

[This question paper contains **16** printed pages.]

Your Roll No_____

Sr. No. of Question Paper : **1391**
 Unique Paper Code : 2273100027
 Name of the Paper : Introduction to Causal Inference
 Name of the Course : Common Pool of DSE
 Semester : VI/VIII

Duration : 3 hours

Maximum Marks : 90

Instructions

There are two Sections, A & B with a total of 6 questions.

Question 1 from Section A is compulsory.

Answer any 4 out of the 5 questions from Section B.

Marks are given at the end of each question.

Use of simple calculators is allowed.

Section A

1. (a) Show how the Simple Difference in Outcomes (SDO) can be written as a function of the Average Treatment Effect on the Treated (ATT), a selection bias term, and a heterogeneous treatment effect bias term. Define each component carefully.

(b) How does unit randomization of treatment assignment eliminate both the selection bias term and the heterogeneous treatment effect bias term from the SDO?

(c) In a well-randomized experiment, why might a researcher still prefer to cluster standard errors at a level higher than the unit of randomization? Provide an example to illustrate.

(d) Explain the two assumptions that together constitute the Stable Unit Treatment Value Assumption (SUTVA). Why is each necessary for the potential outcomes framework to yield well-defined causal estimands?

(e) A government runs a school deworming program. Tablets are distributed only to children in treated schools. However, treated children share tablets with younger siblings at home, some of whom attend control schools. Additionally, in some treated schools the tablets distributed are a higher-dose formulation than in others, though all schools are classified simply as 'treated'.

Identify which component(s) of SUTVA are violated and explain each violation clearly. What are the implications for the estimated treatment effect? Suggest one design-based remedy for each violation.

(6 + 3 + 2 + 4 + 3)

Section-B

2. The Pradhan Mantri Ujjwala Yojana (PMUY) scheme provided free LPG connections to BPL households to reduce dependence on solid biomass cooking fuels. State M implemented the scheme aggressively in 2021, while comparable State N did not implement it until 2024. A researcher uses a DiD design to estimate the scheme's effect on indoor air pollution, measured as average PM_{2.5} concentration ($\mu\text{g}/\text{m}^3$) in BPL households, using data from 2020 (pre-treatment) and 2022 (post-treatment).

Variable	Coefficient	Std. Error	p-value
Intercept	95.0	2.5	0.000
Post (2022 = 1)	-4.0	2.8	0.155
Treated (State M = 1)	6.0	2.6	0.021
Treated $\times$ Post	-22.0	3.8	0.000

(a) Provide the estimating equation for this DiD model. What is the estimated effect of the Ujjwala scheme on indoor PM_{2.5} concentration? Interpret the DiD coefficient in plain language.

(b) Is the effect statistically significant at the 5% level? Justify using the appropriate values.

(c) What is the predicted PM_{2.5} concentration in 2022 for BPL households in State M?

- (d) What is the key identifying assumption of the DiD approach here?
- (e) Suggest one empirical strategy to test the plausibility of the parallel trends assumption in this setting.
- (f) Suppose the central government ran a nationwide clean cooking fuel awareness campaign in 2021 that had a stronger impact in State M due to higher baseline media penetration. How would this affect the DiD estimate? $(4 + 2 + 2 + 3 + 3 + 4)$

3. A state government runs a remedial coaching programme for students who score above 50 on a standardised qualifying examination. Students scoring above 50 are automatically enrolled in the coaching programme ($D = 1$); those scoring 50 or below are not enrolled ($D = 0$). A researcher uses a Regression Discontinuity Design (RDD) to estimate the causal effect of the coaching programme on the final examination score (Y). The running variable is the qualifying score (X), recentered at the cutoff as $\tilde{X} = X - 50$.

The following local linear regressions are estimated separately on each side of the cutoff, using only students with qualifying scores within a 20-point bandwidth (i.e., $X \in [30, 70]$).

