

# Econ 271/Mgtecon 604 - Section 1

Amar Venugopal

Stanford University

Winter 2026

1. General Info & Logistics
2. Statistical Background
3. Identification & Estimation
4. Properties of Estimators
  - Finite Sample
  - Large Sample
5. Practice Problems

1. General Info & Logistics

2. Statistical Background

3. Identification & Estimation

4. Properties of Estimators

- Finite Sample
- Large Sample

5. Practice Problems

- Email: `amar.venugopal@stanford.edu`
- Section: Fridays 1:30-3:20pm in 200-303
- Office Hours: Tuesdays 12pm-1pm, Wednesdays 3:30-4:30pm, in Econ 149
- Section Questions: [Question Form](#)
- Lecture/Section Feedback: [Feedback Form](#)
- Note: please use office hours to ask questions about the material, rather than over email.

For the entirety of the course, section materials are drawn heavily from those prepared by Lea Bottmer last year, as well as from some materials produced by Merrill Warnick in previous years.

In order to foster a healthy, inclusive classroom environment, we will all abide by the following **community norms**:

- One microphone: we will allow everyone to finish their thought without interruption
- Everyone got accepted to Stanford and is here for a reason. There is no need to prove anything to anyone.
- Asking questions is a great way to deepen your understanding and resolve any confusion you may have. It is not a signaling device.
- There is no such thing as a stupid or dumb question
- Questions asked during section will be answered by the TA

- First problem set due **Thursday Jan 15, 11:59pm**
- You may work in groups, but you must turn in your own submission and note who you collaborated with on any given problem
- For easy grading, please submit only typed or scanned/electronic handwritten submissions (no photos of paper please)
- For problem set-related questions, come to my office hours or reach out to me
  - I **will not** respond to problem set emails sent on the due date

## Hints for Problem Set 1

- For problem 1, make sure you remember how to use matrices and matrix multiplications (e.g. matrix derivations if you use matrix notation for some of the subparts or generally how to take the inverse of a  $2 \times 2$ ). In the last part, make sure you use the fact that  $x$  is binary.
- For problem 2, recall from the lecture to please not use Excel or Stata for this problem or even any packages in Matlab/R/etc. We want you to implement the OLS estimator **from scratch!**

1. General Info & Logistics

2. Statistical Background

3. Identification & Estimation

4. Properties of Estimators

- Finite Sample
- Large Sample

5. Practice Problems

# Properties of Random Variables

Let  $X, Y, Z$  be random variables and  $M$  a (constant) matrix. We then have:

- **Linearity of Expectation:**  $\mathbb{E}_{XY}[X + Y] = \mathbb{E}_X[X] + \mathbb{E}_Y[Y]$
- **Covariance:**  $\text{Cov}(X, Y) = \mathbb{E}[XY] - \mathbb{E}[X]\mathbb{E}[Y]$ 
  - $\text{Cov}(aX, Y) = a \text{Cov}(X, Y)$
  - $\text{Cov}(X + Z, Y) = \text{Cov}(X, Y) + \text{Cov}(Z, Y)$
- **Variance:**  $\text{Var}(X) = \text{Cov}(X, X) = \mathbb{E}[X^2] - (\mathbb{E}[X])^2$ 
  - Given a matrix  $M$ , we have:  $\text{Var}(MX) = M \text{Var}(X) M^T$
- **Law of Iterated Expectations (Tower Rule):**  $\mathbb{E}_X[X] = \mathbb{E}_Y[\mathbb{E}_X[X|Y]]$
- **Law of Total Variance:**  $\text{Var}(Y) = \text{Var}(\mathbb{E}[Y|X]) + \mathbb{E}[\text{Var}(Y|X)]$

# Convergence Concepts

Let  $X_n$  be a sequence of random variables, with corresponding CDFs  $F_{X_n}$ .

## ■ Almost-Sure Convergence:

$$X_n \xrightarrow{\text{a.s.}} X$$

(aka w.p. 1)

$$\forall \varepsilon > 0, P\left(\lim_{n \rightarrow \infty} |X_n - X| < \varepsilon\right) = 1$$

## ■ Convergence in Probability:

$$X_n \xrightarrow{P} X$$

$$\forall \varepsilon > 0, \lim_{n \rightarrow \infty} P(|X_n - X| \geq \varepsilon) = 0$$

## ■ Convergence in Distribution:

$$X_n \xrightarrow{d} F$$

$$\lim_{n \rightarrow \infty} F_{X_n}(x) = F_X(x)$$

(at all  $x \in \mathbb{R}$  where  $F_X(x)$  continuous)

Lemma:

$$(X_n \xrightarrow{d} a)$$

$$\Rightarrow (X_n \xrightarrow{P} a)$$

$\Rightarrow$  = pointwise convergence of CDFs

## ■ Weak Law of Large Numbers (WLLN)

- Assumptions:  $X_i \stackrel{iid}{\sim} F$ ,  $E[X_i] = \mu$ ,  $Var(X_i) = \sigma^2 < \infty$

- Result:  $\bar{X} = E_n[X_i] = \frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} \mu$

(Strong LLN:  $\bar{X} \xrightarrow{a.s.} \mu$ , extra assumption:  $E[X^4] < \infty$ )

## ■ Continuous Mapping Theorem

- Assumptions:  $X_n \xrightarrow{P} X$ ,  $g(x)$  is continuous

- Result:  $g(X_n) \xrightarrow{P} g(X)$

## ■ Central Limit Theorem (CLT)

- Assumptions:  $X_1, \dots, X_n$  iid,  $E[X_i] = \mu$ ,  $\text{Var}(X_i) = \sigma^2$ ,  $\bar{X} = E_n[X_i]$

- Result:  $\sqrt{n} \left( \frac{\bar{X} - \mu}{\sigma} \right) \xrightarrow{d} N(0, 1)$

→ "asymptotic normality"

## ■ Slutsky's Theorem

- Assumptions:  $X_n \xrightarrow{d} X$ ,  $Y_n \xrightarrow{p} a$

- Result:  $X_n + Y_n \xrightarrow{d} X + a$ ,  $X_n Y_n \xrightarrow{d} Xa$

1. General Info & Logistics

2. Statistical Background

3. Identification & Estimation

4. Properties of Estimators

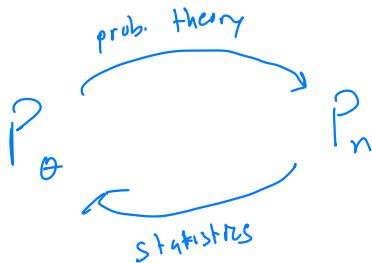
- Finite Sample
- Large Sample

5. Practice Problems

## Population vs. Sample Distributions

$P_{\theta}$  = population distribution  
↳ population parameter of interest

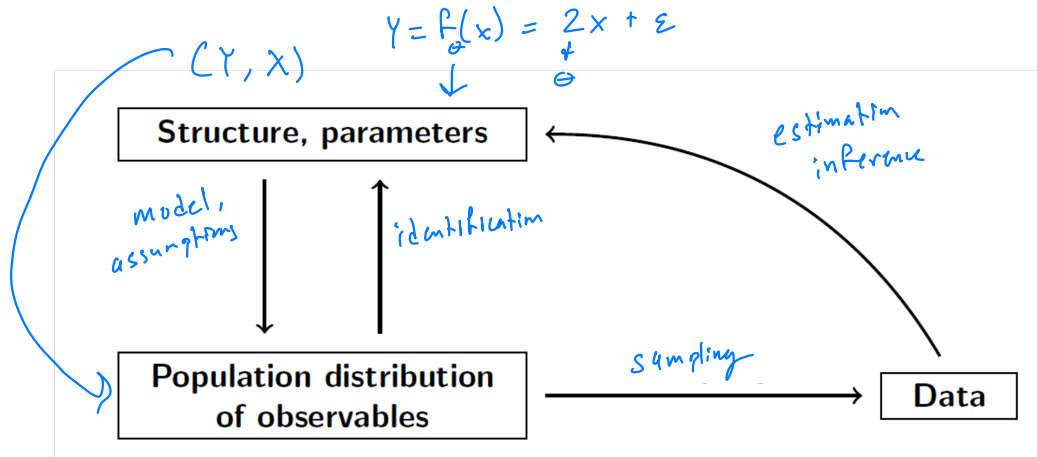
$P_n$  = sampling distribution  
↳ sample size



## Definitions

- **Estimand:** Object of interest of estimation as a population quantity  
(function of  $P_\theta$ )
- **Estimator:** Sample equivalent of the estimand  
(function of  $P_n$ )
- **Estimate:** Numeric value of estimator when we plug in data

# Big Picture



1. General Info & Logistics

2. Statistical Background

3. Identification & Estimation

4. Properties of Estimators

- Finite Sample
- Large Sample

5. Practice Problems

- **Average Treatment Effect:**

$$\tau = E[Y(1) - Y(0)]$$

- **Difference-in-Means Estimator:**

$$\hat{\tau} = \frac{1}{n_1} \sum_{d_i=1} Y_i - \frac{1}{n_0} \sum_{d_i=0} Y_i$$

1. General Info & Logistics
2. Statistical Background
3. Identification & Estimation
4. Properties of Estimators
  - Finite Sample
  - Large Sample
5. Practice Problems

**Definition:**  $\hat{\mu}$  is unbiased for  $\mu \iff E[\hat{\mu}] = \mu$

Is  $\hat{\tau}$  unbiased?

$$E[\hat{\tau} | d_1, \dots, d_n] = \frac{1}{n_1} \sum_{d_i=1} E[Y_i | d_1, \dots, d_n] - \frac{1}{n_0} \sum_{d_i=0} E[Y_i | d_1, \dots, d_n]$$

(SUTVA)

$$= \frac{1}{n_1} \sum_{d_i=1} E[Y_i | d_i] - \frac{1}{n_0} \sum_{d_i=0} E[Y_i | d_i]$$

(defn. of P.O.)

$$= \frac{1}{n_1} \sum_{d_i=1} E[Y_i(1)] - \frac{1}{n_0} \sum_{d_i=0} E[Y_i(0)]$$

(identical dist. of P.O.)

$$= \left( \frac{1}{n_1} \cdot n_1 \right) E[Y(1)] - \left( \frac{1}{n_0} \cdot n \right) E[Y(0)]$$

(linearity of  $E[\cdot]$ )

$$= E[Y(1) - Y(0)]$$

$$= \tau$$

$$\text{Var}(\hat{\mu}) = E[(\hat{\mu} - E[\hat{\mu}])^2]$$

**Definition:**  $MSE(\hat{\mu}) = E[(\hat{\mu} - \mu)^2]$

Lemma:  $MSE(\hat{\mu}) = \text{Bias}^2(\hat{\mu}) + \text{Var}(\hat{\mu})$

Proof:  $E[(\hat{\mu} - \mu)^2]$

What is the MSE of  $\hat{\mu}$ ?

$$\begin{aligned}
 &= E[(\hat{\mu} - E[\hat{\mu}])^2] + E[(\mu - E[\hat{\mu}])^2] - 2E[(\hat{\mu} - E[\hat{\mu}])(\mu - E[\hat{\mu}])] \\
 &\quad \underbrace{\hspace{1cm}}_{\text{Var}(\hat{\mu})} \qquad \qquad \qquad \underbrace{\hspace{1cm}}_{=0} \\
 &= E[\mu^2 + E[\hat{\mu}]^2 - 2\mu E[\hat{\mu}]] \\
 &= \mu^2 + E[\hat{\mu}]^2 - 2\mu E[\hat{\mu}] \\
 &= (E[\hat{\mu}] - \mu)^2 \\
 &= \text{Bias}^2(\hat{\mu})
 \end{aligned}$$

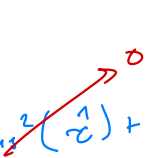
Tower rule

## Mean Squared Error (MSE)

Definition: See above

What is the MSE of  $\hat{\tau}$ ?

