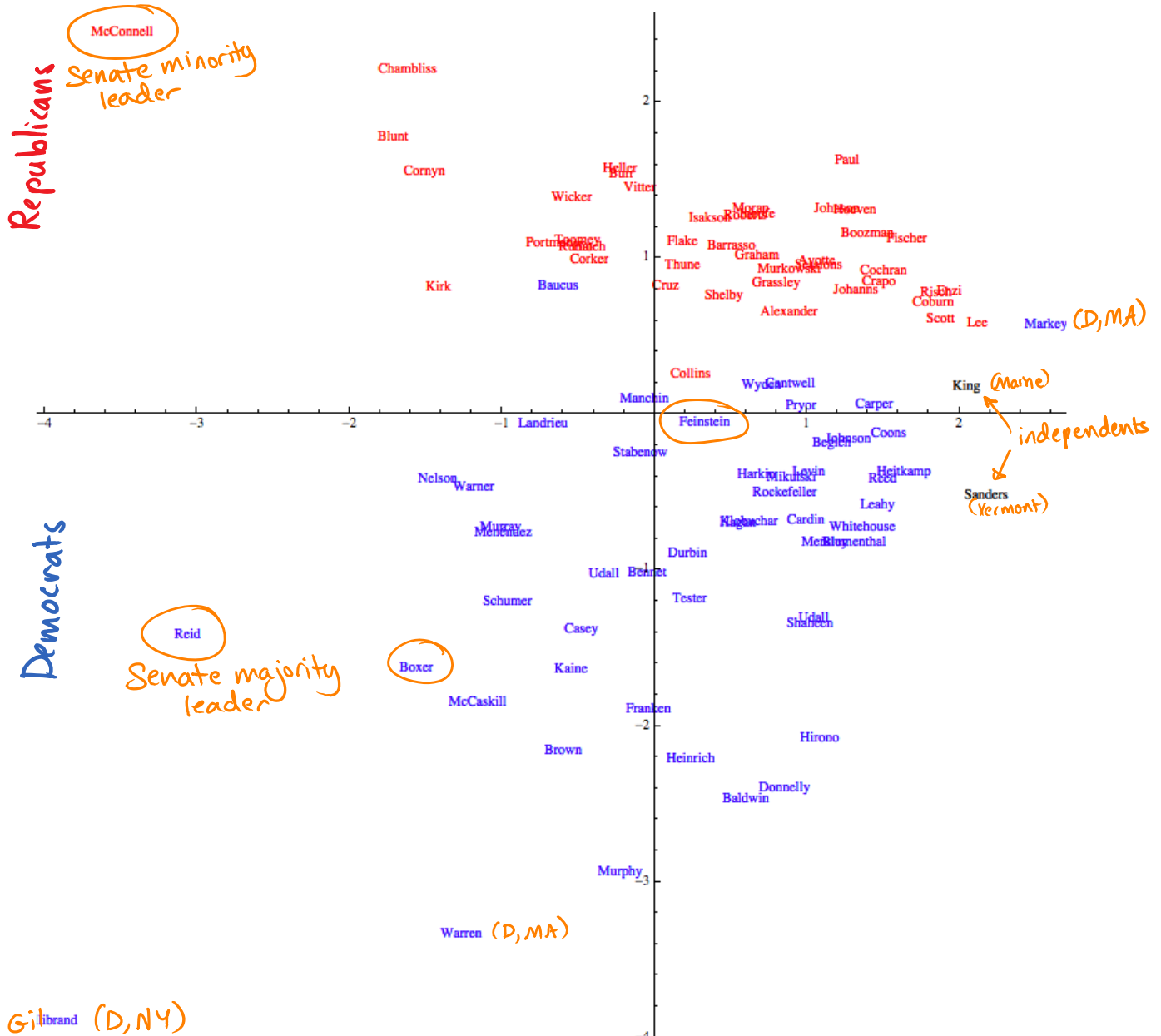
```

$$\begin{aligned}
 \text{Since } \text{data} &= U M V^T, \quad \text{data} * \begin{pmatrix} \text{first two} \\ \text{columns of } V \end{pmatrix} = \sum_{i=1}^2 \sigma_i \vec{u}_i \vec{e}_i^T \\
 &= \sum_{i=1}^{17} \sigma_i \vec{u}_i \vec{v}_i^T \\
 &= \vec{v}_1 \vec{e}_1^T + \vec{v}_2 \vec{e}_2^T \\
 &= \begin{pmatrix} | & | \\ \sigma_1 \vec{u}_1 & \sigma_2 \vec{u}_2 \\ | & | \end{pmatrix}
 \end{aligned}$$

Of course, we could also just have picked out the first two left singular vectors directly.

