

Linear Regression

C_p , AIC, BIC, and Adjusted R^2

Prof. Asim Tewari
IIT Bombay

$$MSE = \frac{RSS}{n}$$

$$(MSE)_{Training} = \frac{RSS_{Training}}{n}$$

$$(MSE)_{Testing} = \frac{RSS_{Testing}}{n}$$

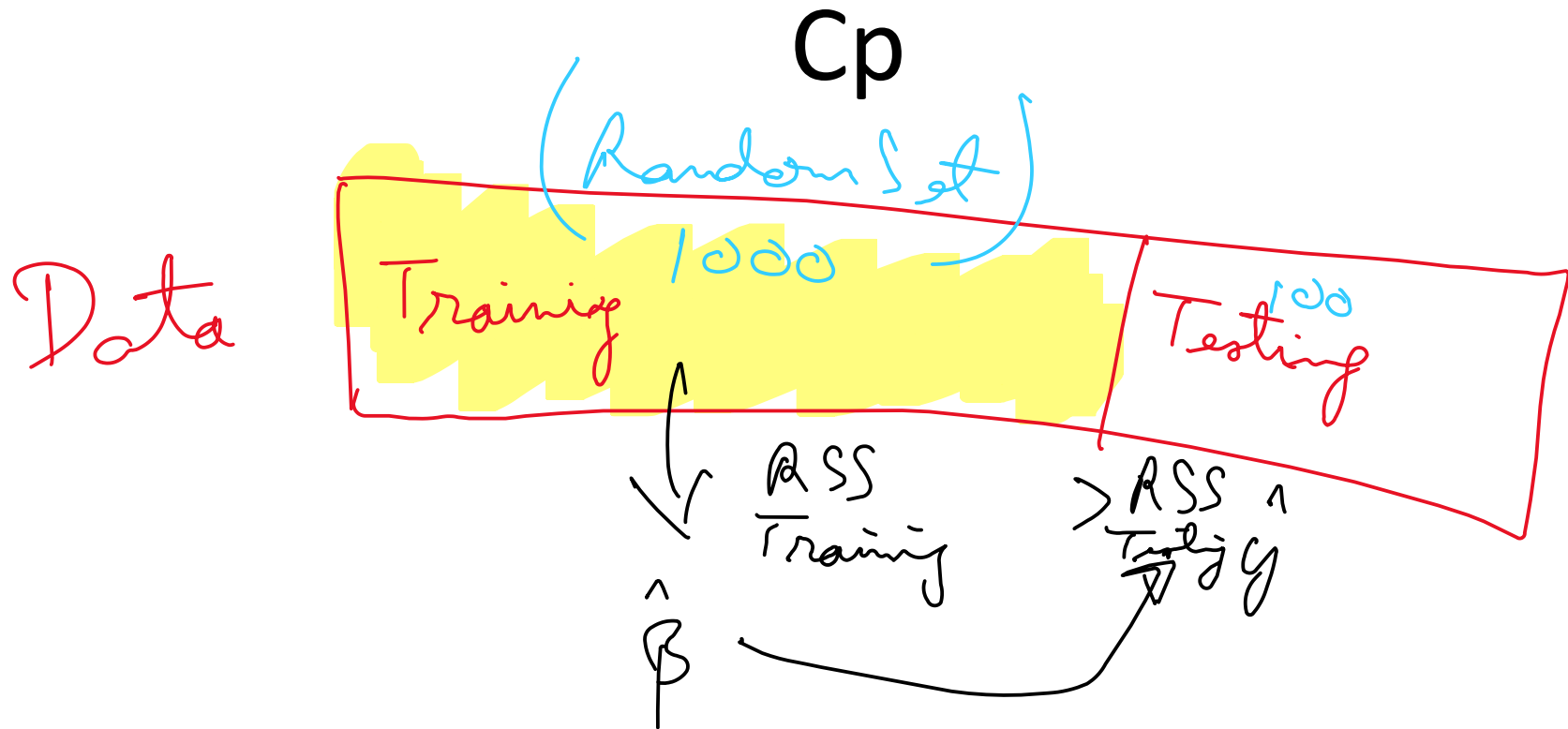
Cp

- For a fitted least squares model containing d predictors, the C_p estimate of test MSE is computed using the equation

$$C_p = \frac{1}{n} (RSS + 2d\hat{\sigma}^2)$$

$\swarrow$ $\nwarrow$ $\nearrow$ $\nwarrow$
 $(MSE)_{Training}$ $+ 2d\hat{\sigma}^2$ $\rightarrow$ Estimator of MSE of Testing

- where $\hat{\sigma}^2$ is an estimate of the variance of the error associated with each response measurement



AIC

- The AIC criterion is defined for a large class of models fit by maximum likelihood. In the case of the model

$$Y = \beta_0 + \beta_1 X_1 + \cdots + \beta_p X_p + \epsilon$$

with Gaussian errors, maximum likelihood and least squares are the same thing.

$$\text{AIC} = \frac{1}{n\hat{\sigma}^2} (\text{RSS} + 2d\hat{\sigma}^2)$$

BIC

- BIC is derived from a Bayesian point of view, but ends up looking similar to C_p (and AIC) as well. For the least squares model with d predictors, the BIC is, up to irrelevant constants, given by

$$\log(n) > 2 \Rightarrow n > 7$$

$$\text{BIC} = \frac{1}{n} (\text{RSS} + \underbrace{\log(n)}_{\text{blue}} d \hat{\sigma}^2)$$

$$C_p = \frac{1}{n} (\text{RSS} + \underbrace{2}_{\text{blue}} d \hat{\sigma}^2)$$

Adjusted R^2

$$R^2 \equiv 1 - \frac{RSS}{TSS} ; \quad TSS = \sum (y_i - \bar{y})^2$$

$$\text{Adjusted } R^2 \equiv 1 - \frac{RSS/(n-d-1)}{TSS/(n-1)}$$

Choosing the Optimal Model

For a fitted least squares model containing d predictors

- C_p

$$C_p = \frac{1}{n} (\text{RSS} + \underline{2d\hat{\sigma}^2})$$

Handwritten notes:
• $\frac{\text{RSS}}{n} = \text{MSE of training}$
• $\underline{2d\hat{\sigma}^2}$ is a measure for MSE for testing

- Akaike information criterion (AIC)

$$\text{AIC} = \frac{1}{n\hat{\sigma}^2} (\text{RSS} + 2d\hat{\sigma}^2)$$

- Bayesian information (BIC)

$$\text{BIC} = \frac{1}{n} (\text{RSS} + \log(n)d\hat{\sigma}^2)$$

- Adjusted R^2

where $\hat{\sigma}^2$ is an estimate of the variance of the error associated with each response measurement

$$\text{Adjusted } R^2 = 1 - \frac{\text{RSS}/(n - d - 1)}{\text{TSS}/(n - 1)}$$

Regression

$$y = f(x) + \epsilon$$

↑
True 'f'

Data

$$D = \{(\bar{x}_i, y_i)\}_{i=1}^n$$

$$\hat{y} = h(x)$$

↳ Parametric model/Function $h_{\beta}(x) = \hat{f}(x, \beta)$

$$\text{E.g. } \hat{f}(x, \beta) = \beta_0 + \beta_1 x^{(1)} + \beta_2 x^{(2)}$$

ML is to use data to find β .

Minimize Cost Function $L(\beta)$

$$\hat{\beta} = \arg \min_{\beta} L(\beta)$$

Null Hypothesis

$\Rightarrow x_1, x_2, \dots$ are
true features or not?

Is $\beta_i = 0$?

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \epsilon$$

Which of the X_i 's are actually important?

Obj: Choose the most imp. 'd' variable from the given 'p' variables.

- 1.) Subset Selection:
- 2.) Shrinkage Methods
- 3.) Dimension Reduction

Subset Selection

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_p x_p$$

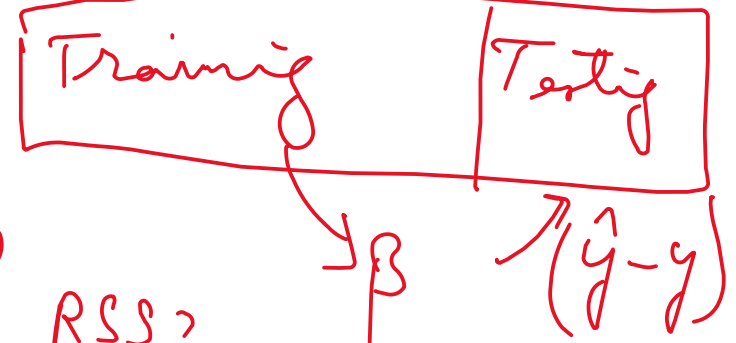
$$h_{\beta}(x, \beta)$$

$$\left\{ \begin{array}{l} x_1, x_2, \dots, x_p \\ 0 \text{ variable: } 1 \text{ way } (y = \beta_0) \\ 1 \text{ variable: } p \text{ ways } (y = \beta_0 + \beta_1 x) \\ 2 \text{ Variables: } \frac{p(p-1)}{2} = {}^p C_2 \\ d \text{ variables: } {}^p C_d = \frac{{}^p P_d}{d!} \\ p \text{ variable: } 1 \text{ way} \end{array} \right.$$

$$1 + {}^p C_1 + {}^p C_2 + \dots + {}^p C_d + \dots + {}^p C_p = (1+1)^p = 2^p$$

Subset Selection ^{Data}

Which one is better?



① $\hat{y} = \beta_0 + \beta_1 x_1$

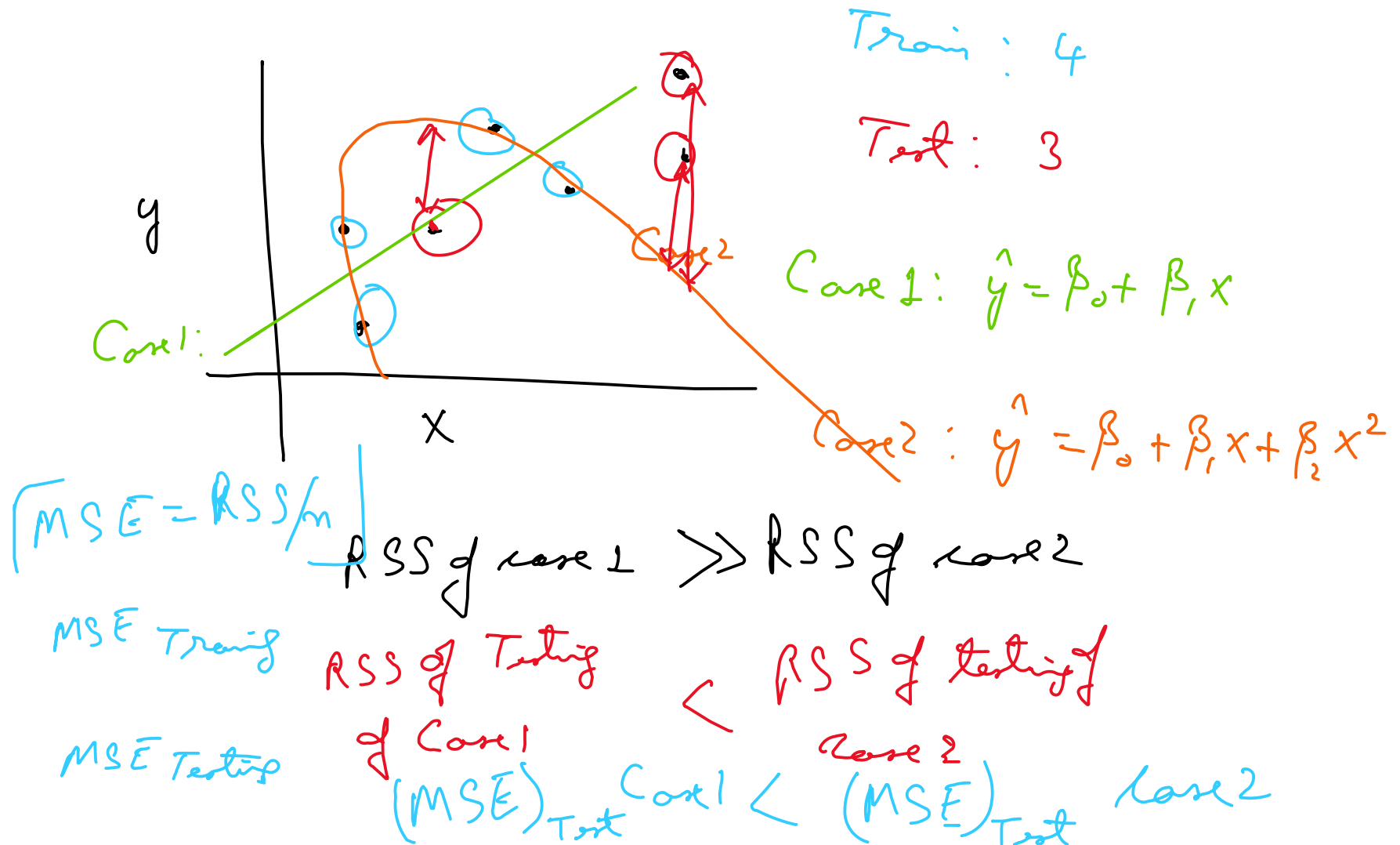
② $\hat{y} = \beta_0 + \beta_1 x_1 + \beta_2 x_2$

