

100x

9

More about fitted values;

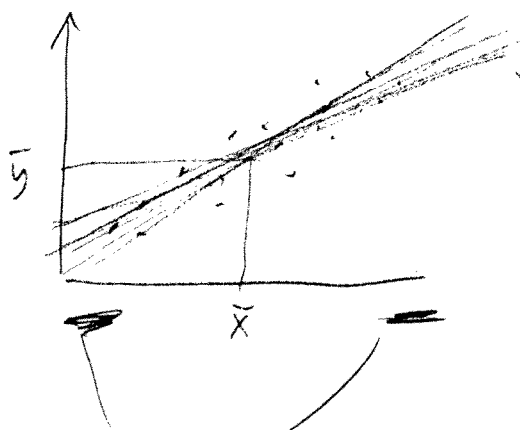
$$V(\hat{y}_i) = \sigma^2 h_{ii}$$

$$V(\hat{y}_i) = V(\hat{\beta}_0 + \hat{\beta}_1 x_i) = V(\bar{y} + \hat{\beta}_1 (x_i - \bar{x})) = V(\bar{y}) + (x_i - \bar{x})^2 V(\hat{\beta}_1)$$

$$\left(\begin{aligned} \text{Since } \text{Cov}(\bar{y}, \hat{\beta}_1) &= \text{Cov}\left(\sum_{j=1}^n \frac{1}{n} y_j, \sum_{j=1}^n k_j y_j\right) \\ &= \sum_j \frac{1}{n} k_j V(y_j) = \frac{\sigma^2}{n} \sum k_j = 0 \end{aligned} \right)$$

$$= \frac{\sigma^2}{n} + \sigma^2 \frac{(x_i - \bar{x})^2}{\sum_j (x_j - \bar{x})^2} = \sigma^2 h_{ii}$$

Meaning; At locations of high leverage, the regression line can fluctuate a lot from sample to sample



h_{ii} large here

What about the residuals e_i ?

$$V(e_i) = \sigma^2(1-h_{ii})$$

$$\left[\begin{aligned} e_i &= y_i - \hat{y}_i = y_i - \sum h_{ij} y_j \\ \text{Cov}(e_i, e_i) &= \text{Cov}(y_i - \hat{y}_i, y_i - \hat{y}_i) = \text{Cov}(y_i, y_i) + \text{Cov}(\hat{y}_i, \hat{y}_i) - 2\text{Cov}(y_i, \hat{y}_i) \\ &= \sigma^2 + \sigma^2 h_{ii} - 2\sigma^2 h_{ii} = \sigma^2(1-h_{ii}) \end{aligned} \right]$$

$\swarrow$
 $\perp \text{Cov}(e_i, \hat{y}_i) = 0$
 LS properties

$$\begin{aligned} \Rightarrow \text{Cov}(e_i, e_j) &= \text{Cov}(y_i - \hat{y}_i, y_j - \hat{y}_j) = \text{Cov}(y_i, y_j) + \text{Cov}(\hat{y}_i, \hat{y}_j) - \text{Cov}(y_i, \hat{y}_j) \\ &\quad - \text{Cov}(y_j, \hat{y}_i) \\ &= \text{Cov}\left(\sum h_{ik} y_k, \sum h_{jl} y_l\right) - \text{Cov}\left(y_i, \sum h_{jl} y_l\right) - \text{Cov}\left(y_j, \sum h_{ik} y_k\right) \\ &= \sum h_{ik} h_{jl} \sigma^2 - h_{ji} \sigma^2 - h_{ij} \sigma^2 = \sigma^2 h_{ij} - 2\sigma^2 h_{ij} = -\sigma^2 h_{ij} \end{aligned}$$

MEANING :

- ① Residuals are correlated (but errors ε_i are not)
- ② residuals have nonconstant variance $V(e_i) = \sigma^2(1-h_{ii})$ (error $V(\varepsilon_i) = \sigma^2$)

③ Residual variance is small in high leverage regions, since observations here drive the fit

✱

(11)

Now: $y_i = \beta_0 + \beta_1 x_i + \varepsilon_i \Rightarrow E(y_i) = \beta_0 + \beta_1 x_i$ since $E(\varepsilon_i) = 0$
(assumption 2)

$E(\hat{\beta}_1) = E(\sum k_i y_i) = \sum k_i E(y_i)$ (by uncorrelated ε_i , assumption 3)