Let's plot it!

```
ListPlot[{{#} & /@ dataprojectedmanually, AspectRatio -> 1, PlotMarkers -> data[[All, 1, 1]],
PlotStyle -> (data[[All, 1, 3]]) /. {"D" -> Blue, "R" -> Red, "I" -> Black}]
{97, 17}
{17, 2}
```



Interpretation: 2nd principal component → political party
(maybe liberal vs. conservative?)
1st principal component → ???

Since all entries of the first principal direction are < 0, the x-direction above says who raised more money (with the most to the left). The y-direction (second principal direction) says how that money is

distributed away from average. Negative means more Labor. Positive means more Energy/Nat. Resources.

Notice that there is a slight funnel effect; candidates who raise more money (on the left) also get it from a more biased distribution of donors.

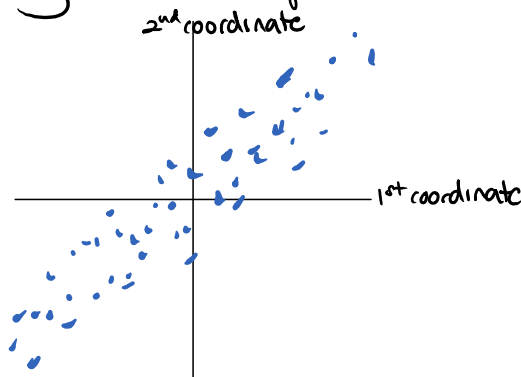