$\left. \begin{array}{l} \text{RSS of } \textcircled{1} \\ \text{RSS of } \textcircled{2} \end{array} \right\} \begin{array}{l} \text{Testing} \\ \text{Training} \end{array}$

$$\text{RSS of } \textcircled{1} \geq \text{RSS of } \textcircled{2}$$

Better $\Rightarrow$ Lower MSE of Testing

Subset Selection



Subset Selection

- We cannot try all the 2^p models that contain subsets of p variables.
 - Thus we have strategies for choosing the subset.
 - Various statistics can be used to judge the quality of a model. These include
 - Mallow's C_p , Akaike information criterion (AIC),
 - Bayesian information criterion (BIC), and
 - Adjusted R^2
- Handwritten red notes:* } MSE
y_{Test}

Subset Selection

- We cannot try all the 2^p models that contain subsets of p variables.
- Thus we have strategies for choosing the subset.
- Various statistics can be used to judge the quality of a model. These include
 - Mallow's C_p , Akaike information criterion (AIC),
 - Bayesian information criterion (BIC), and
 - Adjusted R^2

Strategies for choosing the subset

1. Best-Subset Selection:

Best subset regression finds for each $k \in \{0, 1, 2, \dots, p\}$ the subset of size k that gives smallest MSE for testing

2. Forward- and Backward-Stepwise Selection:

its selection starts with the intercept, and then sequentially adds into the model the predictor that most improves the fit.

Subset Selection

$k=1$, $\mathcal{M}_1$
 $k=2$, $\mathcal{M}_2$
 $\mathcal{M}_3$
:
 $\mathcal{M}_p$

- *Best Subset Selection*

Algorithm

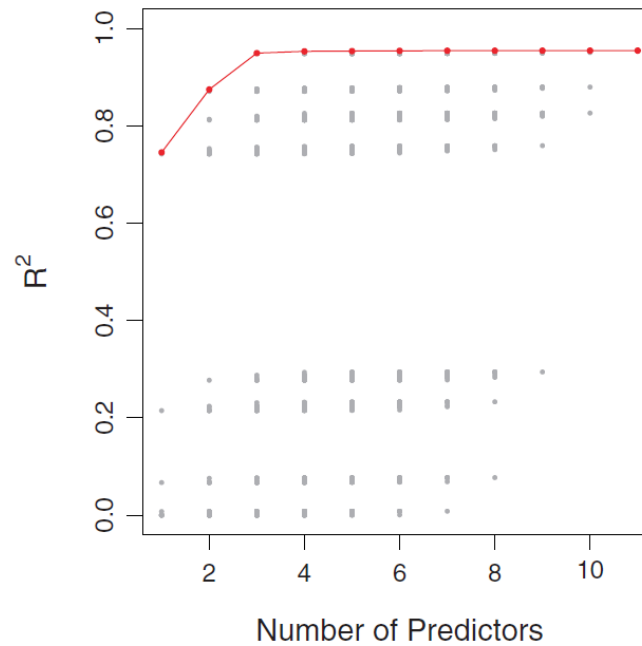
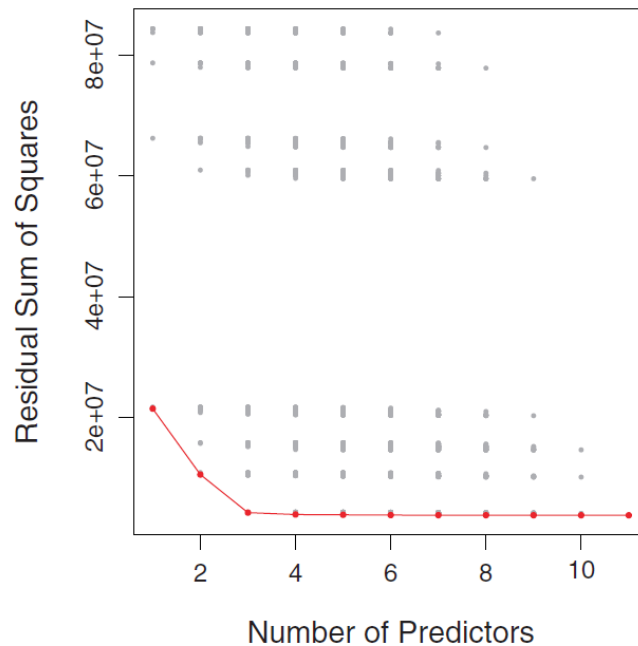
1. Let $\mathcal{M}_0$ denote the *null model*, which contains no predictors. This model simply predicts the sample mean for each observation.
2. For $k = 1, 2, \dots, p$:
 - (a) Fit all $\binom{p}{k}$ models that contain exactly k predictors.
 - (b) Pick the best among these $\binom{p}{k}$ models, and call it $\mathcal{M}_k$. Here *best* is defined as having the smallest RSS, or equivalently largest R^2 .
3. Select a single best model from among $\mathcal{M}_0, \dots, \mathcal{M}_p$ using cross-validated prediction error, C_p (AIC), BIC, or adjusted R^2 .

$\binom{p}{k}$



Subset Selection

- Best Subset Selection*



For each possible model containing a subset of the ten predictors in the Credit data set, the RSS and R^2 are displayed. The red frontier tracks the best model for a given number of predictors, according to RSS and R^2 . Though the data set contains only ten predictors, the x-axis ranges from 1 to 11, since one of the variables is categorical and takes on three values, leading to the creation of two dummy variables.

Strategies for choosing the subset

1. Forward selection

- We begin with the null model—a model that contains an intercept but no predictors.
- We then fit p simple linear regressions and add to the null model the variable that results in the lowest RSS.
- We then add to that model the variable that results in the lowest RSS for the new two-variable model.
- This approach is continued until some stopping rule is satisfied.

Subset Selection

- *Stepwise Selection*
 - Forward Stepwise Selection

-
1. Let $\mathcal{M}_0$ denote the *null* model, which contains no predictors.
 2. For $k = 0, \dots, p - 1$:
 - (a) Consider all $p - k$ models that augment the predictors in $\mathcal{M}_k$ with one additional predictor.
 - (b) Choose the *best* among these $p - k$ models, and call it $\mathcal{M}_{k+1}$. Here *best* is defined as having smallest RSS or highest R^2 .
 3. Select a single best model from among $\mathcal{M}_0, \dots, \mathcal{M}_p$ using cross-validated prediction error, C_p (AIC), BIC, or adjusted R^2 .
-

Strategies for choosing the subset

2. Backward selection

- We start with all variables in the model, and remove the variable with the largest p-value—that is, the variable
- that is the least statistically significant.
- The new $(p - 1)$ -variable model is fit, and the variable with the largest p-value is removed.
- This procedure continues until a stopping rule is reached.
- For instance, we may stop when all remaining variables have a p-value below some threshold.

Subset Selection

- *Stepwise Selection*

- Backward Stepwise Selection

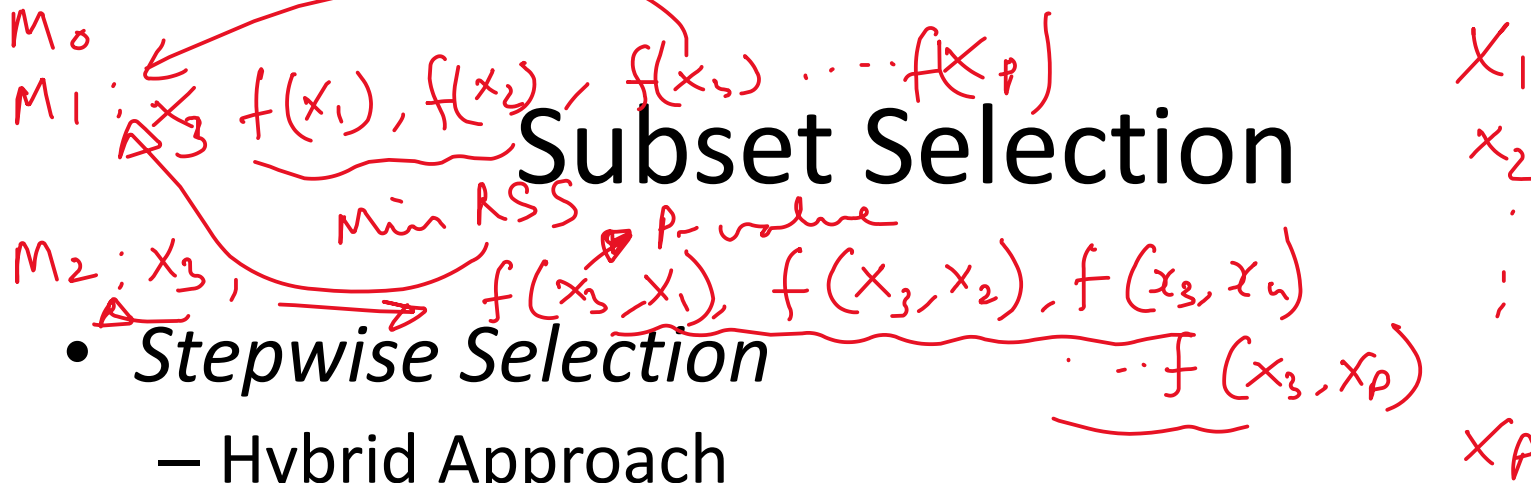
-
1. Let $\mathcal{M}_p$ denote the *full* model, which contains all p predictors.
 2. For $k = p, p - 1, \dots, 1$:
 - (a) Consider all k models that contain all but one of the predictors in $\mathcal{M}_k$, for a total of $k - 1$ predictors.
 - (b) Choose the *best* among these k models, and call it $\mathcal{M}_{k-1}$. Here *best* is defined as having smallest RSS or highest R^2 .
 3. Select a single best model from among $\mathcal{M}_0, \dots, \mathcal{M}_p$ using cross-validated prediction error, C_p (AIC), BIC, or adjusted R^2 .
-

Strategies for choosing the subset

3. Mixed selection

- This is a combination of forward and backward selection.
- We start with no variables in the model, and as with forward selection, we add the variable that provides the best fit.
- We continue to add variables one-by-one.
- It is possible sometimes that the p-values for variables can become larger as new predictors are added to the model.
- Hence, if at any point the p-value for one of the variables in the model rises above a certain threshold, then we remove that variable from the model.
- We continue to perform these forward and backward steps until all variables in the model have a sufficiently low p-value, and all variables outside the model would have a large p-value if added to the model.

Subset Selection



• Stepwise Selection

– Hybrid Approach

- Variables are added to the model sequentially, in analogy to forward selection.
- However, after adding each new variable, the method may also remove any variables that no longer provide an improvement in the model fit.
- Such an approach attempts to more closely mimic best subset selection while retaining the computational advantages of forward and backward stepwise selection.

Linear Model Selection and Regularization

- **Shrinkage**. This approach involves fitting a model involving all p predictors. However, the estimated coefficients are shrunk towards zero relative to the least squares estimates. This shrinkage (also known as regularization) has the effect of reducing variance. Depending on what type of shrinkage is performed, some of the coefficients may be estimated to be exactly zero. Hence, shrinkage methods can also perform variable selection.

Find β s by min RSS

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_p x_p$$

Shrinkage

- By retaining a subset of the predictors and discarding the rest, subset selection produces a model that is interpretable and has possibly lower prediction error than the full model.
- However, because it is a discrete process variables are either retained or discarded—it often exhibits high variance, and so doesn't reduce the prediction error of the full model.
- Shrinkage methods are more continuous, and don't suffer as much from high variability.

