

Handout 2

①

RECAP

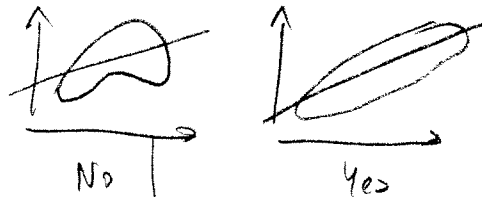
Summary $y \sim x$ relationship $y_i = \beta_0 + \beta_1 x_i + \epsilon_i$, ϵ_i uncorrelated, $E(\epsilon_i) = 0$

LS criterion $\sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_i))^2$ and $V(\epsilon_i) = \sigma^2$

equal
contribution of
all terms

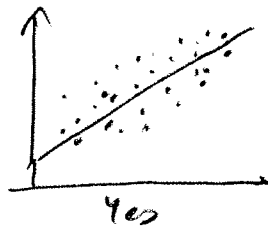
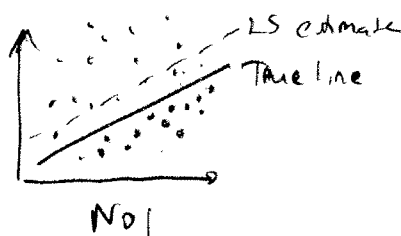
squared error

BASIC ASSUMPTIONS ① Model sufficiency



Check? Plot, plot, plot!

② Symmetry of error scatter



Solution? Transformation may help
→ now use a different fitting
criterion

Check? • QQplot of residuals
from LS fit.

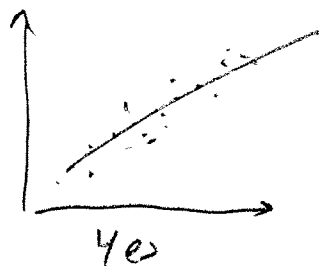
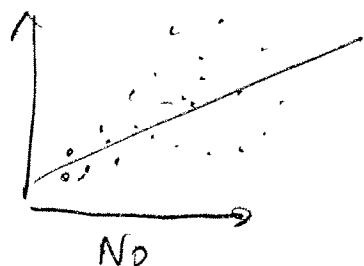
• Scatter plot y on x
may already show a
clear asymmetry of course.

③ Uncorrelated errors

Check? Think about the sampling process.

- Is the data sampled sequentially over time? ⇒ Time Series
- Or samples nested? (e.g. students from the same class)
Solution ⇒ Mixed effects or patients from the same hospital
Hierarchical Linear Model

④ Constant variance of error scatter



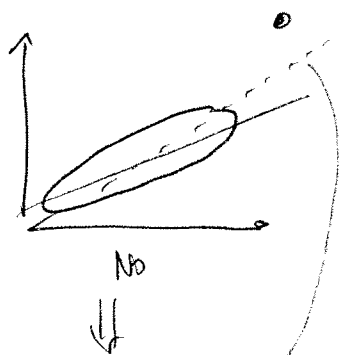
Check? Plot, plot, plot ; both y on x and residuals

Solution $\Rightarrow$ Try some data transformation $\log(y), \sqrt{y}, \sqrt[3]{y}, \frac{1}{y}, \dots$

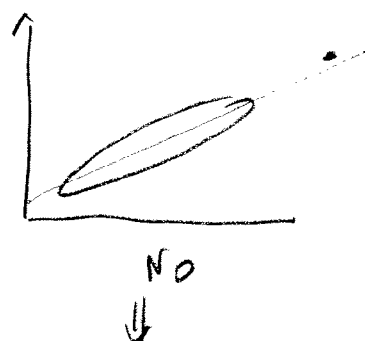
$\Rightarrow$ if that doesn't work $\Rightarrow$ Weighted Least Squares

$$\sum_{i=1}^n w_i (y_i - (\beta_0 + \beta_1 x_i))^2$$

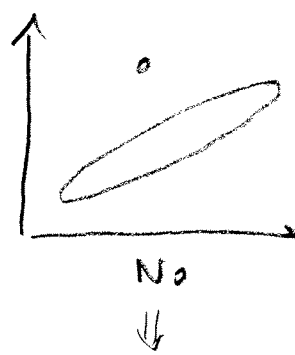
⑤ No outliers



Effect: bad fit



Effect: fit looks "too good"



Effect: you will overestimate the error scatter (more on this later)

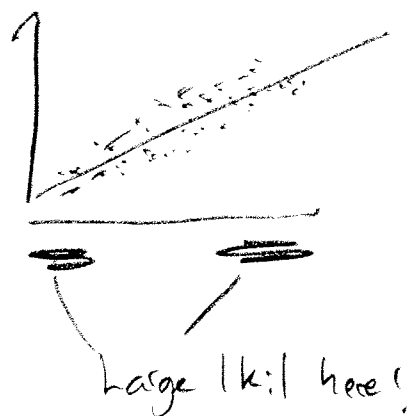
LS estimates $\rightarrow$ Properties

3

Have $\hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$, $\hat{\beta}_1 = \frac{\sum_i (x_i - \bar{x})(y_i - \bar{y})}{\sum_j (x_j - \bar{x})^2} = \sum_i h_i y_i$

where $h_i = \frac{(x_i - \bar{x})}{\sum_j (x_j - \bar{x})^2}$

Which observations contribute? The ones w. x_i far from $\bar{x}$!



