

EVALUACIÓN DE IMPACTO USANDO STATA: UN ENFOQUE APLICADO DE POLÍTICA PÚBLICA EN EL ESTADO DE PUEBLA

EMMANUEL OLIVERA PÉREZ



Gobierno de Puebla
Hacer historia. Hacer futuro.



**Secretaría
de Educación**

CONCYTEP
Consejo de Ciencia
y Tecnología del Estado
de Puebla

EVALUACIÓN DE IMPACTO USANDO STATA:

UN ENFOQUE APLICADO DE POLÍTICA PÚBLICA EN EL ESTADO DE PUEBLA

EMMANUEL OLIVERA PÉREZ



Gobierno de Puebla
Hacer historia. Hacer futuro.



**Secretaría
de Educación**

CONCYTEP
Consejo de Ciencia
y Tecnología del Estado
de Puebla

Miguel Barbosa Huerta

GOBERNADOR CONSTITUCIONAL DEL
ESTADO DE PUEBLA

María del Rosario Orozco Caballero

PRESIDENTA DEL SISTEMA ESTATAL PARA EL
DESARROLLO INTEGRAL DE LA FAMILIA

Ana Lucía Hill Mayoral

SECRETARIA DE GOBERNACIÓN DEL ESTADO
DE PUEBLA

José Luis Sorcia Ramírez

SECRETARIO DE EDUCACIÓN DEL ESTADO DE PUEBLA

Sergio Salomón Céspedes Peregrina

PRESIDENTE DE LA JUNTA DE GOBIERNO Y
COORDINACIÓN POLÍTICA DEL H. CONGRESO
DEL ESTADO LIBRE Y SOBERANO DE PUEBLA

Margarita Gayosso Ponce

PRESIDENTA DEL TRIBUNAL SUPERIOR DE
JUSTICIA DEL ESTADO DE PUEBLA

Victoriano Gabriel Covarrubias Salvatori

DIRECTOR GENERAL DEL CONSEJO DE
CIENCIA Y TECNOLOGÍA DEL ESTADO DE
PUEBLA

Luis Gerardo Aguirre Rodríguez

RESPONSABLE DEL ÁREA DE PUBLICACIONES

Karla Rodríguez Rosas

REVISORA EDITORIAL

María Angélica Hernández Hernández

REVISORA DE ESTILO

Vanessa Salamanca Xaltenco

DISEÑADORA EDITORIAL Y DE PORTADA

Primera edición, México, 2022

Publicado por el Consejo de Ciencia y
Tecnología de Puebla (CONCYTEP) B Poniente
de La 16 de Sept. 4511, Col. Huexotitla, 72534.
Puebla, Pue.

ISBN: 978-607-8839-96-4

CÓDIGO IDENTIFICADOR CONCYTEP:

c-l-2022-11-119

La información contenida en este documento puede
ser reproducida total o parcialmente por cualquier
medio, indicando los créditos y las fuentes de origen
respectivas.



índice

Parte I

Métodos y prácticas	1
1. Introducción	2
2. Principios básicos de evaluación	3
Monitoreo vs evaluación	3
Monitoreo	4
Establecimiento de indicadores en un marco de seguimiento y evaluación	4
Monitoreo basado en resultados	5
Principales desafíos en la creación de un sistema de monitoreo	6
Evaluación operativa	6
Evaluación operativa vs evaluación de Impacto	7
Evaluaciones de impacto cuantitativas y cualitativas	7
Evaluación cuantitativa del impacto: evaluación ex post vs ex ante	8
Teoría básica de la evaluación de impacto: el problema del sesgo de selección	8
Diferentes enfoques de evaluación de impacto ex post	11
Descripción general: diseño e implementación de evaluaciones de impacto	12
3. Asignación aleatoria	13
Establecer el contrafactual	14
Cálculo de los efectos del tratamiento	15
Tipos de aleatorización	16
La aleatorización en el diseño de la evaluación: diferentes métodos de aleatorización	18
4. Diferencias en diferencias	21
Abordar el sesgo de selección desde una perspectiva diferente: el uso de las diferencias como contrafactualidad	21
Método DD: teoría y aplicación	22
Modelo de efectos fijos de panel	23
Implementando DD	24
Ventajas y desventajas de usar DD	25
Modelos DD alternativos	26



5. Métodos de regresión discontinua y de canalización	26
Introducción	26
Discontinuidad de la regresión en la teoría	27
Pasos para aplicar el enfoque de RD	28
Ventajas y desventajas del enfoque RD	30
6. Estimación de los efectos distributivos de los programas	30
Examen de los impactos heterogéneos de los programas: marco de regresión lineal	31
Enfoques de regresión cuantil	32
Efectos cuantiles del tratamiento utilizando intervenciones de programas aleatorios	32
Enfoques de regresión cuantil utilizando datos no experimentales	33
7. Evaluación de políticas mediante modelos económicos	35
Enfoques estructurales versus de forma reducida	35
Caso de modelado de los efectos de las políticas	36
8. Conclusiones	37
 Parte II	
Aplicaciones prácticas en stata 16	39
9. Introducción a Stata	40
Descripción de los datos	40
Ejercicio inicial: introducción a Stata	41
Trabajando con la base de datos	44
Cambio de conjuntos de datos	60
Combinación de conjuntos de datos	70
Trabajar con archivos .log y .do	73
10. Evaluación de impacto aleatoria	78
Impactos del programa	78
11. Método diferencias en diferencias	84
Estimación	85
12. Método de regresión discontinua	94
Estimación y aspectos adicionales	95
13. Bibliografía	104



Métodos y prácticas

1. Introducción

La evaluación de impacto es un tema de interés para las políticas públicas, ya que es necesario realizar dicha evaluación para verificar los verdaderos resultados de los programas implementados en materia de salud, educación, agrícola, entre otros. Los métodos para comprender si realmente funcionan, así como el nivel y la naturaleza de los impactos sobre los beneficiarios previstos, son temas que se abordan en el presente trabajo.

Existen diferentes organizaciones públicas y privadas, incluyendo el Gobierno, que diseñan una serie de programas encaminados a mejorar el bienestar social y contribuir al desarrollo económico. Sin embargo, a pesar de lo prometedor que pueden parecer los programas, cuando se aplican no necesariamente se obtienen los resultados previstos, puesto que existen múltiples factores a lo largo del tiempo que se pueden intervenir en los resultados del programa. Es por ello que en el presente trabajo se analizan diferentes metodologías que pueden dar un panorama más completo sobre los resultados que se obtienen de los programas implementados, siendo sumamente objetivo en encontrar los resultados reales que generaron estos, considerando la temporalidad, la zona objetivo, las características de los beneficiarios, las características del grupo control, e incluir una serie de parámetros que se pueden modificar para simular posibles situaciones. Esta serie de características son esenciales para los evaluadores y los agentes responsables de las políticas, pues aunque existen muchos programas que pudieran ser muy optimistas, los recursos son escasos y sólo se pueden elegir aquellos programas que tengan mejores resultados, mismos que se obtienen a través de estas metodologías.

La evaluación de impacto abarca métodos cualitativos y cuantitativos. El primero busca medir los impactos potenciales que el programa puede generar, es decir, proporciona información que puede ser crítica, para entender los mecanismos a través de los cuales el programa ayuda a los beneficiarios. Por su parte, el análisis cuantitativo abarca dos enfoques: ex ante y ex post. El diseño ex ante predice los resultados dadas ciertas suposiciones sobre el comportamiento individual y los mercados, y nos ayuda a determinar los posibles efectos positivos o negativos de una intervención a través de simulación o modelos económicos. Este análisis permite definir o redefinir los programas con mayor precisión antes de que se implementen. La evaluación de impacto ex post se centra en los datos reales que fueron recopilados después



de la intervención del programa. Estas miden los impactos reales acumulados por los beneficiarios.

Sería recomendable aplicar tanto métodos cualitativos como cuantitativos para lograr una evaluación de impacto efectiva, que permita evaluar con precisión los mecanismos mediante los cuales los beneficiarios están respondiendo a la intervención.

Este manual está organizado en dos secciones. En la parte I se distinguen los principios básicos de evaluación y se analiza la asignación aleatoria. Posteriormente, se estudian las diferencias en diferencias, siguiendo los métodos de regresión discontinua y de canalización, y la estimación de los efectos distributivos de los programas. Finalmente, se analiza la evaluación de políticas mediante modelos económicos y se establecen las principales conclusiones.

En la parte II se llevan a cabo una serie de ejercicios prácticos en el software Stata 16 en el contexto de evaluación empleando datos de la Encuesta de Movilidad Social, Ciudadana y Cohesión Social (ENBIENCOMÚN, 2015) para el estado de Puebla.

2. Principios básicos de evaluación

Monitoreo vs evaluación

Un sistema de monitoreo tiene como objetivo establecer metas e indicadores. Los resultados de este sistema se emplean para evaluar el desempeño de las intervenciones de un programa. En varios países se han estado implementando y desarrollando sistemas de monitoreo para analizar distintas iniciativas y su impacto en diferentes problemáticas, tal es el caso de la pobreza. Al momento de comparar los resultados del programa con metas específicas el monitoreo puede contribuir a mejorar el diseño y la implementación de políticas. Asimismo, promueve la rendición de cuentas y el diálogo entre los responsables de la formulación de políticas y las partes interesadas.

El Coneval señala que las evaluaciones de impacto permiten medir, mediante el uso de metodologías rigurosas, los efectos que un programa puede tener sobre su población beneficiaria y conocer si dichos efectos son en realidad atribuibles a su intervención. El principal reto de una evaluación de impacto es determinar qué habría pasado con los beneficiarios si el programa no hubiera existido. También es un instrumento que contribuye a la toma de decisiones y a la rendición de cuentas, es decir, aporta información tanto para actores a nivel gerencial como para los ciudadanos sobre la efectividad de los programas a los cuales se destina un presupuesto público.

Por su parte, el seguimiento y la evaluación (M&E) es una función de gestión continua para evaluar si se avanza en el logro de los resultados esperados, para detectar cuellos de botella en la implementación y para resaltar si hay efectos no deseados (positivos o negativos) de un plan, programa o proyecto (véase los casos de Tebani & Mederbal, 2018; Dos Santos & Raupp, 2015; Hernandez & Hornero, 2021 y Von *et al.*, 2022).

Monitoreo

Para analizar el progreso de una intervención se requiere lo siguiente:

- Identificar los objetivos del programa
- Identificar los indicadores clave que permitan monitorear el progreso de dicho programa en relación con los objetivos del mismo
- Establecer objetivos que permitan cuantificar el nivel de los indicadores
- Establecer un sistema de monitoreo que permita dar el seguimiento pertinente del progreso del programa hacia el logro de los objetivos
- Establecer reportes del seguimiento cada cierto periodo a los agentes responsables de la formulación de políticas

Establecimiento de indicadores en un marco de seguimiento y evaluación

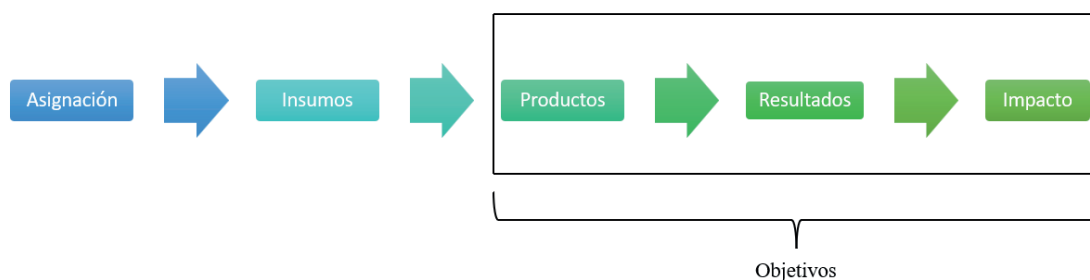
Se define indicador como una unidad de medida, la cual permite el monitoreo y evaluación de las variables clave de un sistema organizacional, mediante su comparación en el tiempo, con referentes externos e internos. De acuerdo al M&E, los indicadores se clasifican en dos grupos:

1. Los indicadores finales: Miden los resultados de los programas una vez establecidos, así como su impacto en términos de bienestar.
2. Los indicadores intermedios: Miden los insumos y los productos de un programa.



3. El marco de seguimiento y evaluación cubre todo el proceso de desempeño del programa, por lo que es importante seleccionar adecuadamente los indicadores que permitan dar el seguimiento adecuado, o sea, los que posibiliten una medición adecuada y oportuna.

Figura 2.1 Marco de seguimiento y evaluación



Monitoreo basado en resultados

El monitoreo basado en los resultados es una condición importante para el análisis de los programas y, por consiguiente, para la implementación de políticas. Hu en 2022 muestra la construcción de un sistema de monitorización. En la siguiente parte se resume los diez pasos para el monitoreo, basado en resultados, como parte de un marco de seguimiento y evaluación:

1. Se debe realizar una evaluación de preparación. Esto es que se debe tener claridad de las necesidades, características y actores clave que serán responsables de la implementación del programa.
2. Establecer los indicadores clave y especificar los resultados para monitorear y evaluar el desempeño del programa.
3. Los evaluadores deben seleccionar cómo medirán las tendencias de los resultados del programa. Dicha medición puede realizarse de manera cualitativa o cuantitativa.
4. Se deben determinar los instrumentos que se utilizarán para recopilar la información necesaria para evaluar el impacto del programa.
5. Se deben establecer objetivos que permitan monitorear los resultados considerando el espacio temporal.
6. Recopilar datos de buena calidad.
7. Reconocer cuando los indicadores se están alejando de las metas establecidas. Los evaluadores deben poder reaccionar con inmediatez si esto sucede para realizar los ajustes correspondientes.
8. Presentar los reportes correspondientes cada cierto periodo de tiempo.

9. Retroalimentar de acuerdo a los resultados obtenidos.
10. Esta serie de pasos debe efectuarse dentro del marco del sistema M&E.

Principales desafíos en la creación de un sistema de monitoreo

Los principales desafíos a enfrentar en la creación de un sistema de monitoreo efectivo son:

1. No definir claramente los indicadores
2. Existencia de ambigüedad
3. No identificar oportunamente cuando los indicadores se desvían de los objetivos
4. Información limitada o incompleta
5. Los evaluadores no efectuaron adecuadamente sus roles
6. Deficiencias en el análisis de datos

Cabe señalar que siempre se han de enfrentar diferentes obstáculos. Sin embargo, debemos saber cómo enfrentar dichas situaciones. Para ello se sugiere:

- a. Comprender los insumos y productos del programa
- b. Exhaustiva especificación de los indicadores
- c. Recopilar información detallada de la zona y los beneficiarios, manteniendo actualizada dicha información
- d. Considerar y pronosticar los efectos en caso de una externalidad
- e. Considerar la presencia de la heterogeneidad entre los beneficiarios

Evaluación operativa

La evaluación operativa busca realizar un análisis de los procesos implementados para la entrega de los productos establecidos por determinada política o programa, y determina de qué manera estos aportan o no al cumplimiento y el logro de los objetivos de la intervención. En otras palabras, permite identificar y retroalimentar de acuerdo con la brecha encontrada entre lo planeado y lo entregado. Esta información es importante para identificar cuáles fueron los errores, qué se puede modificar y qué se puede implementar para mejorar en futuros proyectos.

Cabe mencionar que es importante diseñar correctamente medidas de calidad de la implementación. Además se debe considerar la parte financiera que



involucra los costos y la manera en que se realizó la asignación de recursos, así como sus correspondientes efectos directos e indirectos. Esta información es imprescindible dado que nos ayuda a identificar posibles sesgos en la medición de los impactos del programa.

Evaluación operativa vs evaluación de Impacto

La cantidad de recursos públicos para la implementación de un programa se basa en los resultados que presenta el evaluador. Por consiguiente, es necesario medir eficientemente el impacto o los efectos de una intervención para que los responsables de la formulación de políticas puedan tener la certeza de decidir si vale la pena apoyar el programa o si este debe continuar, ampliarse o disolverse.

Como se mencionó con anterioridad, la evaluación operativa nos encamina a asegurar la implementación efectiva de un programa. La evaluación de impacto nos permite identificar si dicha intervención tuvo realmente impactos significativos en términos de bienestar. De acuerdo con la Figura 1, la evaluación de impacto se enfoca en tres etapas: productos, resultados e impacto.

La evaluación operativa y la evaluación de impacto son complementarias y partes integrales. Al considerar estas dos evaluaciones el estudio tiene mayor rigor, ya que ambas proporcionan la información óptima que permite la correcta interpretación de los resultados e impacto por parte de los evaluadores.

Evaluaciones de impacto cuantitativas y cualitativas

El análisis cuantitativo y cualitativo no son excluyentes, al contrario, se complementan. Proporcionan información relevante sobre la intervención, por lo que los evaluadores necesitan de ambos análisis para que puedan ser presentados a los agentes responsables de las políticas puesto que ellos necesitan información sólida. Se necesita el análisis cualitativo que describe la zona y los beneficiarios objetivo de dicho programa, y la parte cuantitativa proporciona los análisis estadísticos que fundamentan o respaldan los objetivos del programa. Por consiguiente, se recomienda a los evaluadores emplear un enfoque mixto (cuantitativo y cualitativo).

Evaluación cuantitativa del impacto: evaluación *ex post* vs *ex ante*

Para una evaluación cuantitativa es necesario y recomendable considerar una evaluación *ex ante* y *ex post*. La evaluación *ex ante* nos proporciona información de posibles escenarios dada una intervención, y es pronosticada mediante modelos con la información recabada; dichas simulaciones predicen los impactos del programa. Por otra parte, la evaluación *ex post* nos proporciona la información real de los impactos que son atribuibles a la intervención del programa. Estas evaluaciones tienen beneficios inmediatos y reflejan la realidad (véase Mendez, Everaert & Valcke, 2020).

¿Y si hubiera sido lo que no fue? Podría ser interesante saber cuál sería el impacto sobre los beneficiarios si el programa no hubiera existido, pero la realidad es que un individuo no puede tener dos existencias simultáneas. El resultado de un beneficiario en ausencia de la intervención sería su contrafactual. Sin embargo, es posible considerar la evaluación de los beneficiarios antes y después del tratamiento, o bien considerar una comparación entre el grupo tratado y no tratado (grupo control). Estas situaciones son consideradas su contrafactual falsificado, dado que podría generar una sobreestimación o una subestimación del efecto del programa. Esto es que los impactos resultantes no necesariamente son generados por el programa, puesto que es posible que estén interviniendo otros factores externos. Por ello, es recomendable tener mucha cautela a la hora de considerar estas comparaciones reflexivas.

Las comparaciones reflexivas son muy útiles, pero se debe tener mucha precisión y ser exhaustivo a la hora de recabar los datos antes de la intervención, ya que es necesario considerar múltiples variables para poder controlar factores que podrían estar cambiando con el tiempo.

Teoría básica de la evaluación de impacto: el problema del sesgo de selección

Dado que no se tiene información sobre el contrafactual, la mejor alternativa es comparar los resultados de las personas u hogares tratados con un grupo que no ha sido tratado. Es importante que ambos grupos sean similares, de tal forma que aquellos que recibieron tratamiento hubieran tenido resultados similares a los del otro grupo. El éxito de una evaluación de impacto depende considerablemente de encontrar un buen grupo de comparación. Para esto,



los investigadores recurren a dos enfoques para imitar el contrafactual de un grupo tratado:

- Crear un grupo de comparación a través de un diseño estadístico
- Modificar la estrategia de focalización del programa para eliminar las diferencias que han existido entre los grupos tratados y no tratados antes de comparar los resultados entre los dos grupos.

La ecuación 2.1 presenta el problema de evaluación básico que compara los resultados (Y) entre individuos tratados y no tratados (i):

$$Y_i = \alpha X_i + \beta T_i + E_i \quad (2.1)$$

T: Variable dummy, igual a 1 para los que participan y 0 para los que no participan

X: Es un conjunto de otras características observadas del individuo, de su hogar y entorno local

β : Estimación del efecto del programa

E: Término de error (refleja características no observadas que también afectan a Y)

La ecuación 2.1 refleja un enfoque comúnmente utilizado en las evaluaciones de impacto, que consiste en medir el efecto directo del programa T sobre los resultados Y. Cabe señalar que los efectos indirectos del programa (aquellos que no están directamente relacionados con la participación) también pueden ser de interés (se retoma en el tema 7).

La ecuación 2.1 puede presentar un sesgo. Esto implica que el término de error en la ecuación de estimación contiene variables que están correlacionadas con el tratamiento ficticio T y no se pueden medir, y, por lo tanto, explicar. Estas características no observadas en la ecuación 2.1 conducen a un sesgo de selección no observado, por lo que se viola uno de los supuestos clave de los mínimos cuadrados ordinarios, puesto que la cov (T, E) $\neq$ 0.

Por ejemplo, supongamos que se está evaluando un programa antipobreza, como una intervención crediticia, destinada a aumentar los ingresos de los hogares. Sea Y_i el ingreso per cápita del hogar i. Para los participantes $T_i=1$, y el valor de Y_i bajo tratamiento se representa como $Y_i(1)$. Para los no participantes $T_i = 0$ y Y_i se puede representar como $Y_i(0)$. Si se usa $Y_i(0)$ entre los hogares no participantes como resultado de comparación para

los resultados de los participantes $Y_i(1)$, el efecto promedio del programa podría representarse de la siguiente manera:

$$D = (E(T_1=1) - E(T_1=0)) \quad (2.2)$$

El problema que se presenta es que los grupos tratados y no tratados pueden no ser los mismos antes de la intervención, por lo que la diferencia esperada entre esos grupos puede no deberse por completo a la intervención del programa. Si en la ecuación 2.2 uno suma y resta el resultado esperado para los no participantes en el programa $E(Y_i(0) \mid T_i = 1)$, u otra forma de especificar el contrafactual, se obtiene:

$$D = E(T_1 = 1) - E(T_1 = 0) + E(T_1 = 1) - E(T_1 = 1) \quad (2.3)$$

$$D = ATE + E(T_1 = 1) - E(T_1 = 0) \quad (2.4)$$

$$D = ATE + B \quad (2.5)$$

En estas ecuaciones, ATE es el efecto promedio del tratamiento $[E(Y_i(1) \mid T_i = 1) - E(Y_i(0) \mid T_i = 1)]$. Es decir, la ganancia promedio en los resultados de los participantes en relación con los no participantes, como si los hogares no participantes también fueran tratados. La ATE corresponde a una situación en la que un hogar elegido al azar de la población es asignado para participar en el programa, por lo que los hogares participantes y no participantes tienen la misma probabilidad de recibir el tratamiento T .

El término $B [E(Y_i(0) \mid T_i = 1) - E(Y_i(0) \mid T_i = 0)]$ es el alcance del sesgo de selección que surge al usar D como una estimación del ATE. Como no se conoce $E(Y_i(0) \mid T_i = 1)$, no se puede calcular la magnitud del sesgo de selección. Como resultado, al no saber hasta qué punto el sesgo de selección constituye D , es posible que nunca se sepa la diferencia exacta en los resultados entre los grupos tratados y de control.

El objetivo esencial de una evaluación de impacto sólida es encontrar las formas de eliminar el sesgo de selección ($B = 0$), o bien encontrar formas de explicarlo. Un enfoque, discutido en el tema 3, es asignar aleatoriamente el programa. También se ha argumentado que el sesgo de selección desaparecería si se pudiera suponer que el hecho de que los hogares o las personas reciban o no tratamiento (condicionado a un conjunto de covariables X) será independiente de los resultados que tengan. Esta suposición se denomina suposición de



ausencia de confusión, igualmente conocida como suposición de independencia condicional (ver Lechner 1999; Rosenbaum y Rubin 1983).

$$(Y_i(1), Y_i(0)) \perp T_i | X_i \quad (2.6)$$

La solidez de las estimaciones del impacto depende de cuán justificables sean los supuestos sobre la comparabilidad de los grupos de participantes y de comparación, así como la exogeneidad de la focalización del programa en áreas tratadas y no tratadas. Sin ningún enfoque o suposición, no se podrá evaluar el alcance del sesgo B.

Diferentes enfoques de evaluación de impacto *ex post*

La teoría de la evaluación del impacto considera diferentes métodos para abordar la cuestión fundamental del contrafactual faltante. Cada uno de estos métodos conllevan sus propios supuestos sobre la naturaleza del posible sesgo de selección en los objetivos y la participación en el programa; y los supuestos son cruciales para desarrollar el modelo apropiado para determinar los impactos del programa. Estos métodos incluyen:

1. Evaluaciones aleatorias

Estos métodos varían según sus suposiciones subyacentes a cómo resolver el sesgo de selección al estimar el efecto del tratamiento del programa. Las evaluaciones aleatorias involucran una iniciativa asignada al azar a través de una muestra de sujetos. El progreso de los sujetos de tratamiento y de control que exhiben características similares previas al programa, se rastrea a lo largo del tiempo. Estos experimentos tienen la ventaja de evitar el sesgo de selección al nivel de la aleatorización.

2. Métodos de emparejamiento, específicamente emparejamiento por puntuación de propensión (PSM)

Los métodos de PSM comparan los efectos del tratamiento entre las unidades participantes y no participantes emparejadas, y se realizan en función de una serie de características observadas. Por lo tanto, los métodos de PSM asumen que el sesgo de selección se basa sólo en las características observadas; no pueden explicar los factores no observados que afectan la participación.

3. Métodos diferencias en diferencias (DD)

Los métodos DD asumen que existe una selección no observada y que es invariable en el tiempo: el efecto del tratamiento se determina tomando

la diferencia en los resultados entre las unidades de tratamiento y control, antes y después de la intervención del programa. Los métodos DD se pueden utilizar tanto en entornos experimentales como no experimentales.

4. **Métodos de variables instrumentales (IV)**

Los modelos IV pueden utilizarse con datos de corte transversal o de panel y, en este último caso, permiten que el sesgo de selección de las características no observadas varía con el tiempo. En este enfoque, el sesgo de selección en las características no observadas se corrige encontrando una variable (o instrumento) que esté correlacionada con la participación, pero no con las características no observadas que afectan al resultado. Este instrumento se utiliza para predecir la participación

5. **Métodos de diseño de discontinuidad de regresión (RD) y métodos de canalización**

Los métodos RD son una extensión de IV y métodos experimentales. Aprovechan las reglas exógenas del programa (como los requisitos de elegibilidad) para comparar participantes y no participantes en una zona cercana al límite de elegibilidad. Los métodos de canalización construyen un grupo de comparación de sujetos que son elegibles para el programa, pero que aún no lo han recibido.

6. **Impactos distributivos**

7. **Enfoques de modelado estructural y de otro tipo**

Los impactos distributivos de los programas y los enfoques de modelado pueden resaltar los mecanismos (las fuerzas intermedias del mercado) mediante los cuales los programas tienen un impacto.

Descripción general: diseño e implementación de evaluaciones de impacto

Para lograr una evaluación de impacto efectiva es necesario precisar la importancia y los objetivos de evaluación durante la identificación y preparación del proyecto. Para aislar el efecto del programa en los resultados, es elemental programar y estructurar las evaluaciones de impacto de antemano para ayudar a los funcionarios del programa a evaluar y actualizar la focalización, así como otras pautas para la implementación durante el curso de la intervención. Por



otra parte, también se debe considerar la disponibilidad y la calidad de los datos, puesto que son parte integral de la evaluación de los efectos del programa. El tipo de datos requerido dependerá del enfoque que se aplique, cuantitativo o cualitativo, o ambos; de si el marco es ex ante, ex post o ambos; o si se requieren nuevos datos se debe respetar el tiempo, el diseño y la selección de la muestra, y la selección de los instrumentos de encuesta apropiados. Para este caso, será recomendable realizar encuestas piloto en campo para que las preguntas de las entrevistas puedan revisarse y redefinirse en caso de ser necesario. La recopilación de datos sobre características socioeconómicas relevantes, tanto a nivel de beneficiario como a nivel comunitario, puede ayudar a comprender mejor el comportamiento de los encuestados dentro de sus entornos económicos y sociales.

Es importante que el personal involucrado y los evaluadores desempeñen sus roles correspondientes con profesionalismo y reflejen su experiencia, ética y moral. Cada miembro debe brindar los reportes correspondientes en tiempo y forma, así como mantener una constante actualización de los datos; esto permitirá que los agentes responsables de las políticas cuenten con la información oportuna sobre el progreso, así como de cualquier parámetro del programa que deba adaptarse a las circunstancias o tendencias cambiantes para permitir una retroalimentación potencialmente valiosa. Este aporte, además de los hallazgos de la propia evaluación, también puede ayudar a guiar el diseño de políticas futuras.

3. Asignación aleatoria de evaluación

La asignación aleatoria es un procedimiento utilizado en experimentos para crear múltiples grupos que incluyan participantes con características similares, de forma que los grupos sean equivalentes al inicio del estudio. El procedimiento implica asignar personas al tratamiento o programa experimental de manera al azar (como cuando echamos un volado). Esto significa que cada individuo tiene la misma oportunidad de ser asignado a uno de los grupos. Este procedimiento es una solución para evitar el sesgo de selección, siempre que los impactos del programa se examinen al nivel de la aleatorización. Es

importante realizar la selección de forma cuidadosa, en este tipo de tratamiento se distinguen dos efectos:

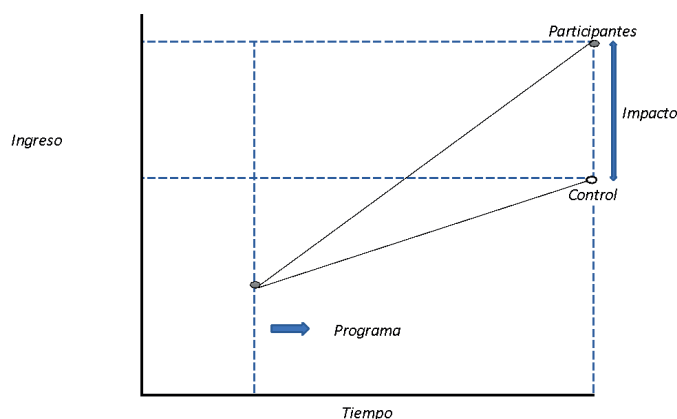
- Average Treatment Effect (ATE)
Es la diferencia esperada de los resultados de participar y no participar.
- Treatment Effect on the Treated (TOT)
Es la diferencia entre los valores del resultado esperado con y sin tratamiento para aquellos que participan. Los efectos de aquellos donde el programa es efectivamente aplicado.

Establecer el contrafactual

La asignación aleatoria es un método para resolver el problema del contrafactual, sin embargo, es importante que los investigadores formulen los procedimientos minuciosamente. Se debe de tener cuidado al seleccionar el grupo de control para asegurar la comparabilidad, y de esta manera obtener resultados más certeros (véase Seúl & Booyuel, 2016; Komro *et al.*, 2022 y Pinto *et al.*, 2021).