Shrinkage Methods

1. Ridge Regression

- Ridge regression shrinks the regression coefficients by imposing a penalty on their size. The ridge coefficients minimize a penalized residual sum of squares,

$$\hat{\beta}^{\text{ridge}} = \underset{\beta}{\operatorname{argmin}} \left\{ \sum_{i=1}^N (y_i - \beta_0 - \sum_{j=1}^p x_{ij} \beta_j)^2 + \lambda \sum_{j=1}^p \beta_j^2 \right\}.$$

- Here $\lambda \geq 0$ is a complexity parameter that controls the amount of shrinkage. It's directly proportional to shrinkage .

Ridge Regression

$$RSS(\beta)_{\text{Ridge}} = \left\| X\beta - Y \right\|_2^2 + \lambda \left\| \beta \right\|_2^2$$

$$RSS(\beta) = \left\| X\beta - Y \right\|_2^2$$

↑

min

constraint

$$\left\| \beta \right\|_2^2 < t$$

↑
+ve

Ridge Regression

$$RSS(\beta) = \underbrace{\quad}_{\text{Ridge}} \quad ||X\beta - Y||_2^2 + \lambda ||\beta||_2^2$$

$$RSS = ||X\beta - Y||_2^2$$

↑ min

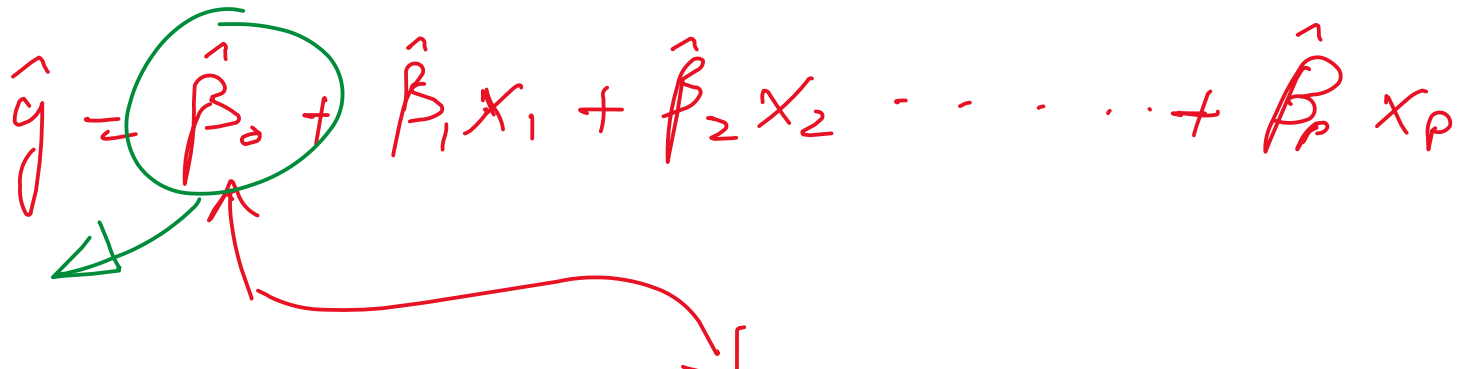
Constrain
 β

$$\sum \beta_i^2 < t \rightarrow$$

↑
+ve user defined

1000 kg
 $10^3 \times 10^3 g$

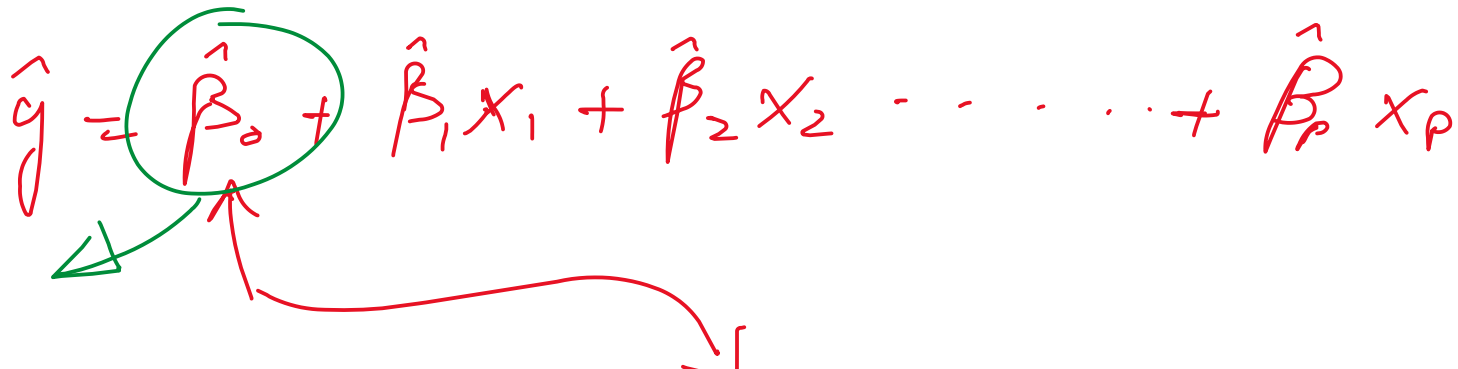
Ridge Regression

$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 x_2 + \dots + \hat{\beta}_p x_p$$


$$\sum \beta_i^2 < k$$

$$RSS = \sum_{i=1}^n (\hat{y} - y)^2 = \sum (y - \underbrace{\hat{\beta}_0}_{\text{green}} + \hat{\beta}_1 x_1 + \dots)^2$$

Ridge Regression

$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 x_2 + \dots + \hat{\beta}_p x_p$$


$$\sum \beta_i^2 < k$$

$$RSS = \sum_{i=1}^n (\hat{y} - y)^2 = \sum (y - \underbrace{\hat{\beta}_0}_{\text{green}} + \hat{\beta}_1 x_1 + \dots)^2$$

Ridge Regression

$$\underset{\text{min}}{\text{RSS}}(\beta) = \underset{\text{Ridge}}{\|X\beta - y\|_2^2 + \lambda \|\beta\|_2^2}$$

Tends to 0 as $\lambda \uparrow$

$$\equiv \underset{\text{min}}{\text{RSS}}(\beta) = \|X\beta - y\|_2^2$$

Constraint $\|\beta\|_2^2 < t$

- Has no β_0
- x are normalized

Lasso Regression

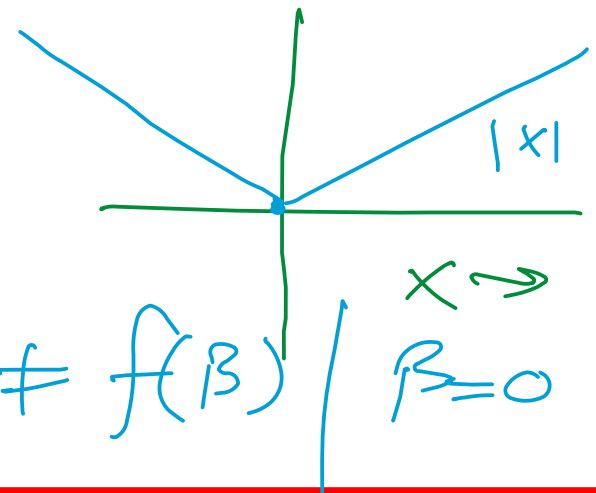
$$RSS(\beta) \underset{\text{Ridge}}{=} \|X\beta - y\|_2^2 + \lambda \|\beta\|_1$$

$$\min RSS(\beta) = \|X\beta - y\|_2^2$$

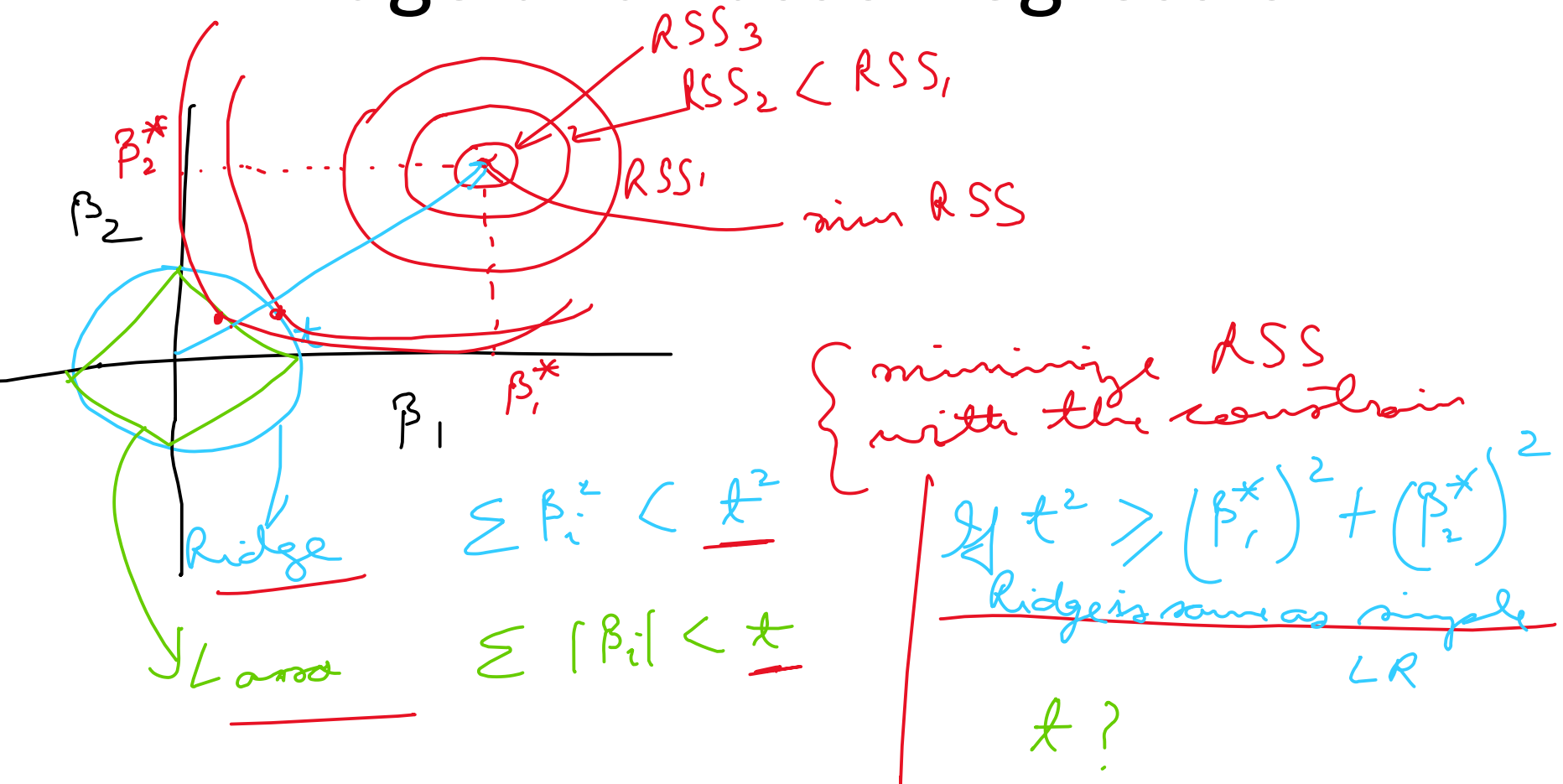
constraint $\sum |\beta| < t$

$$\lambda \frac{\partial \beta^2}{\partial \beta} = \underline{\underline{2\lambda\beta}}$$

