

[This question paper contains 12 printed pages.]

Your Roll No _____

Sr. No. of Question Paper : 7547
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) Express the Simple Difference in Outcomes (SDO) in terms of the Average Treatment Effect on the Treated (ATT), selection bias, and heterogeneous treatment effect bias. Carefully define each component.
- (b) Explain how unit randomization eliminates both selection bias and heterogeneous treatment effect bias from the SDO. What property of randomization is essential for this result?
- (c) Suppose selection into treatment is independent of Y^0 but not of Y^1 . What does this imply for unbiased estimation of the ATE? Of the ATT?
- (d) A researcher runs a randomized experiment to test the effect of a tutoring intervention on student scores. They find perfect covariate balance but still choose to include baseline test scores as a control variable in the regression. Explain the econometric rationale for this decision.

(e) A government evaluates the effect of a cash transfer program for farmers. Treated farmers, upon receiving cash, purchase fertilizers from local shops, which are also patronized by control farmers — allowing control farmers to indirectly benefit from increased fertilizer availability and lower prices.

Does this situation violate SUTVA? If so, which assumption and which component? What are the implications for causal inference, and suggest one design-based remedy?

(6 + 3 + 2 + 4 + 3)

Section B

2. A state government in India introduced a free textbook distribution scheme for government schools in 2022 to improve learning outcomes. Schools in Districts A, B, and C received the intervention; schools in Districts D, E, and F (similar in demographics) did not. You have data on average student test scores from 2021 (pre-treatment) and 2023 (post-treatment).

The regression output for a DiD model is given below:

Variable	Coefficient	Std. Error	p-value
Intercept	58.0	0.9	0.000
Post (2023 = 1)	1.8	1.1	0.105
Treated (Districts A–C = 1)	-1.2	1.0	0.235
Treated × Post	6.5	1.6	0.000

(a) Provide the estimating equation for this DiD setup. What is the estimated effect of the free textbook scheme on student test scores? Interpret the coefficient in plain language.

(b) Is the effect statistically significant at the 5% level? Justify your answer.

(c) What is the predicted test score in 2023 for schools in the treated districts?

(d) What is the key identifying assumption of the DiD method? Explain its importance in this context.

(e) Describe an event-study approach that could be used to assess the plausibility of the parallel trends assumption before treatment.

(f) Suppose the central government simultaneously launched a teacher training program in Districts A–C in 2022. How could this confound the DiD estimate and what does it represent?

(4 + 2 + 2 + 3 + 3 + 4)

3. A researcher studies the effect of access to credit on small business revenue using observational data. Since access to formal credit is not randomly assigned, the researcher uses bank branch density in a district as an instrumental variable for credit access. The regression results are:

Dependent Variable	Independent Variables	Coefficient	Std. Error	t-stat
First Stage (Credit Access = 1)	Branch density (per 100 sq km)	0.08	0.02	4.00
	Constant	0.30	0.05	
Reduced Form (Log Revenue)	Branch density (per 100 sq km)	0.048	0.016	3.00
	Constant	10.2	0.3	
IV Estimate (2SLS)	Credit Access	0.60	0.22	2.73

- Interpret the coefficients from the first-stage and reduced-form regressions.
- Use the Wald estimator to manually compute the IV estimate. Verify that it matches the 2SLS result given.
- Discuss the three key assumptions required for bank branch density to be a valid instrument. Where is the exclusion restriction most likely to be challenged?
- What additional information would you need to perform an F-test for weak instruments? What are the consequences of using a weak instrument in 2SLS?
- What explains a larger IV estimate compared to OLS in this setting? Who are the compliers in this IV framework?

(4 + 3 + 5 + 3 + 3)

4. The Central Government announced a Direct Benefit Transfer (DBT) scheme for cooking gas subsidies in 2022, eliminating a previous in-kind subsidy. States that adopted DBT early (by Q1 2022) form the treatment group, while states that adopted it later (after 2023) form the control group. You have annual data for both groups for 2021 (pre-treatment) and 2023 (post-treatment). The outcome variable is Average Monthly LPG Expenditure (in ₹) per household.

