

Econ 272/Mgtecon 607 - Section 1

Amar Venugopal

Stanford University

Spring 2025

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

- Email: `amar.venugopal@stanford.edu`
- Section: Thursdays 4:30-6:20pm in STLC 114
- Office Hours: Wednesdays 2-4pm in Econ 149
- Section Questions: [Question Form](#)
- 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 in past years.

- Remember to email me about whether you wish to take the final exam, and if you require any special accommodations for it
 - Unless you are an Econ PhD, in which case you have no choice :)
 - If you are a predoc and wish to waive the class next year, you must take the exam as well

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 **Sunday April 6, 11pm**. No late submissions are accepted.
- All submissions should consist of 2 files:
 - A **typed** writeup, which should be self-contained (the grader should not have to view any code)
 - A .zip file with all your code and user-created programs
- Please make sure your writeup is easy to read, and code should be well-commented
- For empirical questions: try to make economic sense of your results, or explain why they may not make sense
- For theoretical questions: make sure to explain steps clearly (cite theorems, references, etc) for full credit
- For problem set-related questions, come to my office hours or reach out to me (cc Guido on all problem set-related emails)
 - I **will not** respond to problem set emails sent on the due date

- The lecture notes (on Canvas) may be helpful if you get stuck at estimating homoskedastic variance

Guido and I are both committed to making lecture and section as helpful/instructive as possible.

Feel free to discuss or email either of us with any feedback you might have.

Can also you section feedback google form if you prefer anonymous submission of feedback (just note in your response that the feedback pertains to lecture).

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

Key question: what is the effect of changing someone's treatment status on some outcome?

3 key areas of interest:

- 1 Point estimator ($\hat{\tau}$)
- 2 Inference ($\hat{V}$)
- 3 Hypothesis testing

Fundamental problem of causal inference:

we never observe counterfactuals

i denotes unit

- Y_i : Outcome
- W_i : Treatment $\rightarrow W_i \in \{0, 1\}$
- $Y_i(w)$: Potential outcome at $W_i = w$
- X_i : Covariates

Note: Via Neymans's repeated sampling approach, PO are assumed to be *fixed* (unlike observed outcomes)

$$\begin{aligned} \text{Var}(Y_0) &= \text{Var}(W_0 Y_0 + (1 - W_0) Y_0) = \underbrace{\text{Var}(W_0 Y_0)}_{\text{Var}} + \text{Var}((1 - W_0) Y_0) \\ &= \text{Var}(W_0 Y_0(1)) \\ &= Y_0^2(1) \text{Var}(W_0) \end{aligned}$$

- Known treatment assignment $(Y_i(0), Y_i(1)) \perp W_i$

- Stable Unit Treatment Value Assumption $Y_i(\vec{w}) = Y_i(w_i)$

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

- Estimand: $\tau^{ATE} = E[Y_i(1) - Y_i(0)]$
 $\tau^{ATT} = E[Y(1) - Y(0) | v=1]$
- Estimator: $\hat{\tau}^{DiM} = \bar{Y}_1^{obs} - \bar{Y}_0^{obs}$
- Properties: Unbiased $E[\hat{\tau}^{DiM}] = \tau^{ATE}$

$$\hat{\tau}^{ATE}$$

- True Variance:

$$Var(\tau) = \frac{S_1^2}{n_1} + \frac{S_0^2}{n_0} - \frac{S_{01}^2}{N} \quad \sum (Y_i(1) - Y_i(0) - \tau)^2$$

inestimable

$$\tilde{Var}(\tau^{ATE}) = \frac{\tilde{S}_0^2}{n_0} + \frac{\tilde{S}_1^2}{n_1}$$

- Estimator:

$$\tilde{\tilde{Var}}(\tau^{ATE}) = \frac{\hat{S}_1^2}{n_1} + \frac{\hat{S}_0^2}{n_0}$$

- Properties:

Unbiased + consistent for $\tilde{Var}(\tau)$

Bias + inconsistent for $Var(\tau)$

- Confidence Intervals:

Asymptotically normal CLT

$$\left[\hat{\tau} - 1.96 \sqrt{\hat{Var}}, \hat{\tau} + 1.96 \sqrt{\hat{Var}} \right]$$

Hypothesis Testing: Fisher Exact p -Values

4-step procedure:

$$H_1: Y_i(1) = aY_i(0) + b$$

1 Define a sharp null H_0
↳ $H_0: Y_i(0), Y_i(1)$ under H_0

2 Define test statistic, T ($T = \bar{Y}_1 - \bar{Y}_0$)

3 Calculate distribution of T (exact numbers)



4 Calculate p -value $p = \frac{\sum_{a \in A} \mathbb{1}_{|T^a| \geq |T^{obs}|}}{|A|}$

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

What is the problem?

Assumed correctly



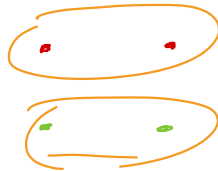
Chris:



Stouffer:



Pared:



■ Estimand: $\tau^{\text{strat}} = E[Y_0(T) - Y_0(C)] = E[\tau(X_i)]$

■ Estimator: $\hat{\tau}^{\text{strat}} = \sum_{g=1}^G \frac{N_g}{N} \hat{\tau}_g$

$$Y_i^{\text{obs}} = \alpha + \tau W_i + \beta \mathbb{1}\{G_i = g\} + \epsilon_i$$

$$\tau^{\text{RM}} = \bar{Y}_1^{\text{obs}} - \bar{Y}_0^{\text{obs}}$$

■ True Variance:
$$V_{\text{cr}}(\hat{\tau}^{\text{struc}}) = \sum_g \left(\frac{n_g}{n}\right)^2 V_{\sigma}(\hat{\tau}_g)$$

■ Estimator:
$$\widehat{V_{\text{cr}}(\hat{\tau}^{\text{struc}})} = \sum_g \left(\frac{n_g}{n}\right)^2 \widehat{V_{\sigma}(\hat{\tau}_g)}$$

	CRD	Strat	Paired
Pro	Single, easy	Lower MSE vs CRD Heterogeneity	<u>Lowest</u> MSE in finite sample Higher power in large samples
Cons	Doesn't use potentially useful covariates		Can't estimate variance of ATE (within strata)

$$V_C - V_S = A \left(\underbrace{n(x) \dots}_{>0} - \underbrace{\frac{m(\tilde{x})^2}{n}}_{>0} \right) \geq 0 \rightarrow >0$$

If sample large + predictive covs. $\rightarrow$ stable or par

Generally, small shrink $>$ par

1. General Info & Logistics
2. Discussion Questions
3. Intro to Causality
4. Completely Randomized Experiments
5. Stratified/Paired Experiments
6. Power Analyses
7. Practice Problems

Power of a test (β): $P(\text{reject } H_0 \mid H_0 \text{ false})$

Size of a test (α): $P(\text{reject } H_0 \mid H_0 \text{ true})$

General procedure to get required sample size for given α, β :

1 Specify rejection condition (e.g. $|T| \geq \Phi^{-1}(1 - \alpha/2)$)

2 Specify probability = β

3 Solve for N

$$\begin{array}{c} \downarrow \\ P\left(\frac{\bar{X} - \mu_0}{\sigma/\sqrt{N}} \geq \Phi^{-1}(1 - \alpha/2) \right) = \beta \end{array}$$

To determine minimum sample size when...

- Testing a mean against 0 (assuming random sample from normal distribution):

$$N = \left(\frac{\Phi^{-1}(\beta) + \Phi^{-1}(1 - \alpha/2)}{\mu_0/\sigma} \right)^2$$

- Testing DiM with unequal sample sizes and equal variances:

$$N = \frac{(\Phi^{-1}(\beta) + \Phi^{-1}(1 - \alpha/2))^2}{(\tau_0^2/\sigma^2)\gamma(1 - \gamma)}$$

where γ = share of treated units

Econ 272/Mgtecon 607 - Section 2

Amar Venugopal

Stanford University

Spring 2025

1. Announcements
2. Discussion Questions
3. Clustered Experiments
4. Clustering in Sampling and in Assignment
5. Practice Problems

1. Announcements

2. Discussion Questions

3. Clustered Experiments

4. Clustering in Sampling and in Assignment

5. Practice Problems

- Remember to email me about whether you wish to take the final exam, and if you require any special accommodations for it
 - Unless you are an Econ PhD, in which case you have no choice :)
 - If you are a predoc and wish to waive the class next year, you must take the exam as well

Hints for Problem Set 2 (Due April 13, 11pm)

- Within-group correlations of a single variable try to answer the question: does *my* value going up mean that *your* value is more likely to go up/down?
 - $\text{Corr}(Y_i, Y_j)$ for $i \neq j$ can be calculated by constructing pairs (combinatorially) of observations from the data

- Recall the definition/equation for correlation - how can it be simplified in this context?

$$\text{Corr}(Y_i, Y_j) = \frac{\tilde{Y}_i = \frac{Y_i - \bar{Y}}{\sigma_Y}}{\frac{\text{Cov}(Y_i, Y_j)}{\sqrt{\text{Var}(Y_i)\text{Var}(Y_j)}}} = \frac{\text{Cov}(\tilde{Y}_i, \tilde{Y}_j)}{\text{Var}(\tilde{Y}_i)} = \underbrace{\bar{E}[\tilde{Y}_i \tilde{Y}_j] - \bar{E}[\tilde{Y}_i] \bar{E}[\tilde{Y}_j]}_{\underbrace{\frac{1}{\binom{N_g}{2}} \sum_{i \neq j} \tilde{Y}_i \tilde{Y}_j}}$$

$\begin{array}{c|c} g & P_g \\ \hline \vdots & \vdots \end{array}$

- Simplified equations for LZ/EHW variance estimators on slides (Lecture 4, Slide 4) are missing a factor of $1/N$

A note on terminology

- In the course, we will often use the term “population” to refer to the experimental sample since we are mainly thinking about design-based inference
- We will therefore refer to the broader population from which the experimental sample is drawn as the “super-population”
- We may occasionally refer to our experimental sample as the “sample”, for example when we want to distinguish between sample and population ATEs
- It is confusing, I will try to be as explicit as possible

From last time: nontrivial sharp nulls

Test statistic for sharp null: $Y_i(1) = a \cdot Y_i(0) + b$

$$\begin{aligned} T &= \bar{\tilde{Y}}_i(1) - \bar{\tilde{Y}}_i(0) \\ &= \bar{Y}_i(1) - a \bar{Y}_i(0) - b \end{aligned}$$

1. Announcements

2. Discussion Questions

3. Clustered Experiments

4. Clustering in Sampling and in Assignment

5. Practice Problems

1. Announcements

2. Discussion Questions

3. Clustered Experiments → Clustered Assignment

4. Clustering in Sampling and in Assignment

5. Practice Problems

Why cluster?

