

Multiple Linear Regression using Matrices - II I

- Geometrically, the vector $\mathbf{y}$ is a point in the n dimensional sample space.
- The vector $\boldsymbol{\beta}$ is a point in the p dimensional *parameter space*., where $p = k + 1$.
- For each $\boldsymbol{\beta}$, $\mathbf{X}\boldsymbol{\beta}$ is a point in the sample space (i.e. the $n \times p$ matrix $\mathbf{X}$ is a mapping from the parameter space to the sample space).
- The values $\mathbf{X}\boldsymbol{\beta}$ form a p dimensional linear subspace or plane within the sample space. It is also assumed that $\mathbf{X}$ is of rank p .
- This subspace is called the *expectation surface*, because $E(\mathbf{y}) = \mathbf{X}\boldsymbol{\beta}$.

Multiple Linear Regression using Matrices - II II

- The sum of squares

$$S(\beta) = (\mathbf{y} - \mathbf{X}\beta)^T (\mathbf{y} - \mathbf{X}\beta)$$

is the squared length of the vector $\mathbf{y} - \mathbf{X}\beta$, or the squared distance between the observed $\mathbf{y}$ and the point $\mathbf{X}\beta$.

- The method of least squares finds the value of β , called $\hat{\beta}$, for which the point $\mathbf{X}\hat{\beta}$ on the expectation surface is closest to the point $\mathbf{y}$.
- From geometric considerations, $\mathbf{X}\hat{\beta}$ must be the perpendicular projection of $\mathbf{y}$ on the expectation surface.

Hat matrix

- The 'hat' matrix, $\mathbf{H} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$, projects $\mathbf{y}$ onto the expectation surface at a point $\hat{\mathbf{y}}$ closest to $\mathbf{y}$.
- So $\hat{\mathbf{y}} = \mathbf{H}\mathbf{y} = \mathbf{X}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}] = \mathbf{X}\hat{\boldsymbol{\beta}}$, where $\hat{\boldsymbol{\beta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$ is the least squares estimate of $\boldsymbol{\beta}$.
- Note that $\mathbf{H}\mathbf{X} = \mathbf{X}$. This means, for example, that any vector $\mathbf{v}$ in the column space of $\mathbf{X}$, $\mathbf{H}\mathbf{v} = \mathbf{v}$. As an example, $\mathbf{H}\mathbf{1} = \mathbf{1}$.
- Note that $\mathbf{H}$ is symmetric, and that $\mathbf{H}^2 = \mathbf{H}$.

Sampling distribution of $\mathbf{y}$

We assume that:

$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$, where

- elements of $\boldsymbol{\epsilon}$ are normally distributed with
- $\mathbf{E}(\boldsymbol{\epsilon}) = \mathbf{0}$
- $\mathbf{Cov}(\boldsymbol{\epsilon}) = \sigma^2 \mathbf{I}$

equivalently

- elements of $\mathbf{y}$ are normally distributed with
- $\mathbf{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\beta}$ and
- $\mathbf{Cov}(\mathbf{y}) = \sigma^2 \mathbf{I}$

Sampling distribution of $\hat{\beta}$

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

- elements of $\hat{\beta}$ are normally distributed, because linear combinations of normal random variables are normally distributed, with
- $E(\hat{\beta}) = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T E(\mathbf{y}) = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} \beta = \beta$ and
-

$$\begin{aligned} \text{Cov}(\hat{\beta}) &= \text{Cov}((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}) \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \text{Cov}(\mathbf{y}) (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \sigma^2 \mathbf{I} \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \\ &= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \end{aligned}$$

Sampling distribution of $\hat{\mathbf{y}}$

$$\hat{\mathbf{y}} = \mathbf{H}\mathbf{y}$$

- elements of $\hat{\mathbf{y}}$ are normally distributed, because linear combinations of normal random variables are normally distributed, with