Note, LS estimates are unbiased

$$E(\hat{\beta}_1) = E\left(\sum_i h_i y_i\right) = \sum_i h_i E(y_i) = \sum_i h_i (\beta_0 + \beta_1 x_i) \quad \left[\text{Since } E(y_i) = \beta_0 + \beta_1 x_i + E(\varepsilon_i) \right]$$

$$= \underbrace{\left(\sum_i h_i\right)}_{=0} \cdot \beta_0 + \beta_1 \underbrace{\left(\sum_i h_i x_i\right)}_{=1} = \beta_1 \quad \checkmark$$

$$\sum_i h_i = \frac{\sum_i (x_i - \bar{x})}{\sum_j (x_j - \bar{x})^2} = 0$$

$$\sum_i h_i x_i = \frac{\sum_i (x_i - \bar{x})^2}{\sum_j (x_j - \bar{x})^2} = 1$$

(Similarly $E(\hat{\beta}_0) = \beta_0$)

So, if we obtain many new data sets (x, y) , on average we'll obtain the true regression line.

What about the variance?

(4)

$$V(\hat{\beta}_1) = V(\sum k_i y_i) = \sum k_i^2 V(y_i) = \sum k_i^2 \sigma^2$$

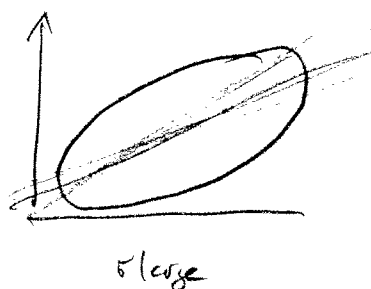
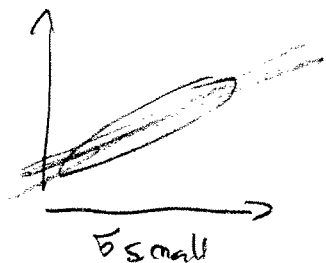
since $V(y_i) = V(\varepsilon_i) = \sigma^2$

↑
since y_i 's
are independent
(because ε_i 's are)

$$= \frac{\sum_i \frac{(x_i - \bar{x})^2}{(\sum_j (x_j - \bar{x})^2)^2} \sigma^2}{\sum_j (x_j - \bar{x})^2} = \frac{\sigma^2}{\sum_j (x_j - \bar{x})^2}$$

Meaning? ① $V(\hat{\beta}_1)$ increases with the noise level

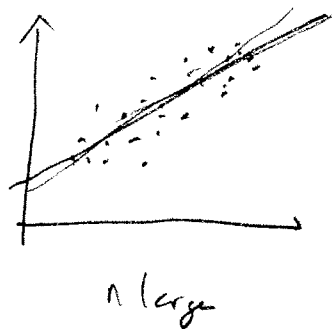
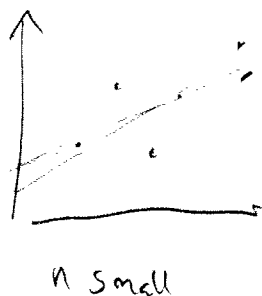
→ more uncertainty, more difficult to identify the slope



② $V(\hat{\beta}_1)$ decreases $\sim \frac{1}{n}$ w. the sample size n

$\sum_{i=1}^n (x_i - \bar{x})^2 = \text{sum of } n \text{ terms} \approx n \cdot c$, c some constant

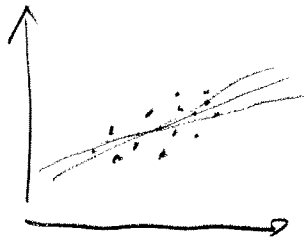
→ More data, easier to identify the slope.