$$\lambda \frac{\partial |\beta|}{\partial \beta} \neq f(\beta) \mid \beta=0$$

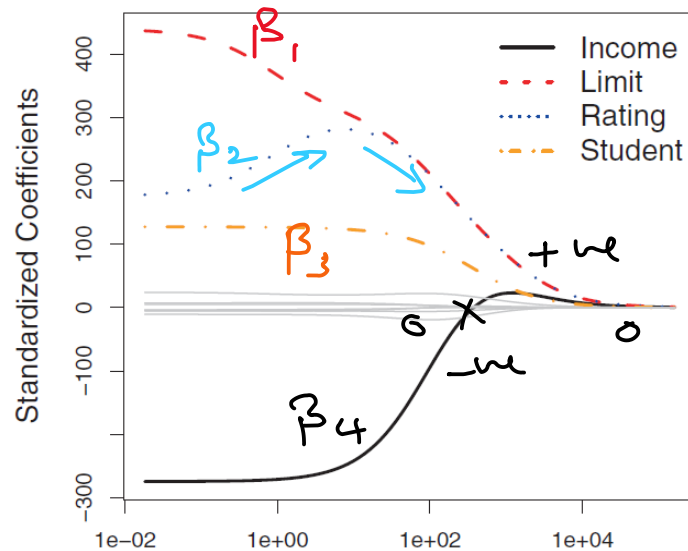


Ridge and Lasso Regression

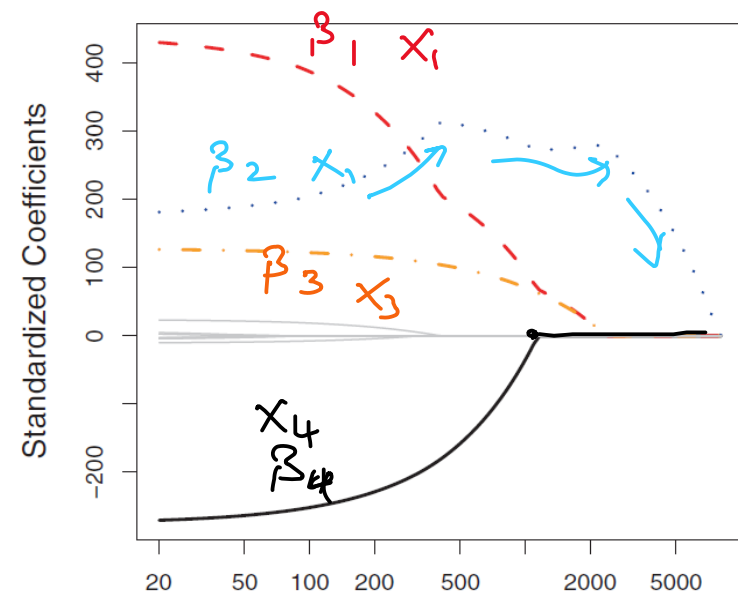


Ridge and Lasso Regression

$$RSS(\beta)$$



Ridge Regression



Lasso Regression

Ridge and Lasso Regression



$$RSS_{min} < RSS_1 < RSS_2 < RSS_3$$

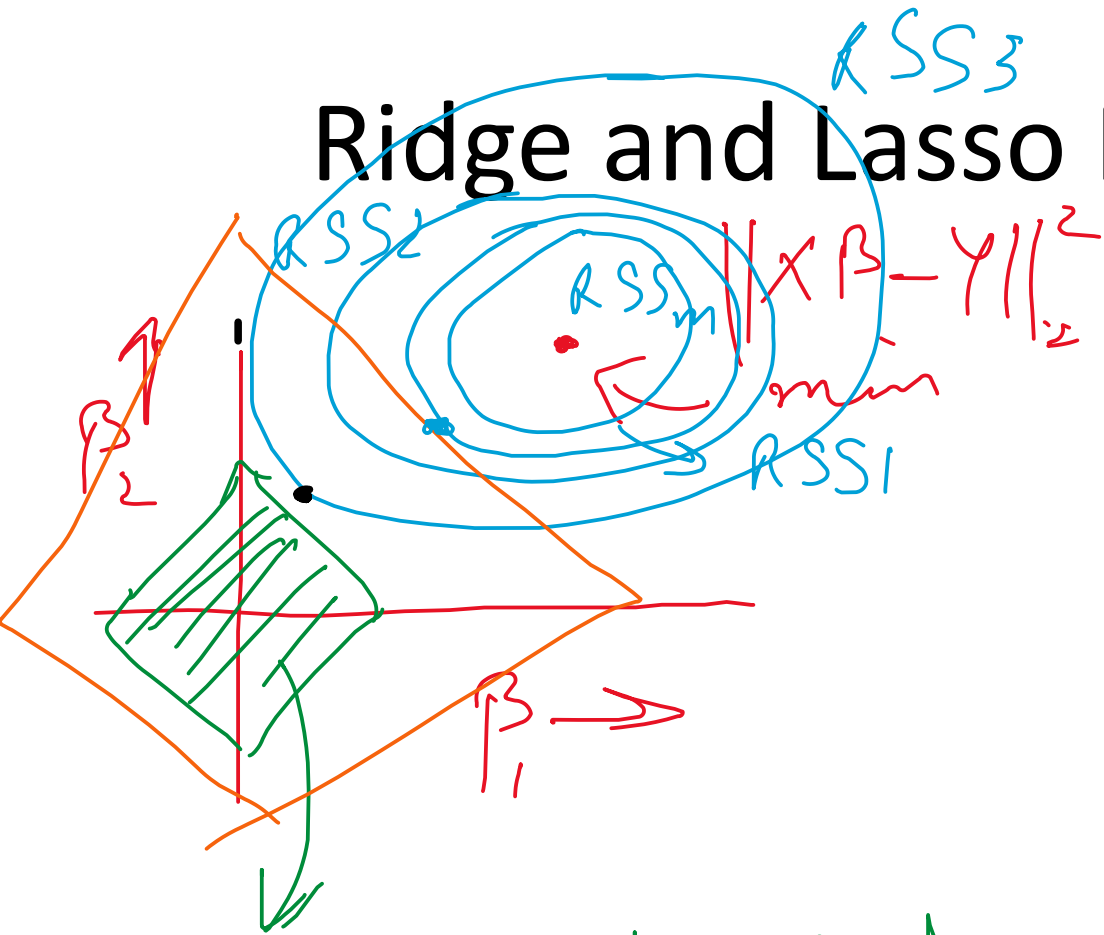
Tangent Point

Ridge $\sqrt{\sum \beta_i^2} < t$

$$\beta_1^2 + \beta_2^2 < t$$

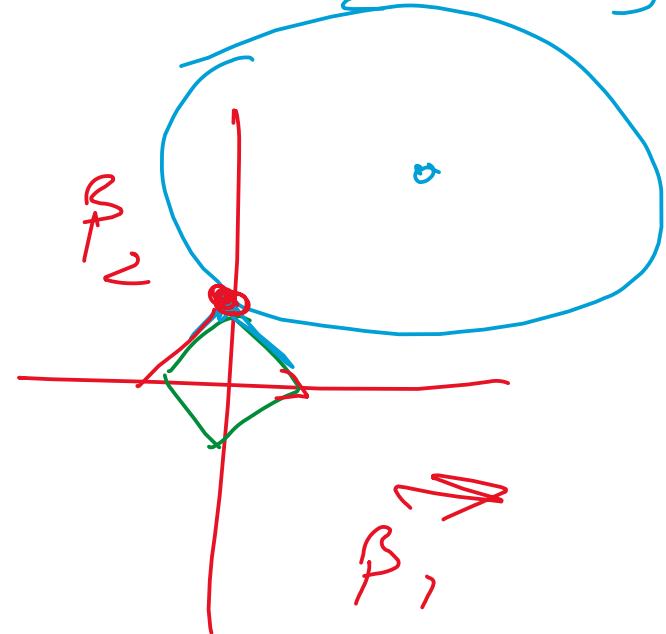
Circle

Ridge and Lasso Regression



Lasso: $|\beta_1| + |\beta_2| < t$

$$RSS_{min} < RSS_1 < RSS_2 < RSS_3$$



Ridge and Lasso Regression

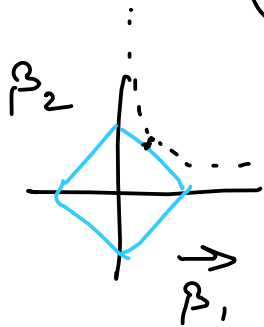
$$RSS(\beta) + \lambda \sum |\beta_i|^q$$

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2$$

$$q = 0 \rightarrow OLR$$

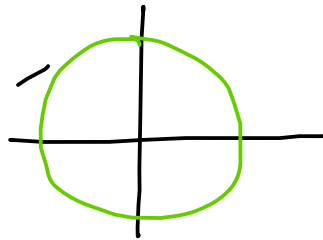
$$q = 1 \text{ Lasso}$$

$$q = 2 \text{ Ridge}$$



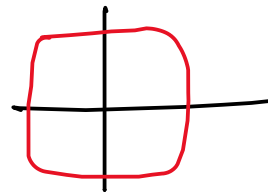
$$q = 1$$

Lasso

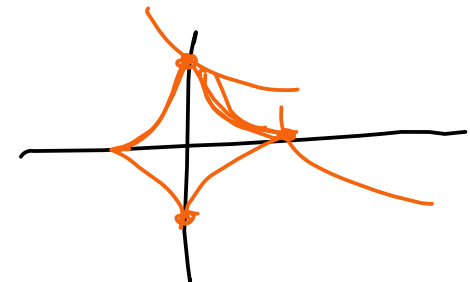


$$q = 2$$

Ridge



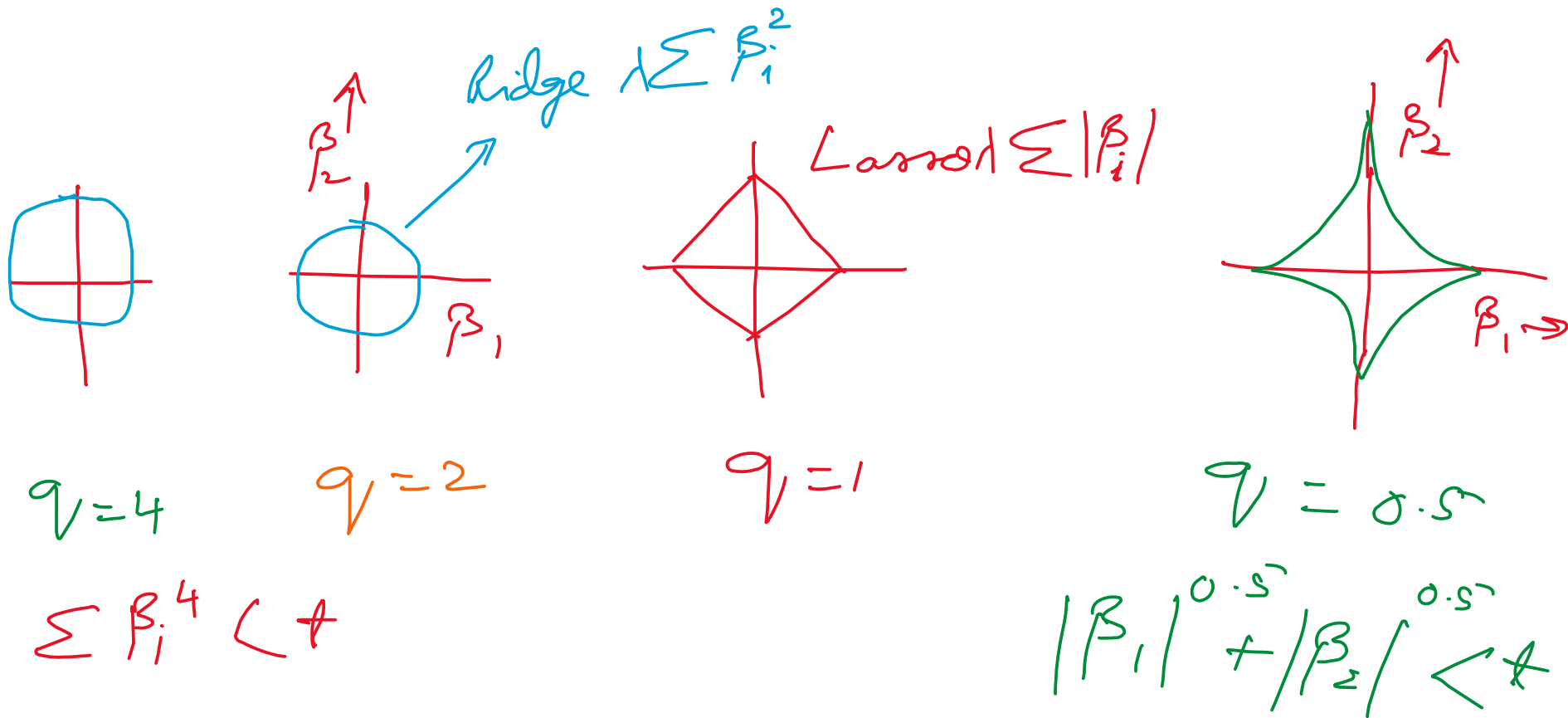
$$q = 4$$



$$q = 0.5$$

Ridge and Lasso Regression

General Lasso $\lambda \sum |\beta_i|^q$



Shrinkage Methods

- An equivalent way to write the ridge problem is

$$\hat{\beta}^{\text{ridge}} = \underset{\beta}{\operatorname{argmin}} \sum_{i=1}^N \left(y_i - \beta_0 - \sum_{j=1}^p x_{ij} \beta_j \right)^2,$$
$$\text{subject to } \sum_{j=1}^p \beta_j^2 \leq t,$$