$= \sum k_i (\beta_0 + \beta_1 x_i) = \beta_0 \sum k_i + \beta_1 \frac{\sum (x_i - \bar{x}) x_i}{\sum (x_i - \bar{x})^2} = \beta_1 \frac{\sum (x_i - \bar{x})(x_i - \bar{x})}{\sum (x_i - \bar{x})^2} = \beta_1$ ✓

//
0 since
 $\sum k_i = \frac{\sum (x_i - \bar{x})}{\sum (x_i - \bar{x})^2}$

$= \frac{\sum x_i - n\bar{x}}{\sum (x_i - \bar{x})^2} = 0$

(and similarly for β_0)

Meaning, on average LS
give us the correct β_0, β_1 .

Facts

- residuals sum to 0, $\sum \varepsilon_i = 0$ (i) β_0 estimated) to use

Follows from 1st normal equation

- $\sum x_i \varepsilon_i = 0$, "Orthogonal projection"

Meaning, we used up all the (linear) information
that x has on y with our LS regression fit.

Follows from 2nd normal equation (NE).

- $\sum \hat{y}_i \varepsilon_i = 0$, fitted values $\perp$ to residuals

"Decomposition"

$y = \hat{y} + \varepsilon$

explainable
..

orthogonal
to x .

Let's look closer at the fitted values $\hat{y}$.

$\hat{y}_i$ (fitted value for the i th observation y_i)

$$= \hat{\beta}_0 + \hat{\beta}_1 x_i = \bar{y} - \hat{\beta}_1 \bar{x} + \hat{\beta}_1 x_i = \bar{y} + \hat{\beta}_1 (x_i - \bar{x})$$

$$= \bar{y} + \left(\sum_j k_j y_j \right) (x_i - \bar{x}) = \sum_j \left(\frac{1}{n} y_j + k_j (x_i - \bar{x}) y_j \right)$$

$$= \sum_j \underbrace{\left(\frac{1}{n} + k_j (x_i - \bar{x}) \right)}_{h_{ij}} y_j = \sum_j h_{ij} y_j$$

That is, $\hat{y}_i$ is a weighted average of all y_j 's

(a linear model, linear in the data y)

Question: how much does y_i impact the overall fit @ x_i ?

Weight for y_i in $\hat{y}_i$ is $\boxed{h_{ii}}$, also called the leverage

$$h_{ii} = \left(\frac{1}{n} + k_i (x_i - \bar{x}) \right) = \left(\frac{1}{n} + \frac{(x_i - \bar{x})^2}{\sum_j (x_j - \bar{x})^2} \right)$$

Note • $h_{ii} \sim O\left(\frac{1}{n}\right)$

- if x_i "extreme", far from $\bar{x}$, h_{ii} is large
- if x_i close to $\bar{x}$, h_{ii} is small.

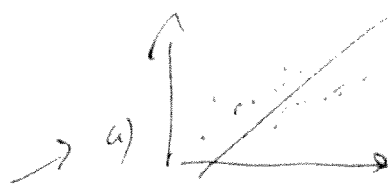
(1) $h_{ii} \gg h_{jj}$ $\hat{y}_i$ is dominated by ... 1 11

13

What does this mean?



These observations drive the fit of the regression line!



In both a) & b) observation i has high leverage h_{ii} , and could potentially influence the fit. In case a), the i th observation is an ~~influencer~~ influential outlier, producing a poor fit to the rest of the data. In case b), the slope estimate is not much affected, but we get an exaggerated sense of the goodness of fit of the model. (more later).

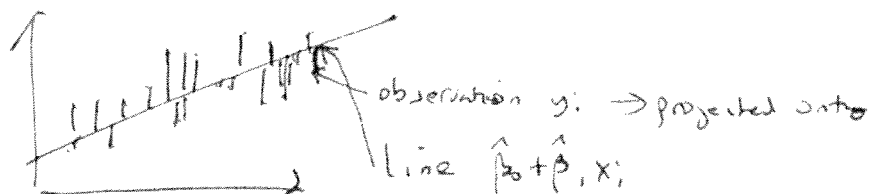
We collect the h_{ii}, h_{ij} in a matrix called H , or the hat-matrix

$$\hat{\mathbf{y}} = \mathbf{H} \mathbf{y} : \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_n \end{pmatrix} = \begin{pmatrix} h_{11} & h_{12} & h_{13} & \dots & h_{1n} \\ h_{21} & h_{22} & h_{23} & \dots & h_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ h_{n1} & \dots & \dots & \dots & h_{nn} \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$$

$n \times 1$ $n \times n$ $n \times 1$

(14)

- This H matrix is a projection matrix.
- It takes the data y , and projects it onto a function of $\bar{X}$
e.g. simple linear regression



- The projection is orthogonal to $X \underline{h}_x$

- Note, it can be precomputed using only X -values.