La Figura 3.1 ilustra gráficamente el caso de la aleatorización. Por ejemplo, considere una distribución aleatoria de dos grupos “similares” de hogares o individuos: un grupo recibe tratamiento y el otro grupo no. Son equivalentes en que ambos grupos antes de la intervención tienen el mismo nivel de ingresos (en este caso Y_0). Después de que se lleva a cabo el tratamiento, se encuentra que el ingreso observado del grupo tratado es Y_2 , mientras que el nivel de ingreso del grupo de control es Y_1 . Por lo tanto, el efecto de la intervención del programa se puede describir como $(Y_2 - Y_1)$, al igual se indica a continuación.

Figura 3.1 El experimento ideal con un grupo de control equivalente





A pesar de los esfuerzos por parte de los investigadores para encontrar grupos que sean realmente similares o equivalentes, en la realidad puede ser muy difícil, es por ello que los estadísticos han sugerido realizar este modelo en dos etapas:

- Etapa 1: Se selecciona aleatoriamente una muestra de posibles participantes de la población correspondiente. Esta muestra debe ser representativa de la población, dentro de un cierto error de muestreo. Esta etapa asegura la validez externa del experimento.
- Etapa 2: Los individuos de esta muestra se asignan al azar a grupos de tratamiento y comparación, lo que garantiza la validez interna en el sentido de que los cambios posteriores en los resultados medidos se deben al programa y no a otros factores.

Cálculo de los efectos del tratamiento

La aleatorización puede corregir el sesgo de selección B, discutido en el tema 2, mediante la asignación aleatoria de individuos o grupos a los grupos de tratamiento y control. Retomando el ejemplo (del tema 2), considere el problema clásico de medir los efectos del tratamiento: sea el tratamiento T_i igual a 1 si el sujeto i es tratado y 0 si no. Sea $Y_i(1)$ el resultado bajo tratamiento y $Y_i(0)$ si no hay tratamiento.

Observe Y_i y T_i , donde $Y_i = [T_i \cdot Y_i(1) + (1 - T_i) \cdot Y_i(0)]$, el efecto del tratamiento para la unidad i es $Y_i(1) - Y_i(0)$, y el ATE es $ATE = E[Y_i(1) - Y_i(0)]$, o la diferencia en resultados de estar en un proyecto en relación con el área de control para una persona o unidad extraída aleatoriamente de la población. Esta formulación supone, por ejemplo, que todos los miembros de la población tienen la misma probabilidad de ser objetivo.

$E[Y_i(1)|T_i = 1]$. Esta ecuación muestra los resultados promedio de los tratados, aquellos que están condicionados a estar en un área tratada. $E[Y_i(0)|T_i = 0]$. Esta ecuación muestra los resultados promedio de los no tratados, aquellos que están condicionados a no estar en un área tratada. Con focalización no aleatoria y observaciones en sólo una submuestra de la población, $E[Y_i(1)]$ no es necesariamente igual a $E[Y_i(1)|T_i = 1]$, y $E[Y_i(0)]$ no es necesariamente igual a $E[Y_i(0)|T_i = 0]$.

Por lo tanto, los efectos de tratamientos alternativos se observan en forma de TOT: $TOT = E[Y_i(1) - Y_i(0)|T_i = 1]$, o la diferencia en los resultados de recibir el

programa en comparación con estar en un área de control para una persona o sujeto extraído al azar de la muestra tratada. El TOT refleja las ganancias promedio de los participantes, condicionadas a que estos reciban el programa. Por ejemplo, supongamos que el área de interés es TOT, $E[Y_i(1) - Y_i(0)|T_i = 1]$. Si T_i no es aleatorio, una simple diferencia entre las áreas tratadas y de control $D = E[Y_i(1)|T_i = 1] - E[Y_i(0)|T_i = 0]$, no será igual al TOT. La discrepancia entre el TOT y este D será $E[Y_i(0)|T_i = 1] - E[Y_i(0)|T_i = 0]$, que es igual al sesgo B en la estimación del efecto del tratamiento (tema 2) :

$$TOT = E[Y_i(1) - Y_i(0)|T_i = 1] \quad (3.1)$$

$$= E[Y_i(1)|T_i = 1] - E[Y_i(0)|T_i = 1] \quad (3.2)$$

$$= D = E[Y_i(1)|T_i = 1] - E[Y_i(0)|T_i = 0] \quad \text{if} \quad E[T_i = 0] = E[Y_i(0)|T_i = 1] \quad (3.3)$$

$$\Rightarrow TOT = D \quad \text{if} \quad \beta = 0 \quad (3.4)$$

Tipos de aleatorización

1. Aleatorización pura

Si el tratamiento se llevara a cabo de forma puramente aleatoria siguiendo el procedimiento de dos etapas descrito anteriormente, entonces los hogares tratados y no tratados tendrían el mismo resultado esperado en ausencia del programa. Entonces, $E[Y_i(0)|T_i = 1]$ es igual a $E[Y_i(0)|T_i = 0]$. Dado que el tratamiento sería aleatorio, y no una función de las características no observadas (como la personalidad u otros gustos) entre los individuos, no se esperaría que los resultados hubieran variado para los dos grupos si no hubiera existido la intervención. Así, el sesgo de selección se convierte en cero en el caso de la aleatorización.

Considere el caso de aleatorización pura, donde una muestra de individuos u hogares se extrae al azar de la población de interés. Luego, la muestra experimental se divide aleatoriamente en dos grupos: (a) el grupo de tratamiento que está expuesto a la intervención del programa, y (b) el grupo de control que no recibe el programa. En términos de una regresión se puede expresar como:

$$Y_i = a + \beta T_i + E_i \quad (3.5)$$

Donde T_i es la variable ficticia de tratamiento igual a 1 si la unidad i se trata aleatoriamente y 0 en caso contrario. Y_i se define como:

$$Y_i = [Y_i(1) \cdot T_i] + [Y_i(0) \cdot (1 - T_i)] \quad (3.6)$$



Si el tratamiento es aleatorio (entonces T y E son independientes). La ecuación 3.5 se puede estimar usando mínimos cuadrados ordinarios (OLS), y el efecto del tratamiento β^{OLS} estima la diferencia en los resultados del grupo tratado y el de control. Si una evaluación aleatoria se diseña e implementa correctamente, se puede encontrar una estimación imparcial del impacto de un programa.

2. Aleatorización parcial

Una aleatorización pura es extremadamente rara de llevar a cabo, es por ello que comúnmente se utiliza la aleatorización parcial, en el que las muestras de tratamiento y de control se eligen al azar, condicionadas a algunas características observables X (por ejemplo, tenencia de la tierra o ingresos). Si se puede hacer una suposición llamada exogeneidad condicional de la colocación del programa, se puede encontrar una estimación no sesgada de la estimación del programa.

En este caso el modelo se basa en el caso de Zhu & Savla (2022). Denotando por simplicidad $Y_i(1)$ como Y_i^T $Y_i(0)$ como Y_i^C , la ecuación 3.5 podría aplicarse a una submuestra de participantes y no participantes de la siguiente manera:

$$Y_i^T = a^T + X_i \beta^T + \mu_i^T \text{ if } T_i = 1, i = 1, \dots, n \quad (3.7)$$

$$Y_i^C = a^C + X_i \beta^C + \mu_i^C \text{ if } T_i = 0, i = 1, \dots, n \quad (3.8)$$

Es una práctica común estimar lo anterior como una sola regresión, agrupando los datos de los grupos de control y tratamiento. Uno puede multiplicar la ecuación 3.7 por T_i y multiplicar la ecuación 3.8 por $(1 - T_i)$, y usar la identidad en la ecuación 3.6 para obtener:

$$Y_i = a^C + (a^T - a^C)T_i + X_i \beta^C + X_i(\beta^T - \beta^C) T_i + E_i \quad (3.9)$$

Donde $E_i = T_i (\mu_i^T - \mu_i^C) + \mu_i^C$. El efecto del tratamiento de la ecuación 3.9 se puede escribir como $ATT = E(T_i = 1, X) = E[a^T - a^C + X_i(\beta^T - \beta^C)]$. Aquí, ATT es sólo el efecto del tratamiento en el TOT tratado, discutido anteriormente.

Para la ecuación 3.9, se puede obtener una estimación consistente del efecto del programa con OLS si se puede suponer $E(\mu_i^T | X, T = t) = E(\mu_i^C | X, T = t) = 0$, $t = \{0, 1\}$. Es decir, no hay sesgo de selección debido a la aleatorización. En la práctica, a menudo se usa un modelo de impacto común que supone $\beta^T = \beta^C$. El ATE es entonces simplemente T-C.

La aleatorización en el diseño de la evaluación: Diferentes métodos de aleatorización

Los evaluadores tienen que decidir sobre qué tipo de aleatorización emplear, pues se tienen distintos enfoques. Los autores Noonan & Zhigljavsky (2022) presentan algunos de estos. Se resaltan los siguientes:

- **Suscripción excesiva**

Si los recursos limitados son una carga para el programa, la implementación puede asignarse aleatoriamente a un subconjunto de participantes elegibles, y los restantes que no reciben el programa pueden considerarse controles. Se debe hacer un examen del presupuesto, evaluando cuántos temas podrían ser encuestados versus aquellos realmente seleccionados, para dibujar un grupo de control lo suficientemente grande para la muestra de beneficiarios potenciales.

- **Incorporación aleatoria**

Este enfoque escalona gradualmente el programa en un conjunto de áreas elegibles, de modo que los controles representen áreas elegibles que aún esperan recibir el programa. Este método ayuda a aliviar los problemas de equidad y aumenta la probabilidad de que las áreas del programa y de control sean similares en las características observadas.

- **Aleatorización dentro del grupo**

En un enfoque de introducción aleatoria, si el retraso entre la génesis del programa y la recepción real de los beneficios es grande, puede surgir una mayor controversia sobre qué área o áreas deberían recibir el programa primero. En ese caso, todavía se puede introducir un elemento de aleatorización proporcionando el programa a algunos subgrupos en cada área objetivo. Por lo tanto, este enfoque es similar a la aleatorización gradual en una escala más pequeña. Un problema es que los efectos indirectos pueden ser más probables en este contexto.

- **Diseño de estímulos**

En lugar de aleatorizar el tratamiento, los investigadores asignan aleatoriamente a los sujetos un anuncio o un incentivo para participar en el programa. Se da algún aviso del programa por adelantado (ya sea durante el tiempo de la línea de base para conservar los recursos o, en general, antes de que se implemente el programa) a un subconjunto



aleatorio de beneficiarios elegibles. Este aviso se puede utilizar como un instrumento para la aceptación en el programa. Los efectos indirectos también podrían medirse bien en este contexto, si también se recopilan datos en las redes sociales de los hogares que reciben el aviso, para ver cómo la aceptación puede diferir entre los hogares que están conectados o no. Sin embargo, tal experimento requeriría una recopilación de datos más intensiva.

A pesar de que este método es muy utilizado en el área de la salud, los investigadores se enfrentan a diferentes preocupaciones, tales como los problemas éticos, la validez externa, el cumplimiento parcial o la falta de cumplimiento, la deserción selectiva y los efectos indirectos.

• Problemas éticos

Considera el programa PROGRESA (Programa de Educación, Salud y Alimentación, o Programa de Educación, Salud y Nutrición de México). Este se ha implementado mediante el método de asignación aleatoria, puesto que las personas o los beneficiarios que reciben el programa y los que no lo reciben pueden contar con las mismas características. Es decir, encontrarse en una situación de pobreza. De esta forma, cómo es posible justificar qué personas entran o no al programa si ambos grupos deberían de recibir el programa porque se encuentran en la misma situación. Sin embargo, dado que los recursos son limitados no es posible beneficiar a todos, y esta asignación no se realiza de manera arbitraria sino que se utilizan distintos softwares que permiten realizar este estudio de forma transparente y ética. Una cuestión importante es respaldar este método con información sólida, de manera que los agentes responsables de las políticas puedan tomar las decisiones más justas y éticas posibles, de la misma manera es necesario identificar con qué tipo de agentes se deben dirigir los evaluadores. En este caso, podemos encontrar a socios potenciales tales como el Gobierno (INPI, SAGARPA), Asociaciones Civiles (Save the Children, Ayuda en Acción) y empresas del sector privado (UPAEP).

• Validez interna y validez externa

Esta problemática se aborda en dos etapas:

Etapas 1:

Los agentes responsables de la formulación de políticas deben establecer con claridad la muestra aleatoria que se seleccionará para el análisis y la

población de la que se extrae. De esta manera, el experimento tendría validez externa, lo que implica que los resultados obtenidos podrían generalizarse a otros grupos o entornos. Este enfoque corresponde a la siguiente condición $E[Y_i(0)|T_i = 1] = E[Y_i(0)|T_i = 0]$ y $E[Y_i(1)|T_i = 1] = E[Y_i(1)|T_i = 0]$.

Etapa 2:

Es importante garantizar que el efecto del tratamiento sea una función únicamente de la intervención y no causado por otros elementos. Para alcanzar esta garantía se deben tomar medidas al asignar aleatoriamente esta muestra entre las condiciones de tratamiento y control. Este criterio se conoce como validez interna y refleja la capacidad de controlar los problemas que afectarían la interpretación causal del impacto del tratamiento. El sesgo sistemático (asociado con la selección de grupos que no son equivalentes, el desgaste selectivo de muestras, la contaminación de áreas objetivo por la muestra de control y los cambios en los instrumentos utilizados para medir el progreso y los resultados en el transcurso del experimento), así como el efecto de enfocarse en opciones relacionadas y resultados de los participantes dentro de la muestra objetivo, son ejemplos de tales problemas. La variación aleatoria en otros eventos que ocurren mientras el experimento está en progreso, aunque no representa una amenaza directa para la validez interna, también debe monitorearse dentro de la recopilación de datos porque una variación aleatoria muy grande puede representar una amenaza para la previsibilidad de la medición de datos.

• Efectos secundarios

Es de suma importancia tener cuidado con que no se dé una contaminación. Es decir, que el área de control y el área de tratamiento no se mezclen, de lo contrario se vería afectado el impacto del programa. Estas áreas deben estar ubicadas lo suficientemente separadas para evitar dicha contaminación del área del proyecto. Es una situación complicada pero no imposible, sobre todo en proyectos realizados a mayor escala. Enfrentar esta situación representa un gran reto, pues hay factores externos que no podemos controlar, si consideramos un proyecto de gran escala pueden darse diferentes situaciones como que los beneficiarios o los no beneficiarios se muden mezclándose entre las áreas, los beneficiarios pueden abandonar el programa y los no beneficiarios pueden verse



afectados indirectamente por el programa; estas situaciones afectarían los resultados del impacto del programa.

Estos efectos parciales del tratamiento pueden ser de interés, independiente del investigador, en especial porque es probable que sean significativos si la política se aplica a gran escala. Estos pueden abordarse midiendo los impactos por intención de tratar (ITT) o instrumentando la participación real en el programa mediante la estrategia de asignación aleatoria.

4. Diferencias en diferencias

Abordar el sesgo de selección desde una perspectiva diferente: el uso de las diferencias como contrafactualidad

El método de diferencias en diferencias se puede utilizar en datos de corte transversal y datos panel, o bien datos transversales repetidos, manteniendo la composición de los grupos de los participantes y del grupo control lo más estable posible a lo largo del tiempo.

Con datos panel, la estimación de DD resuelve el problema de los datos faltantes al medir los resultados y las covariables, tanto para los participantes como para los no participantes en los períodos previos y posteriores a la intervención. DD esencialmente compara los grupos de tratamiento y de comparación en términos de cambios de resultado a lo largo del tiempo en relación con los resultados observados para una línea de base previa a la intervención. Es decir, dada un escenario de dos períodos donde $t = 0$ antes del programa y $t = 1$ después de la implementación del programa, dejando que Y_i^T y Y_i^C sean los resultados respectivos para un beneficiario del programa y de las unidades no tratadas en el tiempo t . El método DD estimará el impacto promedio del programa de la siguiente manera:

$$DD = E(Y_i^T - Y_0^T | T_1 = 1) - E(Y_1^C - Y_0^C | T_1 = 0) \quad (4.1)$$

En la ecuación 4.1, $T_1 = 1$ denota tratamiento o la presencia del programa en $t = 1$, mientras que $T_1 = 0$ denota áreas no tratadas. El estimador DD permite

la heterogeneidad no observada (la diferencia no observada en los resultados contrafactuales medios entre las unidades tratadas y las no tratadas) que puede conducir a un sesgo de selección. Por ejemplo, es posible que se desee tener en cuenta factores no observados por el investigador, como las diferencias en la capacidad innata o la personalidad entre los sujetos tratados y los sujetos de control, o los efectos de la ubicación no aleatoria del programa en el nivel de formulación de políticas. DD asume que esta heterogeneidad no observada es invariable en el tiempo, por lo que el sesgo se cancela mediante la diferenciación. En otras palabras, los cambios de resultado de los no participantes revelan los cambios de resultado contrafactuales, como se muestra en la ecuación 4.1.

Método DD: teoría y aplicación

El estimador de diferencias en diferencias es un método de estimación de la inferencia causal estadística apropiado en el contexto de estudios observacionales. Su utilización en la evaluación de programas públicos de formación ha mostrado resultados ciertamente interesantes (véase el caso de Saldivar *et al.*, 2022; Kct, Devanaboyina, & Shaik, 2022; Na & Kim, 2022 y Jiao & Sun, 2021).

Cuando se dispone de datos de referencia se pueden estimar los impactos, suponiendo que la heterogeneidad no observada no varía en el tiempo y no está correlacionada con el tratamiento a lo largo del tiempo. Esta suposición es más débil que la exogeneidad condicional y hace que los resultados cambien para un grupo comparable de no participantes, es decir, $E(Y_1^c - Y_0^c | T_1) = 0$ como el contrafactual apropiado, o bien, igual $E(Y_1^c - Y_0^c | T_1) = 1$.

La estimación de DD también se puede calcular dentro de un marco de regresión. Esto se puede ponderar para tener en cuenta los posibles sesgos en DD. En particular, la ecuación de estimación se especificaría de la siguiente manera:

$$Y_{it} = a + \beta T_{i1}t + \rho T_{i1} + \gamma t + E_{it} \quad (4.2)$$

En la ecuación 4.2, el coeficiente β de la interacción entre la variable de tratamiento posterior al programa (T_{i1}) y el tiempo ($t = 1 \dots T$) da el efecto DD promedio del programa, usando la notación de la ecuación 4.1, $\beta = DD$. Además de este término de interacción, las variables T_{i1} y t se incluyen por separado para recoger cualquier efecto medio separado del tiempo, así



como el efecto de ser objetivo frente a no serlo. Una vez más, siempre que se disponga de datos sobre cuatro grupos diferentes para comparar, los datos de panel no son necesarios para implementar el enfoque DD.

Para comprender mejor la intuición detrás de la ecuación 4.2, se puede escribir en detalle en forma de expectativas (suprimiendo el subíndice i por el momento):

$$E(T_1=1) = (a + DD + p + y) - (a + p) \quad (4.3a)$$

$$E(T_1=0) = (a + p) - a \quad (4.3b)$$

Siguiendo la ecuación 4.1, restando 4.3b de 4.3a se obtiene DD. El método DD es imparcial sólo si la fuente potencial de sesgo de selección es aditiva e invariante en el tiempo. Utilizando el mismo enfoque, si se calcula un impacto simple antes y después de la estimación en la muestra de participantes (un diseño reflexivo), el impacto calculado del programa sería $DD + y$, y el sesgo correspondiente sería y . También se podría intentar comparar sólo la diferencia posterior al programa en los resultados entre las unidades de tratamiento y control. Sin embargo, en este caso, el impacto estimado de la política sería $DD + p$, y el sesgo sería p . Las diferencias sistemáticas no medidas que podrían estar correlacionadas con el tratamiento no pueden separarse fácilmente.

Para que el estimador DD anterior se interprete correctamente, se debe cumplir lo siguiente:

1. El modelo en la ecuación (resultado) está correctamente especificado
2. El término de error no está correlacionado con las otras variables en la ecuación:

$$\text{Cov}(E_{it}, T_{it}) = 0$$

$$\text{Cov}(E_{it}, t) = 0$$

$$\text{Cov}(E_{it}, T_{it}t) = 0$$

El último de estos supuestos, también conocido como supuesto de tendencia paralela, es el más crítico. Significa que las características no observadas que afectan la participación en el programa no varían a lo largo del tiempo con el estado del tratamiento.

Modelo de efectos fijos de panel

Esta posibilidad es particularmente importante para un modelo que controla no sólo la heterogeneidad invariable en el tiempo no observada, sino también la heterogeneidad en las características observadas en un entorno de

múltiples períodos (véase Pigni & Bartolucci, 2022; Piper, 2022; Sharma & Aggarwal, 2022 y Ramirez *et al.*, 2022). Específicamente, se puede hacer una regresión de Y_{it} sobre T_{it} , un rango de covariables variables en el tiempo X_{it} y heterogeneidad individual invariable en el tiempo no observada n_i que puede estar correlacionada, tanto con el tratamiento como con otras características no observadas E_{it} . Considere la siguiente revisión de la ecuación 4.2:

$$Y_{it} = \phi T_{it} + \delta X_{it} + n_i + E_{it} \quad (4.4)$$

Diferenciando el lado derecho del izquierdo de la ecuación 4.4 a lo largo del tiempo, se obtendría la siguiente ecuación diferenciada:

$$(Y_{it} - Y_{it-1}) = \phi (T_{it} - T_{it-1}) + \delta (X_{it} - X_{it-1}) + (n_i - n_i) + (E_{it} - E_{it-1}) \quad (4.5a)$$

$$\Rightarrow \Delta Y_{it} = \phi \Delta T_{it} + \delta \Delta X_{it} + \Delta E_{it} \quad (4.5b)$$

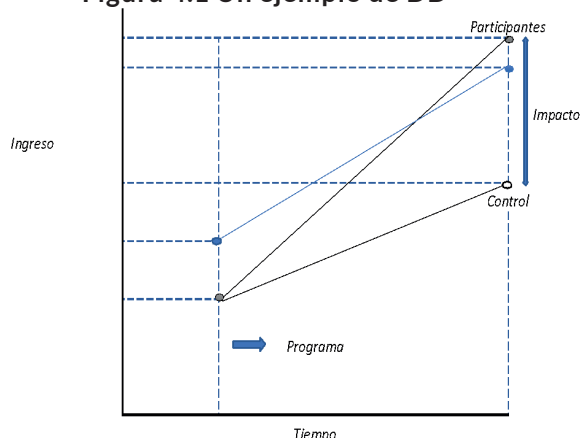
En este caso, debido a que la fuente de endogeneidad (es decir, las características individuales no observadas n_i) se elimina de la diferenciación, se pueden aplicar mínimos cuadrados ordinarios (OLS) a la ecuación 4.5b para estimar el efecto no sesgado del programa (ϕ). Con dos periodos de tiempo, ϕ es equivalente a la estimación DD en la ecuación 5.2, controlando por las mismas covariables X_{it} . Sin embargo, es posible que sea necesario corregir los errores estándar para la correlación serial (Menke & Creel, 2022). Con más de dos periodos de tiempo, la estimación del impacto del programa divergirá de DD.

Implementando DD

Para aplicar un enfoque DD utilizando datos de panel, se deben recopilar datos de referencia en las áreas del programa y de control antes de la implementación del programa. Es importante contar con información, tanto cuantitativa como cualitativa sobre estas áreas, ya que será útil para determinar quién es probable que participe. Las encuestas de seguimiento después de la intervención del programa también deben realizarse en las mismas unidades. Calcular la diferencia promedio en los resultados por separado para participantes y no participantes durante los periodos, y luego tomar una diferencia adicional entre los cambios promedio de los resultados para estos dos grupos obtendrá el impacto de la DD. En la Figura 4.1 se muestra un ejemplo: $DD = (Y_4 - Y_0) - (Y_3 - Y_1)$.



Figura 4.1 Un ejemplo de DD



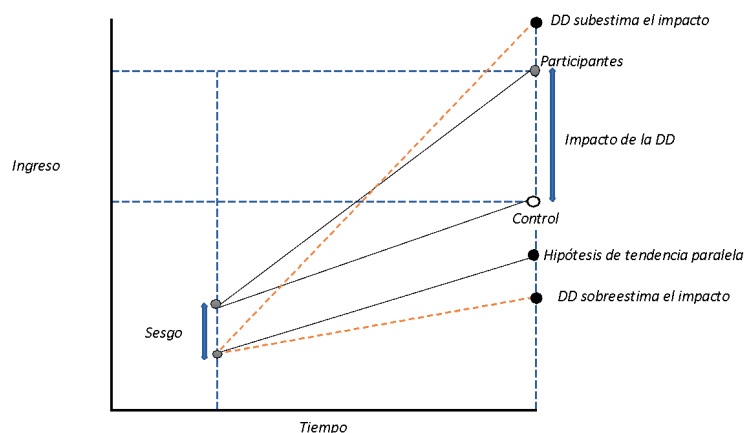
La línea inferior de la Figura 4.1 también representa los verdaderos resultados contrafácticos, que nunca se observan (ver capítulo 2). Bajo el enfoque DD, se supone que las características no observadas que crean una brecha entre los resultados de control medidos y los resultados contrafactuales verdaderos no varían en el tiempo, de modo que la brecha entre las dos tendencias es lo mismo durante el período. Esta suposición implica que $(Y_3 - Y_2) = (Y_1 - Y_0)$. Usando esta igualdad en la ecuación DD anterior, se obtiene $DD = (Y_4 - Y_2)$.

Ventajas y desventajas de usar DD

Tabla 4.1 Ventajas y desventajas de usar DD

Ventajas	Desventajas
Relaja la suposición de exogeneidad condicional o selección solo en las características observadas.	La noción de sesgo de selección invariable en el tiempo es inverosímil para muchos programas específicos en países en desarrollo.
Proporciona una forma manejable e intuitiva de dar cuenta de la selección de características no observadas.	Es probable que un método DD sobreestime el efecto del programa, en algunos casos.
La heterogeneidad no observada ex ante que varía en el tiempo podría explicarse con un diseño adecuado del programa, lo que incluye garantizar que las áreas del proyecto y de control compartan características previas al programa similares.	Es posible encontrar un sesgo cuando la diferencia entre los resultados de los no participantes y los contrafácticos cambia con el tiempo; la heterogeneidad no observada variable en el tiempo podría conducir a un sesgo hacia arriba o hacia abajo (vease la figura 4.2).
Si hay suficientes datos disponibles sobre otros factores exógenos o independientes del comportamiento que afectan a los participantes y no participantes a lo largo del tiempo, esos factores pueden explotarse para identificar impactos cuando la heterogeneidad no observada no es constante. Se podría implementar, por ejemplo, un enfoque de efectos fijos de panel de variables instrumentales.	Si las áreas de comparación no son similares a los participantes potenciales en términos de características observadas y no observadas, entonces los cambios en el resultado a lo largo del tiempo pueden ser una función de esta diferencia. Este factor también sesgaría el DD.

Figura 4.2 Heterogeneidad no observada variable en el tiempo



Modelos DD alternativos

Esta metodología es eficiente siempre y cuando la comunidad no observada y la heterogeneidad individual no varían en el tiempo. Sin embargo, en la realidad nos enfrentamos a constantes cambios, a una serie de externalidades que no podemos controlar, no se tiene conocimiento perfecto y hay incertidumbre. Al enfrentar una situación o situaciones en que las características no observadas de una población pueden cambiar con el tiempo, se han propuesto algunas variantes del método DD para controlar los factores que afectan estos cambios no observables, tales como PSM (Métodos de emparejamiento de puntuación de propensión) con DD, Método de triple diferencia y Ajuste de tendencias de tiempo diferencial.

5. Métodos de regresión discontinua y de canalización

Introducción

La regresión discontinua es un diseño cuasi-experimental que permite realizar inferencia causal en ausencia de aleatorización. Se puede aplicar cuando la exposición de interés se asigna, al menos parcialmente, según el valor de una variable aleatoria continua si esta cae por encima o por debajo de un determinado valor umbral. Por ejemplo, en muchos programas públicos, la decisión sobre quién recibe el programa y quién no se toma en función



de algún puntaje o variable continua, como lo puede ser el puntaje en una prueba o el nivel de ingresos de una familia. En estos casos resulta imposible implementar una selección aleatoria para medir el impacto del programa, y el analista deberá recurrir a técnicas cuasi-experimentales. Cuando el acceso al beneficio de un programa ocurre sólo para quienes tengan un puntaje que sobrepase algún umbral, la metodología por excelencia para evaluar impacto es el diseño de regresión discontinua (RD) (véase los casos de Steinmann & Olsen, 2022; Soliman, 2022; Tolentino, 2022 y Zhang, Liu & Wan, 2022)

Discontinuidad de la regresión en la teoría

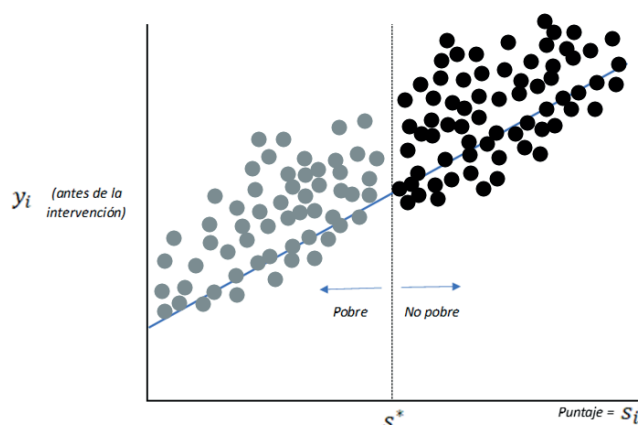
Para modelar el efecto de un programa en particular sobre los resultados individuales y_i a través de un enfoque RD, se necesita una variable S_i que determine la elegibilidad del programa (como la edad, la tenencia de activos, entre otros) con un límite de elegibilidad de s^* . La ecuación de estimación es $y_i = \beta S_i + E_i$, donde las personas con $s_i > s^*$. Por ejemplo, reciben el programa, y las personas con $s_i < s^*$ no son elegibles para participar. Los individuos en una banda estrecha por encima y por debajo de s^* deben ser “comparables” en el sentido de que se esperaría que logran resultados similares antes de la intervención del programa. La Figura 5.1 da un ejemplo de esta propiedad, donde los individuos por debajo de s^* se consideran pobres y los que están por encima del umbral se consideran no pobres.

Si se supone que existen límites a ambos lados del umbral s^* , el estimador de impacto para un $E > 0$ arbitrariamente pequeño alrededor del umbral sería el siguiente:

$$E[y_i | s^* - E] - E[y_i | s^* + E] = E[\beta S_i | s^* - E] - E[\beta S_i | s^* + E] \quad (5.1)$$

Tomando el límite de ambos lados de la ecuación 5.1 como $E \rightarrow 0$ identificaría a β como el cociente de la diferencia en los resultados de los individuos justo por encima y por debajo del umbral, ponderado por la diferencia en sus realizaciones de S_i .