Ridge Regression

- Ridge regression is very similar to least squares, except that the coefficients ridge regression are estimated by minimizing a slightly different quantity.

$$\sum_{i=1}^n \left(y_i - \beta_0 - \sum_{j=1}^p \beta_j x_{ij} \right)^2 + \lambda \sum_{j=1}^p \beta_j^2 = \text{RSS} + \lambda \sum_{j=1}^p \beta_j^2,$$

- where $\lambda \geq 0$ is a tuning parameter. The second term, is called a shrinkage penalty, it is small when $\beta_1, \dots, \beta_p$ are close to zero

Shrinkage Methods

2.The lasso

Lasso is a shrinkage method like ridge, with subtle but important differences. The lasso estimate is defined by

$$\hat{\beta}^{\text{lasso}} = \underset{\beta}{\operatorname{argmin}} \sum_{i=1}^N \left(y_i - \beta_0 - \sum_{j=1}^p x_{ij} \beta_j \right)^2$$

subject to $\sum_{j=1}^p |\beta_j| \leq t.$

Shrinkage Methods

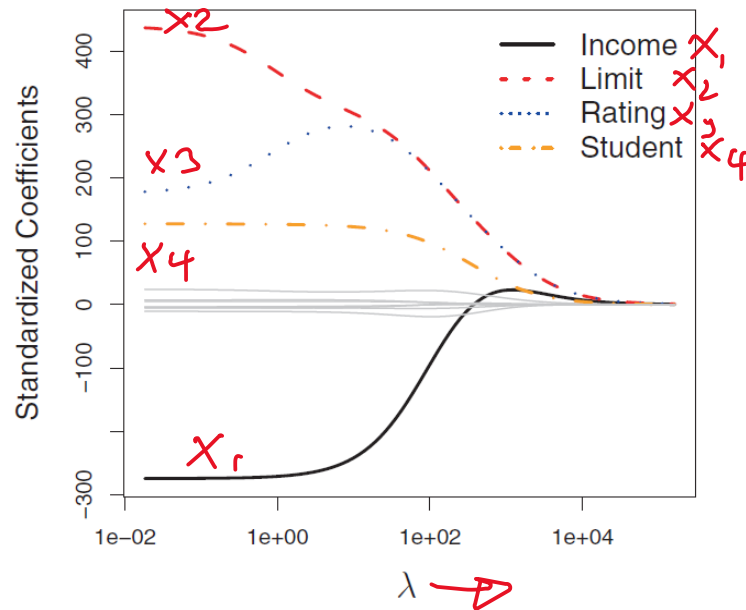
- We can also write the lasso problem in the equivalent Lagrangian form

$$\hat{\beta}^{\text{lasso}} = \underset{\beta}{\operatorname{argmin}} \left\{ \frac{1}{2} \sum_{i=1}^N (y_i - \beta_0 - \sum_{j=1}^p x_{ij} \beta_j)^2 + \lambda \sum_{j=1}^p |\beta_j| \right\}.$$

$$\frac{\partial (RSS)_{\text{Lasso}}}{\partial \beta_1} = \dots + \lambda \quad \text{Lasso}$$

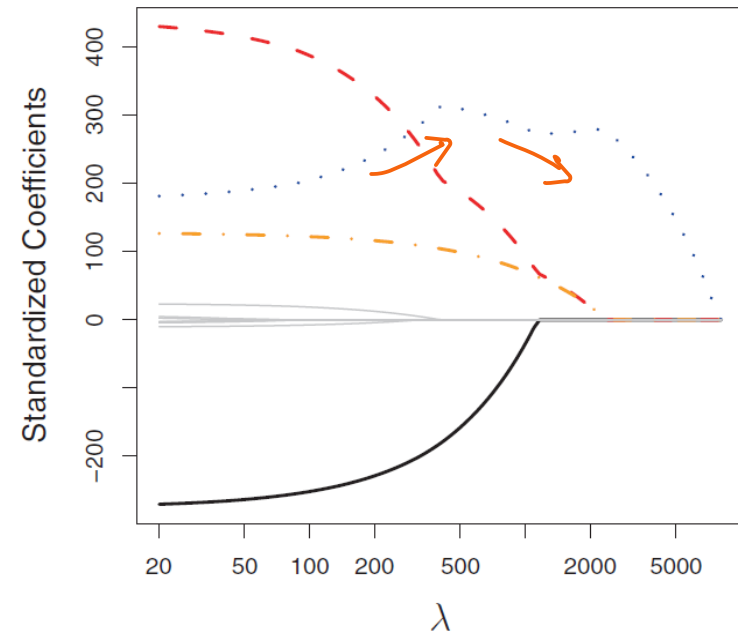
$$\frac{\partial (RSS)_{\text{Ridge}}}{\partial \beta_1} = \dots + 2\lambda \beta_1$$

Ridge and Lasso Regression



Ridge Regression

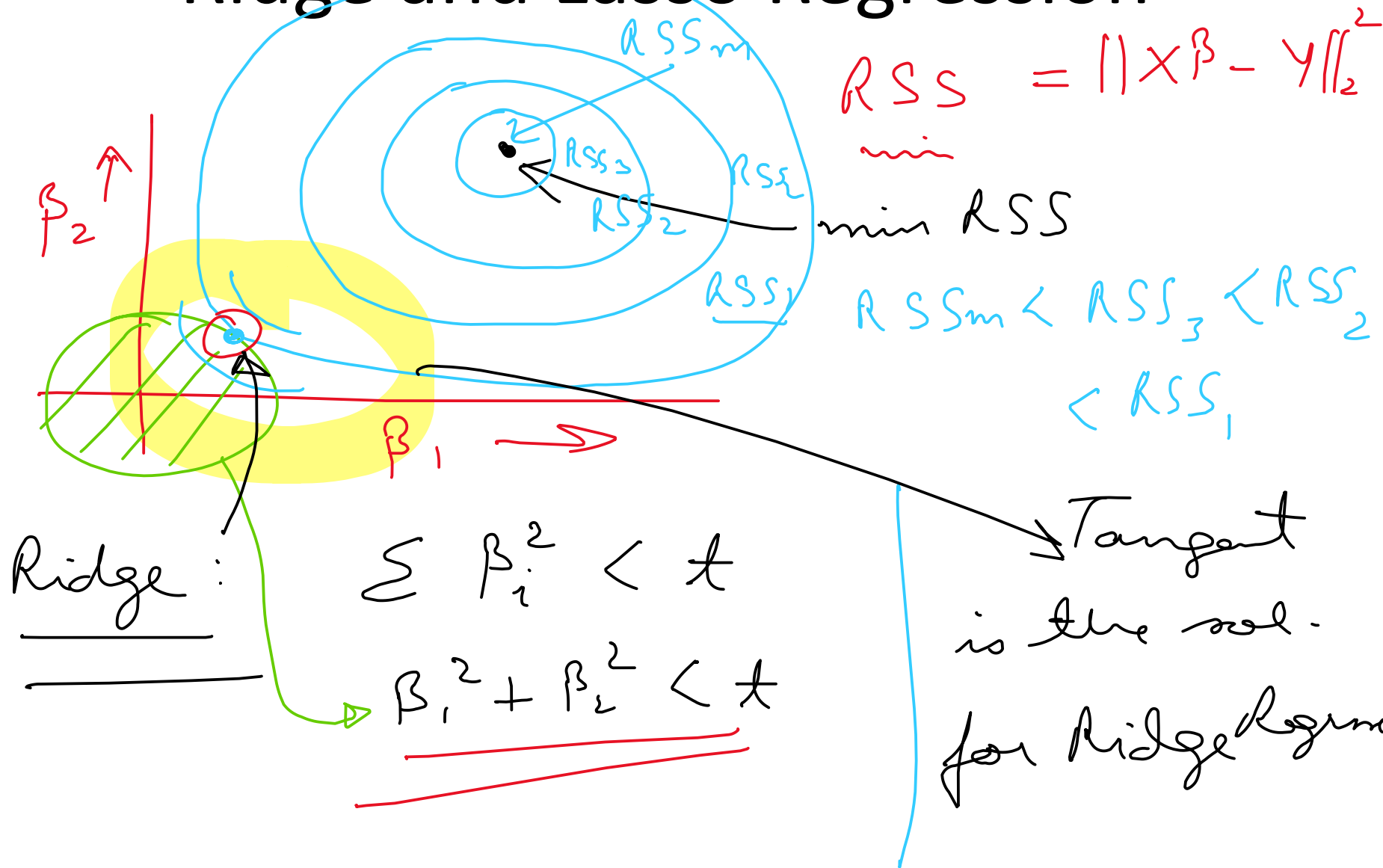
$$\lambda \sum \beta_i^2$$



Lasso Regression

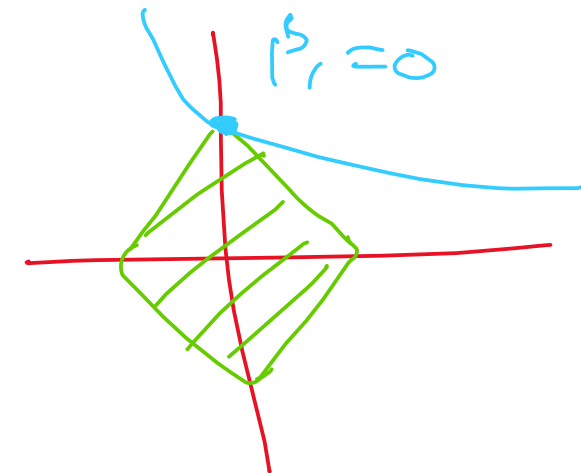
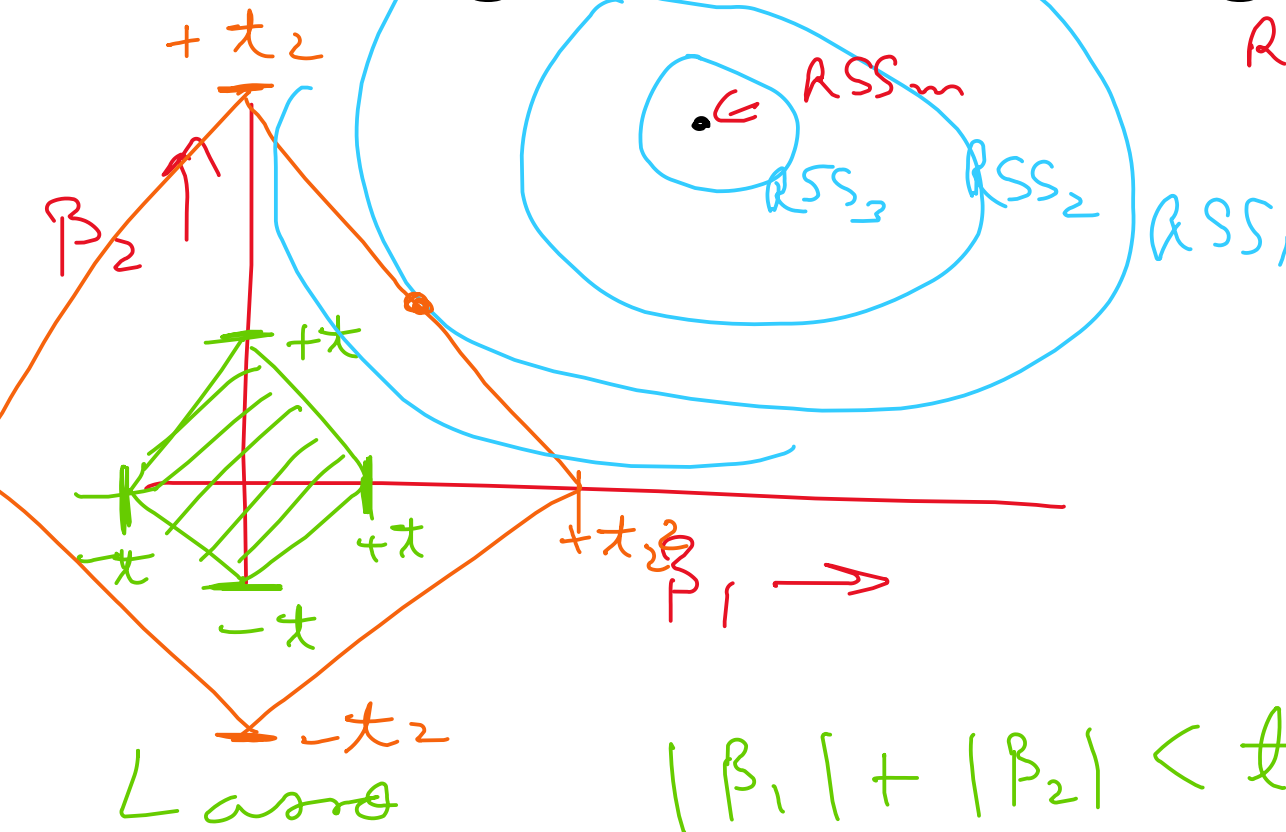
$$\lambda \sum |\beta_i|$$

Ridge and Lasso Regression



Ridge and Lasso Regression

$$RSS = \|X\beta - y\|_2^2$$



$$|\beta_1| + |\beta_2| < t$$

Summary

- Ridge regression does a proportional shrinkage.
- Lasso translates each coefficient by a constant factor, truncating at zero. This is called “soft thresholding,”.
- Best-subset selection drops all variables with coefficients smaller than the M th largest; this is a form of “hard-thresholding.”