Side of Cutoff	Estimated Regression Equation	SE on Intercept
Left ($D = 0$, score ≤ 50)	$\hat{Y} = 100.4 + 1.52 \times (X - 50)$	0.85
Right ($D = 1$, score > 50)	$\hat{Y} = 139.8 + 1.49 \times (X - 50)$	0.82

The following conditional mean final scores are observed in selected score bins:

Score Bin (X)	Mean Qualifying Score	Mean Final Score (Y)	N	D
[40, 45)	42.3	88.5	68	0
[45, 50)	47.4	96.1	73	0
(50, 55]	52.6	140.3	71	1
(55, 60]	57.5	148.1	65	1

- (a) State the sharp RDD treatment assignment rule formally. Using the regression results, compute the predicted final score at the cutoff ($X = 50$) from the left side and from the right side. Compute the estimated Local Average Treatment Effect (LATE) of the coaching programme. Show all steps clearly.

- (b) The estimated slopes on both sides of the cutoff are 1.52 (left) and 1.49 (right). Write the full estimating equation for a sharp RDD that allows the slope of the regression line to differ on each side of the cutoff. Interpret each coefficient in your equation. What does the near-equality of the two slopes suggest about the data-generating process?
- (c) Using only the nearest bins on each side of the cutoff from the table, compute a simple ‘difference-at-the-cutoff’ estimate of the LATE, adjusting for the slope using a value of 1.5 per score point. Compare this result with the local linear estimate from part (a). Explain why local linear regression is generally preferred over a simple comparison of nearest-bin averages.
- (d) State the continuity assumption for the sharp RDD in the context of this study. Explain in plain language why this assumption is required for the LATE to have a causal interpretation. Now, suppose the government also secretly provided extra study materials only to students who scored just above 50. (i) Would this violate the continuity assumption? (ii) In which direction would the RDD estimate be biased, and why?
- (e) A researcher estimates a simpler model without controlling for any polynomial in X , just regressing the final score Y on the treatment indicator D for all students. The estimated coefficient on D is 6,580 ($SE = 305.88$), even though the true causal effect is zero (in a hypothetical scenario where coaching has no effect but scores follow a nonlinear trend). (i) Compute the t -statistic. (ii) Explain why this estimate is spurious. (iii) What is the correct fix, and what general lesson does this provide about specification choice in RDD?
- (f) The McCrary density test is commonly used as a diagnostic check in RDD studies. (i) Explain what the McCrary test checks for and why it matters for RDD validity. (ii) Describe what pattern in the data would indicate that students are ‘manipulating’ their qualifying score to cross the 50-point threshold. (iii) If such manipulation were detected, explain how it would threaten the causal interpretation of the estimated LATE.

(3 + 3 + 3 + 3 + 3 + 3)

4. (a) Explain the assumptions involved in SUTVA. Why is each necessary for valid causal inference in the potential outcomes framework?

- (b) A public health program distributes insecticide-treated mosquito nets to households in randomly selected villages. Untreated villages in the vicinity benefit from reduced mosquito populations due to the herd protection effect created by treated villages. Additionally, the researcher notices that some households in treated villages did not use the nets at all.

Does this situation violate SUTVA? If so, identify which component(s) are violated and explain. What are the implications for causal inference, and suggest one design-based remedy for the spillover problem?

- (c) Researchers are studying the effect of class size on student test scores at the primary school level. Since class size is potentially endogenous (it may be correlated with school quality and parental choices), they use the number of births in a school's catchment area as an instrument for class size. The following results are obtained:

Dep. Variable	Indep. Variable	Coefficient	Std. Error	t-stat
First Stage (Class Size)	Births in cohort	0.45	0.12	3.75
	Constant	22.0	1.5	
	$R^2 = 0.22$			
Reduced Form (Test Score)	Births in cohort	-1.35	0.55	-2.45
	Constant	85.0	2.0	
	$R^2 = 0.08$			
IV Estimate (2SLS)	Class Size	-3.0	1.3	-2.31
	Constant	88.0	3.0	

- Interpret the coefficients from the first-stage and reduced-form regressions in the context of this study.
- Use the Wald estimator to manually compute the IV estimate. Show your working.
- Suppose the OLS estimate of the effect of class size on test scores is -1.5 . What can explain the difference between the OLS and IV estimates in this context?