By Lemma:  $MSE(\hat{\tau}) = Bias^2(\hat{\tau}) + Var(\hat{\tau})$  (unbiased)



**Definition:**  $\hat{\mu}$  linear  $\leftrightarrow \exists a_i : \hat{\mu} = \sum_{i=1}^n a_i X_i$

**Is  $\hat{\tau}$  linear?**  $\hat{\tau} = \frac{1}{n_1} \sum_{d_i=1} Y_i - \frac{1}{n_0} \sum_{d_i=0} Y_i$

$$\rightarrow a_i = \begin{cases} \frac{1}{n_1} & \text{if } d_i=1 \\ -\frac{1}{n_0} & \text{if } d_i=0 \end{cases}$$

1. General Info & Logistics
2. Statistical Background
3. Identification & Estimation
4. Properties of Estimators
  - Finite Sample
  - Large Sample
5. Practice Problems

**Definition:**  $\hat{\mu}$  consistent for  $\mu \iff \hat{\mu} \xrightarrow{P} \mu$

Is  $\hat{\tau}$  consistent?  $\hat{\tau} = E_n[Y_i | d_i=1] - E_n[Y_i | d_i=0]$

$$\begin{aligned}
 & \left( \text{WLLN} + \Delta\text{-inequality} \right) \xrightarrow{P} E[Y | d=1] - E[Y | d=0] \\
 & \left( \begin{array}{l} \text{linearity of E(\cdot)} \\ \text{defn of P.O.} \end{array} \right) = E[Y(1) - Y(0)] \\
 & = \tau
 \end{aligned}$$

## Asymptotic Normality ( $\sqrt{n}$ - consistent)

**Definition:**  $\hat{\mu}$  asymptotically normal  $\Leftrightarrow \sqrt{n}(\hat{\mu} - \mu) \xrightarrow{d} N(0, \sigma^2)$

Is  $\hat{\tau}$  asymptotically normal?

$$\sqrt{n} \left( \frac{1}{\hat{\tau}} - \tau \right) = \sqrt{n} \left( \frac{1}{\hat{\tau}_1} - \tau \right)$$

1. General Info & Logistics

2. Statistical Background

3. Identification & Estimation

4. Properties of Estimators

- Finite Sample
- Large Sample

5. Practice Problems

# Econ 271/Mgtecon 604 - Section 2

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Identification Revisited
3. Review: Matrices
4. OLS: Things are Nice
5. OLS: Things Fall Apart (?)
6. Practice Problems

1. Announcements

2. Identification Revisited

3. Review: Matrices

4. OLS: Things are Nice

5. OLS: Things Fall Apart (?)

6. Practice Problems

- Midterm exam: finalized for **Monday Feb 9** (during lecture)
- If you have OAE accommodations, please upload them to the form announced on canvas so that they can be taken into account for the exams
- Optional **Lecture on Friday Jan 30** (same room, time as regularly scheduled section)
  - Lecture content (econometrics with incentives) will not be on problem sets or exams
  - Makeup section earlier that week

## Hints for Problem Set 2

$$\rightarrow (x \rightarrow \varepsilon)$$

- For problem 1, think about what possible values  $\epsilon_n$  could take to obtain the full probability distribution for  $\epsilon_i$ . Then use important properties of the expectation and variances of random variables (and their sums) for (b). To show that the variance is minimum for an estimator, there are two ways: (1) show it directly or (2) show or argue that it is the OLS estimator.
- For problem 2, think about our general idea for solving for estimands: 1) use the optimization problem directly and plug in your true model somewhere or 2) use what we know the estimator converges to and plug in the true model. You will need use the fact that  $x \sim N(0, 1)$  somewhere.
- For problem 3, remember that there is a difference between deriving the asymptotic variance  $\hat{\Sigma}$  and the asymptotic variance of the estimator  $\hat{\beta}$  (namely a factor of  $N$ ).

1. Announcements

2. Identification Revisited

3. Review: Matrices

4. OLS: Things are Nice

5. OLS: Things Fall Apart (?)

6. Practice Problems

## What is identified?

Ex. Data  $(Y_i, W_i)$

$$\Theta = E[1 \{Y(1) > Y(0)\}]$$

$\Theta$  is not identified

Even given joint dist'n of  $(Y, v)$ ,  
can't get joint of  $(Y(1), Y(0))$

Note:  $\tau = E[Y(1) - Y(0)]$

$$\hat{\tau} \xrightarrow{P} \underbrace{E[Y|W=1]}_{= E[Y(1)]} - \underbrace{E[Y|W=0]}_{= E[Y(0)]}$$

→ requires randomization  
of  $W$

Ex.  $Y_i | X_i \sim N(\alpha + \beta X_i + \gamma(\frac{X_i}{2}), \sigma^2)$

→ can only ID  $(\beta + \frac{\gamma}{2})$



1. Announcements
2. Identification Revisited
3. Review: Matrices
4. OLS: Things are Nice
5. OLS: Things Fall Apart (?)
6. Practice Problems

# Matrix Notation

Data:  $(y_i, \vec{x}_i)$   $i=1, \dots, n$   
 $\downarrow$   $\downarrow$   
 $\in \mathbb{R}$   $\mathbb{R}^k$

$x_i \in \mathbb{R}^{k \times 1}$

$$y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \quad X = \begin{pmatrix} \text{---} \vec{x}_1 \text{---} \\ \vdots \\ \text{---} \vec{x}_n \text{---} \end{pmatrix} \in \mathbb{R}^{n \times k}$$

$$y_i = x_i^T \beta + \varepsilon_i \rightarrow y = X\beta + \varepsilon$$

# Matrix Properties

Let  $\mathbf{A}$  be a matrix.

■  $\mathbf{A}$  symmetric  $\iff \mathbf{A} = \mathbf{A}^T$

■  $\mathbf{A}$  left invertible  $\iff \exists \mathbf{A}^- : \mathbf{A}^- \mathbf{A} = \mathbf{I}$

■  $\mathbf{A}$  right invertible  $\iff \exists \mathbf{A}^- : \mathbf{A} \mathbf{A}^- = \mathbf{I}$

■  $\mathbf{A}$  invertible  $\iff \exists \mathbf{A}^- : \mathbf{A}^- \mathbf{A} = \mathbf{A} \mathbf{A}^- = \mathbf{I}$   
*A square*

■  $\mathbf{A}^\dagger$  pseudoinverse of  $\mathbf{A}$   $\iff$   $\textcircled{1} \mathbf{A} \mathbf{A}^\dagger \mathbf{A} = \mathbf{A}$   $\textcircled{2} \mathbf{A}^\dagger \mathbf{A} \mathbf{A}^\dagger = \mathbf{A}^\dagger$   
 $\textcircled{3} (\mathbf{A} \mathbf{A}^\dagger)^T = \mathbf{A} \mathbf{A}^\dagger$   $\textcircled{4} (\mathbf{A}^\dagger \mathbf{A})^T = \mathbf{A}^\dagger \mathbf{A}$

■  $\mathbf{A}$  is positive semi-definite  $\iff \forall \mathbf{x} \geq 0, \mathbf{x}^T \mathbf{A} \mathbf{x} \geq 0$

Properties of  $\mathbf{A}^\dagger$ :

① If columns of  $\mathbf{A}$  linearly indep:

$$\mathbf{A}^\dagger = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T$$

② If rows linearly indep

$$\mathbf{A}^\dagger = \mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1}$$

③ If  $\mathbf{A}$  invertible,  
 $\mathbf{A}^\dagger = \mathbf{A}^{-1}$

# Matrix Rules

Let  $A, B, C$  be matrices.

- Addition is:

- Commutative  $A+B = B+A$
- Associative  $(A+B)+C = A+(B+C)$

- Multiplication is:

- Distributive over addition  $(A+B)C = AC + BC$
- Associative  $(AB)C = A(BC)$

- Transposition is:

- Distributive over addition  $(A+B)^T = A^T + B^T$
- Distributive(ish) over multiplication  $(AB)^T = B^T A^T$

- Inversion is:

- Distributive(ish) over multiplication  $(AB)^{-1} = B^{-1} A^{-1}$
- Swappable with transposition  $(A^T)^{-1} = (A^{-1})^T$

# Matrix Derivatives

Let  $\mathbf{y}, \mathbf{x}$  be (column) vectors and  $\mathbf{A}$  be a matrix (and assume all the shapes work out).

$$\blacksquare \mathbf{y} = \mathbf{Ax} \implies \frac{\partial \mathbf{y}}{\partial \mathbf{x}} = \mathbf{A}$$

$$d = \mathbf{y}^T \mathbf{x} \\ \rightarrow \frac{\partial d}{\partial \mathbf{x}} = \mathbf{y}$$

$$\blacksquare \alpha = \mathbf{y}^T \mathbf{Ax} \implies \frac{\partial \alpha}{\partial \mathbf{x}} = \mathbf{y}^T \mathbf{A}$$

$$\blacksquare \alpha = \mathbf{x}^T \mathbf{Ax} \implies \frac{\partial \alpha}{\partial \mathbf{x}} = \mathbf{x}^T (\mathbf{A} + \mathbf{A}^T)$$

1. Announcements
2. Identification Revisited
3. Review: Matrices
4. OLS: Things are Nice
5. OLS: Things Fall Apart (?)
6. Practice Problems

■ **Data:**  $(y_i, x_i)$  for  $i=1, \dots, n$ ,  $x_i \in \mathbb{R}^k$

■ **Baseline model:**  $y_i = x_i^T \beta + \varepsilon_i$ ,  $\beta \in \mathbb{R}^k$

■ **Modeling assumptions:** correct specification:  $E[\varepsilon|x] = 0$  a.s.  
homoskedasticity:  $\text{Var}(\varepsilon|x) = \sigma^2$  a.s.

## Estimand: "Analytic" Solution

Note that we can write  $\beta$  in the population as:

$$\begin{aligned}\beta &= E^{-1} [xx'] E[xy] \\ &= E^{-1} [xx'] E[x(x'\beta + \varepsilon)] \\ &= E^{-1} [xx'] (E[xx'\beta] + E[x\varepsilon]) \\ &= \beta\end{aligned}$$

Handwritten notes in red ink:

- A red arrow points from the  $\beta$  in  $E[xx'\beta]$  to the  $\beta$  in the final result.
- A red arrow points from the  $\varepsilon$  in  $E[x\varepsilon]$  to the  $\varepsilon$  in the first  $E[x\varepsilon]$  term.
- Below  $E[x\varepsilon]$ , the following steps are written in red:
  - $= 0$
  - $= E[E(x\varepsilon|x)]$
  - $= E[x E(\varepsilon|x)]$
  - $= 0$

Recall that OLS solves the following optimization problem:

$$\beta = \underset{b \in \mathbb{R}^k}{\operatorname{argmin}} E[(y - x^T b)^2]$$

We can then write:

$$\text{FOC: } \frac{\partial}{\partial b} = 2 E[-x(y - x^T b^*)] = 0$$

$$\rightarrow E[xy] = E[xx' b^*]$$

$$\rightarrow b^* = E[xx']^{-1} E[xy]$$

## Estimator: Analogy Principle

Population:  $\beta = E^{-1}[xx^T] E[xy]$

↓

Sample:  $\hat{\beta} = E_n^{-1}[xx^T] E[xy]$

$$= \left( \frac{1}{n} \sum_{i=1}^n x_i x_i^T \right)^{-1} \left( \frac{1}{n} \sum_{i=1}^n x_i y_i \right)$$
$$= (X^T X)^{-1} X^T y$$

# Gauss-Markov Theorem

Under these assumptions: Model  $y = X\beta + \varepsilon$ ,  $E[\varepsilon|x] = 0$ ,  $\text{Var}(\varepsilon|x) = \sigma^2 I_n$   
 $X^T X$  invertible

OLS is:

- **B**est (lowest variance)  $\text{Var}(A) \geq \text{Var}(B) \iff \text{Var}(A) - \text{Var}(B) \text{ PSD}$
- **L**inear
- **U**nbiased
- **E**stimator

# Proof of Gauss-Markov

Let  $\tilde{\beta} = Cy$ , wlog  $C = (X^T X)^{-1} X^T + D$

Unbiased:  $E[\tilde{\beta}|x] = E[Cy|x] = E[(X^T X)^{-1} X^T + D](X\beta + \varepsilon|x)$   
 $= (X^T X)^{-1} X^T + D)(X\beta + E[\varepsilon|x])$   
 $= \underbrace{(X^T X)^{-1} X^T X}_{=I_k} \beta + D X \beta = (I_k + DX) \beta \rightarrow DX = 0$

Variance:  $Var(\tilde{\beta}|x) = Var(Cy|x) = C Var(y|x) C^T = \sigma^2 C C^T$   
 $= \sigma^2 ((X^T X)^{-1} X^T + D)(X(X^T X)^{-1} + D^T)$   
 $\vdots$   
 $= \sigma^2 ((X^T X)^{-1} + DD^T) = Var(\hat{\beta}|x) + \underbrace{\sigma^2 DD^T}_{\geq 0, \text{ PSD}}$

1. Announcements
2. Identification Revisited
3. Review: Matrices
4. OLS: Things are Nice
5. OLS: Things Fall Apart (?)
6. Practice Problems

# Relaxing Assumptions

We made 3 key assumptions before:

1  $E[\varepsilon|x] = 0 \rightarrow$  correct specification

2  $\text{Var}(\varepsilon|x) = \sigma^2 \rightarrow$  homoskedasticity

3  $E[xx^T]$  invertible

2 of these can be relaxed to allow for:

1 Misspecification:  $E[y|x] \neq x^T \beta$

2 Heteroskedasticity:  $\text{Var}(\varepsilon|x) \neq \sigma^2$

$\hookrightarrow$  new variance: Sandwich Formula:  $H^{-1} S H^{-1} \rightarrow E^{-1}(xx^T)$

$\nearrow E[x \varepsilon^2 x^T]$

## Theorem: Best Linear Predictor

$$\beta = \underset{b}{\operatorname{argmin}} E[(E[y|x] - b^T x)^2] = E^{-1}[xx^T] E[xy]$$

What changed?

