

INTRODUCCIÓN A LA EPIDEMIOLOGÍA

Rebeca Millán Guerrero

Benjamín Trujillo Hernández

José Ramiro Caballero Hoyos



$$n = \frac{(3.84 \times 36)}{9} = \frac{138.2}{9} = 15$$

$$n = \frac{(3.84) \times (2275)}{(6.5)^2} = 842.$$

$$n = \frac{n z^2 (pq)}{e^2 (n-1) + z^2 (q)}$$

$$n = \frac{(3.84) \times (2275)}{(6.5)^2} = \frac{8,736}{42.25} = 206$$

UNIVERSIDAD DE COLIMA

INTRODUCCIÓN A LA
EPIDEMIOLOGÍA
CLÍNICA Y ESTADÍSTICA

UNIVERSIDAD DE COLIMA

Mtro. José Eduardo Hernández Nava, Rector

Mtro. Christian Torres Ortiz Zermelo, Secretario General

Licda. Ma. Guadalupe Carrillo Cárdenas, Coordinadora General de Comunicación Social

Mtra. Gloria Guillermina Araiza Torres, Directora General de Publicaciones

INTRODUCCIÓN A LA
EPIDEMIOLOGÍA
CLÍNICA Y ESTADÍSTICA

Rebeca Millán Guerrero
Benjamín Trujillo Hernández
José Ramiro Caballero Hoyos



UNIVERSIDAD DE COLIMA

© UNIVERSIDAD DE COLIMA, 2015
Avenida Universidad 333
Colima, Colima, México
Dirección General de Publicaciones
Teléfonos: (312) 31 61081 y 31 61000, ext. 35004
Correo electrónico: publicaciones@ucol.mx
www.ucol.mx

ISBN: 978-607-8356-43-0
Derechos reservados conforme a la ley
Impreso en México / *Printed in Mexico*

Proceso editorial certificado con normas Iso desde 2005
Dictaminación y edición registradas en el Sistema Editorial Electrónico PRED
Registro: LI-021-13
Recibido: Noviembre de 2013
Publicado: Mayo de 2015

Índice

Prólogo	7
Introducción	9
CAPÍTULO I	
Epidemiología y estadística clínica	15
Epidemiología clínica	17
Comprensión de la literatura médica	19
Clasificación de la epidemiología	20
<i>Benjamín Trujillo Hernández</i>	
CAPÍTULO II	
El método científico y su aplicación en la investigación	23
Introducción	23
Aplicación del método en el proceso de investigación	25
Fases deductivas	26
Fases inductivas	30
<i>José Ramiro Caballero Hoyos</i>	
CAPÍTULO III	
Cómo elaborar un proyecto de investigación	41
Introducción	41
Estructura de un proyecto	42
<i>Rebeca Millán Guerrero</i>	
CAPÍTULO IV	
Riesgo, causalidad y mediciones de riesgo	51
Mediciones de riesgo	53
Pruebas diagnósticas	58
<i>Benjamín Trujillo Hernández</i>	

CAPÍTULO V

Diseños de estudio	65
Introducción	65
Tipos de diseño	67
<i>Rebeca Millán Guerrero</i>	

CAPÍTULO VI

Variables y escalas de medición	81
Introducción	81
Variable	82
Nociones de estadística	86
<i>Benjamín Trujillo Hernández</i>	

CAPÍTULO VII

Muestreo y tamaño de muestra	91
Conceptos generales	91
Tipos de muestreo	92
<i>Benjamín Trujillo Hernández</i>	

CAPÍTULO VIII

Cómo elegir la prueba estadística	105
Introducción	105
Elección de la prueba estadística adecuada	105
Otras pruebas estadísticas	111
<i>Benjamín Trujillo Hernández</i>	

CAPÍTULO IX

Análisis de datos con el paquete estadístico SPSS	113
Introducción	113
Introducción al análisis estadístico con el paquete SPSS	120
Síntesis	249
<i>José Ramiro Caballero Hoyos</i>	

CAPÍTULO X

Medicina basada en evidencias aplicada a la neurología	255
La medicina basada en evidencias	255
La utilidad de la MBE en la neurología	256
Diferentes escalas de gradación para medir la MBE	259
<i>Rebeca Millán Guerrero</i>	

Prólogo

He tenido el emocionante privilegio de ser testigo de la gestación y desarrollo de la investigación científica en el Instituto Mexicano del Seguro Social (IMSS) de Colima. El trabajo de investigación científica se vio cristalizado con la creación de la Unidad de Investigación en Epidemiología Clínica en el Hospital General de Zona N° 1 en la ciudad de Colima y con otra Unidad en el municipio de Cuauhtémoc, Colima. Tres investigadores forman esta Unidad, la Dra. Rebeca Millán Guerrero, el Dr. Benjamín Trujillo Hernández y el Dr. José Ramiro Caballero Hoyos cuya productividad queda de manifiesto y son investigadores nivel 2 del Sistema Nacional de Investigadores (SNI).

Perdura en mí el recuerdo de su desempeño en la cátedra que impartían como profesores de las asignaturas de Epidemiología Clínica y Estadística en la maestría en Ciencias Médicas que se imparte en la Universidad de Colima. También me consta cómo comunican su entusiasmo a los jóvenes estudiantes que han estado en contacto académico con ellos. Toda esta experiencia y conocimientos que poseen, por fortuna, lo han cristalizado en un excelente libro: *Introducción a la epidemiología clínica y estadística*.

Este libro llena un hueco académico en el estado de Colima y su región, para fortalecer la formación de los estudiantes que se dedicarán a la investigación clínica y a la medicina asistencial ya que, tanto la epidemiología clínica y la estadística, son útiles y necesarias para la formación de un investigador clínico en este campo. Asimismo, tiene enorme utilidad para los estudiantes de posgrado y pregrado en ciencias médicas, para los mé-

dicos y especialistas, para las enfermeras y los psicólogos, ya que la incorporación de la investigación y la perspectiva de la investigación científica en la prevención y asistencia en la medicina es un punto de inflexión en el desarrollo de la medicina que se ejerce en Colima.

La estructura del libro es sencilla y accesible. Cada capítulo está escrito por uno de los autores expertos en su disciplina y constituye una herramienta indispensable para el ejercicio de la medicina. Considero que es un texto completo, ameno y dinámico sobre la interpretación de la estadística y de la información epidemiológica para la atención al paciente.

Por último, debo resaltar que los tres autores tienen un profundo cariño y una dedicación completa a las actividades particulares que desarrollan como científicos; en el caso de la Dra. Rebeca Millán y el Dr. Benjamín Trujillo, como médicos y el Dr. José Ramiro Caballero como experto en estudios epidemiológicos. Con este libro se constata que las transformaciones más profundas en el área de la formación que hacen una diferencia en el antes y después, parte de los propios investigadores y su generosidad para invertir tiempo en elaborar la edición de este libro, el cual será un puente entre la investigación de laboratorio y la investigación clínica.

*Dr. Miguel Huerta Viera**

* Centro Universitario de Investigaciones Biomédicas de la Universidad de Colima; ex Coordinador académico, fundador de la Maestría y el Doctorado en Ciencias Médicas y profesor de Seminarios de Investigación.