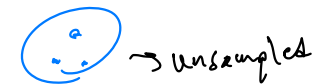
Spillovers, network effects $\rightarrow$ SUTVA violations

My outcome depends on your treatment

Idea: If spillovers occur on specific clusters
treating clusters as units.

Clustered Sampling vs. Assignment

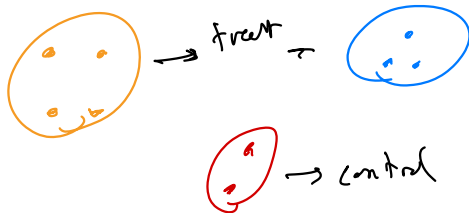
■ Clustered Sampling:



C_1, C_2, C_3
Pool 2

■ Clustered Assignment:

helps
spillovers



- G : Clusters
- $G_i \in \{1, \dots, G\}$: cluster assignment for unit i
- N_g : # units in cluster g ($N = \sum_g N_g$)

■ Estimands:

$$\tau_{\text{ITT}} = \frac{1}{n} \sum_i Y_i(1) - Y_i(0)$$

overall 2x2 ATB

= if $N_g = n \forall g$
if ITT cluster

$$\tau_{\text{cluster}} = \frac{1}{G} \sum_g \left(\frac{1}{n_g} \sum_i Y_i(1) - Y_i(0) \right)$$

overall cluster-level ATB

■ Estimators:

$$\hat{\tau}_{\text{ITT}} = \bar{Y}_1 - \bar{Y}_0$$

had to estimate γ directly

$$\hat{\tau}_{\text{cluster}} = \frac{1}{G_1} \sum_g \left(\frac{1}{n_g} \sum_i w_i Y_i \right) - \frac{1}{G_0} \sum_g \left(\frac{1}{n_g} \sum_i (1-w_i) Y_0 \right)$$

Both estimators are unbiased

Variance: $\hat{\tau}_{cluster}$



$$V_{regression} = \frac{S_0^2}{N_0} + \frac{S_1^2}{N_1} - \frac{S_{01}^2}{N}$$

■ True Variance:

$$Var(\hat{\tau}_{cluster}) = \frac{S_0^2}{G_0} + \frac{S_1^2}{G_1} - \frac{S_{01}^2}{G}$$

■ Estimator:

$$\hat{Var}(\hat{\tau}_{cluster}) = \frac{\hat{S}_0^2}{G_0} + \frac{\hat{S}_1^2}{G_1}$$

$$S_v^2 = \frac{1}{G-1} \sum_g (\bar{Y}_g(v) - \bar{Y}_{g(v)})^2$$

Variance: $\hat{\tau}_{pop}$

Challenge: treatment assignment not indep.

■ True Variance: Asymptotics: sequence of pops. $k=1, 2, \dots$

In pop. k , sample all units, assign each g to treatment v.e. θ

For N_g pops, let $G_k, N_k \rightarrow \infty$:

$$\sqrt{N_k} (\hat{\tau}_{gk} - \tau_{gk}) \xrightarrow{d} N(0, \underbrace{V_k}_{\substack{\text{eqn on} \\ \text{slide} \\ 24}})$$

■ Estimator: Conservative estimate var: $Y_i = \alpha + \beta W_i + \epsilon_i$

$$\underbrace{\hat{V}_{\text{cluster}}}_{\hat{V}_{L2}} = \left(\frac{1}{n} \sum_i \tilde{w}_i^2 \right)^{-1} \left(\frac{1}{n} \sum_i \sum_j (\epsilon_i \tilde{w}_j)^2 \right) \left(\frac{1}{n} \sum_i \tilde{w}_i^2 \right)^{-1}$$

$$(\geq V_{\text{EHW}})$$

$$\tilde{w}_i = w_i - \frac{1}{n} \sum_{j=1}^n w_j$$

Inference based on rand. data induced by
clustered assignments rather than indiv.

Typically rely asymptotically.

But, Fisher exact & STM work, but none of
analyses are clusters

1. Announcements
2. Discussion Questions
3. Clustered Experiments
4. Clustering in Sampling and in Assignment
5. Practice Problems

2 "Flavors" of Error / Variability

1 Sampling - how sample constructed

2 Assignment - how treatment assigned

Don't needed for inference in causal settings
→ Abadie et. al. (2003) (general framework)

(Even More) Notation

Sequences of populations indexed by $k = 1, 2, 3, \dots$

■ g_k : # clusters in k

■ $n_{g,k}$: # units in cluster g in popn k

■ $q_k \in (0, 1]$: cluster sampling probability in popn k $\left\{ \begin{array}{l} q_k = 1 : \text{random sampling} \\ q_k < 1 : \text{clustered sampling} \end{array} \right.$

■ p_k : prob. a unit is sampled from cluster

■ $A_{k,g}$: prob. unit in cluster g in popn k is included: $A_{k,g} \sim (\mu_k, \sigma_k^2)$

not all clusters need be sampled

not all units need same treatment within cluster

} Nest C.S.
+
C.A.

Clustering Regimes

q_k, σ_k^2, μ_k
 $\downarrow$ Sampling
 $\underbrace{\sigma_k^2, \mu_k}_{\text{assignment}}$

		Assignment, $A_{k,g}$		
		Random	Partially Clustered	Clustered
Sampling q_k	Random	$q_k = 1, \sigma_k^2 \geq 0$	$q_k \geq 1,$	$q_k = 1, \sigma_k^2 \geq \mu_k(1+\mu_k)$
	Clustered	$q_k < 1, \sigma_k^2 \geq 0$	$0 < \sigma_k^2 < \mu_k(1-\mu_k)$ $q_k < 1,$	$\downarrow$ $q_k < 1,$

- Random sampling of clusters from a large population of clusters $\rightarrow p_k$ small
- Random sampling from a large population $\rightarrow p_k=1, p_k$ small
 - Completely random assignment $A_{k,g} = A_k$ ($\sigma_k^2 = 0$)
 - Clustered random assignment $A_{k,g} \in \{0, 1\}$
- What can we test?

Can test $p_k = 1$ vs. < 1

Can test $\sigma_k^2 = 0$ vs. ≥ 0

Causal Cluster Variance (CCV) - Big Idea

Asymptotic DM converges: $\sqrt{N_k} (\hat{\tau}_k - \tau_k) \xrightarrow{d} N(0, V_k^{CCV})$

$$V_k^{CCV} = V_k^{CCV}(q_k, \sigma_k^2, q_k, \mu_k)$$

Special cases:

$$V_k^{BHV} \approx V_k^{CCV}(q_k \geq 1, \sigma_k^2 = 0)$$

$$V_k^{LZ} - V_k^{CCV} \geq 0 \rightarrow LZ \text{ conservative}$$

$$V_k \hat{\tau}^{DIM} \text{ to est. } \tau_{PT}$$

$$R_k \quad \hat{\Sigma}^{\text{Dim}} \rightarrow \Sigma_{\text{reg}}$$

		Assignment, $A_{k,g}$		
		Random	Partially Clustered	Clustered
Sampling q_k	Random	L2+W	CCV	L2
	Clustered	CCV	CCV	CCV

$$\downarrow$$

$$\approx L2 \text{ or } q_k \text{ small}$$

$$\text{If } q_k \text{ small, } E(L+W) \approx L2 \approx CCV$$

$$\sigma_z^2 \rightarrow 0, \hat{V}^{ccv} \rightarrow \hat{V}^{\text{EHV}} \quad \sigma_z^2 \rightarrow \mu_k(v\mu_k), \hat{V}^{ccv} \rightarrow \hat{V}^{L2}$$

$$\textcircled{1} \text{ Analytic: } \hat{V}_k^{ccv} = \underbrace{\hat{q}_k}_{\substack{\text{knowledge} \\ \text{on the \# of classes in } \mathcal{Z}_{\text{em}}}} \cdot \hat{V}_k^{ccv}(1) + \overset{L_k \rightarrow 0}{(1 - \hat{q}_k)} \hat{V}_k^{L2}$$

$$\textcircled{2} \text{ Bootstrap: } \text{TSCB}$$

Variance Estimation by Regime

		Assignment, $A_{k,g}$		
		Random	Partially Clustered	Clustered
Sampling q_k	Random	$\hat{V}^{ELW}$	$\hat{V}^{CCV} \rightarrow \hat{V}^{LZ} \rightarrow \hat{V}^{CCV}(1)$	$\hat{V}^{LZ}$
	Clustered	$\hat{V}^{CCV}$	$\hat{V}^{CCV}$	$\hat{V}^{CCV}$

Note: the choice of estimator really matters here!! Standard errors can vary widely. TSCEB

LZ: if wanted unbiased class on trying to estimate

ELW: if drawing clusters local to some

CCV: if imbalanced, or if you use clustered sampling

not available

Econ 272/Mgtecon 607 - Section 3

Amar Venugopal

Stanford University

Spring 2025

1. Discussion Questions
2. Multi-Armed Bandits
3. Identification
4. Practice Problems

1. Discussion Questions

2. Multi-Armed Bandits

3. Identification

4. Practice Problems

1. Discussion Questions

2. Multi-Armed Bandits

3. Identification

4. Practice Problems

What's different from traditional experiments?

