

$$Y = X\tilde{b} + \tilde{e}$$

$n \times 1$

$$\mathbb{E} \tilde{e} = \mathbf{0}$$

$$\text{Cov} \tilde{e} = \sigma^2 \mathbf{I}$$

uncorrelated

STAT 714 fa 2025 Lec 03

Gauss-Markov model, BLUE, Aitken model, generalized least-squares

Want to estimate $\tilde{e}^T \tilde{b}$
when $\tilde{e}$ is estimable.

Karl B. Gregory

University of South Carolina

$$\text{Var} \left(\tilde{e}^T \hat{\tilde{b}} \right) = ?$$

$$\tilde{x}_i = \begin{bmatrix} x_{i1} - \bar{x}_{i.} \\ \vdots \\ x_{in} - \bar{x}_{i.} \end{bmatrix}$$

$$\mathbf{1}^T \tilde{x}_i = 0$$

$$\sum_{j=1}^n (x_{ij} - \bar{x}_{i.}) = 0$$

These slides are an instructional aid; their sole purpose is to display, during the lecture, definitions, plots, results, etc. which take too much time to write by hand on the blackboard. They are not intended to explain or expound on any material.

- 1 Gauss-Markov model
- 2 Best linear unbiased estimator
- 3 Variance estimation
- 4 Underfitting and overfitting
- 5 Aitken model and generalized least squares
- 6 Best linear unbiased estimation in the restricted model

Gauss-Markov model

The *Gauss-Markov model* assumes $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$, where $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = (\sigma^2)\mathbf{I}_n$.

We consider the mean and variance of the LS estimator $\mathbf{c}^T \hat{\mathbf{b}}$ in this model.

We also consider how to estimate the error term variance σ^2 .

Result (Must know)

For a random vector $\mathbf{y}$, vectors $\mathbf{a}$ and $\mathbf{b}$, and a matrices $\mathbf{A}$ and $\mathbf{B}$, we have

$$① \mathbb{E} \mathbf{a}^T \mathbf{y} = \mathbf{a}^T \mathbb{E} \mathbf{y}. \quad \mathbb{E} \sum_{i=1}^n a_i y_i = \sum_{i=1}^n a_i \mathbb{E} y_i = \mathbf{a}^T \mathbb{E} \mathbf{y}$$

$$② \text{Var} \mathbf{a}^T \mathbf{y} = \mathbf{a}^T (\text{Cov} \mathbf{y}) \mathbf{a}.$$

$$③ \text{Cov}(\mathbf{a}^T \mathbf{y}, \mathbf{b}^T \mathbf{y}) = \mathbf{a}^T (\text{Cov} \mathbf{y}) \mathbf{b}.$$

$$④ \text{Cov} \mathbf{A} \mathbf{y} = \mathbf{A} (\text{Cov} \mathbf{y}) \mathbf{A}^T.$$

$$⑤ \text{Cov}(\mathbf{A} \mathbf{y}, \mathbf{B} \mathbf{y}) = \mathbf{A} (\text{Cov} \mathbf{y}) \mathbf{B}^T.$$

IP $\mathbf{a} \sim N(0, \mathbf{I}_n \sigma^2)$
 $\hat{\beta} \sim N(\mathbf{X}^T \mathbf{X}^{-1} \mathbf{X}^T \beta, \mathbf{X}^T \mathbf{X}^{-1} \sigma^2)$

Result (Variance of estimable contrast)

Let $\mathbf{c}^T \mathbf{b}$ be an estimable contrast in $\mathbf{y} = \mathbf{X} \mathbf{b} + \mathbf{e}$ with LS estimator $\mathbf{c}^T \hat{\mathbf{b}}$. Then

$$\text{Var} \mathbf{c}^T \hat{\mathbf{b}} = \sigma^2 \mathbf{c}^T (\mathbf{X}^T \mathbf{X})^{-} \mathbf{c}.$$

Exercise: Derive and show invariance to the choice of generalized inverse.

$$\text{Var } X = \mathbb{E}[(X - \mathbb{E}X)^2]$$

$$\begin{aligned} \text{Var } \tilde{a}^T \tilde{y} &= \text{Var} \left(\sum_{i=1}^n a_i y_i \right) \\ &= \mathbb{E} \left(\left(\sum_{i=1}^n a_i y_i - \sum_{i=1}^n a_i \mathbb{E} y_i \right)^2 \right) \\ &= \mathbb{E} \left(\sum_{i=1}^n a_i (y_i - \mathbb{E} y_i) \right)^2 \\ &= \mathbb{E} \left[\sum_{i=1}^n \sum_{j=1}^n a_i a_j (y_i - \mathbb{E} y_i) (y_j - \mathbb{E} y_j) \right] \\ &= \sum_{i=1}^n \sum_{j=1}^n a_i a_j \underbrace{\mathbb{E} (y_i - \mathbb{E} y_i) (y_j - \mathbb{E} y_j)}_{\text{Cov}(y_i, y_j)} \\ &= \sum_{i=1}^n \sum_{j=1}^n a_i a_j \text{Cov}(y_i, y_j) \\ &= \tilde{a}^T (\text{Cov } \tilde{y}) \tilde{a} \end{aligned}$$

let $\tilde{c}^T \tilde{b}$ be estimable. Means $\tilde{c} \in \text{Col } X^T$.

let $\hat{\tilde{b}}$ satisfy $X^T X \hat{\tilde{b}} = X^T y$.

Then $\text{Var} \left(\tilde{c}^T \hat{\tilde{b}} \right) = \tilde{c}^T (X^T X)^{-} \tilde{c} \sigma^2$

PROOF

Define all solutions to Normal equations as

$$\tilde{b}(\tilde{z}) = (X^T X)^{-} X^T y + \left(I - (X^T X)^{-} X^T X \right) \tilde{z}$$

$$\forall \tilde{z} \in \mathbb{R}^p$$

$$\hat{\underline{b}}_{\underline{z}}^T(\underline{z}) = \underline{c}^T \left[(\underline{x}^T \underline{x})^{-1} \underline{x}^T \underline{y} + (\underline{I} - (\underline{x}^T \underline{x})^{-1} \underline{x}^T \underline{x}) \underline{z} \right]$$

$$= \underline{d}^T \underline{x} \left[\dots \right]$$

Write $\underline{c} = \underline{x}^T \underline{d}$
for some $\underline{d}$.