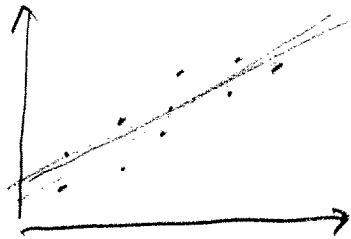
③ $V(\hat{\beta}_1)$ decreases w. the variance of x

⑤

$\Rightarrow$ The more spread in x we have, the easier it will be to identify the y - x relationship



$\sum (x_i - \bar{x})^2$ small



$\sum (x_i - \bar{x})^2$ large

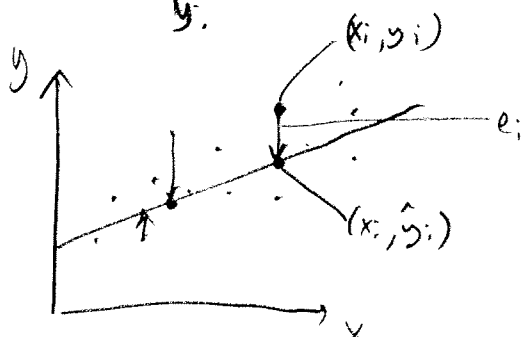
More properties

• Residuals sum to 0, $\sum e_i = 0$ (if β_0 included in the model)

• $\sum x_i e_i = 0$, "orthogonal projection"

Meaning $\Rightarrow$ We used up all the linear information in x about y w. our LS fit

$$\left[e_i = y_i - \bar{y} - \underbrace{\left(\frac{\sum (x_j - \bar{x}) y_j}{\sum (x_k - \bar{x})^2} \right)}_{\hat{y}_i} (x_i - \bar{x}) \right], \quad \sum x_i e_i = \sum x_i y_i - n \bar{x} \bar{y} - \underbrace{\frac{\sum (x_j - \bar{x}) y_j}{\sum (x_k - \bar{x})^2}}_{\hat{y}_i} \sum x_i (x_i - \bar{x})^2 = 0$$



⊗ projection $\downarrow$ is orthogonal ($\perp$) to the x -axis

More properties

(6)

$\sum \hat{y}_i e_i = 0$, fitted values $\perp$ the residuals

$\hat{y}_i$ = linear function of x_i , the part that's explainable from x

e_i = orthogonal to x_i , the part of y_i that's not explainable from x

Side note $V(\hat{\beta}_0) = \sigma^2 \left(\frac{1}{n} + \frac{\bar{x}^2}{\sum (x_i - \bar{x})^2} \right)$ Increases if $\bar{x}^2$ large compared w. spread $\sum (x_i - \bar{x})^2$
→ meaning if x near constant

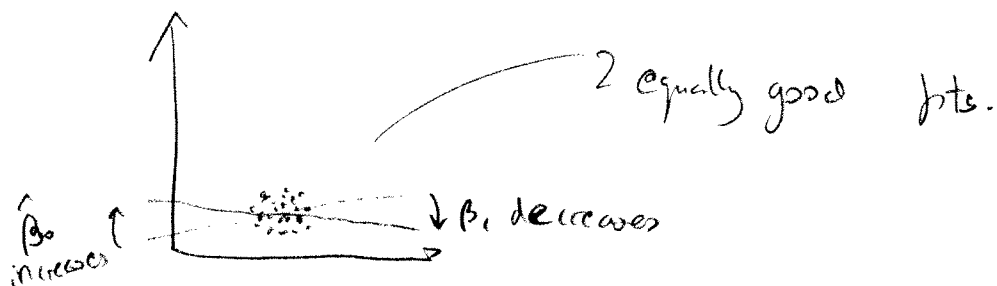
Also $\text{Cor}(\hat{\beta}_0, \hat{\beta}_1) = \frac{-\bar{x} \sigma^2}{\sum (x_i - \bar{x})^2}$, large negative if x near constant

So - i) x near constant, the model

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i \approx \mu + \varepsilon_i$$

and we don't need 2 parameters β_0, β_1 .

Equally good fits are obtained w. lots of regression line, meaning there is a lot of uncertainty in the estimates of β_0, β_1 .



What about the fitted values?

$$\begin{aligned}\hat{y}_i &= \hat{\beta}_0 + \hat{\beta}_1 x_i = \bar{y} + \hat{\beta}_1 (x_i - \bar{x}) = \bar{y} + \left(\sum_j h_{ij} y_j \right) (x_i - \bar{x}) \\ &= \sum_j \left(\frac{1}{n} y_j + h_{ij} (x_i - \bar{x}) y_j \right) = \sum_j \left(\frac{1}{n} + h_{ij} (x_i - \bar{x}) \right) y_j = \sum_j h_{ij} y_j\end{aligned}$$

= a weighted average of y -values.