- **Setting:** 1 experimental sample $\rightarrow$ sequential samples ("online")

Inference is hard

- **Treatment Arms:** 2 $\rightarrow$ K (> 2)

- **Goal:** Estimation $\rightarrow$ right decision

"Dual Mandate" for Bandits

- Exploration : get a good sense of each arm's efficacy

- Exploitation : give units best possible treatment

→ balance these to minimize regret (suboptimal assignments)
give as many units as possible best treatment possible

- **Action/Arm:** $k_t \in \{1, \dots, K\}$ (a_t)
- **Outcome/Reward:** $Y_k \in \mathbb{R}$ (r_t)
- **Policy:** $\pi : \text{history} \rightarrow \text{action}$
 $\pi(k_1, Y_1, \dots, k_{t-1}, Y_{t-1}) = k_t$

Naive Approach

Given: N units, K arms

- 1 Assign N/K units per arm, observe avg. outcome
- 2 Choose highest avg. arm: $k^* = \arg \max_k \bar{Y}_k$

Drawbacks:

- risk averse
 - no priors
 - equal confidence in each arm $\rightarrow$ could be wrong about best
- no exploiting

Big idea: "learn as we go"

Dynamically assign units to arms to balance explore vs. exploit

2 common methods/algorithms:

- 1 Thompson Sampling (TS)
- 2 Upper Confidence Bound (UCB)

Quick Detour: Bayesian Updating

1 Form prior : $p(\theta)$

2 Observe data : $p(x|\theta)$

3 Update via Bayes' Rule : $p(\theta|x) = \frac{p(\theta, x)}{p(x)} = \frac{p(x|\theta) p(\theta)}{p(x)}$

$\underbrace{\hspace{1cm}}_{\text{posterior}} \quad \underbrace{\hspace{1cm}}_{\text{likelihood}} \quad \underbrace{\hspace{1cm}}_{\text{prior}}$

Incalculable analytically $\rightarrow$ favor conjugate priors

distributions posterior, prior same class
ex. Normal, Beta (uniform)

Thompson Sampling: Idea

- 1 **Select:** according to probability arm is optimal
 - 2 **Update:** observe Y_t , update posterior for selected arm
- **Exploration:** Each arm has nonzero probability of being selected
 - **Exploitation:** Choosing arms we think are better more often
- Modification: Exploration Sampling - choose according to variance
↑ Exploration, ↓ Exploitation

Thompson Sampling: Example

$$\mu_k = P(Y_k = 1)$$

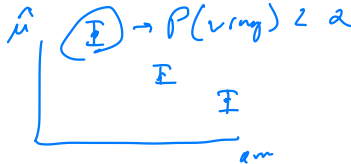
$K = 3$ arms, binary rewards, $\mu_1 = 0.5, \mu_2 = 0.2, \mu_3 = 0.9$, uniform prior: $B(1, 1)$

1 Sample from prior
 $\hat{\mu}_1 = 0.43, \hat{\mu}_2 = 0.75, \hat{\mu}_3 = 0.11 \rightarrow B(1, 1) \quad \underbrace{U(0, 1)}$

2 Assign to highest: $k_1 = 2$, observe failure ($Y = 0$)
 Update: $\beta_2 = \beta_2 + 1 \rightarrow B(1, 2) \rightarrow \text{supp} = [0, 1]$

3 Sample posterior: $\underbrace{\hat{\mu}_1 = 0.52}_{\text{prior}}, \underbrace{\hat{\mu}_2 = 0.27}_{\text{posterior}}, \underbrace{\hat{\mu}_3 = 0.86}_{\text{prior}}$

4 $k_2 = 3$, observe, update, ...
 $\vdots$



Upper Confidence Bound: Idea

1 **Select:** Calculate CI for μ_k , choose arm with highest UB

2 **Update:** $UB_t^{(k)} = \hat{\mu}_{k_t} + \sqrt{\frac{2 \log(1/\epsilon)}{T(k, t)}}$

$\frac{2 \log(1/\epsilon)}{T(k, t)}$ → "confidence level" ($= 1/\epsilon$)
→ # pulls
→ wider CI over time
↑ Exploration
↓ Exploitation

■ **Exploration:** More uni's → narrower CI

■ **Exploitation:** As we get better estimates of means, CIs shift accordingly



Upper Confidence Bound: Example

$K = 3$ arms, binary rewards, $\mu_1 = 0.5, \mu_2 = 0.2, \mu_3 = 0.9$

- $t=1$ **1** Pull each once to get initial estimates: $\hat{\mu}_1 = 0, \hat{\mu}_2 = 0, \hat{\mu}_3 = 1$
- $t=2$ **2** Calculate UCB: $UB_1 = 0 + \sqrt{\frac{2 \ln(2)}{1}} = ?$ $UB_3 = 1 + \underbrace{\sqrt{\frac{2 \ln(2)}{1}}}_{\text{larger}} = ?$
($q = 1/t$)
- 3** Pull arm 3, observe, update UB, ...
- ;

■ **Regret:**

$$R(\pi) = \mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T Y_t \mid \pi^A \right] - \mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T Y_t \mid \pi \right]$$

■ **Lower Bound:**
$$\mathbb{E}[R(\pi)] \geq \frac{\ln(T)}{T} \sum_{k \neq k^A} \frac{\mu^A - \mu_k}{KL(\mu_k, \mu^A)} + o(1)$$

① Decreasing in T

② Small if μ_k v. diff. μ^A

③ Hard to learn if $\mu_k \sim \mu^A$

$\uparrow$ close

■ **TS Regret:**
$$\mathbb{E}[R(\pi^{TS})] = O \left(\frac{\ln(T)}{T} \sum_{k \neq k^A} \frac{1}{(\mu^A - \mu_k)^2} \right)$$

- What if we observe covariates?

→ Lin UCB

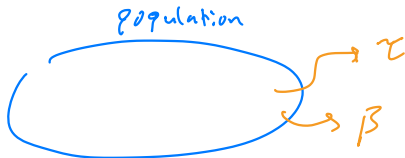
- Inference?

1. Discussion Questions

2. Multi-Armed Bandits

3. Identification

4. Practice Problems



- **Question:** Given "enough data", can we learn it?
- **Challenge:** What is "enough data"?

3 Types of Identification

1 Point

2 Partial / Set

3 None

Focus on cross-sectional data. Consider sequence of populations $k = 1, \dots, \infty$

Point Identification - Example

OLS, $\theta = \beta$.

Assume: $y_i | x_i \sim \mathcal{N}(\underbrace{\alpha + \beta x_i}_{\theta}, \sigma^2)$

$$\beta = \frac{\text{Cov}(x_i, y_i)}{\text{Var}(x_i)} \rightarrow \text{estimable from data}$$

$$\rightarrow \mathcal{N}(\alpha + \beta x_i + \underbrace{\gamma\left(\frac{x_i}{2}\right)}_{\text{not ID}}, \sigma^2)$$
$$(\beta + \frac{\gamma}{2})$$

Partial Identification - Example

$W_i, Y_i \in \{0, 1\}$, observe outcome only for treated units. $\theta = E[Y_i] = P(Y_i = 1)$.

Never observe: $P(Y_i | W_i = 0) \in [0, 1]$

$$LB \leq \theta \leq UB$$

Econ 272/Mgtecon 607 - Section 4

Amar Venugopal

Stanford University

Spring 2025

1. Announcements
2. Directed Acyclic Graphs (DAGs)
3. Practice Problems
4. eCONometrics? Evolution of Views

1. Announcements

2. Directed Acyclic Graphs (DAGs)

3. Practice Problems

4. eCONometrics? Evolution of Views

- Final Exam: Wednesday June 11, 8:30-11:30am (room TBD)
- Remember to email me about whether you wish to take the final exam, and if you require any special accommodations for it
 - Unless you are an Econ PhD, in which case you have no choice :)
 - If you are a predoc and wish to waive the class next year, you must take the exam as well
- Assignment 5 - Due May 11 (skipping a week, good luck with midterms!)

1. Announcements

2. Directed Acyclic Graphs (DAGs)

3. Practice Problems

4. eCONometrics? Evolution of Views

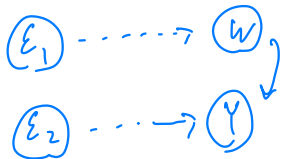
Consider a simple randomized experiment.

■ Structural Equation:

$$W = f_1(\varepsilon_1) \qquad Y = f_2(W, \varepsilon_2)$$



■ DAG:



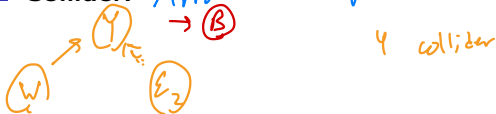
Why DAGs? Help (visually) argue for identification



- **Nodes:** Quantities/variables (A, B, C)
- **Parent/Child:** Immediate predecessor/successor
- **Ancestor:** (Non)immediate predecessor
Descendant
- **(Dashed) Arrows:** (Unobserved) links
- **Directed Path:** Path with all arrows in same direction
 $C \rightarrow A \rightarrow B$ dir. path
 $C \leftarrow A \rightarrow B$ path

Types of Nodes (on a given path)

- **Collider:** Arrows coming in, but no out



- **Mediator:** 1 arrow in, 1 arrow out

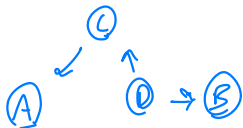


- **Confounder:** Arrows going out, but not in



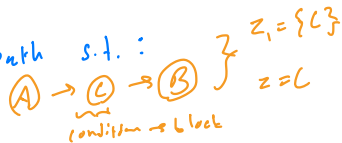
Types of Paths (from A to B...)