Define:  $\varepsilon = y - x^T \beta$  . FOL gives:  $E[x\varepsilon] = 0$

→ assumption by construction  
→ weaker than  
 $E[\varepsilon|x] = 0$

## Estimand Under Misspecification: Excluded Variable

What if there is some unobserved variable  $u$  that impacts  $y$  linearly?

$$\text{True model: } y = \beta_0 + x\beta_1 + u\beta_2 + \varepsilon \quad E[\varepsilon | x, u] = 0$$

$$\text{Misspecified reg.: } y = \tilde{\beta}_0 + x\tilde{\beta}_1 + \delta$$

$$\tilde{\beta}_1 = \beta_1 + \beta_2 \frac{\text{Cov}(u, x)}{\text{Var}(x)}$$

$$\rightarrow \text{OVB} \text{ unless } \text{Cov}(u, x) = 0 \text{ or } \beta_2 = 0$$

# Frisch-Waugh-Lovell Theorem

Motivation: Multivariate regression, but only care about effect of 1 regressor  
Want to control for others

Controlling: Let  $X_A$  is regressor of interest.  $X_B$  collects all controls  
By controlling for  $X_B$ , account for variance in  $y$  explainable by  $X_B$   
and see how much of residual explainable by  $X_A$ .

May also be true  $X_B \rightarrow X_A \rightarrow y$ , so need to residualize  $X_A$  too

What I want:  $\hat{y} = \hat{\gamma}_1 X_B$ ,  $\hat{X}_A = \hat{\gamma}_2 X_B \rightarrow (y - \hat{y}) = \beta(X_A - \hat{X}_A) + \varepsilon$   
↳ param. of interest

If I run 1 big regression:  $y = \tilde{\beta} X_A + \gamma_3 X_B + \varepsilon$

By FWL:  $\tilde{\beta} = \beta$

1. Announcements
2. Identification Revisited
3. Review: Matrices
4. OLS: Things are Nice
5. OLS: Things Fall Apart (?)
6. Practice Problems

# Econ 271/Mgtecon 604 - Section 3

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Linear Regression: Cluster-Robust Inference
3. Intro to Maximum Likelihood Estimation (MLE)
4. Practice Problems

1. Announcements

2. Linear Regression: Cluster-Robust Inference

3. Intro to Maximum Likelihood Estimation (MLE)

4. Practice Problems

- Optional **Lecture** on **Friday Jan 30** (same room, time as regularly scheduled section)
  - Lecture content (econometrics with incentives) will not be on problem sets or exams
- Makeup section on **Wednesday Jan 28, 4-6pm, GSB M104**
- Office hours next week: Tuesday 11:30am-1pm (Landau 149) (Wednesday OH canceled for section)
- No problem set this week!

1. Announcements

2. Linear Regression: Cluster-Robust Inference

3. Intro to Maximum Likelihood Estimation (MLE)

4. Practice Problems

**Motivation:** Suppose observe same individuals at multiple points in time  
Individuals  $i$ , each observed  $T_i$  times

$(x_{i0}, x_{i1}, \dots, x_{iT_i}) \rightarrow$  correlated! no longer iid

$$\begin{pmatrix} x_{11}, \dots, x_{1T} \\ \vdots \\ x_{n1}, \dots, x_{nT} \end{pmatrix}$$

$\underbrace{\hspace{10em}}_{n \times T}$

$\rightarrow$  effectively just have  $n$  "real" data points  
the rest are correlated "versions"

# Assumptions

■ **Model:**  $y_{it} = x_{it}' \beta + \varepsilon_{it}$

$$E[x_{it} \varepsilon_{it}] = 0$$

Assuming balanced panel,  $T_i = T \quad \forall i$

■ **Sample:** Sampling individuals, not time periods

$(y_{01}, \dots, y_{0T}, x_{01}, \dots, x_{0T})$

$(y_{11}, \dots, y_{1T}, x_{11}, \dots, x_{1T})$

$\vdots$

$(y_{n1}, \dots, y_{nT}, x_{n1}, \dots, x_{nT})$

sample

## Estimation: Pooled OLS

Big Idea: ignore time component  $\rightarrow T_i$  in general

$$\hat{\beta} = \left( \sum_{i=1}^n \sum_{t=1}^T x_{it} x_{it}' \right)^{-1} \left( \sum_{i=1}^n \sum_{t=1}^T x_{it} y_{it} \right)$$

Assume iid  $\downarrow$  across  $i$ , but not  $t$   
sampling

$$\hat{\beta} = \left( \frac{1}{n} \sum_{i=1}^n \sum_{t=1}^T x_{it} x_{it}' \right)^{-1} \left( \frac{1}{n} \sum_{i=1}^n \sum_{t=1}^T x_{it} y_{it} \right)$$

(LLN:  
 $E_n[x] \xrightarrow{P} E[x]$ )

$$= E_n^{-1} \left[ \sum_{t=1}^T x_t x_t' \right] E_n \left[ \sum_{t=1}^T x_t y_t \right]$$

$$\xrightarrow{P} E^{-1} \left[ \sum_{t=1}^T x_t x_t' \right] \xrightarrow{P} E \left[ \sum_{t=1}^T x_t y_t \right]$$

by LLN + LMT                      by LLN

$$\xrightarrow{P} E^{-1} \left[ \sum_{t=1}^T x_t x_t' \right] E \left[ \sum_{t=1}^T x_t y_t \right]$$

by LMT (  $g(x,y) = xy$  is cts )

# Asymptotic Normality

$$\sqrt{n} (\hat{\beta} - \beta) = \underbrace{\left( \frac{1}{n} \sum_{i=1}^n \sum_{t=1}^T x_{it} x_{it}' \right)^{-1}}_{\substack{\xrightarrow{P} E^{-1} \left[ \sum_{t=1}^T x_t x_t' \right] \\ \text{by LLN, } nT}} \underbrace{\left( \sqrt{n} \cdot \frac{1}{n} \sum_{i=1}^n \sum_{t=1}^T x_{it} \underbrace{(y_{it} - x_{it}' \beta)}_{\varepsilon_{it}} \right)}_{\substack{\xrightarrow{d} N(0, \Sigma) \\ \text{by CLT}}}$$

$$\rightarrow \neq E[\varepsilon \varepsilon' x]$$

Define  $Q = E \left[ \sum_{t=1}^T x_t x_t' \right]$ ,  $\Sigma = E \left[ \left( \sum_{t=1}^T x_t \varepsilon_t \right) \left( \sum_{t=1}^T x_t \varepsilon_t \right)' \right]$

By Slutsky:  $\sqrt{n} (\hat{\beta} - \beta) \xrightarrow{d} N(0, Q^{-1} \Sigma Q^{-1})$

Did not assume  $\varepsilon_{it}$  uncorrelated across  $t$

1. Announcements

2. Linear Regression: Cluster-Robust Inference

3. Intro to Maximum Likelihood Estimation (MLE)

4. Practice Problems

## Big Idea (Normal Example)

■ **Model:**  $z \sim \mathcal{N}(\mu, \sigma^2)$

■ **Sample:**  $z_1, \dots, z_n \stackrel{\text{iid}}{\sim} \mathcal{N}(\mu, \sigma^2)$

■ **Likelihood:**  $L(\mu, \sigma^2) = P(\text{data} | \text{model}) = \prod_{i=1}^n f_{(\mu, \sigma^2)}(z_i)$

*iid sampling*  
*normal pdf*

■ **Estimator:**  $(\hat{\mu}, \hat{\sigma}^2) = \operatorname{argmax}_{\mu, \sigma^2} L(\mu, \sigma^2)$

*generally write  $\theta = (\mu, \sigma^2)$ ,  $L(\theta)$*

## Conditional Log-Likelihood

Let  $z = (y, x)$ ,  $y|x \sim f_\beta(y|x)$ ,  $x \sim g_\phi(x)$ ,  $\theta = (\beta, \phi)$

Definition of conditional probability:  $f_\theta(y, x) = f_\beta(y|x) g_\phi(x)$

$$L(\theta) = \prod_{i=1}^n f_\theta(y_i, x_i) = \prod_{i=1}^n f_\beta(y_i|x_i) g_\phi(x_i)$$

Identification condition:  $P_{\theta_0}(f_\theta(z) \neq f_{\theta_0}(z)) > 0 \quad \forall \theta \neq \theta_0$

Jensen's Inequality gives us Information Inequality:  $E_{\theta_0}[\log f_\theta(z)] > E_{\theta_0}[\log f_{\theta_0}(z)] \quad \forall \theta \neq \theta_0$

Take logs of everything:

$$\ell(\theta) = \log L(\theta) = \sum_{i=1}^n \underbrace{\log f_\beta(y_i|x_i)}_{\perp \theta} + \underbrace{\log g_\phi(x_i)}_{\perp \beta}$$

If  $\beta, \phi$  don't overlap, then  $\hat{\beta} = \arg \max_{\beta} \sum_{i=1}^n \log f_\beta(y_i|x_i)$

## An Algorithm for MLE

- 1 Transform population model into an equation for which you know the distribution if needed
- 2 Write down  $L(\theta) = \prod_{i=1}^n f_{\theta}(x_i|x_i)$ , where we know distribution of  $y|x$ , and  $\theta$  is parameter of interest
- 3 Transform to log-likelihood,  $l(\theta)$
- 4 Simplify objective function (ignore positive scalar multiples, extraneous terms, etc.)
- 5 Take FOC and solve for  $\hat{\theta}$
- 6 Check second derivative to verify curvature

## Estimation Example: OLS

Let  $y = x'\beta + \epsilon$ ,  $\epsilon \sim \mathcal{N}(0, \sigma^2)$ ,  $\sigma^2$  known.

1 What is the likelihood?

$$\epsilon = y - x'\beta \rightarrow y|x \sim \mathcal{N}(x'\beta, \sigma^2) \rightarrow L(\beta) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y_i - x_i'\beta)^2}{2\sigma^2}}$$

2 What is the log-likelihood?

$$\ell(\beta) = \sum_{i=1}^n \log \left( \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y_i - x_i'\beta)^2}{2\sigma^2}} \right)$$

3 Simplify

$$= \sum_{i=1}^n \log \left( \frac{1}{\sqrt{2\pi\sigma^2}} \right) - \frac{(y_i - x_i'\beta)^2}{2\sigma^2} \rightarrow \text{positive scalar}$$

$$\rightarrow \hat{\beta} = \arg \min_{\beta} \sum_{i=1}^n (y_i - x_i'\beta)^2 = \hat{\beta}^{\text{OLS}}$$

OLS is MLE (in this setting)

## Estimation Example: Logit

Let  $y \in \{0, 1\}$ ,  $P(y = 1|x) = (1 + \exp(-x^\top \beta))^{-1}$ .

