

Handout 2

①

RECAP

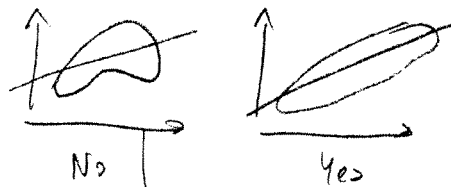
Summary $y \sim x$ relationship $y_i = \beta_0 + \beta_1 x_i + \epsilon_i$, ϵ_i : uncorrelated, $E(\epsilon_i) = 0$ and $V(\epsilon_i) = \sigma^2$

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

equal contribution of all terms

squared error

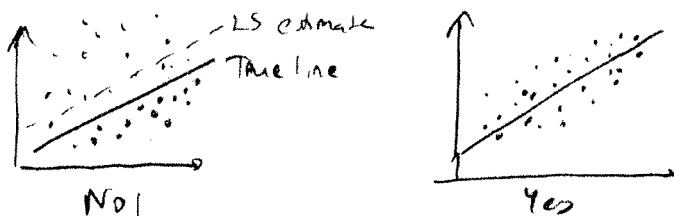
BASIC ASSUMPTIONS ① Model sufficiency



(check? Plot, plot, plot!)

Solution? Try data transformations (e.g. x^2 instead of x , $\log(y)$ instead of y)

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.

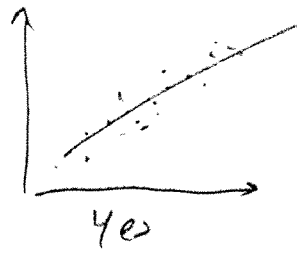
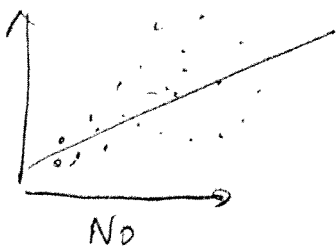
- Is the data sampled sequentially over time? ⇒ Time Series

- Or samples nested? (e.g. students from the same class)

Solution ⇒ Mixed effects or patients from the same hospital
Hierarchical Linear Model

②

④ Constant variance of error scatter



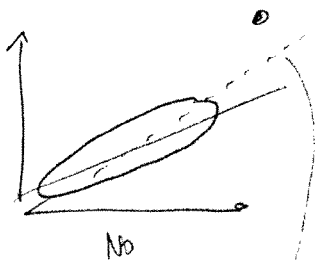
(Check? Plot, plot, plot ; both y on x and residuals)

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

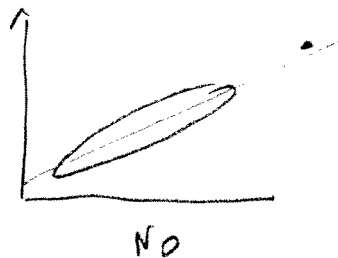
$\Rightarrow$ if that doesn't work $\Rightarrow$ Weighted Least Squares

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

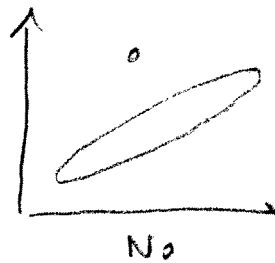
⑤ No outliers



Effect: bad fit



Effect: fit looks "too good"



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

(9)

Now, can also write errors $e_i = (y_i - (\beta_0 + \beta_1 x_{i1} + \dots + \beta_{p-1} x_{i,p-1}))$

in vector form
$$\underline{e} = \begin{pmatrix} e_1 \\ e_2 \\ \vdots \\ e_n \end{pmatrix} = \underline{y} - \underline{X}\underline{\beta}$$

$$\Rightarrow Q = \underline{e}'\underline{e} = (\underline{y} - \underline{X}\underline{\beta})'(\underline{y} - \underline{X}\underline{\beta})$$

We want to minimize Q w.r.t respect to $\underline{\beta}$.