$$= \underline{d}^T \left[\underline{x} (\underline{x}^T \underline{x})^{-1} \underline{x}^T \underline{y} + \underbrace{\underline{x} \underline{z} - \underline{x} (\underline{x}^T \underline{x})^{-1} \underline{x}^T \underline{x} \underline{z}}_{\underline{0}} \right]$$

$$= \underline{d}^T \underline{P}_x \underline{y}$$

$\Rightarrow$

$$\text{Var} \left(\hat{\underline{b}}_{\underline{z}}^T(\underline{z}) \right) = \text{Var} \left(\underline{d}^T \underline{P}_x \underline{y} \right)$$

$$\underline{y} = \underline{x} \underline{b} + \underline{e}$$

$$= \text{Var} \left((\underline{P}_x \underline{d})^T \underline{y} \right)$$

$$\text{Var} \left(\underline{a}^T \underline{y} \right) = \underline{a}^T (\text{Cov} \underline{y}) \underline{a}$$

$$= (\underline{P}_x \underline{d})^T (\text{Cov} \underline{y}) \underline{P}_x \underline{d}$$

$$= \underline{d}^T \underline{P}_x (\sigma^2 \underline{I}) \underline{P}_x \underline{d}$$

$$= \sigma^2 \underline{d}^T \underline{P}_x \underline{d}$$

$$= \sigma^2 \underline{d}^T \underline{x} (\underline{x}^T \underline{x})^{-1} \underline{x}^T \underline{d}$$

$$= \sigma^2 \underline{c}^T (\underline{x}^T \underline{x})^{-1} \underline{c}$$

For $\mathbf{c} \in \mathbb{R}^n$,

$$\text{Var}(\mathbf{c}^T \hat{\mathbf{b}}) = \mathbf{c}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{c} \sigma^2$$

$$\mathbb{E} \mathbf{c}^T \hat{\mathbf{b}} = \mathbf{c}^T \mathbf{b} \quad \forall \mathbf{b}$$

- 1 Gauss-Markov model
- 2 Best linear unbiased estimator
- 3 Variance estimation
- 4 Underfitting and overfitting
- 5 Aitken model and generalized least squares
- 6 Best linear unbiased estimation in the restricted model

Best linear unbiased estimator (BLUE)

A *best linear unbiased estimator* (BLUE) of an estimable contrast $\mathbf{c}^T \mathbf{b}$ is an estimator of the form $a_0 + \mathbf{a}^T \mathbf{y}$ such that $\mathbb{E}[a_0 + \mathbf{a}^T \mathbf{y}] = \mathbf{c}^T \mathbf{b}$ for all $\mathbf{b}$ and such that no other such estimator has smaller variance for any $\mathbf{b}$.

A BLUE achieves the smallest possible variance among linear unbiased estimators.

Theorem (Gauss-Markov Theorem)

Let $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$ with $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2 \mathbf{I}_n$ and let $\mathbf{c}^T \mathbf{b}$ be an estimable contrast. If $\hat{\mathbf{b}}$ satisfies $\mathbf{X}^T \mathbf{X} \hat{\mathbf{b}} = \mathbf{X}^T \mathbf{y}$ then $\mathbf{c}^T \hat{\mathbf{b}}$ is the BLUE for $\mathbf{c}^T \mathbf{b}$.

See Thm 4.1 of Monahan (2008).

Prove the result.

let $\underline{c}^T \underline{b}$ be estimable, so $\underline{c} \in \text{Col } X^T$

$$\Rightarrow \underline{c} = X^T \underline{d} \text{ for some } \underline{d}.$$

Suppose $\exists a_0, \underline{a}$ such that

$$E [a_0 + \underline{a}^T \underline{y}] = \underline{c}^T \underline{b} \quad \forall \underline{b}.$$

$$\Rightarrow a_0 + \underline{a}^T X \underline{b} = \underline{c}^T \underline{b} \quad \forall \underline{b}.$$

$$\Rightarrow a_0 = 0, \quad \underline{c}^T = \underline{a}^T X.$$

$$\Leftrightarrow \underline{c} = X^T \underline{a}$$

$$\text{Var}(\underline{a}^T \underline{y}) = \text{Var}(\underline{a}^T \underline{y} - \underline{c}^T \hat{\underline{b}} + \underline{c}^T \hat{\underline{b}})$$

$$= \text{Var}(\underline{a}^T \underline{y} - \underline{c}^T \hat{\underline{b}}) + \text{Var}(\underline{c}^T \hat{\underline{b}})$$

$$+ 2 \text{Cov}(\underline{a}^T \underline{y} - \underline{c}^T \hat{\underline{b}}, \underline{c}^T \hat{\underline{b}})$$

$$\text{Cov}(\underline{a}^T \underline{y} - \underline{c}^T \hat{\underline{b}}, \underline{c}^T \hat{\underline{b}})$$

$$= \text{Cov}(\underline{a}^T \underline{y} - \underline{a}^T X \hat{\underline{b}}, \underline{a}^T X \hat{\underline{b}})$$

$$= \text{Cov}(\underline{a}^T \underline{y} - \underline{a}^T P_X \underline{y}, \underline{a}^T P_X \underline{y})$$

$$= \text{Cov}([\underline{I} - P_X] \underline{a})^T \underline{y}, (P_X \underline{a})^T \underline{y})$$

$$\begin{aligned} \text{Cov}(\underline{a}^T \underline{y}, \underline{b}^T \underline{y}) \\ = \underline{a}^T (\text{Cov } \underline{y}) \underline{b} \end{aligned}$$

$$= \left[(\mathbf{I} - \mathbf{P}_X) \tilde{\mathbf{a}} \right]^T \begin{pmatrix} \mathbf{C} & \mathbf{v} \\ \tilde{\mathbf{y}} & \tilde{\mathbf{y}} \end{pmatrix} \mathbf{P}_X \tilde{\mathbf{a}}$$

$$= \sigma^2 \tilde{\mathbf{a}}^T (\mathbf{I} - \mathbf{P}_X) \mathbf{P}_X \tilde{\mathbf{a}}$$

$$= 0$$

$\Rightarrow$

$$\text{Var}(\tilde{\mathbf{a}}^T \tilde{\mathbf{y}}) = \underbrace{\text{Var}(\tilde{\mathbf{a}}^T \tilde{\mathbf{y}} - \tilde{\mathbf{a}}^T \hat{\tilde{\mathbf{b}}})}_{\neq 0} + \text{Var}(\tilde{\mathbf{a}}^T \hat{\tilde{\mathbf{b}}})$$

$$\Rightarrow \text{Var}(\tilde{\mathbf{a}}^T \tilde{\mathbf{y}}) \geq \text{Var}(\tilde{\mathbf{a}}^T \hat{\tilde{\mathbf{b}}}).$$

1 Gauss-Markov model

2 Best linear unbiased estimator

3 Variance estimation

$$\hat{\sigma}^2 =$$

4 Underfitting and overfitting

5 Aitken model and generalized least squares

6 Best linear unbiased estimation in the restricted model

Result (Expected value of a quadratic form)

Let $\mathbf{z}$ be a random vector with $\mathbb{E}\mathbf{z} = \boldsymbol{\mu}$ and $\text{Cov } \mathbf{z} = \boldsymbol{\Sigma}$. Then

$$\mathbb{E}\mathbf{z}^T \mathbf{A} \mathbf{z} = \boldsymbol{\mu}^T \mathbf{A} \boldsymbol{\mu} + \text{tr}(\mathbf{A} \boldsymbol{\Sigma}).$$

See Lem 4.1 of Monahan (2008).

Can be used to prove the following result.