1 What is the likelihood?

$$L(\beta) = \prod_{i=1}^n \left( \frac{1}{1 + e^{-x_i^\top \beta}} \right)^{y_i} \left( 1 - \frac{1}{1 + e^{-x_i^\top \beta}} \right)^{1-y_i}$$

2 What is the log-likelihood?

$$\ell(\beta) = \sum_{i=1}^n \log \left[ \left( \frac{1}{1 + e^{-x_i^\top \beta}} \right)^{y_i} \left( 1 - \frac{1}{1 + e^{-x_i^\top \beta}} \right)^{1-y_i} \right]$$

3 Simplify

$$\begin{aligned} \ell(\beta) &= \sum_{i=1}^n y_i \log \left( \frac{1}{1 + e^{-x_i^\top \beta}} \right) + (1-y_i) \log \left( 1 - \frac{1}{1 + e^{-x_i^\top \beta}} \right) \\ &= \sum_{i=1}^n -y_i \log(1 + e^{-x_i^\top \beta}) + (1-y_i) \log \left( \frac{e^{-x_i^\top \beta}}{1 + e^{-x_i^\top \beta}} \right) \\ &\vdots \\ &\rightarrow \hat{\beta}_{\text{logit}} = \underset{\beta}{\operatorname{argmin}} \sum_{i=1}^n [(1-y_i)(x_i^\top \beta) + \log(1 + e^{-x_i^\top \beta})] \end{aligned}$$

1. Announcements

2. Linear Regression: Cluster-Robust Inference

3. Intro to Maximum Likelihood Estimation (MLE)

4. Practice Problems

# Econ 271/Mgtecon 604 - Section 4

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. MLE: Theoretical Properties
3. Panel Data
  - Fixed Effects
  - Random Effects
  - Difference in Differences
4. Practice Problems

## 1. Announcements

## 2. MLE: Theoretical Properties

## 3. Panel Data

- Fixed Effects
- Random Effects
- Difference in Differences

## 4. Practice Problems

- Optional **Lecture** on **Friday Jan 30** (same room, time as regularly scheduled section)
- Next week, schedule goes back to normal
- **Midterm Monday Feb 9** - Section 5 will spend (at least) some time on review

1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects
- Random Effects
- Difference in Differences

4. Practice Problems

## Extremum Estimators (M-Estimators)

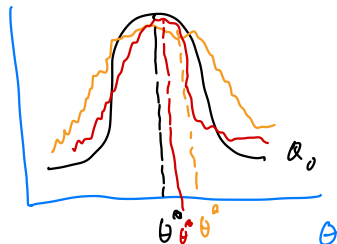
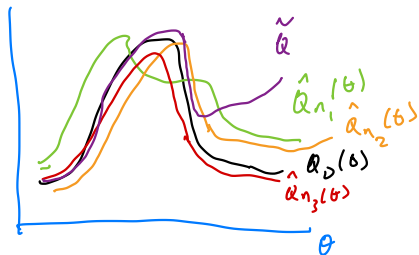
$$\hat{\theta} = \underset{\theta \in \Theta}{\operatorname{argmax}} \underbrace{\hat{Q}_n(\theta)}_{\hookrightarrow \hat{Q}_n(\theta) = \sum_{i=1}^n m(z_i; \theta)}$$

Ex. MLE:

$$m(z_i; \theta) = \log f_{\theta}(z_i)$$

# Convergence of Functions

**Uniform Convergence in Probability:**  $\hat{Q}_n(\theta)$  converges uniformly in probability to  $Q_0(\theta)$  if:  $\sup_{\theta \in \Theta} |\hat{Q}_n(\theta) - Q_0(\theta)| \xrightarrow{P} 0$



## Useful Lemma

$$\text{Let } \hat{Q}_n(\theta) = \frac{1}{n} \sum_{i=1}^n m(z_i; \theta) \quad , \quad Q_0(\theta) = E[m(z; \theta)]$$

Assume:

①  $\{z_1, \dots, z_n\}$  iid

②  $\Theta$  compact

③  $m(z; \theta)$  is continuous at each  $\theta \in \Theta$  w.p. 1  
and  $E[M(y)] < \infty$

④  $\exists M(y)$  s.t.  $|m(y; \theta)| \leq M(y) \quad \forall \theta \in \Theta$

"M-estimators"

Then:  $Q_0(\theta)$  continuous and  $\hat{Q}_n(\theta)$  converges uniformly in probability to  $Q_0(\theta)$

## Theorem: Consistency

Assume there is a function  $Q_0(\theta)$  and a vector  $\theta_0 \in \Theta \cup \mathbb{R}^k$  such that:

- 1  $Q_0(\theta)$  is uniquely maximized at  $\theta_0$
- 2  $\Theta$  compact
- 3  $Q_0(\theta)$  continuous
- 4  $\hat{Q}_n(\theta)$  converge unif. in prob. to  $Q_0(\theta)$

Then,  $\hat{\theta} \xrightarrow{P} \theta_0$ .

## Influence Functions + Asymptotic Linearity

**Asymptotic Linearity:**  $\sqrt{n}(\hat{\theta} - \theta_0) = \frac{1}{\sqrt{n}} \sum_{i=1}^n \underbrace{\psi(z_i)}_{\text{influence function}} + \underbrace{o_p(1)}_{\text{noise}}$

$\underbrace{\text{avg. over small increments}} \quad \hookrightarrow \begin{aligned} E[\psi(z)] &= 0 \\ E[\|\psi(z)\|^2] &< \infty \end{aligned}$

Then, CLT gives:

$$\sqrt{n}(\hat{\theta} - \theta_0) \xrightarrow{d} N(0, V)$$

$$V = \text{Var}(\psi(z)) = E[\psi(z)\psi(z)']$$

## MLE Influence Function

**Score Function:**  $s(z; \theta) = \frac{\partial \log f_\theta(z)}{\partial \theta}$

**MLE FOC:**  $\frac{1}{n} \sum_{i=1}^n s(z_i; \theta) = 0$

**Taylor Expansion around  $\theta_0$ :**  $s(z_i; \theta_0) + \frac{\partial}{\partial \theta} s(z_i; \theta_0)' (\hat{\theta} - \theta_0) \approx 0$

$\rightarrow \frac{1}{n} \sum s(z_i; \theta_0) + \frac{1}{n} \sum \frac{\partial}{\partial \theta} s(z_i; \theta_0)' (\hat{\theta} - \theta_0) \approx 0$

$H = \frac{1}{n} \sum_{i=1}^n \frac{\partial}{\partial \theta} s(z_i; \theta)$

$\rightarrow \frac{1}{n} \sum_{i=1}^n s(z_i; \theta_0) + H(\hat{\theta} - \theta_0) \approx 0$

$\rightarrow \sqrt{n}(\hat{\theta} - \theta_0) \approx -H^{-1} \frac{1}{\sqrt{n}} \sum_{i=1}^n s(z_i; \theta_0)$

$\psi(z_i) = -H^{-1} s(z_i; \theta_0)$

$\rightarrow \sqrt{n}(\hat{\theta} - \theta_0) \xrightarrow{d} N(0, V)$

$V = H^{-1} J H^{-1}$

$J = E[s(z; \theta) s(z; \theta)']$

**Kullback-Liebler Divergence:**

$$D_{KL}(p \parallel q) = \mathbb{E}_p \left[ \log \frac{p(x)}{q(x)} \right]$$

$p$  = true underlying distribution  
 $q$  = our approx.

**Maximizing Likelihood = Minimizing KL Divergence:**

$$\begin{aligned} \operatorname{argmin}_{\theta} D_{KL}(p \parallel q(\theta)) &= \operatorname{argmin}_{\theta} \mathbb{E} \left[ \log \frac{p(x)}{q(x; \theta)} \right] \\ &= \operatorname{argmin}_{\theta} \mathbb{E} [\log p(x) - \log q(x; \theta)] \\ &= \operatorname{argmin}_{\theta} -\mathbb{E} [\log q(x; \theta)] \\ &= \operatorname{argmax}_{\theta} \mathbb{E} [\log q(x; \theta)] \\ &= \hat{\theta}_{MLE} \end{aligned}$$



1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects
- Random Effects
- Difference in Differences

4. Practice Problems

1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects

- Random Effects

- Difference in Differences

4. Practice Problems

## Modeling Assumptions + Estimator

**Modeling Assumptions:**  $y_{it} = x_{it}'\beta + \eta_i + \varepsilon_{it}$   
↳  $\eta_i$  is a common or individual variable

Data iid across  $i$ ,  $E[\varepsilon_{it} | x_{i1}, \dots, x_{iT}, \eta_i] = 0$

May have  $\text{Cov}(x_{it}, \eta_i) \neq 0$

**Estimator:**  $(\hat{\beta}_{FE}, \hat{\eta}_1, \dots, \hat{\eta}_n) = \underset{\beta, \{\eta_i\}_{i=1}^n}{\text{argmin}} \sum_{i=1}^n \sum_{t=1}^T (y_{it} - x_{it}'\beta - \eta_i)^2$

$$\text{FOC: } \frac{\partial}{\partial \beta} : \sum_i \sum_t -2x_{it} (y_{it} - x_{it}'\hat{\beta} - \hat{\eta}_i) = 0 \quad \frac{\partial}{\partial \eta_i} : \sum_t -2(y_{it} - x_{it}'\hat{\beta} - \hat{\eta}_i) = 0$$

$$\rightarrow \hat{\beta}_{FE} = \left( \sum_i \sum_t (x_{it} - \bar{x}_i)(x_{it} - \bar{x}_i)' \right)^{-1} \left( \sum_i \sum_t (x_{it} - \bar{x}_i)(y_{it} - \bar{y}_i) \right) = \bar{y}_i - \bar{x}_i' \hat{\beta}$$

Can also apply FGLS, get same thing.

## Asymptotic Properties

Small  $T$ , large  $n$  asymptotics

OLS machinery applies:

$$\hat{\beta}_{FE} \xrightarrow{P} \beta_{FE}$$

$$\sqrt{n}(\hat{\beta}_{FE} - \beta_{FE}) \xrightarrow{d} N(0, H^{-1} J H^{-1})$$

$$H = E \left[ \sum_i (x_{it} - \bar{x}_i)(x_{it} - \bar{x}_i)' \right]$$

$$J = E \left[ \left( \sum_i (x_{it} - \bar{x}_i)(\varepsilon_{it} - \bar{\varepsilon}_i) \right) \left( \sum_i (x_{it} - \bar{x}_i)(\varepsilon_{it} - \bar{\varepsilon}_i) \right)' \right]$$

1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects

- **Random Effects**

- Difference in Differences

4. Practice Problems

**Modeling Assumptions:**  $y_{it} = x_{it}'\beta + \eta_{it} + \varepsilon_i$

$$E[\varepsilon_{it} | x_{it}, \dots, x_{it}, \eta_i] = 0$$

$$E[\eta_i | x_{it}, \dots, x_{it}] = \alpha$$

diff. b/c  $i$  are random

**Estimator:** Treat  $\eta_i$  like error term,  $u_{it} + \alpha = \varepsilon_{it} + \eta_i$

$$y_{it} = \alpha + x_{it}'\beta + u_{it}$$

$$E[u_{it} | x_{it}, \dots, x_{it}] = 0$$

→ looks just like pooled regression!  
inherits asymptotic properties

1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects
- Random Effects
- Difference in Differences

4. Practice Problems

## Model: Two Way Fixed Effects

Suppose also want time FE

Consider causal setting,  $d_{it}$  = treatment indicator

TwFE:  $y_{it} = \alpha_i + \beta_t + \tau d_{it} + \varepsilon_{it}$

---

2 units, 2 time periods,  $d_{A2} = 1$

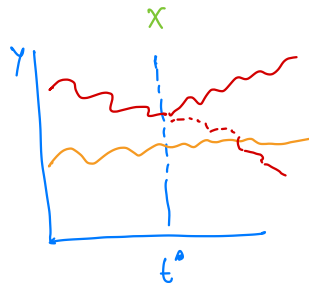
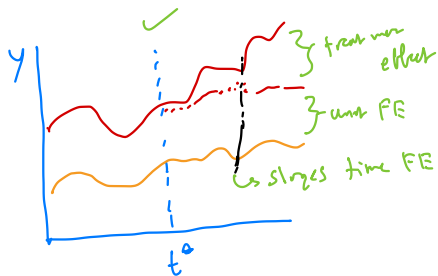
|             |          |          |
|-------------|----------|----------|
| $i \quad t$ | 1        | 2        |
| A           | $y_{A1}$ | $y_{A2}$ |
| B           | $y_{B1}$ | $y_{B2}$ |

$$(y_{A2} - y_{A1}) - (y_{B2} - y_{B1}) = \tau^{DiD}$$

When does  $\tau^{DiD} = \tau^{TwFE}$

## Key Assumption: Parallel Trends

**Big Idea:** If treatment never happened, then the 2 groups' outcomes would have remained parallel



Testable? No

But, can get at suggestive evidence  $\rightarrow$  pre-period

1. Announcements

2. MLE: Theoretical Properties

3. Panel Data

- Fixed Effects
- Random Effects
- Difference in Differences

4. Practice Problems

# Econ 271/Mgtecon 604 - Section 5

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. DiD Extensions
3. Midterm Practice

1. Announcements

2. DiD Extensions

3. Midterm Practice

- **Midterm Monday Feb 9**
  - Usual lecture room/time
  - One-sided formula sheet permitted