Figura 5.1 Resultados antes de la intervención del programa



$$E[y_i | s^-] - E[y_i | s^+ + \epsilon] = y^- - y^+ = \beta(S^- - S^+) \\ \Rightarrow \beta = \frac{y^- - y^+}{S^- - S^+} \quad (5.2)$$

De acuerdo con la Figura 5.1, los resultados después de la intervención del programa medidos por el modelo de discontinuidad se reflejan en la Figura 5.2.

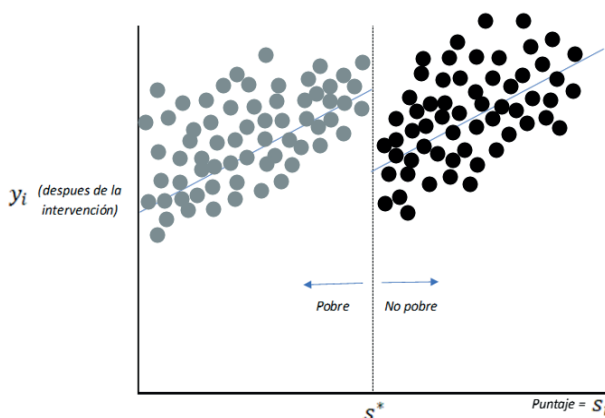
En la práctica, la determinación o aplicación de la elegibilidad puede no ser "aguda" (como en un experimento aleatorio), sino que puede reemplazarse con una probabilidad de participación $P(S) = E(T|S)$, donde $T = 1$ si el tratamiento se recibe, y $T = 0$ si no. En este caso, la discontinuidad es estocástica o "difusa" (véase los casos de Mi, Chen & He, 2022; Wang *et al.*, 2021 y Giovanis, Ozdamar & Ozdas, 2021), y en lugar de medir las diferencias en los resultados por encima y por debajo de s^* ; el estimador de impacto mediría la diferencia alrededor de un entorno de s^* . Este resultado puede ocurrir cuando las reglas de elegibilidad no se cumplen estrictamente o cuando se apunta a ciertas áreas geográficas pero los límites no están bien definidos y la movilidad es común. Si el umbral de elegibilidad está determinado exógenamente por el programa y está altamente correlacionado con el tratamiento, también se podría usar s^* como un IV para la participación.

Pasos para aplicar el enfoque de RD

La regresión no paramétrica estándar se puede utilizar para estimar el efecto del tratamiento en la configuración de discontinuidad de regresión aguda o difusa. Para un diseño de discontinuidad pronunciada, el efecto del tratamiento se puede estimar mediante una simple comparación de los resultados medios de los individuos a la izquierda y a la derecha del umbral.



Figura 5.2 Resultados después de la intervención del programa



Específicamente, se deben ejecutar regresiones lineales locales sobre el resultado y , dado un conjunto de covariables x , para personas en ambos lados del umbral, para estimar la diferencia $y^- - y^+$:

$$y^- - y^+ = E(y_i | S_i = S^*) - E(y_i | S_i = S^*) \quad (5.3)$$

Como ejemplo, y^- y y^+ se pueden especificar a través de estimaciones de Kernel:

$$y^- = \frac{\sum_{i=1}^n y_i \cdot \alpha_i \cdot K(u_i)}{\sum_{i=1}^n \alpha_i \cdot K(u_i)}$$

$$y^+ = \frac{\sum_{i=1}^n y_i \cdot (1 - \alpha_i) \cdot K(u_i)}{\sum_{i=1}^n (1 - \alpha_i) \cdot K(u_i)} \quad (5.4)$$

Para un diseño de discontinuidad difusa es necesario un proceso de dos pasos. Se puede aplicar una regresión lineal local sobre el resultado para las personas a ambos lados del umbral y determinar la magnitud de la diferencia (o discontinuidad) para el resultado. Del mismo modo, la regresión lineal local puede aplicarse al indicador de tratamiento y así llegar a una diferencia o discontinuidad para el indicador de tratamiento. La relación entre la discontinuidad del resultado y la discontinuidad del tratamiento es el efecto del tratamiento para un diseño de discontinuidad difusa.

Ventajas y desventajas del enfoque RD

Tabla 5.1 : Ventajas y desventajas del enfoque RD

Ventajas	Desventajas
Produce una estimación no sesgada del efecto del tratamiento de la discontinuidad.	Produce efectos de tratamiento locales promedio que no siempre son generalizables.
Puede aprovechar una regla conocida para asignar el beneficio que es común en los diseños de políticas sociales.	El efecto se estima en la discontinuidad, por lo que, en general, existen menos observaciones que en un experimento aleatorio con el mismo tamaño de muestra.
Un grupo de hogares o individuos elegibles no necesitan ser excluidos del tratamiento.	La especificación puede ser sensible a la forma funcional, incluidas las relaciones e interacciones no lineales.

Las comparaciones de canalización se pueden usar con diseños de discontinuidad si un tratamiento se asigna sobre la base de alguna característica exógena y los participantes potenciales (quizás para un programa relacionado) están esperando la intervención. Este enfoque podría usarse en el contexto de un programa que espera una expansión presupuestaria. Por ejemplo, donde las personas que esperan tratamiento pueden usarse como grupo de comparación. En esta situación, se usaría el mismo enfoque de DR, pero con un subconjunto (dinámico) adicional de no participantes. Otro ejemplo podría ser cuando una inversión local, como una carretera. Es una fuente de mejoras adicionales en el mercado, de modo que las personas alrededor del límite de la carretera se beneficien de las intervenciones futuras; la variación en la exposición potencial en función de la distancia desde la carretera podría explotarse como fuente de identificación (véase el caso de Lu, & Wang, 2022).

6. Estimación de los efectos distributivos de los programas

Los formuladores de políticas también deben de considerar la importancia de comprender cómo los programas han afectado a los hogares o individuos a lo largo de la distribución de los resultados, sobre todo en los programas sociales, aquellos encaminados a disminuir o erradicar la pobreza. El efecto medio o promedio de un programa o política, basado en el supuesto de un efecto común en todos los hogares o individuos a los que se dirige, es una forma concisa de evaluar su rendimiento. La finalidad de los programas es contribuir a un incremento del bienestar de la población, por ello es esencial analizar los



efectos distributivos de los programas para que los agentes puedan decidir oportunamente si el programa merece continuar, expandirse o disolverse, como el caso de PROGRESA, aunque cambió de nombre durante cada sexenio, el objetivo era el mismo. No cabe duda de que impactó positivamente en el bienestar de la población (véase los casos de Seth & Tutor, 2021; Campbell & Gaddis, 2017; Silverman, Holtyn & Jarvis, 2016 y Asadullah & Ara, 2016)

Examen de los impactos heterogéneos de los programas: marco de regresión lineal

Siguiendo el objetivo de analizar los efectos distributivos de los programas, existen varias formas de presentar los impactos distributivos de un programa. En un programa contra la pobreza el impacto puede ser tan directo como la proporción de personas objetivo que salieron de la pobreza. Analizar estos impactos es elemental para los formuladores de políticas, puesto que disminuir o erradicar la pobreza siempre ha sido el objetivo de cualquier país a lo largo del tiempo, por lo que existe un interés permanente para dar un seguimiento a las disparidades, la pobreza, la justicia y la desigualdad.

Para examinar cómo varía el impacto del programa entre diferentes individuos u hogares, es posible considerar un marco basado en una regresión lineal (véase los casos de Fedorova, Druchok & Drogovoz, 2022 y Bondonio, 2022). Los impactos heterogéneos del programa pueden representarse variando el intercepto α , el coeficiente β , o ambos en la variable T_i del programa, a través de los individuos $i = 1, \dots, n$:

$$Y_i = \alpha_i + \beta_i T_i + \gamma X_i + \varepsilon_i \quad (6.1)$$

Por ejemplo, se podría dividir la muestra de hogares e individuos en diferentes grupos demográficos (hombres y mujeres, o diferentes cohortes de edad) y ejecutar la misma regresión de T sobre Y por separado en cada grupo. Interactuar el tratamiento con diferentes características socioeconómicas del hogar X (como el género o la propiedad de la tierra) es otra forma de capturar las diferencias en los efectos del programa. Cabe señalar que agregar demasiados términos de interacción en la misma regresión puede generar problemas de multicolinealidad.

Para comprender la incidencia de las ganancias de un programa en un entorno más descriptivo, con datos antes y después de una intervención, los gráficos

pueden ayudar a resaltar los impactos distributivos del programa en muestras tratadas y de control, variando los resultados Y para las dos muestras frente a una covariable X_k dada. Las regresiones ponderadas localmente no paramétricas de Y en X_k se pueden graficar junto con diagramas de dispersión para brindar también una tendencia más suave de los patrones.

Enfoques de regresión cuantil

La regresión por cuantiles es otro enfoque para estimar los efectos del programa para un cuantil dado t en la distribución del resultado Y , condicionado a las covariables observadas X . Siguiendo el caso de los autores Bossman, Umar & Teplova, (2022), suponga que Y_i es una muestra de observaciones sobre el resultado y que X_i es un vector $K \times 1$ (que comprende el proyecto o tratamiento T , así como otras covariables). El modelo de regresión por cuantiles se puede expresar como:

$$Y_i = X_i + \beta_t X_i + \varepsilon_{ti}, Q_t(X_i) = \beta_t X_i, t \in (0,1) \quad (6.2)$$

Donde $Q_t(Y_i | X_i)$ denota el cuantil t del resultado Y (logaritmo del gasto per cápita), condicionado al vector de covariables (X). Específicamente, los coeficientes del cuantil pueden interpretarse como la derivada parcial del cuantil condicional de Y con respecto a uno de los regresores, como el programa T .

Efectos cuantiles del tratamiento utilizando intervenciones de programas aleatorios

En este método se encuentra un problema contrafactual al medir los impactos distributivos de un programa: el evaluador no sabe dónde aparecería la persona o el hogar i en la distribución de tratamiento, en la distribución de control o sin tratamiento. Sin embargo, si el programa es aleatorio, se puede calcular el efecto de tratamiento por cuantiles (QTE) de un programa. El QTE es la diferencia en el resultado y entre los hogares de tratamiento (T) y control (C) que, dentro de sus respectivos grupos, caen en el cuantil t de Y :

$$QTE = Y^T(t) - Y^C(t) \quad (6.3)$$

QTE refleja cómo cambia la distribución de los resultados si el programa se asigna al azar. Por ejemplo, para $t = 0,25$, QTE es la diferencia en el percentil 25 de Y entre los grupos de tratamiento y control. Sin embargo, QTE no identifica la distribución de los efectos del tratamiento, ni identifica el



impacto para los individuos en cuantiles específicos. Este problema también está relacionado con el desconocimiento del contrafactual. El QTE se puede derivar de las distribuciones marginales $F_T(Y) \equiv \Pr[Y_i^T \leq y]$ y $F_C(Y) \equiv \Pr[Y_i^C \leq y]$, ambos conocidos; pero los cuantiles del efecto del tratamiento $Y_i^T - Y_i^C$ no puede escribirse como una función sólo de las distribuciones marginales. Se necesitan suposiciones acerca de la distribución conjunta de Y_i^T y Y_i^C . Dado el caso que se supiera que la posición de un hogar en la distribución del gasto per cápita sería la misma independientemente de si el hogar estaba en un grupo tratado o de control, el QTE daría el efecto del tratamiento para un hogar en el cuantil t en la distribución; aunque esta suposición es muy fuerte.

Enfoques de regresión cuantil utilizando datos no experimentales

Calcular los impactos distributivos de una intervención no aleatoria requiere supuestos adicionales sobre el contrafactual y el sesgo de selección. Empleando datos de dos períodos sobre grupos tratados y no tratados antes y después de la introducción del programa, se puede construir una estimación de diferencia en diferencia cuantil (QDD). En el enfoque QDD de la distribución contrafactual se calcula primero el cambio en Y a lo largo del tiempo en el q -ésimo cuantil del grupo de control y luego agregando este cambio al q -ésimo cuantil de Y (observado antes del programa) al grupo de tratamiento.

$$QDD_{Y(t)} = Y_0^T(t) + (Y_1^C(t) - Y_0^C(t)) \quad (6.4)$$

Uno de los supuestos subyacentes de QDD es que la distribución contrafáctica de Y para el grupo tratado es igual a $(0,1)$. Sin embargo, esto se basa en suposiciones potencialmente cuestionables sobre la distribución de Y , incluido que la distribución subyacente de las características no observadas que pueden afectar la participación es la misma para todos los subgrupos. Bajo los supuestos de QDD, el modelo estándar de diferencias en diferencias (DD) es un caso especial de QDD.

Considere las siguientes ecuaciones de regresión por cuantiles para datos de dos períodos (variantes de la ecuación 6.2) para estimar los efectos distributivos del crecimiento sobre el ingreso per cápita o el gasto de consumo $y_{it}, t=\{1,2\}$:

$$Q_t(\log y_{i1} | X_{i1}, \mu_i) = Y_t, X_{i1} + \mu_i, TE(0,1) \quad (6.5a)$$

$$Q_t(\log y_{i2} | X_{i2}, \mu_i) = Y_t, X_{i2} + \mu_i, TE(0,1) \quad (6.5b)$$

En estas ecuaciones:

X_{it} : Representa el vector de otras covariables observadas, incluido el tratamiento T.

μ_i : Es el efecto fijo del hogar no observado.

$Q_t(\log y_{i1} | X_{i1}, \mu_i)$: Denota el cuantil t del logaritmo del ingreso per cápita en el período 1, condicionado al efecto fijo y otras covariables en el período uno.

$Q_t(\log y_{i2} | X_{i2}, \mu_i)$: Es el cuantil condicional t del logaritmo del ingreso per cápita en el período dos.

A diferencia del modelo DD lineal, un cuantil condicional no se puede restar del otro para diferenciar μ_i , porque los cuantiles no son operadores lineales:

$$Q_t(\log y_{i2} - \log y_{i1} | X_{i1}, X_{i2}, \mu_i) \neq Q_t(\log y_{i2} | X_{i2}, \mu_i) - Q_t(\log y_{i1} | X_{i1}, \mu_i) \quad (6.6)$$

El efecto fijo μ_i puede especificarse como una función lineal de las covariables en los períodos 1 y 2 de la siguiente manera:

$$\mu_i = \Phi + \lambda_1 X_{i1} + \lambda_2 X_{i2} + \omega_i \quad (6.7)$$

Donde:

Φ : Es un escalar.

ω_i : Es un término de error no correlacionado con X_{it} , $t = \{1, 2\}$.

La sustitución de la ecuación 6.7 en cualquiera de los cuantiles condicionales de las ecuaciones 6.5a y 6.5b permite estimar los impactos distributivos en el gasto per cápita utilizando este procedimiento de estimación de cuantiles ajustados:

$$Q_t(\log y_{i1} | X_{i1}, \mu_i) = \Phi + (Y_t + \lambda_1^1) X_{i1} + \lambda_2^2 X_{i2}, TE(0,1) \quad (6.8a)$$

$$Q_t(\log y_{i2} | X_{i2}, \mu_i) = \Phi + (Y_t + \lambda_1^2) X_{i2} + \lambda_2^1 X_{i1}, TE(0,1) \quad (6.8b)$$

Las ecuaciones 6.8a y 6.8b utilizan una regresión de cuantiles lineales combinados, donde las observaciones correspondientes al mismo hogar se apilan como un par.



7. Evaluación de políticas mediante modelos económicos

Enfoques estructurales versus de forma reducida

El sesgo de selección y el problema del contrafactual no observado son los principales problemas de identificación abordados a través de diferentes vías, experimentales o no. Este efecto unidireccional es un ejemplo de un enfoque de estimación de forma reducida. Un enfoque de este tipo especifica un resultado Y_i del hogar o individual como una función de un programa T_i y otras variables exógenas X_i :

$$Y_i = \alpha + \beta T_i + \gamma X_i + \varepsilon_i \quad (7.1)$$

En la ecuación 7.1 se supone que el programa y otras variables X_i son exógenas. El enfoque del efecto del tratamiento es un caso especial de estimación de forma reducida, en un contexto donde T_i es apropiado para un subconjunto de la población y Y_i y X_i también se observan para grupos de comparación separados.

Los modelos estructurales pueden ayudar a crear un esquema para interpretar los efectos de las políticas a partir de las regresiones, particularmente cuando intervienen múltiples factores. Estos modelos especifican las interrelaciones entre las variables endógenas (como los resultados Y) y las variables o factores exógenos.

Un ejemplo de un modelo estructural es el siguiente sistema de ecuaciones simultáneas (Limphaibool *et al.*, 2021; Buzeti, 2022; Mutonyi *et al.*, 2022; Al-Sharafi *et al.*, 2022 y Rafiq & McNally, 2022).

$$Y_{1i} = \alpha_1 + \delta_1 Y_{2i} + p_1 Z_{1i} + \varepsilon_{1i} \quad (7.2a)$$

$$Y_{2i} = \alpha_2 + \delta_2 Y_{1i} + p_2 Z_{2i} + \varepsilon_{2i} \quad (7.2b)$$

Las ecuaciones 7.2a y 7.2b son ecuaciones estructurales para las variables endógenas Y_{1i} , Y_{2i} . En el caso de Z_{2i} y Z_{1i} son vectores de variables exógenas que abarcan, por ejemplo, las características individuales y del hogar, con $E[Z'\varepsilon] = 0$. La política T_i en sí podría ser exógena y, por lo tanto, estar incluida en uno de los vectores Z_{ki} , $k = 1, 2$. La imposición de restricciones de exclusión

de ciertas variables de cada ecuación permite resolver las estimaciones α_k , δ_k , p_k en el modelo.

Para generar una ecuación de forma reducida se pueden resolver las ecuaciones estructurales 7.2a y 7.2b. Si la ecuación 7.2b se reorganiza de modo que Y_{1i} esté en el lado izquierdo y Y_{2i} en el lado derecho, se puede tomar la diferencia de la ecuación 7.2a y la ecuación 7.2b para que Y_{2i} pueda escribirse como una función de las variables exógenas Z_{1i} y Z_{2i} :

$$Y_{2i} = \pi_0 + \pi_1 Z_{1i} + \pi_2 Z_{2i} + v_{2i} \quad (7.3)$$

En la ecuación 7.3, π_0 , π_1 y π_2 son funciones de los parámetros estructurales α_k , δ_k , p_k , y v_{2i} es un error aleatorio que es función de ε_{1i} , ε_{2i} y δ_k . La principal distinción entre el enfoque estructural y de forma reducida es que si se parte de un modelo de forma reducida como el de la ecuación 7.3, se pierden relaciones potencialmente importantes entre Y_{1i} y Y_{2i} descritas en las ecuaciones estructurales.

Caso de modelado de los efectos de las políticas

Para modelar las opciones de oferta laboral de los hogares frente a los cambios impositivos (véase los casos de Williams & Gashi, 2022; Apunyo *et al.*, 2022; Rabaté & Rochut, 2020; Johnsen & Reiso, 2020 y Tong *et al.*, 2020). En términos generales, las preferencias de los hogares (o individuos) se representan como una función de utilidad U , que depende de sus elecciones sobre el consumo (c) y la oferta laboral (L): $U = U(c, L)$. El problema del hogar o del individuo es maximizar la utilidad sujeta a una restricción presupuestaria: $p_c \leq y + wL + t$. Es decir, la restricción presupuestaria requiere que los gastos p_c (donde p son los precios de mercado para el consumo) no excedan el ingreso no laboral (y), otras ganancias del trabajo (wL , donde w es la tasa salarial de mercado para el trabajo) y una transferencia propuesta t de una intervención del programa. La solución a este problema de maximización son las opciones óptimas c^* y L^* (el nivel óptimo de consumo y oferta de mano de obra). Estas elecciones, a su vez, son una función del programa t , así como de las variables exógenas w , y y p .

La estimación del modelo y la derivación de los cambios en c y L del programa requieren datos sobre c , L , w , y , p y t en los hogares i . Es decir, a partir del modelo anterior, se puede construir una especificación econométrica de L ,



por ejemplo, entre los hogares i , en función de los precios y la intervención del programa:

$$L_i = f(w_i, y_i, p_i; t) \quad (7.7)$$

En el enfoque de modelado, también se podría hacer que la utilidad U dependiera de las características socioeconómicas exógenas del hogar X , así como $U = U(c, L; X)$, de modo que las opciones óptimas de oferta y consumo de mano de obra en la especificación econométrica también sean una función de X .

Los modelos económicos son mecanismos que pueden ayudar a los evaluadores a realizar los pronósticos *ex ante* estableciendo los supuestos y las formas funcionales de las ecuaciones cuidadosamente. Mediante los parámetros es posible comprender cómo podría funcionar el programa o la política en un entorno económico diferente si se cambiaran algunos parámetros, además pueden ayudar a reducir los costos puesto que mediante estos modelos se obtienen los resultados de los impactos que podrían generar los programas. Es una visión al futuro para los agentes responsables de las políticas.

8. Conclusiones

Los métodos de evaluación de impacto son esenciales para mejorar el bienestar social. Los programas y las políticas de desarrollo suelen estar diseñados para mejorar la calidad de vida de la sociedad, por lo que es necesario saber si los programas implementados orientados a diferentes sectores impactan realmente en la sociedad. Para esto se debe tener la mayor certeza posible de que los impactos se deben al programa implementado y no a otros factores externos, es por ello que existen diferentes metodologías *ex ante* y *ex post* que permiten diseñar y mejorar los programas, mantener un seguimiento constante, mejorar la rendición de cuentas, definir las asignaciones presupuestarias, orientar el diseño o rediseño del programa y las decisiones de políticas. Estos modelos, que implican un análisis cualitativo y cuantitativo, nos permiten verificar y mejorar la calidad, la eficiencia y efectividad de los programas en las diferentes etapas de la implementación.

La evaluación de impacto genera evidencia sobre qué programas o intervenciones funcionan, cuáles no lo hacen y la manera en qué podemos

mejorarlos para alcanzar mejores resultados en materia de desarrollo. En esta primera parte se plantearon diferentes metodologías de evaluación de impacto, en las cuales es posible contrastar los resultados entre los grupos de tratamiento y los grupos que no reciben tratamiento (grupo control). Las evaluaciones de impacto nos permiten explotar alternativas e implementación de un programa para probar innovaciones o analizar diferentes programas con el fin de evaluar el desempeño comparativamente.

En realidad, ningún método es perfecto o infalible, siempre existe un margen de error, por ello es aconsejable verificar los hallazgos con métodos alternativos, y si es posible, combinar diferentes métodos de evaluación ex ante y ex post, así como enfoques cuantitativos y cualitativos. Los evaluadores deben comprender en su totalidad el diseño y la implementación de una intervención, las metas y los mecanismos mediante los cuales se pueden lograr los objetivos del programa y las características detalladas de las áreas objetivo y no objetivo. Al realizar una buena evaluación de impacto a lo largo del programa y al comenzar temprano en el diseño e implementación del proyecto, es posible cambiar o ajustar ciertos aspectos del programa para mejorarlo, y de esta manera maximizar su impacto.



Aplicaciones prácticas en stata 16

9. Introducción a Stata

Descripción de los datos

Para realizar los ejercicios propuestos en este manual se utilizaron los datos obtenidos de la Encuesta de Movilidad Social, Ciudadanía y Cohesión Social en la Zona Metropolitana y el Estado de Puebla (En Bien Común 2015). La población objetivo fueron mujeres y hombres de entre 30 y 59 años de edad, pertenecientes a zonas urbanas y rurales del estado de Puebla. La finalidad de esta encuesta era conocer los patrones de movilidad social de personas que tienen ocupaciones laborales, así como de subgrupos poblacionales que pertenecen al grupo de 30-59 años. La muestra se eligió mediante la selección de unidades por aproximación sucesiva cada vez más específicas; es decir, se realizó muestreo de conglomerados, que se eligieron por etapas en unidades cada vez más específicas. Primero se eligieron municipios, luego áreas geoestadísticas básicas (AGEBs), se continuó con manzanas o caseríos de zonas rurales, de ahí viviendas y, por último, residentes que cumplan con el perfil requerido en la población objetivo. La encuesta se aplicó entre los meses de noviembre y diciembre de 2015 en 49 municipios del estado de Puebla aleatoriamente seleccionados, y 19 municipios que integran la zona metropolitana del estado, los cuales son: Acatzingo, Acteopan, Atempan, Atlixco, Coyomeapan, Chiautla, Chignahuapan, Chignautla, Honey, General Felipe Ángeles, Guadalupe, Huaquechula, Huauchinango, Izúcar de Matamoros, Jalpan, Cañada Morelos, Naupan, Nopalucan, Oriental, Pahuatlán, Palmar de Bravo, Pantepec, Petlalcingo, Puebla, Quecholac, Quimixtlán, Rafael Lara Grajales, San Salvador el seco, San Sebastián Tlacotepec, Tecali de Herrera, Tecamachalco, Tehuacán, Tehuitzingo, Tepeaca, Tepemaxalco, Tepeyahualco, Tetela de Ocampo, Teziutlán, Tlacotepec de Benito Juárez, Tuzamapan de Galeana, Venustiano Carranza, Vicente Guerrero, Xochitlán de Vicente Suárez, Zacapoaxtla, Zacatlán, Zapotitlán, Zaragoza, Zautla y Zoquitlán. La encuesta cuenta con 3 003 casos levantados de manera polietápica y aleatoria para cada etapa. Se recolectó información socioeconómica para el entrevistado y retrospectiva para sus condiciones de origen. También cuenta con información sobre valores ciudadanos y ciudadanía, patrones de escolarización, condición laboral, situación económica de la vivienda, información sobre los padres del entrevistado, hábitos parentales, capital social, felicidad y religión. Es representativa para la Zona Metropolitana y para el estado de Puebla.



Ejercicio inicial: introducción a Stata

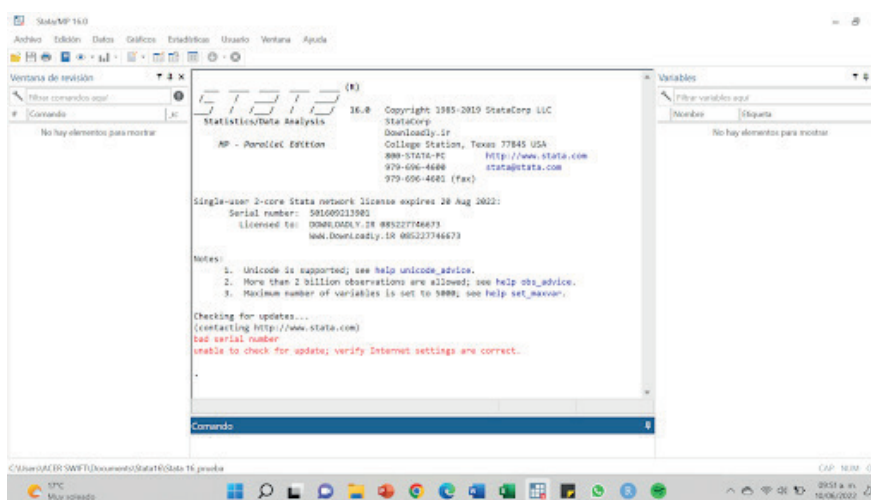
Stata es un software estadístico que provee todas las herramientas para la gestión, análisis y visualización de datos asociados a una interfaz gráfica potente y a la vez fácil de usar. Cuenta con una gran diversidad de procedimientos estadísticos que pueden ser aplicados en diferentes áreas y sectores, por lo que es ampliamente utilizado en investigación académica (Masini & Medeiros, 2022 y Larsen, Meng & Kendall, 2019), así como en entidades gubernamentales (Raza, 2022 y Wunsch & Gourbin, 2020), financieras (Ramos & Prats, 2020), comerciales y de servicios.

Ejecución de Stata

Esta sección ofrece una visión general de lo que ocurre en una sesión típica de Stata.

Ventanas de Stata

Cuando se inicia Stata, se abre una pantalla como la que se muestra en la Figura 9, que contiene cuatro ventanas etiquetadas:



El tamaño y la forma de estas ventanas se pueden cambiar y mover en la pantalla, ya que la interfaz de Stata tiene un menú y una barra de herramientas en la parte superior para realizar operaciones. Puede usar el menú y la barra de herramientas para ejecutar diferentes comandos, como abrir y guardar archivos de datos, o bien puede usar la ventana de comandos para realizar esas tareas.

Abrir una base de datos

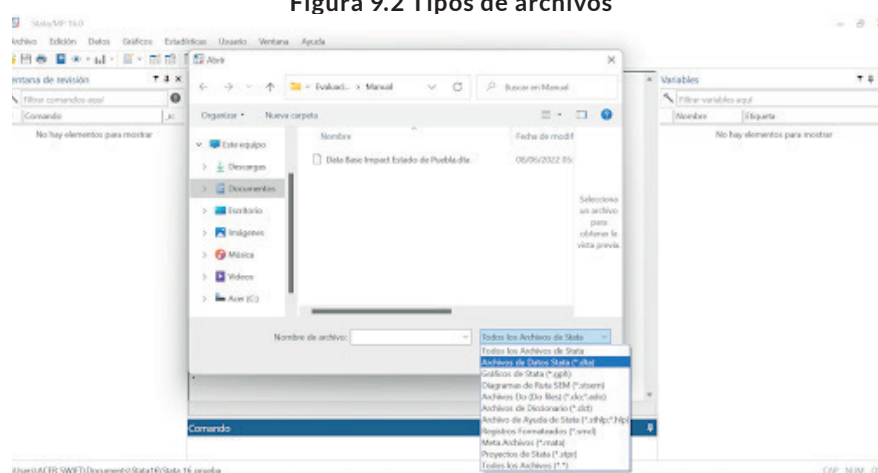
Para abrir una base de datos debe hacer clic en “Archivo” y luego en “Abrir”, y buscar el archivo que necesita, o bien usando el comando use y colocando entre comillas la dirección de ubicación de su archivo. Stata responde mostrando lo siguiente en la ventana “Resultados de Stata”:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Imp
> act Estado de Puebla.dta"
```

Tipos de archivos

Stata puede abrir diferentes tipos de archivos como se muestra en la Figura 9.2.

Figura 9.2 Tipos de archivos



Guardar una base de datos

Si realiza cambios en su base de datos y desea guardar estos, puede hacerlo utilizando el comando “save” de Stata. Por ejemplo:

```
. save "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Imp
> act Estado de Puebla.dta", replace
file C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact
> Estado de Puebla.dta saved
```

Replace le dice inequívocamente a Stata que sobrescribe la versión original preexistente con la nueva versión. Si no desea perder la versión original, debe especificar un nuevo nombre de archivo.

Salir de Stata

Una manera fácil de salir de Stata es emitir el comando “exit”. Sin embargo, si tiene abierto un conjunto de datos sin guardar, Stata emitirá el siguiente mensaje de error:

```
. exit
no; dataset in memory has changed since last saved
Save the data or specify option clear to exit anyway.
```



Para solucionar este problema, puede guardar el archivo de datos y luego emitir el comando "exit". Si realmente desea salir sin guardar el archivo de datos, primero puede borrar la memoria, usando el comando "clear" o "drop_all" y ejecutar el comando "exit". O bien puede simplificar el proceso combinando los dos comandos:
exit.clear

Recursos de memoria

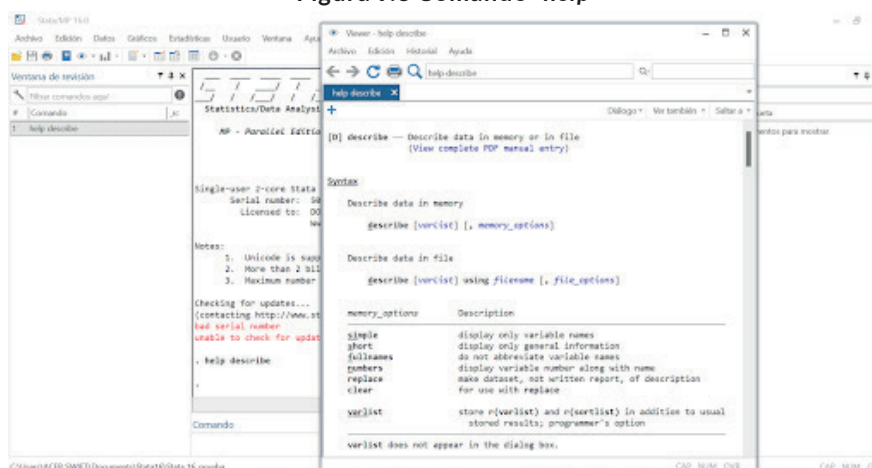
Por defecto, la memoria con la que trabajamos en un PC en Stata es de 1024K. Pero si la base de datos que vamos a manejar es muy grande, es posible que no tengamos memoria suficiente. Podemos incrementar los recursos de memoria con el comando "set memory". Afortunadamente, ya no es necesario ajustar la memoria en los Stata modernos; los ajustes de memoria se realizan sobre la marcha de forma automática.