(4 + 5 + 3 + 3 + 3)

5. In 2023, a state government launched a nutrition supplementation programme for children under age 5 in selected districts. Participation was voluntary. To estimate the programme's effect on child weight (in kg), a researcher uses two approaches:

Approach 1 — Propensity Score Matching (PSM): Using 2024 data, treated children are matched with control children on pre-treatment characteristics (baseline weight, mother's education, household income). The post-matching results are:

Group	Mean Child Weight (kg)
Treated	12.8
Matched Control	11.5
ATT	1.3

Approach 2 — Difference-in-Differences (DiD): The researcher uses panel data from 2022 (pre-treatment) and 2024 (post-treatment):

$$Y_{it} = \beta_0 + \beta_1 \cdot \text{Post}_t + \beta_2 \cdot \text{Treated}_i + \beta_3 \cdot (\text{Treated}_i \times \text{Post}_t) + \varepsilon_{it}$$

Variable	Coefficient	Std. Error	p-value
Intercept	10.2	0.3	0.000
Post (2024 = 1)	0.8	0.2	0.000
Treated	0.3	0.4	0.450
Treated × Post	0.9	0.3	0.003

- Interpret the ATT from PSM and the DiD coefficient separately. Why do they differ (1.3 vs. 0.9)?
- Explain the key identifying assumption that DiD relies on but PSM does not. How would you test its plausibility?
- Which estimate do you consider more credible in this context? Justify your answer.
- What is the predicted child weight in 2024 for a treated child based on the DiD model? Show your calculation.

- (e) Suppose the programme was rolled out in three phases across different districts. How would you modify the DiD framework to handle staggered treatment adoption? What additional assumption would be needed?

(4 + 4 + 3 + 3 + 4)

6. The PM-KISAN scheme provides ₹6,000 per year in direct income support to landholding farmer families. Since uptake is not fully universal — some eligible farmers did not register — a researcher uses Propensity Score Matching (PSM) to estimate the causal effect of PM-KISAN receipt on annual agricultural input expenditure (in ₹). The propensity score is estimated using a logit model with the following pre-treatment covariates: land holding size, household income, access to credit, and district fixed effects.

After 1:1 nearest-neighbour matching with replacement, the researcher estimates the following on the matched sample:

$$Y_i = \alpha + \tau \cdot D_i + \varepsilon_i$$

where Y_i is annual input expenditure (₹) and $D_i = 1$ if the farmer received PM-KISAN.

Variable	Coefficient	Std. Error	p-value
Intercept	8,200	350	0.000
Treated (D)	1,650	580	0.005

- (a) Interpret the coefficient on the Treated variable. What does this estimate represent?
- (b) Why is a logit model preferred over a Linear Probability Model (LPM) in constructing the propensity score?
- (c) What is the Conditional Independence Assumption (CIA) in this context? Why is it critical, and why is it fundamentally untestable?
- (d) Describe a method to assess whether the common support assumption is satisfied in this study.
- (e) Suppose a number of treated farmers have no close matches in the control group. What issue does this create, and how might the researcher address it?
- (f) What are the potential limitations of using nearest-neighbour matching with replacement in this context?

(3 + 3 + 3 + 3 + 3 + 3)

खंड क

1.(a) दर्शाइए कि परिणामों में सरल अंतर (एसडीओ) को उपचारित व्यक्तियों पर औसत उपचार प्रभाव (एटीटी), चयन पूर्वाग्रह पद और विषम उपचार प्रभाव पूर्वाग्रह पद के फलन के रूप में कैसे लिखा जा सकता है। प्रत्येक घटक को सावधानीपूर्वक परिभाषित कीजिए।

(b) उपचार $\square$ वॉटन के इकाई यादृच्छिकीकरण से एसडीओ से चयन पूर्वाग्रह पद और विषम उपचार प्रभाव पूर्वाग्रह पद दोनों कैसे समाप्त हो जाते हैं?