1. Announcements

2. DiD Extensions

3. Midterm Practice

# Staggered Adoption

**Big Idea:** Different units start receiving treatment at different times

$e_i = 1^{\text{st}}$  time period unit  $i$  is treated

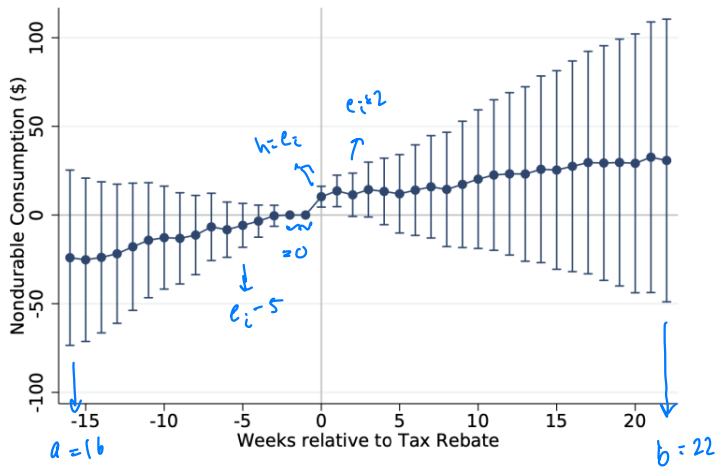
Absorbing treatment

**Ex.  $2 \times 2$  DiD with staggered adoption**

$$\hat{\tau} = (E_n[y|t=2, e_i=2] - E_n[y|t=1, e_i=2]) - (E_n[y|t=2, e_i=\infty] - E_n[y|t=1, e_i=\infty])$$

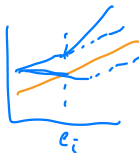
**TWFE:**

$$y_{it} = \alpha_i + \beta_t + \underbrace{\sum_{h=-1}^{b-1} \tau_h 1\{t=e_i+h\}}_{\text{pre/post-treatment effects per period}} + \underbrace{\tau_b 1\{t \geq e_i+b\}}_{\text{final period (lumping all future together)}} + \varepsilon_{it}$$



# What if we are misspecified?

Relied heavily on 2 assumptions:



1 Parallel trends

2 Homogeneous TE  $\rightarrow$  relax  
seems plausible this must not  
hold under staggered adoption...

# Effect Heterogeneity - Simple Example

**Specification:**  $y_i = d_i \tau + x_i' \beta + \varepsilon_i$   
 (only  $\tau, \beta$ )

**Estimator:**  $\hat{\tau} = \frac{\sum_i (d_i - x_i' \hat{\gamma}) y_i}{\sum_i (d_i - x_i' \hat{\gamma})^2}$   
 $\hat{\gamma}$  = coeff. of reg.  $d_i \sim x_i$   
 $\text{Var}(d - \hat{\gamma})$

**Alternative Model:**  $y_i = d_i \tilde{\tau}_i + x_i' \tilde{\beta} + \tilde{\varepsilon}_i$

**Estimand (Under Alternative):**

$$E \left[ \hat{\tau} \mid \begin{matrix} d_1, \dots, d_n \\ x_1, \dots, x_n \\ \tilde{\tau}_1, \dots, \tilde{\tau}_n \end{matrix} \right] = E \left[ \frac{\sum_i (d_i - x_i' \hat{\gamma}) (d_i \tilde{\tau}_i + x_i' \tilde{\beta} + \tilde{\varepsilon}_i)}{\sum_i (d_i - x_i' \hat{\gamma})^2} \right] = \frac{\sum_i (d_i - x_i' \hat{\gamma}) d_i \tilde{\tau}_i}{\sum_i (d_i - x_i' \hat{\gamma}) d_i} = \frac{\sum_i w_i \tilde{\tau}_i}{\sum_i w_i}$$

**Specification:**  $y_{it} = \alpha_i + \beta_t + d_{it} \tau + \varepsilon_{it}$

**Estimator:**  $\hat{\tau}$  = very complicated

**Alternative Model:**  $y_{it} = \alpha_i + \beta_t + d_{it} \tau_{it} + \varepsilon_{it}$

**Estimand (Under Alternative):**  $E[\hat{\tau} | (d_{it}, \tau_{it})_{it}] = \sum_{d_{it}=1} \tau_{it} v_{it}$

(Borusyak, Jarman, Spiess 2024)

→ spawned massive literature on  
"fixes"

$\sum_{d_{it}=1} v_{it} \geq 0$ , but can be negative

1. Announcements

2. DiD Extensions

3. Midterm Practice

# Econ 271/Mgtecon 604 - Section 6

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Synthetic Control
3. Machine Learning Intro
4. Practice Problems

1. Announcements

2. Synthetic Control

3. Machine Learning Intro

4. Practice Problems

- No class Monday Feb 16 (Presidents' Day)

1. Announcements

2. Synthetic Control

3. Machine Learning Intro

4. Practice Problems

What kind of panel must we have?

Proper panel

$$W = \begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix} \left. \vphantom{\begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}} \right\}^N$$

$\underbrace{\hspace{10em}}_T$

Estimand:

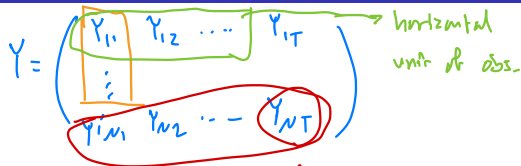
ATT

$$\tau^{ATT} = \frac{\sum_{i,t} W_{it} (Y_{it}(1) - Y_{it}(0))}{\sum_{i,t} W_{it}}$$

$$= \underbrace{Y_{N+}(1)}_{\text{observed}} - \underbrace{Y_{NT}(0)}_{\text{need to estimate}} = \tau^N$$

# Vertical vs. Horizontal Regression

**Goal:** Impute / estimate  $Y_{NT}(0)$



**Vertical Regression:**  $T-1$  obs.,  $N-1$  regressors  
eg.  $SC$

Main idea: similar units behave similarly

**Horizontal Regression:**  $N-1$  obs.,  $T-1$  regressors  
eg. unconfoundedness

Main idea: history explains all

is unit of obs.  $\uparrow$   
vertical  $Y_{11}$   
 $\hat{Y}(0)$

$$Y_{NT} = \sum_{j \neq N} w_j Y_{jT}$$

$$w_j \in \Delta^J$$

## Synthetic Control - Assumption

**Factor Model:** Assume movement in outcome summarized by shared latent factors

$$Y_{it}(0) = \mu_i^T \lambda_t + \varepsilon_{it}$$

↓      ↓  
loadings    factors

With covariates:

$$Y_{it}(0) = \delta_t + \theta_t z_i + \mu_i^T \lambda_t + \varepsilon_{it}$$

↓      obs cov  
common factor    unaffected  
constant load    by treatment

## Synthetic Control - Estimation (No Covariates)

**Ideal:** Match on factor loadings

$$\sum_{j=1}^{N-1} \hat{w}_j \mu_j = \mu_N \rightarrow \mu_i \text{ unobserved}$$

**In Practice:** Match on pre-treatment outcomes

$$\hat{Y}_{NT}(0) = \sum_{j=1}^{N-1} \hat{w}_j Y_{jT}$$

$$\hat{w} = \underset{w \in \Delta^{N-1}}{\operatorname{argmin}} \sum_{t=1}^{T-1} \left( Y_{Nb} - \sum_{j=1}^{N-1} w_j Y_{jt} \right)^2$$

**Challenge:** With 1 treated unit, no averaging  $\rightarrow$  no CVT

Open question

Placebo analyses generally used, but still have a choice to make:

leave out true  
treated unit

**1 Placebo on Units:** Iteratively pretend each unit is treated, estimate SC for treated period

$$\hat{V} = \frac{1}{N-1} \sum_{j=21}^{N-1} (\hat{Y}_{jT} - Y_{jT})^2$$

except for true treated period

**2 Placebo on Time Periods:** Iteratively pretend each period is the treatment period, estimate SC for true treated unit

$$\hat{V} = \frac{1}{T-T_0-1} \sum_{t=T_0+1}^{T-1} (\hat{Y}_{Nt} - Y_{Nt})^2$$

1. Announcements

2. Synthetic Control

3. Machine Learning Intro

4. Practice Problems

**Training Data:**  $(y_1, x_1), \dots, (y_n, x_n)$  (assume iid)

**Prediction:** Suppose true  $f: X \rightarrow Y$

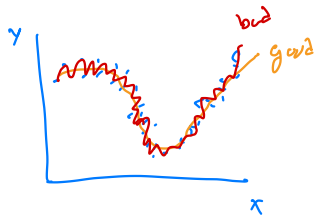
Goal:  $\hat{f} = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \mathcal{L}(f(x), y)$

loss function  
e.g. MSE (regression)  
cross-entropy (classification)

**(Main?) Concern:** Overfitting

Given new  $(y', x')$  want  $\hat{f}(x') \approx y'$

If I do "too well" in-sample, I may not generalize



**Big Idea:** Prevent models from getting too complex by penalizing complexity

$$\hat{\beta} = \underset{\beta}{\operatorname{argmin}} \sum_{i=1}^n (y_i - x_i' \beta)^2 + \lambda \operatorname{Penalty}(\beta)$$

$\hookrightarrow$  penalty "strength"

**Ridge ( $L_2$ ):**  $\operatorname{Penalty}(\beta) = \|\beta\|_2^2 = \sum_{j=1}^k \beta_j^2$

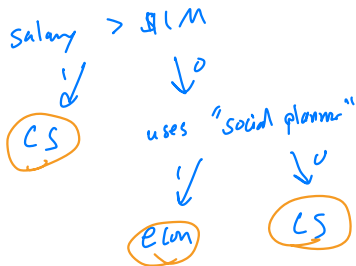
Shrink towards 0, but never all the way. Distribution

**Lasso ( $L_1$ ):**  $\operatorname{Penalty}(\beta) = \|\beta\|_1 = \sum_{j=1}^k |\beta_j|$

Shrink to 0. Selective  $\rightarrow$  yields sparse solutions

**Big Idea:** Come up with "split" rules to fit the data

e.g. economist or CS?



Choose splits to minimize  
entropy / maximize info. gain  
at each step

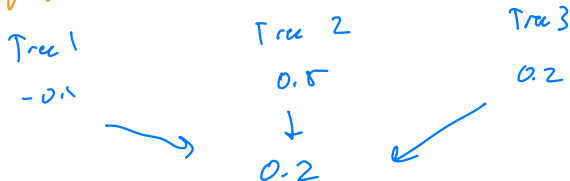
**How to regularize?** Limit depth, min # of units in node before split, ...

## Random Forests Ensemble method

**Big Idea:** Instead of believing just 1 tree, "grow" a whole "forest" and have them "vote"

↳ each gets random subset of features, observations

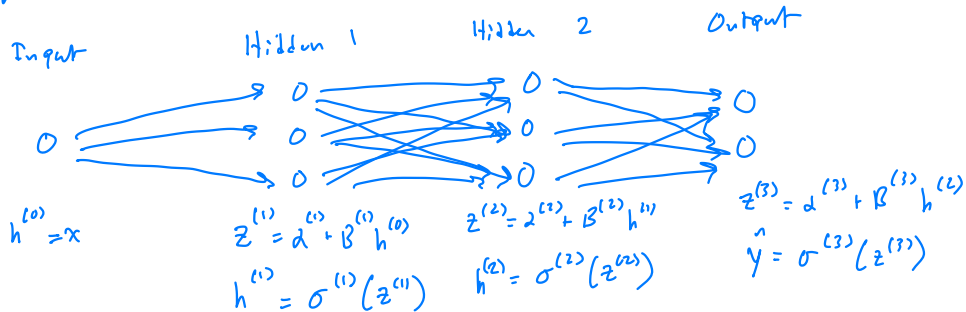
↳ aggregate via "bagging"



**How to regularize?** Change # trees, # features given to each, ...

# Neural Networks

**Big Idea:** Mimic "brain" structure via lots of layers (with nonlinear components) allowing for very flexible functional forms



$\sigma^{(l)} \in \{ \text{ReLU}, \text{sigmoid}, \text{tanh}, \dots \}$

**How to regularize?** # layers, # neurons/layer, penalty terms, dropout

1. Announcements

2. Synthetic Control

3. Machine Learning Intro

4. Practice Problems

# Econ 271/Mgtecon 604 - Section 7

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Unconfoundedness, Propensity Scores, & IPW
3. Model Tuning
4. Multimodal Data
5. Practice Problems

1. Announcements

2. Unconfoundedness, Propensity Scores, & IPW

3. Model Tuning

4. Multimodal Data

5. Practice Problems

- Problem Set 4 due Thursday Feb 26
- Jann's last lecture Monday Feb 23
- Luigi starts Wednesday Feb 25 (time series!)