- **Backdoor Path:** $A \leftarrow \dots \rightarrow B$

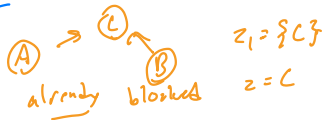


- **Blocked Path:**

Z_1 blocks path if $\exists z \in Z_1$ on path s.t. :
① z is a non-collider that is conditioned on
OR
② z is a collider that is not conditioned on, and all of its descendants not in Z_1



- **Open Path:** Any path that is not blocked



Blocking and Conditioning

Mediators, confounders don't block (by default)

Colliders do " " " "

Conditioning reverses blocking rules, may be necessary for ID

↕
including in regression



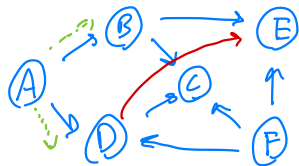
- **Definition:** A set of variables " d -separates" 2 other variables if it $\vee$ blocks all paths between them conditionally

- **Implication:** If A is d -separated from B by set Z_1 ,
 $A \perp B \parallel Z_1$

Identification Criteria (for causal effect of A on B)

- Back door Criterion: Can we block all ^{backdoor} paths by conditioning on some set of nodes? If yes, ID causal effect of $A \rightarrow B$ (conditional)

- Front door Criterion: Set of nodes Z , that satisfies:
 - ① Z intercepts all dir. paths from $A \rightarrow B$
 - ② No unblocked backdoor paths from $A \rightarrow Z$
 - ③ All backdoor paths from $Z \rightarrow B$ blocked by AIf all 3 hold, ID causal effect $A \rightarrow B$



What if A unobserved?

Yes, still have ID

$\{B, F\}$

Cannot condition on unobserved

Paths with unobserved nodes
still need to be blocked!

Backdoor Paths:

$D \leftarrow A \rightarrow B \rightarrow E$

$D \leftarrow F \rightarrow E$ collider

$D \leftarrow F \rightarrow C \leftarrow B \rightarrow E$

$D \leftarrow A \rightarrow B \rightarrow C \leftarrow F \rightarrow E$

Block by:

$\{A\}, \{B\}, \{A, B\}$

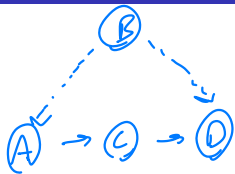
$\{F\}$

$\emptyset$

$\emptyset$

Options:

$\{A, F\}, \{B, F\}$, $\{A, B, F\} \rightarrow$ condition on
prefer fewer



① Intercept all directed paths

$$A \rightarrow C \rightarrow D \quad \{C\}$$

② No backdoor path $A \leftarrow \dots \rightarrow Z_1$

$$\text{No backdoor } A \leftarrow \dots \rightarrow C$$

③ All backdoor $Z_1 \leftarrow \dots \rightarrow D$ blocked by A

$$C \leftarrow \dots \rightarrow D$$

$$C \leftarrow A \leftarrow B \rightarrow D \quad (C \text{ blocked by } A) \checkmark$$

All 3 criteria satisfied by $Z_1 = \{C\}$

1. Announcements

2. Directed Acyclic Graphs (DAGs)

3. Practice Problems

4. eCONometrics? Evolution of Views

1. Announcements

2. Directed Acyclic Graphs (DAGs)

3. Practice Problems

4. eCONometrics? Evolution of Views

Econ 272/Mgtecon 607 - Section 5

Amar Venugopal

Stanford University

Spring 2025

1. Discussion Questions
2. Unconfoundedness
3. Semiparametric Efficiency Bounds
4. Practice Problems

1. Discussion Questions

2. Unconfoundedness

3. Semiparametric Efficiency Bounds

4. Practice Problems

1. Discussion Questions

2. Unconfoundedness

3. Semiparametric Efficiency Bounds

4. 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, we are “as-if” random.

- **Motivation:** With continuous covariates $P(\exists i, j : x_i = x_j) = 0$
- **Propensity Score:** $e(x) = P(W_i = 1 \mid X_i = x)$
→ "sufficient statistic" for selection on observables
- **(Alternative) Unconfoundedness:** $W_i \perp (Y_i(0), Y_i(1)) \mid e_i$
→ match on propensity scores, rather than raw covariates

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

2 assessment strategies:

- 1 Estimate effect on pseudo outcome
- 2 Estimate effects of pseudo treatments

Strategy #1: Estimating Effect on Pseudo Outcomes

Big Idea: Treat some covariates as outcomes, test independence.
$$X_i = (X_i^p, X_i^r) \left. \vphantom{\begin{matrix} X_i^p \\ X_i^r \end{matrix}} \right\} \text{ If } X_i^p \approx Y_i(\omega), \text{ then under } \textcircled{Q}: \\ \downarrow \text{pseudo outcome} \\ X_i^p \perp W_i \mid X_i^r \rightarrow \text{testable!}$$

Ex. Estimate avg. diff. in X_i^p by treatment status, after adjusting for differences in X_i^r as you choose.

One option: $X_i^p = \text{lagged outcome (if available)}$

Strategy #2: Estimating Effects of Pseudo Treatments

Big Idea: If we have multiple control groups, we can test if their observed outcomes differ.

$$\left. \begin{array}{l} G_i \in \{c, c_1, c_2\} \\ W_i = \begin{cases} 0, & G_i \in \{c_1, c_2\} \\ 1, & G_i = c \end{cases} \end{array} \right\} \begin{array}{l} \text{stronger version of (2):} \\ (Y_i(0), Y_i(1)) \perp G_i \mid X_i \end{array}$$

Testable implication:

$$Y_i(0) \perp G_i \mid X_i, G_i \in \{c_1, c_2\} \iff Y_i^{obs} \perp G_i \mid X_i, G_i \in \{c_1, c_2\}$$

Ex. As before, test whether Y_i differs on avg. b/w c_1, c_2 after adjusting for X_i

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:
$$\begin{aligned} \mathbb{E}[\tau(X_i)] &= \mathbb{E}[\mathbb{E}[Y_i(1) - Y_i(0) | X_i = x]] && \text{(definition of } \tau(x)) \\ &= \mathbb{E}[\mathbb{E}[Y_i(1) | X_i = x] - \mathbb{E}[Y_i(0) | X_i = x]] && \text{(linearity of expectation)} \\ &= \mathbb{E}[\mathbb{E}[Y_i(1) | U_i = 1, X_i = x] - \mathbb{E}[Y_i(0) | U_i = 0, X_i = x]] && \text{(unconfoundedness)} \\ &= \mathbb{E}[\underbrace{\mathbb{E}[Y_i | U_i = 1, X_i = x]}_{\text{?}} - \underbrace{\mathbb{E}[Y_i | U_i = 0, X_i = x]}_{\text{?}}] && \text{(definition of PO)} \\ &\stackrel{?}{=} \tau && \begin{array}{l} \rightarrow \text{need something more} \\ \searrow \text{may not be separately estimable!} \end{array} \end{aligned}$$

Definition: $P(w_i=1 | x_i) = e_i \in (0,1)$

For any x , we can find both treated and control

This assumption is *testable*. Some strategies:

- 1 Normalized differences
- 2 Plotting
- 3 Log odds ratio

Strategy #1: Normalized Differences

Big Idea: Check if covariate distributions vary across treatment status

$$nd = \frac{\bar{x}_t - \bar{x}_c}{\sqrt{\frac{s_t^2 + s_c^2}{2}}}$$

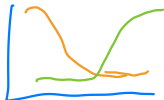
→ better than t statistic, which will increase in sample size

Rule of Thumb: If $nd > 0.25$, problem!

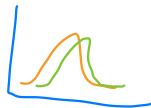
Strategy #2: Plotting

Big Idea: Check visually if distributions "look" overlapping

"Rule of Thumb":



Bad!



Good!

Strategy #3: Log Odds Ratio

Big Idea: Propensity scores should be similar between the 2 groups.

$$LOR = \log\left(\frac{e(x)}{1-e(x)}\right) \rightarrow \text{calculate means for each group and compare}$$

Rule of Thumb: If means > 1 s.d. apart, problem!
std. dev.

What if we have only limited overlap?

Challenge: We cannot get good estimates of the separate conditional outcomes.
Need a "better" sample.

2 strategies:

1 Matching

2 Subsampling

Strategy #1: (Propensity Score-) Matched Samples

Big Idea: Construct a new sample of similar units in each treatment group.

Better to match on propensity scores than raw covariates.

Algorithm:

- 1 Estimate $e(x_i)$ $\forall i$
- 2 Order treated units by decreasing propensity
- 3 For each, find control unit with closest propensity (without replacement)
- 4 Now have balanced sample of $2 \cdot N_T$ units
- 5 Proceed as usual, with this as your new sample

Strategy #2: Subsamples

Big Idea: Only consider units with covariate values where we have good overlap

$$\tau_c(A) = \frac{1}{|A|} \sum_{i: X_i \in A} \tau(X_i) \rightarrow \text{subpopulation ATE}$$

Choosing Subsample: Choose A to minimize variance of $\tau_c(A)$.