Find correlations, clusters, voids, study more principal components, etc.

Note: Matlab/Mathematica/R all have built-in PCA functions:

```
principalcomponents = PrincipalComponents[puredatastandardized];
```

```
"To keep just the first two principal components (first two columns), use:"  
principalcomponents[[All, {1, 2}]]
```

Question: Why can't the plot look like this?



Answer: We are plotting $(x_1, y_1), \dots, (x_m, y_m)$, where $(x_1, \dots, x_m)$ is the first left singular vector and $(y_1, \dots, y_m)$ is the second left singular vector. These vectors are orthogonal:

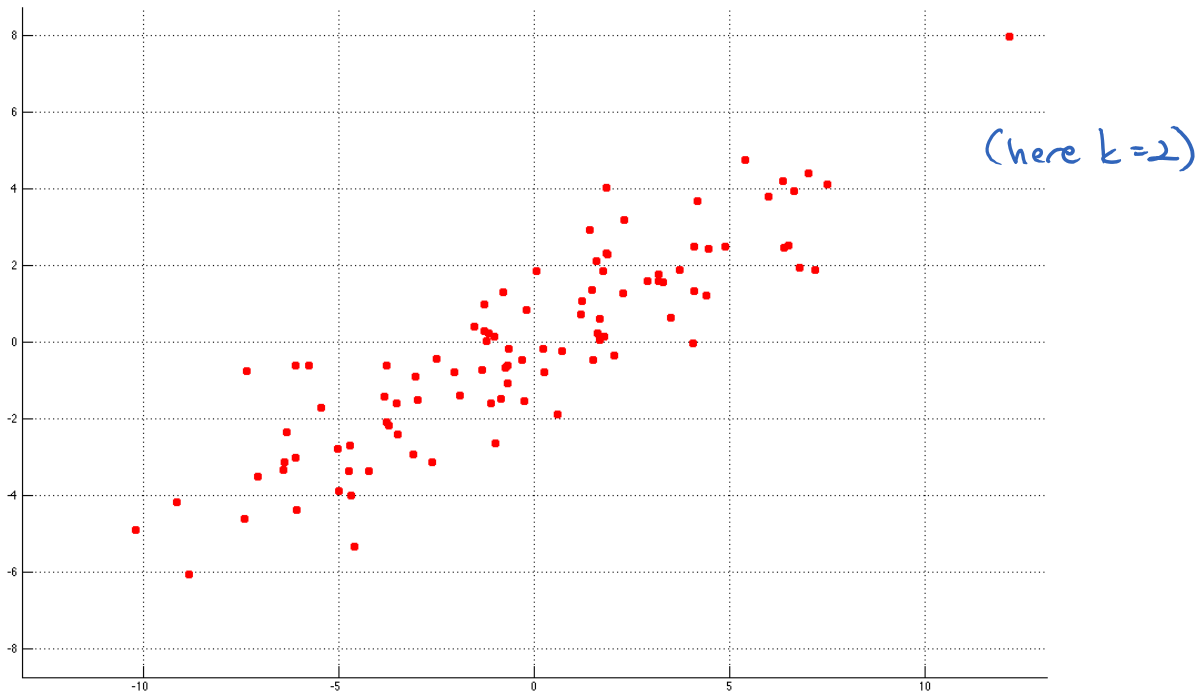
$$\sum_{j=1}^m x_j y_j = 0.$$

Informal motivation for PCA:

"The first k principal directions are the k factors that explain the data best."

Formal mathematical motivation:

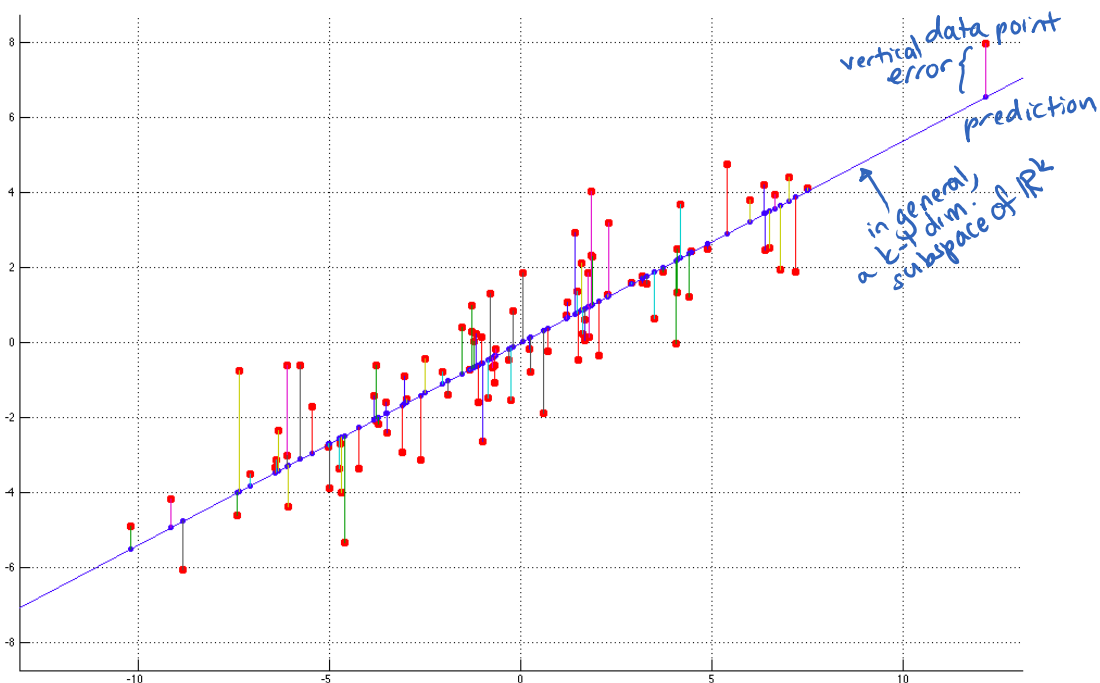
