

## Chap 13

### Fancy Regression

— SLR need not be so simple.

① log & log-log transformations.

$$\text{model: } E\{\log Y\} = \beta_0 + \beta_1 x$$

$$E\{Y\} \approx e^{\beta_0} e^{\beta_1 x}$$

non-linear fit with a SLR  
through log transformation of  $Y$

After you have fit a model with  
a transformation, it's helpful to  
back-transform, especially for  
visuals.

Consider  $E\{Y\} = ax^b$  a multiplicative model.

log-log model:  $\underbrace{\log E\{Y\}}_{\log} = \underbrace{\log a}_{b_0} + \underbrace{b \log x}_{b_1 \log x}$

A log-log model is appropriate whenever quantities are linearly related on a multiplicative or percentage scale.

estimated slope  $\hat{b}_1$  in a log-log  
SUR is known as elasticity

$\downarrow$   
% change in  $Y$  given  
1% change in  $X$

# Ch 13

## ⑥ Polynomial Regression

Taylor's theorem: any curve can be approximated locally by a low-degree polynomial

Weierstrass' Theorem : any function can be approximated globally by a high-degree polynomial

→ the sweet spot is somewhere in between.

Model :  $y_i = \beta_0 + \beta_1 x_i + \beta_2 x_i^2 + \dots + \epsilon_i$ ,  
 $\downarrow$   
 Technique a MLR  $\epsilon_i \sim N(0, \sigma^2)$

You can choose to "power up" as much as you'd like, but be careful:

- when  $n \leq d-1$  ( $d$  degree of poly)  
you get a "perfect" fit, because  
the polynomial interpolates all points
- but it is not "ideal" since  $\epsilon_i = 0$   
for all  $i$ , which isn't Gaussian

→ "overfit"

you can check if you need degree  $d$   
by testing if  $\epsilon_d = 0$  or not.

## Chapter 13

MLR

### ① Multiple Linear Regression

extends SLR to more than one input  $x$ .

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_d x_{id} + \varepsilon_i$$
$$\varepsilon_i \stackrel{\text{iid}}{\sim} N(0, \sigma^2)$$

$$y_i | x_{i1}, \dots, x_{id} \stackrel{\text{iid}}{\sim} N(\underbrace{\beta_0 + \beta_1 x_{i1} + \dots + \beta_d x_{id}}_{\mu_i}, \sigma^2)$$

$$Y \sim N_n(X\beta, \sigma^2 I_n)$$

$\hat{=}$   $n$ -variate  $MVN$

$$Y = (y_1, \dots, y_n)^T \quad n\text{-vector}$$

$$\beta = (\beta_0, \beta_1, \dots, \beta_d) \quad p = d+1\text{-vector}$$

$$I_n = n \times n \text{ identity matrix}$$

$X = n \times p$  design matrix with  
 $x_i^T = (1, x_{i1}, \dots, x_{id})$

SLR is a special case where  $d=1$ .

$d=2 \rightarrow$  linear surface is a plane

$d \geq 3 \rightarrow$  we say hyperplane

$$\hat{\beta}_j = \frac{\partial E\{Y | x_1, \dots, x_d\}}{\partial x_j}$$

Calculating the MLE is the same as in Ch 7  
 but via MVN with  $X\beta \equiv \mu$  &  $\Pi_{\sigma^2} \equiv \Sigma$

$$L(\beta, \sigma^2) = (2\pi\sigma^2)^{-\frac{n}{2}} \exp \left\{ -\frac{1}{2\sigma^2} \underbrace{(Y - X\beta)^T (Y - X\beta)}_n \right\}$$

$$\begin{aligned} \ell(\beta, \sigma^2) &= c - \frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} (Y - X\beta)^T (Y - X\beta) \\ &\downarrow \\ &= c - \frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} (Y^T Y - 2Y^T X\beta + \beta^T X^T X\beta) \end{aligned}$$

$$\phi_{\text{sex}} = \frac{\partial \ell}{\partial \beta} = \frac{1}{\sigma^2} (x^T y + x^T x \beta)$$

$$x^T x \beta = x^T y$$

$$\hat{\beta} = (x^T x)^{-1} x^T y$$

looks like

$$b_1 = \frac{S_{yx}}{S_{xx}}$$

Important Projection result for mvns

If  $y \sim N_n(\mu, \Sigma)$  then  $Py \sim N_p(P\mu, P\Sigma P^T)$  &  $P$  is a  $p \times n$  projection matrix

