

Properties of Exp. families

(7)

— need this to come up with a fitting scheme.

① loglikelihood $l(\theta, \varphi | y) = \sum_i \left(\frac{y_i \theta_i - b(\theta_i)}{a(\varphi)} + c(y_i, \varphi) \right)$

② score function $l'_i(\theta) = U[\theta_i] = \frac{\partial l}{\partial \theta_i} = \frac{y_i - b'(\theta_i)}{a(\varphi)}$

Since $l(\theta) = \log f_y(\theta) \Rightarrow$ chain rule

$$l'(\theta) = \frac{f'_y(\theta)}{f_y(\theta)}$$

$$\Rightarrow E(l'(\theta)) = \int f_y(\theta) \frac{f'_y(\theta)}{f_y(\theta)} dy$$

"nice"
+ $\frac{d}{d\theta} \int f_y(\theta) dy = \frac{d}{d\theta} 1 = 0$

since f
is a density

So $E(\text{Score}) = 0 \Rightarrow E\left(\frac{y_i - b'(\theta_i)}{a(\varphi)}\right) = 0$

$E(y_i) = b'(\theta_i)$

(From 2nd derivative)
 $V(y_i) = a(\varphi) b''(\theta_i)$

This $E \sim V$ relationship is assumed known and follows from choice of family as Pearson or Binomial

So - how to fit GLMs to data?

⑧

Two problems $\longrightarrow E(y) = \mu_i = g'(\theta_i) = g'(x_i' \beta)$

is often non-linear in β

①

Since $V(y_i)$ depends

on $E(y_i)$ through

● $b''(\theta_i)$ we have

non-constant error

②

● Variance

① $\rightarrow$ TLS like we did
for nonlinear regression

② $\rightarrow$ Weighted Least Squares (WLS)

①+② called

Iterative re-weighted Least Squares

or

IWLS

WLS

(4)

Let's say $V(y_i) = \sigma_i^2$ not constant σ^2

$\Rightarrow V(y) = \sigma^2$ ~~✓~~ varying components. y ~~✓~~ non-diagonal — y 's are correlated as well.

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

$$E(\hat{\beta}) = \beta, V(\hat{\beta}) = \sigma^2 (X'X)^{-1} (X' \cancel{V} X) (X'X)^{-1}$$

~~WLS is better than OLS~~

WLS

$$\min_{\beta} \sum_{i=1}^n w_i (y_i - x_i' \beta)^2 = \sum_{i=1}^n (w_i^{1/2} y_i - x_i' \beta)^2 = \|W^{1/2} y - W^{1/2} X \beta\|^2$$

$$= \|\tilde{y} - \tilde{X} \beta\|^2 \quad \text{transformed variables}$$

$$\begin{cases} \tilde{y} = W^{1/2} y \\ \tilde{X} = W^{1/2} X \end{cases}$$

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

$$\hat{y} = Hy = \tilde{X}' \hat{\beta}$$

$$= (X' W X)^{-1} X' W y$$

$$= \underbrace{(W^{1/2} X (X' W X)^{-1} X' W^{1/2})}_H y$$

Now if $V(y) = \sigma^2 V$ and we use $W = V^{-1}$

$$\Rightarrow V(\hat{\beta}) = \sigma^2 (X' W X)^{-1} \quad \text{best possible}$$

That is, use weights $w_i \sim \frac{1}{V(y_i)}$

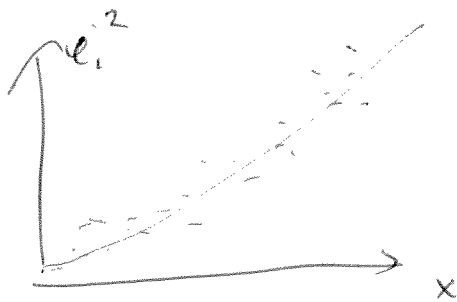
What if I don't know $V(y_i)$

(10)

(can learn from data)



→ get first regression fit w/o weight



regression of e_i^2 on $e_i x$

$$\Rightarrow V(y_i) = x_i' \delta$$

For GLMs

$$V(y_i) = \underbrace{a(\eta) b''(\theta_i)}_{\text{depends on unknown mean parameters } \theta_i}$$

depends on unknown mean parameters θ_i

So what we will do is - in each iterative step
- use previous values of $\hat{\theta}_i \Rightarrow \hat{V}(y_i) \Rightarrow w_i = \frac{1}{\hat{V}(y_i)}$

Slip? IWLS

(11)