Given data points $\{(x_1^{(1)}, \dots, x_k^{(1)}), (x_1^{(2)}, \dots, x_k^{(2)}), \dots, (x_1^{(m)}, \dots, x_k^{(m)})\}$,



The "method of least squares" lets us predict the last component/attribute in terms of the first $k-1$ components, minimizing the squared error

$$\left\| \begin{pmatrix} 1 & x_1^{(1)} & \dots & x_{k-1}^{(1)} \\ \vdots & \vdots & & \vdots \\ 1 & x_1^{(m)} & \dots & x_{k-1}^{(m)} \end{pmatrix} \begin{pmatrix} \alpha_0 \\ \alpha_1 \\ \vdots \\ \alpha_{k-1} \end{pmatrix} - \begin{pmatrix} x_k^{(1)} \\ \vdots \\ x_k^{(m)} \end{pmatrix} \right\|^2$$

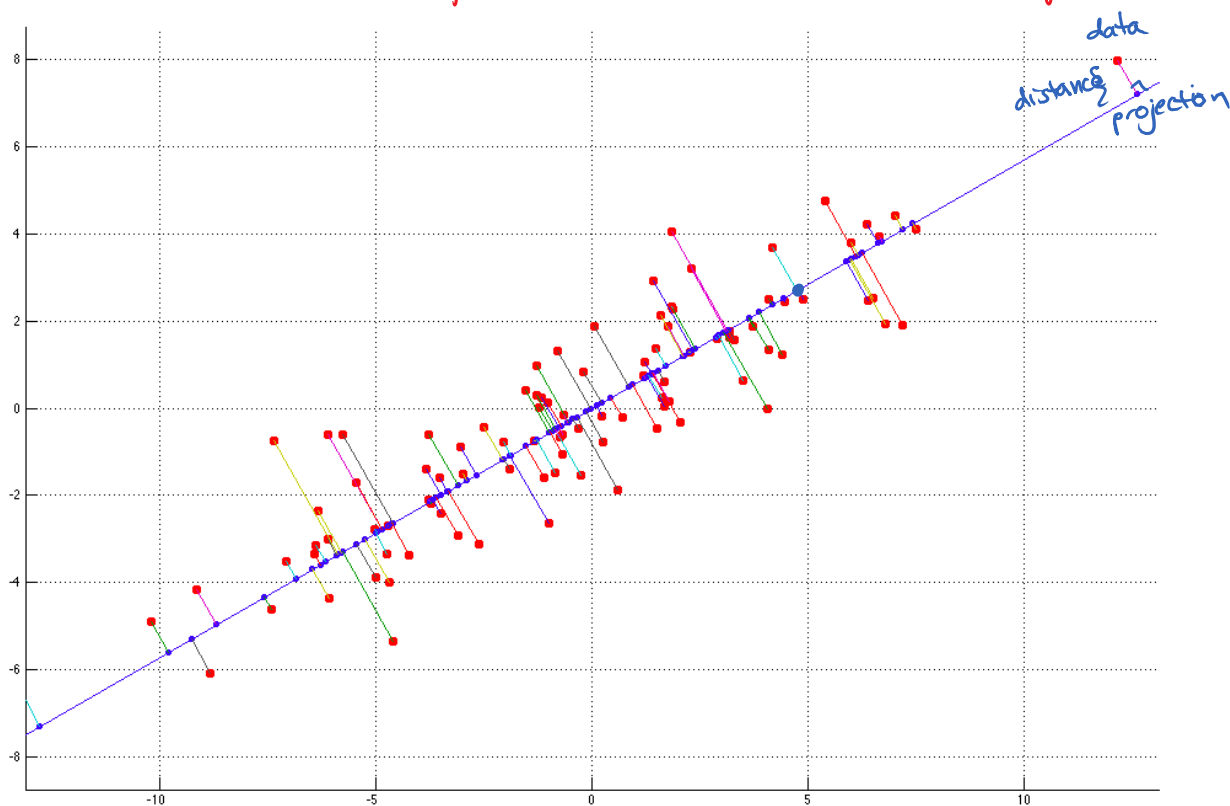
$$= \sum_{j=1}^m \left(\alpha_0 + \alpha_1 x_1^{(j)} + \dots + \alpha_{k-1} x_{k-1}^{(j)} - x_k^{(j)} \right)^2$$



This is most useful if the first $k-1$ attributes are exact, and the last one imprecise or noisy.