More LS properties

- Variance $(\hat{\beta}_1) = V(\hat{\beta}_1) = V(\sum k_i y_i) = \sum k_i^2 V(y_i)$ (uncorrelated ϵ_i)
 $= \sum k_i^2 \sigma^2 = \frac{\sum (x_i - \bar{x})^2}{(\sum_j (x_j - \bar{x})^2)^2} \sigma^2$
 $= \frac{\sigma^2}{\sum_j (x_j - \bar{x})^2}$

($V(y_i) = \sigma^2$ since the only thing random in y_i is ϵ_i)

Note

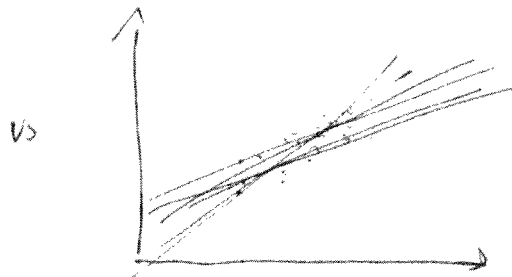
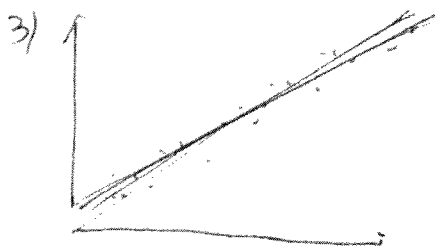
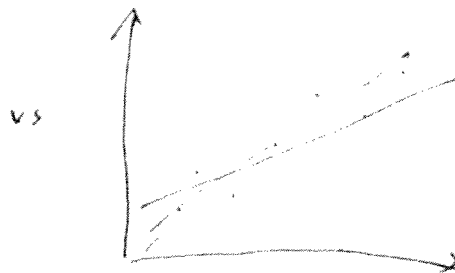
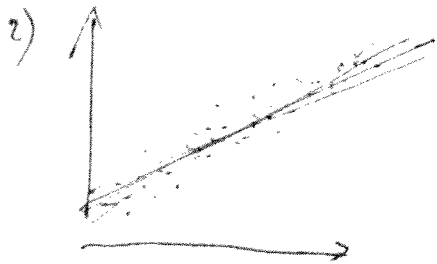
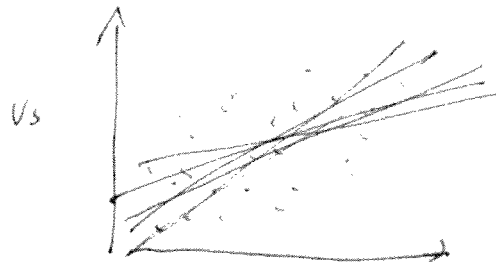
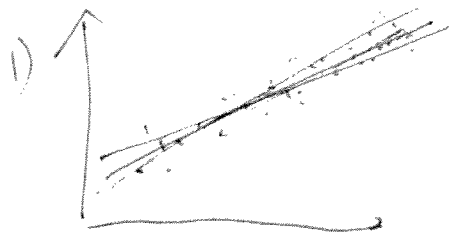
- LS is the minimum variance unbiased linear estimate, any other weights $c_i \neq k_i \Rightarrow$ larger variance.

- 3 factors impact the variance of $\hat{\beta}_1$.

1) Random scatter variance σ^2
(more noise $\rightarrow$ less precision in $\hat{\beta}_1$ estimate)

