

Why would anyone ever use a linear model?

Saturday, January 12, 2013
1:14 PM

① test hypothesis (test theory)

② predict future observations

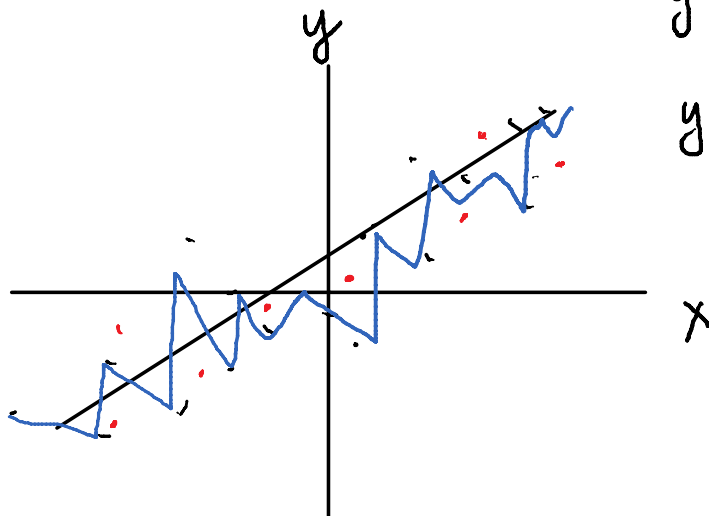
— policy initiative

— forecasting

— validation (out-of-sample)

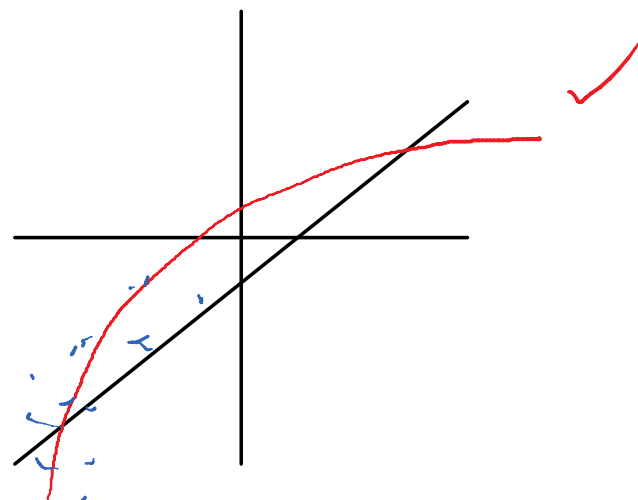
statistical
model

inference + predictions = data + model assumptions



$$y = mx + b + \varepsilon$$

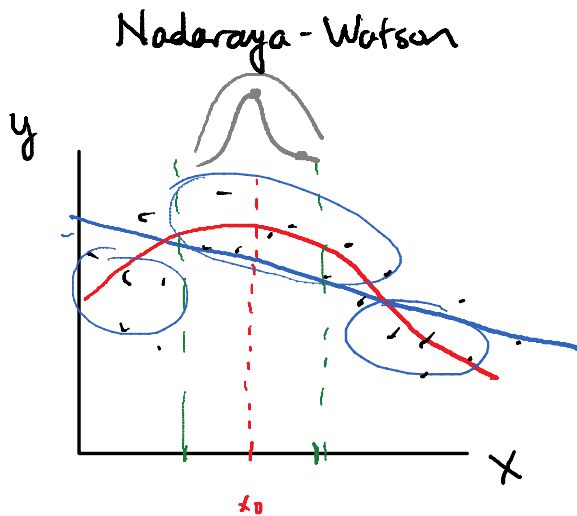
$$y = \beta_0 + \beta_1 \cdot x + \varepsilon$$



Low-assumption approaches

Saturday, January 12, 2013
1:20 PM

non-parametric models (no formal mathematical structures imposed on the data that are flexible only through a discrete set of tunable parameters)



$$\hat{y}(x_0) = \frac{\sum y_i \cdot k(x_i - x_0)}{\sum k(x_i - x_0)}$$

when assumptions are good, they add value to an analysis.