```
. set memory 20000
set memory ignored.
Memory no longer needs to be set in modern Statas; memory adjustments are performed on the fly automatically.
```

Ayuda de Stata

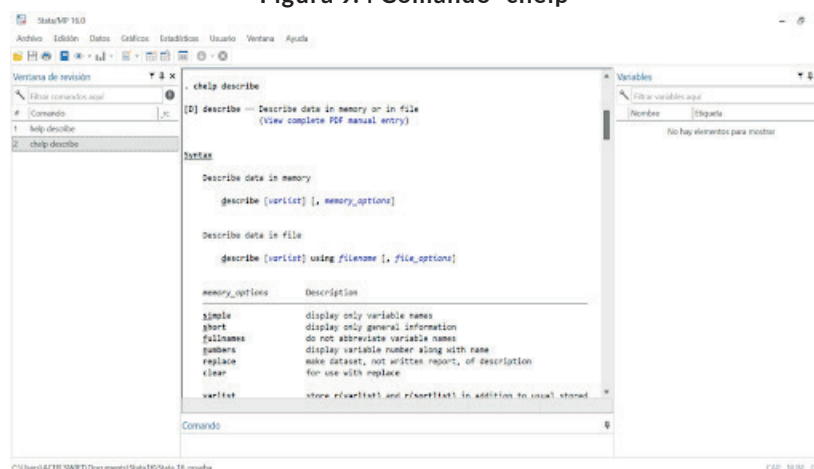
Stata viene con un excelente conjunto de manuales en varios volúmenes. Sin embargo, la ayuda en el ordenador es extensa y muy útil. Si tiene acceso a la Web, hay disponible un conjunto aún mayor de macros e información útil. El sistema de ayuda para los comandos de Stata es una de las herramientas que más rápidamente puede familiarizar al usuario con el manejo del programa. Alternativamente al sistema de ventanas, el usuario puede digitar en el cuadro de comandos "help", seguido del comando del cual desea información, como se muestra en la Figura 9.3

Figura 9.3 Comando "help"



O bien puede utilizar “chelp” como se muestra en la Figura 9.4

Figura 9.4 Comando “chelp”



En ambos casos se está pidiendo ayuda a Stata sobre el comando “describe”. La diferencia entre “help” y “chelp” radica en que en el primer caso la ayuda se despliega en una ventana nueva conocida como “viewer”, mientras que en el segundo la ayuda aparece en el cuadro de resultados.

Si no puede recordar el nombre completo del comando, o si no está seguro de qué comando desea, puede usar el comando “lookup” o “search”, seguido del nombre del comando o la palabra clave.

```
. search memoria
```

```
. lookup memoria
```

Este comando enumerará todos los comandos asociados con esta palabra clave y mostrará una breve descripción de cada uno de estos. Luego, puede elegir el que considere relevante y usar la ayuda para obtener la referencia específica. El sitio web de Stata (<http://www.stata.com>) tiene excelentes servicios de ayuda, como un tutorial en línea y preguntas frecuentes (FAQ).

Trabajando con la base de datos

Para los siguientes ejercicios se utiliza la base de datos “Data Base Impact Estado de Puebla.dta”.

Lista de variables

Para ver todas las variables en el conjunto de datos, use el comando “describe” (en su totalidad o abreviado):

```
. describe
```



Parte 2

Este comando proporciona información sobre el conjunto de datos, tales como nombre, tamaño, número de observaciones y enumera todas las variables con nombre, formato de almacenamiento, formato de visualización y etiqueta). Para ver sólo una variable o una lista de variables, use el comando describe seguido del nombre o nombres de la variable:

```
. describe SEXO_01 EDAD_01 P2_11_01 P2_12_01 P6_3 P6_4 P6_10A P6_10B P6_10C P6_10D P6_10E
+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+
variable name    storage    display    value    variable label
type            format    label
-----
SEXO_01          double    %10.0g    SEXO_01    P2.4. ¿Es hombre o mujer? (Integrante 1)
EDAD_01          double    %10.0g    EDAD_01    P2.5. ¿Cuántos años cumplidos tiene?
P2_11_01         double    %10.0g    P2_11_01   P2.11. ¿Cuántas horas trabajó en promedio
P2_12_01         double    %10.0g    P2_12_01   P2.12. ¿Cuánto ganó el mes pasado por el
P6_3             double    %10.0g    P6_3       P6.3. ¿Cuántas horas trabaja en promedio a
P6_4             double    %10.0g    P6_4       P6.4. ¿Recibe prestaciones de servicios
P6_10A           double    %10.0g    P6_10A     P6.10.a. ¿Recibe usted ingresos de
P6_10B           double    %10.0g    P6_10B     P6.10.b. ¿Recibe usted ingresos del
P6_10C           double    %10.0g    P6_10C     P6.10.c. ¿Recibe usted ingresos del
P6_10D           double    %10.0g    P6_10D     P6.10.d. ¿Ha recibido usted o algún
P6_10E           double    %10.0g    P6_10E     P6.10.e. ¿Ha recibido usted ingresos de
P7_6             double    %10.0g    P7_6       P7.6. ¿Cuál es el total de ingresos que
P10_5B           double    %10.0g    P10_5B     P10.5.b. ¿Qué tanta confianza le tiene
P10_5C           double    %10.0g    P10_5C     P10.5.c. ¿Qué tanta confianza le tiene
P10_5E           double    %10.0g    P10_5E     P10.5.e. ¿Qué tanta confianza le tiene
P10_5H           double    %10.0g    P10_5H     P10.5.h. ¿Qué tanta confianza le tiene
P6_10A           double    %10.0g    P6_10A     P6.10.a. ¿Recibe usted ingresos de
P6_10B           double    %10.0g    P6_10B     P6.10.b. ¿Recibe usted ingresos del
P6_10C           double    %10.0g    P6_10C     P6.10.c. ¿Recibe usted ingresos del
P6_10D           double    %10.0g    P6_10D     P6.10.d. ¿Ha recibido usted o algún
P6_10E           double    %10.0g    P6_10E     P6.10.e. ¿Ha recibido usted ingresos de
P7_6             double    %10.0g    P7_6       P7.6. ¿Cuál es el total de ingresos que
P10_5B           double    %10.0g    P10_5B     P10.5.b. ¿Qué tanta confianza le tiene
P10_5C           double    %10.0g    P10_5C     P10.5.c. ¿Qué tanta confianza le tiene
P10_5E           double    %10.0g    P10_5E     P10.5.e. ¿Qué tanta confianza le tiene
P10_5H           double    %10.0g    P10_5H     P10.5.h. ¿Qué tanta confianza le tiene
```

Como puede ver, el comando “describe” también muestra el tipo y la longitud de la variable, así como una breve descripción de esta (si está disponible). Deben tenerse en cuenta los siguientes puntos:

Puede abreviar una lista de variables escribiendo sólo el nombre de la primera y la última, separados por un guión (-). Por ejemplo, para ver todas las variables desde "P2_1" hasta "P2_4", puede escribir

```
. describe P2_1-P2_4
```

El símbolo de comodín (*) es útil para ahorrar algo de escritura. Para ver todas las variables que comienzan con "P2", puede escribir

```
. describe P2*
```

Listado de datos

Para ver los datos reales almacenados en las variables, use el comando "list" o "l". Si escribe el comando "list" solo, Stata mostrará valores para todas las variables y observaciones. Sin embargo, puede requerir sólo los datos de ciertas variables y observaciones, para ello escribe el comando "list" con una lista de variables y con ciertas condiciones.

El siguiente comando enumera todas las variables de las primeras tres observaciones:

```
. list in 1/3
```

Aquí Stata muestra todas las observaciones comenzando con la observación 1 y terminando con la observación 3. También puede mostrar datos como una hoja de cálculo. Para ello, utilice los dos iconos de la barra de herramientas denominados Editor de datos y Navegador de datos (véase las figuras 9.5 a 9.8). Al hacer clic en uno, aparecerá una nueva ventana emergente donde los datos se mostrarán como una tabla, con las observaciones en filas y las variables en columnas. El Explorador de datos sólo mostrará los datos, mientras que puede editarlos con el Editor de datos. Los comandos "edit" y "browse" también abrirán las ventanas correspondientes al editor y navegador de datos.

Figura 9.5 Editor de datos



Figura 9.6 Ventana Editor de datos

	P2_1	P2_2	P2_3	P2_4	SECC_05	SECC_06	P3_11_05	P3_12_05	P3_13_05	P3_15	P3_1	P3_2
1	2	18	.	.	Indice	39	40	4000	En casado	30	No	
2	4	18	.	.	Indice	40	51	4500	En casado	52	No	
3	2	18	.	.	Indice	40	30	3000	Vive en U.	40	No	
4	4	18	.	.	Indice	55	40	4000	En casado	.	.	
5	4	18	.	.	Indice	55	50	2500	Vive en U.	50	No	
6	4	18	.	.	Indice	40	45	3000	En casado	44	No	
7	2	18	.	.	Indice	50	40	2100	En casado	42	No	
8	2	18	.	.	Indice	42	24	3000	En casado	34	No	
9	2	18	.	.	Indice	30	40	3000	En casado	40	No	
10	4	18	2	2	Indice	53	30	2500	En casado	30	No	
11	4	18	.	.	Indice	40	50	15000	En casado	60	No	
12	2	18	.	.	Indice	50	50	7500	En casado	40	No	
13	4	18	.	.	Indice	25	40	4000	Vive en U.	40	No	
14	2	18	.	.	Indice	50	30	2500	En casado	30	No	
15	2	18	.	.	Indice	33	60	1000	Vive en U.	50	No	
16	4	18	.	.	Indice	55	8	000	Vive en U.	45	No	
17	2	18	.	.	Indice	40	40	4000	En casado	40	No	
18	2	18	.	.	Indice	50	50	2000	En casado	50	No	
19	4	18	.	.	Indice	40	8	0	Vive en U.	25	No	
20	3	18	.	.	Indice	53	50	3000	En casado	30	No	
21	4	18	.	.	Indice	50	40	100000	En casado	54	No	
22	4	18	.	.	Indice	43	70	2500	En casado	70	No	
23	2	18	.	.	Indice	50	50	4000	En casado	50	No	
24	4	18	.	.	Indice	33	50	3000	Vive en U.	60	No	



Figura 9.7 Navegador de datos

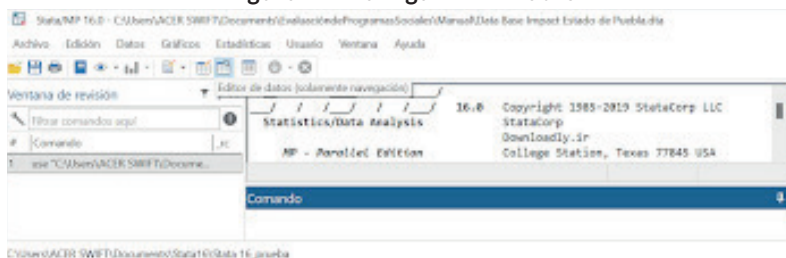
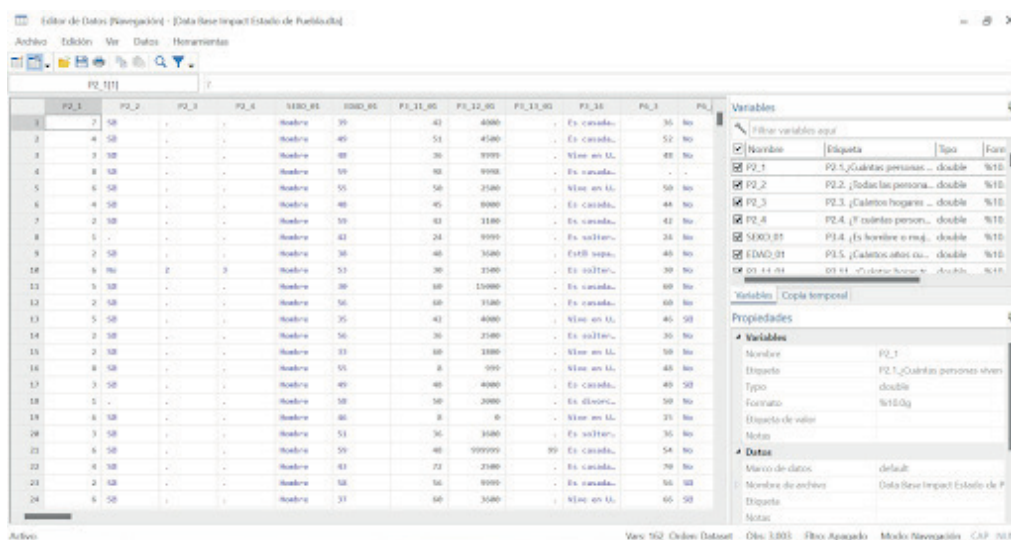


Figura 9.8 Navegador de datos



El siguiente comando enumera el tamaño del hogar y la educación del jefe de familia para hogares encabezados por una mujer menor de 45 años:

```
. list P2_1 educ_padre if (SEX0_01==1 & EDAD_01<45)
```

La declaración anterior utiliza dos operadores relacionales (== y <) y un operador lógico (&). Los operadores relacionales imponen una condición sobre una variable, mientras que los operadores lógicos combinan dos o más operadores relacionales. La tabla 9.1 muestra los operadores relacionales y lógicos utilizados en Stata.

Tabla 9.1 Operadores relacionales y lógicos utilizados en Stata

Operadores relacionales	Operadores lógicos
> (mayor que)	~ (no)
< (menor que)	(o)
== (igual)	& (y)
>= (mayor o igual)	
<= (menor o igual)	
!= o ~= (diferente)	

Resumen de datos

El comando “summarize” es muy útil (puede abreviarse como “sum”), pues calcula y muestra algunas estadísticas de resumen, incluidas las medias y las desviaciones estándar. Si no se especifica ninguna variable, se calcula el resumen de estadísticas para todas las variables del conjunto de datos. El siguiente comando resume el tamaño del hogar y el ingreso total que recibe el hogar en un mes normal:

```
. sum P2_4 ingreso
```

Variable	Obs	Mean	Std. Dev.	Min	Max
P2_4	124	3.580645	1.73953	1	10
ingreso	2,331	4723.092	5306.748	0	90000

Stata excluye cualquier observación que tenga un valor faltante para las variables que se resumen en este cálculo. Si desea conocer la mediana y los percentiles de una variable, agregue la opción “detail” (abreviado “d”):

```
. sum P2_4 ingreso, d
```

P2.4. ¿Y cuántas personas viven normalmente en el hogar donde usted reside?

Percentiles		Smallest		
1%	1	1		
5%	1	1		
10%	2	1	Obs	124
25%	2	1	Sum of Wgt.	124
50%		3	Mean	3.580645
			Std. Dev.	1.73953
			Largest	8
75%	4	8		
90%	5	9	Variance	3.025964
95%	8	9	Skewness	1.264509
99%	9	10	Kurtosis	5.142669

P7.6. ¿Cuál es el total de ingresos que recibe este hogar en un mes normal?

Percentiles		Smallest		
1%	500	0		
5%	999	0		
10%	1000	0	Obs	2,331
25%	2000	0	Sum of Wgt.	2,331
50%		3500	Mean	4723.092
			Std. Dev.	5306.748
			Largest	62000
75%	6000	62000		
90%	9999	80000	Variance	2.82e+07
95%	10000	80000	Skewness	6.954623
99%	22000	90000	Kurtosis	85.42544



Parte 2

Una gran ventaja de Stata es que permite el uso de ponderaciones. La opción de ponderación es útil si la probabilidad de muestreo de una observación es diferente de la de otra. En la mayoría de las encuestas de hogares el marco de muestreo es estratificado, donde se muestrean las primeras unidades de muestreo primarias, y condicionado a la selección de la unidad de muestreo primaria, se extraen las unidades de muestreo secundarias (a menudo hogares). Las encuestas de hogares generalmente proporcionan ponderaciones para corregir las diferencias en el diseño del muestreo y, a veces, los problemas de recopilación de datos. La implementación en Stata es sencilla:

```
. summarize P2_4 educ [aweight = SEXO_01]
```

Variable	Obs	Weight	Mean	Std. Dev.	Min	Max
P2_4	86	86	3.697674	1.946761	1	10
educ	1,764	1764	2.22449	1.247823	0	5

```
. summarize clases educ [aweight = ing_mens ]
```

Variable	Obs	Weight	Mean	Std. Dev.	Min	Max
clases	1,383	5341661	3.059615	1.478113	1	6
educ	1,472	5725170	2.866102	1.322553	0	5

Una alternativa útil al comando "summary" es el comando "tabstat", que le permite especificar la lista de estadísticas que desea mostrar en una sola tabla. Puede estar condicionado por otra variable. El siguiente comando muestra la media y la desviación estándar del tamaño de la familia y la educación del entrevistado del hogar por la variable sexo:

```
. tabstat P2_4 educ, statistics(mean sd) by(SEXO_01)
```

Summary statistics: mean, sd

by categories of: SEXO_01 (P3.4. ¿Es hombre o mujer? (Integrante 1))

SEXO_01	P2_4	educ
Hombre	3.315789 1.117557	2.557536 1.255132
Mujer	3.697674 1.946761	2.22449 1.247823
Total	3.580645 1.73953	2.361574 1.261326

Distribución de frecuencia (tabulaciones)

A menudo se necesitan distribuciones de frecuencia y tabulaciones cruzadas. El comando "tabulate" (abreviado "tab") se usa para hacer esto:

```
. tab educ_padre
```

Educación del Padre del Encuestado	Freq.	Percent	Cum.
Sin estudios	1,289	51.13	51.13
Primaria incompleta	606	24.04	75.17
Primaria	417	16.54	91.71
Secundaria	120	4.76	96.47
Preparatoria	64	2.54	99.01
Profesional	25	0.99	100.00
Total	2,521	100.00	

El siguiente comando da la distribución de la ocupación principal del entrevistado si es mujer:

```
. tab P6_6 if SEX0_01==1
```

P6.6. ¿En su ocupación principal es...?	Freq.	Percent	Cum.
Patrón o empleador	62	9.75	9.75
Trabajador por cuenta propia	359	56.45	66.19
Empleado u obrero del sector público (g	39	6.13	72.33
Empleado u obrero de empresas públicas	19	2.99	75.31
Empleado u obrero del sector privado	99	15.57	90.88
Servicio doméstico	46	7.23	98.11
Quehaceres del hogar (Ama de casa en el	3	0.47	98.58
Familiar sin pago	5	0.79	99.37
Ns (esp.)	3	0.47	99.84
Nc (esp.)	1	0.16	100.00
Total	636	100.00	

El comando "tabulate" también se puede utilizar para mostrar una distribución bidireccional. Por ejemplo, uno podría querer verificar si existe algún sesgo de género en la educación de los encuestados.



```
. tabulate educ SEX0_01
```

Educación del Encuestado	P3.4. ¿Es hombre o mujer? (Integrante 1)		Total
	Hombre	Mujer	
Sin estudios	52	140	192
Primaria incompleta	199	385	584
Primaria	359	527	886
Secundaria	354	428	782
Preparatoria	173	219	392
Profesional	97	65	162
Total	1,234	1,764	2,998

Tabla 9.2 Las opciones más frecuentes dentro de contents son:

n (número de observaciones)
mean (media)
sd (desviación típica)
median (mediana)
max (máximo)
min (mínimo)
p1 (primer percentil)
p2 (segundo percentil)
...
p98 (percentil 98)
p99 (percentil 99)
iqr (rango intercuartílico)

Para ver porcentajes por fila o columna, agregue opciones al comando "tabulate", puede considerar realizar tablas cruzadas:

```
. tabulate educ SEX0_01, row col
```

educ	SEX0_01
Sin estudios	52 140
Primaria incompleta	199 385
Primaria	359 527
Secundaria	354 428
Preparatoria	173 219
Profesional	97 65

Puede calcular el chi-cuadrado con el siguiente comando:

Educación del Encuestado	Sexo (Integrante 1)		Total
	Hombre	Mujer	
Sin estudios	52 27.08 4.21	140 72.92 7.94	192 100.00 6.40
Primaria incompleta	199 34.08 16.13	385 65.92 21.83	584 100.00 19.48
Primaria	359 40.52 29.09	527 59.48 29.88	886 100.00 29.55
Secundaria	354 45.27 28.69	428 54.73 24.26	782 100.00 26.08
Preparatoria	173 44.13 14.02	219 55.87 12.41	392 100.00 13.08
Profesional	97 59.88 7.86	65 40.12 3.68	162 100.00 5.40
Total	1,234 41.16 100.00	1,764 58.84 100.00	2,998 100.00 100.00

Puede calcular el chi-cuadrado con el siguiente comando:

```
. tabulate educ SEX0_01, row chi2
```

Key
frequency
row percentage



Educación del Encuestado	P3.4. ¿Es hombre o mujer? (Integrante 1)		Total
	Hombre	Mujer	
Sin estudios	52 27.08	140 72.92	192 100.00
Primaria incompleta	199 34.08	385 65.92	584 100.00
Primaria	359 40.52	527 59.48	886 100.00
Secundaria	354 45.27	428 54.73	782 100.00
Preparatoria	173 44.13	219 55.87	392 100.00
Profesional	97 59.88	65 40.12	162 100.00
Total	1,234 41.16	1,764 58.84	2,998 100.00

Pearson chi2(5) = 58.2756 Pr = 0.000

Dado que en algunos casos se pueden tener celdas de clasificación cruzada con pequeños recuentos, es posible que queramos utilizar la prueba exacta de Fisher en lugar de la prueba de chi-cuadrado. El comando necesario es el siguiente (empleando sólo 2003 observaciones):

```
. tabulate educ SEXO_01, exact nofreq

Enumerating sample-space combinations:
stage 6: enumerations = 1
stage 5: enumerations = 23
stage 4: enumerations = 579
stage 3: enumerations = 15138
stage 2: enumerations = 416235
stage 1: enumerations = 0

Fisher's exact = 0.004
```

Puede obtener la tabla correspondiente tanto para la prueba de chi-cuadrado como las pruebas exactas de Fisher utilizando:

```
. tabulate educ SEXO_01, row chi2 exact
```

Key
frequency
row percentage

```
Enumerating sample-space combinations:
stage 6: enumerations = 1
stage 5: enumerations = 23
stage 4: enumerations = 579
stage 3: enumerations = 15138
stage 2: enumerations = 416235
stage 1: enumerations = 0
```

Educación del Encuestado	P3.4. ¿Es hombre o mujer? (Integrante 1)		Total
	Hombre	Mujer	
Sin estudios	12 7.89	140 92.11	152 100.00
Primaria incompleta	40 9.41	385 90.59	425 100.00
Primaria	72 12.02	527 87.98	599 100.00
Secundaria	58 11.93	428 88.07	486 100.00
Preparatoria	30 12.05	219 87.95	249 100.00
Profesional	22 25.29	65 74.71	87 100.00
Total	234 11.71	1,764 88.29	1,998 100.00

```
Pearson chi2(5) = 19.9281 Pr = 0.001
Fisher's exact = 0.004
```

Distribuciones de estadísticas descriptivas (comando "table")

Otro comando muy conveniente es "table", que combina características de los comandos "sum" y "tab". Además, muestra los resultados de una forma más presentable. El siguiente comando de "table" muestra la media del tamaño de la familia y de la educación del entrevistado del hogar, por su participación en programas de oportunidades:

```
. table P6_10A, c(media P2_4 media educ)
```

P6.10.a. ¿Recibe usted ingresos de Prospera (Programa Oportunidades)?		
	med(P2_4)	med(educ)
Si	4	2
No	3	2

Como se observa, los resultados se presentan como números enteros y no como fracción. Esto se debe a que las variables "P2_4" y "educ" se almacenan como números enteros y Stata simplemente trunca los números después del decimal. Observe la descripción de las variables:



```
. d P2_4
```

variable name	storage type	display format	value label	variable label
P2_4	double	%10.0g	P2_4	P2.4. ¿Y cuántas personas viven normalmente en el hogar donde usted reside?

```
. d educ
```

variable name	storage type	display format	value label	variable label
educ	float	%19.0g	educ_padre	Educación del Encuestado

El comando "table" puede mostrar hasta cinco estadísticas y variables distintas de la media (como la suma, el mínimo o el máximo). Se pueden mostrar tablas bidireccionales, tridireccionales o incluso de mayor dimensión. Aquí hay un ejemplo de una tabla de doble entrada que desglosa las siguientes variables.

```
. table P6_10A SEXO_01 , c(media P2_4 media educ)
```

P6.10.a. ¿Recibe usted ingresos de Prospera (Programa Oportunidades)?	P3.4. ¿Eshombre (Integrante 1)	
	Hombre	Mujer
Si	4 2	4 2
No	3 2.5	4 2

Puede comparar estadísticas, tales como el promedio y la desviación estándar:

```
. table P11_3, contents(mean ingtot_mens sd ingtot_mens)
```

P11.3. ¿Qué tan satisfecho se encuentra usted con su calidad de vida?	mean(ingtot~s)	sd(ingtot~s)
Muy satisfecho	9957.2414	20314.92
Satisfecho	10901.079	24119.47
Ni satisfecho, ni insatisfecho (esp.)	16661.725	60391.2
Insatisfecho	8437.9106	20962.57
muy insatisfecho	3369.6429	2071.282
Ns	19433.579	36024.97

Valores faltantes en Stata

En Stata, un valor faltante se representa con un punto (.). Un valor faltante se considera mayor que cualquier número. El comando "summarize" ignora las

observaciones con valores faltantes, y el comando "tabulate" hace lo mismo, a menos que se le obligue a incluir valores faltantes.

Aquí se muestra los valores faltantes, aquellos encuestados que no tienen respuesta respecto a si reciben ingresos del programa prospera.

```
. tab P6_10A, m
```

P6.10.a. ¿Recibe usted ingresos de Prospera (Programa Oportunidad es)?	Freq.	Percent	Cum.
Si	711	23.68	23.68
No	78	2.60	26.27
.	2,214	73.73	100.00
Total	3,003	100.00	

El siguiente comando ilustra cómo se manejan los valores faltantes en las sentencias de asignación.