$\Rightarrow$ Matrix derivative

$$\frac{\partial Q}{\partial \underline{\beta}} = -2\underline{X}'(\underline{y} - \underline{X}\underline{\beta}) = 0$$

inner
derivative

$$\Rightarrow \boxed{(\underline{X}'\underline{X})\underline{\beta} = \underline{X}'\underline{y}}$$

Normal equations in matrix form.

"Interpretation" $(\underline{X}'\underline{X}) \sim \text{Cov}(\underline{X})$, where diagonal of $\underline{X}'\underline{X}$ matrix is the $\text{Var}(X_1), \text{Var}(X_2), \dots$ and off-diagonal elements are the covariances $\text{Cov}(X_i, X_j), \dots$

So if some X 's are linearly dependent, or close to
 $\rightarrow (\underline{X}'\underline{X})^{-1}$ may not exist, or be numerically unstable

$(\underline{X}'\underline{y}) \sim \text{Cov}(\underline{X}, \underline{y})$ — measure of linear dependence between the

① $(X'X)^{-1}$ exists $\Rightarrow \hat{\beta} = \underline{(X'X)^{-1} X'y}$

②

Problems: If $(X'X)^{-1}$ numerically unstable, meaning X 's are related ~~it~~ difficult to interpret the meaning of the slopes $\hat{\beta}_1, \dots, \hat{\beta}_{p-1}$ - they are not just measuring the contribution of one x , if x 's are related.

Back to simple linear regression $\Rightarrow$ ③

Properties

• unbiased estimates

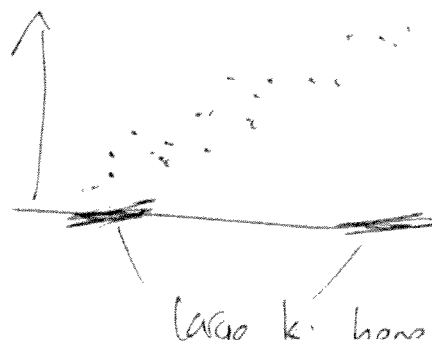
$$\hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$$

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

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

That is, $\hat{\beta}_1$ is a weighted average of y -values!

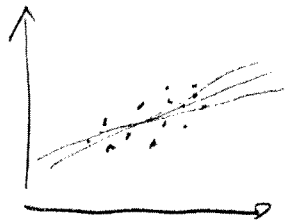
Note, extreme x_i 's, far from $\bar{x}$, contribute the most to the slope estimate



(5)

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

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



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



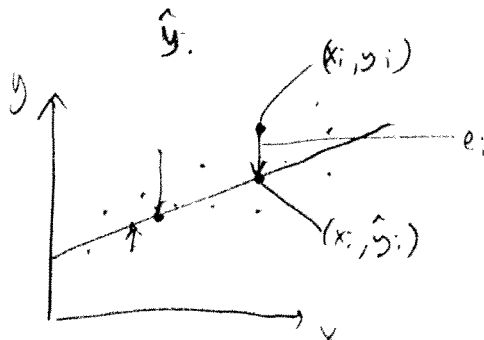
$\sum (x_i - \bar{x})^2$ large

More properties

- Residuals sum to 0, $\sum e_i = 0$ (i) 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

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



⊗ 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

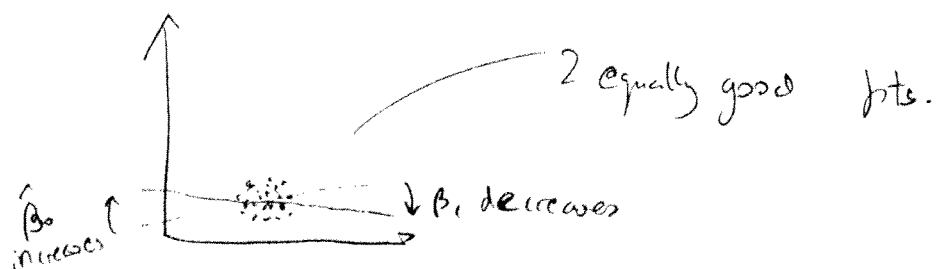
Also $\text{Cor}(\hat{\beta}_0, \hat{\beta}_1) = \frac{-\bar{x} \sigma^2}{\sum (x_i - \bar{x})^2}$, large negative if x near constant

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 .