1. Announcements

2. Unconfoundedness, Propensity Scores, & IPW

3. Model Tuning

4. Multimodal Data

5. Practice Problems

How can we identify causal effects?

■ **Experiments:** Randomization of treatment

○ **Observational Data:** Treatment no longer random

Goal: "as-if" randomization  
unconfoundedness  
+  
overlap

# Unconfoundedness Assumption

Also known as “selection on observables” or “exogeneity”.

**Definition:**  $W_i \perp (Y_i(0), Y_i(1)) \mid X_i$

Observables may impact treatment assignment  
But, conditional on them, “as-if” random

- **Motivation:** *even continuous  $X$ 's,  $P(\exists i, j: X_i = X_j) = 0$*
- **Propensity Score:**  $e(x) = P(V_i = 1 | X_i = x)$   
*→ "sufficient statistic" for selection on observables*
- **(Alternative) Unconfoundedness:**  $W_i \perp (Y_i(0), Y_i(1)) | e_i$

While unconfoundedness is not *testable*, its plausibility can be *assessed*.

2 assessment strategies:

1 Estimate effect on *pseudo outcome*

2 Estimate effect on *pseudo treatment*

# Identification?

**Conditional Average Treatment Effect (CATE):**  $\tau(X) = \mathbb{E}[Y_i(1) - Y_i(0) | X_i = x]$

Goal for ATE ID:  $\tau = \mathbb{E}[\tau(X_i)]$

Let's see:  $E[\tau(X_i)] = E[E[Y_i(1) - Y_i(0) | X_i = x]]$  (defn of  $\tau(x)$ )

$$= E[E[Y_i(1) | X_i = x] - E[Y_i(0) | X_i = x]] \quad (\text{linearity of } E[\cdot])$$

$$= E[E[Y_i(1) | W_i = 1, X_i = x] - E[Y_i(0) | W_i = 0, X_i = x]] \quad (\text{unconfoundedness})$$

$$= E[E[Y_i | W_i = 1, X_i = x] - E[Y_i | W_i = 0, X_i = x]] \quad (\text{defn of } PO)$$

?  
 $\tau$

may not be separately  
estimable

need something  
more

# Overlap Assumption

**Definition:**  $P(W_i=1 | X_i) = e_i \in (0,1)$

For any  $X$ , we can find both treated and control units

This assumption is *testable*. Some strategies:

1 Normalized differences

2 Plotting



3 Log odds ratio

$$\log \left( \frac{e(x)}{1-e(x)} \right)$$

## What if we have only limited overlap?

**Challenge:** *cannot get good estimates of separate conditional outcomes*  
*→ need a "better" sample*

2 strategies:

1 *Matching*

2 *Subsampling*

$$\begin{aligned}\mathbb{E}[\tau(X_i)] &= \mathbb{E}[\mathbb{E}[Y_i(1) - Y_i(0)|X_i = x]] && \text{by definition} \\ &= \mathbb{E}[\mathbb{E}[Y_i(1)|X_i = x] - \mathbb{E}[Y_i(0)|X_i = x]] && \text{by linearity} \\ &= \mathbb{E}[\mathbb{E}[Y_i(1)|W_i = 1, X_i = x] - \mathbb{E}[Y_i(0)|W_i = 0, X_i = x]] && \text{by unconfoundedness} \\ &= \mathbb{E}[\mathbb{E}[Y_i|W_i = 1, X_i = x] - \mathbb{E}[Y_i|W_i = 0, X_i = x]] && \text{by definition} \\ &= \tau && \text{by overlap}\end{aligned}$$

# Inverse Propensity Weighting (IPW)

Starting from our standard difference-in-means estimator:

$$\hat{\tau} = \mathbb{E}_n[Y_i | W_i = 1] - \mathbb{E}_n[Y_i | W_i = 0]$$

$$= \frac{1}{n_1} \sum_{W_i=1} W_i Y_i - \frac{1}{n_0} \sum_{W_i=0} (1-W_i) Y_i$$

$$= \frac{\sum_i W_i Y_i}{\sum_i W_i} - \frac{\sum_i (1-W_i) Y_i}{\sum_i (1-W_i)}$$

$$\hat{\tau} \leftarrow \frac{\mathbb{E}_n[W_i Y_i]}{\mathbb{E}_n[W_i]} - \frac{\mathbb{E}_n[(1-W_i) Y_i]}{\mathbb{E}_n[1-W_i]}$$

$$= \mathbb{E}_n \left[ \frac{W_i Y_i}{\hat{p}} - \frac{(1-W_i) Y_i}{1-\hat{p}} \right]$$

→ IPW estimator (non-random assignment):

$$\hat{\tau}_{IPW} = \mathbb{E}_n \left[ \frac{W_i Y_i}{\hat{e}(x_i)} - \frac{(1-W_i) Y_i}{1-\hat{e}(x_i)} \right]$$

↓ ↓  
probability of treatment  
depends on covariates

→ each group weighted  
by inverse probability

1. Announcements

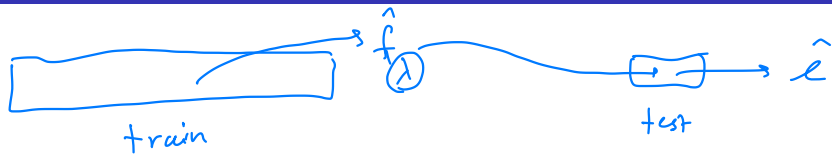
2. Unconfoundedness, Propensity Scores, & IPW

3. Model Tuning

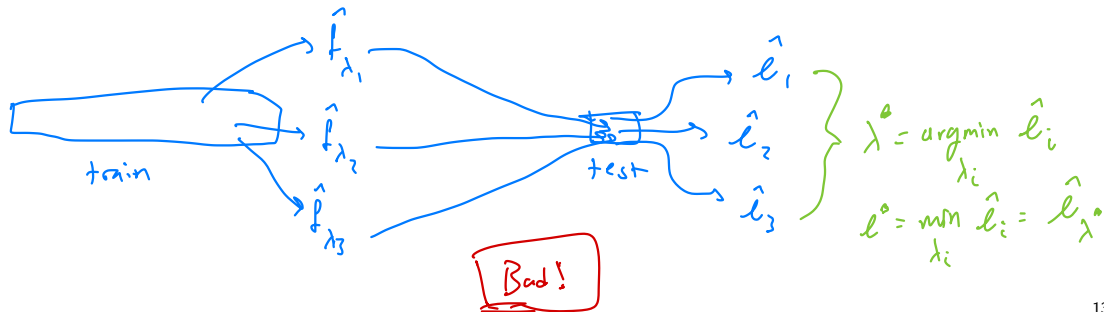
4. Multimodal Data

5. Practice Problems

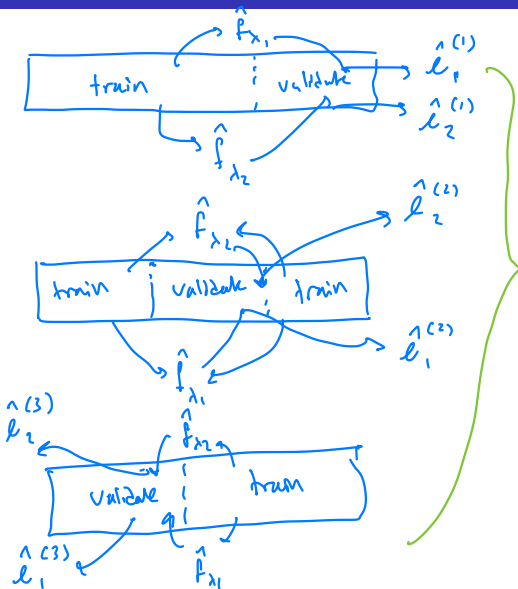
How do we choose a value for our (regularization) parameters?  $\lambda$



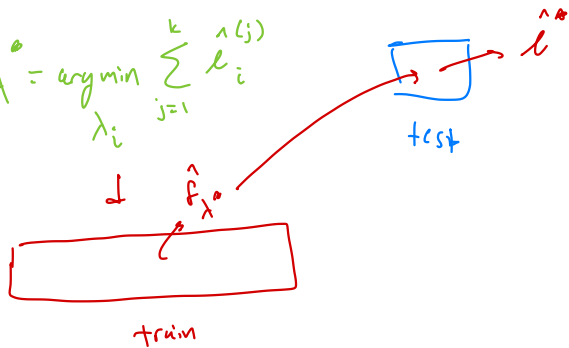
One strategy:



# Cross-Validation



$$\lambda^* = \arg \min_{\lambda_i} \sum_{j=1}^k \hat{l}_i^{(j)}$$



1. Announcements

2. Unconfoundedness, Propensity Scores, & IPW

3. Model Tuning

4. Multimodal Data

5. Practice Problems

**Big Idea:** Model images as collections of pixels, each of which has an RGB value

→ learn to predict pixel values from neighboring pixels

**Convolutional Neural Networks (CNNs):**

Use "sliding window" (assign neurons to specific regions of space in image)

→ takes into account "grid topology" of image

→ allows for dense representation ("embedding") of image based  
learned inter-pixel relationships

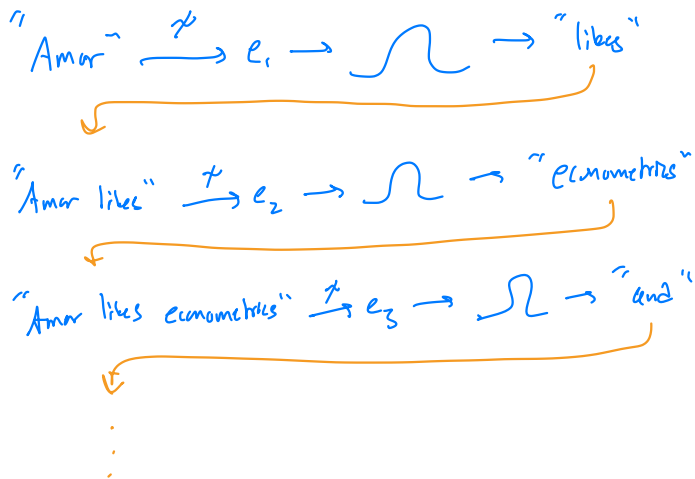
**Big Idea:** Model text as sequence of words  
→ learn to predict next word from previous word(s)

Progression of methods:

- 1 Bag of words - word count matrix
  - 2 Topic model (LDA, ...) - probabilistic "clustering" in text space
  - 3 Embeddings (word2vec) - dense vector for predicting words
  - 4 Transformers - fancy architecture that captures complex intra-word relationships (attention)
  - 5 LLMs - transformer-based models that generate text (iteratively)
- 

$$\mathcal{T}: \text{text} \rightarrow \text{embedding}$$

Iterative Sampling:



At each step, model's embedding is "sufficient statistic" of next token/word.

But, generating that embedding requires knowledge of all prior tokens/words.