```
. gen sum_1 = P6_10A + P6_10B  
(2,214 missing values generated)
```

2986.	3
2987.	.
2988.	.
2989.	.
2990.	.
2991.	2
2992.	3
2993.	.
2994.	.
2995.	.
2996.	3
2997.	3
2998.	3
2999.	.
3000.	2
3001.	.
3002.	.
3003.	.



Parte 2

Puede usar las funciones “*rowmiss*” y “*rownomiss*” para determinar la cantidad de valores que faltan y la cantidad de valores que no faltan, respectivamente, en una lista de variables. Esto se ilustra a continuación:

```
. egen miss = rowmiss( P6_10A - P6_10B )
. egen nomiss = rownonmiss( P6_10A - P6_10B )
```

El comando “*inspect*” también nos proporciona indica el número de valores positivos, negativos, cero y faltantes (missing values), así como el número de valores enteros y no enteros de las variables var1 y var2. Los “missing values” se señalan en Stata mediante un punto (.). Se considera que un “missing value” es mayor que cualquier valor. Por ejemplo:

```
. inspect P3_12_01
```

P3_12_01: P3.122 ¿Cuánto ganó el mes pa		Number of Observations		
		Total	Integers	Nonintegers
#	Negative	-	-	-
#	Zero	18	18	-
#	Positive	1,831	1,831	-
#				
#	Total	1,849	1,849	-
#	Missing	1,154		
0 999999		3,003		
(More than 99 unique values)				

Estos errores de introducción de datos también pueden detectarse mediante el comando “*codebook*”. Use las siguientes variables:

```
. codebook P6_10a P6_10c P6_10d P6_10e P6_10f
```

P6_10a		P6.10.a. ¿Recibe usted ingresos del Programa de Adultos mayores?	
type: numeric (double)			
label: P6_10a			
range: [1,2]		width: 1	
unique values: 3		missing : 1,214/3,003	
tabulation: Freq. Numeric Label			
	1	1	Yes
	2	2	No
	3	3	No
	2,214		

P6_10c		P6.10.c. ¿Recibe usted ingresos del Programa de empleo temporal?	
type: numeric (double)			
label: P6_10c			
range: [1,2]		width: 1	
unique values: 3		missing : 1,214/3,003	
tabulation: Freq. Numeric Label			
	1	1	Yes
	2	2	No
	3	3	No
	2,214		

P6_100	P6_100.d. ¿Ha recibido usted o algún integrante de este hogar, ayuda de Conedores?		
type: numeric (double)			
label: P6_100			
range: [1,3]		units: 1	
unique values: 3		missing : 2,214/3,003	
tabulation:			
Freq.	Numeric	Label	
26	1	SI	
763	2	No	
1	3	No	
2,214	.		

P6_105	P6_105.g. ¿Ha recibido usted ingresos de Ayuda por algún desastre natural?		
type: numeric (double)			
label: P6_105			
range: [1,2]		units: 1	
unique values: 2		missing : 2,214/3,003	
tabulation:			
Freq.	Numeric	Label	
1	1	SI	
763	2	No	
2,214	.		

P6_106	P6_106.i. ¿Recibe usted ingresos de Vales para compra de bienes de duraderos desp		
type: numeric (double)			
label: P6_106			
range: [1,3]		units: 1	
unique values: 3		missing : 2,214/3,003	
tabulation:			
Freq.	Numeric	Label	
4	1	SI	
764	2	No	
1	3	No	
2,214	.		

Alternativamente, podemos detectar errores utilizando el comando “assert”:

```

. assert P6_108==1|P6_108==2|P6_108==.
3 contradictions in 3,003 observations
assertion is false

```

Además, es posible reemplazar un valor por “missing”:

```

. replace P6_108=. if P6_108==2
(703 real changes made, 703 to missing)

```

Conteo de observaciones

El comando “count” se usa para contar el número de observaciones en el conjunto de datos:

```

. count
3,003

```



El comando "count" se puede usar con condiciones. El siguiente comando da el número de hogares donde el encuestado es mayor de 50 años:

```
count if EDAD_01>50
840
```

Uso de ponderaciones

En la mayoría de las encuestas de hogares, las observaciones se seleccionan mediante un proceso aleatorio y pueden tener diferentes probabilidades de selección. Por lo tanto, se deben utilizar ponderaciones que sean iguales a la inversa de la probabilidad de ser muestreado. Una ponderación de w_j para la j -ésima observación significa, en términos generales, que la j -ésima observación representa w_j elementos en la población de la que se extrajo la muestra. La omisión de las ponderaciones muestrales en el análisis suele generar estimaciones sesgadas, que pueden estar lejos de los valores reales.

Suelen ser necesarios varios ajustes posteriores al muestreo de las ponderaciones. Stata tiene cuatro tipos de ponderaciones:

- Las ponderaciones de frecuencia ("*fweight*") indican cuántas observaciones en la población están representadas por cada observación en la muestra. Deben tomar valores enteros.
- Las ponderaciones analíticas ("*aweight*") son especialmente apropiadas cuando se trabaja con datos que contienen promedios (por ejemplo, el ingreso per cápita promedio en un hogar). La variable de ponderación es proporcional al número de personas sobre las que se calculó el promedio (por ejemplo, número de miembros de un hogar). Técnicamente, las ponderaciones analíticas están en proporción inversa a la varianza de una observación (es decir, una ponderación más alta significa que la observación se basó en más información y, por lo tanto, es más confiable en el sentido de que tiene menos varianza).
- Las ponderaciones de muestreo ("*pweight*") son el inverso de la probabilidad de selección debido al diseño de la muestra.
- Las ponderaciones de importancia ("*iweight*") indican la importancia relativa de la observación. Los más utilizados son "*pweight*" y "*aweight*". Se puede obtener más información sobre las ponderaciones escribiendo "*help weight*".

Cambio de conjuntos de datos

Trabajar con bases de datos en Stata a menudo implica realizar cambios en los datos, tales como crear nuevas variables o cambiar valores de variables existentes. Los siguientes ejercicios demuestran cómo se pueden incorporar esos cambios.

Generación de nuevas variables

En Stata, el comando "generate" (abreviado "gen") crea nuevas variables, mientras que el comando "replace" cambia los valores de una variable existente. Los siguientes comandos crean una nueva variable llamada "anciano" y luego establecen su valor en uno si el cabeza de familia tiene más de 50 años, y en cero de lo contrario:

```
. gen oldhead=1 if EDAD_01>50
(2,163 missing values generated)

. replace oldhead=0 if EDAD_01<=50
(2,163 real changes made)
```

Lo que ocurre aquí es que, para cada observación, el comando "gen" comprueba la condición (si el encuestado tiene más de 50 años) y establece el valor de la variable "oldhead" en uno para esa observación si la condición es verdadera y un valor de cero en caso contrario. El comando "replace" funciona de forma similar. Después del comando "generate", Stata indica que 2 163 observaciones no cumplieron la condición, y después del comando "replace" Stata indica que esas 2 163 observaciones tienen nuevos valores (cero en este caso). Cabe destacar los siguientes puntos:

- Si se emite un comando "gen" o "replace" sin ninguna condición, ese comando se aplica a todas las observaciones en el archivo de datos.
- Al usar el comando generar, uno debe tener cuidado de manejar los valores faltantes correctamente.
- El lado derecho del signo = en los comandos "gen" o "replace" puede ser cualquier expresión que involucre nombres de variables, no sólo un valor. Así, por ejemplo, el comando "gen young = (EDAD_01<=32)" crearía una variable llamada "young" que tomaría el valor de uno si tiene 32 años o menos (es decir, si la expresión entre paréntesis es verdadera), y un valor de cero en caso contrario.



Parte 2

- El comando *“replace”* se puede utilizar para cambiar los valores de cualquier variable existente, independientemente del comando *“generate”*.

Una extensión del comando "generate" es "egen". Al igual que el comando "gen", el comando "egen" puede crear variables para almacenar estadísticas descriptivas, como la media, la suma, el máximo y el mínimo. La principal característica del comando "egen" es su capacidad para crear estadísticas que involucren múltiples observaciones. Por ejemplo, el siguiente comando crea una variable "promedio" que contiene la edad promedio de los encuestados de cada hogar para todos los datos:

```
. egen avgage=mean(EDAD_01)
```

Todas las observaciones en el conjunto de datos obtienen el mismo valor "promedio". El siguiente comando crea las mismas estadísticas, pero esta vez para hogares encabezados por hombres y mujeres por separado:

```
. egen avgagemf=mean(EDAD_01), by(SEX0_01)
```

Etiquetado de variables

Puede poner etiquetas a las variables para darles una descripción. Por ejemplo, la variable "oldhead" no tiene ninguna etiqueta ahora. Puede asignar una etiqueta a esta variable:

```
. label variable oldhead "Mayor de 50: 1=SI, 0=NO"
```

En el comando "label", la variable se puede acortar a "var". Ahora para ver la nueva etiqueta, escribe lo siguiente:

```
. des oldhead
```

variable name	storage type	display format	value label	variable label
oldhead	float	%9.0g		Mayor de 50: 1=SI, 0=NO

Etiquetado de datos

Se pueden crear otros tipos de etiquetas. Para adjuntar una etiqueta a todo el conjunto de datos que aparece en la parte superior de la lista "describe", escribe lo siguiente:

```
. label data "Encuesta de Movilidad Social 2015"  
. des
```

Etiquetado de valores de variables

Las variables que son categóricas, como la de “oldhead (1 = Si, 0 = No)”, pueden tener etiquetas que ayuden a recordar cuáles son las categorías. Al tabular la variable “oldhead” sólo se muestran los valores cero y uno:

```
. tab oldhead
```

Mayor de 50: 1=SI, 0=NO	Freq.	Percent	Cum.
0	2,163	72.03	72.03
1	840	27.97	100.00
Total	3,003	100.00	

Para adjuntar etiquetas a los valores de una variable, se deben hacer dos cosas. Primero, defina una etiqueta de valor. Luego asigne esta etiqueta a la variable. Usando las nuevas categorías para “oldhead”, escriba:

```
. label define etiqueta_oldhead 0 "no" 1 "si"
```

```
. label values oldhead etiqueta_oldhead
```

```
. tab oldhead
```

Mayor de 50: 1=SI, 0=NO	Freq.	Percent	Cum.
no	2,163	72.03	72.03
si	840	27.97	100.00
Total	3,003	100.00	

Si desea ver los valores reales de la variable “oldhead”, que siguen siendo ceros y unos, puede agregar una opción para no mostrar las etiquetas asignadas a los valores de la variable.

```
. tab oldhead, nolabel
```

Mayor de 50: 1=SI, 0=NO	Freq.	Percent	Cum.
0	2,163	72.03	72.03
1	840	27.97	100.00
Total	3,003	100.00	



Parte 2

Otra alternativa para asignar etiquetas del mismo tipo es:

```
. label define agua_entubada 0 No 1 Si  
. label values electricidad agua_entubada  
  
. label values banio agua_entubada  
  
. label values refrigerador agua_entubada
```

Los tres últimos comandos también podrían haberse realizado con el siguiente comando:

```
. foreach x in electricidad banio refrigerador {  
  2. label values 'x' agua_entubada  
  3. }
```

Mantener y eliminar variables y observaciones

Las variables y observaciones de un conjunto de datos se pueden seleccionar usando los comandos "keep" o "drop". Suponga que tiene un conjunto de datos con seis variables: var1, var2, ..., var6. Le gustaría mantener un archivo con sólo tres de ellos (por ejemplo, var1, var2 y var3). Puede utilizar cualquiera de los siguientes dos comandos:

- “keep var1 var2 var3” (o “keep var1-var3” si las variables están en este orden)
- “drop var4 var5 var6” (o “drop var4-var6” si las variables están en este orden)

Tenga en cuenta el uso de un guión (-) en ambos comandos. Es una buena práctica usar el comando que involucra menos variables o menos tipos, y por lo tanto, menos riesgo de error). También puede utilizar operadores relacionales o lógicos. Por ejemplo, el siguiente comando descarta o elimina aquellas observaciones en las que el encuestado del hogar tiene 58 años o más:

```
drop if EDAD_01>=58  
(27 observations deleted)
```

Y este comando mantiene aquellas observaciones donde el tamaño del hogar es de seis miembros o menos:

```
keep if P2_4<=6  
(672 observations deleted)
```

Los dos comandos anteriores eliminan o mantienen todas las variables según las condiciones. No puede incluir una lista de variables en un comando de "drop" o "keep". Por ejemplo, el siguiente comando fallará:

```
. keep educ P2_4 if P2_4<=6
invalid syntax
r(198);
```

Tienes que usar dos comandos para hacer el trabajo:

```
. keep if P2_4<=6
. keep educ P2_4
```

También puede usar la palabra clave en un comando "drop" o "keep". Por ejemplo, para descartar las primeras veinte observaciones:

```
. drop in 1/20
(20 observations deleted)
```

Otras estadísticas

- `sdtest`

También podemos comprobar si las varianzas difieren significativamente utilizando (P6_10: ¿Recibe ingresos de Prospera?)

```
. sdtest ingtot_mens, by(P6_10A)
```

Variance ratio test

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
Si	595	11425.39	1083.53	26430.14	9297.374	13553.41
No	70	7520.286	2005.342	16777.9	3519.738	11520.83
combined	665	11014.33	992.9073	25604.69	9064.71	12963.94

```
ratio = sd(Si) / sd(No)                                f = 2.4816
Ho: ratio = 1                                           degrees of freedom = 594, 69
```

Ha: ratio < 1
Pr(F < f) = 1.0000

Ha: ratio != 1
2*Pr(F > f) = 0.0000

Ha: ratio > 1
Pr(F > f) = 0.0000



Parte 2

- ttest

```
. ttest ingtot_mens, by(P6_10A)
```

Two-sample t test with equal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
Si	595	11425.39	1083.53	26430.14	9297.374	13553.41
No	70	7520.286	2005.342	16777.9	3519.738	11520.83
combined	665	11014.33	992.9073	25604.69	9064.71	12963.94
diff		3905.104	3234.247		-2445.496	10255.7

```
diff = mean(Si) - mean>No)
Ho: diff = 0
t = 1.2074
degrees of freedom = 663
```

```
Ha: diff < 0
Pr(T < t) = 0.8862
Ha: diff != 0
Pr(|T| > |t|) = 0.2277
Ha: diff > 0
Pr(T > t) = 0.1138
```

- corr

Correlaciones entre las variables

```
. corr ing_mens hrs_lab educ
(obs=1,490)
```

	ing_mens	hrs_lab	educ
ing_mens	1.0000		
hrs_lab	0.1556	1.0000	
educ	0.2966	0.0635	1.0000

- pwcorr:

Un enfoque alternativo es utilizar el comando pwcorr para incluir, en cada correlación, todas las observaciones que tienen datos completos para el correspondiente par de variables, lo que da lugar a diferentes tamaños de muestra para diferentes correlaciones.

```
. pwcorr ing_mens hrs_lab educ, obs sig
```

	ing_mens	hrs_lab	educ
ing_mens	1.0000		
	1492		
hrs_lab	0.1565	1.0000	
	0.0000		
	1492	1849	
educ	0.2966	0.0515	1.0000
	0.0000	0.0269	
	1490	1846	2998

- ktau

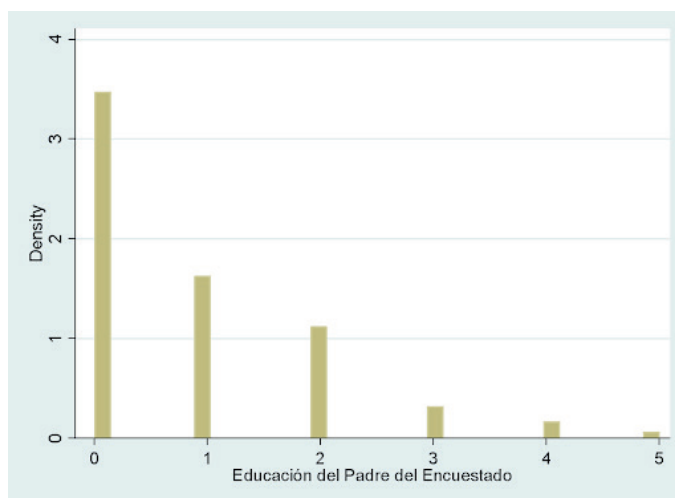
```
. ktau ing_mens hrs_lab educ
(obs=1498)
```

	ing_mens	hrs_lab	educ
ing_mens	0.9626		
hrs_lab	0.1840	0.9423	
educ	0.2343	0.0367	0.7886

Producción de gráficos

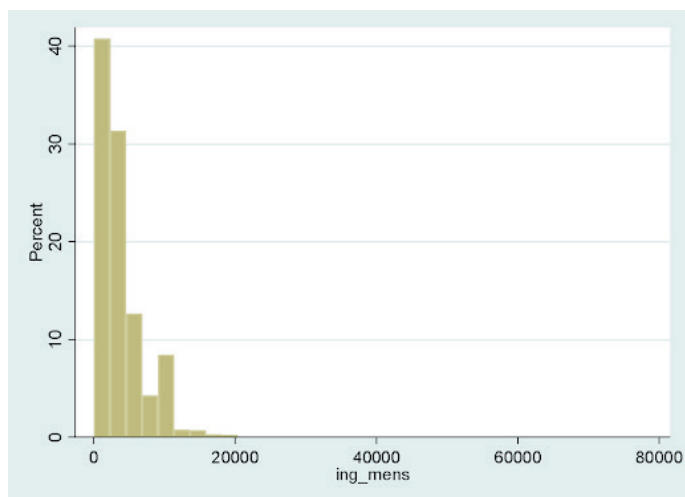
Stata es bastante bueno en la producción de gráficos básicos, aunque puede ser necesaria una experimentación considerable para producir magníficos gráficos. El siguiente comando muestra la distribución de la edad del jefe de hogar en un gráfico de barras (histograma):

```
. histogram educ_padre
(bin=34, start=0, width=.14705882)
```



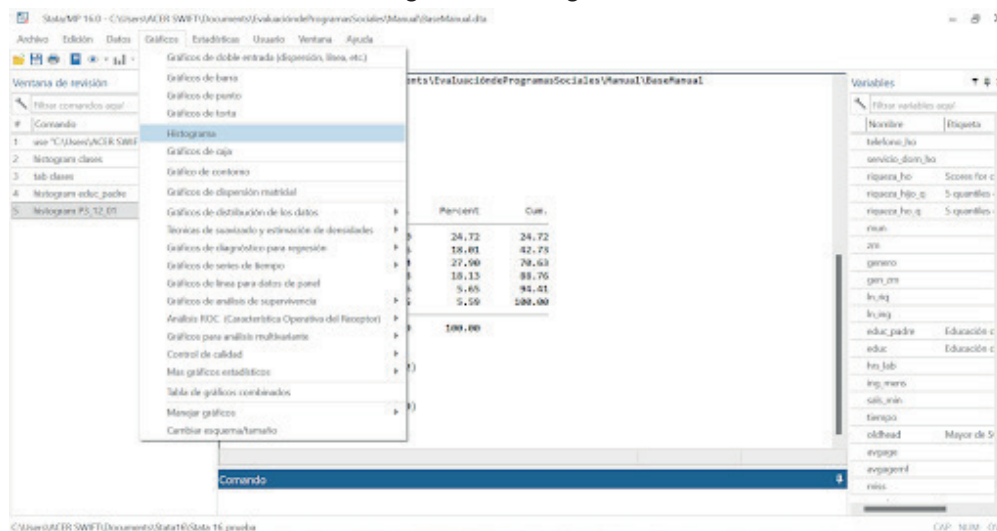


```
. histogram ing_mens, percent
(bin=31, start=0, width=2258.0645)
```



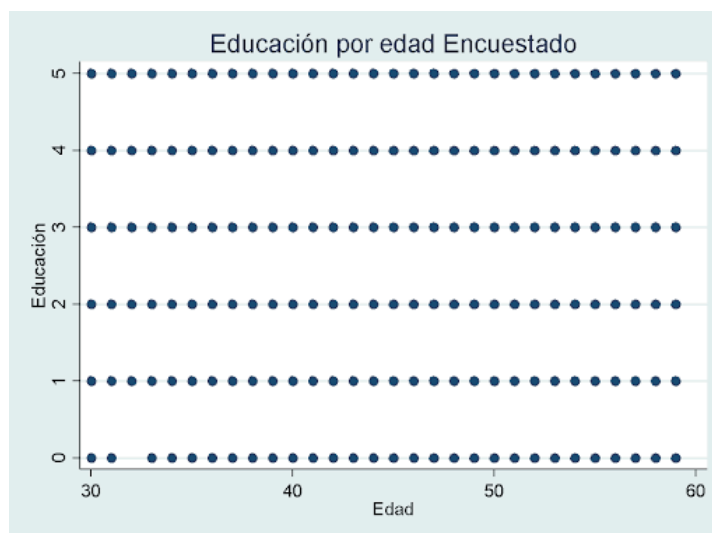
En muchos casos, la forma más fácil de producir gráficos es usando los menús. En este caso, haga clic en “Gráficos” y luego en “Histograma” y siga las indicaciones (véase Figura 9.5). Una manera fácil de guardar un gráfico es hacer clic derecho sobre él y copiarlo para pegarlo en un documento de Microsoft Word o Excel. Además, también es posible cambiar los colores y sombreado en la ventana del gráfico.

Figura 9.9 Histograma

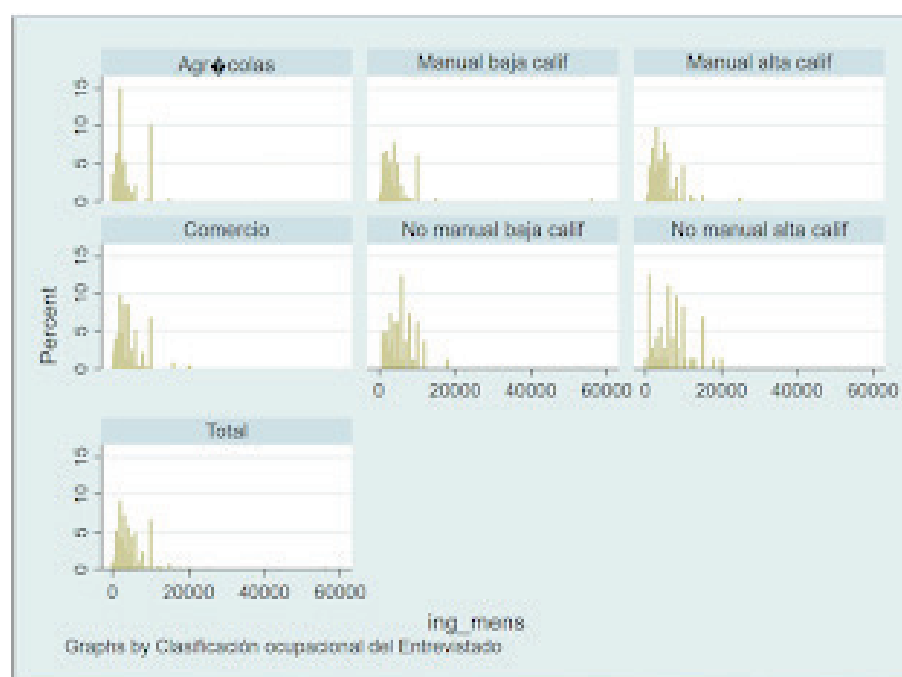


Para un diagrama de dispersión de dos variables, puede escribir:

```
. twoway (scatter educ EDAD_01 ), ytitle(Educación) xtitle(Edad) title(Educación por edad Encues
> tado)
```

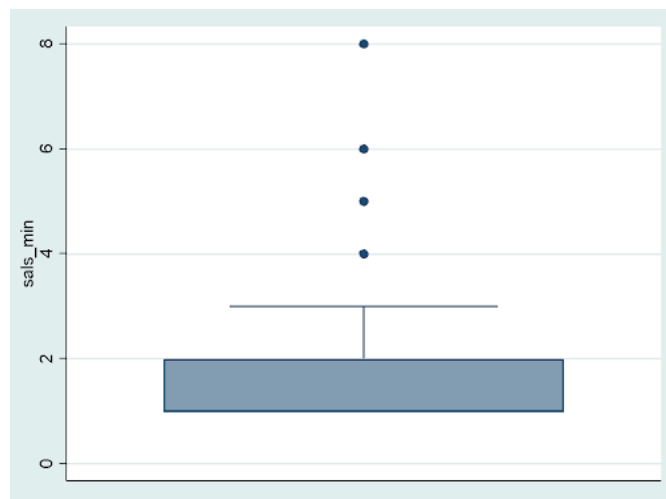


```
. histogram ing_mens, percent start(0) width(10) by(clases, total)
```

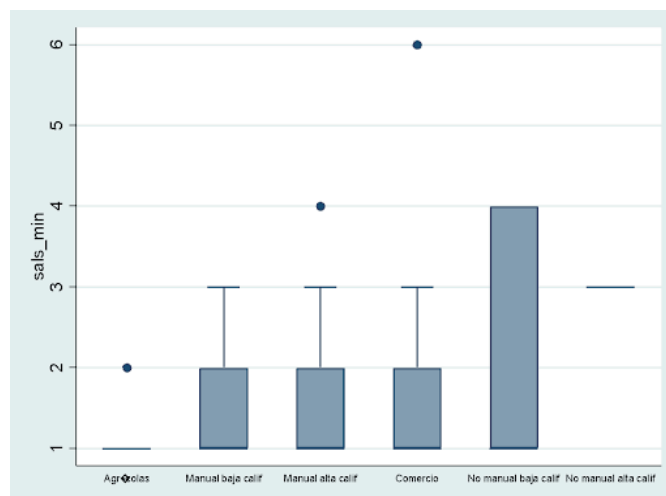


Box plots

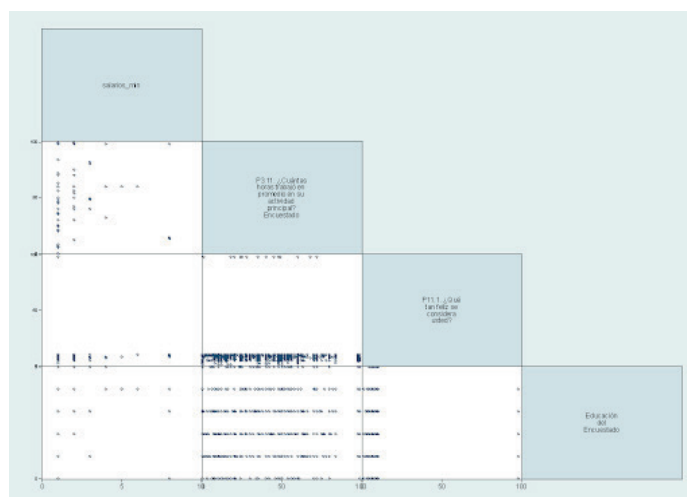
Los gráficos de caja transmiten información de un vistazo sobre el centro, la dispersión, la simetría y los valores atípicos. Por ejemplo, el salario mínimo:



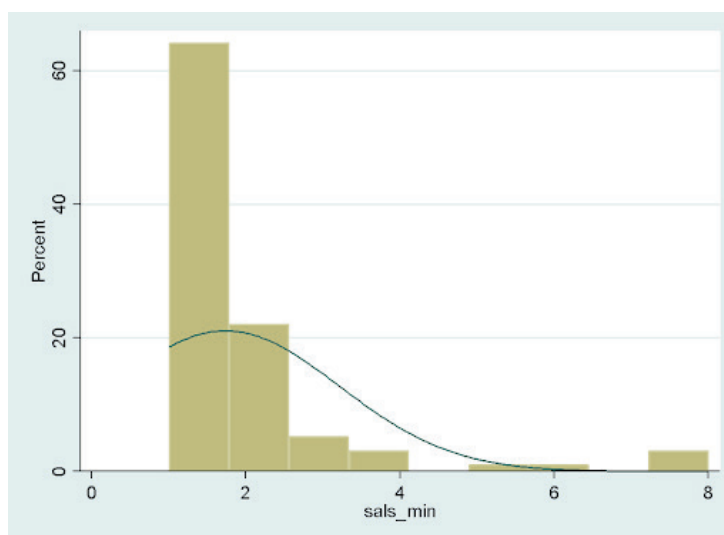
```
. graph box sals_min , over(clases)
```



```
. graph matrix sals_min P3_11_01 P11_1 educ , half msymbol(oh)
```



```
. histogram sals_min, norm percent
(bin=9, start=1, width=.77777778)
```



Combinación de conjuntos de datos

En Stata sólo se puede trabajar con un archivo de datos a la vez, es decir, únicamente puede haber un conjunto de datos en la memoria al mismo tiempo. Sin embargo, la información útil a menudo se distribuye en varios archivos de datos a los que es necesario acceder simultáneamente. Para usar dicha información, Stata tiene comandos que combinan esos archivos. Dependiendo de cómo se distribuya dicha información entre los archivos, se pueden fusionar o agregar varios archivos.

Combinar conjuntos de datos (merging)

La combinación de archivos de datos se realiza cuando se necesitan variables que se distribuyen en dos o más archivos. Como ejemplo de merging, la base original se dividirá en dos conjuntos de datos, de tal manera que uno contenga unas variables que el otro no, y luego los conjuntos de datos se combinarán (merged) para obtener la base original. Abra el archivo Data Base Impact Estado de Puebla_1.dta, elimine las variables de participación de los programas y guarde el archivo de datos como *Data Base Impact Estado de Puebla_01.dta*.

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad
> o de Puebla_1.dta", clear
. drop P6_10A P6_10B P6_10C P6_10D P6_10G
. save "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Est
> ado de Puebla_01.dta", replace
file C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad
> o de Puebla_01.dta saved
```



Debe dar a este archivo un nuevo nombre (*Data Base Impact Estado de Puebla_01.dta*), porque no se quiere cambiar la base original (*Data Base Impact Estado de Puebla_1.dta*) de forma permanente. Ahora abra la base original nuevamente. Esta vez, mantenga sólo las variables de participación de los programas. Guarde este archivo como *Data Base Impact Estado de Puebla_02.dta*

```
. keep EST_MUN_LOC_AGEb P6_10A P6_10B P6_10C P6_10D P6_10G
. save "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estado de Puebla_02.dta", replace
```

Observe que mantuvo la identificación del hogar *EST_MUN_LOC_AGEb* además de las variables de participación en los programas. Esto es necesario porque el merge requiere al menos una variable de identificación común entre los dos archivos que se van a fusionar. La variable “*EST_MUN_LOC_AGEb*” es común entre los dos archivos. Ahora tiene dos conjuntos de datos: uno tiene las variables de participación de los programas del hogar (*Data Base Impact Estado de Puebla_02.dta*) y el otro no tiene esas variables (*Data Base Impact Estado de Puebla_01.dta*). Si necesita usar las variables de ambos archivos, deberá unir los dos archivos. Sin embargo, antes de realizar el merge de los dos archivos, debe asegurarse de que ambos archivos están ordenados por la variable de identificación. Esto se puede hacer rápidamente de la siguiente manera:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estado de Puebla_01.dta", clear
. sort EST_MUN_LOC_AGEb
. save, replace
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estado de Puebla_02.dta", clear
. sort EST_MUN_LOC_AGEb
. save, replace
```

Ahora está listo para realizar el merge de los dos archivos. Uno de los archivos tiene que estar abierto (no importa cuál). Abra el archivo *Data Base Impact Estado de Puebla_01.dta* y luego use el comando merge al archivo *Data Base Impact Estado de Puebla_02.dta*:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estado de Puebla_01.dta", clear
. merge EST_MUN_LOC_AGEb using "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estado de Puebla_02.dta"
```

En este contexto, *Data Base Impact Estado de Puebla_01.dta* se denomina archivo maestro (el archivo que permanece en la memoria antes de la operación de fusión) y *Data Base Impact Estado de Puebla_02.dta* se denomina archivo esclavo. Para ver cómo fue la operación de merge, escriba el siguiente comando:

```
. tab _merge
```

_merge	Freq.	Percent	Cum.
3	3,003	100.00	100.00
Total	3,003	100.00	

Estos resultados se interpretan como:

- Un valor de 1 muestra el número de observaciones provenientes del archivo maestro únicamente
- Un valor de 2 muestra el número de observaciones provenientes del archivo esclavo solamente.
- Un valor de 3 muestra el número de observaciones comunes en ambos archivos

El número total de observaciones en el conjunto de datos resultante es la suma de estas tres frecuencias "merge". En este ejemplo, sin embargo, cada observación (hogar) en el archivo *Data Base Impact Estado de Puebla_01.dta* tiene una coincidencia exacta en el archivo *Data Base Impact Estado de Puebla_02.dta*, por lo que obtuvo "_merge=3" y no se obtuvo 1 ni 2. Pero en ejemplos de la vida real, los 1 y los 2 pueden permanecer después del merge. La mayoría de las veces uno quiere trabajar con las observaciones que son comunes en ambos archivos (es decir, "_merge=3"). Eso se hace emitiendo el siguiente comando después de la operación:

```
. keep if _merge==3
(0 observations deleted)
```

Agregar conjuntos de datos (appending)

Es necesario agregar conjuntos de datos cuando necesita combinar dos conjuntos de datos que tienen las mismas (o casi las mismas) variables, pero las unidades de observación (hogares, por ejemplo) son mutuamente excluyentes. Para demostrar la operación de agregar, volverá a dividir *Data Base Impact Estado de Puebla_1.dta*. Esta vez, en lugar de descartar variables, eliminará algunas observaciones. Abra el archivo *Data Base Impact Estado de Puebla_1.dta*, elimine las observaciones 1 a 700 y guarde este archivo



Parte 2

como Data Base Impact Estado de Puebla_01.dta:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad  
> o de Puebla_1.dta", clear  
  
. drop in 1/700  
(700 observations deleted)  
  
. save "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Esta  
> do de Puebla_01.dta", replace
```

A continuación, vuelva a abrir Data Base Impact Estado de Puebla_1.dta pero mantenga las observaciones 1 a 700, y guarde este archivo como Data Base Impact Estado de Puebla_02.dta.

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad  
> o de Puebla_1.dta", clear  
  
. keep in 1/700  
(2,303 observations deleted)  
  
. save "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Esta  
> do de Puebla_02.dta", replace
```

Ahora, tiene dos conjuntos de datos. Ambos tienen variables idénticas, pero conjuntos de hogares diferentes. En esta situación, necesita añadir dos archivos. Una vez más, uno de los archivos tiene que estar en la memoria (no importa cuál). Abra Data Base Impact Estado de Puebla_01.dta, y luego use el comando append Data Base Impact Estado de Puebla_02.dta.