Consider  $P = (x^T x)^{-1} x^T$  so  $\hat{\beta} = Py$

observe  $Py\beta = (x^T x)^{-1} x^T x \beta = \beta$

$$PP^T = (x^T x)^{-1} \cancel{x^T x} \cancel{(x^T x)^{-1}} = (x^T x)^{-1}$$

$$\text{so } \hat{\beta} = Py \sim N_p(\beta, \sigma^2 (x^T x)^{-1})$$

↑  
MLE

↓  
sampling distribution

What to use for  $\sigma^2$ ?

$$l(\hat{\beta}, \sigma^2) = -\frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} (Y - X\hat{\beta})^T (Y - X\hat{\beta})$$

$$\phi^{st} = \frac{dl}{d\sigma^2} = -\frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} (Y - X\hat{\beta})^T (Y - X\hat{\beta})$$

$$\hat{\sigma}^2 = \frac{1}{n} (Y - X\hat{\beta})^T (Y - X\hat{\beta})$$

Biased estimate, but you  
can correct that bias

$$s^2 = \frac{1}{n-p} (Y - X\hat{\beta})^T (Y - X\hat{\beta})$$

Can show that  $\sigma^2 \sim (n-p)s^2 / \chi^2_{n-p}$

so you can use

Student-t for hypothesis testing

1 C5.



Prediction for MLR involves another projection since

$$\hat{\underline{y}} = \underline{X} \hat{\underline{\beta}} = \underline{H} \underline{y}$$

where  $\underline{H} = \underline{X}(\underline{X}^T \underline{X})^{-1} \underline{X}^T$

↙  
"hat matrix"

Details of how to use  $\underline{H}$  for prediction at new  $\underline{x}_p$  in book.

## Chapter 13

### ① Dummies & Interactions

A dummy variable allows you to use a categorical grouping in a MLR.

if  $w_i \in \{\text{Dodge, Ford, GMC}\}$

Then  $x_{i2} = \mathbb{I}_{\{w_i = \text{Ford}\}}$

$x_{i3} = \mathbb{I}_{\{w_i = \text{GMC}\}}$

Dodge is  $(x_{i2}, x_{i3}) = (0, 0) \rightarrow$  in the "Intercept"

Ford  $= (1, 0)$

GMC  $= (0, 1)$

Model:  $y = b_0 + b_1 x_1 + b_2 x_2 + b_3 x_3 + \varepsilon$   
 $\downarrow$   
year

An MLR with entirely dummy variables  
is an ANOVA

→ see book for ice cream  
example

→ testing, wait to  $\textcircled{E}$

---

An interaction effect allows the contribution  
of one input  $x_j$  to depend on the level  
of another input  $x_k$

$$\text{Model: } y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \underbrace{\beta_3 x_{i1} x_{i2}}_{\text{interaction term}} + \dots + \varepsilon_i$$

so that

$$\frac{dE\{y \mid x_1, x_2\}}{dx_1} = \beta_1 + \beta_3 x_2$$

↑

## Chapter 13

### (E) Model Selection

How to select from two "nested" MLR models?

→ good stat is  $R^2 \equiv \text{cor}(Y, \hat{Y})$   
 $\Delta R^2$

Take inspiration from ANOVA

$$H_0: \beta_{\text{Red}+1} = \dots = \beta_{\text{full}} = 0$$

v.  $H_1$ : at least one of  $\beta_j \neq 0 \quad j > p_{\text{Red}}$

"easy" to visualize an  $R^2$  comparison with these hypotheses.

The math version is also similar to ANOVA: look at ratios of MSE, which can be expressed as  $R^2$ -values

$$f = \frac{(SSE_{full} - SSE_{base}) / p_{\Delta}}{SSE_{full} / (n - p_{full})} \rightarrow p_b - p_f$$

$$= \frac{R_{\Delta}^2 / p_{\Delta}}{(1 - R_{full}^2) / (n - p_{full})} \quad R_{\Delta}^2 = R_{full}^2 - R_{base}^2$$

Can show that  $F \sim F_{p_{\Delta}, n - p_{full}}$