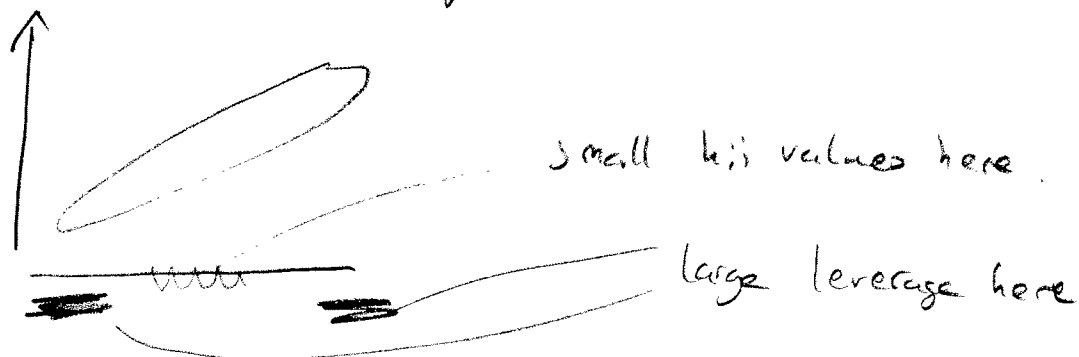
Note, if x is constant $h_{ij} = \frac{1}{n}$ since $x_i = \bar{x} \forall i$
 $\Rightarrow \hat{y}_i = \bar{y}$ if x has no information about y

Observation i contributes h_{ii} to its own fitted value

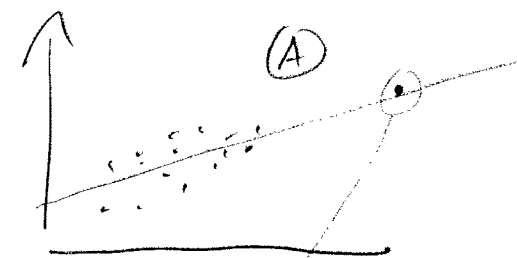
$\rightarrow$ Extreme $h_{ii} = 1$ $h_{ij} = 0 \forall j \neq i \Rightarrow \hat{y}_i = y_i$

Now we use no part of the other observations to help learn the $y \sim x$ relationship

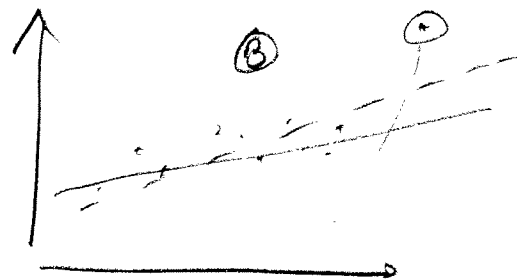
$$h_{ii} \text{ "the leverage"} = \left(\frac{(x_i - \bar{x})^2}{\sum_{j=1}^n (x_j - \bar{x})^2} + \frac{1}{n} \right)$$



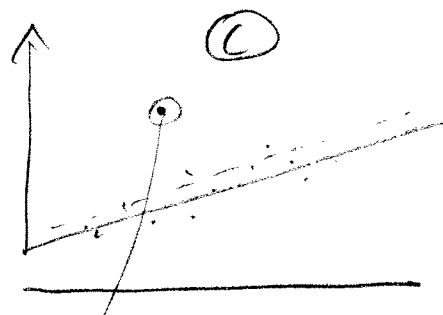
Problem: If observation i has high leverage, a contamination (or outlier) at this location can have a huge impact on the analysis. (8)



High leverage, but
not influential



High leverage, and
influential



Low leverage, but
creates a bias in the
fit. "Pure outlier"

Case A: Do the analysis with and without the observation i
& B included $\rightarrow$ discuss impact

Case C: Do the analysis with & without $\dots$ (9)

$\rightarrow$ check if inferences are affected

Go back to the data source and see if you can figure out why the outlier is present.

More about fitted values ;

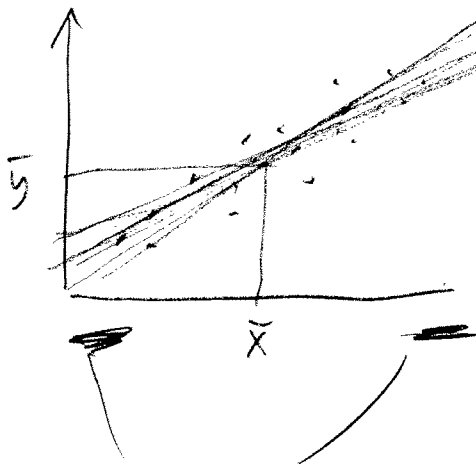
$$V(\hat{y}_i) = \sigma^2 h_{ii}$$

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

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

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

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



Different $(\hat{\beta}_0, \hat{\beta}_1)$ estimates
obtained w. samples
from $(x_i, y_i = \beta_0 + \beta_1 x_i + \varepsilon_i)$

h_{ii} large here

What about the residuals e_i ?

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

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

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

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

MEANING :

① Residuals are correlated (but errors ε_i are not)

② Residuals have nonconstant variance $V(e_i) = \sigma^2(1-h_{ii})$ (errors $V(\varepsilon_i) = \sigma^2$)

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

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

Note, all the variances involve one extra parameter σ^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.