Least-squares regression is not useful for analyzing the campaign finance data, because all the components are imprecise/noisy (equally noisy after scaling the columns), and we aren't just trying to predict one of them. Basically, the data are less structured.

In contrast, principal component analysis (PCA) finds the subspace (with a dimension you can choose) that **minimizes the total squared distances from the subspace**.



Observe: For any subspace S , and any data point $\vec{x}$,

$$\|\vec{x}\|^2 = \|\text{Proj}_S(\vec{x})\|^2 + \|\vec{x} - \text{Proj}_S(\vec{x})\|^2$$

$$\Rightarrow \sum_{i=1}^m \|\vec{x}^{(i)}\|^2 = \sum_j \|\text{Proj}_S(\vec{x}^{(j)})\|^2 + \sum_j \|\vec{x} - \text{Proj}_S(\vec{x})\|^2$$

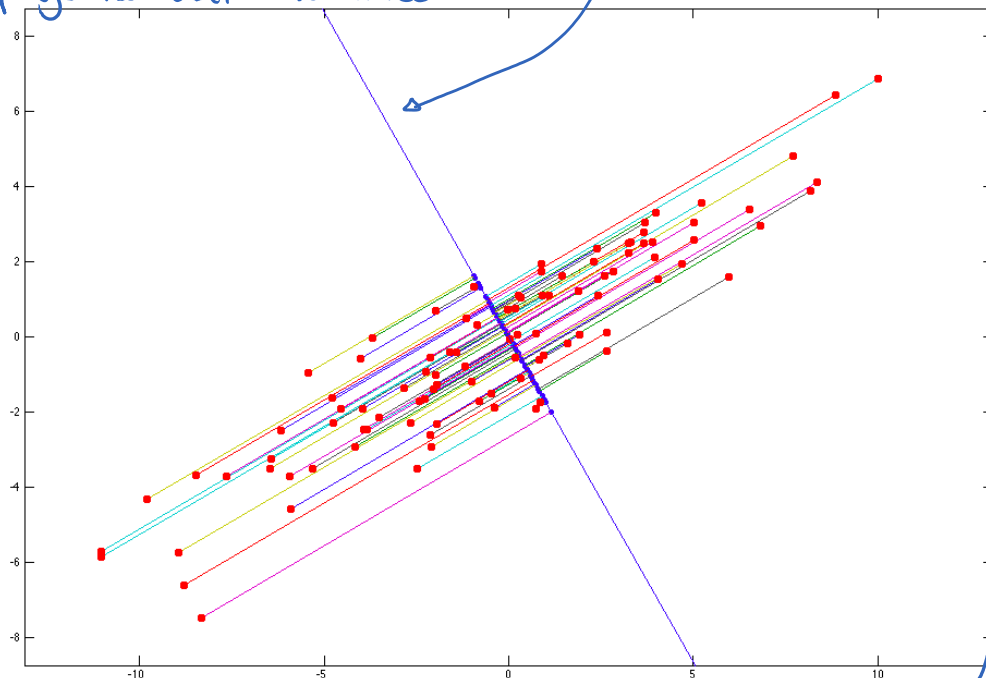
$\therefore$ minimizing sum of squared distances to S

$\iff$ maximizing $\sum_j \|\text{Proj}_S(\vec{x}^{(j)})\|^2$
ie., maximizing the projected data's variance

squared distances to S

ie., maximizing the projected data's variance

In contrast, projecting down to this line minimizes the projected data's variance:



Let's solve this optimization problem...