#

More about fitted values ;

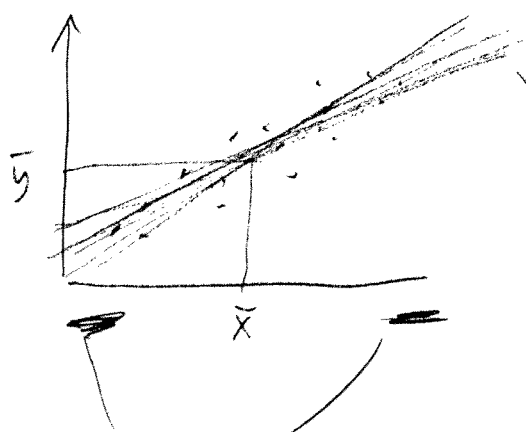
$$V(\hat{y}_i) = \sigma^2 h_{ii}$$

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

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

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

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 $(x_i, y_i: \beta_0 + \beta_1 x_i + \epsilon_i)$

h_{ii} large here

What about the residuals e_i ?

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

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

$\swarrow$
 $\left(\begin{array}{l} \perp \text{Cov}(e_i, \hat{y}_i) = 0 \\ \text{LS properties} \end{array} \right)$

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

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 there 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 .

Fact

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

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" $y = \hat{y} \oplus e$

/ \
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.

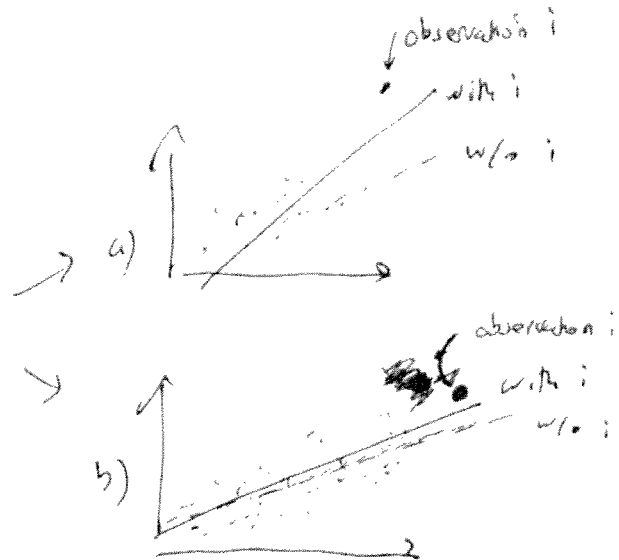
(1) $h_{ii} \gg h_{is}$ $\hat{y}_i$ is dominated by ... 1. 11

12

What does this mean?



Flex 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 ~~influence~~ influential outlier, producing a poor fit to the rest of the data.

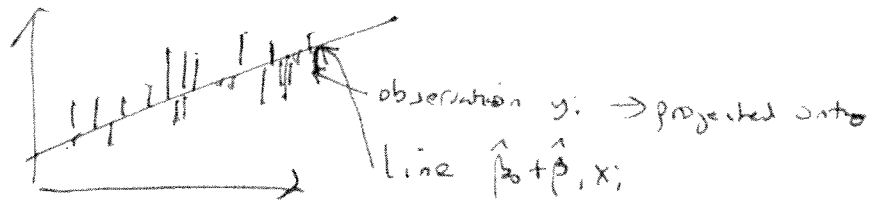
In case b), the slope estimate is not much affected, but we get an exaggerated sense of the goodness of fit of the model. (more later).