2) Sample size n

3) spread of x



- Variance of fitted values, variance of the regression line

$$V(\hat{y}_i) = \sigma^2 h_{ii}$$

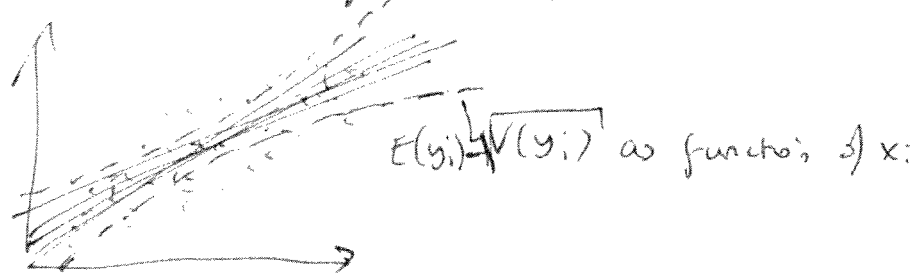
$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i \Rightarrow V(\hat{y}_i) = V(\hat{\beta}_0 + \hat{\beta}_1 x_i) = V(\bar{y} + \hat{\beta}_1 (x_i - \bar{x}))$$

$$\left(\bar{y} \perp \hat{\beta}_1 \text{ since } \text{cov}(\bar{y}, \hat{\beta}_1) = \text{cov}\left(\frac{1}{n} \sum y_i, \sum k_j y_j\right) = \sum \frac{1}{n} k_i \text{cov}(y_i, y_i) \right. \\ \left. = \frac{\sigma^2}{n} \sum k_i = 0 \right) \quad (\text{by uncorrelated } \varepsilon_i)$$

$$\Rightarrow V(\bar{y} + \hat{\beta}_1 (x_i - \bar{x})) = V(\bar{y}) + (x_i - \bar{x})^2 V(\hat{\beta}_1) = \frac{\sigma^2}{n} + \frac{\sigma^2 (x_i - \bar{x})^2}{\sum (x_j - \bar{x})^2} \\ = \sigma^2 h_{ii}$$

(16)

This makes sense: A small change for extreme x_i observations will change the slope. However, since the line always goes through the center of the data $(\bar{x}, \bar{y})$, the observations in the center will have stable fitted values



- $V(e_i) = \sigma^2(1 - h_{ii})$

residual $e_i = y_i - \hat{y}_i = y_i - \sum h_{ij} y_j$

Since $E(e_i) = 0$ $V(e_i) = E(e_i^2)$ or $\text{cov}(e_i, e_i)$

$$\text{cov}(e_i, e_i) = \text{cov}(y_i - \hat{y}_i, y_i - \hat{y}_i) = \text{cov}(y_i, y_i) + \text{cov}(\hat{y}_i, \hat{y}_i) - 2\text{cov}(y_i, \hat{y}_i)$$

$$= \sigma^2 + \sigma^2 h_{ii} - 2\text{cov}(y_i, \hat{y}_i)$$

/

$\hat{y}_i, e_i$

$$= \sigma^2 - \sigma^2 h_{ii} - 2\text{cov}(e_i, \hat{y}_i)$$

$$= \sigma^2(1 - h_{ii})$$

= 0 by orthogonal projection

Meaning: residual variance is small for high leverage observations.

That makes sense: for high leverage observations, the regression line is pulled towards the observation, so errors are forced to be small.

Note residuals are not directly comparable since their variances differ

$\Rightarrow$ better to look @ and use

Standardized residuals

$$\tilde{e}_i = \frac{e_i}{\sqrt{1-h_{ii}}}$$

Note $V(\tilde{e}_i) = \sigma^2$

So far $\left\{ \begin{array}{l} E(\hat{\beta}_i) = \beta_i, \quad V(\hat{\beta}_i) = \frac{\sigma^2}{\sum (x_i - \bar{x})^2} \\ E(\hat{y}_i) = \beta_0 + \beta_1 x_i, \quad V(\hat{y}_i) = \sigma^2 h_{ii} \\ E(e_i) = 0, \quad V(e_i) = \sigma^2 (1-h_{ii}) \end{array} \right.$

Note, all the variances involve one extra parameter

$$\sigma^2$$

If we don't know it, we can't make inferences

$\Rightarrow$ use the data to estimate this "nuisance" parameter as well.

• RSS = residual sum of squares (or SSE error sum of squares)

$$= \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

• MSE = mean squared error $= \frac{RSS}{n-2} = \hat{\sigma}^2$

is an unbiased estimate of the error variance σ^2

Note denominator $n-2$; we've already used the data for two parameter estimates ($\hat{\beta}_0$ & $\hat{\beta}_1$) so only $n-2$ degrees of freedom remains

Next lecture $\Rightarrow R^2$, and the F & t-tests.