Result (Unbiased estimator of the variance)

Let $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$ with $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2 \mathbf{I}_n$ and let $\mathbf{X}$ have rank r . Then

$$\hat{\sigma}^2 = \frac{\|\hat{\mathbf{e}}\|^2}{n - r}, \quad \text{where } \hat{\mathbf{e}} = (\mathbf{I} - \mathbf{P}_X)\mathbf{y},$$

is an unbiased estimator of σ^2 .

$$\begin{aligned} \tilde{\mathbf{y}} &= \hat{\tilde{\mathbf{y}}} + \hat{\tilde{\mathbf{e}}} \\ \hat{\tilde{\mathbf{y}}} &= \mathbf{P}_X \tilde{\mathbf{y}} = \mathbf{X} \hat{\tilde{\mathbf{b}}} \\ \hat{\tilde{\mathbf{e}}} &= (\mathbf{I} - \mathbf{P}_X) \tilde{\mathbf{y}} \end{aligned}$$

See Res 4.2 of Monahan (2008).

Prove both results.

For a random vector $\underline{z}$ with $\mathbb{E} \underline{z} = \underline{\mu}$, $\text{Cov} \underline{z} = \Sigma$,
we have

$$\mathbb{E} \underline{z}^T A \underline{z} = \underline{\mu}^T A \underline{\mu} + \text{tr}(A \Sigma).$$

Proof:

Write

$$\begin{aligned} \mathbb{E} \left[(\underline{z} - \underline{\mu})^T A (\underline{z} - \underline{\mu}) \right] &= \mathbb{E} \left[\underline{z}^T A \underline{z} + \underline{\mu}^T A \underline{\mu} - \overset{\downarrow}{2} \underline{z}^T A \underline{\mu} \right] \\ &= \mathbb{E} \left[\underline{z}^T A \underline{z} \right] - \underline{\mu}^T A \underline{\mu}. \end{aligned}$$

$\Rightarrow$

$$\mathbb{E} \underline{z}^T A \underline{z} = \mathbb{E} (\underline{z} - \underline{\mu})^T A (\underline{z} - \underline{\mu}) + \underline{\mu}^T A \underline{\mu}.$$

$$= \mathbb{E} \text{tr}(\underline{z} - \underline{\mu})^T A (\underline{z} - \underline{\mu}) + \underline{\mu}^T A \underline{\mu}.$$

$$= \mathbb{E} \text{tr} \left[(\underline{z} - \underline{\mu}) (\underline{z} - \underline{\mu})^T A \right] + \underline{\mu}^T A \underline{\mu}$$

$$= \text{tr} \left[\underbrace{\mathbb{E} (\underline{z} - \underline{\mu}) (\underline{z} - \underline{\mu})^T}_{n \times n} A \right] + \underline{\mu}^T A \underline{\mu}$$

$$= \text{tr}(\Sigma A) + \underline{\mu}^T A \underline{\mu}.$$

$$\begin{aligned} \text{tr}(AB) &= \text{tr}(BA) \end{aligned}$$

$$\mathbb{E} \hat{\sigma}^2 = \mathbb{E} \left[\frac{\|\hat{\underline{e}}\|^2}{n-r} \right]$$

$$\|\underline{a}\|^2 = \underline{a}^T \underline{a}$$

$$= \mathbb{E} \left\| (\mathbf{I} - \mathbf{P}_X) \underline{y} \right\|^2 / (n-r)$$

$$= \mathbb{E} \left[\underline{y}^T (\mathbf{I} - \mathbf{P}_X) \underline{y} \right] / (n-r)$$

$$= \left[(\mathbb{E} \underline{y})^T (\mathbf{I} - \mathbf{P}_X) \mathbb{E} \underline{y} + \text{tr} \left[\overset{\sigma^2 \mathbf{I}}{\downarrow} (\text{Cov} \underline{y}) (\mathbf{I} - \mathbf{P}_X) \right] \right] / (n-r)$$

$$= \left[\underbrace{(\underline{x} \underline{b})^T (\mathbf{I} - \mathbf{P}_X) \underline{x} \underline{b}}_{=0} + \sigma^2 \text{tr} (\mathbf{I} - \mathbf{P}_X) \right] / (n-r)$$

$$\begin{aligned} \mathbf{P}_X \underline{x} \underline{b} &= \frac{\sigma^2}{n-r} \text{tr} (\mathbf{I} - \mathbf{P}_X) \\ &= \sigma^2. \end{aligned}$$

$$\text{tr} (\mathbf{I} - \mathbf{P}_X) = \text{tr} \mathbf{I} - \text{tr} \mathbf{P}_X = n - \text{tr} \mathbf{P}_X = n-r.$$

$\mathbf{P}_X$: idempotent, so eigenvalues = 0 or 1

symm., so rank = # non zero eigenvalues

$\text{tr} \mathbf{P}_X$ = sum of eigenvalues = # nonzero

$\mathbf{P}_X$ projects onto its own column space, and this is

$$\text{Col } X. \quad \text{so} \quad \text{Col } P_X = \text{Col } X.$$

$$\Rightarrow \dim \text{Col } P_X = \dim \text{Col } X$$

$$\text{rank } P_X = \text{rank } X = r.$$

1 Gauss-Markov model

2 Best linear unbiased estimator

3 Variance estimation

4 Underfitting and overfitting



5 Aitken model and generalized least squares

6 Best linear unbiased estimation in the restricted model

A setup for “underfitting” or omitting important covariates

Suppose $\mathbf{y}$ is truly given by

$$\mathbf{y} = \mathbf{X}\mathbf{b} + \boldsymbol{\eta} + \mathbf{e}, \quad \mathbb{E}\mathbf{e} = 0, \quad \text{Cov } \mathbf{e} = \sigma^2 \mathbf{I}_n,$$

but one assumes the model $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$ and finds $\hat{\mathbf{b}}$ satisfying $\mathbf{X}^T \mathbf{X} \hat{\mathbf{b}} = \mathbf{X}^T \mathbf{y}$.

So $\boldsymbol{\eta}$ represents the effects of omitted covariates.

Exercise: In the above setup find the expected value of

- ① $\mathbf{c}^T \hat{\mathbf{b}}$, where $\mathbf{c} \in \text{Col } \mathbf{X}^T$
- ② $\hat{\sigma}^2 = \|\mathbf{y} - \mathbf{X}\hat{\mathbf{b}}\|^2 / (n - r)$, where $r = \text{rank } \mathbf{X}$.

In each case give a condition on $\boldsymbol{\eta}$ under which the estimator will be unbiased.

LOOK AT O.Y.O. *on your own.*

A setup for “overfitting” or including unimportant covariates

Suppose $\mathbf{y}$ is truly given by

$$\mathbf{y} = \mathbf{X}_1 \mathbf{b}_1 + \mathbf{e}, \quad \mathbb{E} \mathbf{e} = 0, \quad \text{Cov } \mathbf{e} = \sigma^2 \mathbf{I}_n,$$

but one assumes $\mathbf{y} = \mathbf{X}_1 \mathbf{b}_1 + \mathbf{X}_2 \mathbf{b}_2 + \mathbf{e}$ and finds $\hat{\mathbf{b}} = [\mathbf{b}_1^T \ \mathbf{b}_2^T]^T$ satisfying

$$\mathbf{X}^T \mathbf{X} \hat{\mathbf{b}} = \mathbf{X}^T \mathbf{y}, \quad \text{where} \quad \mathbf{X} = [\mathbf{X}_1 \ \mathbf{X}_2].$$

In truth $\mathbf{b}_2 = \mathbf{0}$, so $\mathbf{X}_2$ contains unimportant (extra) covariate information.