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad  
> o de Puebla_01.dta", clear  
  
. append using "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Imp  
> act Estado de Puebla_02.dta"
```

Tenga en cuenta que no es necesario ordenar los archivos individuales para la operación append, y Stata no crea ninguna variable nueva como "_merge" después de la operación de agregar. Puede verificar que la operación de agregar se ejecutó correctamente emitiendo el comando "count" de Stata, que muestra el número de observaciones en el conjunto de datos resultante, que debe ser la suma de las observaciones en los dos archivos individuales.

```
. count  
3,003
```

Trabajar con archivos .log y .do

Esta sección analiza el uso de dos tipos de archivos que son extremadamente eficientes en las aplicaciones de Stata. Uno almacena comandos y resultados de Stata para revisión posterior (archivos .log), y el otro almacena comandos para ejecuciones repetidas. Los dos tipos de archivos pueden funcionar de

forma interactiva, lo que es muy útil para depurar comandos y obtener una buena "intuición" de los datos.

.log files

A menudo uno quiere guardar los resultados de los comandos de Stata y tal vez imprimirlos; para ellos es útil un archivo .log file. Dicho archivo se crea mediante la emisión de un comando de "log using" y se cierra con un comando de "log close". Todos los comandos emitidos en el medio, así como las salidas correspondientes (excepto los gráficos) se guardan en el archivo .log file. Utilice la base original. Suponga que desea guardar solo el resumen de educación del encuestado por sexo del hogar. Aquí están los comandos:

```
. log using educ.log
```

```

      name: <unnamed>
      log:  C:\Users\ACER SWIFT\Documents\Stata16\Stata 16_prueba\educ.log
  log type: text
opened on: 11 Jun 2022, 17:42:29

. by SEXO_01, sort:sum educ
```

```
-> SEXO_01 = Hombre
```

Variable	Obs	Mean	Std. Dev.	Min	Max
educ	1,234	2.557536	1.255132	0	5

```
-> SEXO_01 = Mujer
```

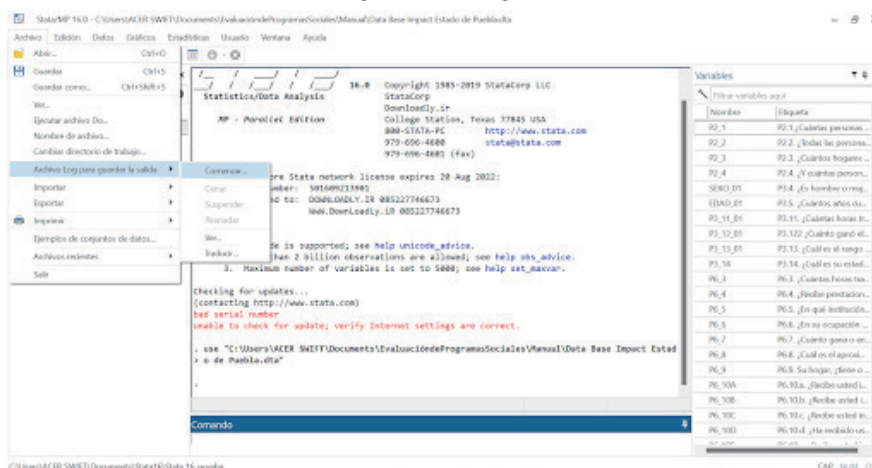
Variable	Obs	Mean	Std. Dev.	Min	Max
educ	1,764	2.22449	1.247823	0	5

```
. log close
      name: <unnamed>
      log:  C:\Users\ACER SWIFT\Documents\Stata16\Stata 16_prueba\educ.log
  log type: text
closed on: 11 Jun 2022, 17:44:45
```

Lo que sucede aquí es que Stata crea un archivo de texto llamado educ.log en la carpeta actual y guarda el resultado del resumen en este archivo. Si desea que el archivo .log file se guarde en una carpeta que no sea la carpeta actual, puede especificar la ruta completa de la carpeta en el comando de creación de .log. O bien puede guardar el archivo log desde el menú como se observa en la Figura 9.8:



Figura 9.10 .log file



Si ya existe un archivo .log, puede reemplazarlo con "log using educ.log" o reemplazarlo o agregarle una nueva salida con "log using educm.log, append". Si realmente desea mantener el archivo .log file existente sin cambios, puede cambiar el nombre de este archivo o el archivo en el comando de creación de .log. Si desea suprimir una parte de un archivo .log file, puede emitir un comando de "log off" antes de esa parte, seguido de un comando de "log on" para la parte que desea guardar. Debe cerrar un archivo .log file antes de abrir uno nuevo; de lo contrario, obtendrá un mensaje de error.

.do files

Hasta ahora ha visto el uso interactivo de los comandos de Stata, lo cual es útil para depurar comandos y obtener una buena información de los datos. Usted escribe una línea de comando cada vez, y Stata procesa ese comando, muestra el resultado (si lo hay), y espera el siguiente. Aunque este enfoque tiene sus propios beneficios, un uso más avanzado de Stata implica la ejecución de comandos en un lote (rutina), es decir, los comandos se agrupan y se envían juntos a Stata en lugar de uno a la vez.

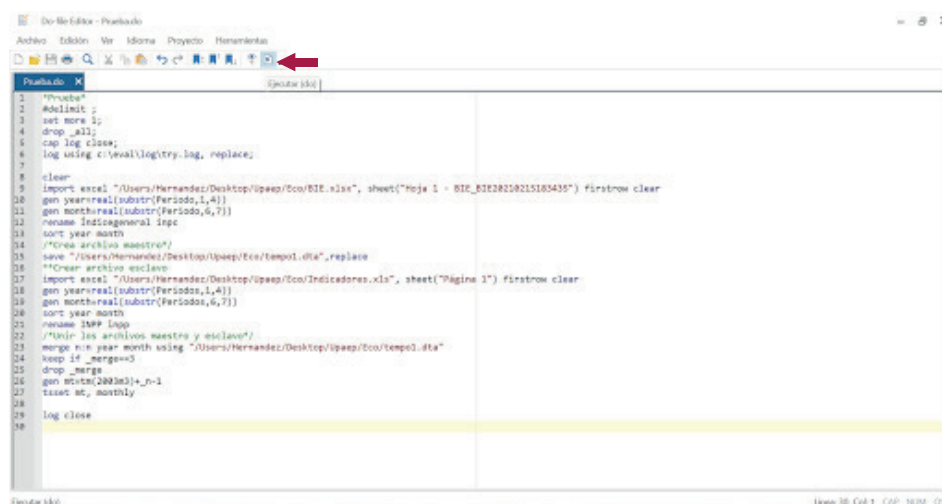
Si se encuentra usando el mismo conjunto de comandos repetidamente, puede guardar esos comandos en un archivo y ejecutarlos juntos cuando los necesite. Estos se denominan archivos .do file, y son el equivalente de Stata de las macros. Puede crear archivos .do al menos de tres maneras:

1. Simplemente escriba los comandos en un archivo de texto, etiquételo como "educ.do" (el sufijo .do es importante) y ejecute el archivo usando "do educ" en la ventana de comandos de Stata.

- Haga clic con el botón derecho en cualquier parte de la ventana "Revisar" para guardar todos los comandos que se usaron de forma interactiva. El archivo en el que se guardaron se puede editar, etiquetar y usar como un archivo .do file.
- Utilice el editor .do file integrado de Stata. Se invoca haciendo clic en el icono en la parte superior de la página. Los comandos se pueden escribir en el editor.
Ejecute estos comandos resaltándolos y usando el ícono apropiado (ejecutar do) dentro del editor .do file. Con la práctica, este procedimiento se convierte en una forma muy rápida y conveniente de trabajar con Stata.

Véase un ejemplo de la ventana .do file mostrando una rutina en la Figura 9.9:

Figura 9.11 .do file



```

1  /*Comando*/
2  #delimit ;
3  set more 1;
4  drop _all;
5  cap log close;
6  log using c:\eval\log\try.log, replace;
7
8  clear
9  import excel "Users\Hernandez\Desktop\Upaep\Eco\BIE.xls", sheet("Hoja 1 - BIE_BIE2010215103435") firstrow clear
10 gen yearreal(substr(Periodo,1,4))
11 gen monthreal(substr(Periodo,6,7))
12 rename Indicegeneral lnpe
13 sort year month
14 /*Crear archivo maestro*/
15 save "Users\Hernandez\Desktop\Upaep\Eco\temp1.dta", replace
16 /*Crear archivo esclavo*/
17 import excel "Users\Hernandez\Desktop\Upaep\Eco\Indicadores.xls", sheet("Página 1") firstrow clear
18 gen yearreal(substr(Periodo,1,4))
19 gen monthreal(substr(Periodo,6,7))
20 sort year month
21 rename lnpe lnpe
22 /*Unir los archivos maestro y esclavo*/
23 merge m=1 year month using "Users\Hernandez\Desktop\Upaep\Eco\temp1.dta"
24 keep if _merge==3
25 drop _merge
26 gen stc=(2003m3)+_n-1
27 tssat stc, monthly
28
29 log close
30

```

Las principales ventajas de usar archivos .do file en lugar de escribir comandos línea por línea son la replicabilidad y la repetibilidad. Con un archivo .do file, uno puede replicar los resultados en los que se trabajó semanas o meses antes. Además, los archivos .do file son especialmente útiles cuando es necesario repetir conjuntos de comandos, por ejemplo, con diferentes conjuntos o grupos de datos.

La primera línea del archivo es un comentario. Stata trata cualquier línea que comience con un asterisco (*) como un comentario y la ignora. Puede escribir un comentario de varias líneas usando una barra inclinada y un asterisco (/*) como inicio del comentario, y terminar el comentario con un asterisco y una barra inclinada (/). Algunos comandos de la Figura 9.9 nos indica:

#delimitar: De manera predeterminada, Stata asume que cada comando finaliza



con un retorno de carro (es decir, presionando la tecla Intro). Sin embargo, si un comando es demasiado largo para caber en una línea, puede distribuirlo en más de una línea. Para ello, informe a Stata cuál es el delimitador de comando. El comando del ejemplo dice que un punto y coma (;) finaliza un comando. Cada comando que sigue al comando "delimitar" debe terminar con un punto y coma. Aunque para este archivo .do en particular no se necesita el comando "#delimit" (todos los comandos son lo suficientemente cortos), se hace para explicar el comando.

set more 1: Normalmente, Stata muestra los resultados de una en una pantalla y espera a que el usuario pulse alguna tecla. Pero este proceso pronto se convertiría en una molestia si, después de dejar que se ejecute un archivo .do, hay que pulsar una tecla para cada pantalla hasta que el programa termine. Este comando muestra toda la salida, saltando página tras página automáticamente.

drop _all: Este comando borra la memoria.

cap log close: Este comando cierra cualquier archivo .log file abierto. Si no hay ningún archivo de registro abierto, Stata simplemente ignora este comando.

.ado files

Los archivos .ado files son programas de Stata destinados a realizar tareas específicas. Muchos comandos de Stata se implementan como archivos .ado files (por ejemplo, el comando "summarize"). Para ejecutar dicho programa, simplemente escriba el nombre del programa en la línea de comando. Los usuarios pueden escribir sus propios programas .ado para cumplir con requisitos especiales. De hecho, los usuarios y desarrolladores de Stata escriben continuamente dichos programas, que a menudo se ponen a disposición de la gran comunidad de usuarios de Stata en Internet. Stata tiene comandos incorporados para descargar e incorporar tales comandos. Por ejemplo, la técnica de coincidencia de puntuación de propensión se implementa mediante un archivo .ado denominado pscore.ado. Para descargar la última versión de este comando, escriba el siguiente comando en la línea de comandos:

```
. findit pscore
```

Stata responde con una lista de implementaciones .ado del programa. Al hacer clic en uno de ellos, dará sus detalles y presentará la opción para instalarlo. Cuando Stata instala un programa .ado, también instala los archivos de ayuda asociados con él.

10. Evaluación de Impacto aleatoria

La aleatorización funciona en el escenario ideal donde los individuos o los hogares se asignan al azar, eliminando el sesgo de selección. En un intento de obtener una estimación del impacto de un determinado programa, la comparación de las personas tratadas a lo largo del tiempo no brinda una estimación consistente del impacto del programa, porque otros factores, además del programa, pueden afectar los resultados. Sin embargo, comparar el resultado de los individuos tratados con el de un grupo de control similar puede proporcionar una estimación del impacto del programa. Esta comparación funciona bien con la aleatorización, porque la asignación de individuos o grupos a los grupos de tratamiento y comparación es aleatoria. Se obtendrá una estimación imparcial del impacto del programa en la muestra cuando el diseño y la implementación de la evaluación aleatoria sean apropiados. Este ejercicio demuestra la estimación aleatoria del impacto con diferentes escenarios.

Impactos del programa

En este apartado analizaremos los impactos del programa Oportunidades. Supongamos que este programa se asigna al azar. Se genera una variable “zm” zona metropolitana que toma el valor de 1 si la persona está en el grupo de tratamiento y toma el valor de 0 si la persona está en el grupo control:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad
> o de Puebla.dta"
```

```
. tab zm
```

zm	Freq.	Percent	Cum.
0	1,973	65.70	65.70
1	1,030	34.30	100.00
Total	3,003	100.00	



Parte 2

Use el método más simple para calcular el efecto de tratamiento promedio usando el comando “ttest” de Stata, que compara el resultado entre el grupo de tratamiento y grupo control por ingreso mensual.

```
. ttest ing_mens, by(zm)
```

Two-sample t test with equal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
0	905	3264.509	98.38641	2959.78	3071.417	3457.602
1	587	4746.826	204.4543	4953.536	4345.274	5148.379
combined	1,492	3847.7	101.8593	3934.46	3647.897	4047.502
diff		-1482.317	205.0141		-1884.464	-1080.17

diff = mean(0) - mean(1) t = -7.2303
Ho: diff = 0 degrees of freedom = 1490

Ha: diff < 0 Ha: diff != 0 Ha: diff > 0
Pr(T < t) = 0.0000 Pr(|T| > |t|) = 0.0000 Pr(T > t) = 1.0000

Observe los resultados empleando el ingreso mensual en logaritmos:

```
. ttest ln_ing, by(zm)
```

Two-sample t test with equal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
0	959	.8476458	.0184756	.5721477	.8113884	.8839031
1	599	1.067879	.0209949	.5138398	1.026646	1.109111
combined	1,558	.9323181	.0142035	.5606321	.9044582	.9601781
diff		-.220233	.0286678		-.2764646	-.1640013

diff = mean(0) - mean(1) t = -7.6822
Ho: diff = 0 degrees of freedom = 1556

Ha: diff < 0 Ha: diff != 0 Ha: diff > 0
Pr(T < t) = 0.0000 Pr(|T| > |t|) = 0.0000 Pr(T > t) = 1.0000

Alternativamente, puede aplicar mínimos cuadrados ordinarios:

```
. reg ing_mens zm
```

Source	SS	df	MS	Number of obs	=	1,492
Model	782347941	1	782347941	F(1, 1490)	=	52.28
Residual	2.2298e+10	1,490	14965297.7	Prob > F	=	0.0000
				R-squared	=	0.0339
				Adj R-squared	=	0.0332
Total	2.3081e+10	1,491	15479974.2	Root MSE	=	3868.5

ing_mens	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]
zm	1482.317	205.0141	7.23	0.000	1080.17 1884.464
_cons	3264.509	128.5933	25.39	0.000	3012.266 3516.753

Observe los resultados empleando el ingreso mensual en logaritmos:

```
. reg ln_ing zm
```

Source	SS	df	MS	Number of obs	=	1,558
Model	17.8830961	1	17.8830961	F(1, 1556)	=	59.02
Residual	471.494935	1,556	.30301731	Prob > F	=	0.0000
				R-squared	=	0.0365
				Adj R-squared	=	0.0359
Total	489.378031	1,557	.314308305	Root MSE	=	.55047

ln_ing	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]
zm	.220233	.0286678	7.68	0.000	.1640013 .2764646
_cons	.8476458	.0177756	47.69	0.000	.8127791 .8825124

Aplicando modelo logit:

```
. logit ln_ing zm
```

```
Iteration 0: log likelihood = -787.05473
Iteration 1: log likelihood = -773.35187
Iteration 2: log likelihood = -773.17061
Iteration 3: log likelihood = -773.17053
Iteration 4: log likelihood = -773.17053
```

```
Logistic regression               Number of obs   =   1,558
                                LR chi2(1)          =   27.77
                                Prob > chi2           =   0.0000
Log likelihood = -773.17053       Pseudo R2       =   0.0176
```

ln_ing	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
zm	.7161177	.1405912	5.09	0.000	.440564 .9916715
_cons	1.125206	.0750769	14.99	0.000	.9780579 1.272354



Parte 2

Obtenga los impactos marginales:

```
. margins, dydx(*)
```

```
Average marginal effects      Number of obs    =      1,558
Model VCE      : OIM
```

```
Expression      : Pr(ln_ing), predict()
dy/dx w.r.t.    : zm
```

	Delta-method					
	dy/dx	Std. Err.	z	P> z	[95% Conf. Interval]	
zm	.1140772	.0220803	5.17	0.000	.0708005	.1573538

Aplicando modelo logit considerando género:

```
. logit ln_ing gen_zm
```

```
Iteration 0:    log likelihood = -787.05473
Iteration 1:    log likelihood = -787.03869
Iteration 2:    log likelihood = -787.03869
```

```
Logistic regression      Number of obs    =      1,558
                        LR chi2(1)      =        0.03
                        Prob > chi2     =      0.8579
Log likelihood = -787.03869      Pseudo R2      =      0.0000
```

	ln_ing	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
	gen_zm	.031573	.1766971	0.18	0.858	-.314747	.377893
	_cons	1.360026	.0682222	19.94	0.000	1.226313	1.49374

Nuevamente obtenga los impactos marginales:

```
. margins, dydx(*)
```

```
Average marginal effects      Number of obs    =      1,558
Model VCE      : OIM
```

```
Expression      : Pr(ln_ing), predict()
dy/dx w.r.t.    : gen_zm
```

	Delta-method					
	dy/dx	Std. Err.	z	P> z	[95% Conf. Interval]	
gen_zm	.0051168	.0286358	0.18	0.858	-.0510082	.0612419

La regresión anterior estima el impacto general de los programas implementados en el estado de Puebla en el ingreso mensual de los hogares. Los resultados pueden ser diferentes del impacto en el ingreso mensual después de mantener constantes otros factores, es decir, especificar el modelo ajustado para las

covariables que afectan los resultados de interés. Ahora, haga una regresión empleando el logaritmo del ingreso mensual más otros factores que pueden influir en el ingreso.

```
. reg ing_mens zm SEXO_01 EDAD_01 educ electricidad agua_entubada banio refrigerador estufa
```

Source	SS	df	MS	Number of obs	=	1,488
Model	3.0426e+09	9	338069556	F(9, 1478)	=	25.00
Residual	1.9989e+10	1,478	13524060.9	Prob > F	=	0.0000
				R-squared	=	0.1321
				Adj R-squared	=	0.1268
Total	2.3031e+10	1,487	15488357.8	Root MSE	=	3677.5

ing_mens	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
zm	1089.183	207.8359	5.24	0.000	681.4985	1496.868
SEXO_01	-1263.563	199.9917	-6.32	0.000	-1655.861	-871.2652
EDAD_01	10.11161	11.3413	0.89	0.373	-12.13514	32.35836
educ	794.0312	87.48158	9.08	0.000	622.4299	965.6325
electricidad	-565.9056	891.2677	-0.63	0.526	-2314.19	1182.379
agua_entubada	299.1753	292.3021	1.02	0.306	-274.1958	872.5463
banio	-380.364	224.2389	-1.70	0.090	-820.2244	59.49639
refrigerador	522.2035	240.3649	2.17	0.030	50.7109	993.696
estufa	10.55505	317.1162	0.03	0.973	-611.4907	632.6008
_cons	1607.401	1027.323	1.56	0.118	-407.7652	3622.567

Consideramos la variable “zm” por género:

```
. tab gen_zm
```

gen_zm	Freq.	Percent	Cum.
0	2,404	80.05	80.05
1	599	19.95	100.00
Total	3,003	100.00	

```
. reg ing_mens gen_zm EDAD_01 educ electricidad agua_entubada banio refrigerador estufa
```

Source	SS	df	MS	Number of obs	=	1,488
Model	2.3187e+09	8	289840205	F(8, 1479)	=	20.70
Residual	2.0712e+10	1,479	14004372.1	Prob > F	=	0.0000
				R-squared	=	0.1007
				Adj R-squared	=	0.0958
Total	2.3031e+10	1,487	15488357.8	Root MSE	=	3742.2

ing_mens	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
gen_zm	-903.712	274.0282	-3.30	0.001	-1441.237	-366.1867
EDAD_01	8.561326	11.56122	0.74	0.459	-14.11681	31.23946
educ	876.3367	88.33002	9.92	0.000	703.0713	1049.602
electricidad	-686.0085	907.1348	-0.76	0.450	-2465.416	1093.399
agua_entubada	221.4977	297.2945	0.75	0.456	-361.666	804.6614
banio	-189.3781	225.1184	-0.84	0.400	-630.9633	252.2071
refrigerador	620.1252	244.2479	2.54	0.011	141.0161	1099.234
estufa	202.764	321.2547	0.63	0.528	-427.3994	832.9274
_cons	1418.705	1043.69	1.36	0.174	-628.5647	3465.975



Parte 2

Impactos de la participación en el programa

Aunque la asignación del programa de oportunidades es aleatoria entre las zonas, la participación puede no serlo. Sólo aquellos hogares que viven en situación de pobreza son los llamados grupos objetivo. Comience con el método más simple para calcular el efecto promedio del tratamiento de la participación en el programa. Usando el comando “ttest”, que compara el resultado entre el grupo de tratamiento y el de control:

```
. ttest P6_10A, by(zm)
```

Two-sample t test with equal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
0	630	1.084127	.0110678	.2777987	1.062393	1.105861
1	159	1.157233	.0289599	.3651702	1.100034	1.214431
combined	789	1.098859	.0106327	.2986624	1.077988	1.119731
diff		-.0731057	.0263949		-.1249184	-.021293

```
diff = mean(0) - mean(1)                                t = -2.7697
Ho: diff = 0                                           degrees of freedom = 787
```

```
Ha: diff < 0                                Ha: diff != 0                                Ha: diff > 0
Pr(T < t) = 0.0029                        Pr(|T| > |t|) = 0.0057                        Pr(T > t) = 0.9971
```

Nuevamente alternativamente puede aplicar mínimos cuadrados ordinarios:

```
. reg P6_10A zm
```

Source	SS	df	MS	Number of obs	=	789
Model	.678521153	1	.678521153	F(1, 787)	=	7.67
Residual	69.6104522	787	.088450384	Prob > F	=	0.0057
				R-squared	=	0.0097
				Adj R-squared	=	0.0084
Total	70.2889734	788	.089199205	Root MSE	=	.29741

P6_10A	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
zm	.0731057	.0263949	2.77	0.006	.021293	.1249184
_cons	1.084127	.0118489	91.50	0.000	1.060868	1.107386

Considerando factores adicionales:

```
. reg P6_10A zm SEXO_01 EDAD_01 educ electricidad agua_entubada banio refrigerador estufa
```

Source	SS	df	MS	Number of obs	=	787
Model	6.74114788	9	.749016431	F(9, 777)	=	9.16
Residual	63.5282295	777	.081760913	Prob > F	=	0.0000
				R-squared	=	0.0959
				Adj R-squared	=	0.0855
Total	70.2693774	786	.089401243	Root MSE	=	.28594

P6_10A	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]
zm	.0301189	.0267882	1.12	0.261	-.0224669 .0827046
SEXO_01	-.0238694	.0225135	-1.06	0.289	-.0680639 .0203251
EDAD_01	.0062134	.0012411	5.01	0.000	.0037771 .0086498
educ	.047583	.0103953	4.58	0.000	.0271767 .0679893
electricidad	.0450472	.0874425	0.52	0.607	-.1266043 .2166986
agua_entubada	-.0229827	.0282091	-0.81	0.415	-.0783578 .0323923
banio	.0373108	.0224577	1.66	0.097	-.0067742 .0813958
refrigerador	.0640123	.0225924	2.83	0.005	.0196629 .1083617
estufa	.0203868	.0261711	0.78	0.436	-.0309876 .0717612
_cons	.6568088	.1075126	6.11	0.000	.4457593 .8678584

Considere todas las variables en un sólo modelo:

```
. reg ln_ing P6_10A zm SEXO_01 EDAD_01 educ electricidad agua_entubada banio refrigerador estufa
```

Source	SS	df	MS	Number of obs	=	314
Model	9.35198426	10	.935198426	F(10, 303)	=	2.71
Residual	104.579031	303	.345145317	Prob > F	=	0.0034
				R-squared	=	0.0821
				Adj R-squared	=	0.0518
Total	113.931015	313	.363996854	Root MSE	=	.58749

ln_ing	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]
P6_10A	-.0699669	.107399	-0.65	0.515	-.2813092 .1413753
zm	.1826526	.0816214	2.24	0.026	.0220362 .3432691
SEXO_01	-.1008966	.0686713	-1.47	0.143	-.2360296 .0342364
EDAD_01	-.0047283	.0042221	-1.12	0.264	-.0130367 .0035801
educ	.0450383	.0343432	1.31	0.191	-.022543 .1126196
electricidad	-.5390378	.3061592	-1.76	0.079	-1.141505 .0634298
agua_entubada	.098078	.0942978	1.04	0.299	-.0874835 .2836395
banio	-.03154	.0751487	-0.42	0.675	-.1794195 .1163395
refrigerador	.2072795	.0780874	2.65	0.008	.0536173 .3609417
estufa	-.1601002	.0909327	-1.76	0.079	-.3390397 .0188393
_cons	1.471919	.3703446	3.97	0.000	.7431463 2.200693

11. Método diferencias en diferencias

Los métodos de emparejamiento discutidos en ejercicios anteriores están destinados a reducir el sesgo al elegir los grupos de tratamiento y comparación



sobre la base de las características observables. Por lo general, se implementan después de que el programa ha estado funcionando durante algún tiempo y se han recopilado los datos de la encuesta. Otra forma poderosa de medir el impacto de un programa es mediante el uso de datos de panel, recopilados a partir de una encuesta de referencia antes de que se implementara el programa y después de que este haya estado funcionando durante algún tiempo. Estas dos encuestas deben ser comparables en cuanto a las preguntas y los métodos, y deben administrarse tanto a los participantes como a los no participantes. El uso de datos de panel permite eliminar el sesgo de variables no observadas, siempre que no cambie con el tiempo.

Este enfoque, el método de doble diferencia (DD, también conocido comúnmente como diferencia en diferencia), ha sido popular en las evaluaciones no experimentales. El método DD estima la diferencia en el resultado durante el período posterior a la intervención entre un grupo de tratamiento y un grupo de comparación en relación con los resultados observados durante una encuesta de referencia previa a la intervención.