(0) Remember $E(y_i) = \mu_i$, $g(\mu_i) = \eta_i = x_i' \beta$

$$g(\mu) \approx g(\mu_0) + (\mu - \mu_0) \underbrace{\frac{\partial g(\mu)}{\partial \mu}}_{\left. \frac{\partial \eta}{\partial \mu} \right|_{\mu_0}} \bigg|_{\mu_0}$$

linearize at value μ_0
(or $\eta_0 = x_i' \beta_0$)

η_0

$\frac{\partial \eta}{\partial \mu} \bigg|_{\mu_0}$

(1) Initial values $\beta^0 \rightarrow \eta_i^0 = x_i' \beta^0 \rightarrow \mu_0 = g^{-1}(\eta_0)$

(2) "Working values" $Z_0 = \eta^0 + (y - \mu_0) \frac{\partial \eta}{\partial \mu} \bigg|_{\mu_0}$

"linear approx of data near η^0 "

"Working weights" $W_0^{-1} = \left(\frac{\partial \eta}{\partial \mu} \bigg|_{\mu_0} \right)^2 V_0$

$$b''(\theta) \bigg|_{\mu_0} = \frac{\partial \mu}{\partial \theta} \bigg|_{\mu_0}$$

known variance for model

family e.g. μ_i for Poisson

$\pi_i(1-\pi_i)$ for Binomial

(3) Regress Z_0 on X

$$\rightarrow \hat{\beta}^{\text{new}}$$

$$\rightarrow \hat{\eta}^{\text{new}} \rightarrow \text{Go to (1)}$$

Validation & Inference

(12)

LM & NLH $\Rightarrow$ RSS $\Rightarrow$ F-tests, R^2 etc

Here, everything is based on the likelihood

Residual Deviance of a model $M = -2 \times \log \text{likelihood}(y, \hat{\theta})$

plays the role of ~~the~~ RSS of a model in LM and NLH

- — denote by $D(M)$
- $\frac{D(M)}{a(p)}$ is called the scaled deviance.

η normal error case

$$\frac{D(M)}{a(p)} = \frac{\text{RSS}(M)}{\sigma^2} \sim^d \chi^2_{n-p}$$

- Here, that is only true for very large data sets (asymptotically)

- Large n $\left| \frac{D(M)}{a(p)} \sim^d \chi^2_{n-p} \right|$

Instead of an F-test for goodness-of-fit, we use

a χ^2 -test. $\eta \frac{D(M)}{a(p)} \gg \chi^2_{n-p}(1-\alpha)$ we reject the fit.

$\rightarrow$ Try transformations, other links, other model families

We can also use the deviance to check if our model is appropriate in terms of scale

$$\hat{a}(\psi) = \frac{D(\mu)}{n-p}$$

if $\hat{a}(\psi) \gg$ assumed $a(\psi)$ (e.g. 1 for Poisson and Binomial)

→ we call the data over-dispersed → other model family needed.

Model selection

Simple model M_0 , Complex model M — which to choose

e.g. x_1, x_2, x_3
3 parameters

x_1, x_2, x_3, x_4, x_5
5 parameters

① if ψ known → χ^2 -test Analysis of Deviance

$$\frac{D(M_0) - D(M)}{p} \sim \chi^2_{p-p_0} \quad \text{if } M_0 \text{ sufficient}$$

if $\frac{D(M_0) - D(M)}{p} >> \chi^2_{p-p_0}(1-\alpha)$ reject M_0 .

(b) γ ρ not known $\rightarrow$ estimate it as above $\hat{\gamma} = \frac{D(M)}{n-p}$ (14)

$\rightarrow$ use an approximate F-test

$$\frac{D(M_0) - D(M)}{\hat{\gamma} (p - p_0)} \sim F_{p-p_0, n-p} \quad \text{if } M_0 \text{ sufficient.}$$

(This is pretty ad-hoc, but works ok in practice)

(c) Can always use CV, bootstrap, BIC etc

Diagnosis

Can be tricky because the residuals look kind of weird.

$\rightarrow$ Make sure that residuals on average

≈ 0 as a function of $\hat{y}$, x 's.

Just like NHM $\rightarrow$ LT's based on last line

(15)

approximation

— now use Z-intervals (based on large n)

— check profile function and if not linear

$\rightarrow$ bootstrap.

Common problems to be aware of

Over-dispersion

$$\frac{D(M)}{n-p} \gg \psi \text{ assumed}$$



• Inflate all SE's by $\hat{\psi}$

• ~~bootstrap~~

• Other models

Binomial data w.
clear 0/1 separation

$\rightarrow$ can get really weird
 $\hat{\beta}'s \rightarrow \pm \infty$ ~~values~~

$\rightarrow$ can trust $\hat{\pi}_i$, but not
 $\hat{\beta}$ values.

Like in LM, be aware of high
leverage / influential observations.