Generally, this means:

$$A = \{X \in X : \alpha < e(X) < 1 - \alpha\} \rightarrow \text{throw out extrema of propensity score}$$

$$\alpha = 2 \cdot E \left[\frac{1}{e(X)(1-e(X))} \mid \frac{1}{e(X)(1-e(X))} < \alpha \right] \rightarrow \alpha \approx 0.1$$

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

1. Discussion Questions

2. Unconfoundedness

3. Semiparametric Efficiency Bounds

4. Practice Problems

Types of Models

- **Parametric:** Well-defined models, finite dimensional space
e.g. OLS, $y \sim N(\alpha + \beta x, \sigma^2)$
- **Semiparametric:** Parameter of interest is finite-dimensional, with another that is infinite-dimensional
e.g. Cox Proportional Hazards, $y \sim f(y | \beta, \lambda_0)$
PI → function (∞-dim)
"nuisance parameter"
- **Nonparametric:** All parameters can span an infinite space
e.g. decision trees

Parametric Efficiency

Given: $f_X(x|\theta)$ pdf
 $\hat{\theta}(x)$ unbiased for θ

Fisher Information: $I(\theta) = -E \left[\frac{\partial^2 \log(f)}{\partial \theta \partial \theta} (x|\theta) \right]$

Cramer-Rao Lower Bound: $V(\hat{\theta}(x)) \geq I(\theta_0)^{-1}$

MLE generally
unbiased + efficient

$\hat{\theta}$ is asymptotically efficient if:

$$\sqrt{n}(\hat{\theta} - \theta_0) \xrightarrow{d} N(0, I(\theta_0)^{-1})$$

Parametric Submodels & Fisher Information

Semiparametric model: $f(x|\theta, h(\gamma)) = f(x|\theta, \gamma) \rightarrow$ parametric sub-model
 ∞ -dim $h(\gamma)$ known up to finite parameter γ

CRLB tells us: $\underbrace{ASV(\hat{\theta})}_{\text{asymptotic variance}} \geq \underbrace{I^{\theta\theta'}}_{\substack{\text{inverse} \\ \text{Fisher info.} \\ \text{for } \theta}} = (I_{\theta\theta} - I_{\theta\gamma} I_{\gamma\gamma}^{-1} I_{\gamma\theta})^{-1}$

Goal: rewrite $I^{\theta\theta'}$ in simpler terms

Score Function = 1st deriv of log. max. likelihood

Let $\beta = (\theta, \gamma)$

$$S_{\beta}(x) = \begin{pmatrix} s_{\theta}(x) \\ s_{\gamma}(x) \end{pmatrix} = \begin{pmatrix} \frac{\partial \log f}{\partial \theta}(x; \theta, \gamma) \\ \frac{\partial \log f}{\partial \gamma}(x; \theta, \gamma) \end{pmatrix}$$

Can then write Fisher info:

$$I(\theta, \gamma) = \begin{pmatrix} I_{\theta\theta} & I_{\theta\gamma} \\ I_{\gamma\theta} & I_{\gamma\gamma} \end{pmatrix} = E \left[\begin{pmatrix} s_{\theta}(x) s_{\theta}(x)' & s_{\theta}(x) s_{\gamma}(x)' \\ s_{\gamma}(x) s_{\theta}(x)' & s_{\gamma}(x) s_{\gamma}(x)' \end{pmatrix} \right]$$

CLB then gives:

$$(\mathbb{E}[\psi(x)\psi(x)'])^{-1} = \mathbb{E}[\psi(x)\psi(x)']^{-1} - \mathbb{E}[\psi(x)\psi(x)']^{-1} \mathbb{E}[\psi(x)\psi(x)'] \mathbb{E}[\psi(x)\psi(x)']^{-1}$$

$$= \mathbb{E}[\psi(x)\psi(x)']^{-1}$$

$$\eta(x) = S_0(x) - \mathbb{E}[S_0(x)S_0'(x)]^{-1} \mathbb{E}[S_0(x)S_0'(x)] S_0(x)$$

↳ residual from projection of S_0 onto S_r = "efficient score"

In practice, we do not know S_r , so we look over all possible values, called the tangent set → generally tricky to find

(Efficient) Influence Function

Efficient Influence Function = influence function with smallest variance

EIP orthogonal to nuisance tangent space, so can be obtained by rescaling efficient score:

$$\phi(x) = E[\psi(x)\psi(x)']^{-1} \psi(x)$$

Intuition: tells you how much a single data point "influences" your point estimate

If $\hat{\theta}$ is efficient, we say it is asymptotically linear:

$$\sqrt{n}(\hat{\theta} - \theta) = \frac{1}{n} \sum_{i=1}^n \underbrace{\phi(x_i)}_{\phi \text{ gives "impact" of } x_i \text{ on } \hat{\theta}} + o_p(n^{-1/2}) \xrightarrow{d} N(0, I^{\theta\theta})$$

Semiparametric Variance Bound

Bound given by variance of influence function:

$$E[\psi(x_i) \psi(x_i)'] = \left(E[\psi(x_i) \psi(x_i)'] \right)^{-1} \quad (\text{SVB})$$

$\forall \hat{\theta}$ semiparametric, $\text{Var}(\hat{\theta}) \geq \text{SVB}$

if $\hat{\theta}$ is "semiparametrically efficient"

(many such estimators, coming soon)

Riesz Representer

Simpler way to calculate variance bound.

Differentiable Parameter: θ differentiable if $\exists \alpha(\cdot)$ s.t.

$$\frac{\partial \theta}{\partial \gamma}(\gamma) = E[\alpha(z_i)' S(z_i; \gamma)] \quad \text{Riesz Representer}$$

If differentiability holds, we can characterize the Cramer-Rao bound via the Riesz representer:

$$E[\phi(z_i; \gamma) \phi(z_i; \gamma)'] = E[\alpha(z_i)' S(z_i; \gamma)] \underbrace{(E[S(z_i; \gamma) S(z_i; \gamma)'])^{-1}}_{\text{SVB}} E[S(z_i; \gamma) \alpha(z_i)]$$

Where IF given by projection of RR on full score function:

$$\phi(z; \gamma) = S(z; \gamma)' (E[S(z; \gamma) S(z; \gamma)'])^{-1} E[S(z; \gamma)' \alpha(z)]$$

Can estimate SVB by just estimating each of the unknown components

Econ 272/Mgtecon 607 - Section 6

Amar Venugopal

Stanford University

Spring 2025

1. Announcements
2. Discussion Questions
3. Doubly Robust Methods
4. Linear IV
5. Practice Problems

1. Announcements

2. Discussion Questions

3. Doubly Robust Methods

4. Linear IV

5. Practice Problems

- Makeup Lectures
 - May 14, 21 → 8-9:50am GSB C105
 - June 2, 4 no lecture (maybe section?) → May 22, 4:30-6:20pm STLC 114
- EOY Pizza Party: Date/Time TBD (June 5-12?)

Hints for Problem Set 5 (Due May 11, 11pm)

- Imbens (2004)¹ may be helpful in finding ATT estimators
- When bootstrapping, you must re-estimate relevant propensity scores for each bootstrap sample!
- If you have any covariates that are constant/perfectly colinear within a given stratum, you can just drop them
- For theoretical calculations: think about what we can add/subtract to the propensity score to express it as the difference of the covariate densities. Would doing so change the variance estimate?

¹NONPARAMETRIC ESTIMATION OF AVERAGE TREATMENT EFFECTS UNDER EXOGENEITY: A REVIEW

1. Announcements

2. Discussion Questions

3. Doubly Robust Methods

4. Linear IV

5. Practice Problems

1. Announcements

2. Discussion Questions

3. Doubly Robust Methods

4. Linear IV

5. Practice Problems

Quick Recap: Semiparametric Efficiency

- **Goal:**

Given a model with finite θ and ∞ -dim nuisance param, $h(\cdot)$, find the best variance possible for θ

- **Big Idea:** Considering all values for $h(\cdot)$, find best var. for θ

- **Asymptotic Linearity:** Can express asymptotic dist. of θ as a linear combo of influence function

$$\hat{\theta} = \tau + \frac{1}{n} \sum_{i=1}^n \phi(Y_i, W_i, X_i) + o_p(n^{-1/2})$$

- **ATE Bound:**

$$V^{SPVB} = E \left[\frac{\sigma_1^2(X_i)}{e(X_i)} + \frac{\sigma_0^2(X_i)}{1-e(X_i)} + (\tau(X_i) - \tau)^2 \right]$$

Given: outcome models $\mu_1(x), \mu_0(x)$, propensity score $\hat{e}(x)$

1 Influence Function:

↓
AIPW

$$\tilde{\tau}^{IF} = \frac{1}{N} \sum_{i=1}^N \left[\mu_1(x_i) - \mu_0(x_i) + w_i \frac{y_i - \mu_1(x_i)}{e(x_i)} - (1-w_i) \frac{y_i - \mu_0(x_i)}{1-e(x_i)} \right]$$

2 Regression:

$$\tilde{\tau}^{reg} = \frac{1}{N} \sum_{i=1}^N \mu_1(x_i) - \mu_0(x_i)$$

Infeasible because
all rely on
 μ_1, μ_0, e

3 IPW:

$$\tilde{\tau}^{IPW} = \frac{1}{N} \sum_{i=1}^N \left[w_i \frac{y_i}{e(x_i)} - (1-w_i) \frac{y_i}{1-e(x_i)} \right]$$

1 Influence Function: Asymptotic linearity $\rightarrow$ efficient

$$\sqrt{N}(\tilde{\tau}^{IF} - \tau) \xrightarrow{d} N(0, V^{SPVB})$$

2 Regression: Lower than SPVB

(assume knowledge of $\mu \rightarrow$ more info $\rightarrow$ lower variance)

3 IPW: Higher than SPVB

$$\tilde{\tau}^{IPW} = \frac{1}{N} \sum_{i=1}^N \left[w_i \frac{y_i}{e(x_i)} - (1-w_i) \frac{y_i}{1-e(x_i)} \right]$$