- $E(\hat{\mathbf{y}}) = \mathbf{H}E(\mathbf{y}) = \mathbf{H}\mathbf{X}\boldsymbol{\beta} = \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T\mathbf{X}\boldsymbol{\beta} = \mathbf{X}\boldsymbol{\beta}$



$$\begin{aligned} \text{Cov}(\hat{\mathbf{y}}) &= \text{Cov}(\mathbf{H}\mathbf{y}) = \mathbf{H}\text{Cov}(\mathbf{y})\mathbf{H}^T = \mathbf{H}\sigma^2\mathbf{I}\mathbf{H}^T \\ &= \sigma^2\mathbf{H}\mathbf{H}^T = \sigma^2\mathbf{H} = \sigma^2\mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T \end{aligned}$$

Sampling distribution of residuals

$$\mathbf{r} = \mathbf{y} - \hat{\mathbf{y}} = (\mathbf{I} - \mathbf{H})\mathbf{y}$$

- elements of $\mathbf{r}$ are normally distributed, because linear combinations of normal random variables are normally distributed, with



$$\mathbf{E}(\mathbf{r}) = (\mathbf{I} - \mathbf{H})\mathbf{E}(\mathbf{y}) = (\mathbf{I} - \mathbf{H})\mathbf{X}\boldsymbol{\beta} = \mathbf{0}$$



$$\begin{aligned}\text{Cov}(\mathbf{r}) &= \text{Cov}((\mathbf{I} - \mathbf{H})\mathbf{y}) = (\mathbf{I} - \mathbf{H})\sigma^2\mathbf{I}(\mathbf{I} - \mathbf{H}) \\ &= \sigma^2(\mathbf{I} - \mathbf{H})(\mathbf{I} - \mathbf{H})^T = \sigma^2(\mathbf{I} - \mathbf{H})\end{aligned}$$

- The variance of the i 'th residual is $\sigma^2(1 - h_{ii})$ where h_{ii} is the i 'th diagonal element of $\mathbf{H}$. This means that the residuals will typically have different variances.

Covariance of residuals and fitted values



$$\begin{aligned}\text{Cov}(\mathbf{r}, \hat{\mathbf{y}}) &= \text{Cov}((\mathbf{I} - \mathbf{H})\mathbf{y}, \mathbf{H}\mathbf{y}) \\ &= \sigma^2(\mathbf{I} - \mathbf{H})\mathbf{H} = \mathbf{0}\end{aligned}$$

(Why?)

- which means that residuals and fitted values are uncorrelated
- which means that residuals and fitted values are independent (as uncorrelated normal random variables are independent).

Estimation of σ^2

- The residuals are contained in the vector $\mathbf{e} = \mathbf{y} - \hat{\mathbf{y}}$.
- The residual, or error, sum of squares is $SSE = \sum e_i^2 = \mathbf{e}^T \mathbf{e}$
- With k predictor variables, the error mean square is

$$MSE = \frac{SSE}{n - (k + 1)} = \frac{SSE}{n - p}$$

- $\hat{\sigma}^2 = MSE$ is an unbiased estimate of σ^2 because

$$\begin{aligned} E(SSE) &= E(\mathbf{e}^T \mathbf{e}) = \mathbf{0}^T \mathbf{0} + \sigma^2 \text{tr}(\mathbf{I}(\mathbf{I} - \mathbf{H})) = \sigma^2(\text{tr} \mathbf{I} - \text{tr}(\mathbf{H})) \\ &= \sigma^2(n - p) \end{aligned}$$

(Why?)

Confidence intervals I

- We will see that SSE is independent of $\hat{\beta}$, and
- where C_{jj} is the $j + 1$ 'st diagonal entry of $(\mathbf{X}^T \mathbf{X})^{-1}$, the standard error of $\hat{\beta}_j$ is $\hat{\sigma} \sqrt{C_{jj}}$, and
- $t = \frac{\hat{\beta}_j - \beta_j}{\hat{\sigma} \sqrt{C_{jj}}}$ has a t distribution with $n - (k + 1) = n - p$ degrees of freedom.
- It follows that a $100(1 - \alpha)\%$ confidence interval for β_j is given by

$$\hat{\beta}_j \pm t_{\alpha/2, n-p} \hat{\sigma} \sqrt{C_{jj}}$$

- where $\mathbf{x}_0$ is a vector of covariate values, a $100(1 - \alpha)\%$ confidence interval for the mean of y when $\mathbf{x} = \mathbf{x}_0$, $\mu_{y|\mathbf{x}_0} = E(y|\mathbf{x}_0)$, is given by

$$\hat{y}_0 \pm t_{\alpha/2, n-p} \hat{\sigma} \sqrt{\mathbf{x}_0^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{x}_0}$$