- ① extract better information ...
- ② ... with greater confidence.



these advantages grow as the # of predictor variables grows.

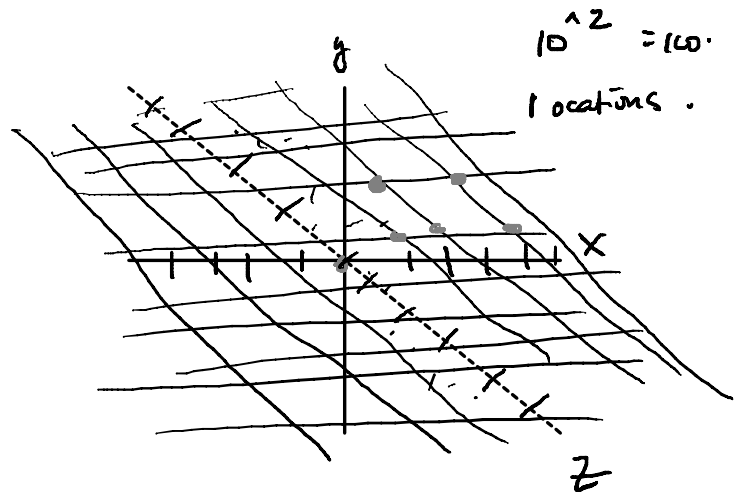
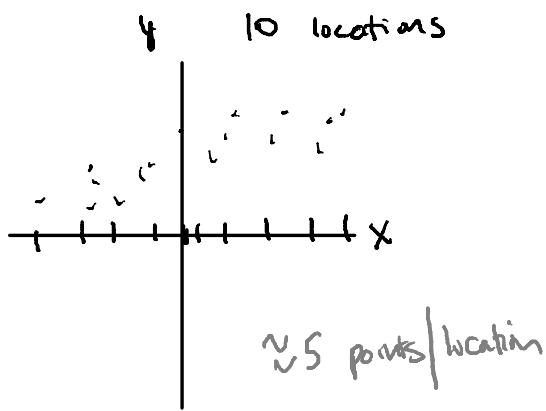
$N = 50$

$\frac{1}{2}$ point/location

4 10 locations

8

$10^2 = 100$



"Curse of dimensionality."

A first take on OLS: assume correct specification

Saturday, January 12, 2013

1:37 PM

$$\begin{matrix} k \times 1 \\ \text{red} \end{matrix} x' y = \begin{matrix} n \times k \\ \text{blue} \end{matrix} x' X \beta + x' u$$

↓

$$(x'x)^{-1} x' y = \underline{(x'x)^{-1}} \underline{x'x} \beta + (x'x)^{-1} x' u$$

$$(x'x)^{-1} x' y = \beta + \underbrace{(x'x)^{-1} x' u}_{=0} \rightarrow$$

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

$$= 0$$

↳ coefficients of x on the noise term, u .

assume that x and u are uncorrelated, this coef = 0.

Assumptions

① correct specification

→ questionable?

② x and u are uncorrelated.

The invertibility of $X'X$

Saturday, January 12, 2013
1:35 PM

$$\boxed{\begin{matrix} (X'X)^{-1} \\ \begin{matrix} k \times k & n \times k \\ k & k & k \end{matrix} \end{matrix}}$$

X be of full rank (column rank).
 $n \times k$

Theorem: $\text{rank}(X'X) = \text{rank}(X)$.

Pf. : omitted.

Invertible Matrix Theorem:

- ① $(X'X)^{-1}$ exists.
- ② $(X'X)^{-1}$ has full rank.

X has rank k when all its k variables (columns) are linearly independent (not perfectly collinear).

Regression only works when independent variables are not linear functions of one another.

$$\left[\begin{array}{cccc} \text{const} & = & c_1 & + & c_2 & + & c_3 \\ 1 & & 1 & & 0 & & 0 \end{array} \right]$$

$X'X$ not invertible.

$$\begin{bmatrix} 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 \end{bmatrix}$$

$X'X$ not invertible.