Introducción

Los métodos de la epidemiología desarrollados para el estudio de poblaciones, como incidencia, prevalencia, factores de riesgo, evolución natural y pronóstico de la enfermedad, se usan para responder a problemas que el investigador encuentra en su práctica diaria.

Como se sabe, la información clínica se basa en datos inciertos y valiosos, pero subjetivos. Las observaciones clínicas las efectúan médicos que emiten sus propios juicios, lo cual propicia la existencia de posibles errores; por lo tanto, es necesario unificar criterios que, basados en los principios científicos, se hagan objetivos y puedan generalizarse.

Por tal razón, un clínico con experiencia y capacidad de resumir un documento de dos volúmenes, a menudo tiene dificultad para presentar un artículo de investigación de cuatro páginas. Del mismo modo que aprendió los pasos para la exploración clínica, deberá aprender los pasos para realizar una investigación con el método científico aceptado y comprender los conceptos básicos de la ciencia.

La presente edición de *Introducción a la epidemiología clínica y estadística* tiene como objetivo proporcionar una herramienta didáctica para el aprendizaje de la investigación en epidemiología clínica. En sus páginas se exponen algunos contenidos teóricos básicos para que el estudiante pueda incorporar las lógicas deductiva

e inductiva del método científico a su práctica cotidiana, a fin de contar con un instrumento que le ayude a responder las preguntas médicas que con frecuencia surgen en el ámbito clínico.

A fin de lograr el objetivo propuesto, los contenidos del libro intentan dar respuesta a preguntas como las que siguen: ¿Qué es la epidemiología clínica? ¿Cómo se relaciona ésta con la estadística para constituir una disciplina científica? ¿Qué es el método y cómo se aplica al proceso de investigación científica? ¿Cómo se traducen las lógicas del método científico en proyectos de investigación que pretenden resolver problemas clínicos concretos? ¿Cómo se relacionan la causalidad y el riesgo? ¿Cuáles son las medidas estadísticas que permiten generar distintas aproximaciones a los factores de riesgo? ¿Cuáles son los principales tipos de diseño en la investigación clínica? ¿Cómo se construyen variables y escalas de medición adecuadas para contrastar hipótesis estadísticas? ¿Qué son las estrategias de muestreo y cómo se calculan tamaños de muestra suficientes para realizar inferencias válidas? ¿Cómo se seleccionan las pruebas estadísticas apropiadas a las variables disponibles en una investigación? ¿Cómo se procesan los análisis estadísticos computarizados? ¿Qué es la medicina basada en evidencias y cuáles son sus principales aplicaciones?

En tal sentido, la estructura del libro consta de diez capítulos organizados bajo una lógica deductiva que responde, primero, a las preguntas teórico-conceptuales y, luego, a las que requieren precisiones técnicas de aplicaciones estadísticas en la planeación y ejecución del diseño de una investigación clínica. A continuación se bosquejan algunos de los aportes de dichos capítulos:

En el capítulo 1, Benjamín Trujillo Hernández presenta una concepción general sobre la relación entre la epidemiología clínica y la estadística. Plantea los antecedentes históricos de la epidemiología como una disciplina intuitiva que, gracias a las técnicas del conteo y la teoría de la probabilidad, devino en el siglo XIX como una disciplina científica de las enfermedades infecto-contagiosas que se amplió, en el siglo XX, a los procesos de salud

y enfermedad. En ese contexto, la epidemiología clínica es una rama de la epidemiología que estudia a los individuos en contextos clínicos con el propósito de elegir los métodos diagnósticos más certeros, determinar la eficacia terapéutica para escoger los mejores tratamientos, pronosticar el curso de las enfermedades y calcular los costos del proceso de atención, entre otros.

En el capítulo II, José Ramiro Caballero Hoyos define el método científico como el conjunto de principios y normas que orientan la planeación e implementación técnica de las distintas fases que integran el proceso de investigación. Sostiene que su historia consistió en un lento aprendizaje cultural del método experimental inductivo que estuvo acompañado de algunos cambios económicos y políticos profundos en las sociedades, lo cual generó desde el siglo XVII una revolución mental sin precedentes en la forma de explicar y resolver los problemas sociales. Describe también la forma en que el método se aplica a las distintas fases del proceso de investigación, bajo los razonamientos lógicos deductivo e inductivo y el cumplimiento de algunos supuestos filosóficos que dan sustento a los procedimientos metodológicos que se siguen.

En el capítulo III, Rebeca Millán Guerrero define qué es un proyecto de investigación y cuáles son sus características. Concibe dicho proyecto como la versión escrita de un plan de estudio factible para el contexto clínico donde se quiera desarrollar, a fin de generar un conocimiento válido y generalizable que contribuya a mejorar la práctica clínica. Señala también que el documento del proyecto incluye los distintos elementos del método científico, al abarcar: la pregunta a investigar, el diseño, los individuos que tomarán parte en él, la asignación, el tamaño de la muestra requerido en el tiempo fijado, el análisis e interpretación, aspectos éticos, etcétera. En la parte medular del texto, la autora plantea una guía de los elementos a incluir cuando se elabora un proyecto de investigación clínica. Ilustra los mismos con ejemplos clínicos que incluyen valiosas sugerencias de aplicación a la práctica.

En el capítulo iv, Benjamín Trujillo Hernández aborda el asunto de la relación entre la causalidad y el riesgo, dos conceptos claves de la epidemiología clínica. Primero, define la causa como la condición que sola o acompañada produce una enfermedad y describe los distintos tipos de causalidad. Luego, define el riesgo como la probabilidad de desarrollar una enfermedad por la exposición a ciertos factores de riesgo físicos, químicos, biológicos, sociales y psicológicos que la preceden o favorecen su ocurrencia. Tales factores pueden tener una asociación estadísticamente significativa con una enfermedad, pero para poder ser causas deben cumplir distintos criterios de causalidad como la demostración experimental, la plausibilidad biológica, la coherencia temporal y otros. Basado en tales definiciones, el autor hace una descripción detallada de las diferentes medidas estadísticas que permiten evaluar la asociación causal entre el riesgo y la enfermedad en el marco de los distintos diseños de investigación.

En el capítulo v, Rebeca Millán Guerrero desarrolla el tema de los diseños de estudio en la epidemiología clínica. Señala que los diseños se clasifican en función de cómo el investigador controle sus variables y les asigne una dirección temporal, lo cual le permite proponer la posibilidad de realizar estudios observacionales, transversales, longitudinales y de intervención con sus aproximaciones específicas a la realidad clínica (casos y controles, cohortes, ensayos clínicos, experimentos, etcétera). La autora enfatiza que la elección de los diseños debe basarse más en la pregunta que se quiera responder en la investigación que en criterios técnicos e instrumentales.

En el capítulo vi, Benjamín Trujillo Hernández introduce el tema referente al manejo de variables y escalas de medición en la investigación clínica. Advierte que la identificación del tipo de variables es un tema que suscita muchas controversias en alumnos y maestros, por las propuestas dispares incluidas en los libros de metodología. Debido a ello, se propone aclarar el asunto mediante una clasificación de variables basada en cuatro característi-

cas: 1) su naturaleza, 2) su nivel de medición, 3) su interrelación y 4) su reversibilidad. En dicha clasificación, distingue distintos tipos de variables y sus formas de medición correspondientes. El autor acompaña la clasificación con varios ejemplos del contexto clínico y con sugerencias prácticas.