Exercise: In the above setup, assume (for simplicity) that $\mathbf{X}$ is full-rank. Then:

- ① Find $\mathbb{E} \mathbf{c}_1^T \hat{\mathbf{b}}_1$ and $\text{Var } \mathbf{c}_1^T \hat{\mathbf{b}}_1$
- ② In what situation does adding “extra” covariates have no effect on $\text{Var } \mathbf{c}_1^T \hat{\mathbf{b}}_1$?
- ③ Find $\mathbb{E} \|(\mathbf{I} - \mathbf{P}_{\mathbf{X}}) \mathbf{y}\|^2 / (N - \text{rank } \mathbf{X})$ and $\mathbb{E} \|(\mathbf{I} - \mathbf{P}_{\mathbf{X}_1}) \mathbf{y}\|^2 / (N - \text{rank } \mathbf{X}_1)$.

G-M:

$$y_{\sim} = x_{\sim} b_{\sim} + e_{\sim} \quad , \quad \mathbb{E} e_{\sim} = 0 \quad , \quad \boxed{\text{Cov } e_{\sim} = \sigma^2 I_n}$$

1 Gauss-Markov model

Aitken :

$$y_{\sim} = x_{\sim} b_{\sim} + e_{\sim} \quad , \quad \mathbb{E} e_{\sim} = 0 \quad , \quad \text{Cov } e_{\sim} = \sqrt{\sigma^2}$$

2 Best linear unbiased estimator

3 Variance estimation

4 Underfitting and overfitting

5 Aitken model and generalized least squares

6 Best linear unbiased estimation in the restricted model

Aitken model

The *Aitken model* assumes $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$, where $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2 \mathbf{V}$, $\mathbf{V}$ pd.

The ordinary least-squares estimator $\mathbf{c}^T \mathbf{b}$ may not be the BLUE anymore...

Generalized least-squares estimator of an estimable contrast

Under the Aitken model, the *GLS estimator* of a contrast is $\mathbf{c}^T \hat{\mathbf{b}}_{\text{GLS}}$, where $\hat{\mathbf{b}}_{\text{GLS}}$ is any vector such that $(\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X}) \hat{\mathbf{b}}_{\text{GLS}} = \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y}$.

Exercise: Show that $\hat{\mathbf{b}}_{\text{GLS}}$ minimizes $Q_{\mathbf{V}}(\mathbf{b}) = (\mathbf{y} - \mathbf{X}\mathbf{b})^T \mathbf{V}^{-1} (\mathbf{y} - \mathbf{X}\mathbf{b})$.

Generalized Normal $\mathbb{E}_{\mathbf{gls}}$

$$Q(\mathbf{b}) = (\mathbf{y} - \mathbf{X}\mathbf{b})^T (\mathbf{y} - \mathbf{X}\mathbf{b})$$

$$Q_V(\tilde{b}) = \tilde{y}^T V^{-1} \tilde{y} + \tilde{b}^T X^T V^{-1} X \tilde{b} - 2 \tilde{y}^T V^{-1} X \tilde{b}$$

$$\frac{\partial}{\partial \tilde{b}} Q_V(\tilde{b}) = 2 X^T V^{-1} X \tilde{b} - 2 X^T V^{-1} \tilde{y} \stackrel{\text{set}}{=} 0$$

$\Rightarrow$

$$X^T V^{-1} X \tilde{b} = X^T V^{-1} \tilde{y}$$

gen. Norm. eqs.

We will need these results for establishing some properties of the GLS estimator.

Result ("Generalized cool result" and another gen. inverse of $\mathbf{X}$)

- 1 If $\mathbf{V}$ is positive definite there exists nonsingular $\mathbf{R}$ such that $\mathbf{R}\mathbf{R} = \mathbf{V}^{-1}$.
Symmetric
- 2 $\text{Nul } \mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} = \text{Nul } \mathbf{X}$. ✓
- 3 $\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} \mathbf{A} = \mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} \mathbf{B} \iff \mathbf{X} \mathbf{A} = \mathbf{X} \mathbf{B}$
- 4 $(\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{V}^{-1}$ is a generalized inverse of $\mathbf{X}$.

Prove the results.

① $\mathbf{V}$ p.d. $\Rightarrow \exists \mathbf{R}$ nonsing, symm, st. $\mathbf{R}\mathbf{R} = \mathbf{V}^{-1}$.

proof $\mathbf{V}$ p.d. $\Rightarrow \mathbf{V} = \mathbf{P} \mathbf{D} \mathbf{P}^T$ where $\mathbf{P}^T = \mathbf{P}^{-1}$ s.t. $\mathbf{P}^T \mathbf{P} = \mathbf{P} \mathbf{P}^T = \mathbf{I}$.
 $\mathbf{D}$ diag. with positive diagonals.

Why do we want Cov. matrices to be p.d.?

Suppose
 $\text{Cov } \underline{y} = V$.

Then

$$\text{Var}(\underline{a}^T \underline{y}) = \underline{a}^T V \underline{a} \stackrel{\text{want}}{> 0}$$

$$D = \begin{bmatrix} d_1 & & \\ & \ddots & \\ & & d_n \end{bmatrix}$$

$$D^{-1} = \begin{bmatrix} \frac{1}{d_1} & & \\ & \ddots & \\ & & \frac{1}{d_n} \end{bmatrix}$$

$$R^{-1} = P D^{-1/2} P^T$$

$$\text{because } P D^{1/2} P^T P D^{-1/2} P^T = P D^{1/2} D^{-1/2} P^T = P P^T = I.$$

$$(2) \quad \text{Nul } X^T V^{-1} X = \text{Nul } X$$

$$\text{"C"} \quad \underline{x} \in \text{Nul } X^T V^{-1} X$$

$$\Rightarrow X^T V^{-1} X \underline{x} = \underline{0}$$

$$\Rightarrow \underline{x}^T X^T V^{-1} X \underline{x} = 0$$

$$\begin{aligned} P D^{-1} P^T V &= P D^{-1} \overbrace{P^T P}^I D P^T \\ &= P D^{-1} D P^T \\ &= P P^T \overbrace{I}^I \\ &= I \end{aligned}$$

↑
take reciprocal of diag

$\Rightarrow$

$$V^{-1} = P D^{-1} P^T$$

$$= P D^{-1/2} D^{-1/2} P^T$$

$$= \underbrace{P D^{-1/2} P^T P D^{-1/2} P^T}_I$$

$$= R R$$

$$D^{-1/2} = \begin{bmatrix} \frac{1}{\sqrt{d_1}} & & \\ & \ddots & \\ & & \frac{1}{\sqrt{d_n}} \end{bmatrix}$$

$$R = P D^{-1/2} P^T$$

symmetric, nonsing.,
 square root matrix
 for V^{-1} .