The DiD regression is:

$$\text{LPG_Exp}_{it} = \beta_0 + \beta_1 \cdot \text{Post}_t + \beta_2 \cdot \text{Treated}_i + \beta_3 \cdot \text{Treated}_i \cdot \text{Post}_t + \varepsilon_{it}$$

Variable	Coefficient	Std. Error	p-value
Intercept	480	18	0.000
Post	42	22	0.060
Treated	15	20	0.460
Post × Treated	-75	28	0.008

(a) What is the estimated effect of the DBT scheme on average monthly LPG expenditure? Interpret the DiD coefficient in plain language. Is the effect statistically significant at the 5% level? Justify your answer.

(b) What is the predicted average monthly LPG expenditure in 2023 for households in the treatment states?

(c) What is the identifying assumption of the DiD method? Why is it important for this study?

(d) Suppose treatment states were already experiencing a declining trend in LPG expenditure before the policy change due to rising household incomes. How does this violate the DiD assumption, and what direction of bias would result?

(e) Explain one robustness check or supplementary approach that could validate the DiD assumption in this context.

(4 + 2 + 4 + 4 + 4)

5. A state health department launched a community health worker (CHW) program in rural areas to improve maternal healthcare utilization. Participation was voluntary at the village level. The researcher suspects systematic selection — villages with pre-existing health infrastructure may be more likely to participate. To estimate the causal effect on institutional delivery rates (%), the researcher uses Propensity Score Matching (PSM). The propensity score is estimated using a logistic regression on baseline covariates: distance to the nearest PHC, baseline delivery rate, literacy rate, and village population.

After 1:1 nearest-neighbor matching (with replacement), the following regression is run on the matched sample:

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

where Y_i is institutional delivery rate (%) and $D_i = 1$ if the village has a CHW.

Variable	Coefficient	Std. Error	p-value
Intercept	52.0	3.5	0.000
Treated (D)	11.4	4.2	0.007

- (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 the context of PSM? Why is it untestable, and what can researchers do to make it more credible?
- (d) Describe a method to assess whether the common support assumption is satisfied in this study.
- (e) Suppose several treated villages have no comparable control matches. What issue arises, and how might the researcher address it without discarding all unmatched treated units?
- (f) What are the potential limitations of using nearest-neighbor matching with replacement in this context?

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

6. In 2022, a state government introduced a performance-linked incentive for rural primary health centres (PHCs) that scored below 50 on a composite health service quality index. PHCs with a score of 49 or below received additional funding and staff support; those scoring 50 or above did not. A researcher estimates the impact of this incentive on the percentage of children fully immunized using two methods:

Regression Discontinuity Design (RDD): PHCs with index scores between 40 and 60 are used in a local linear regression around the cutoff.

Difference-in-Differences (DiD): The researcher compares PHCs below and above the cutoff in 2021 (pre-policy) and 2023 (post-policy).

RDD Results (Local Linear Regression, Bandwidth = 10)

$$Y_i = \alpha + \tau \cdot D_i + f(\text{Score}_i - 50) + \varepsilon_i$$

Variable	Coefficient	Std. Error	p-value
Intercept	72.0	1.8	0.000
Treated (Score < 50)	4.5	1.6	0.006
Score (centered)	0.5	0.12	0.000

DiD Results

$$Y_{it} = \alpha + \delta \cdot \text{Post}_t + \gamma \cdot \text{Treated}_i + \tau \cdot (\text{Post}_t \cdot \text{Treated}_i) + \varepsilon_{it}$$

Variable	Coefficient	Std. Error	p-value
Intercept	71.5	1.1	0.000
Post	1.3	0.7	0.064
Treated	-0.8	1.0	0.420
Post × Treated	2.8	1.0	0.006

- (a) Assuming the RDD is valid, why might the RDD estimate be larger than the DiD estimate?
- (b) Explain how the comparison groups differ between the RDD and DiD methods in this context.
- (c) Discuss a situation in this context where the RDD estimate would be preferred over DiD.
- (d) What are the limitations of using only PHCs near the cutoff in RDD? How might this affect the external validity of the findings?
- (e) Under what conditions would DiD provide biased estimates, even if the RDD design appears valid?
- (f) Discuss the advantages of combining RDD with DiD (local DiD around the cutoff) to strengthen causal claims.