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

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

(14)

- This H matrix is a projection matrix.
- It takes the data y , and projects it onto a function of 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)
 $= \sum k_i^2 \sigma^2 = \frac{\sum (x_i - \bar{x})^2}{(\sum_j (x_j - \bar{x})^2)^2} \sigma^2$ ($V(y_i) = \sigma^2$ since the only thing random in y_i is ε_i)

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

$$\sum_j (x_j - \bar{x})^2$$

Note

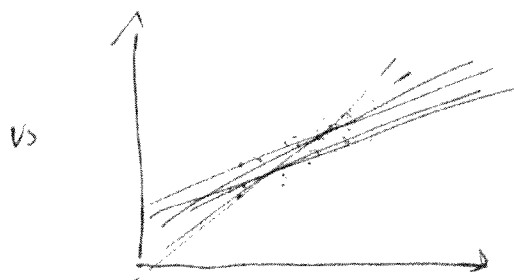
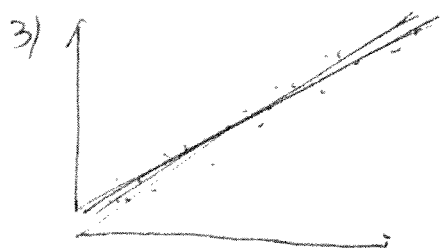
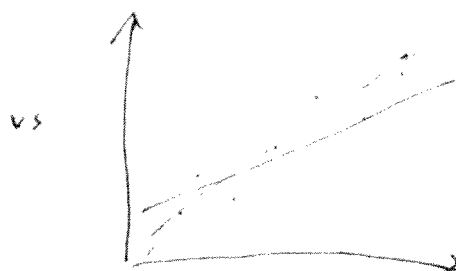
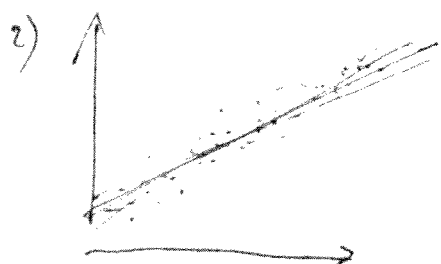
- LS is the minimum variance & unbiased linear estimate, any other weights $c_i \neq k_i \Rightarrow$ larger variance.

- 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

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

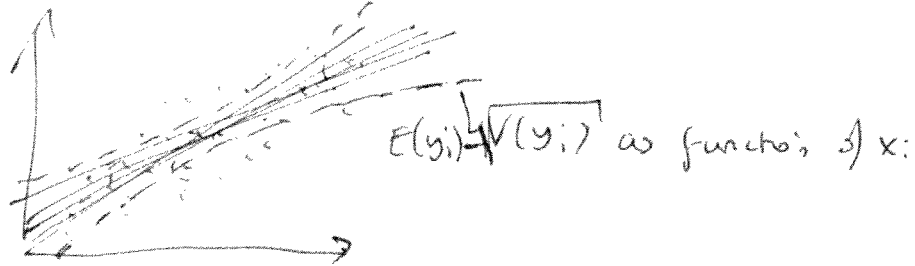
$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i \Rightarrow V(\hat{y}_i) = V(\hat{\beta}_0 + \hat{\beta}_1 x_i) = V(\bar{y} + \hat{\beta}_1 (x_i - \bar{x}))$$

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

$$\begin{aligned} \Rightarrow V(\bar{y} + \hat{\beta}_1 (x_i - \bar{x})) &\stackrel{\Downarrow}{=} V(\bar{y}) + (x_i - \bar{x})^2 V(\hat{\beta}_1) = \frac{\sigma^2}{n} + \frac{\sigma^2 (x_i - \bar{x})^2}{\sum_j (x_j - \bar{x})^2} \\ &= \sigma^2 h_{ii} \end{aligned}$$

(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 $\text{cov}(e_i, e_i)$

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

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

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

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

= 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.