Estimación

Use la base de datos original:

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad
> o de Puebla.dta"
```

Considere la variable tiempo (dummy) que toma el valor de 0 antes del tratamiento y 1 después del tratamiento:

```
. tab tiempo
```

tiempo	Freq.	Percent	Cum.
0	1,235	41.13	41.13
1	1,768	58.87	100.00
Total	3,003	100.00	

Considere la variable zona metropolitana “zm” (dummy) que toma el valor de 0 para grupo control y 1 para el grupo de tratamiento:

```
. label define etiqueta_zm 0 "Control" 1 "Treated"
```

```
. label values zm etiqueta_zm
```

```
. tab zm
```

zm	Freq.	Percent	Cum.
Control	1,973	65.70	65.70
Treated	1,030	34.30	100.00
Total	3,003	100.00	

Genere una variable llamada “did” que muestre la interacción entre la intervención y el tiempo:

```

* Crear la variable did
gen did = tiempo * zm

```

Aplicando diferencias en diferencias (DD)

Observe los impactos sobre el ingreso mensual:

```
. reg ing_mens tiempo zm did, r
```

Linear regression	Number of obs	=	1,492
	F(3, 1488)	=	22.84
	Prob > F	=	0.0000
	R-squared	=	0.0730
	Root MSE	=	3792

ing_mens	Coef.	Robust Std. Err.	t	P> t	[95% Conf. Interval]	
tiempo	-700.804	208.5192	-3.36	0.001	-1109.827	-291.7812
zm	2172.221	328.2531	6.62	0.000	1528.333	2816.109
did	-1699.974	407.3541	-4.17	0.000	-2499.023	-900.9246
_cons	3503.015	119.1941	29.39	0.000	3269.209	3736.821



También puede reescribir el comando de la siguiente manera:

```
. reg ing_mens tiempo zm did, vce(robust)
```

Linear regression	Number of obs	=	1,492
	F(3, 1488)	=	22.84
	Prob > F	=	0.0000
	R-squared	=	0.0730
	Root MSE	=	3792

ing_mens	Coef.	Robust Std. Err.	t	P> t	[95% Conf. Interval]	
tiempo	-700.804	208.5192	-3.36	0.001	-1109.827	-291.7812
zm	2172.221	328.2531	6.62	0.000	1528.333	2816.109
did	-1699.974	407.3541	-4.17	0.000	-2499.023	-900.9246
_cons	3503.015	119.1941	29.39	0.000	3269.209	3736.821



Parte 2

Estimación del estimador DD (utilizando el método de la etiqueta, sin necesidad de generar la interacción):

```
. reg ing_mens tiempo##zm, r
```

```
Linear regression              Number of obs   =    1,492
                              F(3, 1488)      =    22.84
                              Prob > F        =    0.0000
                              R-squared       =    0.0730
                              Root MSE    =    3792
```

ing_mens	Coef.	Robust Std. Err.	t	P> t	[95% Conf. Interval]	
1.tiempo	-700.804	208.5192	-3.36	0.001	-1109.827	-291.7812
zm						
Treated	2172.221	328.2531	6.62	0.000	1528.333	2816.109
tiempo#zm						
1#Treated	-1699.974	407.3541	-4.17	0.000	-2499.023	-900.9246
_cons	3503.015	119.1941	29.39	0.000	3269.209	3736.821

Considere otros factores adicionales:

```
. reg ing_mens tiempo zm did agua_entubada banio electricidad refrigerador educ, r
```

```
Linear regression              Number of obs   =    1,488
                              F(8, 1479)      =    21.24
                              Prob > F        =    0.0000
                              R-squared       =    0.1400
                              Root MSE    =    3659.4
```

ing_mens	Coef.	Robust Std. Err.	t	P> t	[95% Conf. Interval]	
tiempo	-641.954	206.2046	-3.11	0.002	-1046.439	-237.4694
zm	1672.832	304.5319	5.49	0.000	1075.471	2270.192
did	-1538.995	392.4379	-3.92	0.000	-2308.789	-769.2007
agua_entubada	306.9028	233.967	1.31	0.190	-152.0397	765.8452
banio	-432.1062	168.1198	-2.57	0.010	-761.8848	-102.3275
electricidad	-628.9861	664.4076	-0.95	0.344	-1932.268	674.2955
refrigerador	536.1485	161.0695	3.33	0.001	220.1995	852.0976
educ	767.669	81.18407	9.46	0.000	608.4209	926.9172
_cons	1986.503	685.1456	2.90	0.004	642.5428	3330.464

Para instalar el comando en Stata, escriba:

```
. ssc install diff
checking diff consistency and verifying not already installed...
installing into C:\Users\ACER SWIFT\ado\plus\...
installation complete.
```

```
. diff ing_mens, t(zm) p(tiempo)
```

DIFFERENCE-IN-DIFFERENCES ESTIMATION RESULTS

Number of observations in the DIFF-IN-DIFF: 1492

	Before	After		
Control:	597	308	905	
Treated:	360	227	587	
	957	535		
Outcome var.	ing_mens	S. Err.	t	P> t
Before				
Control	3503.015			
Treated	5675.236			
Diff (T-C)	2172.221	253.037	8.58	0.000***
After				
Control	2802.211			
Treated	3274.458			
Diff (T-C)	472.247	331.707	1.42	0.155
Diff-in-Diff	-1.7e+03	417.201	4.07	0.000***

R-square: 0.07

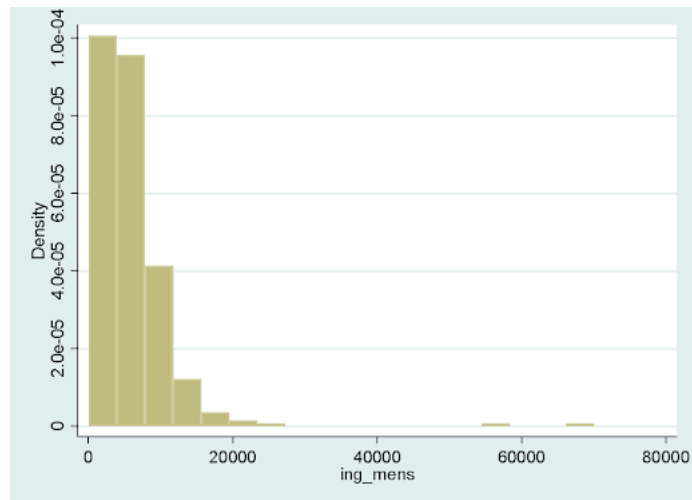
* Means and Standard Errors are estimated by linear regression

Inference: * p<0.01; ** p<0.05; * p<0.1

Análisis gráfico

Impacto del ingreso mensual antes de la intervención al grupo treated:

```
. twoway histogram ing_mens if tiempo==0 & zm == 1
```

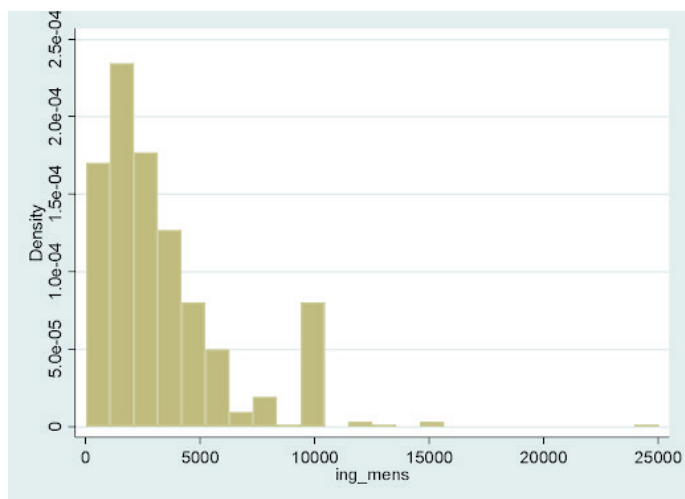




Parte 2

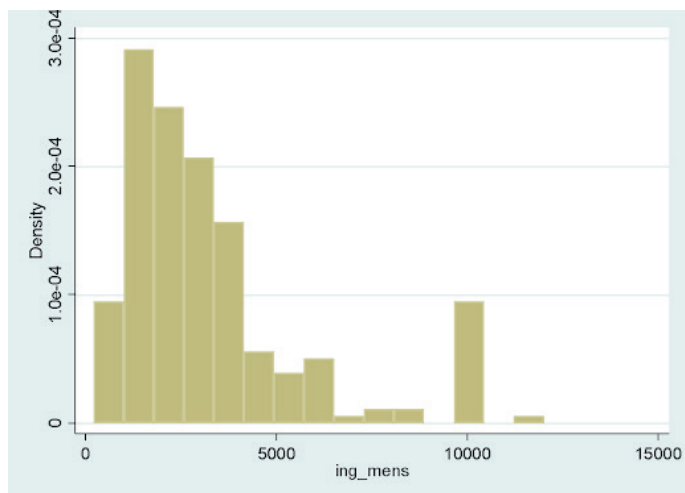
Impacto del ingreso mensual antes de la intervención al grupo control:

```
. twoway histogram ing_mens if tiempo==0 & zm == 0
```



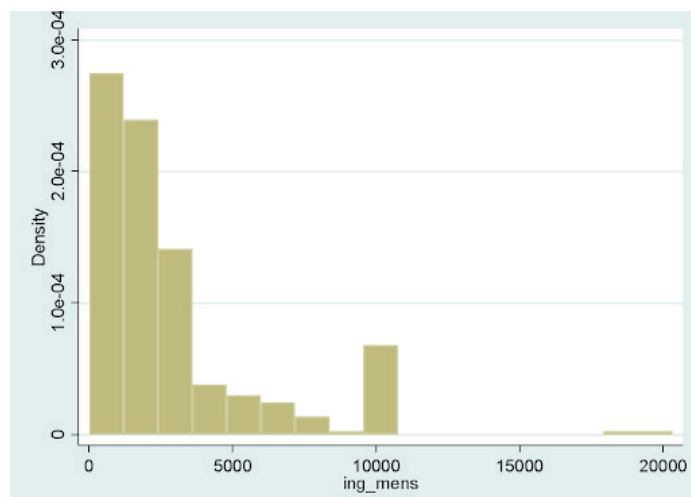
Impacto del ingreso mensual después de la intervención al grupo treated:

```
. twoway histogram ing_mens if tiempo==1 & zm == 1
```



Impacto del ingreso mensual después de la intervención al grupo control:

```
. twoway histogram ing_mens if tiempo==1 & zm == 0
```



Considere otros factores adicionales, empleando el comando “diff”:

```
. diff ing_mens, t(zm) p(tiempo) cov( agua_entubada banio electricidad refrigerador educ) test  
TWO-SAMPLE T TEST
```

Number of observations (baseline): 957

	Before	After	
Control:	597	-	597
Treated:	360	-	360
	957	-	

t-test at period = 0:

Variable(s)	Mean Control	Mean Treated	Diff.	t	Pr(T > t)
ing_mens	3503.015	5675.236	2172.221	7.68	0.0000***
agua_entub~a	0.843	0.894	0.052	2.26	0.0241**
banio	0.475	0.811	0.336	10.90	0.0000***
electricidad	0.983	0.994	0.011	1.51	0.1319
refrigerador	0.635	0.844	0.210	7.12	0.0000***
educ	2.262	2.950	0.688	8.63	0.0000***

*** p<0.01; ** p<0.05; * p<0.1



Parte 2

O bien:

```
. diff ing_mens, t(zm) p(tiempo) cov( agua_entubada banio electricidad refrigerador educ) robus
> t
```

DIFFERENCE-IN-DIFFERENCES WITH COVARIATES

DIFFERENCE-IN-DIFFERENCES ESTIMATION RESULTS

Number of observations in the DIFF-IN-DIFF: 1488

	Before	After		
Control:	595	307	902	
Treated:	360	226	586	
	955	533		
Outcome var.	ing_mens	S. Err.	t	P> t
Before				
Control	1986.503			
Treated	3659.335			
Diff (T-C)	1672.832	304.532	5.49	0.000***
After				
Control	1344.549			
Treated	1478.386			
Diff (T-C)	133.837	240.006	0.56	0.577
Diff-in-Diff	-1.5e+03	392.438	3.92	0.000***

R-square: 0.14

* Means and Standard Errors are estimated by linear regression

**Robust Std. Errors

Inference: * p<0.01; ** p<0.05; * p<0.1

Estimación alternativa

Como otra opción, también puede emplear el comando “ttest”:

```
. ttest tiempo, by(zm)
```

Two-sample t test with equal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
Control	1,973	.5924987	.0110651	.4914941	.5707982	.6141992
Treated	1,030	.5815534	.0153783	.4935439	.5513771	.6117297
combined	3,003	.5887446	.0089808	.4921433	.5711355	.6063537
diff		.0109453	.0189206		-.0261534	.048044

diff = mean(Control) - mean(Treated) t = 0.5785
Ho: diff = 0 degrees of freedom = 3001

Ha: diff < 0
Pr(T < t) = 0.7185

Ha: diff != 0
Pr(|T| > |t|) = 0.5630

Ha: diff > 0
Pr(T > t) = 0.2815

Comprobación de la robustez de la DID con una regresión de efectos fijos:

```
. xtreg ln_ing tiempo did, fe i(zm)
```

```
Fixed-effects (within) regression      Number of obs   =    1,558
Group variable: zm                    Number of groups =      2

R-sq:                                Obs per group:
    within = 0.0312                      min =      599
    between = 1.0000                     avg =     779.0
    overall = 0.0171                      max =      959

corr(u_i, Xb) = -0.2049                F(2,1554)       =    25.04
                                         Prob > F        =    0.0000
```

ln_ing	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
tiempo	-.1652095	.0369594	-4.47	0.000	-.237705	-.0927139
did	-.0834838	.0584913	-1.43	0.154	-.198214	.0312463
_cons	1.004558	.0171938	58.43	0.000	.9708326	1.038284
sigma_u	.1853007					
sigma_e	.54215851					
rho	.10459719	(fraction of variance due to u_i)				

```
F test that all u_i=0: F(1, 1554) = 53.90                Prob > F = 0.0000
```

Puede incluir otras variables en la regresión. Así, el modelo de efectos fijos se puede extender de la siguiente manera:

```
. xtreg ln_ing tiempo agua_entubada banio electricidad refrigerador educ, fe i(zm)
```

```
Fixed-effects (within) regression      Number of obs   =    1,554
Group variable: zm                    Number of groups =      2

R-sq:                                Obs per group:
    within = 0.0961                      min =      598
    between = 1.0000                     avg =     777.0
    overall = 0.1093                     max =      956

corr(u_i, Xb) = 0.1551                F(6,1546)       =    27.41
                                         Prob > F        =    0.0000
```

ln_ing	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
tiempo	-.1787501	.0278059	-6.43	0.000	-.2332914	-.1242088
agua_entubada	.0617263	.0421834	1.46	0.144	-.0210163	.144469
banio	-.0604094	.0311748	-1.94	0.053	-.1215588	.0007401
electricidad	-.3431018	.1258913	-2.73	0.006	-.5900376	-.0961661
refrigerador	.1270196	.0322521	3.94	0.000	.0637571	.1902821
educ	.0939246	.0116956	8.03	0.000	.0709836	.1168656
_cons	.9923248	.1271739	7.80	0.000	.7428732	1.241776
sigma_u	.11947429					
sigma_e	.52374379					
rho	.04946297	(fraction of variance due to u_i)				

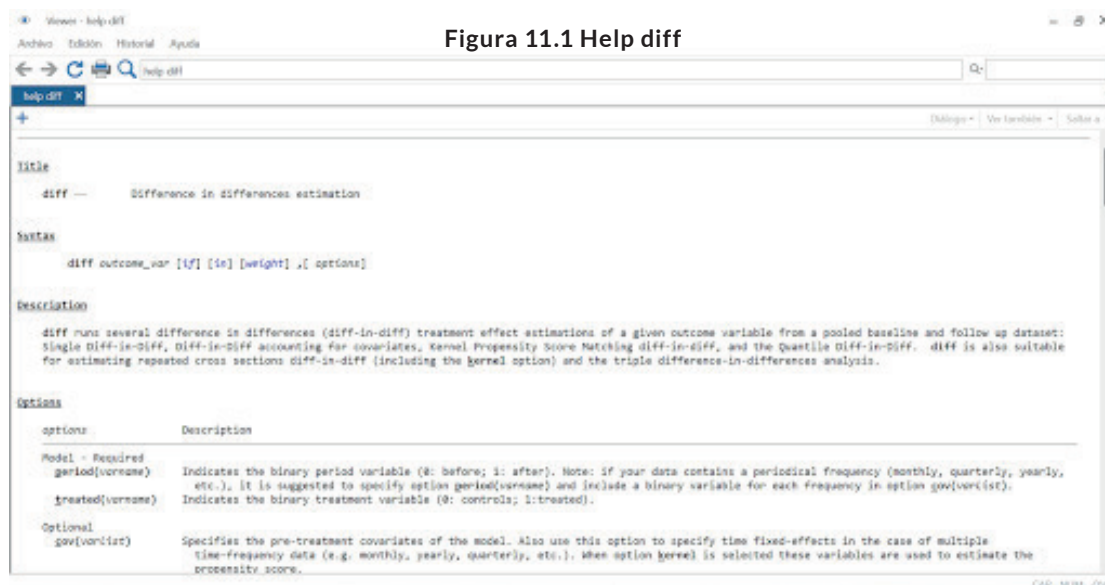
```
F test that all u_i=0: F(1, 1546) = 34.09                Prob > F = 0.0000
```



Aunque la implementación de DD a través de la regresión (OLS o efectos fijos) controla las covariables a nivel de hogar y comunidad, las condiciones iniciales durante la encuesta de línea de base pueden tener una influencia separada en los cambios posteriores en el resultado o la asignación al tratamiento. Por lo tanto, ignorar el efecto separado de las condiciones iniciales puede sesgar las estimaciones de DD.

Sin embargo, incluir las condiciones iniciales en la regresión es complicado. Debido a que las observaciones de la línea de base en la muestra del panel ya contienen características iniciales, no se pueden agregar directamente variables adicionales para las condiciones iniciales. Una forma de agregar estas es tener en cuenta una implementación alternativa de la regresión de efectos fijos. En esta implementación, las variables de diferencia se crean para todas las variables entre los años, y luego estas se utilizan en la regresión en lugar de las variables originales. En este conjunto de datos modificado, las variables de condición inicial se pueden agregar como regresores adicionales sin un problema de colinealidad. Aunque se sugiere ampliamente utilizar el comando “diff” para mejores resultados. La siguiente figura 11.1 muestra los amplios usos del comando que se pueden emplear de acuerdo con el tipo de datos que estén trabajando.

Figura 11.1 Help diff



<code>kernel</code>	Performs the Kernel-based Propensity Score Matching diff-in-diff. This option generates the variable <code>_weights</code> containing the weights derived from the Kernel Propensity Score Matching, and <code>_ps</code> when the Propensity Score is not supplied in <code>pscore(varname)</code> , following Leisen and Sianesi (2014) . This option requires the <code>id(varname)</code> of each unit or individual except under the repeated cross section <code>rca</code> setting.
<code>id(varname)</code>	Option kernel requires the supply of the identification variable.
<code>bw(#)</code>	Supplied bandwidth of the Kernel function. The default bandwidth is 0.05.
<code>ktype(kernel)</code>	Specifies the type of the Kernel function. The types are <code>spannecheisov</code> (the default), <code>gaussian</code> , <code>biweight</code> , <code>uniform</code> and <code>tricube</code> .
<code>rca</code>	Indicates that the kernel is set for repeated cross section. This option does not require option <code>id(varname)</code> . Option <code>rca</code> strongly assumes that covariates in <code>ps(varlist)</code> do not vary over time.
<code>qdid(quantile)</code>	Performs the Quantile Difference in Differences estimation at the specified quantile from 0.1 to 0.9 (quantile 0.5 performs the QDD at the median). You may combine this option with <code>kernel</code> and <code>gov</code> . <code>qdid</code> does not support weights nor robust standard errors. This option uses <code>[R] ereg</code> and <code>[R] lmerreg</code> for bootstrapped standard errors.
<code>pscore(varname)</code>	Supplied Propensity Score.
<code>logit</code>	Specifies logit estimation of the Propensity Score. The default is <code>probit</code> .
<code>zscore</code>	Performs diff on the common support of the propensity score given the option <code>kernel</code> .
<code>addcov(varlist)</code>	Indicates additional covariates in addition to those specified in the estimation of the propensity score. Also use this option to specify time fixed-effects in the case of multiple time-frequency data (e.g. monthly, yearly, quarterly, etc.).
<code>ddd(varname)</code>	Additional category for triple difference estimation. <code>treated(varname)</code> is deemed as the first category and <code>ddd(varname)</code> the second category. This option is not compatible with options <code>kernel</code> , <code>test</code> or <code>qdid(quantile)</code> .
<code>SE/Robust</code>	
<code>cluster(varname)</code>	Calculates clustered Std. Errors by <code>varname</code> .
<code>robust</code>	Calculates robust Std. Errors.
<code>bs</code>	Performs a Bootstrap estimation of coefficients and standard errors.
<code>reps(int)</code>	Specifies the number of repetitions when the <code>bs</code> is selected. The default are 50 repetitions.
Balancing test	
<code>test</code>	Performs a balancing t-test of the difference in the means of the covariates between the control and treated groups in period <code>= 0</code> . The option <code>test</code> combined with <code>kernel</code> performs the balancing t-test with the weighted covariates. See [R] ttest .
Reporting	
<code>coef1</code>	Displays the inference of the included covariates or the estimation of the Propensity Score when option <code>kernel</code> is specified.
<code>aster</code>	Removes the inference stars from the z-values.

GSP: NIRM: 0208

12. Método de regresión discontinua

Cuando el tratamiento se asigna exclusivamente sobre la base de un valor de corte, el diseño de discontinuidad de regresión (RD) es una alternativa adecuada a los experimentos aleatorios u otros diseños cuasi-experimentales. A diferencia del diseño aleatorizado, un grupo elegible no necesita ser excluido del tratamiento sólo por el bien de la evaluación del impacto.

La evaluación del impacto de la RD se basa en la idea de que la muestra en la zona del punto de corte (por encima y por debajo) representa las características del diseño aleatorio, porque los hogares de los grupos de tratamiento y de control son muy similares en sus características y sólo varían en su estado de tratamiento. Por lo tanto, una diferencia en los resultados medios de los grupos de tratamiento y de control restringida a la zona del punto de corte (es decir, local a la discontinuidad) proporciona el impacto de la intervención. La RD tiene dos versiones. En una, denominada discontinuidad aguda, el punto de corte establece de forma determinista el estado de tratamiento. O sea, todos los elegibles reciben el tratamiento, y ninguno de los no elegibles lo recibe. En el otro tipo de discontinuidad, denominada discontinuidad difusa, el estado de tratamiento no salta bruscamente de cero a uno a medida que los hogares pasan de ser elegibles a no serlo. Este escenario es más realista, porque algunos hogares elegibles deciden (por una u otra razón) no participar



Parte 2

en el programa, mientras que algunos hogares no elegibles sí participan. En un buen diseño de DR, puede venir dado por la siguiente expresión:

$$I = (y^+ - y^-) - (s^+ - s^-) \quad (12.1)$$

Estimación y aspectos adicionales

Use la base de datos original, considere la variable de ingreso total mensual y utilice el logaritmo del ingreso total del hogar y el logaritmo de la riqueza del hogar para fines prácticos en la aplicación de Stata.

```
. use "C:\Users\ACER SWIFT\Documents\Evaluación de Programas Sociales\Manual\Data Base Impact Estad
> o de Puebla.dta"
. sum ln_ing ln_riq
```

Variable	Obs	Mean	Std. Dev.	Min	Max
ln_ing	1,558	.9323181	.5606321	0	1.609438
ln_riq	1,192	.373505	1.077537	-4.475435	1.846319

Es necesario que instale `rdrobust`, como se muestra a continuación:

```
. net install rdrobust, from(https://raw.githubusercontent.com/rdpackages/rdrobust/master/stata)
> replace
checking rdrobust consistency and verifying not already installed...
installing into C:\Users\ACER SWIFT\ado\plus\...
installation complete.
```

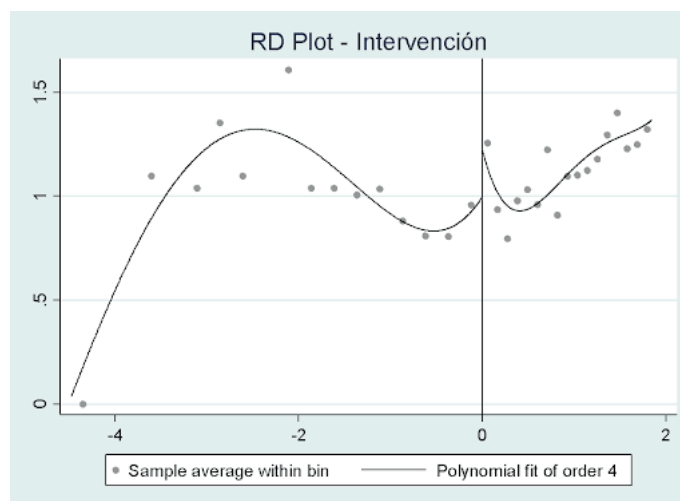
Usamos el comando “`rdplot`”:

```
. rdplot ln_ing ln_riq, graph_options(title(RD Plot - Intervención))
RD Plot with evenly spaced mimicking variance number of bins using polynomial regression.
```

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
			Kernel =	Uniform
Number of obs	198	492		
Eff. Number of obs	198	479		
Order poly. fit (p)	4	4		
BW poly. fit (h)	4.475	1.846		
Number of bins scale	1.000	1.000		

Outcome: `ln_ing`. Running variable: `ln_riq`.

	Left of c	Right of c
Bins selected	18	17
Average bin length	0.249	0.109
Median bin length	0.249	0.109
IMSE-optimal bins	8	7
Mimicking Var. bins	18	17
Rel. to IMSE-optimal:		
Implied scale	2.250	2.429
WIMSE var. weight	0.081	0.065
WIMSE bias weight	0.919	0.935



Construye un diagrama de RD alternativo utilizando intervalos uniformes seleccionados para trazar la función de regresión subyacente (es decir, selectores óptimos de IMSE) e implementados utilizando estimadores de espaciado, como se muestra a continuación:

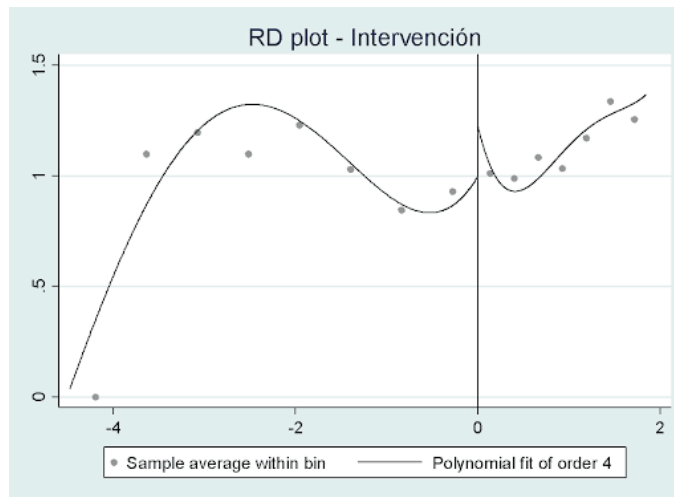
```
. rdplot ln_ing ln_riq, binselect(es) graph_options(title(RD plot - Intervención))
```

RD Plot with evenly spaced number of bins using polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
Number of obs	198	492	Kernel =	Uniform
Eff. Number of obs	198	479		
Order poly. fit (p)	4	4		
BW poly. fit (h)	4.475	1.846		
Number of bins scale	1.000	1.000		

Outcome: ln_ing. Running variable: ln_riq.

	Left of c	Right of c
Bins selected	8	7
Average bin length	0.559	0.264
Median bin length	0.559	0.264
IMSE-optimal bins	8	7
Mimicking Var. bins	18	17
Rel. to IMSE-optimal:		
Implied scale	1.000	1.000
WIMSE var. weight	0.500	0.500
WIMSE bias weight	0.500	0.500



El siguiente comando puede ilustrar algunas de las otras características del comando “rdplot”. En este caso, el gráfico RD se construye utilizando el enfoque de agrupamiento cuantílico-espaciado.

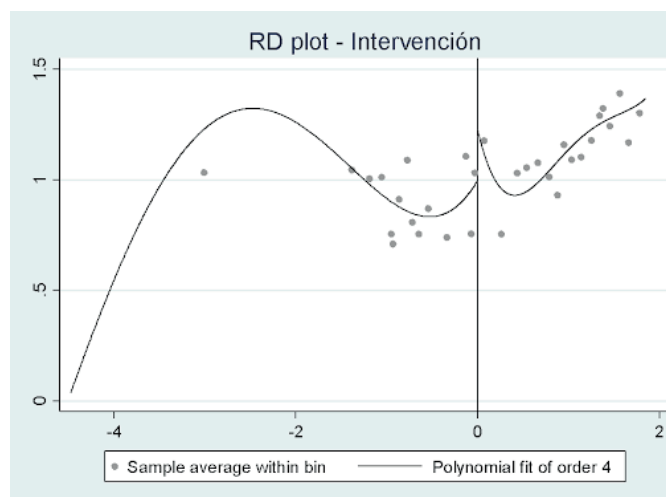
```
. rdplot ln_ing ln_riq, binselect(qsmv) graph_options(title(RD plot - Intervención))
```

RD Plot with quantile spaced mimicking variance number of bins using polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
			Kernel =	Uniform
Number of obs	198	492		
Eff. Number of obs	198	479		
Order poly. fit (p)	4	4		
BW poly. fit (h)	4.475	1.846		
Number of bins scale	1.000	1.000		

Outcome: ln_ing. Running variable: ln_riq.

	Left of c	Right of c
Bins selected	17	18
Average bin length	0.263	0.103
Median bin length	0.075	0.102
IMSE-optimal bins	22	10
Mimicking Var. bins	17	18
Rel. to IMSE-optimal:		
Implied scale	0.773	1.000
WIMSE var. weight	0.684	0.146
WIMSE bias weight	0.316	0.854



Llevamos a cabo una estimación e inferencia de los efectos del tratamiento de la DR basada en los datos. El comando “rdrobust”, utilizando sus opciones por defecto, da la siguiente salida:

```
. rdrobust ln_ing ln_rfq
```

Sharp RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
Number of obs	198	492	BW type	= mserd
Eff. Number of obs	64	99	Kernel	= Triangular
Order est. (p)	1	1	VCE method	= HH
Order bias (q)	2	2		
BW est. (h)	0.525	0.525		
BW bias (b)	0.860	0.860		
rho (h/b)	0.610	0.610		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_rfq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
Conventional	.1847	.13103	1.4096	0.159	-.072116 .441521
Robust	-	-	1.4038	0.160	-.086425 .522732

Estimates adjusted for mass points in the running variable.

El comando “rdrobust” también ofrece una salida más detallada, que incluye todos los estimadores puntuales, los estimadores de varianza y los Cis. Estos resultados se obtienen incluyendo la opción “all”, de la siguiente manera:

```
. rdrobust ln_ing ln_rfq, all
```



Sharp RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs = 690	
Number of obs	198	492	BW type	= mserd
Eff. Number of obs	64	99	Kernel	= Triangular
Order est. (p)	1	1	VCE method	= NN
Order bias (q)	2	2		
BW est. (h)	0.525	0.525		
BW bias (b)	0.860	0.860		
rho (h/b)	0.610	0.610		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_riq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	.1847	.13103	1.4096	0.159	-.072116	.441521
Bias-corrected	.21815	.13103	1.6649	0.096	-.038665	.474972
Robust	.21815	.1554	1.4038	0.160	-.086425	.522732

Estimates adjusted for mass points in the running variable.

Exploramos todos los procedimientos de selección de ancho de banda contenidos en este paquete. También podemos emplear `rdbwselect`:

```
. rdbwselect ln_ing ln_riq, all
```

Bandwidth estimators for sharp RD local polynomial regression.

Cutoff c =	Left of c	Right of c	Number of obs = 690	
Number of obs	198	492	Kernel	= Triangular
Min of ln_riq	-4.475	0.013	VCE method	= NN
Max of ln_riq	-0.007	1.846		
Order est. (p)	1	1		
Order bias (q)	2	2		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_riq.

Method	BW est. (h)		BW bias (b)	
	Left of c	Right of c	Left of c	Right of c
mserd	0.525	0.525	0.860	0.860
msetwo	0.712	0.348	1.218	0.733
msum	0.428	0.428	0.844	0.844
msecomb1	0.428	0.428	0.844	0.844
msecomb2	0.525	0.428	0.860	0.844
cernd	0.379	0.379	0.860	0.860
certwo	0.514	0.251	1.218	0.733
cersum	0.309	0.309	0.844	0.844
cercomb1	0.309	0.309	0.844	0.844
cercomb2	0.379	0.309	0.860	0.844

Estimates adjusted for mass points in the running variable.

Existen otros comandos que pueden utilizarse para llevar a cabo la inferencia en otras configuraciones de diseño de RD. Por ejemplo:

Estimación de la curva aguda RD

```
. rdrobust ln_ing ln_riq, deriv(1)
```

Sharp Kink RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
Number of obs	198	492	BW type =	mserd
Eff. Number of obs	59	86	Kernel =	Triangular
Order est. (p)	2	2	VCE method =	NN
Order bias (q)	3	3		
BW est. (h)	0.506	0.506		
BW bias (b)	0.786	0.786		
rho (h/b)	0.643	0.643		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_riq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	-5.8655	2.7944	-2.0990	0.036	-11.3424	-.388543
Robust	-	-	-1.9727	0.049	-14.1392	-.045897

Estimates adjusted for mass points in the running variable.

Estimación para RD difusa

```
. rdrobust ln_ing ln_riq, fuzzy(tiempo)
```

Fuzzy RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
Number of obs	198	492	BW type =	mserd
Eff. Number of obs	77	122	Kernel =	Triangular
Order est. (p)	1	1	VCE method =	NN
Order bias (q)	2	2		
BW est. (h)	0.664	0.664		
BW bias (b)	1.196	1.196		
rho (h/b)	0.555	0.555		
Unique obs	87	217		

First-stage estimates. Outcome: tiempo. Running variable: ln_riq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	.26106	.12265	2.1285	0.033	.020666	.501457
Robust	-	-	2.1631	0.031	.028017	.568547

Treatment effect estimates. Outcome: ln_ing. Running variable: ln_riq. Treatment Status: tiempo.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	.54976	.54278	1.0129	0.311	-.514072	1.61358
Robust	-	-	0.8430	0.399	-.684963	1.7188

Estimates adjusted for mass points in the running variable.



Parte 2

Estimación de la curva difusa RD