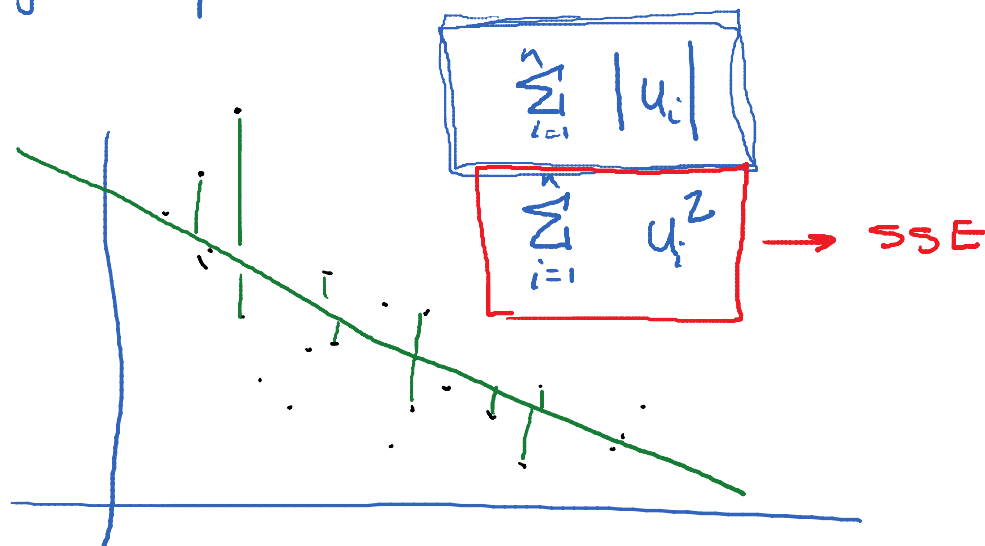
$$2 \text{inc} + \frac{1}{2} \text{tax} = \boxed{51} \text{ inc tax } \times$$

Optimization theory (from calculus)