En el capítulo VII, el mismo autor, Benjamín Trujillo Hernández, trata el tema del muestreo y el tamaño de la muestra en la investigación clínica. Inicialmente, define los conceptos básicos de la teoría del muestreo y clasifica las estrategias muestrales en probabilísticas y no probabilísticas. Luego, propone cuatro fórmulas para calcular el tamaño de la muestra, según el tipo de variables consideradas (proporción o razón) y el número de grupos estudiados (1 o ≥ 2 grupos). Como complemento, incluye ejemplos de procedimientos de selección y cálculos de tamaño de la muestra con las distintas fórmulas. El autor sugiere cuidar y trabajar en detalle el muestreo y la muestra al diseñar un estudio, porque de ello dependerán la validez de sus inferencias y, en buena medida, sus posibilidades de publicación en revistas científicas serias.

En el capítulo VIII, desarrollado también por Benjamín Trujillo Hernández, se expone el tema de la selección de pruebas estadísticas. Señala que dicha selección debe basarse en la identificación de tres elementos básicos: 1) el nivel de medición de la variable dependiente, 2) el número de grupos de estudio, y 3) el número de mediciones al mismo grupo de individuos. Bajo ese enfoque, propone un algoritmo para la selección de pruebas estadísticas paramétricas y no paramétricas que corresponden a las distintas combinaciones de los elementos identificados. En forma adicional, el autor describe otras pruebas estadísticas (asociación causal, correlación, diagnósticas y reducción de riesgo absoluto).

En el capítulo IX, José Ramiro Caballero Hoyos desarrolla el proceso de análisis de datos con el paquete estadístico SPSS. Describe, primero, los supuestos teóricos del análisis computarizado e ilustra, después, los procesos de análisis estadísticos para

distintos tipos de variables y diferentes niveles de complejidad. Para el efecto, emplea bases de datos nacionales e internacionales con variables enfocadas a las prácticas de riesgo de ITS y VIH, y la prevalencia de VIH y la tasa de incidencia de tuberculosis.

Finalmente, en el capítulo x, Rebeca Millán Guerrero trata el tema de la Medicina Basada en Evidencias (MBE). Define a ésta como la integración de las mejores evidencias de la investigación clínica, con la experiencia clínica y los valores del paciente, a fin de orientar con eficiencia la acción médica de diagnóstico, pronóstico y terapéutica. Propone, luego, los pasos a seguir para desarrollar la MBE con un sentido crítico, así como la necesidad de distinguir entre significancia clínica y significancia estadística al interpretar los resultados de la evidencia encontrada.

Con esta obra, no se pretende hacer un estudio exhaustivo de la relación entre la epidemiología clínica y la estadística; más bien, se quiere proporcionar criterios básicos para ubicar al investigador que inicia en esta disciplina, mediante una propuesta de conexión de los conceptos más relevantes con aplicaciones estadísticas descritas en un lenguaje simple que facilite su comprensión.

En su conjunto, los distintos capítulos fueron elaborados bajo el formato de los ensayos científicos. Con el propósito de facilitar el aprendizaje reflexivo de los contenidos, sus textos incluyen diagramas y cuadros conceptuales o empíricos, referencias bibliográficas, glosarios y finalmente una autoevaluación con la que reforzará lo aprendido.

CAPÍTULO I

Epidemiología y estadística clínica

Benjamín Trujillo Hernández

La palabra *epidemiología* se deriva del griego *epi* (*sobre*), *demos* (*pueblo*) y *logos* (*estudio*). Estrictamente, la epidemiología se encarga del estudio relacionado con el proceso salud-enfermedad respecto a una población. Su función es primordial para resolver cualquier problema de salud. La epidemiología surge como una necesidad de comprender las características y causas de las enfermedades infecciosas o epidemias que azotaron en la antigüedad. Los estragos que provocaban estas infecciones ponían en riesgo la estabilidad económica y social de la población que las padecía; había, pues, la necesidad imperiosa de resolver este problema (López, Corcho, y Moreno, 1999).

El primer escrito documentado sobre el estudio de las enfermedades infectocontagiosas, con un enfoque preventivo, fue un libro chino sobre acupuntura titulado *Nei King* (*Canon de Medicina*), el cual fue escrito en el año 2650 a.C. Hipócrates fue quien utilizó por primera vez los vocablos *epidemia* y *endemia*, para referirse a padecimientos propios o extraños en relación con un determinado lugar. La sabiduría de Hipócrates es sorprendente, ya que en su texto “*Aire, aguas, y lugares*” menciona la influencia del medio ambiente sobre los humanos para producir enfermedades. En 1546, Girolamo Fracastoro fue el primero en

describir las enfermedades catalogadas como contagiosas y en establecer por lo menos tres formas posibles de infección (contacto directo, fómites y respiración del aire). Fracastoro estableció que enfermedades específicas resultan de contagios específicos y por esto es considerado como el padre de la epidemiología moderna (López-Moreno, Garrido-Latorre, y Hernández-Ávila, 2000).

La epidemiología, sin embargo, permaneció durante muchos años como una disciplina debido principalmente a que carecía de la utilización del método científico y métodos de medición. Esto es, cualquier rama del saber que se jacte de ser ciencia debe utilizar la física o matemáticas, que son consideradas las ciencias madres.

Por otra parte, en Europa, hasta el siglo xvi, los conteos poblacionales se utilizaron para determinar el número de individuos aptos para ingresar al ejército y la carga de impuestos. No obstante, esta información no fue suficiente y era necesario conocer en forma más precisa las características de la población o fuerza del Estado (actividad que se llamó *estadística*). La estadística utilizada en ese entonces era descriptiva y se sustentaba en operaciones aritméticas sencillas, como suma, resta, multiplicación y división. La estadística contribuyó a determinar las características de la población, el número de nacimientos, de qué se enfermaba la gente, de qué moría, sus condiciones socioeconómicas y también ayudó a la clasificación de las enfermedades.

La utilización de la estadística produjo un avance espectacular en medicina, ya que en este periodo se establecieron las *leyes de enfermedad* (que es la probabilidad de enfermar o morir, utilizando modelos matemáticos de tipo inferencial) o las famosas *tablas de vida* creadas por Edmund Halley (1656-1742) que han servido de sustento para el pago de primas de las aseguradoras modernas. También, la estadística dio paso a la *observación numérica*, que consiste en la observación de un conjunto de individuos con cierta enfermedad y posteriormente aplicarles métodos estadísticos o matemáticos. Al utilizar la observación numérica se lo-

gró determinar que la etiología del escorbuto se relacionaba con un deficiente consumo de cítricos (James Lind, 1747), que la tuberculosis no se transmitía hereditariamente y que la sangría era inútil y aun perjudicial en la mayoría de los casos (Pierre Charles Alexander Louis, 1830). Es aquí donde se inician los conceptos *factor de riesgo y enfermedad*. La epidemiología científica se inicia con el médico inglés John Snow quien estudió y encontró los factores que causaban la epidemia del cólera que afectó a Londres entre los años 1849 y 1854. Snow sustentó que la causa del cólera era el agua contaminada con heces. Debido a esto, Londres fue la primera ciudad en el mundo que implementó el drenaje y el agua potable. A partir de ahí la epidemiología se tornaría en una ciencia, pero su objetivo principal se centró en el estudio de epidemias o enfermedades infectocontagiosas.