(c) एक सुव्यवस्थित यादृच्छिक प्रयोग में, शोधकर्ता यादृच्छिकीकरण की इकाई से उच्च स्तर पर मानक त्रुटियों को समूहित करना क्यों पसंद कर सकता है? इसे स्पष्ट करने के लिए एक उदाहरण दीजिए।

(d) स्थिर इकाई उपचार मूल्य धारणा (एसयूटीवीए) का निर्माण करने वाली दो मान्यताओं की व्याख्या कीजिए। संभावित परिणामों के ढांचे से सुस्पष्ट कारण अनुमान प्राप्त करने के लिए इनमें से प्रत्येक क्यों $\square$ वश्यक है?

(e) सरकार द्वारा स्कूलों में कृमिनाशक कार्यक्रम चलाया जाता है। गोलियां केवल उपचारित स्कूलों के बच्चों को ही वितरित की जाती हैं। हालांकि, उपचारित बच्चे घर पर अपने छोटे भाई-बहनों के साथ गोलियां साझा करते हैं, जिनमें से कुछ नियंत्रण स्कूलों में पढ़ते हैं। इसके अतिरिक्त, कुछ उपचारित स्कूलों में वितरित की जाने वाली गोलियां अन्य स्कूलों की तुलना में अधिक खुराक वाली होती हैं, हालांकि सभी स्कूलों को केवल 'उपचारित' के रूप में वर्गीकृत किया गया है।

SUTVA के किन घटकों का उल्लंघन हुआ है, यह पहचानें और प्रत्येक उल्लंघन की स्पष्ट व्याख्या करें। अनुमानित उपचार प्रभाव पर इसके क्या निहितार्थ हैं? प्रत्येक उल्लंघन के लिए एक डिज़ाइन- $\square$ धारित समाधान सुझाएँ।

(6 + 3 + 2 + 4 + 3)

खंड ख

2. प्रधानमंत्री उज्ज्वला योजना (पीएमयूवाई) के तहत निम्न और मध्यम आय वर्ग (बीपीएल) के परिवारों को खाना पकाने के लिए ठोस बायोमास ईंधन पर निर्भरता कम करने के लिए मुफ्त एलपीजी कनेक्शन प्रदान किए गए। राज्य एम ने 2021 में इस योजना को क्रामक रूप से लागू किया, जबकि तुलनात्मक रूप से समान राज्य एन ने इसे 2024 तक लागू नहीं किया। एक शोधकर्ता ने 2020 (उपचार से पहले) और 2022 (उपचार के बाद) के $\text{PM}_{2.5}$ कणों का उपयोग करते हुए, बीपीएल परिवारों में औसत पीएम_{2.5} सांद्रता ($\mu\text{g}/\text{m}^3$) के रूप में मापे गए इनडोर वायु प्रदूषण पर योजना के प्रभाव का अनुमान लगाने के लिए एक डिड डिज़ाइन का उपयोग किया।

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	95.0	2.5	0.000
बाद में (2022 = 1)	-4.0	2.8	0.155
उपचारित (अवस्था M = 1)	6.0	2.6	0.021
उपचारित × बाद में	-22.0	3.8	0.000

- (a) इस DiD मॉडल के लिए अनुमान समीकरण प्रदान कीजिए। उज्ज्वला योजना का घर के अंदर PM_{2.5} सांद्रता पर अनुमानित प्रभाव क्या है? DiD गुणांक की व्याख्या सरल भाषा में कीजिए।
- (b) क्या यह प्रभाव 5% स्तर पर सांख्यिकीय रूप से महत्वपूर्ण है? उचित मानों का उपयोग करके इसका औचित्य सिद्ध करें।
- (c) राज्य M में निम्न और मध्यम आय वर्ग (BPL) वाले परिवारों के लिए 2022 में PM_{2.5} की अनुमानित सांद्रता क्या है?
- (d) यहाँ DiD दृष्टिकोण की प्रमुख पहचान संबंधी मान्यता क्या है?
- (e) इस संदर्भ में समानांतर प्रवृत्तियों की धारणा की विश्वसनीयता का परीक्षण करने के लिए एक अनुभवजन्य रणनीति सुझाएँ।