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

खंड क

1. (a) उपचारित व्यक्तियों पर औसत उपचार प्रभाव (ATT), चयन पूर्वाग्रह और विषम उपचार प्रभाव पूर्वाग्रह के संदर्भ में परिणामों में सरल अंतर (SDO) को व्यक्त कीजिए। प्रत्येक घटक को सावधानीपूर्वक परिभाषित कीजिए।
- (ख) समझाइए कि इकाई यादृच्छिकीकरण एसडीओ से चयन पूर्वाग्रह और विषम उपचार प्रभाव पूर्वाग्रह दोनों को कैसे समाप्त करता है। इस परिणाम के लिए यादृच्छिकीकरण का कौन सा गुण आवश्यक है?
- (c) मान लीजिए कि उपचार के लिए चयन Y^0 से स्वतंत्र है, लेकिन Y^1 से नहीं। इससे ATE और ATT के निष्पक्ष अनुमान पर क्या प्रभाव पड़ता है?

(d) एक शोधकर्ता छात्रों के अंकों पर ट्यूशन हस्तक्षेप के प्रभाव का परीक्षण करने के लिए एक यादृच्छिक प्रयोग करता है। उन्हें पूर्ण सहचर संतुलन प्राप्त होता है, फिर भी वे आधारभूत परीक्षा अंकों को प्रतिगमन में एक नियंत्रण चर के रूप में शामिल करना चुनते हैं। इस निर्णय के लिए अर्थमितीय तर्क स्पष्ट कीजिए।

(e) सरकार किसानों के लिए नकद हस्तांतरण कार्यक्रम के प्रभाव का मूल्यांकन करती है। नकद प्राप्त करने पर, लाभ प्राप्त किसान स्थानीय दुकानों से उर्वरक खरीदते हैं, जहाँ नियंत्रण समूह के किसान भी खरीदारी करते हैं - जिससे नियंत्रण समूह के किसानों को उर्वरक की बढ़ी हुई उपलब्धता और कम कीमतों से अप्रत्यक्ष रूप से लाभ मिलता है।

क्या यह स्थिति SUTVA का उल्लंघन करती है? यदि हाँ, तो कौन सी मान्यता और कौन सा घटक? कारण-कार्य संबंध निष्कर्ष के लिए इसके क्या निहितार्थ हैं, और एक डिज़ाइन-आधारित उपाय सुझाएँ?

(6 + 3 + 2 + 4 + 3)

खंड ख

2. भारत में एक राज्य सरकार ने 2022 में सरकारी स्कूलों में शिक्षा के परिणामों को बेहतर बनाने के लिए मुफ्त पाठ्यपुस्तक वितरण योजना शुरू की। ज़िला A, B और C के स्कूलों को यह योजना प्राप्त हुई; ज़िला D, E और F (जिनकी जनसांख्यिकी लगभग समान है) के स्कूलों को यह योजना प्राप्त नहीं हुई। आपके पास 2021 (योजना शुरू होने से पहले) और 2023 (योजना समाप्त होने के बाद) के छात्रों के औसत परीक्षा अंकों का डेटा है।

DiD मॉडल के लिए प्रतिगमन आउटपुट नीचे दिया गया है:

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	58.0	0.9	0.000
बाद में (2023 = 1)	1.8	1.1	0.105
उपचारित (जिले ए-सी = 1)	-1.2	1.0	0.235
उपचारित × बाद में	6.5	1.6	0.000

(a) इस DiD सेटअप के लिए अनुमान समीकरण प्रदान कीजिए। निःशुल्क पाठ्यपुस्तक योजना का छात्रों के परीक्षा अंकों पर अनुमानित प्रभाव क्या है? गुणांक की व्याख्या सरल भाषा में कीजिए।

(ख) क्या यह प्रभाव 5% स्तर पर सांख्यिकीय रूप से महत्वपूर्ण है? अपने उत्तर का औचित्य सिद्ध कीजिए।

(ग) उपचारित जिलों के विद्यालयों के लिए 2023 में अनुमानित परीक्षा स्कोर क्या है?