Theorem: For data points $\vec{x}^{(1)}, \vec{x}^{(2)}, \dots, \vec{x}^{(m)}$, let

$$A = \begin{pmatrix} \vec{x}^{(1)} \\ \vec{x}^{(2)} \\ \vdots \\ \vec{x}^{(m)} \end{pmatrix} = \sum_{j=1}^m \vec{e}_j \vec{x}^{(j)T}$$

Then the k -dimensional subspace S that maximizes

$$\sum_{j=1}^m \|\mathcal{P}_S \vec{x}^{(j)}\|^2$$

is given by

$S = \text{Span} \{k \text{ largest e-value eigenvectors of } A^T A\}.$

Proof:

$$\max_S \sum_{j=1}^m \|\mathcal{P}_S(\vec{x}^{(j)})\|^2$$

subspace of dim. k

$$= \max_{\text{orthonormal set } \{u_1, \dots, u_k\}} \sum_{j=1}^m \sum_{i=1}^k (u_i \cdot \vec{x}^{(j)})^2$$

$$\sum_{i=1}^k \sum_{j=1}^m u_i^T \vec{x}^{(j)} \vec{x}^{(j)T} u_i = \sum_{i=1}^k u_i^T \left(\sum_j \vec{x}^{(j)} \vec{x}^{(j)T} \right) u_i$$

$$i=1 \quad \delta^{-1}$$

$$= \sum_{i=1}^k \vec{u}_i^T A^T A \vec{u}_i$$

If $k=1$, $\max_{\vec{u}, \|\vec{u}\|=1} \vec{u}^T A^T A \vec{u}$ is achieved by the principal eigenvector of $A^T A$ (principal right singular vector of A). ✓

For $k > 1$, the proof is trickier. It is similar to the previous claim that the best rank- k approximation to $A^T A$ is given by keeping just the top k eigenvalues.

First of all, choose a basis so that $A^T A$ is diagonal:

$$A^T A = \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & \ddots & \\ & & & \lambda_n \end{pmatrix}, \text{ with } \lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq 0.$$

Note that $\sum_{i=1}^k \vec{u}_i^T A^T A \vec{u}_i = \text{Trace}\left(\sum_{i=1}^k \vec{u}_i^T A^T A \vec{u}_i\right)$ since it is a scalar

$$= \text{Trace}\left(\sum_{i=1}^k A^T A \vec{u}_i \vec{u}_i^T\right) \text{ trace is cyclic}$$

$$= \text{Trace}(A^T A \cdot P)$$

for $P = \sum_{i=1}^k \vec{u}_i \vec{u}_i^T$.

If we knew that P had to be diagonal in the same basis as $A^T A$, then obviously we would choose

$$P = \begin{pmatrix} 1 & & & \\ & 1 & & \\ & & \ddots & \\ & & & 1 & & \\ & & & & 0 & \dots \\ & & & & & \ddots \\ & & & & & & 0 \end{pmatrix}$$

to pick out the k largest eigenvalues of $A^T A$. But we don't know that!

Therefore our proof will have 4 steps:

① Relax the optimization problem $\max_{P \text{ } k\text{-dim proj.}} \text{Tr}(A^T A P)$

② Argue that a solution to the relaxed problem is diagonal.

③ Solve it.

④ Show that the solution solves the original problem, too.

The natural relaxation is to consider not just k -dimensional projections but convex combinations of k -dimensional projections.

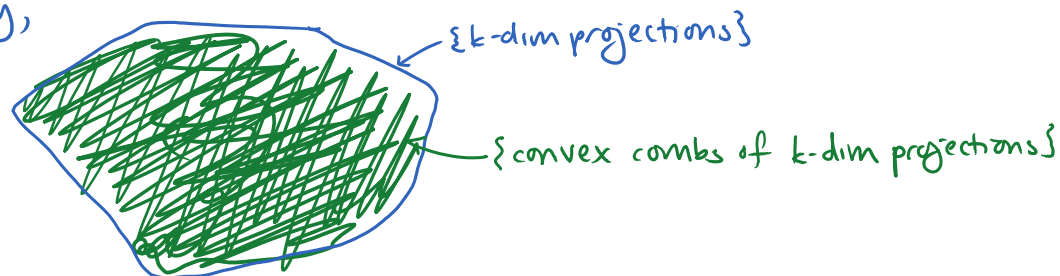
Geometrically,



$\{k\text{-dim projections}\}$

projections are convex combinations of k -dimensional projections.