Dimensionality reduction

- Feature selection: Feature selection approaches try to find a subset of the original variables (also called features or attributes)
 - Filter strategy (e.g. information gain)
 - Wrapper strategy (e.g. search guided by accuracy)
 - Embedded strategy (features are selected to add or be removed while building the model based on the prediction errors)
- Feature projection: Feature projection transforms the data in the high-dimensional space to a space of fewer dimensions. The data transformation may be linear, as in principal component analysis (PCA), but many nonlinear dimensionality reduction techniques also exist. For multidimensional data, tensor representation can be used in dimensionality reduction through multilinear subspace learning.

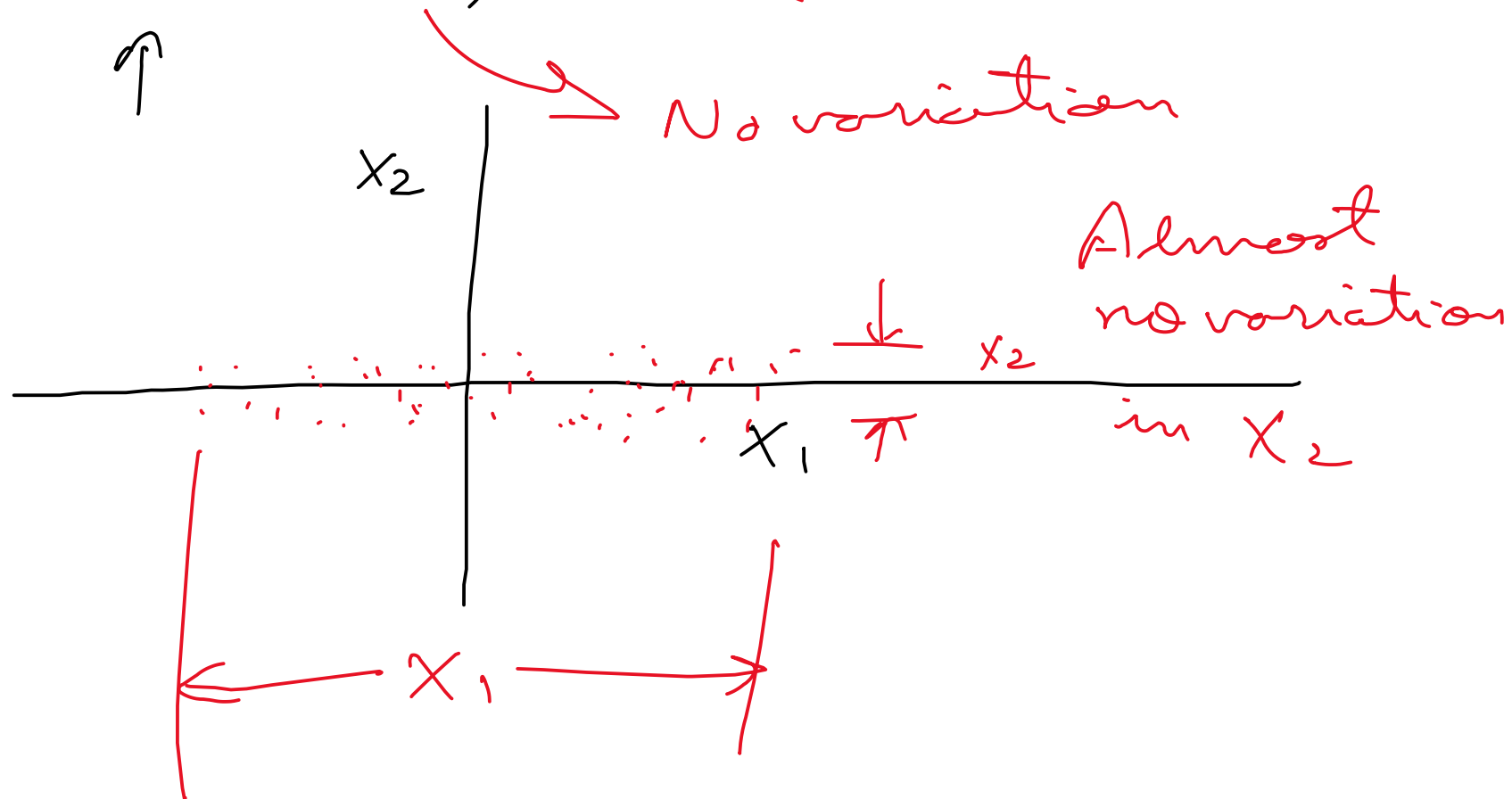
Feature projection

- Principal component analysis (PCA)
- Non-negative matrix factorization (NMF)
- Kernel PCA
- Graph-based kernel PCA
- Linear discriminant analysis (LDA)
- Generalized discriminant analysis (GDA)
- Autoencoder