```
. rdrobust ln_ing ln_riq, fuzzy(tiempo) deriv(1)
```

Fuzzy Kink RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
Number of obs	198	492	BW type =	mserd
Eff. Number of obs	66	99	Kernel =	Triangular
Order est. (p)	2	2	VCE method =	NN
Order bias (q)	3	3		
BW est. (h)	0.544	0.544		
BW bias (b)	0.863	0.863		
rho (h/b)	0.631	0.631		
Unique obs	87	217		

First-stage estimates. Outcome: tiempo. Running variable: ln_riq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	4.0625	2.4261	1.6745	0.094	-.692671	8.81758
Robust	-	-	1.5962	0.110	-1.15849	11.3267

Treatment effect estimates. Outcome: ln_ing. Running variable: ln_riq. Treatment Status: tiempo.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Conventional	-1.3857	.99338	-1.3949	0.163	-3.33265	.561326
Robust	-	-	-1.0312	0.302	-3.90184	1.21153

Estimates adjusted for mass points in the running variable.

Generación de tablas con outreg2

La salida generada por rdrobust puede utilizarse para construir tablas con los paquetes de Stata disponibles, ya que las estimaciones principales y los parámetros elegidos se almacenan como resultados de e(). Esto puede facilitar la creación de tablas y el manejo de resultados numéricos después de la estimación, pero no debe utilizarse para realizar pruebas de hipótesis conjuntas o de otro tipo entre los coeficientes de las matrices almacenadas. Para ilustrar cómo producir tablas de salida con outreg2, consideramos las siguientes dos especificaciones de estimación:

```
. rdrobust ln_ing ln_riq, kernel(uniform)
```

Sharp RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
			BW type =	mserd
			Kernel =	Triangular
			VCE method =	NN
Number of obs	198	492		
Eff. Number of obs	67	106		
Order est. (p)	1	1		
Order bias (q)	2	2		
BW est. (h)	0.570	0.570		
BW bias (b)	0.959	0.959		
rho (h/b)	0.595	0.595		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_rfq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
Conventional	.09109	.12409	0.7341	0.463	-.152116 .3342

. rdrobust ln_ing ln_rfq, p(2) q(3)

Sharp RD estimates using local polynomial regression.

Cutoff c = 0	Left of c	Right of c	Number of obs =	690
			BW type =	mserd
			Kernel =	Triangular
			VCE method =	NN
Number of obs	198	492		
Eff. Number of obs	64	99		
Order est. (p)	2	2		
Order bias (q)	3	3		
BW est. (h)	0.528	0.528		
BW bias (b)	0.786	0.786		
rho (h/b)	0.672	0.672		
Unique obs	87	217		

Outcome: ln_ing. Running variable: ln_rfq.

Method	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
Conventional	.3478	.18786	1.8513	0.064	-.020406 .715999
Robust	-	-	1.7523	0.080	-.044365 .792987

Estimates adjusted for mass points in the running variable.

Debe instalar outreg2

```
. ssc install outreg2
checking outreg2 consistency and verifying not already installed...
installing into C:\Users\ACER SWIFT\ado\plus\...
installation complete.
```

Para utilizar outreg2 de manera correcta debe aplicar este comando después de ejecutar una regresión, porque crea una tabla con los resultados de la regresión. Primero, debe especificar una ubicación para guardar el archivo. Hay dos opciones: puede especificar una ruta en su ordenador o en una



Parte 2

unidad de memoria. Esa ubicación de guardado es parte del comando, que puedes ver a continuación: *Ruta de ordenador*

```
. [outreg2 using C:\Users\kayliec\auto, word replace]
. [outreg2 using C:\Users\kayliec\auto, word append]
```

Unidad de memoria

```
. [ outreg2 using E:\auto, word replace]
. [ outreg2 using E:\auto, word append]
```



13. Bibliografía

1. ———. (1984). “Panel Data.” In *Handbook of Econometrics*, Volume 2, ed. Zvi Griliches and Michael D. Intriligator, 1247–318. Amsterdam: North-Holland.
2. ———. (2005). “Structural Equations, Treatment Effects, and Econometric Policy Evaluation”, *Econometrica*, vol. 73, no. 3, pp. 669–738.
3. ———. (2006). “Using a Social Experiment to Validate a Dynamic Behavioral Model of Child Schooling and Fertility: Assessing the Impact of School Subsidy Program in Mexico”, *American Economic Review*, vol. 96, no. 5, pp. 1384–417.
4. ———. (2008). “Evaluating Anti-Poverty Programs”, In *Handbook of Development Economics*, vol. 4, pp. 3787–846.
5. ———. 1979. “Sample Selection Bias as a Specification Error”, *Econometrica*, vol. 47, no. 1, pp. 153–61.
6. ———. 2001. “Micro Data, Heterogeneity, and the Evaluation of Public Policy: Nobel Lecture”, *Journal of Political Economy*, vol. 109, no. 4, pp. 673–748.
7. Abadie, A., Angrist, J. & Guido W. (2002). “Instrumental Variables Estimates of the Effect of Subsidized Training on the Quantiles of Trainee Earnings”, *Econometrica*, vol. 70, no. 1, pp. 91–117.
8. Abrevaya, J. & Dahl, C. (2008). “The Effects of Birth Inputs on Birthweight: Evidence from Quantile Estimation on Panel Data”, *Journal of Business and Economic Statistics*, vol. 26, no. 4, pp. 379–97.
9. Al-Sharafi, M. *et al.* (2022). “Evaluating the sustainable use of mobile payment contactless technologies within and beyond the COVID-19 pandemic using a hybrid SEM-ANN approach”, *International Journal of Bank Marketing*, vol. 40, no. 5, pp. 1071–1095.
10. Angrist, J. *et al.* (2002). “Vouchers for Private Schooling in Colombia: Evidence from a Randomized Natural Experiment”, *American Economic Review*, vol. 92, no. 5, pp. 1535–58.
11. Apunyo, R. *et al.* (2022). “Interventions to increase youth employment: An evidence and gap map”, *Campbell Systematic Reviews*, vol. 18, no. 1.
12. Araujo, M. *et al.* (2008). “Local Inequality and Project Choice: Theory and Evidence from Ecuador”, *Journal of Public Economics*, vol. 92, no. 5–6, pp. 1022–46.



13. Asadullah, M. & Ara, J. (2016). "Evaluating the long-run impact of an innovative anti-poverty programme: evidence using household panel data", *Applied Economics*, vol. 48, no. 2, pp. 107-120.
14. Athey, S. & Imbens, G. (2006). "Identification and Inference in Nonlinear Difference-inDifferences Models", *Econometrica*, vol. 74, no. 2, pp. 431-97.
15. Banerjee, A. *et al.* (2007). "Remedying Education: Evidence from Two Randomized Experiments in India", *Quarterly Journal of Economics*, vol. 122, no. 3, pp. 1235-64.
16. Banerjee, S., Avjeet, S. & Hussain, S. (2009). *Developing Monitoring and Evaluation Frameworks for Rural Electrification Projects: A Case Study from Nepal*. Washington, DC: World Bank.
17. Behrman, J. & Hoddinott, J. (2005). "Programme Evaluation with Unobserved Heterogeneity and Selective Implementation: The Mexican 'PROGRESA' Impact on Child Nutrition", *Oxford Bulletin of Economics and Statistics*, vol. 67, no. 4, pp. 547-69.
18. Behrman, J. Parker, S. & Todd, P. (2009). "Long-Term Impacts of the Oportunidades Conditional Cash-Transfer Program on Rural Youth in Mexico", In *Poverty, Inequality, and Policy in Latin America*, ed. Stephan Klasen and Felicitas Nowak-Lehmann, 219-70. Cambridge, MA: MIT Press.
19. Bell, B., Blundell, R. and Reenen, J. (1999). "Getting the Unemployed Back to Work: An Evaluation of the New Deal Proposals", *International Tax and Public Finance*, vol. 6, no. 3, pp. 339-60.
20. Bertrand, M., Duflo, E. & Mullainathan, S. (2004). "How Much Should We Trust Differences-in-Differences Estimates?", *Quarterly Journal of Economics*, vol. 119 no. 1, pp. 249-75.
21. Bitler, M. *et al.* (2008). "Distributional Impacts of the Self-Sufficiency Project", *Journal of Public Economics*, vol. 92, no. 3-4, pp. 748-65.
22. Björklund, A. & Moffi, R. (1987). "The Estimation of Wage Gains and Welfare Gains in Self-Selection Models", *Review of Economics and Statistics*, vol. 69, no. 1, pp. 42-49.
23. Blundell, R. & Costa M. (2000). "Evaluation Methods for Non-experimental Data", *Fiscal Studies*, vol. 21, no. 4, pp. 427-68.
24. Bondonio, D. (2022). "Does the Running Variable Matter? A Second Look at Discontinuity Designs for Evaluating Regional Economic Development and Business Incentive Policies", *Evaluation Review*.
25. Bossman, A., Umar, Z. & Teplova, T. (2022). "Modelling the asymmetric effect of COVID-19 on REIT returns: A quantile-on-quantile regression análisis", *Journal of Economic Asymmetries*, vol. 26.



26. Bourguignon, *et al.* (2003). "Conditional Cash Transfers, Schooling, and Child Labor: Micro-simulating Brazil's Bolsa Escola Program", *World Bank Economic Review*, vol. 17, no. 2, pp. 229–54.
27. Bourguignon, F. & Francisco H. (2003). In *The Impact of Economic Policies on Poverty and Income Distribution: Evaluation Techniques and Tools*. Washington, DC: World Bank and Oxford University Press.
28. Bourguignon, F. & Ferreira, F. (2003). "Ex Ante Evaluation of Policy Reforms Using Behavioral Models." In *The Impact of Economic Policies on Poverty and Income Distribution: Evaluation Techniques and Tools*, ed. François Bourguignon and Luiz A. Pereira da Silva, 123–41. Washington, DC: World Bank and Oxford University Press.
29. Buchinsky, M. (1998). "Recent Advances in Quantile Regression Models: A Practical Guide for Empirical Research", *Journal of Human Resources*, vol. 33, no. 1, pp. 88–126.
30. Buddelmeyer, H. & Skoufias, E. (2004). "An Evaluation of the Performance of Regression Discontinuity Design on PROGRESA." Policy Research Working Paper 3386, World Bank, Washington, DC.
31. Buzeti, J. (2022). "The connection between leader behaviour and employee sickness absence in public administration", *International Journal of Organizational Analysis*, vol. 30, no. 7, pp.1-19
32. Campbell, C. & Gaddis, S. (2017). "I Don't Agree with Giving Cash": A Survey Experiment Examining Support for Public Assistance", *Social Science Quarterly*, vol. 98, no. 5, pp. 1352-1373.
33. Chamberlain, G. (1982). "Multivariate Regression Models for Panel Data", *Journal of Econometrics*, vol. 18, no. 1, pp. 5–46.
34. Chaudhury, N. & Parajuli, D. (2006). "Conditional Cash Transfers and Female Schooling: The Impact of the Female School Stipend Program on Public School Enrollments in Punjab, Pakistan." Policy Research Working Paper 4102, World Bank, Washington, DC.
35. Chen, S. & Ravallion, M. (2004). "Welfare Impacts of China's Accession to the World Trade Organization", *World Bank Economic Review*, vol. 18, no. 1, pp. 29–57.
36. Dammert, A. (2007). *Heterogeneous Impacts of Conditional Cash Transfers: Evidence from Nicaragua*. Working Paper, McMaster University, Hamilton, ON, Canada.
37. De Janvry, A. *et al.* (2006). "Can Conditional Cash Transfer Programs Serve as Safety Nets in Keeping Children at School and from Working When Exposed to Shocks?", *Journal of Development Economics*, vol. 79, no. 2, pp. 349–73.



38. Dixit, A. (1990). *Optimization in Economic Theory*. Oxford, U.K.: Oxford University Press.
39. Djebbari, H. and Smith J. (2008). "Heterogeneous Impacts in PROGRESA." IZA Discussion Paper 3362, Institute for the Study of Labor, Bonn, Germany.
40. Dos Santos, G. & Raupp, F. (2015). "Monitoring and evaluation results of government programs outlined in Multiyear Plan", *Revista de Administración Pública*, vol. 49, no. 6, pp. 1429-1451.
41. Duflo, E. (2003). "Grandmothers and Granddaughters: Old Age Pension and Intrahousehold Allocation in South Africa", *World Bank Economic Review*, vol. 17, no. 1, pp. 1-26.
42. Duflo, E., Glennerster R., & Kremer, M. (2008). "Using Randomization in Development Economics Research: A Toolkit", In *Handbook of Development Economics*, vol. 4, pp. 3895-962.
43. Emran, M. Robano, V. & Smith, S. (2009). "Assessing the Frontiers of UltraPoverty Reduction: Evidence from CFPR/TUP, an Innovative Program in Bangladesh." Working Paper, George Washington University, Washington, DC.
44. Essama-Nssah, B. (1997). "Impact of Growth and Distribution on Poverty in Madagascar", *Review of Income and Wealth*, vol. 43, no. 2, pp. 239-52.
45. Essama-Nssah, B. (2005). "Simulating the Poverty Impact of Macroeconomic Shocks and Policies." Policy Research Working Paper 3788, World Bank, Washington, DC.
46. Fedorova, E., Druchok, S. & Drogovoz, P. (2022). "Impact of news sentiment and topics on IPO underpricing: US evidence", *International Journal of Accounting and Information Management*, vol. 30, no. 1, pp. 73-94.
47. Galasso, E. & Ravallion, M. (2004). "Social Protection in a Crisis: Argentina's Plan Jefes y Jefas", *World Bank Economic Review*, vol. 18, no. 3, pp. 367-400.
48. Galasso, E. & Umapathi, N. (2009). "Improving Nutritional Status through Behavioral Change: Lessons from Madagascar", *Journal of Development Effectiveness*, vol. 1, no. 1, pp. 60-85.
49. Gertler, P. (2004). "Do Conditional Cash Transfers Improve Child Health? Evidence from PROGRESA's Control Randomized Experiment", *American Economic Review, Papers and Proceedings*, vol. 94, no. 2, pp. 336-41.
50. Giné, X., Karlan D. & Zinman J. (2008). "The Risk of Asking: Measurement Effects from a Baseline Survey in an Insurance Takeup Experiment" Working Paper, Yale University, New Haven, CT.



51. Giovanis, E., Ozdamar, O. & Ozdas, B. (2021). "The effect of unemployment benefits on health and living standards in Turkey: evidence from structural equation modelling and regression discontinuity design", *International Journal of Manpower*.
52. Gruber, J. (1994). "The Incidence of Mandated Maternity Benefits", *American Economic Review*, vol. 84, no. 3, pp. 622–41.
53. Gugerty, M. & Kremer, M. (2008). "Outside Funding and the Dynamics of Participation in Community Associations", *American Journal of Political Science*, vol. 52, no. 3, pp. 585–602.
54. Hahn, J., Todd, P. & Klaauw, W. (2001). "Identification of Treatment Effects by Regression Discontinuity Design", *Econometrica*, vol. 69, no. 1, pp. 201–9.
55. Heckman, J. & Vytlacil, E. (2000). "Local Instrumental Variables." NBER Technical Working Paper 252, National Bureau of Economic Research, Cambridge, MA.
56. Heckman, J. & Vytlacil, E. (2005). "Structural Equations, Treatment Effects, and Econometric Policy Evaluation", *Econometrica*, vol. 73, no. 3, pp. 669–738.
57. Heckman, J. (1974). "Shadow Prices, Market Wages, and Labor Supply", *Econometrica*, vol. 42, no. 4, pp. 679–94.
58. Heckman, J., & Vytlacil, E. (2005). "Structural Equations, Treatment Effects, and Econometric Policy Evaluation", *Econometrica*, vol. 73, no. 3, pp. 669–738.
59. Heckman, J., Smith, J. & Clements, N. (1997). "Making the Most out of Programme Evaluations and Social Experiments: Accounting for Heterogeneity in Programme Impacts", *Review of Economic Studies*, vol. 64, no. 4, pp. 487–535.
60. Hernandez, R. & Hornero, A. (2021). "Monitoring and assessment of desertification using remote sensing", *Ecosistemas*, vol. 30, no. 4, pp. 2240.
61. Hirano, K. *et al.* (2003). "Efficient Estimation of Average Treatment Effects Using the Estimated Propensity Score", *Econometrica*, vol. 71, no. 4, pp. 1161–89.
62. Hoddinott, J. & Skoufi, E. (2004). "The Impact of PROGRESA on Food Consumption", *Economic Development and Cultural Change*, vol. 53, no. 1, pp. 37–61.
63. Hu, S. (2022). "Research on Monitoring System of Daily Statistical Indexes Through Big Data", *Recent Advances in Computer Science and Communications*, vol. 15, no. 5, pp. 731-740.



64. Ianchovichina, E. & Martin, W. (2004). "Impacts of China's Accession to the World Trade Organization", *World Bank Economic Review*, vol. 18, no. 1, pp. 3-27.
65. Imbens, G. & Angrist, J. (1994). "Identification and Estimation of Local Average Treatment Effects", *Econometrica*, vol. 62, no. 2, pp. 467-76.
66. Jacoby, H. (2002). "Is There an Intrahousehold 'Flypaper Effect'? Evidence from a School Feeding Programme", *Economic Journal*, vol. 112, no. 476, pp. 196-221.
67. Jalan, J. & Ravallion, M. (1998). "Are There Dynamic Gains from a Poor-Area Development Program?", *Journal of Public Economics*, vol. 67, no. 1, pp. 65-85.
68. Jalan, J. & Ravallion, M. (2003). "Estimating the Benefit Incidence for an Antipoverty Program by Propensity-Score Matching", *Journal of Business and Economic Statistics*, vol. 21, no. 1, pp.19-30.
69. Jiao, S. & Sun, Q. (2021). "Digital economic development and its impact on economic growth in china: Research based on the perspective of sustainability", *Sustainability (Switzerland)*, vol. 13, no. 18.
70. Johnsen, J. & Reiso, K. (2020). "Economic Effects of Workfare Reforms for Single Mothers: Benefit Substitution and Labour Supply Responses", *Scandinavian Journal of Economics*, vol. 122, no. 2, pp. 494-523.
71. Kct, S., Devanaboyina, V. & Shaik, T. (2022). "Double difference method with zero and short base length carrier phase measurements for Navigation with Indian Constellation satellites L5 (1176.45 MHz) signal quality analysis", *International Journal of Satellite Communications and Networking*, vol. 40, no. 4, pp. 294-304.
72. Khandker, S., Bakht, Z. & Koolwal, G. (2009). "The Poverty Impacts of Rural Roads: Evidence from Bangladesh", *Economic Development and Cultural Change*, vol. 57, no. 4, pp. 685-722.
73. King, E. & Behrman, J. (2009). "Timing and Duration of Exposure in Evaluations of Social Programs", *World Bank Research Observer*, vol. 24, no. 1, pp. 55-82.
74. Kish, L. (1987). *Statistical Design for Research*. New York: Wiley.
75. Koenker, R. & Gilbert B. (1978). "Regression Quantiles", *Econometrica*, vol. 46, no. 1, pp. 33-50.
76. Komro, A. *et al.* (2022). "Study protocol for a cluster randomized trial of a school, family, and community intervention for preventing drug misuse among older adolescents in the Cherokee Nation", *Trials*, vol. 23, no. 1.



77. Lanjouw, P. (2003). "Estimating Geographically Disaggregated Welfare Levels and Changes." In *The Impact of Economic Policies on Poverty and Income Distribution: Evaluation Techniques and Tools*, ed. François Bourguignon and Luiz A. Pereira da Silva, 85–102. Washington, DC: World Bank and Oxford University Press.
78. Larsen, A., Meng, K. & Kendall, B. (2019). "Causal analysis in control–impact ecological studies with observational data", *Methods in Ecology and Evolution*, vol. 10, no. 7, pp. 924–934.
79. Lechner, M. (1999). "Earnings and Employment Effects of Continuous Off-the-Job Training in East Germany after Unification", *Journal of Business Economic Statistics*, vol. 17, no. 1, pp. 74–90.
80. Limphaibool, W. *et al.* (2021). "Collective Mindfulness and Organizational Resilience in Times of Crisis", *Change Management*, vol. 22, no. 2, pp. 49–60
81. Lokshin, M. & Ravallion, M. (2004). "Gainers and Losers from Trade Reform in Morocco" Policy Research Working Paper 3368, World Bank, Washington, DC.
82. Lu, S. & Wang, H. (2022). "Market-oriented reform and land use efficiency: Evidence from a regression discontinuity design", *Land Use Policy*, vol. 115.
83. *Macroeconomic Policies on Poverty and Income Distribution: Macro-Micro Evaluation Techniques and Tools*. Washington, DC: World Bank and Palgrave Macmillan.
84. Maddala, G. (1986). *Limited-Dependent and Qualitative Variables in Econometrics*. Cambridge, U.K.: Cambridge University Press.
85. Mansuri, G. & Rao, V. (2004). "Community-Based and -Driven Development: A Critical Review", *World Bank Research Observer*, vol. 19, no. 1, pp. 1–39.
86. Masini, R. & Medeiros, M. (2022). "Counterfactual Analysis and Inference With Nonstationary Data", *Journal of Business and Economic Statistics*, vol. 40, no. 1, pp. 227–239.
87. Mendez, M., Everaert, P. & Valcke, M. (2020). "Is the dosed video-vignettes intervention more effective with a longer-lasting effect? A financial literacy Study", *ACM International Conference Proceeding Series*, pp. 193–198.
88. Menke, W. & Creel, R. (2022). "Why Differential Data Work", *Bulletin of the Seismological Society of America*, vol. 112, no. 2, pp. 597–607.
89. Mi, Y., Chen, X. & He, Q. (2022). "Evaluating the effect of China's targeted poverty alleviation policy on credit default: a regression discontinuity approach", *Applied Economics*, vol. 54, no. 31, pp. 3654–3667.
90. Miguel, E. & Kremer, M. (2004). "Worms: Identifying Impacts on Education and Health in the Presence of Treatment Externalities", *Econometrica*, vol. 72, no. 1, pp. 159–217.



91. Moffitt, R. (2003). The Role of Randomized Field Trials in Social Science Research: A Perspective from Evaluations of Reforms from Social Welfare Programs. NBER Technical Working Paper 295, National Bureau of Economic Research, Cambridge, MA.
92. Mutonyi, B. *et al.* (2022). "The impact of organizational culture and leadership climate on organizational attractiveness and innovative behavior: a study of Norwegian hospital employees", BMC Health Services Research, vol. 22, no. 1.
93. Na, Q. & Kim, Y. (2022). "Does the Two-Child Policy Affect the Consumption Upgrade of "Two-Child Families": Based on CFPS Data from the Chinese Family Tracking Survey", Journal of Global Business and Trade, vol. 18, no. 1, pp 97-122.
94. Noonan, J. & Zhigljavsky, A. (2022), "Random and quasi-random designs in group testing", Journal of Statistical Planning and Inference, vol. 221, pp. 29-54.
95. Paxson, C. & Schady N. (2002). "The Allocation and Impact of Social Funds: Spending on School Infrastructure in Peru", World Bank Economic Review, vol. 16, no. 2, pp. 297-319.
96. Pigini, C. & Bartolucci, F. (2022). "Conditional inference for binary panel data models with predetermined covariates", Econometrics and Statistics, vol. 23, pp. 83-104.
97. Pinto, A. *et al.* (2021). "Exploring different methods to evaluate the impact of basic income interventions: a systematic review", International Journal for Equity in Health, vol. 20, no. 1.
98. Piper, A. (2022). "Optimism, pessimism and life satisfaction: an empirical investigation", International Review of Economics, vol. 69, no. 2, pp. 177-208.
99. Platteau, J. (2004). "Monitoring Elite Capture in Community-Driven Development", Development and Change, vol. 35, no. 2, pp. 223-46.
100. Quandt, R. (1972). "Methods for Estimating Switching Regressions." Journal of the American Statistical Association, vol. 67, no. 338, pp. 306-10.
101. Rabaté, S. & Rochut, J. (2020). "Employment and substitution effects of raising the statutory retirement age in France", Journal of Pension Economics and Finance, vol. 19, no. 3, pp. 293-308.
102. Rafiq, R. & McNally, M. (2022). "A structural analysis of the work tour behavior of transit commuters", Transportation Research Part A: Policy and Practice, vol. 160, pp. 61-79.
103. Ramirez, A. *et al.* (2022). "Relationship between the Cost of Capital and Environmental, Social, and Governance Scores: Evidence from Latin America", Sustainability (Switzerland), vol. 14, no. 9.



104. Ramos, M. & Prats, M. (2020). "Fiscal sustainability in the European countries: A panel ardl approach and a dynamic panel threshold model", *Sustainability (Switzerland)*, vol. 12, no. 20, pp. 1-14.
105. Rao, V. & Ibáñez, A. (2005). "The Social Impact of Social Funds in Jamaica: A 'Participatory Econometric' Analysis of Targeting, Collective Action, and Participation in Community-Driven Development", *Journal of Development Studies*, vol. 41, no. 5, pp. 788-838.
106. Ravallion, M. & Chen, S. (2007). "China's (Uneven) Progress against Poverty", *Journal of Development Economics*, vol. 82, no. 1, pp. 1-42.
107. Ravallion, M. (2003). "Assessing the Poverty Impact of an Assigned Program." In *The Impact of Economic Policies on Poverty and Income Distribution: Evaluation Techniques and Tools*, ed. François Bourguignon and Luiz A. Pereira da Silva, 103-22. Washington, DC: World Bank and Oxford University Press.
108. Ravallion, M. (2008). "Evaluating Anti-Poverty Programs", In *Handbook of Development Economics*, vol. 4, pp. 3787-846.
109. Raza, M. (2022). "Towards a sustainable development: econometric analysis of energy use, economic factors, and CO2 emission in Pakistan during 1975-2018", *Environmental Monitoring and Assessment*, vol. 194, no. 2.
110. Rosenbaum, P. & Rubin, D. (1983). "The Central Role of the Propensity Score in Observational Studies for Causal Effects", *Biometrika*, vol. 70, no. 1, pp. 41-55.
111. Roy, A. (1951). "Some Thoughts on the Distribution of Earnings", *Oxford Economic Papers*, vol. 3, no. 2, pp. 135-46.
112. Saldivar, M. *et al.* (2022). "Effect of a conditional cash transference program on food insecurity in Mexican households: 2012-2016", *Public Health Nutrition*, vol. 25, no. 4, pp. 1084-1093.
113. Schady, N. (1999). "Seeking Votes: The Political Economy of Expenditures by the Peruvian Social Fund (FONCODES), 1991-95." Policy Research Working Paper 2166, World Bank, Washington, DC.
114. Schultz, T. (2004). "School Subsidies for the Poor: Evaluating the Mexican PROGRESA Poverty Program", *Journal of Development Economics*, vol. 74, no. 1, pp. 199-250.
115. Seth, S. & Tutor, M. (2021). "Evaluation of Anti-Poverty Programs Impact on Joint Disadvantages: Insights From the Philippine Experience", *Review of Income and Wealth*, vol. 67, no. 4, pp. 977-1004.



116. Seúl, C. & Booyuel, K. (2016). "A beginner's guide to randomized evaluations in development economics", *Seoul Journal of Economics*, vol. 29, no. 4, pp. 529 – 552.
117. Sharma, R. & Aggarwal, P. (2022). "Impact of mandatory corporate social responsibility on corporate financial performance: the Indian experience", *Social Responsibility Journal*, vol. 18, no. 4, pp. 704-722.
118. Silverman, K., Holtyn, A. & Jarvis, B. (2016). "A potential role of anti-poverty programs in health promotion", *Preventive Medicine*, vol. 92, pp. 58-61.
119. Skoufias, E. and Maro, V. (2007). "Conditional Cash Transfers, Adult Work Incentives, and Poverty." Policy Research Working Paper 3973, World Bank, Washington, DC.
120. Soliman, A. (2022). "Prescriptive drought policy and water supplier compliance", *Ecological Economics*, vol. 197.
121. Steinmann, I. & Olsen, R. (2022). "Equal opportunities for all? Analyzing within-country variation in school effectiveness", *Large-Scale Assessments in Education*, vol. 10, no. 1.
122. Tebani, M. & Mederbal, K. (2018). "Monitoring and evaluation of the agricultural and rural renewal program in Algeria: Case of the ouarsenis área", *Revista de Economia e Sociologia Rural*, vol. 56, no. 4, pp. 719-728.
123. Todd, P. & Wolpin, K. (2006). "Ex Ante Evaluation of Social Programs." PIER Working Paper 06-122, Penn Institute for Economic Research, University of Pennsylvania, Philadelphia.
124. Todd, P. & Wolpin, K. (2006). "Assessing the Impact of a School Subsidy Program in Mexico: Using a Social Experiment to Validate a Dynamic Behavioral Model of Child Schooling and Fertility", *American Economic Review*, vol. 96, no. 5, pp. 1384–417.
125. Tolentino, E. (2022). "An evaluation of a mandatory profit-sharing reform in Peru, using quasi-experimental methods", *Journal of Industrial and Business Economics*, vol. 49, no. 2, pp. 313-334.
126. Tong, K. *et al.* (2020). "Aplicación del método difuso de toma de decisiones para la evaluación de programas y políticas de gestión de la educación superior vietnamita", *Revista de finanzas, economía y negocios asiáticos*, vol. 7, no. 9, pp. 719-726.
127. Van, D. (2009). "Impact Evaluation of Rural Road Projects", *Journal of Development Effectiveness*, vol. 1, no. 1, pp. 15–36.



128. Von, A. *et al.* (2022). "Application of a health technology assessment framework to digital health technologies that manage chronic disease: A systematic review", *International Journal of Technology Assessment in Health Care*, vol. 38, no. 1.
129. Wang, X. *et al.* (2021). "Analysis of policies based on the multi-fuzzy regression discontinuity, in terms of the number of deaths in the coronavirus epidemic", *Healthcare (Switzerland)*, vol. 9, no. 2.
130. Williams, C. & Gashi, A. (2022). "Evaluating the wage differential between the formal and informal economy: a gender perspective", *Journal of Economic Studies*, vol. 49, no. 4, pp. 735-750.
131. Wooldridge, J. (2001). *Econometric Analysis of Cross Section and Panel Data*. Cambridge, MA: MIT Press.
132. Wunsch, G. & Gourbin, C. (2020). "Causal assessment in demographic research", *Genus*, vol. 76, no. 1.
133. Zhang, Q., Liu, H. & Wan, Z. (2022). "Evaluation on the effectiveness of ship emission control area policy: Heterogeneity detection with the regression discontinuity method", *Environmental Impact Assessment Review*, vol. 94.
134. Zhu, Y. & Savla, K. (2022). "Information Design in Non-atomic Routing Games with Partial Participation: Computation and Properties", *IEEE Transactions on Control of Network Systems*.