$$\left(\begin{array}{c} \text{Prev.} \\ \text{Nul } X^T X = \text{Nul } X \end{array} \right)$$

$$\Rightarrow \underline{x}^T X^T \overset{R \text{ symm}}{P^T P} X \underline{x} = 0$$

$$\Rightarrow \|PX \underline{x}\|^2 = 0$$

$$\Rightarrow PX \underline{x} = 0$$

(multiply by P^{-1})

$$\Rightarrow X \underline{x} = 0$$

$$\Rightarrow \underline{x} \in \text{Nul } X.$$

" $\supset$ "

$$\underline{x} \in \text{Nul } X \Rightarrow X \underline{x} = \underline{0}$$

$$\Rightarrow X^T V^{-1} X \underline{x} = \underline{0}.$$

$$\Rightarrow \underline{x} \in \text{Nul } X^T V^{-1} X.$$

③

$$X^T V^{-1} X A = X^T V^{-1} X B$$

$$\Leftrightarrow X^T V^{-1} X (A - B) = 0$$

$$\Leftrightarrow \text{Col of } A - B \in \text{Nul } X^T V^{-1} X = \text{Nul } X$$

$$\Leftrightarrow X(A - B) = 0$$

$$\Leftrightarrow XA = XB.$$

4 $(X^T V^{-1} X)^{-1} X^T V^{-1}$ is a generalized inverse of X .

$$(4) \quad X^T V^{-1} X \underbrace{(X^T V^{-1} X)^{-1} X^T V^{-1} X}_A = X^T V^{-1} X \underbrace{I}_B$$

$$\Rightarrow X \underbrace{(X^T V^{-1} X)^{-1} X^T V^{-1}}_{\text{gen-inv of } X} X = X$$

ALSO:

$$Nul \ X^T V^{-1} X = Nul \ X$$

$$\Leftrightarrow (Nul \ X^T V^{-1} X)^\perp = (Nul \ X)^\perp$$

$$\underbrace{Col \ X^T V^{-1} X}_{''} = Col \ X^T$$

$$(Col \ X)^\perp = Nul \ X^T$$

$$(Col \ X^T)^\perp = Nul \ X$$

$$Col \ X^T = (Nul \ X)^\perp$$

$$Col \ X^T V^{-1} X = Col \ X^T$$

Note: Estimability of $\mathbf{c}^T \mathbf{b}$ is still equivalent to $\mathbf{c} \in \text{Col } \mathbf{X}^T$.

Result (Properties of the GLS estimator of an estimable contrast)

Let $\mathbf{c}^T \mathbf{b}$ be an estimable contrast. Then the GLS estimator $\mathbf{c}^T \hat{\mathbf{b}}_{\text{GLS}}$

- 1 is invariant to the choice of $\hat{\mathbf{b}}_{\text{GLS}}$ which satisfies $(\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X}) \hat{\mathbf{b}}_{\text{GLS}} = \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y}$. ✓
- 2 has expected value equal to $\mathbf{c}^T \mathbf{b}$ for all $\mathbf{b}$.
- 3 has variance $\text{Var } \mathbf{c}^T \hat{\mathbf{b}}_{\text{GLS}} = \mathbf{c}^T (\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{c} \sigma^2$

Prove the results.

① $\mathbf{c}^T \mathbf{b}$ estimable $\Rightarrow \mathbf{c} \in \text{Col } \mathbf{X}^T \Rightarrow \mathbf{c} = \mathbf{X}^T \mathbf{d}$ for some $\mathbf{d}$.

let $\hat{\mathbf{b}}_1$ and $\hat{\mathbf{b}}_2$ solve the gen. norm. eqs.

I.e.

$$(X^T V^{-1} X) \hat{b}_{\tilde{1}} = X^T V^{-1} y_{\tilde{}}$$

$$(X^T V^{-1} X) \hat{b}_{\tilde{2}} = X^T V^{-1} y_{\tilde{}}$$

$$\Rightarrow (X^T V^{-1} X) (\hat{b}_{\tilde{1}} - \hat{b}_{\tilde{2}}) = \mathbf{0}_{\tilde{}}$$

$$\Rightarrow \hat{b}_{\tilde{1}} - \hat{b}_{\tilde{2}} \in \text{Nul } X^T V^{-1} X = \text{Nul } X$$

$$\Rightarrow X (\hat{b}_{\tilde{1}} - \hat{b}_{\tilde{2}}) = \mathbf{0}_{\tilde{}}$$

$$\Rightarrow X \hat{b}_{\tilde{1}} = X \hat{b}_{\tilde{2}}.$$

$$\text{Now } \underbrace{\varepsilon_{\tilde{}}^T \hat{b}_{\tilde{1}}} = d_{\tilde{}}^T X \hat{b}_{\tilde{1}} = d_{\tilde{}}^T X \hat{b}_{\tilde{2}} = \underbrace{\varepsilon_{\tilde{}}^T \hat{b}_{\tilde{2}}}.$$

②

$$\begin{aligned} \mathbb{E} \underbrace{\varepsilon_{\tilde{}}^T \hat{b}_{\tilde{1}}}_{\text{}} &= \mathbb{E} d_{\tilde{}}^T X \hat{b}_{\tilde{1}} \\ &= \mathbb{E} d_{\tilde{}}^T X (X^T V^{-1} X)^{-1} \underbrace{X^T V^{-1} X \hat{b}_{\tilde{1}}}_{= X^T V^{-1} y_{\tilde{}}} \\ &= \mathbb{E} \underbrace{d_{\tilde{}}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} y_{\tilde{}}}_{\text{}} \\ &= d_{\tilde{}}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} X \hat{b}_{\tilde{1}} \\ &= d_{\tilde{}}^T X \hat{b}_{\tilde{1}} \end{aligned}$$

$$= \tilde{c}^T \tilde{b}$$

③

$$V_r \left(\tilde{c}^T \hat{\tilde{b}}_{gls} \right) = V_r \left(\tilde{d}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} \underline{y} \right)$$

assume
symmetric

$$= \tilde{d}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} (\cancel{V} \sigma^2) \cancel{V^{-1}} X (X^T V^{-1} X)^{-1} X^T \underline{d}$$

$$= \sigma^2 \tilde{d}^T X \boxed{(X^T V^{-1} X)^{-1} X^T V^{-1}} X (X^T V^{-1} X)^{-1} X^T \underline{d}$$

$$= \sigma^2 \tilde{d}^T X (X^T V^{-1} X)^{-1} X^T \underline{d}$$

$$= \sigma^2 \tilde{c}^T (X^T V^{-1} X)^{-1} \underline{c} .$$

Theorem (Aitken's Theorem)

Let $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$, where $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2\mathbf{V}$, $\mathbf{V}$ pd, and let $\mathbf{c}^T\mathbf{b}$ be an estimable contrast. Then $\mathbf{c}^T\hat{\mathbf{b}}_{\text{GLS}}$ is the BLUE for $\mathbf{c}^T\mathbf{b}$.