1. Announcements

2. Unconfoundedness, Propensity Scores, & IPW

3. Model Tuning

4. Multimodal Data

5. Practice Problems

# Econ 271/Mgtecon 604 - Section 8

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Double Machine Learning
3. CATE Estimation
4. Introduction to Linear Stationary Processes
5. Practice Problems

1. Announcements

2. Double Machine Learning

3. CATE Estimation

4. Introduction to Linear Stationary Processes

5. Practice Problems

- Final Exam: March 18, 3:30-6:30pm
- Last problem set: due March 13, 11:59pm
- Final review session: March 13, 1:30-3:20pm (usual section time/place)

1. Announcements

2. Double Machine Learning

3. CATE Estimation

4. Introduction to Linear Stationary Processes

5. Practice Problems

$$\hat{\tau}^{AIPW} = \frac{1}{N} \sum_{i=1}^N (\hat{\mu}_1(X_i) - \hat{\mu}_0(X_i)) + W_i \frac{Y_i - \hat{\mu}_1(X_i)}{\hat{e}(X_i)} - (1 - W_i) \frac{Y_i - \hat{\mu}_0(X_i)}{1 - \hat{e}(X_i)}$$

**Double Robustness:**

$\hat{\tau}^{AIPW} \xrightarrow{P} \tau$  if:

①  $\hat{e}$  correctly specified (and estimated consistently)

or

②  $\hat{\mu}_1, \hat{\mu}_0$  correctly specified (and estimated consistently)

# Influence Function/AIPW - Correctly Specified Outcomes

$$\hat{\tau}^{AIPW} = \left[ \frac{1}{N} \sum_{i=1}^N (\hat{\mu}_1(X_i) - \hat{\mu}_0(X_i)) \right] + \left[ \frac{1}{N} \sum_{i=1}^N W_i \frac{Y_i - \hat{\mu}_1(X_i)}{\hat{e}(X_i)} - (1 - W_i) \frac{Y_i - \hat{\mu}_0(X_i)}{1 - \hat{e}(X_i)} \right]$$

$\xrightarrow{p} 0$   
 $\xrightarrow{d} 0$   
 $\xrightarrow{p} 0$

$\rightarrow$  something  
 $\rightarrow$  something

$E_n [\hat{\mu}_1(x) - \hat{\mu}_0(x)]$   
 $\rightarrow \mu_1(x) - \mu_0(x)$   
 $\rightarrow E[\mu_1(x) - \mu_0(x)]$   
 $= E[E(Y|1) - E(Y|0) | X]$   
 $= \tau(x)$   
 $= \tau$

# Influence Function/AIPW - Correctly Specified Propensity

$$\hat{\tau}^{AIPW} = \left[ \frac{1}{N} \sum_{i=1}^N W_i \frac{Y_i}{\hat{e}(X_i)} - (1 - W_i) \frac{Y_i}{1 - \hat{e}(X_i)} \right] + \left[ \frac{1}{N} \sum_{i=1}^N (\hat{e}(X_i) - W_i) \left( \frac{\hat{\mu}_1(X_i)}{\hat{e}(X_i)} + \frac{\hat{\mu}_0(X_i)}{1 - \hat{e}(X_i)} \right) \right]$$

$\mathbb{E}PW$   
 $\xrightarrow{P} \tau$   
 $\mathbb{E}[\cdot] = \mathbb{E}[Y_i | V_i = 1] - \mathbb{E}[Y_i | V_i = 0]$

$\mathbb{E}[e(X_i)] = N_1$   
 $\mathbb{E}[\cdot] = 0$   
 $\xrightarrow{P} 0$

## Cross-Fitting/Sample Splitting

**Idea:** Deriving formal properties easier if we estimate nuisances on separate sample from  $z$  (weak independence)

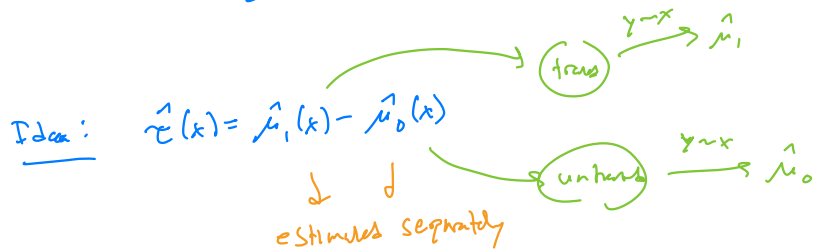
Procedure:

- 1 split sample into  $K$  folds,  $J_i = \{1, \dots, K\}$
- 2 Estimate functions  $\hat{\mu}_w^j(x)$ ,  $\hat{e}^j(x)$  on sample with  $J_i \neq j$
- 3 Estimate  $\hat{z}^j(\hat{\mu}_w^j, \hat{e}^j)$  (AIPW)
- 4 Average over all  $\hat{z}^j$  to get an efficient estimator

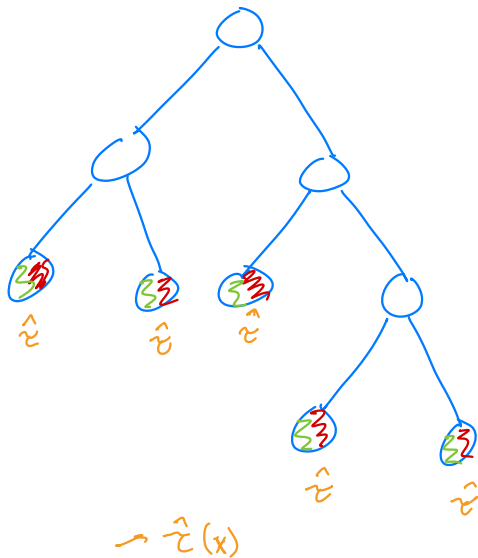
1. Announcements
2. Double Machine Learning
3. CATE Estimation
4. Introduction to Linear Stationary Processes
5. Practice Problems

Big Idea: Maybe observables play some role in TE

$$\tau(x) = E[Y(1) - Y(0) | x]$$



Problem: may not be comparing "alike" individuals



1. Announcements
2. Double Machine Learning
3. CATE Estimation
4. Introduction to Linear Stationary Processes
5. Practice Problems

- **Stochastic Process:** Sequence of RVs ordered by discrete time index

$$\{Y_t\} = [Y_0, Y_1, \dots, Y_T] \rightarrow \text{1 observation}$$

- **Expectation of  $t$ -th observation:**

$$E[Y_t] = \mu_t = \int_{-\infty}^{\infty} y_t f_{Y_t}(y_t) dy_t$$

- **Variance of  $t$ -th observation:**

$$E[(Y_t - \mu_t)^2] = \sigma_{0,t} = \int_{-\infty}^{\infty} (y_t - \mu_t)^2 f_{Y_t}(y_t) dy_t$$

- **Autocovariance (of order  $j$ ):**

$$\gamma_j = E[(Y_t - \mu_t)(Y_{t-j} - \mu_{t-j})] = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \dots \int_{-\infty}^{\infty} (y_t - \mu_t)(y_{t-j} - \mu_{t-j}) f(y_t, y_{t-1}, \dots, y_{t-j}) \times dy_t dy_{t-1} \dots dy_{t-j}$$

A time series  $Y_t$  is **covariance-stationary** (or **weakly stationary**) if the following hold:

1  $\forall t, \mu_t = \mu$

2  $\forall t, \gamma_{j,t} = \gamma_j$

3  $\gamma_0 < \infty$

**Why is this useful?**

Reduces # moments have to estimate

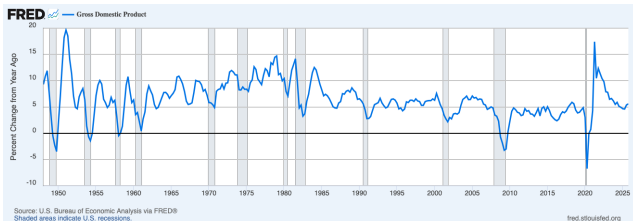
Can now estimate mean from 1 observation

# Stationarity: Examples

Are each of the following covariance-stationary or not?



No!  
Mean changes over time  
(levels)



Yes  
Means roughly constant  
(changes)

- A covariance-stationary process is **ergodic for the mean** if:

$$\lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=1}^T y_t = \mu$$

Sufficient condition for ergodicity:

$$\sum_{j=0}^{\infty} |r_j| < \infty$$

- A covariance-stationary process is **ergodic for second moments** if:

$$\lim_{T \rightarrow \infty} \frac{1}{T-j} \sum_{t=j+1}^T (y_t - \bar{y}(T)) (y_{t-j} - \bar{y}(T)) = r_j$$

↪ ensures  
"diverging over our  
time"

**Why is this useful?** Allows us to swap ensemble averages for time averages.  
Need info. to accumulate over time.

Are each of the following ergodic or not?

- Long-term inflation: inflation rates in an economy where the central <sup>bank</sup> ~~rate~~ has a specific target in mind

Yes! Mean-reversion means things don't "blow up"

- Multiplicative gambling: individuals play a game where each round they win 50% <sup>w>1</sup> of their wealth if the coin is heads and lose 40% of their wealth if the coin is tails

No!

$$\textcircled{1} E[\text{payoff}] = \frac{1}{2} \cdot 1.5 + \frac{1}{2} \cdot 0.6 = 1.05$$

$$\textcircled{2} \text{Payoff}(T) = \prod_{t=1}^T r_t w(t) \rightarrow \bar{r}_T \cdot \underbrace{w_1}_{=1} \rightarrow \bar{r}_T \rightarrow (1.5^{1/2} \cdot 0.6^{1/2}) \approx 0.95$$

A basic building block of time series, white noise has the following properties:

- $E[\varepsilon_t] = 0$
- $E[\varepsilon_t^2] = \text{Var}(\varepsilon_t) = \sigma^2$
- $E[\varepsilon_t \varepsilon_s] = 0 \quad \forall t \neq s$

## Moving Average (MA) Processes: MA(1)

**Definition:**  $Y_t = \mu + \varepsilon_t + \theta \varepsilon_{t-1}$  ,  $\varepsilon_t \sim \mathcal{VN}(0, \sigma^2)$

**Mean:**  $E[Y_t] = \mu$

**Variance:**  $E[(Y_t - \mu)^2] = (1 + \theta^2) E[\varepsilon_t^2] = (1 + \theta^2) \sigma^2$

**jth Autocovariance:**  $\gamma_j = \text{Cov}(Y_t, Y_{t-j}) = \theta \sigma^2 \cdot \mathbb{1}_{\{j \in \{-1, 1\}\}}$

**Definition:** Given  $Y_t$ ,  $L Y_t = Y_{t-1}$

**For higher orders:**  $L^j Y_t = Y_{t-j}$

**Polynomials:**  $\gamma(L) = \gamma_0 + \gamma_1 L + \gamma_2 L^2 + \dots$

**Rewrite MA(1):**  $Y_t = (1 + \theta L) \varepsilon_t + \mu$   
 $= \theta(L) \varepsilon_t + \mu$

## Moving Average (MA) Processes: MA(q)

**Definition:** 
$$Y_t = \mu + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \dots + \theta_q \varepsilon_{t-q} = \mu + \theta(L) \varepsilon_t$$

**Mean:**  $\mu$

**Variance:**  $(1 + \theta_1^2 + \dots + \theta_q^2) \sigma^2 = \gamma_0$

**jth Autocovariance:**

$$\gamma_j = \begin{cases} 0 & \text{if } j \geq q \\ (\theta_j + \theta_{j+1}\theta_1 + \dots + \theta_q\theta_{q-j})\sigma^2 & \text{if } j=1,2,\dots,q \end{cases}$$

**Why is this useful?**

**A necessary assumption for covariance-stationarity:**

$$\sum_{j=1}^{\infty} \theta_j^2 < \infty$$

## Autoregressive (AR) Processes: AR(1)

**Definition:**  $Y_t = c + \phi Y_{t-1} + \varepsilon_t$

**2 strategies for finding MA( $\infty$ ) representation:**

1 Recursive substitution

2 Inverse of polynomial

## Strategy #2: Inverse of Polynomial

Define  $\textcircled{+}_L(L) = 1 + \theta_1 L + \dots + \theta_L L^L$

Inverse:  $\psi(L) = \textcircled{+}_L^{-1}(L) \rightarrow \psi(L) \textcircled{+}_L(L) = 1$

Note:  $Y_t = C + \phi L Y_t + \varepsilon_t \rightarrow (1 - \phi L) Y_t = C + \varepsilon_t$

Need:  $(1 + \gamma_1 L + \gamma_2 L^2 + \dots)(1 - \phi L) = 1$   
 $\rightarrow 1 + (\gamma_1 - \phi)L + (\gamma_2 - \gamma_1 \phi)L^2 + (\gamma_3 - \gamma_2 \phi)L^3 + \dots = 1$

Yields system of equations:

$$\begin{aligned}\gamma_1 - \phi &= 0 \\ \gamma_2 - \gamma_1 \phi &= 0 \\ \gamma_3 - \gamma_2 \phi &= 0\end{aligned}$$

$$\left. \begin{aligned}\gamma_1 - \phi &= 0 \\ \gamma_2 - \gamma_1 \phi &= 0 \\ \gamma_3 - \gamma_2 \phi &= 0\end{aligned} \right\} \gamma_j = \phi^j$$
$$\rightarrow Y_t = \frac{C}{1 - \phi} + \sum_{j=0}^{\infty} \phi^j \varepsilon_{t-j}$$

## Companion Form: Overview

$$\mathbf{z}_t = \begin{pmatrix} y_t \\ \vdots \\ y_{t-p+1} \end{pmatrix}$$

$$F = \begin{pmatrix} \phi_1 & \phi_2 & \dots & \phi_{p-1} & \phi_p \\ 1 & 0 & \dots & 0 & 0 \\ 0 & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & 1 & 0 \end{pmatrix}$$

$$\eta_t = \begin{pmatrix} \varepsilon_t \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

$\rightarrow \mathbf{z}_t = F \mathbf{z}_{t-1} + \eta_t \rightarrow$  our AR(p) in top left corner

$$\rightarrow \mathbf{z}_t = \sum_{j=0}^T F^j \eta_{t-j} + F^{T+1} \mathbf{z}_{t-T-1}$$