(f) मान लीजिए कि केंद्र सरकार ने 2021 में स्वच्छ खाना पकाने के ईंधन के बारे में राष्ट्रव्यापी जागरूकता अभियान चलाया, जिसका राज्य M में उच्च स्तर की मीडिया पहुंच के कारण अधिक प्रभाव पड़ा। इससे DiD अनुमान पर क्या प्रभाव पड़ेगा? $(4 + 2 + 2 + 3 + 3 + 4)$

एक राज्य सरकार मानकीकृत योग्यता परीक्षा में 50 से अधिक अंक प्राप्त करने वाले छात्रों के लिए एक उपचारात्मक कोचिंग कार्यक्रम चलाती है। 50 से अधिक अंक प्राप्त करने वाले छात्रों को कोचिंग कार्यक्रम में स्वतः ही नामांकित कर दिया जाता है ($D = 1$); 50 या उससे कम अंक प्राप्त करने वाले छात्रों को नामांकित नहीं किया जाता है ($D = 0$)। एक शोधकर्ता अंतिम परीक्षा के अंक (Y) पर कोचिंग कार्यक्रम के कारण-कार्य प्रभाव का अनुमान लगाने के लिए रीग्रेशन डिसकॉन्टिन्यूटी डिज़ाइन (RDD) का उपयोग करता है। चर योग्यता अंक (X) है, जिसे कटऑफ पर पुनर्केंद्रित किया गया है। $\tilde{X} = X - 50$

निम्नलिखित स्थानीय रैखिक प्रतिगमन का अनुमान कटऑफ के प्रत्येक पक्ष पर अलग-अलग लगाया जाता है, जिसमें केवल 20-पॉइंट बैंडविड्थ (अर्थात, X) के भीतर योग्यता स्कोर वाले छात्रों का उपयोग किया जाता है। $I \in [30, 70]$).

कटऑफ की तरफ	अनुमानित प्रतिगमन समीकरण	इंटरसेप्ट σ एसई
बाएँ ($D = 0$, स्कोर ≤ 50)	$\hat{Y} = 100.4 + 1.52 \times (X - 50)$	0.85
सही ($D = 1$, स्कोर > 50)	$\hat{Y} = 139.8 + 1.49 \times (X - 50)$	0.82

चयनित स्कोर श्रेणियों में निम्नलिखित सशर्त औसत अंतिम स्कोर देखे गए हैं:

स्कोर बिन (X)	औसत योग्यता स्कोर	अंतिम औसत स्कोर (वर्ष)	एन	डी
[40, 45)	42.3	88.5	68	0
[45, 50)	47.4	96.1	73	0
(50, 55]	52.6	140.3	71	1
(55, 60]	57.5	148.1	65	1

(a) RDD उपचार निर्धारण नियम को स्पष्ट रूप से बताइए। प्रतिगमन परिणामों का उपयोग करते हुए, कटऑफ ($X = 50$) पर अनुमानित अंतिम स्कोर की गणना बाईं और दाईं ओर से कीजिए। कोचिंग कार्यक्रम के अनुमानित स्थानीय औसत उपचार प्रभाव (LATE) की गणना कीजिए। सभी चरणों को स्पष्ट रूप से दर्शाइए।

(b) कटऑफ के दोनों ओर अनुमानित ढलान क्रमशः 1.52 (बाएं) और 1.49 (दाएं) हैं। एक ऐसे सटीक RDD के लिए पूर्ण अनुमान समीकरण लिखिए जो कटऑफ के दोनों ओर प्रतिगमन रेखा के ढलान में अंतर की अनुमति देता है। अपने समीकरण में प्रत्येक गुणांक की व्याख्या कीजिए। दोनों ढलानों की लगभग समानता डेटा-उत्पादन प्रक्रिया के बारे में क्या दर्शाती है?