Comparison: $V^{Reg} \leq V^{IF} = V^{SPVB} \leq V^{IPW}$

Big Idea: Replace μ_1, μ_0, e with estimates $\hat{\mu}_1, \hat{\mu}_0, \hat{e}$

Variance Comparison: $\mathbb{V}^{Reg} = \mathbb{V}^{IF} = \mathbb{V}^{SPVB} = \mathbb{V}^{IPW}$

↓
generally prefer
→ Doubly Robust

Influence Function/AIPW - Double Robustness

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

Double Robustness: $\hat{\tau}^{AIPW} \xrightarrow{P} \tau$ if either μ or e correctly specified
 consistently unbiased estimator

Solves WLS: $\min_{\alpha, \beta, \tau} \sum_{i=1}^N (Y_i - \alpha - \tau W_i - X_i \beta)^2 \underbrace{w(X_i)}_{= \frac{W_i}{e(X_i)} + \frac{1-W_i}{1-e(X_i)}}$

Influence Function/AIPW - Correctly Specified Outcomes

$$\hat{\tau}^{AIPW} = \underbrace{\left[\frac{1}{N} \sum_{i=1}^N (\hat{\mu}_1(X_i) - \hat{\mu}_0(X_i)) \right]}_{E[\mu_1(x) - \mu_0(x)]} + \underbrace{\left[\frac{1}{N} \sum_{i=1}^N w_i \frac{Y_i - \hat{\mu}_1(X_i)}{\hat{e}(X_i)} - (1 - w_i) \frac{Y_i - \hat{\mu}_0(X_i)}{1 - \hat{e}(X_i)} \right]}_{E[\cdot] = 0}$$

$$= E[E[Y_i(1) - Y_i(0) | X=x]]$$

$$= E[\tau(x)]$$

$$= \tau$$

$$\mu_1(x) = E[Y | w=1, X=x]$$

$\hat{\mu}_1: Y \sim X, w$

Alternatively, WLS:

$$\mu_w(x) = \alpha + \tau w + x\beta$$

$$E[Y - \mu_w(x) | X, w] = 0$$

$$WLS \text{ obj} = 0$$

Influence Function/AIPW - Correctly Specified Propensity

Key: in weighted sample $W_i \perp X_i$

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

Diagram illustrating the decomposition of the AIPW estimator into expectation terms:

Under the first bracket, the terms are grouped as $E[\cdot] = \mu_1$ and $E[\cdot] = \mu_0$.

Under the second bracket, the term is grouped as $E[v_i] - E[v_i] = 0$.

The overall expression is simplified to:

$$E[Y|W=1] - E[Y|W=0]$$

$= \tau$

Correct specification:

$$\hat{e} \rightarrow e$$

$$\hat{\mu} \rightarrow \mu$$

Cross-Fitting/Sample Splitting

Idea: Deriving formal properties of estimate μ, c on separate sample from $\mathcal{Z}$
(ensuring independence)

Procedure:

1 Split sample into K folds, $\mathcal{I}_i = \{1, \dots, K\}$

2 Estimate $\hat{\mu}_w^j(x), \hat{e}^j(x)$ on sample $\mathcal{I}_i \neq j$

3 Estimate $\hat{c}^j(\hat{\mu}_w^j, \hat{e}^j)$

4 Average $\frac{1}{K} \sum_{j=1}^K \hat{c}^j$ to get efficient estimator

Goal: Find optimal assignment of units to treatment

■ **Policy Rule:** $\pi; X \rightarrow \{0, 1\}$

■ **Desirable Properties:**

☆ ■ $\pi(x) = \mathbb{1} \{ \mu_1(x) - \mu_0(x) \geq 0 \}$

■ Interpretable, well-behaved assignment rules
eg. $\pi(x) = \mathbb{1} \{ x'\beta \geq 0 \}$

Goal: Maximize gain over treating nobody : $V(\pi) = E[\pi(x_i)(Y_i(1) - Y_i(0))]$
 $= E[\pi(x_i) \tau(x_i)]$

ATE Estimator: $\hat{\tau}^{RF} = \frac{1}{N} \sum_{i=1}^N \hat{\tau}_i \rightarrow$ exactly summation term from before

Effect Estimator: $\hat{V}(\pi) = \frac{1}{N} \sum_{i=1}^N \pi(x_i) \hat{\tau}_i$

Policy Estimator: $\hat{\pi} = \arg \max_{\pi \in \Pi} \underbrace{\frac{1}{N} \sum_{i=1}^N \pi(x_i) \hat{\tau}_i}_{= \hat{V}(\pi)} \rightarrow$ typically estimated via cross-fitting

1. Announcements

2. Discussion Questions

3. Doubly Robust Methods

4. Linear IV

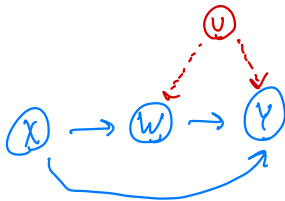
5. Practice Problems

What if unconfoundedness breaks down?

In observational settings, **unconfoundedness** let's us “pretend” that we are dealing with a randomized experiment.

When might unconfoundedness be violated?

We have some unobserved variable that impacts both outcome and treatment

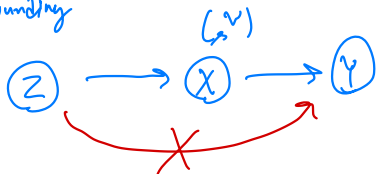


Ex. $\underbrace{wage}_w = \alpha + \beta \underbrace{educ}_e + \varepsilon$
 $U = ability$

- **Endogeneity:** Variable correlated with error term
contain effects of all unobserved components
- **Omitted Variable Bias:** When we cannot control for confounder, estimated coeffc. will be biased
- **Reverse Causality:** Values regressor takes on may be determined by function for outcome
e.g. regressor = price } price set by demand function, so price endogenous
outcome = demand } $d = D(p, U)$

Instrumental Variables

Goal: Find exogenous "instrument" that allows us to remove effect of confounding



To obtain valid causal estimates, IV must satisfy:

- **Exclusion:** Z affects Y only through X $\rightarrow$ not testable (argued)
- **Relevance:** Z has causal effect on X $\rightarrow$ testable ("strong 1st stage")

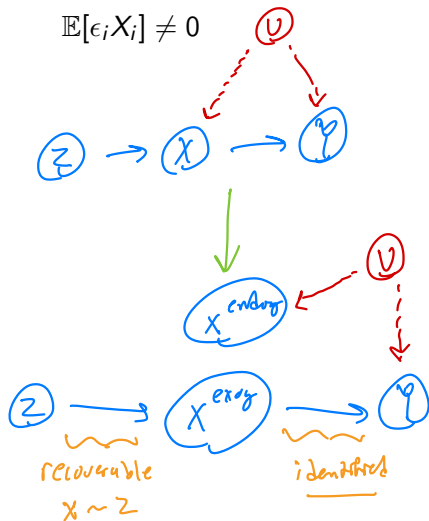
$$Y_i = \alpha + \beta X_i + \epsilon_i$$

$$\mathbb{E}[\epsilon_i X_i] \neq 0$$

Goal: Recover causal effect of X on Y
removing bias in $\hat{\beta}_{naive}$

Estimand:

$$\beta = \frac{\text{Cov}(Y_i, Z_i)}{\text{Cov}(X_i, Z_i)} = \frac{\text{Cov}(Y_i, Z_i) / \text{Var}(Z_i)}{\text{Cov}(X_i, Z_i) / \text{Var}(Z_i)}$$



Estimation Strategy #1 - Two Stage Least Squares (TSLS)

Big Idea:

- ① Regress X on Z , get $\hat{X}$
- ② Regress Y on $\hat{X}$

Estimator:

- ① $X_i = \alpha_0 + \alpha_1 Z_i + u_i$ → contains confounding
- ② $Y_i = \beta_0 + \beta^{TSLS} \hat{X}_i + \epsilon_i$ → $\hat{X} = \hat{\alpha}_0 + \hat{\alpha}_1 Z \perp \epsilon_i$

$$\hat{\beta}^{TSLS} = \frac{\text{Cov}(Y_i, Z_i)}{\text{Cov}(X_i, Z_i)} = (Z'X)^{-1}(Z'Y)$$

Inference: Under homoskedasticity: $\hat{\beta}^{TSLS} - \beta \approx N(0, \sigma^2[\hat{X}\hat{X}']^{-1})$
↳ only from 2nd stage
larger than OLS

Hausman Test

Goal: If X_i exogenous, TSLS and OLS consistent.
But, only OLS efficient $\rightarrow$ only want TSLS if really worth it.