Prediction intervals (not responsible)

- In future, you plan to measure the value of y when the predictor variables equal $\mathbf{x}_0$.
- A point estimate of the future y is given by $\hat{y}_0 = \mathbf{x}_0' \hat{\beta}$.
- The goal is to find an interval in which the future y will lie, with probability $1 - \alpha$.
- A $100(1 - \alpha)\%$ prediction interval for a future value of y when $\mathbf{x} = \mathbf{x}_0$ is given by

$$\hat{y}_0 \pm t_{\alpha/2, n-p} \hat{\sigma} \sqrt{1 + \mathbf{x}_0' (\mathbf{X}' \mathbf{X})^{-1} \mathbf{x}_0}$$

Simultaneous confidence intervals for β_j 's

- When making more than one confidence interval, it is important to ensure that the joint coverage of all of the intervals constructed is $\geq 100(1 - \alpha)$.
- If m intervals are constructed, with each interval having coverage probability $100(1 - \alpha/m)$, it follows from Bonferroni's inequality that the probability that the m intervals will simultaneously cover their associated parameters is at least $1 - \alpha$.
- For example, the $k + 1$ intervals

$$\hat{\beta}_j \pm t_{\alpha/(2(k+1)), n-1-k} \sqrt{MSE} \sqrt{C_{j,j}}, j = 0, 1, 2, \dots, k$$

have joint coverage of at least $100(1 - \alpha)$.

In general, if constructing m simultaneous confidence intervals, just replace α by α/m , and use the usual procedure for the CI.

Bonferroni's inequality (Not responsible for this material)

- Where E_1 and E_2 are events, and E^c is the complement of E , because the probability of the union is less than or equal to the sum of the probabilities, we know that

$$P(E_1^c \cup E_2^c) \leq P(E_1^c) + P(E_2^c)$$

- then use the fact that $P(E) = 1 - P(E^c)$, similarly $P(E_1^c) = 1 - P(E_1)$, and let $E = E_1^c \cup E_2^c$, from which $E^c = E_1 \cap E_2$. Putting all of this together it follows that

$$P(E_1 \cap E_2) = 1 - P(E_1^c \cup E_2^c) \geq 1 - (P(E_1^c) + P(E_2^c))$$

- Letting $P(E_1) = 1 - \alpha/2$ and $P(E_2) = 1 - \alpha/2$, it follows that

$$P(E_1 \cap E_2) \geq 1 - (\alpha/2 + \alpha/2) = 1 - \alpha$$

Application to simultaneous inference - take this as a given.

- Let E_1 be the event that β_j is contained in the interval with endpoints $\hat{\beta}_j \pm t_{(\alpha/2)/2, n-p} \sqrt{MSE} \sqrt{C_{j+1, j+1}}$. The probability of this is $1 - \alpha/2$.
- Let E_2 be the event that β_m is contained in the interval with endpoints $\hat{\beta}_m \pm t_{(\alpha/2)/2, n-p} \sqrt{MSE} \sqrt{C_{m+1, m+1}}$. The probability of this is $1 - \alpha/2$.
- It follows from Bonferroni's inequality that the probability that EACH of β_j and β_m is simultaneously contained in their associated intervals is at least $1 - \alpha$.
- We say that the two intervals form a joint $100(1 - \alpha)\%$ confidence region for (β_j, β_m) .
- What works for two the the β_j 's, works for an arbitrary number, hence the form for the $k + 1$ simultaneous intervals given above, where α is replaced by $\alpha/(k + 1)$.

Simultaneous confidence region for β

- A joint $100(1 - \alpha)\%$ confidence region for the vector β is given by those points β for which

$$(\beta - \hat{\beta})^T \mathbf{X}^T \mathbf{X} (\beta - \hat{\beta}) \leq p \text{MSE } F_{\alpha, p, n-p}$$

- This is a p dimensional ellipse centred at $\hat{\beta}$.