A finales del siglo XIX, en Estados Unidos de Norteamérica, se forma el Servicio de Salud Pública y el Laboratorio de Higiene, este último se convirtió en el principal instituto de investigación epidemiológica del país. Los médicos investigadores de estos institutos se dieron cuenta de que la epidemiología podría utilizarse para resolver cualquier problema de salud y a partir del siglo XX, las herramientas de la epidemiología se utilizaban para el estudio de enfermedades diferentes a las infectocontagiosas, como silicosis, intoxicación por plomo o mercurio, estudio de la pelagra, cáncer, etcétera. En la actualidad, la epidemiología estudia todo lo relacionado con el proceso salud-enfermedad y podemos decir, sin temor a equivocarnos, que se encuentra en la base de todas las ramas de la medicina.

Epidemiología clínica

Este concepto fue utilizado por primera vez en 1938 por John R. Paul, y se define como la utilización de los métodos o estrategias de la epidemiología en el quehacer clínico. Al respecto, es difícil discernir entre epidemiología en general y epidemiología clínica. La epidemiología clínica puede ser utilizada para lo siguiente:

Diagnóstico

Gracias a la epidemiología es posible elegir los métodos diagnósticos más certeros. Los profesionales de la salud y las instituciones de salud se ocupan de utilizar pruebas diagnósticas que sean baratas, inocuas y que detecten en forma temprana las enfermedades. Generalmente todas las enfermedades tienen una prueba diagnóstica idónea o *prueba de oro* (*gold standar*). Sin embargo, la prueba de oro en muchas ocasiones es difícil realizarla porque es riesgosa y, por ello, se utilizan métodos alternos que se acerquen a la certeza del primero. Por ejemplo, para determinar un infarto de miocardio, la prueba de oro sería la biopsia de miocardio, pero esto no es posible por los riesgos que implica, por lo tanto se utilizan pruebas alternativas como son el cuadro clínico, cambios enzimáticos y electrocardiograma; para esto, se vale de la determinación, sensibilidad, especificidad, valor predictivo positivo-negativo y curva característica operativa del receptor (ROC). La determinación de estos parámetros ha permitido utilizar las pruebas idóneas para el diagnóstico más temprano de las enfermedades e iniciar medidas correctivas, siendo una consecuencia de ello la disminución de los costos de la enfermedad. Desafortunadamente algunos profesionales de la salud no dominan estos conceptos y no las utilizan en su bagaje profesional.

Tratamiento

Mediante las herramientas de la epidemiología y estadística es posible determinar la eficacia terapéutica y elegir el mejor tratamiento para la mayoría de las enfermedades. La epidemiología enseña a determinar el diseño metodológico idóneo, el tamaño de muestra, tipo de muestreo, selección de la muestra (criterios de inclusión, exclusión y eliminación), la asignación de los grupos de estudio, las pruebas estadísticas e interpretación de los resultados. El conocimiento de estos conceptos es indispensable para realizar un proyecto de investigación bajo el modelo *ensayo clínico*. La correcta utilización de las herramientas de la epidemio-

logía y estadística, aseguran la calidad y confiabilidad de cualquier ensayo clínico, lo que conduce al avance del conocimiento y la mejora de los sistemas de salud.

Pronóstico

Cuando un paciente enferma pregunta al médico qué tan grave es su enfermedad, cuánto durará y qué secuelas puede dejar. Con el fin de responder a estas preguntas, el médico se vale de instrumentos de evaluación, que anteriormente fueron aceptados o validados por la comunidad médica. Estos instrumentos son modelos matemáticos en donde se correlacionan las variables o factores que en forma aislada o conjunta afectan la evolución o desenlace de una enfermedad. Como ejemplos se encuentran las siguientes escalas: *Glasgow* (traumatismo craneoencefálico), *Apgar* (estado físico del recién nacido), *Apache* (gravedad del paciente hospitalizado). Estas pruebas o escalas también son conocidas como *clinimetría*. Este último término fue acuñado por el médico y matemático Alvan Feinstein, el cual se refería al conjunto de instrumentos o escalas utilizadas con el fin de determinar el grado de gravedad de un paciente.

Costos

La epidemiología es una herramienta que puede ser utilizada para calcular los costos de cualquier evento del proceso salud-enfermedad; por ejemplo, costo por hora-quirófano, día-hospital, medicamentos, rayos x, laboratorio, etcétera.

Comprensión de la literatura médica

Los médicos que posean nociones de epidemiología, les será más sencillo entender la literatura médica ya que podrán diferenciar entre un artículo con y sin fallas metodológicas. Este conocimiento es una ventaja para el médico ya que le permite incorporar en su práctica o conocimiento sólo aquella información que no tenga fallas o errores metodológicos. Desafortunada-

mente es escasa la cantidad de médicos que dominan las herramientas epidemiológicas y frecuentemente incorporan en su bagaje información no confiable. Este fenómeno es más frecuente que suceda en aquellos médicos que son subyugados por la industria farmacéutica a través de estudios clínicos que frecuentemente presentan errores metodológicos graves. Por esto se sugiere que la asignatura de epidemiología en las facultades de medicina sea impartida por profesores que dominen la metodología de la investigación o utilicen con alta regularidad las herramientas de la metodología.

Clasificación de la epidemiología

La epidemiología, de acuerdo con sus objetivos, se divide en *descriptiva* y *analítica*.

Epidemiología descriptiva

Ésta estudia la distribución de las enfermedades en tiempo, espacio y persona; es una descripción lo más fehaciente posible de la realidad. No se ocupa de analizar las causas de la enfermedad o relación causa-efecto. Generalmente utiliza estadística descriptiva como por ejemplo medidas de tendencia central, dispersión y gráficas.

La epidemiología descriptiva es la base de todas las investigaciones futuras, ya que es sumamente difícil que un trabajo de investigación como un ensayo o experimento no se sustente en una investigación descriptiva previa. La epidemiología descriptiva es utilizada por los estudiantes del área de salud para realizar trabajos de investigación en licenciatura, maestría y ocasionalmente doctorado. También se recurre a la epidemiología descriptiva en los departamentos de epidemiología y estadística de las unidades hospitalarias.

Epidemiología analítica

Esta clase de epidemiología busca el porqué o las causas que provocaron la enfermedad; igualmente, trata de buscar asociaciones estadísticas o causales y el peso de las mismas, y pretende que sus

resultados sean utilizados para disminuir la morbilidad de cualquier enfermedad. Los diseños metodológicos analíticos y experimentales utilizan la epidemiología analítica. Una característica de la epidemiología analítica es la determinación de los riesgos relativos (RR), razón de momios o de ventajas (OR), riesgo absoluto (RA), reducción de riesgo relativo (RRR), reducción de riesgo absoluto (RRA), los cuales se determinan en diseños observacionales analíticos (estudios de cohorte o casos y controles) y experimentales (ensayos clínicos).

Conclusión

La epidemiología es una herramienta muy importante que puede ser utilizada en la prevención, diagnóstico, tratamiento y rehabilitación del paciente. El dominio de las herramientas de la epidemiología debe ser imprescindible en epidemiólogos e investigadores clínicos y debería serlo también para todos los profesionales de la salud.