Note: a good, valid IV required

$$H_0: \varepsilon_i \perp X_i \rightarrow \beta^{TSLS} = \beta^{OLS} \rightarrow \hat{\beta}^{TSLS} - \hat{\beta}^{OLS} \approx N(0, \hat{V}^{TSLS} - \hat{V}^{OLS})$$

Test Statistic: $t = \frac{\hat{\beta}^{TSLS} - \hat{\beta}^{OLS}}{\sqrt{\hat{V}^{TSLS} - \hat{V}^{OLS}}}$ $\rightarrow$ if large, reject null and conclude X truly endogenous

Econ 272/Mgtecon 607 - Section 7

Amar Venugopal

Stanford University

Spring 2025

1. Announcements & Discussion Questions
2. Local Average Treatment Effects (LATE)
3. Weak Instruments & Many Instruments
4. Practice Problems

1. Announcements & Discussion Questions

2. Local Average Treatment Effects (LATE)

3. Weak Instruments & Many Instruments

4. Practice Problems

- Next Week Lectures

- May 21 → 8-9:50am GSB C105

- June 2, 4 no lecture (maybe section?) → May 22, 4:30-6:20pm STLC 114

1. Announcements & Discussion Questions

2. Local Average Treatment Effects (LATE)

3. Weak Instruments & Many Instruments

4. Practice Problems

Motivation: Effect Heterogeneity

Goal: Allow for: heterogeneity in TE across units
more flexible functional forms
→ endogeneity $\neq$ correlation with residual

Challenge: Cannot ID population ATE
Can only ID ATE for sub-pop'n → compliers

- (Binary) Z : *Instrument*

- (Binary, endogenous) W : $w_i = w_i(z_i) = \begin{cases} w_i(1) & \text{if } z_i = 1 \\ w_i(0) & \text{if } z_i = 0 \end{cases}$

- Y : $y_i = y_i(w_i) = y_i(w_i(z_i)) = y_i(z_i, w_i)$

Key Assumption: Independence of Instrument

$$Z_i \perp\!\!\!\perp (Y_i(0), Y_i(1), W_i(0), W_i(1))$$

→ IV does not directly affect outcome
(nonparametric assumption)

Conditions:

■ Randomization of Z:

$$Z_i \perp (W_i(0), W_i(1), Y_i(0,0), Y_i(1,0), Y_i(0,1), Y_i(1,1))$$

■ Exclusion: → not testable, must be argued

$$\forall z', z, w : Y_i(z, w) = Y_i(z', w)$$

True Compliance Types

		$W_i(0)$ $Z_i=0$	
		0	1
$W_i(1)$ $Z_i=1$	0	never-taker	defiers
	1	complier	always-takers

Observed Compliance Types

		Z_i	
		0	1
W_i	0	never-taker/complier	never-taker/defier
	1	always-taker/defier	always-taker/complier

Monotonicity

Goal: Disambiguate type of unit from (Z_i, W_i)

Assumption: Monotonicity : $W_i(1) \geq W_i(0) \iff$ no defiers
 $\hookrightarrow = \iff$ always-/never-takers
 $\quad \quad \quad > \iff$ compliers

↓
sometimes testable

more plausible
when IV incentive

		Z_i	
		0	1
W_i	0	never-taker/complier	never-taker
	1	always-taker	always-taker/complier

Conditional Means

Let proportions of compliers, never-takers, and always-takers be given by π_c, π_n, π_a .

- $E[W_i | Z_i = 0] = \pi_a$
 - $E[W_i | Z_i = 1] = \pi_a + \pi_c$
 - $E[Y_i | W_i = 0, Z_i = 0] = \frac{\pi_c}{\pi_c + \pi_n} E[Y_i(0) | c] + \frac{\pi_n}{\pi_c + \pi_n} E[Y_i(0) | n]$
 - $E[Y_i | W_i = 1, Z_i = 0] = E[Y_i(1) | a]$
 - $E[Y_i | W_i = 0, Z_i = 1] = E[Y_i(0) | n]$
 - $E[Y_i | W_i = 1, Z_i = 1] = \frac{\pi_c}{\pi_a + \pi_c} E[Y_i(1) | c] + \frac{\pi_a}{\pi_a + \pi_c} E[Y_i(1) | a]$
- ED $\pi_a, \pi_c, \pi_n = 1 - \pi_a - \pi_c$

Why only compliers? Only group with members in both treatment groups

Estimand (LATE): $LATE = E[Y_i(1) - Y_i(0) | \text{complier}]$

Estimation Strategies:

- Use previous conditional means
- ILS

- Y ~ Z Slope Coefficient:** $E[Y_i | Z_i = 1] - E[Y_i | Z_i = 0]$
 $(E[Y_i(1) | C] \pi_C + E[Y_i(0) | N] \pi_N + E[Y_i(1) | A] \pi_A) = E[Y_i(1) - Y_i(0) | C] \pi_C$
 $= -(E[Y_i(0) | C] \pi_C + E[Y_i(0) | N] \pi_N + E[Y_i(1) | A] \pi_A)$
- W ~ Z Slope Coefficient:** $E[W_i | Z_i = 1] - E[W_i | Z_i = 0] = \pi_C$
- Ratio:** $\tau^{IV} = \frac{E[Y_i | Z_i = 1] - E[Y_i | Z_i = 0]}{E[W_i | Z_i = 1] - E[W_i | Z_i = 0]} = E[Y_i(1) - Y_i(0) | C]$

ATE Bounds (Binary Outcomes)

Goal: Partial ID of full pop'n ATE

$$ATE = E[Y_i(1) - Y_i(0)] = E[\overset{\checkmark}{Y_i(1)} - \overset{\checkmark}{Y_i(0)} | a] \pi_a + E[\overset{\checkmark}{Y_i(1)} - \overset{\checkmark}{Y_i(0)} | n] \pi_n + E[\overset{\checkmark}{Y_i(1)} - \overset{\checkmark}{Y_i(0)} | c] \pi_c$$

→ cannot estimate $E[Y_i(0) | a]$, $E[Y_i(1) | n]$

LB: $E[Y_i(0) | a] = 1$, $E[Y_i(1) | n] = 0$

UB: $E[Y_i(0) | a] = 0$, $E[Y_i(1) | n] = 1$

$$W_i \perp (Y_i(0), Y_i(1)) \mid Z_i$$

What if we (mistakenly) included Z as a control, rather than an instrument?

DiM estimator, after adjusting for Z_i , becomes:

$$E[Y_i \mid V_i=1, Z_i=1] - E[Y_i \mid W_i=0, Z_i=1] = E[Y_i(1) \mid V_i=1, Z_i=1] - E[Y_i(0) \mid V_i=0, Z_i=1]$$

$$= \frac{\pi_u}{\pi_u + \pi_c} E[Y_i(1) \mid u] + \frac{\pi_c}{\pi_u + \pi_c} E[Y_i(1) \mid c] - E[Y_i(0) \mid u]$$

not causally interpretable

1. Announcements & Discussion Questions
2. Local Average Treatment Effects (LATE)
3. Weak Instruments & Many Instruments
4. Practice Problems

Challenge: Standard asymptotics generally rely on nonzero correlation between
IV and endog strong 1st stage

If $\text{correl} \approx 0$, approx inaccurate even in large samples

Goal: Adjust s.e. to allow for this violation

Conventional Asymptotics (Fixed # Instruments, Fixed Parameters)

- **LIML vs. TSLS:** $\sqrt{N}(\hat{\beta}_{\text{LIML}} - \hat{\beta}_{\text{TSLS}}) \rightarrow 0$ *s.e. should be similar*
- **Just-Identified:** # IV = # endog
CI not uniformly valid, but still "decent" coverage in "reasonable" regions

- **Over-Identified:** # IV > # endog
CI potentially very misleading (especially TSLS)
- **Standard Estimator:**
$$\left. \begin{aligned} X_i &= \pi_0 + \pi_1 Z_i + \eta_i \\ Y_i &= \alpha_0 + \alpha_1 Z_i + \nu_i \end{aligned} \right\} \beta_1 = \frac{\alpha_1}{\pi_1} \rightarrow \hat{\beta}^{\text{IV}} = \frac{\frac{1}{N} \sum (Y_i - \bar{Y})(Z_i - \bar{Z})}{\frac{1}{N} \sum (X_i - \bar{X})(Z_i - \bar{Z})}$$
- **Concentration:**
$$\lambda = \pi_1^2 \sum \frac{(Z_i - \bar{Z})^2}{\sigma_\eta^2} \rightarrow \text{strength of instrument}$$

weak: $\lambda \rightarrow 0$

Big Idea: Let 1st stage param. depend on sample size: $\pi_{1n} = \frac{c}{\sqrt{n}}$

Then, $\lambda \rightarrow c^2 \cdot V(z_i)$

Demonstrates that for any sample size, coverage can vary a lot from desired level.

Suggested methods for uniformly valid CIs.

Goal: *Provide uniformly valid CIs*

- **Anderson-Rubin CI:**

- **Kleibergen Test:**

- **Moreira's Similar Test:**

2 problems:

- 1 Bias in standard estimators
 - 2 Misleadingly small s.e.
- } not fixable via bootstrap

Main take-away: LIML works well, but s.e. must be adjusted

2 strategies to adjust standard errors:

- 1 Bekker
- 2 REQML

→ generally recommended

Big Idea:

TSLS is not consistent in large sample, but LIML is
considers new asymptotics to derive uniformly valid CEs

$$V^{\text{Bekker}} = (V^{\text{LIML}}) (\text{increasing \# of instruments})$$

Random Effects Quasi Maximum Likelihood (REQML)

Big Idea: Model 1st stage coeffs. π_{ik} as iid draws from $N(\mu_\pi, \sigma_\pi^2)$

Assume joint normality, normalized Z_i , likelihood function:

$$\mathcal{L}(\beta_0, \beta_1, \pi_0, \mu_\pi, \sigma_\pi^2, \Omega)$$

6-parameter problem

TSLS: posterior mode given flat priors ($\sigma_\pi^2 = \infty$)

LDML: $\mu_\pi = 0$, Ω known

Correction REQML corrections generally larger than Bekker, because ML estimator for σ_π^2 will take into account noise in $\hat{\pi}$

Econ 272/Mgtecon 607 - Section 8

Amar Venugopal

Stanford University

Spring 2025

1. Announcements & Discussion Questions
2. Difference-in-Differences
3. Synthetic Control
4. Alternative/New Panel Methods
5. Practice Problems

1. Announcements & Discussion Questions

2. Difference-in-Differences

3. Synthetic Control

4. Alternative/New Panel Methods

5. Practice Problems

- Last section: Wednesday June 4, 11:30am-1:20pm, GSB Z301 (in place of regular lecture)
 - Focus on exam prep, review + practice problems (?)
- Last Problem Set: HW7 due June 4, 11pm
- Final Exam: June 11, 8:30-11:30am, 200-002
- Pizza Party! June 9 (exact time TBD) - RSVP by June 1