See Thm 4.2 of Monahan (2008).

Prove the result.

$$\begin{aligned} \text{let } a_0 + \tilde{\mathbf{z}}^T \mathbf{y} & \text{ be a L.U.E. so} \\ \mathbb{E} [a_0 + \tilde{\mathbf{z}}^T \mathbf{y}] &= \mathbf{c}^T \mathbf{b} \quad \forall \mathbf{b}. \\ \Rightarrow a_0 + \tilde{\mathbf{z}}^T \mathbf{X} \mathbf{b} &= \mathbf{c}^T \mathbf{b} \quad \forall \mathbf{b} \end{aligned}$$

$$\Rightarrow a_0 = 0, \quad \tilde{a}^T X = \tilde{\varepsilon}^T$$

Now, write

$$\begin{aligned} \text{Var}(X+Y) &= \text{Var} X \\ &\quad + \text{Var} Y \\ &\quad + 2 \text{Cov}(X, Y) \end{aligned}$$

$$\begin{aligned} \underline{\underline{\text{Var}(\tilde{a}^T \tilde{y})}} &= \text{Var}\left(\tilde{a}^T \tilde{y} - \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x} + \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right) \\ &= \underbrace{\text{Var}\left(\tilde{a}^T \tilde{y} - \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right)}_{=0} + \underline{\underline{\text{Var}\left(\tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right)}} \\ &\quad + 2 \text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}, \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right) \end{aligned}$$

Sol. to show $\text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}, \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right) = 0.$

$$\text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}, \tilde{\varepsilon}^T \hat{\tilde{b}}_{y|x}\right)$$

$$\tilde{\varepsilon}^T = \tilde{a}^T X$$

$$= \text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{a}^T X \hat{\tilde{b}}_{y|x}, \tilde{a}^T X \hat{\tilde{b}}_{y|x}\right)$$

$$= \text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{a}^T X \underbrace{(X^T V^{-1} X)^{-1} X^T V^{-1}}_{X^T V^{-1} \tilde{y}} X \hat{\tilde{b}}_{y|x}, \tilde{a}^T X \underbrace{(X^T V^{-1} X)^{-1} X^T V^{-1}}_{X^T V^{-1} \tilde{y}} X \hat{\tilde{b}}_{y|x}\right)$$

$$= \text{Cov}\left(\tilde{a}^T \tilde{y} - \tilde{a}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} \tilde{y}, \tilde{a}^T X (X^T V^{-1} X)^{-1} X^T V^{-1} \tilde{y}\right)$$

$$= \text{Cov}\left(\tilde{a}^T R^{-1} R \tilde{y} - \tilde{a}^T R^{-1} R X (X^T R R X)^{-1} X^T R R \tilde{y},\right.$$

$$\left. \tilde{a}^T R^{-1} R X (X^T R R X)^{-1} X^T R R \tilde{y}\right)$$

$$V^{-1} = R R$$

$$= \text{Cov} \left(\tilde{a}^T R^{-1} \tilde{z} - \tilde{a}^T R^{-1} U (U^T U)^{-1} U^T \tilde{z}, \right. \\ \left. \tilde{a}^T R^{-1} U (U^T U)^{-1} U^T \tilde{z} \right)$$

$$y = X b + e$$

$$R y = R X b + R e$$

$$\tilde{z} = U b + \tilde{\epsilon}$$

$$= \text{Cov} \left(\tilde{a}^T R^{-1} (I - P_U) \tilde{z}, \tilde{a}^T R^{-1} P_U \tilde{z} \right)$$

$$\tilde{z} = R y, U = R X, \tilde{\epsilon} = R e.$$

$$= \tilde{a}^T R^{-1} (I - P_U) \overset{\sigma^2 I}{\text{Cov}(\tilde{z})} P_U R^{-1} \tilde{a}$$

$$y = R^{-1} \tilde{z} \quad X = R^{-1} U$$

$$= \tilde{a}^T R^{-1} \underbrace{(I - P_U) P_U}_{0} R^{-1} \tilde{a} \sigma^2$$

$$P_U = U (U^T U)^{-1} U^T$$

$$\text{Cov}(\tilde{z}) = \text{Cov}(R y)$$

$$= 0.$$

$$= R \text{Cov}(y) R$$

$$= R V R \sigma^2$$

$$= R (V^{-1})^{-1} R \sigma^2$$

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

$$= I \sigma^2$$

Result (Unbiased estimator of the variance in the Aitken model)

Let $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$ with $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2\mathbf{V}$ and let $\mathbf{X}$ have rank r . Then

$$\hat{\sigma}_{\text{gls}}^2 = \frac{1}{n-r} (\mathbf{y} - \mathbf{X}\hat{\mathbf{b}}_{\text{gls}})^T \mathbf{V}^{-1} (\mathbf{y} - \mathbf{X}\hat{\mathbf{b}}_{\text{gls}})$$

is an unbiased estimator of σ^2 .

Prove the result.

$$\begin{aligned} & \mathbb{E} \hat{\sigma}_{\text{gls}}^2 \\ &= \frac{1}{n-r} \mathbb{E} \left[\left(\mathbf{y} - \mathbf{X} \left(\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} \right)^{-1} \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y} \right)^T \mathbf{V}^{-1} \left(\mathbf{y} - \mathbf{X} \left(\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} \right)^{-1} \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y} \right) \right] \end{aligned}$$

$\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X} \hat{\mathbf{b}}_{\text{gls}} = \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y}$

$$= \frac{1}{n-r} \mathbb{E} \left[\left(\tilde{y} - x(\underbrace{x^T V^{-1} x}_{RR})^{-1} \underbrace{x^T V^{-1} y}_{RR} \right)^T \underbrace{V^{-1}}_{RR} \left(\tilde{y} - x(\underbrace{x^T V^{-1} x}_{RR})^{-1} \underbrace{x^T V^{-1} y}_{RR} \right) \right]$$

$$= \frac{1}{n-r} \mathbb{E} \left[\left(R^{-1} \tilde{z} - \underbrace{R^{-1} U (U^T U)^{-1} U^T}_{P_U} \tilde{z} \right)^T R R \left(R^{-1} \tilde{z} - \underbrace{R^{-1} U (U^T U)^{-1} U^T}_{P_U} \tilde{z} \right) \right]$$

$$= \frac{1}{n-r} \mathbb{E} \left[\left(\tilde{z} - P_U \tilde{z} \right)^T \left(\tilde{z} - P_U \tilde{z} \right) \right]$$

$$= \frac{1}{n-r} \mathbb{E} \left[\tilde{z}^T (I - P_U) \tilde{z} \right]$$

$$= \frac{1}{n-r} \left[\underbrace{(U b_{\tilde{z}})^T (I - P_U) U b_{\tilde{z}}}_0 + \text{tr}((I - P_U) \sigma^2) \right]$$

$$= \frac{\sigma^2}{n-r} \text{tr}(I - P_U)$$

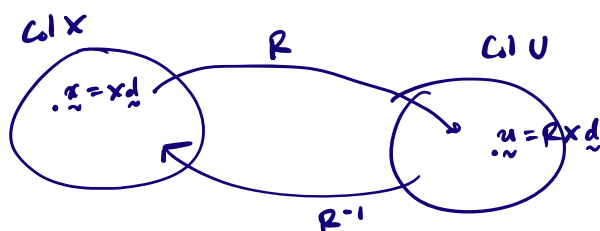
$$= \frac{\sigma^2}{n-r} \left[n - \text{tr}(P_U) \right]$$

$$= \text{rank } P_U = \text{rank } U$$

$$= \frac{\sigma^2}{n-r} (n - r)$$

We know $\text{rank } U = \text{rank } X$ because the transformation R which takes vectors in $\text{Col } X$ and returns vectors in $\text{Col } U$ is 1:1.

