

(8)

Properties of the LS estimators

$$V(\hat{\beta}) = \sigma^2 (X'X)^{-1} = \sigma^2 \underbrace{\sum d_k^{-2} \tilde{u}_k \tilde{u}_k'}_{\text{spectral decomposition}}$$

$$\begin{aligned} E\|\hat{\beta} - \beta\|^2 &= E\left\{(\hat{\beta} - \beta)'(\hat{\beta} - \beta)\right\} = \text{Tr}\{V(\hat{\beta})\} = \sigma^2 \text{Tr}\{(X'X)^{-1}\} \\ &= \sigma^2 \sum d_k^{-2} \quad (*) \end{aligned}$$

So, total MSE of slope estimates depend on eigenvalues $d_1^2, \dots, d_p^2$ of $(X'X)$.

Now, if d_k^2 is small, $(*)$ can get very large!

- That's what's meant by an unstable regression fit — high variance.
- On the other hand, $E(\hat{\beta} - \beta) = 0$, so LS estimators are unbiased.

So, can't have both — trade off of some bias to reduce the variance

⇒ Regularized fits

(9)

How can we achieve this?

Stabilize the estimate — the matrix inversion

$$\hat{\beta}_R = \underbrace{(X'X + rI)}_{G_r}^{-1} X'y$$

$$\bullet E(\hat{\beta}_R) = G_r X'X\beta = \beta - \underbrace{r G_r \beta}_{\text{bias}}$$

$$\bullet V(\hat{\beta}_R) = \sigma^2 G_r (X'X) G_r$$

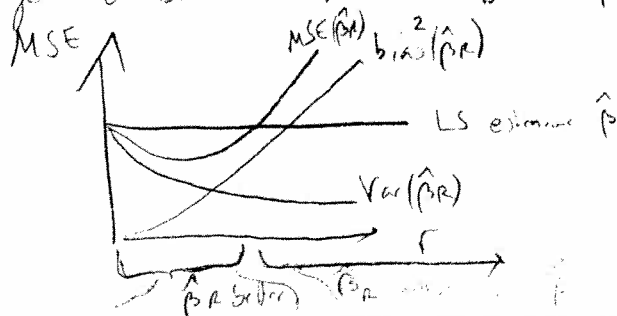
$$\bullet \text{MSE} = \text{Var} + \text{bias}^2 \\ = G_r (\sigma^2 X'X + r^2 \beta\beta') G_r$$

$$\bullet E\|\hat{\beta}_R - \beta\|^2 = \text{Tr}\{\text{MSE}\} = \sum \frac{\sigma^2 d_u^2 + r^2 \beta_u^2}{(d_u^2 + r)^2} \quad \left(r=0 \quad \sum \sigma^2 d_u^2 \right)$$

Now, for some values of r

$$\sum \frac{\sigma^2 d_u^2 + r^2 \beta_u^2}{(d_u^2 + r)^2} < \sum \sigma^2 d_u^{-2}$$

so we get a smaller MSE with $\hat{\beta}_R$ than $\hat{\beta}$



Alternative view : shrinkage estimator

$$\hat{\beta}_S = \frac{\hat{\beta}}{1+c}, \quad c \geq 0$$

$$\Rightarrow \text{bias } \frac{c\beta}{1+c}, \quad V(\hat{\beta}_S) = \sigma^2 (1+c)^{-2} (X'X)^{-1}$$

$$MSE = (1+c)^{-2} (\sigma^2 (X'X)^{-1} + c^2 \beta \beta')$$

• As before - a tradeoff between bias & variance.

• For some c , $MSE(\hat{\beta}_S) < MSE(\hat{\beta})$

• If $(X'X)$ diagonal - same thing as $\hat{\beta}_R = (X'X + rI)^{-1} X'y$
- a shrinkage factor applied to all $\hat{\beta}_i$

We can arrive at the $\hat{\beta}_R$ estimates through another argument

• LS: find β to minimize $(y - X\beta)'(y - X\beta)$

• Penalized LS: find β to minimize $(y - X\beta)'(y - X\beta)$
where $\|\beta\|^2 \leq T$

i.e. we place a constraint on the total magnitude of the coefficients.

(Ridge Regression interpretation)
 $p(\beta) \propto \exp(-\lambda \|\beta\|^2)$

Well - so we want to minimize

$$(y - X\beta)'(y - X\beta)$$

but we also want to make sure

$$\beta'\beta = \|\beta\|^2 = \sum \beta_j^2 \leq T$$

These types of constrained optimization problems are solved through Lagrange methods

$$\textcircled{*} \text{ Minimize } (y - X\beta)'(y - X\beta) + r\beta'\beta$$

where r is such that $\beta'\beta \leq T$

(on trying different values for r until $\beta'\beta \leq T$ (line search))

$$\textcircled{*} \Rightarrow \frac{\partial \mathcal{L}}{\partial \beta} = -2X'(y - X\beta) + 2r\beta = 0$$

$$\Rightarrow (X'X + rI)\beta = X'y$$

$$\Rightarrow \hat{\beta}_r = (X'X + rI)^{-1} X'y$$

This is also called

Ridge regression

(18)