1. Announcements & Discussion Questions

2. Difference-in-Differences

3. Synthetic Control

4. Alternative/New Panel Methods

5. Practice Problems

Key feature: Observe both units ($i=1, \dots, N$) and time periods ($t=1, \dots, T$)

$$Y_{it}$$

3 main types of panel data:

1 Proper Panel: Observe same units in each period

$$Y_{it}, i=1, \dots, N, t=1, \dots, T$$

2 Repeated Cross-Sections: Observe different individuals in each time period

$$Y_i \text{ for } i=1, \dots, N \text{ in period } T_i$$

$G_i = \text{group assignment (treatment/control)}$

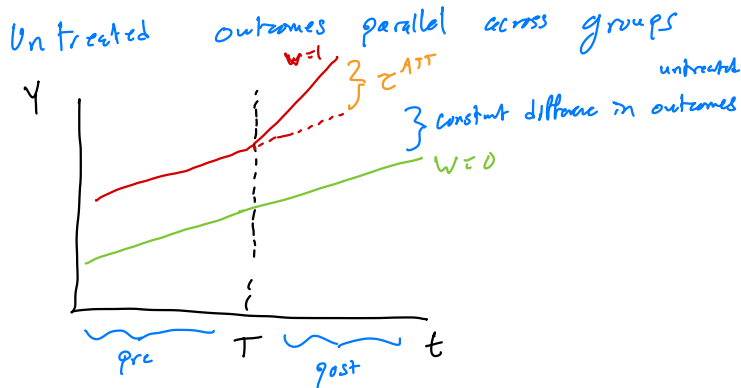
3 Repeated Cross-Sections with Multiple Groups:

$$\text{Same, } G_i = \{1, \dots, G\}$$

Key Assumption: Parallel Trends

(2 x 2)

Big Idea:



Testable?

Not directly

But, can assess by checking if parallel pre-treatment

Repeated Cross-Section, 2 Groups $\times$ 2 Periods - Setup

Observed data: (Y_i, G_i, T_i)
 $\downarrow$ $\hookrightarrow \in \{0,1\}$
 $\in \{0,1\}$

Treatment assignment: Group 1 is treated in period 1

$$W_i = G_i T_i, \quad Y_i = (1 - W_i) Y_i(0) + W_i Y_i(1)$$

Modeling Assumption: $Y_i(0) = \alpha + \beta T_i + \gamma G_i + \varepsilon_i$ ↗ can add covariates $+ \delta^T X_i$
time effect group effect

Typically assume constant τ : $Y_i = \alpha + \beta T_i + \gamma G_i + \tau W_i + \varepsilon_i$

If τ heterogeneous, identifies average effect

Repeated Cross-Section, 2 Groups $\times$ 2 Periods - Estimation

Estimand: $\tau^{DID} = \left(E[Y_i | G_i=1, T_i=1] - E[Y_i | G_i=1, T_i=0] \right) - \left(E[Y_i | G_i=0, T_i=1] - E[Y_i | G_i=0, T_i=0] \right)$

ATT treated group
gr evolution
control
group
evolution

Estimator: $\tau^{DID} = \bar{Y}_{11} - \bar{Y}_{10} - (\bar{Y}_{01} - \bar{Y}_{00})$

$\min_{\alpha, \beta, \gamma, \tau} \sum_{i=1}^N (Y_i - \alpha - \beta T_i - \gamma G_i - \tau W_i)^2$ } both equivalent

Variance:

$$V(\tau)_{\text{robust}} = \frac{\sigma_{11}^2}{N_{11}} + \frac{\sigma_{10}^2}{N_{10}} + \frac{\sigma_{01}^2}{N_{01}} + \frac{\sigma_{00}^2}{N_{00}}$$

Under homoskedasticity: $\sigma_{11}^2 = \sigma_{10}^2 = \sigma_{01}^2 = \sigma_{00}^2 \equiv \sigma^2$

Observed data: $(G_i, Y_{i0}, Y_{i1}, W_{i0}, W_{i1})$

Modeling Assumption: $Y_{it} = \alpha + \beta t + \gamma G_i + \tau W_{it} + \epsilon_{it}$

Estimator: Could do double difference
or

$$Y_{it} = \mu + \alpha_i + \beta_t + \tau W_{it} + \epsilon_{it}$$

↓ ↑
fixed effects TWFE

Unconfoundedness Assumption: Independence conditional on lagged outcomes

$$Y_{i,1} - Y_{i,0} = \alpha + \beta Y_{i,0} + \tau G_i + \eta_i$$

DiD Assumption: $Y_{i,1} - Y_{i,0} = \alpha + \tau G_i + \varepsilon_i$

How to choose? Unconfoundedness allows for flexible estimation of β
But DiD assumption $Y_{i,0}$ may be correlated with ε_i

Nice result: If ④ correct, but estimate DiD, overestimate τ
If DiD right, but estimate ④, underestimate τ } yield bounds

Changes-in-Changes: Continuous Data (2×2)

Motivation: DID sensitive to transformations of data, functional form

Modeling Assumption: $Y_i(0) = h(U_i, T_i)$, $h(\cdot)$ strictly increasing in U
 $U_i \perp T_i \mid G_i \rightarrow$ ^{unobserved} distribution of U_i vary across groups, but not time

$$U_i = G_i T_i$$

Estimand: $\tau^{AT} = \underbrace{E[Y_i(1) \mid G_i=1, T_i=1]}_{\text{observed}} - \underbrace{E[Y_i(0) \mid G_i=1, T_i=1]}_{\text{estimate distribution}}$

Estimation: $F_{Y(0), 11}(y) = F_{Y(0), 10} \left(F_{Y(0), 00}^{-1} \left(F_{Y(0), 01}(y) \right) \right)$
 $\downarrow \quad \quad \downarrow \quad \quad \downarrow$
observed

Modern Approaches: DiD With Heterogeneous Effects, Multiple Groups/Multiple Periods

Big Idea: In these settings, TWFE $\neq$ DiD if effects heterogeneous

3 proposed estimators:

- Chaisemartin & d'Haultfoeuille (2020): Weighted avg. of 1-period ahead DiD estimators

- never-treated may be fundamentally different
 - constant additive effects more plausible in short run

- Callaway & Sant'Anna (2021): Compare evolution of outcomes to never-treated group

1-period ahead
may miss long-run
effects

- Sun & Abraham (2021): Similar to CS, allowing for comparison to last-treated group if never-treated unavailable

All require addition/modification to parallel trends

1. Announcements & Discussion Questions

2. Difference-in-Differences

3. Synthetic Control

4. Alternative/New Panel Methods

5. Practice Problems

What kind of panel must we have? Proper panel

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

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

Estimand:

ATT

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

In case with 1 unit i treated in 1 period T , this becomes:

$$\tau^i = \underbrace{Y_{iT}(1)}_{\text{observed}} - \underbrace{Y_{iT}(0)}_{\text{unobserved}}$$

Vertical vs. Horizontal Regression

Goal: y_{iT} (o)

$$Y = \begin{pmatrix} y_{11} & y_{12} & \dots & y_{1T} \\ y_{21} & y_{22} & \dots & y_{2T} \\ \vdots & \vdots & \ddots & \vdots \\ y_{M1} & y_{M2} & \dots & y_{MT} \end{pmatrix}$$

Vertical Regression: $T-1$ obs., $N-1$ regressors
e.g. SC

Main idea: similar units behave similarly (stable relationship across time over units)

Horizontal Regression: $N-1$ obs., $T-1$ regressors
e.g. unconfoundedness

Main idea: history explains all (stable relationship across units over time)

Takeaway: If T large, N small (fat), go vertical
skinny, go horizontal

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

$$y_{it}(0) = \underbrace{\mu_i}_{\text{loadings}}^T \underbrace{\lambda_t}_{\text{factor values}} + \varepsilon_{it}$$

Synthetic Control - Estimation (No Covariates)

Ideal: Match units based on factor loadings

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

In Practice: Instead match on pre-treatment outcomes

$$\hat{y}_{NT}(0) = \sum_{j=1}^{N-1} \hat{w}_j y_{jT} \quad \hat{w} = \arg \min_{w \in \Delta} \sum_{t=1}^{T-1} \left(y_{Nt} - \sum_{j=1}^{N-1} w_j y_{jt} \right)^2$$

Synthetic Control - Estimation (Covariates)

Big Idea: Use covariates to "impute" missing lagged outcomes
If all lagged outcomes, yields case with no covariates.

2-step optimization:

1 Outer:

2 Inner:

1. Announcements & Discussion Questions
2. Difference-in-Differences
3. Synthetic Control
4. Alternative/New Panel Methods
5. Practice Problems

Synthetic Difference-in-Differences

Big Idea: Merge insights from DiD and SC

Estimation: $\hat{\tau}^{SDID} = \underset{\tau, \gamma, \alpha}{\operatorname{argmin}} \sum_i \sum_t \underbrace{(y_{it} - \gamma_t - \alpha_i - \tau w_{it})^2}_{\text{DiD}} \cdot \underbrace{w_i^{\text{adj}}}_{\substack{\text{SC} \\ \downarrow \\ \text{weight units} \\ \text{by similarity}}} \cdot \underbrace{\lambda_t}_{\substack{\downarrow \\ \text{weight time} \\ \text{periods} \\ \text{by} \\ \text{recency}}}$

Challenge: If I treated i, t , no averaging so no CLT

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

1 Placebo on Units:

2 Placebo on Time Periods:

Both are valid on average
Use version for dimension with less
heteroskedasticity.