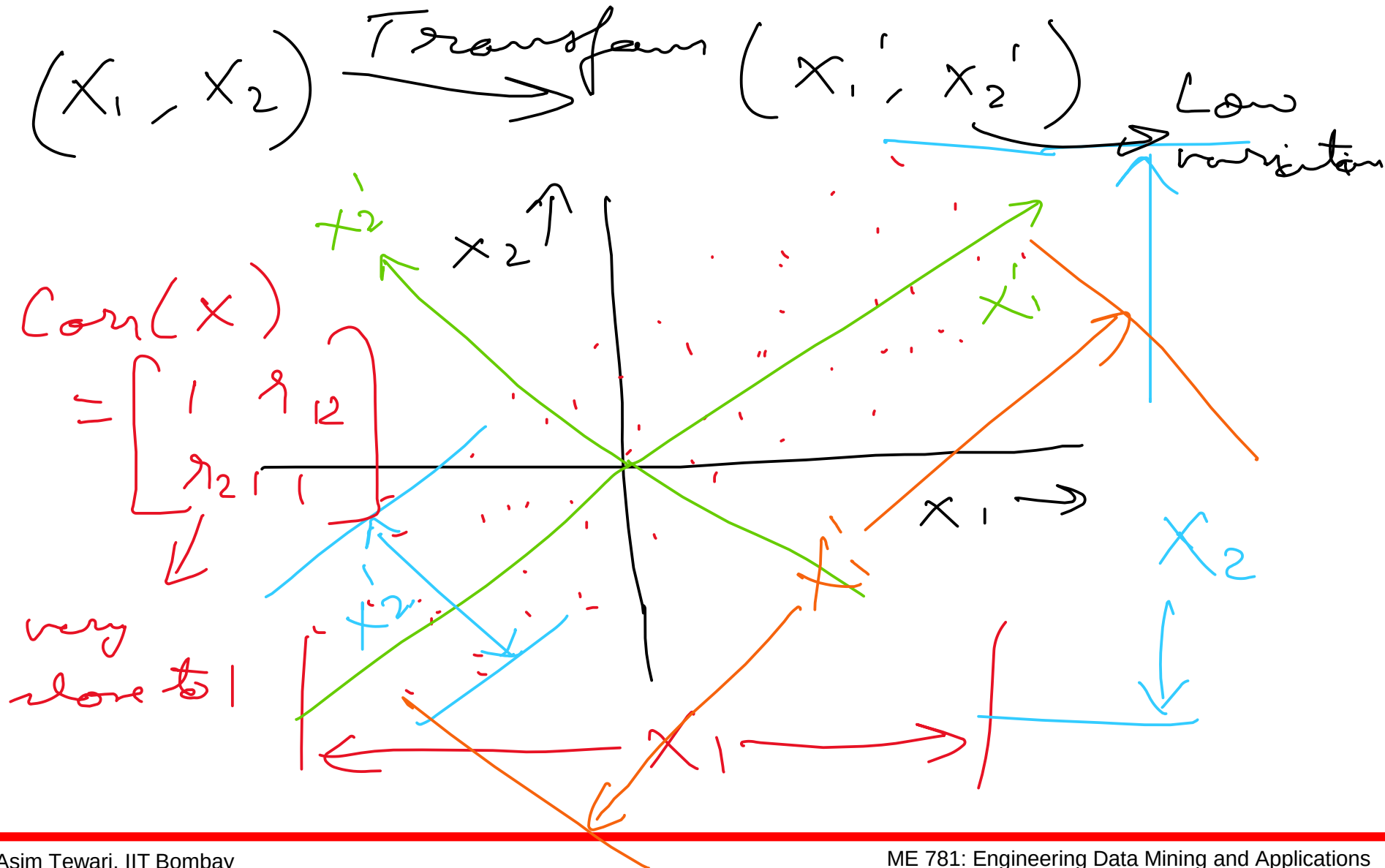
Principal Component Analysis

$$y = f(x_1, x_2) \approx f(x_1)$$

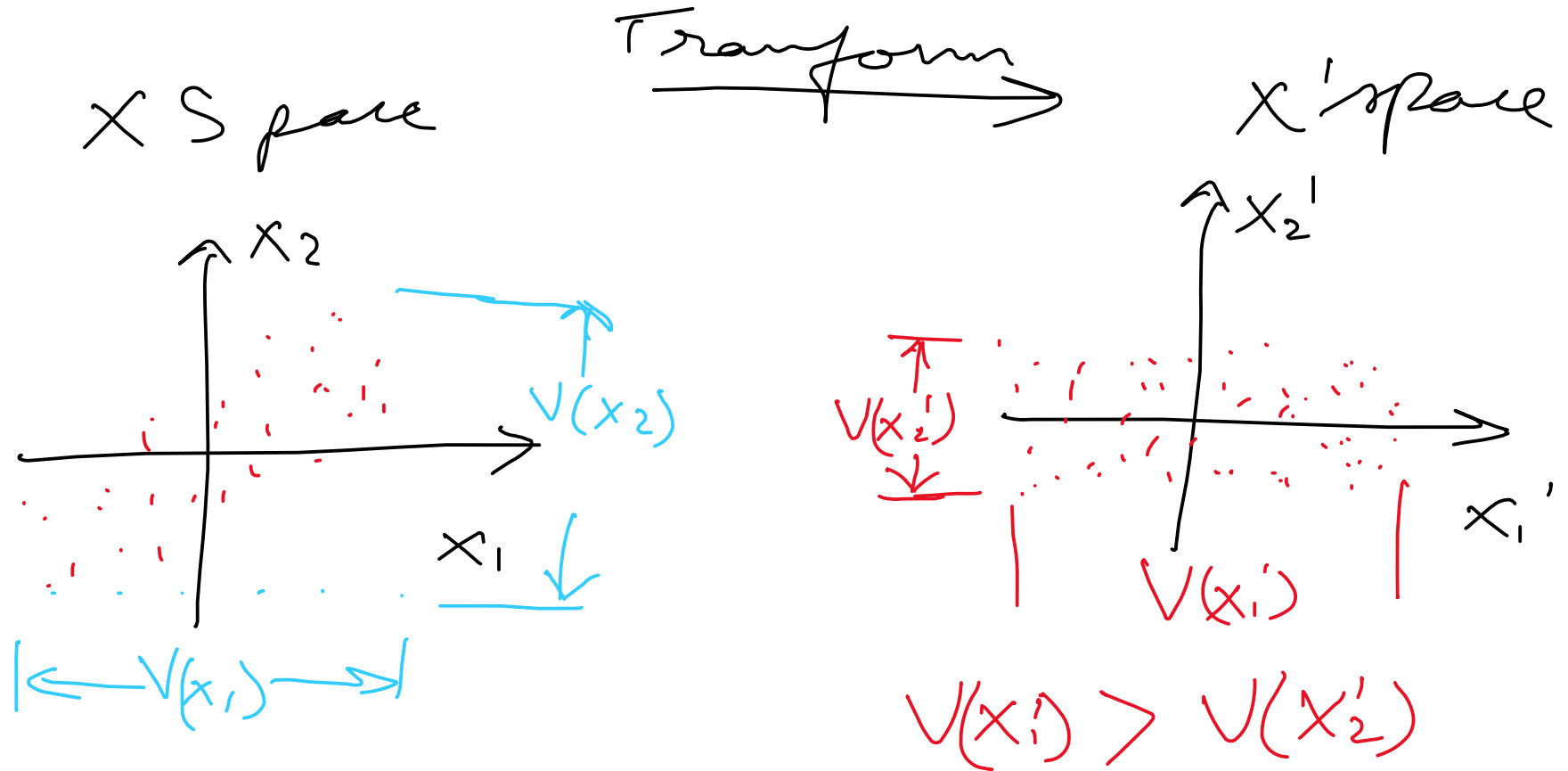
$\uparrow$



Principal Component Analysis



Principal Component Analysis



Principal Component Analysis

In 2D space with n data points

X space $\rightarrow$ X' space

$$X_{n \times 2} = \begin{bmatrix} x_{11} & x_{12} \\ x_{21} & x_{22} \\ \vdots & \vdots \\ x_{n1} & x_{n2} \end{bmatrix} \Rightarrow X'_{n \times 2} = \begin{bmatrix} x'_{11} & x'_{12} \\ x'_{21} & x'_{22} \\ \vdots & \vdots \\ x'_{n1} & x'_{n2} \end{bmatrix}$$

variance

$$V(x'_1) > V(x'_2)$$

$\uparrow$
highest possible variance

Principal Component Analysis

n p -dimensional Space

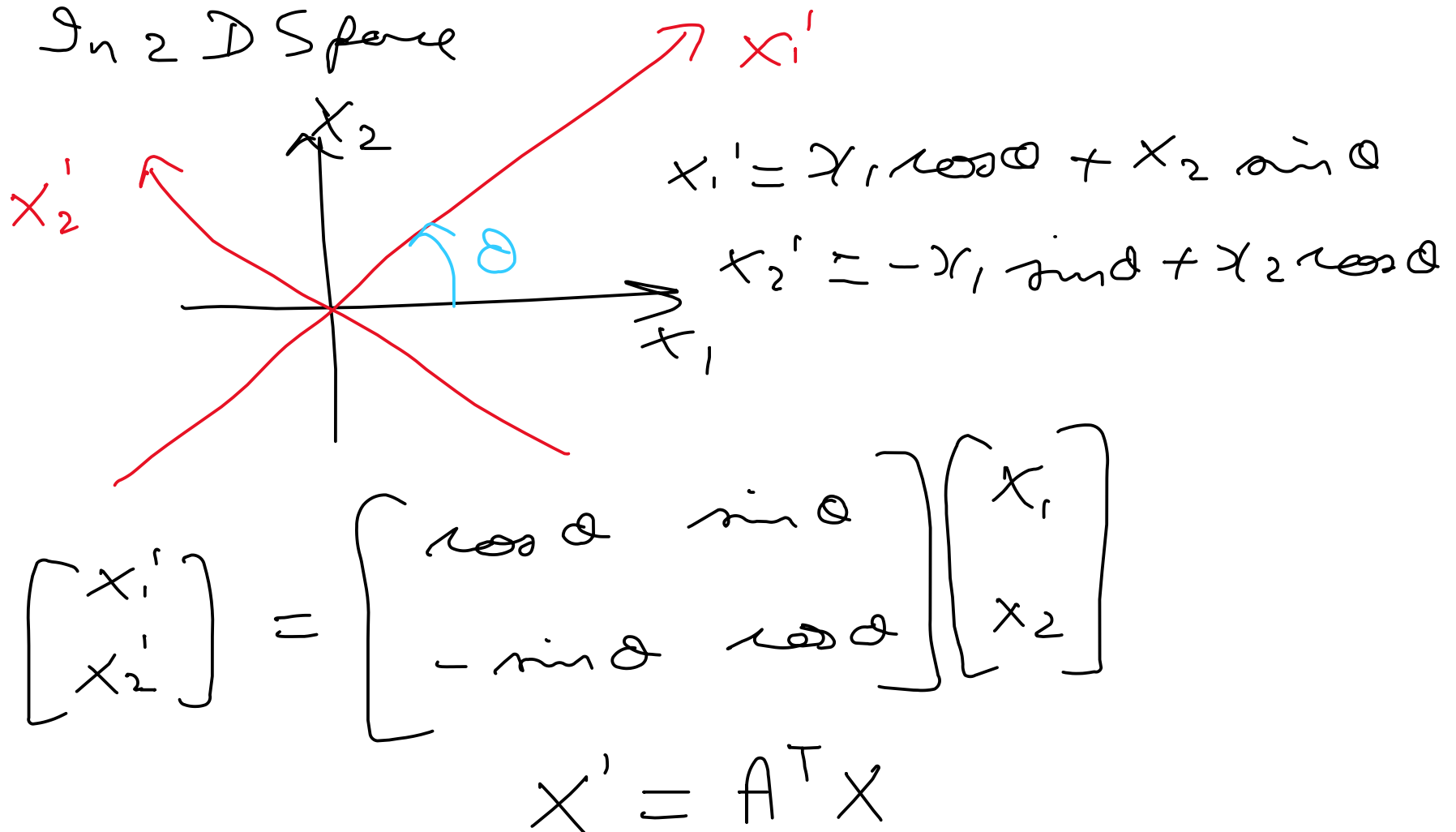
PCA
 x space $\rightarrow$ x' space

$$X_{n \times p} = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{bmatrix}$$

$$X'_{n \times p} = \begin{bmatrix} x'_{11} & x'_{12} & \dots & x'_{1p} \\ x'_{21} & x'_{22} & \dots & x'_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ x'_{n1} & x'_{n2} & \dots & x'_{np} \end{bmatrix}$$

$$V(x'_1) \geq V(x'_2) \geq \dots \geq V(x'_p)$$

Principal Component Analysis



Principal Component Analysis

In 3D

$$X' = A^T X$$

$$\begin{bmatrix} x_1' \\ x_2' \\ x_3' \end{bmatrix} = \begin{bmatrix} \cos \theta & \sin \theta & 0 \\ -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

$R_3(\theta) = A^T$ For rotation about x_3

Principal Component Analysis

$$R_2(\theta) = \begin{bmatrix} \cos \theta & 0 & \sin \theta \\ 0 & 1 & 0 \\ -\sin \theta & 0 & \cos \theta \end{bmatrix}$$

$$R_1(\theta) = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & \sin \theta \\ 0 & -\sin \theta & \cos \theta \end{bmatrix}$$

For a general rotation

$$A^T = R = R_1(\alpha) R_2(\beta) R_3(\gamma)$$

Principal Component Analysis

In general 3D rotation you can always find an axis that is not changed.

$$\begin{array}{ccc} X' & = & R X \\ \uparrow & & \uparrow \\ \text{new} & & \text{old} \\ \text{axis} & & \text{axis} \end{array}$$

Invariant axis

$$\begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

$$\begin{array}{c} X \\ \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} \end{array} = R \begin{array}{c} X \\ \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} \end{array} \Rightarrow (R - I) \underset{\substack{\uparrow \\ \text{Invariant vector}}}{X} = 0$$

Principal Component Analysis

$$\begin{matrix} X' & = & A^T X \\ P \times 1 & & P \times P \quad P \times 1 \end{matrix}$$

In 2D

$$A^T = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}$$

$$\Rightarrow A = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} = [a_1 \quad a_2]$$

$$a_1 = \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix}; a_2 = \begin{bmatrix} -\sin \theta \\ \cos \theta \end{bmatrix}$$

Principal Component Analysis

$$A = \begin{matrix} & \begin{matrix} a_1 & a_2 \end{matrix} \\ \begin{matrix} 2 \times 2 \end{matrix} & \begin{matrix} 2 \times 1 & 2 \times 1 \end{matrix} \end{matrix} \quad \left| \quad \begin{matrix} a_1 = \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix} \\ a_2 = \begin{bmatrix} -\sin \theta \\ \cos \theta \end{bmatrix} \end{matrix} \right.$$

$$a_1^T a_1 = \begin{bmatrix} \cos \theta & \sin \theta \end{bmatrix} \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix} = \cos^2 \theta + \sin^2 \theta = 1$$

$$a_2^T a_2 = 1$$

$$a_1^T a_2 = \begin{bmatrix} \cos \theta & \sin \theta \end{bmatrix} \begin{bmatrix} -\sin \theta \\ \cos \theta \end{bmatrix} = 0$$

Principal Component Analysis

$$A^T A = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = I = A A^T = A A^{-1} \\ = A^{-1} A$$

$$\Rightarrow A^T = A^{-1}$$

$\therefore A$ is an orthogonal Transformation

Principal Component Analysis

In p -dimension

$$A = [a_1 \ a_2 \ \dots \ a_p]$$

$$a_i^T a_j = \delta_{ij} = \begin{cases} 0 & \text{if } i \neq j \\ 1 & \text{if } i = j \end{cases}$$

In PCA we want this transformation to be such that variance of x' are as follows:

$$V(x_1') \geq V(x_2') \geq V(x_3') \dots \geq V(x_p')$$

Principal Component Analysis

p-dimensional Space

$$X_{n \times p} = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ \vdots & \vdots & & \vdots \\ x_{k1} & x_{k2} & \dots & x_{kp} \\ \vdots & \vdots & & \vdots \\ x_{n1} & \dots & \dots & x_{np} \end{bmatrix}_{n \times p}$$

In set form $X = \{x_1, x_2, \dots, x_n\}$
where $x_k = (x_{k1}, x_{k2}, \dots, x_{kp})$

Principal Component Analysis

∴ for the " k^{th} " Set

$$X_k = \begin{bmatrix} x_{k1} \\ x_{k2} \\ \vdots \\ x_{kp} \end{bmatrix}^T$$

Similarly

$$\begin{bmatrix} x'_{k1} \\ x'_{k2} \\ \vdots \\ x'_{kp} \end{bmatrix}^T = \begin{bmatrix} x'_{k1} & x'_{k2} & \dots & x'_{kp} \end{bmatrix} = x'_k$$

$$\begin{bmatrix} x'_{k1} \\ x'_{k2} \\ \vdots \\ x'_{kp} \end{bmatrix} = A^T \begin{bmatrix} x_{k1} \\ x_{k2} \\ \vdots \\ x_{kp} \end{bmatrix}$$