$$\boxed{= \sigma^2}$$



$$\mathbb{E} \tilde{z}^T A \tilde{z} = (\mathbb{E} \tilde{z})^T A \mathbb{E} \tilde{z} + \text{tr}(A \text{Cov } \tilde{z})$$

$$V^{-1} = R R$$

$$\tilde{z} = R \tilde{y} \Leftrightarrow \tilde{y} = R^{-1} \tilde{z}$$

$$U = R X \Leftrightarrow X = R^{-1} U$$

$$\begin{aligned} X^T V^{-1} X &= X^T R R X \\ &= (R X)^T R X \\ &= U^T U \end{aligned}$$

$$\mathbb{E} \tilde{z} = \mathbb{E} R \tilde{y}$$

$$= R X b_{\tilde{y}}$$

$$= U b_{\tilde{y}}$$

$$\text{Cov } \tilde{z} = \text{Cov}(R \tilde{y})$$

$$= R \text{Cov } \tilde{y} R$$

$$= R V R \sigma^2$$

$$(V^{-1} = R R)$$

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

$$= I \sigma^2$$

Let $\underline{x} = X \underline{d}$ some $\underline{d}$. Then $\underline{y} = R \underline{x} = R X \underline{d} = U \underline{d} \in \text{Col } U$.

Also, $R' \underline{y} = \underline{x}$.

Exercise: Let $Y_i = x_i \beta + \varepsilon_i |x_i|$, $i = 1, \dots, n$, where $\varepsilon_1, \dots, \varepsilon_n \stackrel{\text{ind}}{\sim} \text{Normal}(0, \sigma^2)$.

1 Write the model in matrix form as an Aitken model.

2 Give $\hat{\beta}_{\text{gls}}$ as well as $\hat{\beta}_{\text{ols}}$.

~~3 Give $\text{E} \hat{\beta}_{\text{gls}}$ as well as $\text{E} \hat{\beta}_{\text{ols}}$.~~

~~4 Give $\text{Var} \hat{\beta}_{\text{gls}}$ as well as $\text{Var} \hat{\beta}_{\text{ols}}$.~~

5 Give $\text{E} \hat{\beta}_{\text{gls}}$ as well as $\text{E} \hat{\beta}_{\text{ols}}$.

— both unbiased

6 Give $\text{Var} \hat{\beta}_{\text{gls}}$ as well as $\text{Var} \hat{\beta}_{\text{ols}}$.

7 Show that $\text{Var} \hat{\beta}_{\text{ols}} \geq \text{Var} \hat{\beta}_{\text{gls}}$ using the Cauchy-Schwarz inequality.

①

$$\underline{y} = X \underline{\beta} + \underline{e}$$

$$\text{E} \underline{e} = \underline{0}$$

$$\text{Cov} \underline{e} = \sigma^2 V$$

$$\tilde{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \quad X = \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}$$

$$\tilde{b} = \beta$$

$$\tilde{e} = \begin{bmatrix} |x_1| \varepsilon_1 \\ \vdots \\ |x_n| \varepsilon_n \end{bmatrix}$$

$$\text{Cov } \tilde{e} = \sigma^2 \begin{bmatrix} x_1^2 & & \\ & x_2^2 & \\ & & \ddots \\ & & & x_n^2 \end{bmatrix}$$

$n \times n$
✓

②

$$\hat{\beta}_{OLS} = (X^T X)^{-1} X^T \tilde{y} =$$

$$\frac{\sum_{i=1}^n x_i y_i}{\sum_{i=1}^n x_i^2}$$

$$(X^T V^{-1} X) \hat{\beta}_{GLS} = X^T V^{-1} \tilde{y}$$

$$X^T V^{-1} X = [x_1 \dots x_n] \begin{bmatrix} \frac{1}{x_1^2} & & \\ & \ddots & \\ & & \frac{1}{x_n^2} \end{bmatrix} \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} = n$$

$$X^T V^{-1} \tilde{y} = [x_1 \dots x_n] \begin{bmatrix} \frac{1}{x_1^2} & & \\ & \ddots & \\ & & \frac{1}{x_n^2} \end{bmatrix} \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \sum_{i=1}^n \frac{x_i y_i}{x_i^2} = \sum_{i=1}^n \frac{y_i}{x_i}$$