We can also use other penalties on β
 - maybe not constrain sum of squares $\sum \beta_j^2$
 but only sum of absolute values $\sum |\beta_j|$

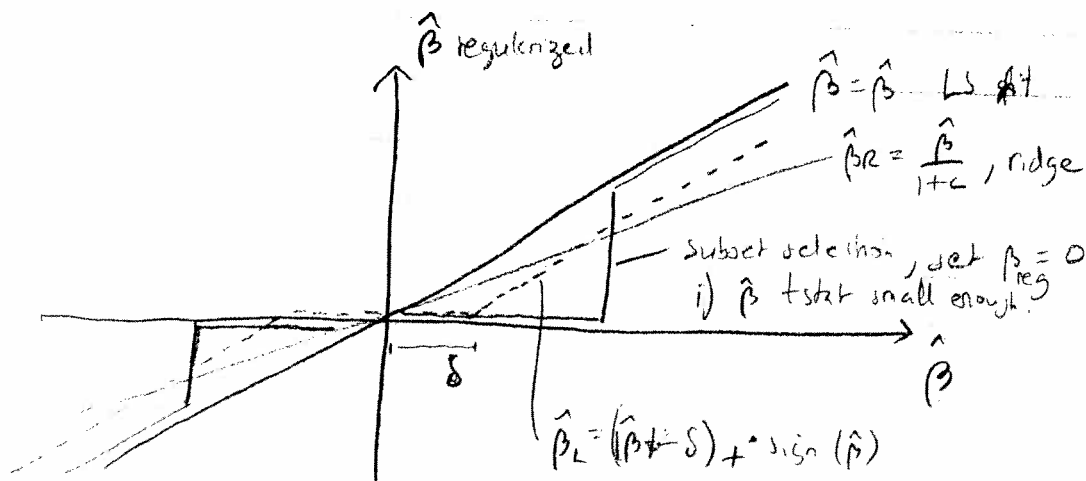
This is called the LASSO

$$\min (y - X\beta)'(y - X\beta) + r \frac{\beta' \beta}{\|\beta\|}$$

• If $(X'X)$ diagonal $\hat{\beta}_{\text{lasso},j} = \text{sign}(\hat{\beta}_j) (|\hat{\beta}_j| - \delta)_+$

where δ s.t. $\sum |\hat{\beta}_{\text{lasso},j}| \leq T$

(compare, $\hat{\beta}_R = \frac{\hat{\beta}}{1+c}$, where c.s.t. $\sum \hat{\beta}_R^2 \leq T$)



• $\hat{\beta}_L$ also called soft shrinkage

Why is it called regularized fit?

LS: $H = X(X'X)^{-1}X'$, $X = UDV'$

$$\hat{y} = X\hat{\beta} = X(X'X)^{-1}X'y = UU'y = \left(\sum_{k=1}^p u_k u_k'\right) y$$

"Complexity" of projection (regression) H

- Look at its trace

$$\text{Tr}\{H\} = \text{Tr}\{X(X'X)^{-1}X'\} = \text{Tr}\{(X'X)(X'X)^{-1}\} \\ = \text{Tr}\{I_p\} = p$$

- p is the degrees of freedom of the LS regression line fit

Ridge: $H_R = X(X'X + rI)^{-1}X'$

$$\begin{aligned} \hat{y}_R &= X\hat{\beta}_R = X(X'X + rI)^{-1}X'y = X(VD^2V' + rI)^{-1}X'y \\ &= X(V(D^2 + rI)V')^{-1}X'y \\ &= UDV'V(D^2 + rI)^{-1}V'V'DU'y \\ &= UD(D^2 + rI)^{-1}DU'y \\ &= \sum_{k=1}^p \left(\frac{d_k^2}{d_k^2 + r}\right) u_k u_k' y \end{aligned}$$

$$\text{Tr}\{H_R\} = \sum_{k=1}^p \frac{d_k^2}{d_k^2 + r} < \text{Tr}\{H\} \quad \text{if } r > 0$$

- so flexibility (df) of ridge model < flexibility of LS fit $\Rightarrow$ regularized fit

* So ridge regression — computes coefficient in coordinate system defined by eigenvectors, but then also shrinks them by a factor $\frac{d_i^2}{d_i^2 + r}$.

* The shrinkage will have most effect in the direction of small variance d_i^2 — alt to ? (reg!!)
 when we drop small eigenvalue vectors

Recap

- ⑥ Collinearity in the data (correlated X 's) can be problematic → either choose a subset of X 's not highly correlated (X)
 → use a regularized fitting technique

* Reminder; other things to consider

- omitted variables
- transformations ($\sqrt{\cdot}$, $\log$, $\frac{1}{\cdot}$, ...)

⑦ Subset selection? — either by knowledge/understanding of the data

selection methods, model search