$\downarrow$   
gets ugly

$$F = T \Delta T^{-1} \rightarrow F^j = T \Delta^j T^{-1}$$

As long as  $\max_i |\lambda_i| < 1$ , stable

1. Announcements
2. Double Machine Learning
3. CATE Estimation
4. Introduction to Linear Stationary Processes
5. Practice Problems

# Econ 271/Mgtecon 604 - Section 9

Amar Venugopal

Stanford University

Winter 2026

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems

1. Announcements

2. Forecasting & Mean Dynamics

3. Wold Representation Theorem

4. VAR and DSGE

5. State Space Representation & Kalman Filter

6. Practice Problems

- Final Exam: March 18, 3:30-6:30pm
- Last problem set: due March 13, 11:59pm
- First half review session: March 11, 4-5:30pm, GSB BC104

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems

**Goal:** Given unemployment so far, what will unemployment be tomorrow?

$$E[Y_{t+1} | Y_t, Y_{t-1}, \dots]$$

**(Linear) Forecasting:**

$$X_t = [Y_t, Y_{t-1}, \dots] \rightarrow Y^t$$

Forecast:  $g(X_t) = \alpha' X_t \rightarrow$  linear restriction

Implied loss:  $E[(Y_{t+1} - g(X_t))^2]$

$$\rightarrow \text{BLP} \quad \alpha = E^{-1}[X_t' X_t] E[X_t' Y_{t+1}]$$

$$Y_t = \mu + \psi(L) \varepsilon_t$$

**Goal:** Since all other models have MA( $\infty$ ) representation, solving this solves many

$$X_t = \begin{pmatrix} \varepsilon_t \\ \varepsilon_{t-1} \\ \vdots \end{pmatrix}$$

$$E^{-1}[X_t' X_t] = (\sigma^2 I)^{-1}$$

$$E[X_t' Y_{t+s}] = E[X_t' (\mu + \varepsilon_{t+s} + \psi_1 \varepsilon_{t+s-1} + \dots)] = \begin{pmatrix} \psi_s \sigma^2 \\ \psi_{s+1} \sigma^2 \\ \vdots \end{pmatrix}$$

$$P(Y_{t+s} | X_t) = \mu + X_t' \alpha = X_t' (\sigma^2 I)^{-1} \begin{pmatrix} \psi_s \sigma^2 \\ \psi_{s+1} \sigma^2 \\ \vdots \end{pmatrix} = \mu + \psi_s \varepsilon_t + \psi_{s+1} \varepsilon_{t-1} + \dots$$

$$E[\varepsilon_t (\mu + \varepsilon_{t+s} + \psi_1 \varepsilon_{t+s-1} + \dots)]$$

~~Goal:~~

$$X_t = \begin{pmatrix} \varepsilon_t \\ \varepsilon_{t-1} \\ \vdots \end{pmatrix}$$

$$Y_{t+s} = \mu + \varepsilon_{t+s} + \gamma_1 \varepsilon_{t+s-1} + \dots$$

$$X_t' Y_{t+s} = \left( \varepsilon_t \left( \mu + \sum_{j=0}^{\infty} \gamma_j \varepsilon_{t+j} \right) \right)$$

# Forecasting: MA( $\infty$ )

Convert  $P(Y_{t+s} | \varepsilon_t, \varepsilon_{t-1}, \dots)$  to  $P(Y_{t+s} | Y_t, Y_{t-1}, \dots)$

Forecast error:  $Y_{t+s} - P(Y_{t+s} | X_t) = \varepsilon_{t+s} + \gamma_1 \varepsilon_{t+s-1} + \dots + \gamma_{s-1} \varepsilon_{t+1}$

Define:  $\frac{\gamma(L)}{L^s} = L^{-s} + \gamma_1 L^{-(s-1)} + \dots$

$$P(Y_{t+s} | X_t) = \mu + \left[ \frac{\gamma(L)}{L^s} \right] \varepsilon_t$$

$$\left( \frac{\gamma(L)}{L^s} \right) \varepsilon_t = L^{-s} \varepsilon_t + \gamma_1 L^{-(s-1)} \varepsilon_t + \dots = \varepsilon_{t+s} + \gamma_1 \varepsilon_{t+s-1} + \dots$$

$\oplus \rightarrow$  sets negative powers of  $L$  to 0

From polynomial inversion:

$$\underbrace{\eta(L)}_{=\gamma^{-1}(L)} (Y_{t-m}) = \varepsilon_t \rightarrow P(Y_{t+s} | Y_t, \dots) = \mu + \left[ \frac{\gamma(L)}{L} \right]_{\oplus} \eta(L) (Y_{t-m})$$

$\hookrightarrow$  linear projection for MA( $\infty$ )

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems

# Wold Representation Theorem

**Theorem:** Let  $\{Y_t\}$  be a covariance-stationary process with  $\mathbb{E}[Y_t] = 0$  and  $\gamma_j = \mathbb{E}[Y_t Y_{t-j}]$ . Then:

$$Y_t = \sum_{j=0}^{\infty} \psi_j \varepsilon_{t-j} + \eta_t$$

Where:

- $\varepsilon_t \sim$  white noise
- $\psi_0 = 1$ ,  $\sum_{j=0}^{\infty} \psi_j^2 < \infty$
- $\forall t, s \quad \mathbb{E}[\varepsilon_t \eta_s] = 0$
- $\eta_{t,s}$  perfectly predictable given  $Y_t$

# Wold Decomposition

A consequence of this is that we can decompose any covariance-stationary process into 2 components:

$$Y_t = \sum_{j=0}^{\infty} \psi_j \varepsilon_{t-j} + \eta_t$$

1  $\sum_{j=0}^{\infty} \psi_j \varepsilon_{t-j}$  : stochastic component, impossible to predict perfectly

2  $\eta_t$  : deterministic, perfectly predictable from history (seasonality)

→ holds for any covariance-stationary process (not just linear ones)

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems

■ **Vector Process:**  $\{Y_t\} = [Y_{1,t}, Y_{2,t}, \dots]$

■ **Mean:**  $\mu_t = E[Y_t]$

→ independent of  $t$  if covariance-stationary

■ **Variance:**  $\Gamma_{0,t} = E[(Y_t - \mu_t)(Y_t - \mu_t)']$

■  $j^{th}$  **Order Autocovariance:**  $\Gamma_{j,t} = E[(Y_t - \mu_t)(Y_{t-j} - \mu_{t-j})']$

**Definition:** 
$$Y_t = \mu + \varepsilon_t + \gamma_1 \varepsilon_{t-1} + \dots$$

Annotations:

- $\mu$ : vector
- $\varepsilon_t$ : vector
- $\gamma_1$ : matrix
- $\varepsilon_{t-1}$ : vector

No restrictions  $\gamma_t \rightarrow$  can have nonzero coefficients on innovations from other sub-processes

**When is it covariance stationary?**

$$\forall i, j, \sum_{s=0}^{\infty} |\gamma_s^{(i,j)}| < \infty$$

**Definition:**  $Y_t = c + \Phi_1 Y_{t-1} + \dots + \Phi_p Y_{t-p} + \varepsilon_t$


  
 multiplies

No restrictions on  $\Phi_p$

When is it covariance stationary?

$$F = \begin{pmatrix} \Phi_1 & \Phi_2 & \dots & \Phi_{p-1} & \Phi_p \\ I_n & 0 & \dots & 0 & 0 \\ 0 & I_n & \dots & 0 & 0 \\ \vdots & \vdots & \dots & \vdots & \vdots \\ 0 & 0 & \dots & I_n & 0 \end{pmatrix}$$

$(n \times p) \times (n \times p)$

→ all eigenvalues must lie in unit circle

## Impulse Response Function

**Goal:** See how other variables respond to an unexpected shock in one variable.

**Consider MA( $\infty$ ) representation:**  $Y_t = \mu + \gamma(L) \varepsilon_t$

Can interpret  $\gamma_s = \frac{\partial Y_{t+s}}{\partial \varepsilon_t^j}$   $\rightarrow$  how change in innovation to  $j$  at  $t$  impacts the whole process  $s$  periods later

**Can we use this to compare across outcome variables?**

No! This holds all other innovations fixed

But  $\varepsilon_{1,t}$  and  $\varepsilon_{2,t}$  may be correlated

$$\varepsilon_t \sim N(0, \Sigma)$$

## Orthogonalizing IRFs

If  $\Sigma$  diagonal, then  $\varepsilon$ 's orthogonal

Cholesky Factorization  $\Sigma = PP'$ ,  $P$  lower triangular

Define  $u_t = P^{-1} \varepsilon_t$

$$E[u_t u_t'] = P^{-1} P P' P^{-1} = \underline{I} \rightarrow \text{orthogonalized}$$

$\hookrightarrow$  implies hierarchy in  
innovations  $\varepsilon_t$   
(causal flow)

$\rightarrow$  order matters

Can now write  $MA(\infty)$  as:  $y_t = \mu + \rho u_t + \gamma_1 \rho u_{t-1} + \dots$   
 $\rightarrow$  here, IRFs orthogonalized

**Does this really help with interpretation?**

Not really ...  $u_t$  have no economic meaning (unless before ordering is true)

**Big Idea:** Many macro models can be represented as VAR  
using laws of motion for state variables as function of  
structural parameters  $\exists$  mapping: structural model  $\rightarrow$  VAR (restricted)

Justifies 2 main uses of VAR in applied macro:

- 1 Reference models
- 2 Structural VARs

**Goal:** Approximate "structure" by imposing restrictions on covariance matrix of VAR innovations

**Identification challenge:**

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems

Let  $\mathbf{z}_t$  be the state vector and  $\mathbf{y}_t$  be the vector of observations.

Gaussian state space model has 2 parts:

- **Law of motion for observables:**

$$\mathbf{y}_t = \boldsymbol{\mu} + \mathbf{A}\mathbf{z}_t + \mathbf{H}\boldsymbol{\eta}_t, \quad \boldsymbol{\eta}_t \sim N(\mathbf{0}, \mathbf{Z})$$

- **Law of motion for state:**

$$\mathbf{z}_t = \boldsymbol{\Phi}\mathbf{z}_{t-1} + \mathbf{R}\boldsymbol{\varepsilon}_t, \quad \boldsymbol{\varepsilon}_t \sim N(\mathbf{0}, \boldsymbol{\Sigma})$$

↳ Markov

# Example: MA(1)

$$z_t = (\epsilon_t, \epsilon_{t-1})$$

$$y_t = \mu + \epsilon_t + \theta \epsilon_{t-1} \quad \epsilon_t \sim \mathcal{WN}(0, \sigma^2)$$

$$y_t = \underline{\mu} + \underline{A} z_t + \underline{H} \eta_t$$

$$\mu = \mu$$

$$\underline{Q} = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$$

$$z_t = \underline{Q} z_{t-1} + \underline{R} \epsilon_t$$

$\begin{matrix} \downarrow \\ \epsilon_t \end{matrix}$

$$\underline{A} = (1 \quad 0)$$

$$\underline{R} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

$$\begin{pmatrix} \epsilon_t \\ \epsilon_{t-1} \end{pmatrix} = \underline{Q} \begin{pmatrix} \epsilon_{t-1} \\ \epsilon_{t-2} \end{pmatrix} + \underline{R} \epsilon_t$$

$$\underline{H} = 0$$

$$\Sigma = \sigma^2$$

$$= \begin{pmatrix} 0 \\ \epsilon_{t-1} \end{pmatrix} + \underline{R} \epsilon_t$$

**Big Idea:** Estimate state from observed time series via Bayesian updating

Consider the following simple example:  $y_t = z_t + \sigma_y \eta_t$ ,  $\eta_t \sim \mathcal{N}(0,1)$

Belief over state:  $p(z_t) \sim \mathcal{N}(\hat{z}_{t|t}, P_{t|t})$   $z_t = \phi z_{t-1} + \sigma_z \epsilon_t$ ,  $\epsilon_t \sim \mathcal{N}(0,1)$

There are 3 steps to produce an update: Given state  $\hat{z}_{t|t-1}$ ,  $P_{t|t-1}$

1 Predict  $y_t \leftarrow$  law of motion for observable

2 Update beliefs  $\hat{z}_{t|t}$ ,  $P_{t|t}$

3 Update beliefs about  $\hat{z}_{t+1|t}$ ,  $P_{t+1|t} \leftarrow$  law of motion for state

1. Announcements
2. Forecasting & Mean Dynamics
3. Wold Representation Theorem
4. VAR and DSGE
5. State Space Representation & Kalman Filter
6. Practice Problems