Glosario

Epidemiología. Es una disciplina que se encarga del estudio relacionado con el proceso salud-enfermedad en una población.

Epidemiología analítica. Estudia el porqué o las causas que provocaron la enfermedad. Trata de buscar asociaciones estadísticas, causales, o ambas, y el peso de las mismas.

Epidemiología clínica. Es una rama que utiliza los métodos o estrategias de la epidemiología en el quehacer clínico.

Epidemiología descriptiva. Estudia la distribución de las enfermedades en tiempo, espacio y persona.

Estadística. Estudia los métodos y procedimientos para recoger, clasificar, resumir y analizar datos y realizar inferencias a partir de los mismos.

Autoevaluación

¿Qué es la epidemiología y cuáles son sus principales antecedentes históricos?

¿Cuáles son algunos aportes importantes de la estadística para el desarrollo de la epidemiología como disciplina científica?

¿Qué es y qué estudia la epidemiología clínica?

¿Cuáles son las principales aplicaciones de la epidemiología clínica para favorecer un servicio más eficiente en los contextos hospitalarios y extra hospitalarios?

Bibliografía

- López, M.S.; Corcho B.A., y Moreno A.A. (1999). Notas históricas sobre el desarrollo de la epidemiología y sus definiciones. En: *Revista Mexicana de Pediatría*, 66(3), pp. 110-114.
- López-Moreno, S.; Garrido-Latorre, F., y Hernández-Ávila, M. (2000). Reseña histórica de la epidemiología. En: *Salud Pública de México*, 42(2), pp. 133-143.
- Guerra de Macedo, C. (1994). *Usos y perspectivas de la epidemiología*. Washington, D.C.: Organización Panamericana de la Salud.
- Frenk, M.J. (1993). *La salud de la población. Hacia una nueva salud pública*. México, D.F.: Fondo de Cultura Económica.
- López-Moreno, S.; Corcho-Berdugo, A., y López-Cervantes, M. (1998). La hipótesis de la comprensión de la morbilidad: un ejemplo de desarrollo teórico en epidemiología. En: *Salud Pública de México*, 40, pp. 442-449.
- Rose, G. (1988). Individuos enfermos y poblaciones enfermas. En: Organización Panamericana de la Salud, *El desafío de la Epidemiología* (pp. 900-909). Washington, D.C.: OPS.

CAPÍTULO II

El método científico y su aplicación en la investigación

José Ramiro Caballero Hoyos

Introducción

El método científico se puede definir como un conjunto de principios y normas que orientan la planeación e implementación técnica de las distintas fases en que se desenvuelve el proceso de la investigación científica (Grawitz, 1984: 291). Su aplicación consiste en proponer reglas mediante las cuales se pretende alcanzar un conocimiento sistemático que goce de un amplio consenso social.

Dicho conocimiento aspira al mayor consenso porque el método con el cual se genera aporta controles técnicos estrictos a sus procedimientos y favorece la discusión crítica de sus argumentos al proponer reglas con un lenguaje de uso común. Por ello, los nuevos descubrimientos tendrán mayor aceptación en la comunidad científica solamente si tienen el soporte de evidencia empírica suficiente y si fueron sometidos a un amplio debate público (Bunge, 1979).

El desarrollo del método científico es el resultado de una historia donde se confrontaron concepciones de mundo divergentes que generó desde el siglo XVII una revolución mental sin precedentes sobre tres grandes asuntos: cómo problematizar la

realidad (mediante formas de cuestionar radicalmente diferentes), cómo proponer respuestas tentativas a los problemas (mediante la formulación de hipótesis estadísticas y formas de control) y cómo generar experiencias no espontáneas para dar respuestas empíricas verificables (mediante el desarrollo del método experimental [Piaget y García, 1987]).

Esa historia consistió en un lento aprendizaje cultural del método experimental inductivo que estuvo acompañado de algunos cambios económicos y políticos profundos en las sociedades, entre los cuales se pueden mencionar los siguientes: a) la desvinculación paulatina del conocimiento orientado por la metafísica; b) la conquista de espacios de libertad de pensamiento y opinión; c) el desarrollo de sistemas de educación pública que favoreció el acceso a la información; d) los adelantos en las concepciones de la física teórica y las matemáticas que permitieron trascender del conocimiento observable del espacio concreto al conocimiento indirecto de espacios abstractos; e) el fermento de las ideas creativas y contestatarias del imaginario renacentista que abonó a la renovación del pensamiento de los siglos xvii y xviii; f) los crecientes descubrimientos técnicos de disciplinas cada vez más especializadas y con herramientas de observación más refinadas; y g) el desarrollo del sistema de producción capitalista que brindó la base material y el sustento ideológico a la ciencia moderna (Koyré, 1977; Grawitz, 1984).

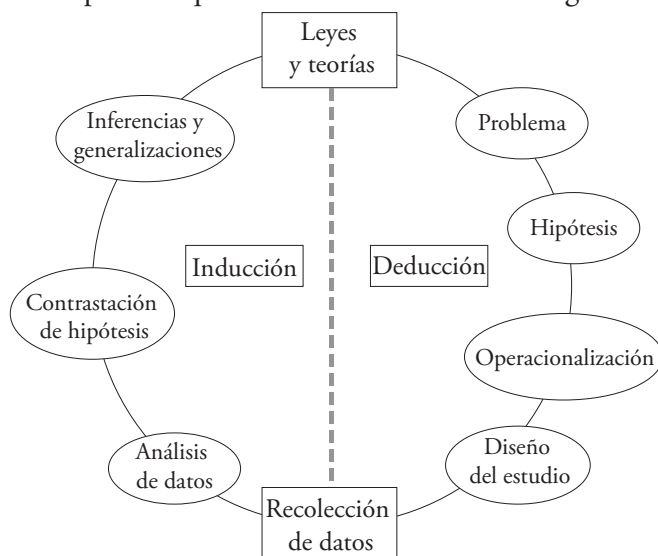
El conocimiento científico, construido a partir de la aplicación del método en el proceso de investigación, ha podido aportar desde el siglo xvii tanto explicaciones a variedad de problemas, como respuestas eficaces para el desarrollo de tecnologías que impulsan el desarrollo socioeconómico y cultural de los países. La misma ha tenido una adaptación muy eficiente a la resolución de problemas emergentes que el conocimiento ordinario basado en mitos y creencias no pudo resolver y ha enriquecido y generado leyes y teorías que han orientado aplicaciones prácticas de incalculable valor para la humanidad.

En el presente capítulo se presentarán las nociones básicas del método científico y una descripción de la forma en que se aplica a las distintas fases que rodean al proceso de investigación. En ese intento, se enfatizarán los razonamientos deductivo e inductivo que orientan la aplicación lógica del método y se señalan algunos supuestos filosóficos que sustentan a los procedimientos más relevantes.

Aplicación del método en el proceso de investigación

En la aplicación del método científico, el investigador genera procesos mentales en los que se elaboran ideas que interactúan con actividades empíricas enfocadas a comprobarlas. Según Bryman (2004) y Wallace (1980), la interacción de estos procesos se desarrolla en las diferentes fases de una investigación empírica convencional que operan bajo los distintos procedimientos de razonamiento deductivo e inductivo como se muestra en la figura 1:

Figura 1
Fases deductivas e inductivas
del proceso que se lleva a cabo en la investigación



Fuente: Adaptado de Wallace (1980).