Principal Component Analysis

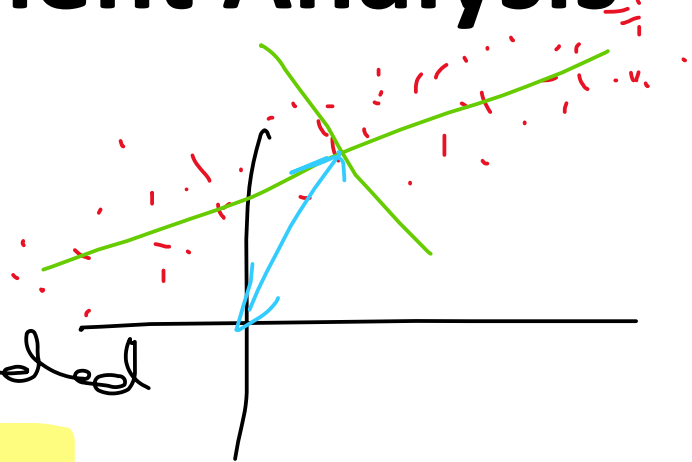
$$\begin{bmatrix} x'_{k1} \\ x'_{k2} \\ \vdots \\ x'_{kp} \end{bmatrix} = A^T \begin{bmatrix} x_{k1} \\ x_{k2} \\ \vdots \\ x_{kp} \end{bmatrix}$$

$$\begin{bmatrix} x'_{k1} \\ x'_{k2} \\ \vdots \\ x'_{kp} \end{bmatrix}^T = \begin{bmatrix} x_{k1} \\ x_{k2} \\ \vdots \\ x_{kp} \end{bmatrix}^T A$$

$$x'_{\alpha k} = x_{\uparrow k} A$$

Principal Component Analysis

$$x'_k = x_k A$$



Translation is also needed

$$x'_k = (x_k - \bar{x}) A$$

$1 \times P$ $1 \times P$ $P \times P$

$$\bar{x} = \frac{1}{n} \sum_{k=1}^n x_k = \frac{1}{n} \begin{bmatrix} \sum x_{k1} & \sum x_{k2} & \dots & \sum x_{kp} \end{bmatrix}$$

$$x_k = x'_k A^T + \bar{x}$$

Principal Component Analysis

$$x_k = x'_k A^T + \bar{x}$$

variance vector of $x' = \begin{bmatrix} \text{var}(x'_1) & \text{var}(x'_2) \\ \dots & \text{var}(x'_p) \end{bmatrix}$

covariance of $x' = \mathcal{V}_{x'} = \frac{1}{n-1} \sum_{k=1}^3 (x'_k)^T (x'_k)$

$\uparrow$ $\uparrow$ $\uparrow$
 $p \times p$ $p \times 1$ $1 \times p$

$$\mathcal{V}_{x'} = \frac{1}{n-1} \sum_{k=1}^3 (x'_k)^T (x'_k)$$

Principal Component Analysis

$$V_{X'} = \frac{1}{n-1} \sum_{k=1}^n (X'_k)^T (X'_k)$$

$$V_{X'} = \frac{1}{n-1} \sum \left[(X_k - \bar{X}) A \right]^T \left[(X_k - \bar{X}) A \right]$$

$X'_k = (X_k - \bar{X}) A$

$$= \frac{1}{n-1} \sum_{k=1}^n A^T (X_k - \bar{X})^T \cdot (X_k - \bar{X}) A$$

$$= A^T \left[\frac{1}{n-1} \sum_{k=1}^n (X_k - \bar{X})^T \cdot (X_k - \bar{X}) \right] A$$

$$V_{X'} = A^T \underbrace{C}_\text{Covariance matrix of } X A$$

Principal Component Analysis

$$\mathcal{V}_{X'} = A^T \underbrace{C}_\text{Covariance matrix of } X A$$

$$\mathcal{C}_{ij} = \frac{1}{n-1} \sum_{k=1}^n \left[x_{ki} - \frac{1}{n} \sum_{l=1}^n x_{li} \right] \left[x_{kj} - \frac{1}{n} \sum_{l=1}^n x_{lj} \right]$$

$$i, j \in \{1, 2, \dots, p\}$$

Choose A such that $\mathcal{V}_{X'}$ is maximized.

$$\text{However } A^T A = I$$

Principal Component Analysis

Maximise V_x given $A^T A = I$

Constrained optimization.

Construct a Lagrange function

$$L = A^T C A - \lambda (A^T A - I)$$

Sol is A which maximises L

$$\begin{aligned} \Rightarrow \frac{\partial L}{\partial A} = 0 &\Rightarrow C A - \lambda A = 0 \\ &\Rightarrow (C - \lambda I) A = 0 \end{aligned}$$

Principal Component Analysis

$$(C - \lambda I)A = 0 \quad \text{Solve for } A.$$

This is an Eigenvalue / Eigenvector Problem.

This started by solving DEs:

$$\frac{dy}{dt} = \underset{\substack{\uparrow \\ n \times n}}{A} y \quad \} \text{ } n \text{ linear DEs}$$

Sol is of the type $y(t) = e^{\lambda t} X$

Principal Component Analysis

$$y(t) = e^{\lambda t} x$$

$$\Rightarrow \lambda e^{\lambda t} x = A e^{\lambda t} x$$

$$\Rightarrow \lambda x = A x$$

$$\underbrace{\left(\underbrace{A}_{n \times n} - \underbrace{\lambda I}_{n \times n} \right)}_{n \times n} \underbrace{x}_{n \times 1} = 0$$

$$\left\{ \frac{dy}{dt} = A y \right\} \begin{matrix} n \text{ linear} \\ \text{DEs} \end{matrix}$$

$$\underbrace{\left(\underbrace{C}_{P \times P} - \underbrace{\lambda I}_{P \times P} \right) \underbrace{A}_{P \times P}}_{P \times P} = 0$$

$$\left. \begin{array}{c} (A - \lambda I) X = 0 \\ \uparrow \\ n \times 1 \end{array} \right\} \begin{array}{l} \text{It gives eigenvalues of } A \\ \text{and eigenvectors of } X \end{array}$$

Ex- $\left. \begin{array}{l} y_1' = 5y_1 + y_2 \\ y_2' = 3y_1 + 3y_2 \end{array} \right\} \frac{dy}{dt} = \begin{bmatrix} 5 & 1 \\ 3 & 3 \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}$

$$\lambda_1 = 6 ; X_1 = \begin{bmatrix} 1 \\ 1 \end{bmatrix} ; \lambda_2 = 2 ; X_2 = \begin{bmatrix} 1 \\ -3 \end{bmatrix}$$

$$y(t) = \begin{bmatrix} y_1(t) \\ y_2(t) \end{bmatrix} = C_1 e^{\lambda_1 t} X_1 + C_2 e^{\lambda_2 t} X_2$$

$$y(t) = C_1 e^{6t} \begin{bmatrix} 1 \\ 1 \end{bmatrix} + C_2 e^{2t} \begin{bmatrix} 1 \\ -3 \end{bmatrix}$$

$\uparrow$ get from B.C. $\uparrow$

Case 1 If we know the eigenvector for
 $Ax = \lambda_1 x$, then we also know
the eigenvector for $A^2x = \lambda_2 x$

$$A^2x = \lambda_2 x \Rightarrow A(Ax) = \lambda_2 x$$

$$\Rightarrow A(\lambda_1 x) = \lambda_2 x$$

$$\Rightarrow \lambda_1^2 x = \lambda_2 x$$

$$\Rightarrow \lambda_2 = \lambda_1^2$$

and they have the same eigenvector

Case 2

$$\text{If } (A + cI)x = \lambda_3 x$$

$$\lambda_3 = (\lambda_1 + c)$$

Similarly $A^n x = \lambda_n x$

then $\lambda_n = (\lambda_1)^n$

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

$\uparrow$
 $P \times P$

This is a concatenation of eigenvectors

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

$\uparrow$
 $n \times 1$

$$A = [a_1 a_2 \dots a_p]$$

$\uparrow$
 $P \times P$

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

$$A = (a_1, a_2, \dots, a_p)$$

↑
↑
eigenvectors

$$\longleftrightarrow (\lambda_1, \lambda_2, \dots, \lambda_p)$$

eigenvalues.

$\Rightarrow$ Variance in X' corresponds to the eigenvalue of $\lambda_1, \lambda_2, \dots, \lambda_p$

$$CA = \lambda A \Rightarrow A^T C A = A^T \lambda A$$

$$\lambda = A^T C A \equiv \mathcal{V}_{X'}$$

$$\lambda_1 \geq \lambda_2 \geq \lambda_3 \dots \geq \lambda_p$$

Principal Component Analysis

Thus, if we want to capture 95% variance,
then we reduce the # of dimensions to q
such that

$$\frac{\sum_{i=1}^q \lambda_i}{\sum_{i=1}^p \lambda_i} \geq \frac{95}{100}$$

Principal Component Analysis

Eg.

$$p = 2, \quad n = 4$$

$$X = \{(1, 1), (2, 1), (2, 2), (3, 2)\}$$

$$\Rightarrow \bar{X} = \left(2, \frac{3}{2}\right); \quad C = \frac{1}{3} \begin{bmatrix} 2 & 1 \\ 1 & 1 \end{bmatrix}$$

The eigenvalues are

$$\lambda_1 = 0.8727, \quad a_1 = \begin{bmatrix} -0.85065 \\ -0.52571 \end{bmatrix}$$

$$\lambda_2 = 0.1273,$$

$$a_2 = \begin{bmatrix} 0.52571 \\ -0.85065 \end{bmatrix}$$

$$A = [a_1, a_2] = \begin{bmatrix} -0.85065 & 0.52571 \\ -0.52571 & -0.85065 \end{bmatrix}$$

Principal Component Analysis

$$X' = A^T (X - \bar{X}) \Rightarrow X'_1 = \begin{bmatrix} -0.85065 & -0.52571 \\ 0.52571 & -0.85065 \end{bmatrix} \begin{bmatrix} 1 & -2 \\ 1 & -1.5 \end{bmatrix}$$

$$\Rightarrow \underline{X'_1} = \begin{bmatrix} 1.1135 \\ -0.10039 \end{bmatrix}$$

$$\underline{X'_2} = \begin{bmatrix} 0.2628 \\ 0.42533 \end{bmatrix} ; \underline{X'_3} = \begin{bmatrix} -0.2628 \\ -0.42533 \end{bmatrix}$$

$$X'_4 = \begin{bmatrix} -1.1135 \\ 0.10039 \end{bmatrix}$$

$X \xrightarrow[A^T]{\text{Transformed}} X'$

Principal Component Analysis

$$X \xrightarrow{A^T} X' \quad ; \quad A^T = \begin{bmatrix} -0.85065 & -0.52571 \\ 0.52571 & -0.85065 \end{bmatrix}$$

$$\{(1,1), (2,1), (2,2), (3,2)\}$$

$$\text{mean } X = (2, 1.5)$$

Covariance of X

$$= \begin{bmatrix} \frac{2}{3} & \frac{1}{3} \\ \frac{1}{3} & \frac{1}{3} \end{bmatrix}$$

$$X' = \{(1.1135, -0.10039), (0.26286, 0.42533), (-0.26286, -0.42533), (-1.1135, -0.10039)\}$$

$$\text{Mean } X' = (0, 0)$$

$$\text{Cov } X' = \begin{bmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{bmatrix} = \begin{bmatrix} 0.87266 & 0 \\ 0 & 0.12732 \end{bmatrix}$$

Principal Component Analysis

$$X \xrightarrow{A^T} X'$$

$$\text{Cov } X' = \begin{bmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_p \end{bmatrix}$$

$$\begin{array}{ccc} X & \xleftarrow{A} & X' \\ X_1, X_2 & \xrightarrow{A^T} & X'_1, X'_2 \end{array}$$

Since X'_1 contains 87% of variance, we can drop X'_2

Principal Component Analysis

Let us drop X_2'

Thus we do not take

$$A = [a_1, a_2]$$

but instead $A' = [a_1, 0]$

$$X \xrightarrow{(A')^T} X'' \xrightarrow{A'} \sim X$$

$$(A')^T = \begin{bmatrix} -0.85065 & -0.52571 \\ 0 & 0 \end{bmatrix}$$

$$X = \begin{cases} (1,1) \\ (2,1) \\ (2,2) \\ (3,2) \end{cases}$$

$$X'' = \begin{cases} (1.11351, 0) \\ (0.26286, 0) \\ (-0.26286, 0) \\ (-1.11351, 0) \end{cases}$$

$$A' \rightarrow \begin{cases} (1.0528, 0.91462) \\ (1.7764, 1.36181) \\ (2.2236, 1.63819) \\ (2.9472, 2.08538) \end{cases}$$

✗

Principal Component Analysis

$$X = \{(1, 1), (2, 1), (2, 2), (3, 2)\}$$

$$\bar{x} = \frac{1}{2} \cdot (4, 3)$$

$$C = \frac{1}{3} \cdot \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$$

$$\lambda_1 = 0.8727$$

$$\lambda_2 = 0.1273$$

$$v_1 = \begin{pmatrix} -0.85065 \\ -0.52573 \end{pmatrix}$$

$$v_2 = \begin{pmatrix} 0.52573 \\ -0.85065 \end{pmatrix}$$

$$E = v_1 = \begin{pmatrix} -0.85065 \\ -0.52573 \end{pmatrix}$$

The projected data

$$Y = \{1.1135, 0.2629, -0.2629, -1.1135\}$$

Inverse PCA yields

$$X' = \{ (1.0528, 0.91459), (1.7764, 1.3618), (2.2236, 1.6382), (2.9472, 2.0854) \} \neq X$$

