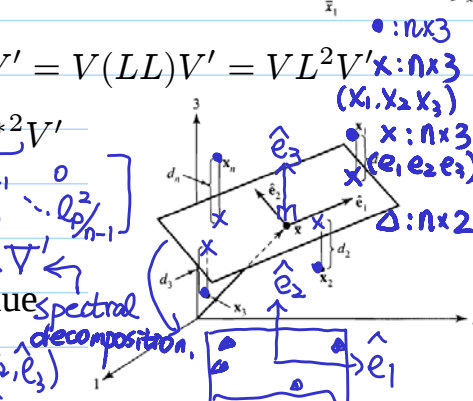
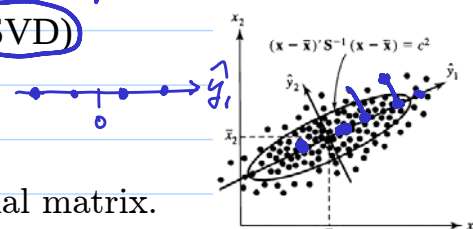
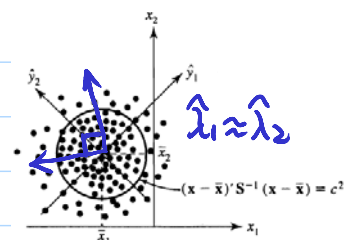


p. 4-11

- equal eigenvalues (i.e., $\lambda_i = \lambda_{i+1} = \dots = \lambda_j$)
 - population principal component
 - sample principal component
- selecting a subset of variables $X_1, \dots, X_p$
 - PCA does not quite reduce the data since PCs are linear combinations of all variables $X_1, \dots, X_p$ so all variables are in some sense still needed. The technique, however, can help us discard some variables by keeping only those with the highest factor loading
- connection to the singular value decomposition (SVD)
 - Let us apply a SVD to the matrix $\mathbf{X} - \mathbf{1}\bar{\mathbf{x}}'$



$$Y_k = \begin{bmatrix} Y_{k1} \\ \vdots \\ Y_{kn} \end{bmatrix}$$

where $U'U = I_p$, $V'V = I_p$, and L is a diagonal matrix.

Then,

$$(n-1)S = (\mathbf{X} - \mathbf{1}\bar{\mathbf{x}}')'(\mathbf{X} - \mathbf{1}\bar{\mathbf{x}}') = V L U' U L V' = V (L L) V' = V L^2 V'$$

$$S = V \left(\frac{1}{n-1} L^2 \right) V' = V \left(\frac{1}{\sqrt{n-1}} L \right)^2 V' \equiv V L^* V'$$

which means

- columns of V : eigenvectors of S
- $\frac{l_i^2}{n-1}$: eigenvalues of S , where l_i : singular value
- $\mathbf{Y} = (\mathbf{X} - \mathbf{1}\bar{\mathbf{x}}')V = UL$: PC scores

spectral decomposition

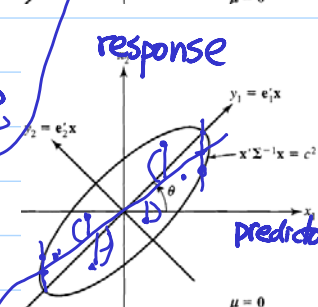
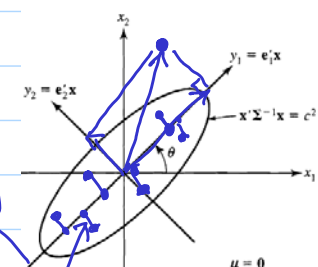
- LNp2-18 result 2A.16
- the matrix $\mathbf{B}$ that minimize $tr[(\mathbf{X} - \mathbf{1}\bar{\mathbf{x}} - \mathbf{B})(\mathbf{X} - \mathbf{1}\bar{\mathbf{x}} - \mathbf{B})'] \in \text{sum of squared difference} = \sum_{i,j} (a_{ij} - b_{ij})^2$ over all $n \times p$ matrices $\mathbf{B}$ having rank no greater than k ($k < p$) is
- $\mathbf{B} = \mathbf{U} \mathbf{L}_k \mathbf{V}'$, minimum value of squared difference $= \hat{\lambda}_{k+1}^2 + \dots + \hat{\lambda}_p^2$
- $\mathbf{U} \mathbf{L}_k$ represents the first k PC scores $[Y_k | 0] \dots x(\hat{e}_1, \hat{e}_2, \hat{e}_3)$

- geometry of PC line (plane)
 - the PC line (plane) minimizes the sum of squared orthogonal distances from each data point to the line (plane)

PC scores

$$\mathbf{x}_j = (\mathbf{x}_j' \hat{\mathbf{e}}_1) \hat{\mathbf{e}}_1 + (\mathbf{x}_j' \hat{\mathbf{e}}_2) \hat{\mathbf{e}}_2 + \dots + (\mathbf{x}_j' \hat{\mathbf{e}}_p) \hat{\mathbf{e}}_p$$

$$= \hat{y}_{j1} \hat{\mathbf{e}}_1 + \hat{y}_{j2} \hat{\mathbf{e}}_2 + \dots + \hat{y}_{jk} \hat{\mathbf{e}}_k + \hat{y}_{j,k+1} \hat{\mathbf{e}}_{k+1} + \dots + \hat{y}_{jp} \hat{\mathbf{e}}_p$$



- (c.f.) the least square line in regression minimizes the sum of vertical distances from the data point to the line

➤ drawbacks of PCA

- PCA only utilizes information contained in the second moments
- nonlinear structure may be missed
- linear combination of variables may not be meaningful especially if the variables do not represent comparable quantities
- outliers may distort the results

Note: PCA does not use the information of

other moments (including 1st moment)

data matrix	dim		cov matrix	sum of variance explained
$X_{n \times p} (\bullet, x_1, x_2, x_3)$	p	$= U \Lambda U'$	$V \Lambda V'$	$\hat{\lambda}_1^2 + \dots + \hat{\lambda}_p^2$
X_k $n \times p$ (x, x_1, x_2, x_3)	k	$= U \Lambda_k U'$	$V \Lambda_k V'$	$\hat{\lambda}_1^2 + \dots + \hat{\lambda}_k^2$
Y $n \times p$ $(\bullet, \hat{e}_1, \hat{e}_2, \hat{e}_3)$	p	$= U \Lambda$	Λ	$\hat{\lambda}_1^2 + \dots + \hat{\lambda}_p^2$
Y_k^* $n \times p$ $(x, \hat{e}_1, \hat{e}_2, \hat{e}_3)$	k	$= U \Lambda_k$	Λ_k	$\hat{\lambda}_1^2 + \dots + \hat{\lambda}_k^2$
Y_k $n \times k$ $(\Delta, \hat{e}_1, \hat{e}_2)$ $\begin{matrix} \uparrow & \uparrow \\ \begin{bmatrix} 1 \\ 0 \end{bmatrix} & \begin{bmatrix} 0 \\ 1 \end{bmatrix} \end{matrix}$	k	$= U \Lambda \begin{bmatrix} I_k \\ 0 \end{bmatrix}$	Λ_k	$\hat{\lambda}_1^2 + \dots + \hat{\lambda}_k^2$

❖ Reading: Textbook, 8.1, 8.2, 8.3, 8.4, 8.5, Supplement 8A