(घ) डीआईडी पद्धति की प्रमुख पहचान संबंधी मान्यता क्या है? इस संदर्भ में इसके महत्व को स्पष्ट कीजिए।

(ङ) ऐसे केस स्टडी दृष्टिकोण का वर्णन कीजिए जिसका उपयोग उपचार से पहले समानांतर रुझानों की धारणा की विश्वसनीयता का आकलन करने के लिए किया जा सकता है।

(f) मान लीजिए कि केंद्र सरकार ने 2022 में जिलों ए-सी में एक साथ शिक्षक प्रशिक्षण कार्यक्रम शुरू किया। इससे डीआईडी अनुमान कैसे प्रभावित हो सकता है और यह क्या दर्शाता है?

(4 + 2 + 2 + 3 + 3 + 4)

3. एक शोधकर्ता अवलोकन डेटा का उपयोग करके छोटे व्यवसायों के राजस्व पर ऋण की उपलब्धता के प्रभाव का अध्ययन करता है। चूंकि औपचारिक ऋण की उपलब्धता यादृच्छिक रूप से निर्धारित नहीं होती है, इसलिए शोधकर्ता ऋण उपलब्धता के लिए एक साधन चर के रूप में जिले में बैंक शाखाओं के घनत्व का उपयोग करता है। प्रतिगमन के परिणाम इस प्रकार हैं:

आश्रित चर	स्वतंत्र प्रभावित करने वाला वस्तुएँ	गुणांक	मानक त्रुटि	टी स्टैट
पहला चरण (क्रेडिट एक्सेस = 1)	शाखा घनत्व (प्रति 100 वर्ग किमी)	0.08	0.02	4.00
	स्थिर	0.30	0.05	
कम किया गया प्रारूप (लॉग राजस्व)	शाखा घनत्व (प्रति 100 वर्ग किमी)	0.048	0.016	3.00
	स्थिर	10.2	0.3	
IV अनुमान (2SLS)	क्रेडिट एक्सेस	0.60	0.22	2.73

(a) प्रथम चरण और सरलीकृत रूप प्रतिगमन से गुणांकों की व्याख्या कीजिए।

(b) वाल्ड एस्टीमेटर का उपयोग करके IV अनुमान की मैनुअल गणना कीजिए। सत्यापित कीजिए कि यह दिए गए 2SLS परिणाम से मेल खाता है।

(ग) बैंक शाखा घनत्व को एक वैध साधन मानने के लिए आवश्यक तीन प्रमुख मान्यताओं पर चर्चा कीजिए। अपवर्जन प्रतिबंध को सबसे अधिक चुनौती कहाँ दी जा सकती है?

(घ) कमजोर इंस्ट्रूमेंट्स के लिए एफ-टेस्ट करने के लिए आपको और कौन सी अतिरिक्त जानकारी की आवश्यकता होगी? 2SLS में कमजोर इंस्ट्रूमेंट का उपयोग करने के क्या परिणाम होते हैं?

(ङ) इस संदर्भ में ओएलएस की तुलना में उच्च आईवी अनुमान का क्या कारण है? इस आईवी ढांचे में संकलक कौन हैं?

(4 + 3 + 5 + 3 + 3)

4. केंद्र सरकार ने 2022 में खाना पकाने की गैस पर सब्सिडी के लिए प्रत्यक्ष लाभ हस्तांतरण (डीबीटी) योजना की घोषणा की, जिससे पहले दी जाने वाली वस्तुगत सब्सिडी समाप्त हो गई। जिन राज्यों ने डीबीटी को जल्दी (2022 की पहली तिमाही तक) अपनाया, वे उपचार समूह में आते हैं, जबकि जिन राज्यों ने इसे बाद में (2023 के बाद) अपनाया, वे नियंत्रण समूह में आते हैं। आपके पास दोनों समूहों के लिए 2021 (उपचार से पहले) और 2023 (उपचार के बाद) का वार्षिक डेटा उपलब्ध है। परिणाम चर प्रति परिवार औसत मासिक एलपीजी व्यय (₹ में) है।