One more parameter

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i \quad \begin{array}{l} \swarrow \text{Symmetry} \\ E(\varepsilon_i) = 0 \end{array} \quad \begin{array}{l} \swarrow \text{Constant variance} \\ V(\varepsilon_i) = \sigma^2 \end{array}$$

LS $\rightarrow$ estimates of β_0, β_1 , but we also want to estimate the amount of scatter around the line $\Rightarrow$ this is a measure of how well the line summarizes the data y .

Scatter variance σ^2 is estimated by the observed residual variance

$$\begin{aligned} \Rightarrow \text{RSS (residual sum of squares)} &= \sum_{i=1}^n (y_i - \hat{y}_i)^2 \\ \text{"} & \\ \text{SSE (Error sum of squares)} & \end{aligned} \quad \underbrace{\quad}_{\text{(= minimized Q)}}$$

$\Rightarrow$ MSE (the mean squared error)

$$= \frac{\text{RSS}}{n-2} = \frac{\sum (y_i - \hat{y}_i)^2}{n-2} = \hat{\sigma}^2$$

|
Why 2?

- Usually, a variance estimate is $\frac{\sum (y_i - \bar{y})^2}{n-1}$, but here our model is not an average value for $\bar{y}$ (a one-parameter model) but the line $\beta_0 + \beta_1 x_i = \hat{y}_i$, which is based on 2 fixed parameters!

Variance Decomposition

Use of $\hat{\sigma}^2$.

Total 'variance' $SS_T = \sum (y_i - \bar{y})^2$ of the data
(Total sum of squares)

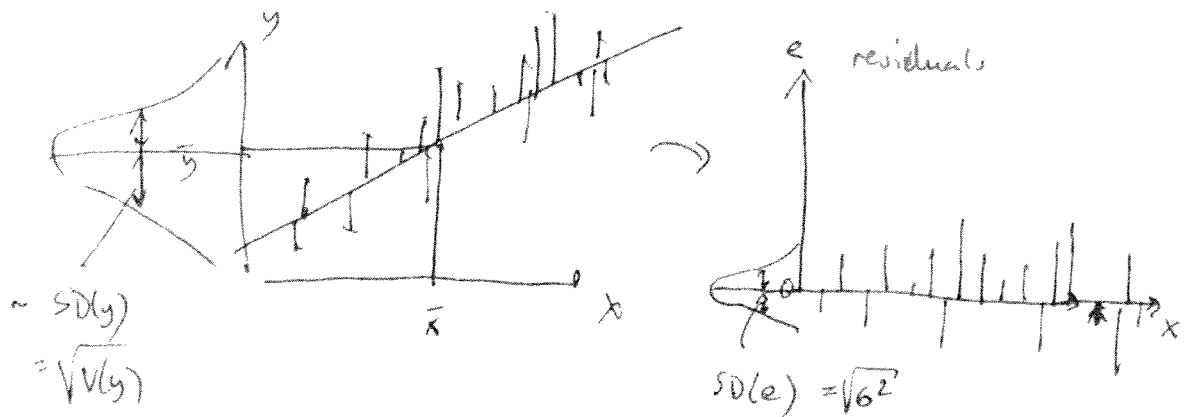
2 parts

$$SS_{reg} = \sum (\hat{y}_i - \bar{y})^2$$

$$RSS \text{ (or } SS_E) = \sum (y_i - \hat{y}_i)^2$$

- explainable part
- Variation of line around the mean

- unexplainable part
- scatter around the line



Effect of regression $\rightarrow$ reducing the unexplained variability $V(y) \rightarrow \sigma^2$

% of variability explained: $R^2 = 1 - \frac{RSS}{SS_T} = \frac{SS_{reg}}{SS_T} \in [0, 1]$

called the Multiple R-squared