Por un lado (parte derecha de la figura 1), las fases deductivas (de lo general a lo particular) surgen desde conceptos generales inspirados en las leyes y teorías vigentes que orientan el planteamiento del problema e hipótesis que predicen posibles explicaciones a dicho problema. Estas hipótesis se *operacionalizan* en medidas que hacen posible la observación empírica de sus conceptos, bajo la guía estratégica de un determinado diseño de estudio (experimental o no experimental).

Por otro lado (parte izquierda de la figura 1), las fases inductivas (de lo particular a lo general) se inician con la recolección de datos mediante la aplicación de procedimientos e instrumentos estandarizados. Los datos obtenidos se transforman y procesan para efectuar distintos niveles de análisis matemáticos enfocados a contrastar las hipótesis estadísticas propuestas. La interpretación de los hallazgos de contrastación permite realizar inferencias de parámetros poblacionales y generalizaciones que enriquecen las leyes y teorías vigentes, lo cual puede orientar el desarrollo de un nuevo proceso de investigación sobre la base de sus hallazgos.

A continuación se describirán por separado y con más detalle las distintas fases correspondientes a los razonamientos lógicos deductivo e inductivo.

Fases deductivas

El conocimiento histórico ha puesto a disposición de los científicos un repertorio acumulado de leyes y teorías que les permiten plantear posibles explicaciones y predicciones para los problemas e hipótesis que deseen investigar.

En términos generales, las leyes científicas son las proposiciones más aceptadas sobre las relaciones y regularidades de la naturaleza. En el planteamiento del problema y de las hipótesis de investigación, estas leyes ayudan a predecir las posibilidades estadísticas de aparición de una amplia serie de eventos. A su vez, dichas leyes pueden ser confirmadas empíricamente en el proceso de investigación.

Por su parte, las teorías son proposiciones generales que fundamentan las leyes, a partir de explicaciones abstractas. Debido a su amplitud, no son susceptibles de comprobación empírica, pero conducen a la formulación de hipótesis específicas que pueden ayudar a predecir la aparición de nuevas leyes.

Esta derivación de lo general (leyes y teorías) a lo particular (problemas e hipótesis) corresponde a un razonamiento deductivo que opera bajo la lógica aristotélica que se refiere a la evaluación de los argumentos, donde se valora si la evidencia de los enunciados (premisas) conduce a conclusiones válidas. Bajo este razonamiento se traducen también los conceptos generales en medidas numéricas que permiten la observación empírica en la fase *recolección de datos* (Richards, 2005).

Las fases que corresponden al razonamiento deductivo son las siguientes:

La definición del problema

Cuando se inicia un proceso de investigación, al científico se le ocurren nuevas ideas sobre un tema y alguna problemática que considera relevante para la sociedad y su campo de conocimiento.

En la revisión de la literatura, el científico, a veces conscientemente y a veces no, examina las premisas y argumentos teóricos y conceptuales. Tras una lectura crítica, plantea un problema de investigación en forma de una o varias preguntas.

Las preguntas planteadas expresan cuáles son las variables del estudio y sus posibles relaciones. A su vez, permiten entrever la posibilidad de realizar pruebas empíricas para establecer si las relaciones son estadísticamente significativas. Con dichas preguntas, el investigador orienta el planteamiento de los objetivos y la justificación del estudio.

La formulación de hipótesis

Las hipótesis son enunciados que predicen la relación o asociación tentativa entre dos o más variables de una población, en la explicación de un fenómeno específico.

Su formulación implica que el investigador ha revisado toda la literatura que incluye teorías y leyes correspondientes al fenómeno que estudia, para poder plantear las explicaciones que propone. A su vez, los enunciados de las hipótesis tratan de dar respuestas tentativas a las preguntas de investigación.

La enunciación de las hipótesis puede proponer relaciones de fuerza o dirección entre las variables, lo cual habilita su verificación mediante procedimientos de observación e inferencia. Esta enunciación se expresa en la forma de una hipótesis estadística bajo la estrategia empleada por las pruebas de hipótesis.

El razonamiento principal de la prueba de hipótesis es que si bien una hipótesis no puede ser aceptada con certeza absoluta, puede más bien ser rechazada por ser falsa o errónea. Así, la prueba de hipótesis conduce a una decisión con respecto a una hipótesis estadística (H), en términos de si hay o no suficiente evidencia para concluir que es falsa:

- Cuando los datos favorecen a la hipótesis nula (H_0), ésta se acepta como cierta.
- Cuando los datos contradicen a la hipótesis nula (H_1), ésta se rechaza como falsa.

Tal proceder se sustenta en la lógica que refiere al ensayo de error de la noción *falsabilidad* propuesta por el filósofo Karl Popper (1967) quien ha tenido una gran influencia en la ciencia contemporánea, a tal punto que algunos consideran que es el elemento clave del método científico. Esta noción sugiere que los científicos no pueden verificar con exactitud las hipótesis científicas, pueden más bien hacer todo lo posible por demostrar su falsedad o nulidad.

La operacionalización de conceptos

La operacionalización consiste en traducir el nivel abstracto de los conceptos enunciados en las hipótesis a un nivel empírico de medidas expresadas en el lenguaje formal de los números.

Dicha traducción implica seguir los siguientes pasos: 1) generar variables que correspondan a los conceptos de las hipótesis; 2) definir el significado de cada variable en términos de lenguaje ordinario (por ejemplo, los que refiere el diccionario); 3) desglosar bajo qué medidas directas o indicadores indirectos se pueden observar las variables; y 4) señalar cómo se combinarán las medidas o indicadores para generar variables complejas sobre la base de escalas o índices numéricos.

La operacionalización es una suerte de aterrizaje del pensamiento abstracto al lenguaje de los números como procedimiento que habilita la observación indirecta; es decir, un proceder tomado de la física teórica para su aplicación en las demás ciencias empíricas, bajo el supuesto filosófico *fenomenalista* que propone lo siguiente: el sujeto conoce el mundo externo no en forma directa, sino a través de *datos sensoriales*. Por ello, la ciencia requiere de un lenguaje formal preciso que medie en la construcción del conocimiento sobre dichos datos sensoriales (Ayer, 1965).

La selección de un diseño de investigación

El diseño de la investigación es una construcción lógica que orienta al investigador a seguir procedimientos sistemáticos de recolección y de análisis de datos para responder sus preguntas de investigación y contrastar sus hipótesis. En esencia, se trata de una herramienta estratégica que lleva a desarrollar inferencias con un mínimo grado de error (De Vaus, 2001).

La tradición de las ciencias empíricas ha desarrollado dos tipos generales de diseños de investigación: el experimental y el no experimental. Cada uno de ellos propone una lógica de acercamiento a la realidad en términos de búsqueda de explicaciones científicas (Twisk, 2004):

Diseños experimentales

En éstos el investigador manipula aspectos de un lugar (laboratorio o situación de campo) y observa los efectos de dicha manipulación sobre los sujetos experimentales en comparación con un

grupo control, bajo el supuesto de que la asignación de sujetos a ambos grupos tuvo un carácter aleatorio. Este procedimiento constituye un tipo de observación empírica controlada que se enfoca a establecer explicaciones causales. Cuando estos diseños se aplican en la epidemiología se denominan ensayos clínicos, cuyo cometido principal es analizar prospectivamente el efecto de una o más intervenciones (variable independiente x) sobre determinada variable dependiente o resultado (y).