DiD रिग्रेशन इस प्रकार है:

$$\text{LPG_Exp}_{it} = \beta_0 + \beta_1 \cdot \text{Post}_t + \beta_2 \cdot \text{Treated}_i + \beta_3 \cdot \text{Treated}_i \cdot \text{Post}_t + \varepsilon_{it}$$

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	480	18	0.000
बाद में	42	22	0.060
उपचार	15	20	0.460
उपचार के बाद ×	-75	28	0.008

(a) डीबीटी योजना का औसत मासिक एलपीजी खर्च पर अनुमानित प्रभाव क्या है? डीआईडी गुणांक की सरल भाषा में व्याख्या कीजिए। क्या यह प्रभाव 5% स्तर पर सांख्यिकीय रूप से महत्वपूर्ण है? अपने उत्तर का औचित्य सिद्ध कीजिए।

(ख) उपचारित राज्यों में परिवारों के लिए वर्ष 2023 में अनुमानित औसत मासिक एलपीजी व्यय क्या है?

(ग) डीआईडी पद्धति की पहचान संबंधी मान्यता क्या है? यह इस अध्ययन के लिए क्यों महत्वपूर्ण है?

(घ) मान लीजिए कि नीति परिवर्तन से पहले ही उपचारित राज्यों में घरेलू आय में वृद्धि के कारण एलपीजी व्यय में गिरावट का रुझान देखा जा रहा था। यह किस प्रकार DiD की मान्यता का उल्लंघन करता है, और इससे किस दिशा में पूर्वाग्रह उत्पन्न होगा?

(ङ) एक सुदृढ़ता जाँच या पूरक दृष्टिकोण की व्याख्या कीजिए जो इस संदर्भ में डी.आई.डी. धारणा को मान्य कर सके।

(4 + 2 + 4 + 4 + 4)

5. राज्य के स्वास्थ्य विभाग ने ग्रामीण क्षेत्रों में मातृ स्वास्थ्य देखभाल के उपयोग को बेहतर बनाने के लिए सामुदायिक स्वास्थ्य कार्यकर्ता (सीएचडब्ल्यू) कार्यक्रम शुरू किया। ग्राम स्तर पर भागीदारी स्वैच्छिक थी। शोधकर्ता को व्यवस्थित चयन का संदेह है - जिन गांवों में पहले से ही स्वास्थ्य संबंधी बुनियादी ढांचा मौजूद है, उनके भाग लेने की संभावना अधिक हो सकती है। संस्थागत प्रसव दरों (%) पर कारण प्रभाव का अनुमान लगाने के लिए, शोधकर्ता प्रोपेंसिटी स्कोर मैचिंग (पीएसएम) का उपयोग करता है। प्रोपेंसिटी स्कोर का अनुमान लॉजिस्टिक रीग्रेशन का उपयोग करके आधारभूत सहचरों पर लगाया जाता है: निकटतम प्राथमिक स्वास्थ्य केंद्र की दूरी, आधारभूत प्रसव दर, साक्षरता दर और ग्राम जनसंख्या।

1:1 निकटतम-पड़ोसी मिलान (प्रतिस्थापन सहित) के बाद, मिलान किए गए नमूने पर निम्नलिखित प्रतिगमन चलाया जाता है:

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

जहाँ Y_i संस्थागत प्रसव दर (%) है और $D_i = 1$ यदि गांव में एक सामुदायिक स्वास्थ्य कार्यकर्ता है।

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	52.0	3.5	0.000
उपचारित (डी)	11.4	4.2	0.007

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

(b) प्रोपेंसिटी स्कोर के निर्माण में लीनियर प्रोबेबिलिटी मॉडल (एलपीएम) की तुलना में लॉजिट मॉडल को क्यों प्राथमिकता दी जाती है?

(ग) पीएसएम के संदर्भ में सशर्त स्वतंत्रता धारणा (सीआईए) क्या है? यह अप्रमाणित क्यों है, और शोधकर्ता इसे अधिक विश्वसनीय बनाने के लिए क्या कर सकते हैं?