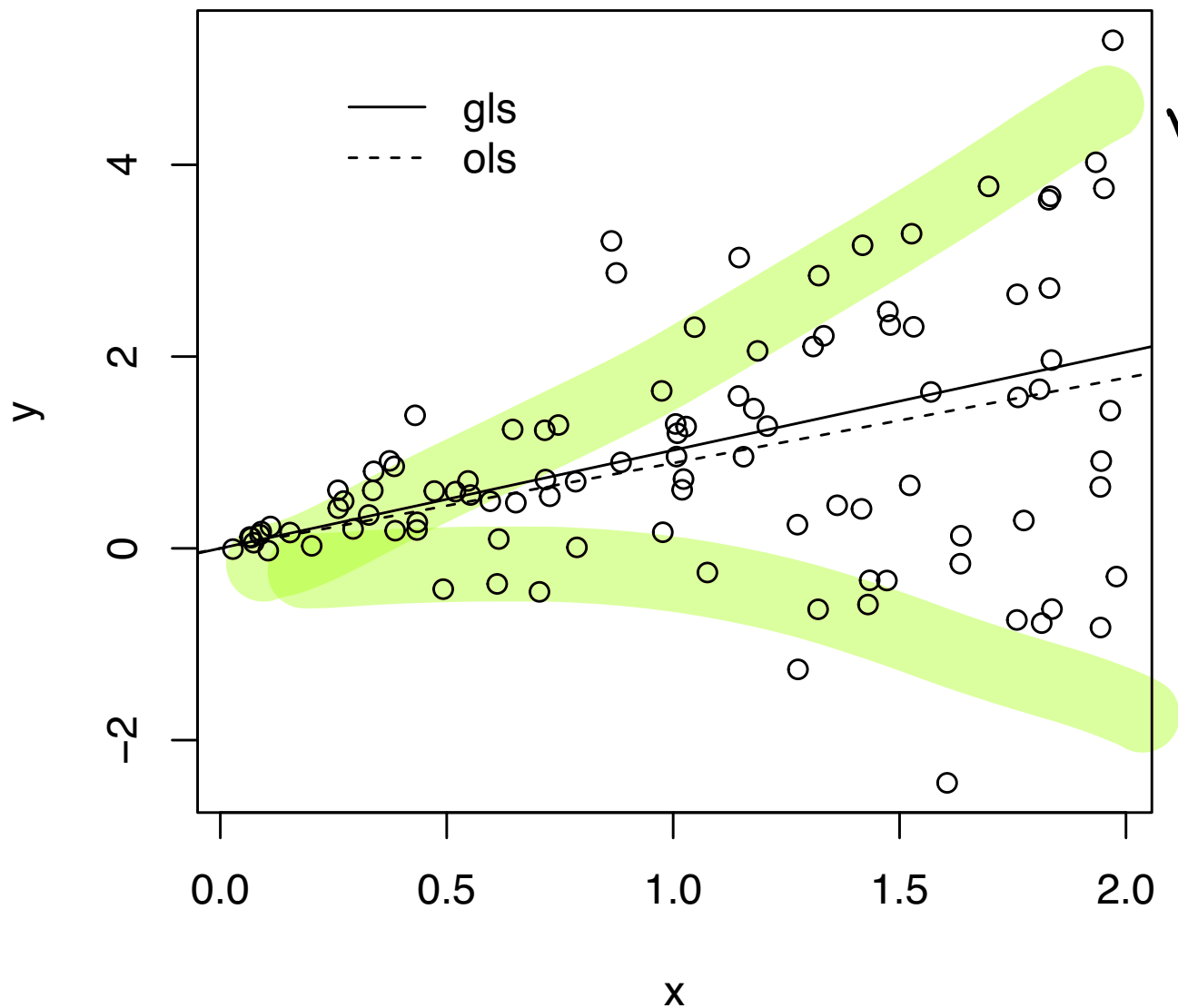
$$\hat{\beta}_{GLS} = (X^T V^{-1} X)^{-1} X^T V^{-1} \tilde{y} = \frac{1}{n} \sum_{i=1}^n \frac{y_i}{x_i}$$

$$\underline{y_i = x_i \beta + |x_i| \varepsilon_i}$$

~~$$\text{Cov } e_i = \sigma^2 I$$~~

$$\text{Cov } e_i = \sigma^2 V$$

$$V = \begin{bmatrix} x_1^2 & & \\ & \ddots & \\ & & x_n^2 \end{bmatrix}$$



$$\text{Var } \hat{\beta}_{OLS} = \text{Var} \left(\frac{\sum_{i=1}^n x_i y_i}{\sum_{i=1}^n x_i^2} \right)$$

$$\varepsilon_1, \dots, \varepsilon_n \stackrel{\text{i.i.d.}}{\sim} \text{Normal}(0, \sigma^2)$$

$$y_i = x_i \beta + |x_i| \varepsilon_i$$

$$\text{Var } y_i = x_i^2 \sigma^2$$

$$= \frac{1}{\left(\sum_{i=1}^n x_i^2 \right)^2} \sum_{i=1}^n x_i^2 \text{Var } y_i$$

$$= \sigma^2 \frac{\sum_{i=1}^n x_i^4}{\left(\sum_{i=1}^n x_i^2 \right)^2}$$

$$\text{Var } \hat{\beta}_{GLS} = \text{Var} \left(\frac{1}{n} \sum_{i=1}^n \frac{y_i}{x_i} \right)$$

Theory:

$$\text{Var } \hat{\beta}_{OLS} \geq \text{Var } \hat{\beta}_{GLS}$$

$$= \frac{1}{n^2} \sum_{i=1}^n \frac{1}{x_i^2} \text{Var } y_i$$

$$= \frac{1}{n^2} \sum_{i=1}^n \frac{x_i^2}{x_i^2} \sigma^2$$

$$\underline{a}^T \underline{b} \leq \|\underline{a}\| \|\underline{b}\|$$

$$= \frac{1}{n} \sigma^2$$

$$\sum_{i=1}^n x_i^2 = \sum_{i=1}^n 1 \cdot x_i^2$$

$$\leq \sqrt{\sum_{i=1}^n 1^2} \sqrt{\sum_{i=1}^n x_i^4}$$

$$\frac{\sum_{i=1}^n x_i^4}{\left(\sum_{i=1}^n x_i^2 \right)^2} \geq \frac{1}{n}$$

$$\Rightarrow \left(\sum_{i=1}^n x_i^2 \right)^2 \leq n \sum_{i=1}^n x_i^4$$

- 1 Gauss-Markov model
- 2 Best linear unbiased estimator
- 3 Variance estimation
- 4 Underfitting and overfitting
- 5 Aitken model and generalized least squares
- 6 Best linear unbiased estimation in the restricted model

Go back to the Gauss-Markov model, impose $\mathbf{P}^T \mathbf{b} = \boldsymbol{\delta}$, and consider BLUE...

Theorem (Best linear unbiased estimator in the restricted G-M model)

Let $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e}$, where $\mathbb{E}\mathbf{e} = \mathbf{0}$ and $\text{Cov } \mathbf{e} = \sigma^2 \mathbf{I}_n$, with the restriction $\mathbf{P}^T \mathbf{b} = \boldsymbol{\delta}$.

Let $\hat{\mathbf{b}}_H$ and $\hat{\mathbf{u}}$ satisfy $\begin{bmatrix} \mathbf{X}^T \mathbf{X} & \mathbf{P} \\ \mathbf{P}^T & \mathbf{0} \end{bmatrix} \begin{bmatrix} \hat{\mathbf{b}}_H \\ \hat{\mathbf{u}} \end{bmatrix} = \begin{bmatrix} \mathbf{X}^T \mathbf{y} \\ \boldsymbol{\delta} \end{bmatrix}$ and let $\mathbf{c}^T \mathbf{b}$ be estimable.

Then $\mathbf{c}^T \hat{\mathbf{b}}_H$ is the BLUE for $\mathbf{c}^T \mathbf{b}$ in the restricted model.

See Res 4.5 of Monahan (2008).

Prove the above in the steps:

- 1 Show consistency of the equations $\begin{bmatrix} \mathbf{X}^T \mathbf{X} & \mathbf{P} \\ \mathbf{P}^T & \mathbf{0} \end{bmatrix} \begin{bmatrix} \mathbf{v}_1 \\ \hat{\mathbf{v}}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{c} \\ \mathbf{0} \end{bmatrix}$. Lemma 4.2 of M.
- 2 Show $\mathbb{E} \mathbf{c}^T \hat{\mathbf{b}}_H = \mathbf{c}^T \mathbf{b}$ for all $\mathbf{b}$. Lemma 4.3 of M.
- 3 Show that any unbiased estimator $a_0 + \mathbf{a}^T \mathbf{y}$ has variance $\geq \text{Var } \mathbf{c}^T \hat{\mathbf{b}}_H$.

Monahan, J. F. (2008). *A primer on linear models*. CRC Press.