(c) तालिका से कटऑफ के दोनों ओर केवल निकटतम बिन्स का उपयोग करते हुए, LATE का एक सरल 'कटऑफ पर अंतर' अनुमान ज्ञात कीजिए, जिसमें प्रति स्कोर पॉइंट 1.5 के मान का उपयोग करके ढलान को समायोजित किया गया हो। इस परिणाम की तुलना भाग (a) से प्राप्त स्थानीय रैखिक अनुमान से कीजिए। समझाइए कि निकटतम बिन औसत की सरल तुलना की तुलना में स्थानीय रैखिक प्रतिगमन को सामान्यतः क्यों प्राथमिकता दी जाती है।

(d) इस अध्ययन के संदर्भ में शार्प आरडीडी के लिए निरंतरता की मान्यता स्पष्ट कीजिए। सरल शब्दों में समझाइए कि एलएटी को कारण-कार्य संबंध के रूप में समझने के लिए यह मान्यता क्यों आवश्यक है। अब मान लीजिए कि सरकार ने गुप्त रूप से केवल उन छात्रों को अतिरिक्त अध्ययन सामग्री प्रदान की जिन्होंने 50 से थोड़ा अधिक अंक प्राप्त किए। (i) क्या इससे निरंतरता की मान्यता का उल्लंघन होगा? (ii) आरडीडी अनुमान किस दिशा में पक्षपाती होगा, और क्यों?

(e) एक शोधकर्ता X में किसी भी बहुपद को नियंत्रित किए बिना एक सरल मॉडल का अनुमान लगाता है, जिसमें सभी छात्रों के लिए अंतिम स्कोर Y को उपचार संकेतक D पर प्रतिगमन किया जाता है। D पर अनुमानित गुणांक 6,580 (SE = 305.88) है, जबकि वास्तविक कारण प्रभाव शून्य है (एक काल्पनिक परिदृश्य में जहां कोचिंग का कोई प्रभाव नहीं है लेकिन स्कोर एक अरेखीय प्रवृत्ति का अनुसरण करते हैं)। (i) t-सांख्यिकी की गणना करें। (ii) समझाएं कि यह अनुमान भ्रामक क्यों है। (iii) सही समाधान क्या है, और इससे RDD में विनिर्देशन चयन के बारे में क्या सामान्य सबक मिलता है?

(f) मैकक्रैरी घनत्व परीक्षण का उपयोग आमतौर पर आरडीडी अध्ययनों में नैदानिक जांच के रूप में किया जाता है। (i) समझाइए कि मैकक्रैरी परीक्षण क्या जांचता है और आरडीडी वैधता के लिए यह क्यों महत्वपूर्ण है। (ii) आंकड़ों में ऐसा कौन सा पैटर्न है जो यह दर्शाता है कि छात्र 50 अंकों की सीमा पार करने के लिए अपने योग्यता स्कोर में हेरफेर कर रहे हैं? (iii) यदि ऐसा हेरफेर पाया जाता है, तो समझाइए कि यह अनुमानित एलएटी की कारण व्याख्या को कैसे खतरे में डाल सकता है।

(3 + 3 + 3 + 3 + 3 + 3)