Diseños no experimentales

En este caso el investigador recolecta datos de las variables de estudio, a veces en forma simultánea (por ejemplo, en la aplicación de un cuestionario con distintos tipos de preguntas, en una sesión de entrevista), y otras en forma previa a la presencia de un efecto previsto (por ejemplo, en las observaciones de seguimiento de cohortes), sin hacer manipulación alguna sobre los aspectos de lugar.

La información recolectada en una muestra de población (seleccionada a veces por muestreo aleatorio y a veces sin aleatoriedad) se procesa y analiza estadísticamente para establecer relaciones de asociación y de diferencia entre las variables, a efectos de contrastar las hipótesis y generar inferencias de los parámetros poblacionales.

Cuando se aplican a la epidemiología, los diseños no experimentales se clasifican en: estudios observacionales de cohorte (retrospectivos, transversales o prospectivos) y estudios de casos y controles (retrospectivos).

Fases inductivas

El razonamiento inductivo nos lleva de un conjunto finito de enunciados singulares a la justificación de un enunciado universal. La evaluación argumentativa de tales enunciados se basa en la regla de realizar la mayor cantidad de observaciones posibles antes de justificar cualquier generalización que pretenda impactar sobre las leyes y teorías aceptadas por la comunidad científica. En con-

traste con el razonamiento deductivo, no tiene el soporte de una lógica de premisas y conclusiones para apreciar si los enunciados son válidos (Richards, 2005).

Las fases que corresponden al razonamiento inductivo son las siguientes:

Recolección de datos

La recolección de datos consiste en la aplicación de los instrumentos de medición estandarizados, a los sujetos de estudio, con el propósito de obtener sus respuestas. La aplicación se realiza bajo ciertas condiciones de observación que contribuyen a minimizar las posibilidades de error en la medición (por ejemplo, en los estudios experimentales se asegura la asignación aleatoria de los sujetos a las condiciones experimental y control para hacer factible la comparación entre grupos).

Análisis de datos

Las respuestas que se obtuvieron de los sujetos se codifican numéricamente bajo diferentes escalas de medición (dicotómica, nominal, ordinal, intervalo o razón) que habilitan la posibilidad de realizar distintos niveles de procesamiento estadístico: a) el análisis univariado donde se describe la distribución de una sola variable (por ejemplo, el cálculo de frecuencias, prevalencias y tasas de incidencia) y se exploran sus propiedades matemáticas (por ejemplo, si las variables numéricas tienen una distribución normal o no); b) el análisis bivariado donde se examina la conexión de dos variables en términos de su distribución (frecuencias y porcentajes) y de sus diferencias o relaciones de asociación (pruebas de significancia y de estimación por intervalo); y c) el análisis multivariado donde se examina la conexión entre más de dos variables en términos de sus diferencias y relaciones de asociación.

La estadística inferencial proporciona métodos para estimar las características o atributos de una población al basarse en datos de una muestra de observaciones y con un determinado margen de error.

Contrastación de hipótesis o verificación

Después de analizar la información, el investigador genera una evaluación sobre la concordancia entre la predicción de sus hipótesis y el comportamiento de la realidad que observó empíricamente; para ello se vale del procedimiento denominado contraste de hipótesis estadísticas.

La contrastación de hipótesis no establece la verdad de la hipótesis, sino un criterio que permite al investigador decidir sobre si una hipótesis nula se acepta o rechaza, o si las muestras observadas difieren significativamente de los resultados esperados.

Cuando una hipótesis es aceptada (por efectos del rechazo de la hipótesis nula), la evidencia para la aceptación es inductivamente conclusiva en términos de significancia estadística, sin la mediación de ingredientes deductivos.

En el proceso de contrastación de hipótesis, el investigador podría cometer dos tipos de errores: a) rechazar una hipótesis nula cuando debiera ser aceptada (error de tipo I); y b) aceptar una hipótesis nula cuando debiera ser rechazada (error de tipo II). Estos errores se originan en el tipo de diseño de investigación que se sigue y por ello es en esa fase que deben tomarse las medidas de control necesarias para minimizarlos (por ejemplo, es central considerar un tamaño de muestra adecuado en función de las características de la hipótesis que se contrasta).

Inferencias estadísticas y generalizaciones

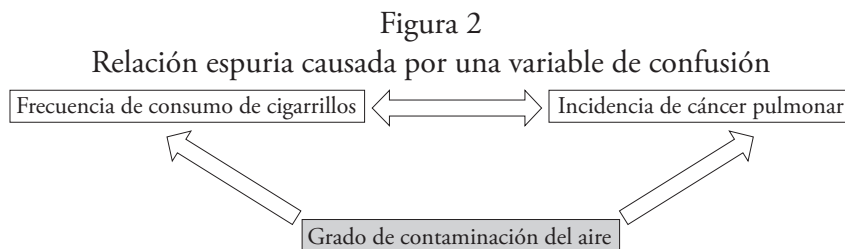
La inferencia estadística es el proceso de usar resultados muestrales para obtener generalizaciones respecto a las características de una población. Comprende un conjunto de procedimientos estadísticos (pruebas de significancia y de estimación por intervalo) en los que interviene la aplicación de modelos de probabilidad y mediante los cuales se estiman los parámetros de una población con cierto grado de error.

La inferencia y la generalización de la investigación cuantitativa se enfocan a generar explicaciones bajo el modelo general

de la causalidad o de las relaciones causa-efecto entre las variables de las hipótesis que se contrastan. Al respecto los investigadores se preguntan: ¿en qué medida la distribución de valores correspondientes a la variable dependiente, es afectada por la variación de la variable independiente?

En la lógica del imaginario causal, se han propuesto algunos criterios que pueden orientar al investigador en su búsqueda de requisitos para establecer si existe o no una relación de causa-efecto en los estudios epidemiológicos (Twisk, 2004: 2):

- Primero, explorar si la distribución de valores relativos a las variables muestra una aparente relación entre éstas y si el grado de relación es fuerte o débil, por ejemplo, si la mayor frecuencia de consumo de cigarrillos (variable independiente) se asocia con una mayor incidencia en cáncer pulmonar (variable dependiente).
- Segundo, confirmar que la relación de variables no sea espuria. Una relación espuria ocurre cuando no hay una relación verdadera entre dos variables que parecen estar conectadas, debido a la presencia de una tercera variable que afecta a ambas. Por ejemplo, en el siguiente esquema (figura 2) se muestra que en la relación entre la frecuencia del consumo de cigarrillos y la tasa de incidencia del cáncer pulmonar, puede mediar el efecto causado por el grado de contaminación del aire en el espacio ambiental donde residan las personas estudiadas (variable de confusión):



Fuente: Elaboración propia.

- Tercero, observar si la relación es consistente en diferentes poblaciones y bajo distintas circunstancias.
- Cuarto, establecer si la variable independiente precede a la variable dependiente, en términos de especificidad (la variable independiente precede a un solo efecto en la variable dependiente) y temporalidad (la variable independiente precede a un efecto en el tiempo).
- Quinto, establecer el gradiente biológico (relación dosis-respuesta).
- Sexto, establecer la plausibilidad biológica de la relación.