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

(e) मान लीजिए कि कई उपचारित गांवों में कोई तुलनीय नियंत्रण इकाई नहीं है। इससे क्या समस्या उत्पन्न होती है, और शोधकर्ता सभी बेमेल उपचारित इकाइयों को हटाए बिना इसका समाधान कैसे कर सकता है?

(f) इस संदर्भ में प्रतिस्थापन के साथ निकटतम-पड़ोसी मिलान का उपयोग करने की संभावित सीमाएँ क्या हैं? (3 + 3 + 3 + 3 + 3 + 3)

6. 2022 में, एक राज्य सरकार ने ग्रामीण प्राथमिक स्वास्थ्य केंद्रों (पीएचसी) के लिए प्रदर्शन-आधारित प्रोत्साहन योजना शुरू की, जिनका समग्र स्वास्थ्य सेवा गुणवत्ता सूचकांक पर स्कोर 50 से कम था। 49 या उससे कम स्कोर वाले पीएचसी को अतिरिक्त धन और कर्मचारी सहायता प्राप्त हुई; 50 या उससे अधिक स्कोर वाले पीएचसी को यह सहायता नहीं मिली। एक शोधकर्ता ने दो विधियों का उपयोग करके पूर्ण रूप से प्रतिरक्षित बच्चों के प्रतिशत पर इस प्रोत्साहन के प्रभाव का अनुमान लगाया है:

रिग्रेशन डिसकॉन्टिन्यूटी डिज़ाइन (आरडीडी): 40 और 60 के बीच सूचकांक स्कोर वाले प्राथमिक स्वास्थ्य केंद्रों का उपयोग कटऑफ के आसपास स्थानीय रैखिक प्रतिगमन में किया जाता है।

अंतर-पर-अंतर (DiD): शोधकर्ता 2021 (नीति से पहले) और 2023 (नीति के बाद) में कटऑफ से नीचे और ऊपर के प्राथमिक स्वास्थ्य केंद्रों की तुलना करता है।

आरडीडी परिणाम (स्थानीय रैखिक प्रतिगमन, बैंडविड्थ = 10)

$$Y_i = \alpha + \tau \cdot D_i + f(\text{स्कोर}_i - 50) + \varepsilon_i$$

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	72.0	1.8	0.000
उपचारित (स्कोर < 50)	4.5	1.6	0.006
स्कोर (केंद्रित)	0.5	0.12	0.000

did परिणाम

$$Y_{it} = \alpha + \delta \cdot \text{पोस्ट}_t + \gamma \cdot \text{इलाज}_i + \tau \cdot (\text{पोस्ट} \cdot \text{इलाज}_i) + \varepsilon_{it}$$

चर	गुणक	मानक त्रुटि	पी-मान
अवरोधन	71.5	1.1	0.000
डाक	1.3	0.7	0.064
इलाज	-0.8	1.0	0.420
उपचार के बाद ×	2.8	1.0	0.006

- (a) यदि RDD मान्य है, तो RDD अनुमान DiD अनुमान से बड़ा क्यों हो सकता है?
- (ख) इस संदर्भ में आर.डी.डी. और डी.आई.डी. विधियों के बीच तुलना समूहों में अंतर को स्पष्ट कीजिए।
- (ग) इस संदर्भ में ऐसी स्थिति पर चर्चा कीजिए जहां आरडीडी अनुमान को डीआईडी पर प्राथमिकता दी जाएगी।
- (घ) आरडीडी में कटऑफ के निकट स्थित पीएचसी का ही उपयोग करने की क्या सीमाएँ हैं? इससे निष्कर्षों की बाह्य वैधता पर क्या प्रभाव पड़ सकता है?
- (ङ) किन परिस्थितियों में डीआईडी पक्षपातपूर्ण अनुमान प्रदान करेगा, भले ही आरडीडी डिजाइन वैध प्रतीत हो?
- (च) कारण संबंधी दावों को मजबूत करने के लिए RDD को DiD (कटऑफ के आसपास स्थानीय DiD) के साथ संयोजित करने के लाभों पर चर्चा कीजिए।

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