Geometrically,



A matrix Q is a convex combination of k -dim. projections if and only if

$$0 \preceq Q \preceq I \quad \text{and} \quad \text{Tr} Q = k.$$

$$\Rightarrow \max_{P \text{ k-dim proj.}} \text{Tr}(A^T A P) \leq \max_{\substack{0 \preceq Q \preceq I \\ \text{Tr} Q = k}} \text{Tr}(A^T A Q). \quad \checkmark_{\text{step ①}}$$

Next we argue that an optimal Q is diagonal. Indeed, let

$$U_\ell = \begin{pmatrix} 1 & & & \\ & \ddots & & \\ & & 1 & \\ & & & \ddots \\ & & & & -1 & \\ & & & & & \ddots \\ & & & & & & -1 & \\ & & & & & & & \ddots \end{pmatrix}.$$

Observe that

$$\begin{aligned} \text{Tr}(A^T A U_\ell Q U_\ell) &= \text{Tr}(U_\ell A^T A U_\ell Q) \quad (\text{trace cyclic}) \\ &= \text{Tr}(A^T A Q) \quad (\text{since } A^T A \text{ is diagonal, } U_\ell A^T A U_\ell = A^T A) \end{aligned}$$

$\Rightarrow$ By linearity of the trace,

$$\text{Tr}\left[A^T A \left(\frac{1}{2} Q + \frac{1}{2} U_\ell Q U_\ell\right)\right] = \text{Tr}(A^T A Q).$$

But expanding $Q = \begin{pmatrix} \overset{\ell}{Q_{11}} & \overset{n-\ell}{Q_{12}} \\ Q_{21} & Q_{22} \end{pmatrix}$, $U_\ell Q U_\ell = \begin{pmatrix} Q_{11} & -Q_{12} \\ -Q_{21} & Q_{22} \end{pmatrix}$.

$$\Rightarrow \frac{1}{2} Q + \frac{1}{2} U_\ell Q U_\ell = \begin{pmatrix} Q_{11} & 0 \\ 0 & Q_{22} \end{pmatrix} \text{ is block-diagonal.}$$

Repeating this argument for different ℓ implies that the optimum to $\max_{0 \preceq Q \preceq I, \text{Tr} Q = k} \text{Tr}(A^T A Q)$ is achieved for Q a diagonal matrix (in the same basis that $A^T A$ is diagonal). $\checkmark_{\text{step ②}}$

Step ③ is trivial. We have k units of mass to distribute along the diagonal of $A^T A = \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix}$. We can put at most 1 on any coordinate. Obviously we should put as much as possible in the upper

left corner: $Q = \begin{pmatrix} 1 & & & \\ & \ddots & & \\ & & \underbrace{1 \dots 1}_k & \\ & & & \underbrace{0 \dots 0}_{n-k} \end{pmatrix}$. ✓

Since the solution we found in Step ⑤ is actually a projection matrix, it is a valid solution to the original optimization problem, and must be optimal for that, too. ✓④

We conclude that $\max_{P \text{ k-dim proj.}} \text{Tr}(ATA P) = \lambda_1 + \lambda_2 + \dots + \lambda_k$,

achieved by the projection onto the span of the top k eigenvectors of ATA . $\square$

Remark: An entirely analogous argument would have worked to solve

$$\min_{B: \text{rank}(B)=k} \|A - B\|,$$

since the function $\|A - B\|$ is convex in B :

$$\|A - \frac{1}{2}B - \frac{1}{2}UBU^T\| \leq \frac{1}{2}\|A - B\| + \frac{1}{2}\|A - UBU^T\|.$$

Applications of Principal Component Analysis (PCA):

Compression to fewer dimensions:

- saves space
- reduces noise
- useful for visualization & finding patterns
- allows for faster data processing
- reduces over-fitting in machine learning
(simpler data means you can use simpler hypothesis classes)