Saturday, January 12, 2013

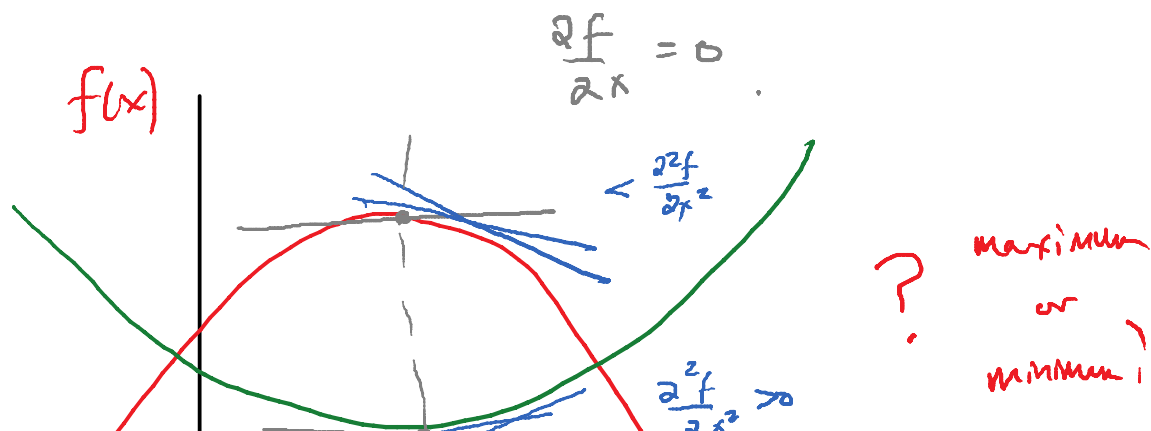
1:21 PM

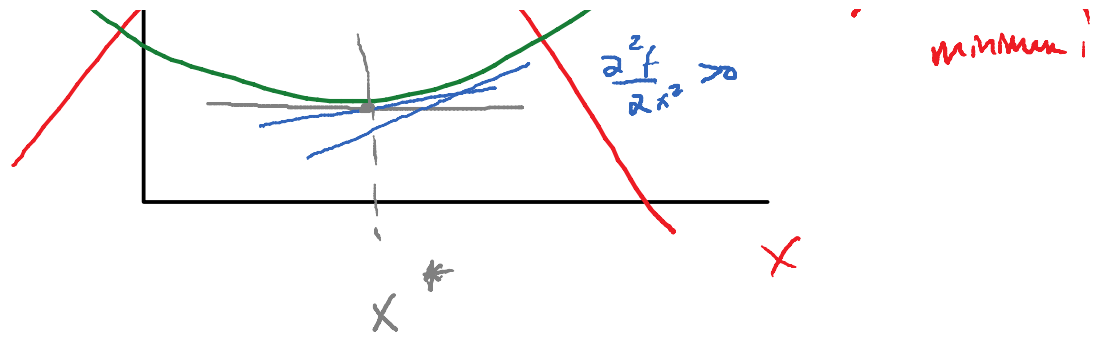
- approximating a DGP w/ linear model
- pick β coefficients that minimize the magnitude of estimated errors, $\hat{u}$



optimization

$f(x)$, pick x to minimize $f(x)$.





$$\min f(\beta_0, \beta_1, \dots, \beta_n) \text{ wrt } \beta_0, \beta_1, \dots, \beta_n$$

$$\left. \begin{array}{l} \frac{\partial f}{\partial \beta_0} = 0 \\ \frac{\partial f}{\partial \beta_1} = 0 \\ \vdots \\ \frac{\partial f}{\partial \beta_n} = 0 \end{array} \right\} \text{ Solve simultaneously.}$$

u

└

$$H_k = \begin{bmatrix} \frac{\partial^2 f}{\partial \beta_0^2} & \frac{\partial^2 f}{\partial \beta_0 \partial \beta_1} & \frac{\partial^2 f}{\partial \beta_0 \partial \beta_2} & \dots & \frac{\partial^2 f}{\partial \beta_0 \partial \beta_k} \\ \frac{\partial^2 f}{\partial \beta_1 \partial \beta_0} & \frac{\partial^2 f}{\partial \beta_1^2} & \frac{\partial^2 f}{\partial \beta_1 \partial \beta_2} & \dots & \frac{\partial^2 f}{\partial \beta_1 \partial \beta_k} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial \beta_k \partial \beta_0} & \frac{\partial^2 f}{\partial \beta_k \partial \beta_1} & \frac{\partial^2 f}{\partial \beta_k \partial \beta_2} & \dots & \frac{\partial^2 f}{\partial \beta_k^2} \end{bmatrix}$$

Yang's theorem: implies H is symmetric.

H must be positive definite

p.d. = a $k \times k$ matrix A is p.d. iff

$$x'Ax > 0 \quad \forall x \in \mathbb{R}^k$$

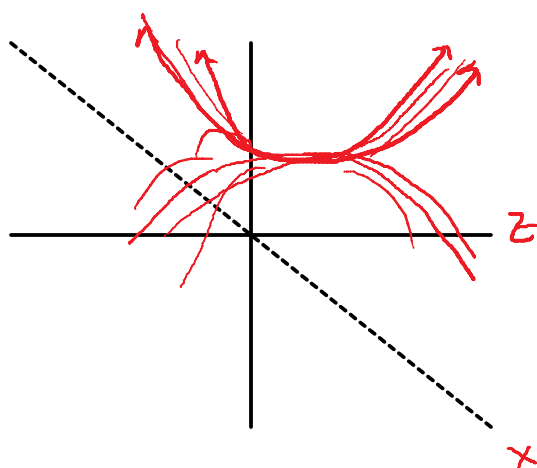
Solving for $\hat{\beta}$ via optimization