5. (a) SUTVA में शामिल मान्यताओं की व्याख्या कीजिए। संभावित परिणामों के ढांचे में वैध कारण-कार्य संबंध स्थापित करने के लिए इनमें से प्रत्येक मान्यता क्यों आवश्यक है?
- (d) एक जन स्वास्थ्य कार्यक्रम के तहत यादृच्छिक रूप से चयनित गांवों के घरों में कीटनाशक-उपचारित मच्छरदानी वितरित की जाती हैं। उपचारित गांवों द्वारा उत्पन्न सामूहिक सुरक्षा प्रभाव के कारण आसपास के अनुपचारित गांवों में मच्छरों की संख्या में कमी आती है। इसके अतिरिक्त, शोधकर्ता ने पाया कि उपचारित गांवों के कुछ घरों ने मच्छरदानी का बिल्कुल भी उपयोग नहीं किया।
- क्या यह स्थिति SUTVA का उल्लंघन करती है? यदि हां, तो उल्लंघन किए गए घटक/घटकों की पहचान करें और व्याख्या करें। कारण-कार्य संबंध निष्कर्ष के लिए इसके क्या निहितार्थ हैं, और इस समस्या के लिए एक डिज़ाइन-आधारित समाधान सुझाएं?
- (e) शोधकर्ता प्राथमिक विद्यालय स्तर पर विद्यार्थियों के परीक्षा अंकों पर कक्षा के आकार के प्रभाव का अध्ययन कर रहे हैं। चूंकि कक्षा का आकार संभावित रूप से अंतर्जात है (यह विद्यालय की गुणवत्ता और अभिभावकों की पसंद से संबंधित हो सकता है), इसलिए वे विद्यालय के कार्यक्षेत्र में जन्मों की संख्या को कक्षा के आकार के लिए एक साधन के रूप में उपयोग करते हैं। निम्नलिखित परिणाम प्राप्त हुए हैं:

□ श्रित चर	स्वतंत्र चर	गुणक	मानक त्रुटि	टी स्टेट
पहला चरण (कक्षा का आकार)	एक ही समूह में जन्म	0.45	0.12	3.75
	स्थिर	22.0	1.5	
	$R^2 = 0.22$			
संक्षिप्त रूप (परीक्षा अंक)	एक ही समूह में जन्म	-1.35	0.55	-2.45
	स्थिर	85.0	2.0	
	$R^2 = 0.08$			
IV अनुमान (2SLS)	क्लास साइज़	-3.0	1.3	-2.31
	स्थिर	88.0	3.0	

- iv. इस अध्ययन के संदर्भ में प्रथम चरण और सरलीकृत रूप प्रतिगमन से प्राप्त गुणांकों की व्याख्या कीजिए।
- v. वाल्ड एस्टीमेटर का उपयोग करके मैक्युअल रूप से IV अनुमान की गणना करें। अपनी गणना का तरीका दिखाएँ।
- vi. मान लीजिए कि कक्षा के आकार का परीक्षा अंकों पर प्रभाव का OLS अनुमान -1.5 है। इस संदर्भ में OLS और IV अनुमानों के बीच अंतर को क्या स्पष्ट कर सकता है?

(4 + 5 + 3 + 3 + 3)

5. 2023 में, एक राज्य सरकार ने चयनित जिलों में 5 वर्ष से कम आयु के बच्चों के लिए पोषण पूरक कार्यक्रम शुरू किया। इसमें भागीदारी स्वैच्छिक थी। बच्चों के वजन (किलोग्राम में) पर कार्यक्रम के प्रभाव का अनुमान लगाने के लिए, एक शोधकर्ता दो दृष्टिकोणों का उपयोग करता है:

दृष्टिकोण 1 — प्रोपेंसिटी स्कोर मचिंग (पीएसएम): 2024 के आंकड़ों का उपयोग करते हुए, उपचारित बच्चों का मिलान उपचार-पूर्व विशेषताओं (आधारभूत वजन, माता की शिक्षा, घरेलू आय) के आधार पर नियंत्रण समूह के बच्चों से किया गया। मिलान के बाद के परिणाम इस प्रकार हैं:

समूह	बच्चों का औसत वजन (किलोग्राम में)
उपचार	12.8
मिलान नियंत्रण	11.5
एटीटी	1.3

दृष्टिकोण 2 — अंतर-में-अंतर (DiD): शोधकर्ता 2022 के पैनल डेटा का उपयोग करता है।

(उपचार से पहले) और 2024 (उपचार के बाद):

$$Y_{it} = \beta_0 + \beta_1 \cdot \text{Post}_t + \beta_2 \cdot \text{Treated}_i + \beta_3 \cdot (\text{Treated}_i \times \text{Post}_t) + \varepsilon_{it}$$

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	10.2	0.3	0.000
बाद में (2024 = 1)	0.8	0.2	0.000
उपचार	0.3	0.4	0.450
उपचारित × पोस्ट	0.9	0.3	0.003