Diagnostics

①

$$y = X\beta + \varepsilon \Rightarrow \text{fit to data} \Rightarrow \hat{\beta} \Rightarrow \hat{y} = X\hat{\beta}$$
$$e = y - \hat{y}$$

Facts

- $V(e_i) = \sigma^2(1 - h_{ii})$ - residuals have non-constant ~~error~~ variance

$\Rightarrow$ standardized residuals

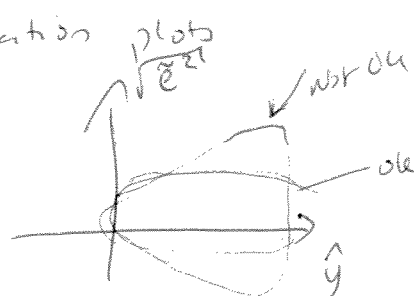
$$\tilde{e}_i = \frac{e_i}{\sqrt{1 - h_{ii}}}, \quad V(\tilde{e}_i) = \sigma^2 \text{ constant}$$

• Plot (standardized) residuals against

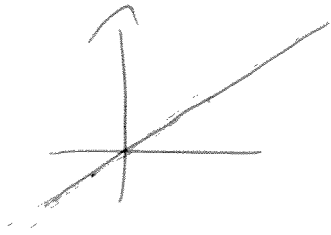
- fitted values
- leverage
- variables included & omitted

} checking basic assumptions
(sufficiency of model, constant error var, no outliers, symm. of errors)

- Scale-location plots



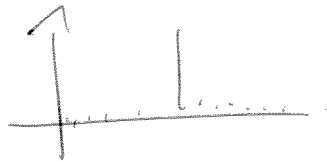
- Transformations
- drop observations
- ...

Σ extras• QQ_{plot}

(check if $e \sim$ Normally distributed)
 → needed for some inference

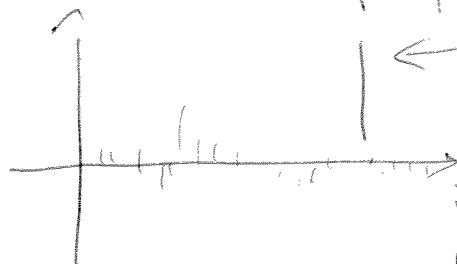
• Contribution by single observation

- leverage



"potential to do harm"

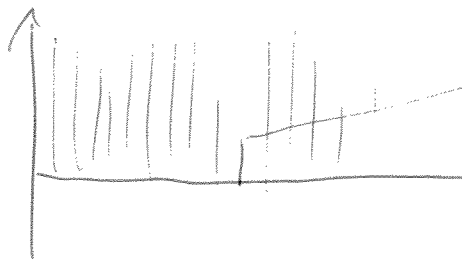
- change in slope

- refit w/o obs i and plot $\hat{\beta} - \hat{\beta}(i)$ vs i 

← changed conclusion/result a lot

index of observation

- change in MSE

- refit w/o obs i and plot $\hat{\sigma}^2(i)$ vs i 

MSE dropped a lot here
 - outlier

Combo of all — Cook's D

◦ Studentized residual — residual for obs i when i is not included in the fit

$$e_i = y_i - \hat{y}_i$$

$$\hat{y}_i = \sum h_{ij} y_j = h_{ii} y_i + \sum_{j \neq i} h_{ij} y_j$$

~~$$\hat{y}_i = \sum h_{ij} y_j = h_{ii} y_i + \sum_{j \neq i} h_{ij} y_j$$~~

$$\hat{y}_{(i)} = \frac{1}{1-h_{ii}} \sum_{j \neq i} h_{ij} y_j \Rightarrow \hat{y}_{(i)} = \sum_{j \neq i} h_{ij} y_j + h_{ii} \hat{y}_{(i)}$$

$(y_i - h_{ii})^2 = (y_i - h_{ii})^2$ ←

$$\Downarrow y_i - \hat{y}_{(i)} = (y_i - \underbrace{\sum_{j \neq i} h_{ij} y_j}_{\hat{y}_{(i)}} + h_{ii} (y_i - \hat{y}_{(i)}))$$

$$\Downarrow \boxed{y_i - \hat{y}_{(i)} = \frac{y_i - \hat{y}_i}{1-h_{ii}}} = e_i^* \text{ or } e(i)$$

$$\Rightarrow \text{Cook's D} = D_i = \frac{e(i)^2}{p \hat{\sigma}^2}, \quad \hat{\sigma}^2 = \frac{\sum e_i^2}{n - df(\text{model})}$$