Saturday, January 12, 2013

1:22 PM

$$\begin{aligned} \min \\ \frac{\partial f}{\partial \beta} &= 0 \quad \forall \beta \\ H &\text{ positive definite.} \end{aligned}$$

$$\begin{aligned} \max \\ \frac{\partial f}{\partial \beta} &= 0 \quad \forall \beta \\ H &\text{ negative definite.} \end{aligned}$$



$$\min \sum \hat{u}_i^2 = \hat{u}'\hat{u}$$

$$(A+B)' = A' + B'$$

$$= (y - X\hat{\beta})' (y - X\hat{\beta})$$

$$[y' X \hat{\beta}]'$$

$$= y'y - \underbrace{y'X\hat{\beta} - (X\hat{\beta})'y}_{\dots \hat{\beta}' \dots \hat{\beta}}$$

$$= y'y - 2(x'\beta)y - (x'\beta) \beta$$

Quadratic form: $x'Ax$

$$\frac{\partial x'Ax}{\partial x} = \begin{bmatrix} \frac{\partial x'Ax}{\partial x_1} \\ \frac{\partial x'Ax}{\partial x_2} \end{bmatrix} =$$

$$= \begin{bmatrix} \underline{2ax_1 + bx_2 + cx_2} \\ \underline{2dx_2 + bx_1 + cx_1} \end{bmatrix} = \underline{x'(A + A')}$$

$$x = \begin{bmatrix} x_1 & x_2 \end{bmatrix} \quad A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$$

$$x'Ax \rightarrow ax_1^2 + bx_1x_2 + cx_1x_2 + dx_2^2$$

$$\frac{\partial}{\partial x_1} = 2ax_1 + bx_2 + cx_2$$

$$\frac{\partial}{\partial x_2} = bx_1 + cx_1 + 2dx_2$$

if A is symmetric, then

$$\frac{\partial x'Ax}{\partial x} = x'(A + A') = \underline{\underline{2x'A'}} = \underline{\underline{2Ax}}$$

back to problem...

$$\hat{u}'\hat{u} = y'y - 2(x'\hat{\beta})y + (x'\hat{\beta})'x\hat{\beta}$$

$$\hat{u}'\hat{u} = y'y - 2(x'\hat{\beta})y + \boxed{\hat{\beta}'x'x\hat{\beta}} \\ - 2\hat{\beta}'x'y$$

$$\boxed{\frac{\partial \hat{u}'\hat{u}}{\partial \hat{\beta}} = -2x'y + 2\hat{\beta}'x'x}$$

$$0 = -2x'y + 2x'x\hat{\beta}$$

$$2x'y = 2x'x\hat{\beta}$$

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

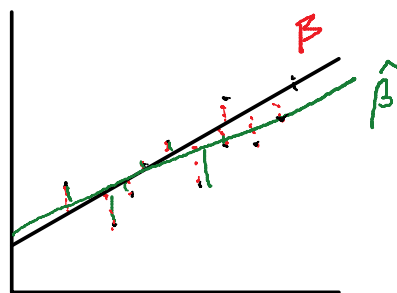
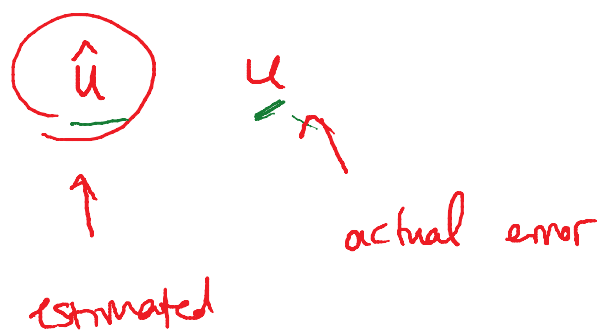
$$(x'x)^{-1} x'y = \hat{\beta} //$$