- (a) पीएसएम सांप्राप्त एटीटी और डीआईडी गुणांक की अलग-अलग व्याख्या करें। इनमें अंतर क्यों है (1.3 बनाम 0.9)?
- (b) वह प्रमुख मान्यता स्पष्ट कीजिए जिस पर DiD आधारित है लेकिन PSM नहीं। आप इसकी सत्यता का परीक्षण कैसे करेंगे?
- (c) इस सदर्भ में आप किस अनुमान को अधिक विश्वसनीय मानते हैं? अपना उत्तर का औचित्य सिद्ध कीजिए।
- (d) DiD मॉडल का आधार पर उपचारित बच्चा का 2024 में अनुमानित वजन कितना होगा? अपनी गणना दर्शाइए।
- (e) मान लीजिए कि कार्यक्रम को अलग-अलग जिलों में तीन चरणों में लागू किया गया। चरणबद्ध उपचार अपनाए गए प्रक्रिया को समझाने के लिए आप DiD फ्रेमवर्क में क्या बदलाव करेंगे? इसका लिए और कौन सी अतिरिक्त मान्यताएँ आवश्यक होंगी?

(4 + 4 + 3 + 3 + 4)

6. पीएम-किसान योजना भूमिधारक किसान परिवारों को प्रति वर्ष 6,000 रुपये की प्रत्यक्ष आय सहायता प्रदान करती है। चूंकि इसका लाभ सभी किसानों को नहीं मिल रहा है (कुछ पात्र किसानों ने प्रमाणीकरण नहीं कराया है), इसलिए एक शोधकर्ता पीएम-किसान योजना का लाभ सांख्यिकी सख्खी वार्षिक व्यय (₹ में) पर पड़ना वाला प्रभाव का अनुमान लगाने के लिए प्रोपेंसिटी स्कोर मैचिंग (PSM) का उपयोग करता है। प्रोपेंसिटी स्कोर का अनुमान लॉजिट मॉडल का उपयोग करके लगाया गया है, जिसमें निम्नलिखित पूर्व-उपचार सहचर शामिल हैं: भूमि का आकार, घर की आय, ऋण तक पहुंच और जिला निश्चित प्रभाव।

प्रतिस्थापन सहित 1:1 निकटतम-पड़ोसी मिलान का बाद, शोधकर्ता मिलान किए गए नमूनों पर

निम्नलिखित का अनुमान लगाता है:

$$Y_i = \alpha + \tau \cdot D_i + \varepsilon_i$$

जहां Y_i वार्षिक इनपुट व्यय (₹) है और $D_i = 1$ यदि किसान को पीएम-किसान प्राप्त हुआ है।

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	8,200	350	0.000
उपचारित (डी)	1,650	580	0.005

(क) उपचारित चर पर गुणांक की व्याख्या कीजिए। यह अनुमान क्या दर्शाता है?

(ख) प्रोपेंसिटी स्कोर का निर्माण में लीनियर प्रोबबिलिटी मॉडल (एलपीएम) की तुलना में लॉजिट मॉडल को क्यों प्राथमिकता दी जाती है?

(ग) इस संदर्भ में सशर्त स्वतंत्रता धारणा (सीआईए) क्या है? यह महत्वपूर्ण क्यों है, और यह मूल रूप से

अप्रमाणित क्यों है?

(घ) इस अध्ययन में सामान्य समर्थन धारणा संतुष्ट है या नहीं, इसका आकलन करना की विधि का वर्णन कीजिए।

(ङ) मान लीजिए कि उपचारित किसानों में साफ़ई का नियंत्रण समूह में कोई करीबी मछ नहीं है। इससे क्या समस्या उत्पन्न होती है, और शोधकर्ता इसका समाधान कैसे कर सकता है?

(च) इस संदर्भ में प्रतिस्थापन का साथ निकटतम-पड़ोसी मिलान का उपयोग करना की संभावित सीमाएँ क्या हैं?

(3 + 3 + 3 + 3 + 3 + 3)