H? Theorem: if any $n \times k$ matrix X has full rank (k), it can be shown that any matrix $X'X$ is positive definite.

$$\frac{\sum \hat{u}^2}{\sum \hat{\beta}^2} = \sum (X'X)$$

Expected value of the error term

Saturday, January 12, 2013
1:37 PM



$$\frac{\partial \hat{u}' \hat{u}}{\partial \hat{\beta}} = 0$$

$$\frac{\partial}{\partial \beta} \sum_{i=1}^n (y_i - x_i \hat{\beta})^2 = 0$$

$$2 \sum (y - x \hat{\beta}) \cdot (-x) = 0$$

$$-2x \sum (y - x \hat{\beta}) = 0$$

$$\underline{-2x \sum \hat{u}_i = 0}$$

$$\boxed{\sum \hat{u}_i = 0}$$

$$\boxed{\mu_{\hat{u}} = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 0}$$

Another take on OLS: the CEF

Saturday, January 12, 2013

1:36 PM

Angrist and Pischke, pp. 28-40

$$E[y|X] = \int y \cdot f(y|x) dy$$

We just proved that $\hat{\beta}_{OLS}$ is the best linear predictor of y , in a least-squares sense.

$$\hat{\beta}_{OLS} = \arg \min_{\beta} E \left[\underbrace{E[y|x]}_{\text{CEF}} - \underbrace{x\beta}_{\text{prediction}} \right]^2$$

Pf: $\hat{\beta}$ minimizes $(y - x\hat{\beta})^2$.

$$\begin{aligned} (y - x\hat{\beta})^2 &= \left(\underbrace{y - E[y|x]}_{\text{residual}} + \underbrace{E[y|x] - x\hat{\beta}}_{\text{prediction error}} \right)^2 \\ &= (y - E[y|x])^2 + (E[y|x] - x\hat{\beta})^2 - \\ &\quad \underline{\underline{2(y - E[y|x])(E[y|x] - x\hat{\beta})}} \end{aligned}$$

$$= \underbrace{(y - E[y|x])^2}_1 + \underbrace{(E[y|x] - x\hat{\beta})^2}_1 - \underline{\underline{2E[(E[y|x] - x\hat{\beta})^2]}}$$

$$= \underbrace{(y - E[y|x]) + (E[y|x] \times \beta)}_{\substack{\downarrow \\ \text{(gap between CEF and } x\beta\text{)}^2}} - \underbrace{\varepsilon (E[y|x])' \beta}$$

CEF decomposition property:

$$y = E[y|x] + \varepsilon$$

(i) $E[\varepsilon|x] = 0$ and (ii) ε is uncorrelated with any function of x .

$$\begin{aligned} \text{(i): } E[\varepsilon|x] &= E[y - E[y|x] | x] = \\ &= E[y|x] - E[E[y|x] | x] \\ &= E[y|x] - E[y|x] = 0 \end{aligned}$$

(ii): for any function of x , $h(x)$:

$$\underbrace{E[h(x)\varepsilon]}_{= 0} = \underbrace{E[E[h(x)\varepsilon|x]]}_{\downarrow} =$$

$$E[\underline{h(x)} \cdot \underline{E[\varepsilon|x]}] = 0 \text{ by (i).}$$

↓
Law of
Iterated
Expectations:
 $E[E[a|b]] = E[a].$

$$\underbrace{\varepsilon}_{h(x)} - 2\varepsilon (E[y|x] - \underline{x\hat{\beta}})$$

$$E\left[\underbrace{(y - E[y|x])^2}_{\text{gap between CEF and } x\hat{\beta}} + \underbrace{(E[y|x] - x\hat{\beta})^2}_{\text{gap between CEF and } x\hat{\beta}} - \cancel{2\varepsilon (E[y|x] - x\hat{\beta})}\right]$$

$$E[(E[y|x] - x\hat{\beta})^2]$$

OLS minimizes this term.

$$\sum_{i=1}^n h(x_i) (E[y|x_i] - x_i \hat{\beta})$$