El establecimiento de las relaciones causa-efecto es fuertemente afectado por la naturaleza del diseño de investigación que se selecciona. Si se opta por el experimento, se debe ser consciente de que la elucidación de causa y efecto es un aspecto explícito del marco teórico de este diseño, bajo el concepto de validez interna (hallazgos con evidencia firme de causalidad por cumplir los distintos requisitos para su establecimiento). En la epidemiología, básicamente es el único diseño que cubre los requisitos de causalidad.

Por el contrario, si se opta por un diseño no experimental, se debe considerar que existe la posibilidad de establecer correlaciones entre variables recolectadas simultáneamente con una limitada posibilidad de inferir causas y efectos entre éstas. En este diseño, el requisito de la temporalidad se ha compensado con la introducción de los estudios longitudinales con dos o

más momentos de recolección de datos a los mismos o a diferentes sujetos, a fin de establecer si el impacto de una variable basal es mayor en el proceso de cambio de una variable dependiente. En este sentido, los estudios longitudinales permiten asomarse al asunto de la explicación causal en términos de probabilidad, pero sin la certeza de los estudios experimentales.

Las generalizaciones causales basadas en las inferencias estadísticas, llevan a los investigadores a generar las conclusiones del estudio, las cuales se confrontan con el repertorio de leyes y teorías acumuladas para valorar su plausibilidad y su relevancia como aportes al conocimiento científico y al desarrollo de nuevas preguntas e hipótesis de investigación. Con ello, el ciclo del proceso de investigación comenzaría otra vez.

Conclusión

El presente capítulo ha descrito el método científico como un conjunto de principios y normas que orientan la planeación e implementación que se lleva a cabo en el proceso de investigación. Bajo esa concepción, ha sintetizado las distintas fases de aplicación del método en dicho proceso.

Un argumento central del ensayo es que el método científico ha sido una construcción histórica que refleja la maduración del pensamiento humano en su afán por conocer su entorno. Antes del siglo xvii predominaron métodos basados en la observación directa, la lógica deductiva aristotélica y la filosofía teñida de ideas teológicas. Estos métodos correspondían a formas de producción no capitalistas con una fuerte influencia del poder religioso y monárquico.

Como resultado de cambios estructurales que originaron el sistema capitalista, la producción masiva y sus efectos demográficos, la secularización de los valores y la modificación de estilos de vida, llegó a imponerse desde el siglo xvii el método científico que complementó la deducción con el razonamiento inductivo, la observación indirecta de los fenómenos, la formu-

lación de hipótesis sustentadas en leyes y teorías, y el diseño de experimentos con auxilio de instrumentos técnicos más precisos.

Lo relevante de la aplicación del método al proceso de investigación es que ha permitido desarrollar un conocimiento científico con leyes y teorías acumulables en el tiempo y ampliamente estables en su intento por explicar la realidad. También ha orientado a construir un repertorio vasto de preguntas e hipótesis como motor que impulsa la renovación de las ideas y la creatividad en el diseño de nuevas investigaciones.

La precisión de sus procedimientos y conceptos, y la claridad de su normatividad han facilitado la evaluación crítica respecto a los procesos de aplicación, con lo cual se ha desarrollado una cultura científica que se comunica con el lenguaje común del método en la difusión de los hallazgos y en el debate de los argumentos, así como del sustento técnico y metodológico.

Al terminar, cabe aclarar que la aplicación del método científico descrita corresponde a la concepción de la filosofía positiva que enfatiza los procedimientos experimentales y la inferencia estadística basada en la contrastación de hipótesis. Esta concepción es la que hegemoniza la cultura científica contemporánea y no está exenta de críticas y cuestionamientos.

Como crítica importante se menciona la de algunas disciplinas de las ciencias sociales y las humanidades que reivindican, desde finales del siglo XIX, la necesidad de comprender las manifestaciones y significados del lenguaje expresado por las personas, en su contexto cultural específico; para ello promueven los métodos cualitativos que profundizan en la observación de los sujetos en su entorno natural y la interpretación de ideas y creencias grupales, sin recurrir a la búsqueda de patrones e inferencias estadísticas basadas en experimentos. El uso de estos métodos es una alternativa que ha ganado un espacio creciente en la comunidad científica internacional y goza cada vez de más legitimidad, por lo cual su descripción sería un asunto a complementar para poder hablar con mayor amplitud sobre el tema del presente capítulo.

Glosario

Conocimiento científico. Saber crítico sobre el mundo real que se genera mediante observación empírica basada en los procedimientos y principios generales por los cuales se compone el método de la ciencia. En su comunicación se emplea un lenguaje con reglas precisas y explícitas que favorecen la crítica de la validez de sus argumentos por los actores de la comunidad científica.

Deducción. Razonamiento lógico que lleva a desarrollar determinados enunciados a partir de otros enunciados dados. En el proceso de investigación, estar al tanto de leyes y teorías universales permite derivar de éstas algunas explicaciones y predicciones.

Falsabilidad. Supuesto filosófico que sustenta la factibilidad del procedimiento que se sigue en las pruebas de hipótesis. Su noción principal indica que aunque las hipótesis científicas no pueden ser verificadas, es factible demostrar su falsedad o nulidad.

Fenomenalismo. Supuesto de la filosofía positiva que sustenta la factibilidad del procedimiento llevado a cabo por la operacionalización de conceptos que se encuadran en las hipótesis. Asimismo postula la idea de que el sujeto no conoce el mundo externo de forma directa, sino a través de *datos sensoriales*, por lo cual la ciencia requiere de un lenguaje formal numérico que medie en la construcción del conocimiento sobre esos datos sensoriales.

Inducción. Razonamiento lógico que lleva de una lista finita de enunciados singulares a la justificación de un enunciado universal.

Inferencia estadística. Es el proceso de usar resultados muestrales para obtener generalizaciones respecto a las características de una población. La misma comprende un conjunto de procedimientos estadísticos (pruebas de significancia y estimaciones de intervalo) en los que interviene la aplicación de modelos de probabilidad y mediante los cuales se estiman los parámetros de una población con determinado grado de error.

Leyes científicas. Definen relaciones y regularidades invariables acerca del mundo, que pueden ser confirmadas empíricamente por cualquier observador.

Método científico. Conjunto de principios y normas que orientan la planeación e implementación relativas a un proceso de investigación.

Teoría. Estructura de relaciones conceptuales que fundamenta las leyes (típicamente como explicaciones). Por su amplitud no es susceptible de comprobación empírica directa, pero brinda elementos para plantear hipótesis concretas sobre algunas relaciones conceptuales que pudieran ser relevantes para responder preguntas de investigación específicas.

Autoevaluación

¿Qué es el método científico y por qué su desarrollo se asocia a una profunda revolución mental en la historia de las sociedades?

¿Por qué se plantea que el método favorece el desarrollo de un conocimiento con amplio consenso social?

¿Cuáles son los razonamientos lógicos que orientan la aplicación de método en todo proceso de investigación? Defina en qué consiste cada una de éstos.

¿Cuáles son las fases deductivas del proceso de investigación?

¿Cuáles son las fases inductivas del proceso de investigación?

Bibliografía

- Ayer, A. J. (1965). *El positivismo lógico*. México, D.F.: Fondo de Cultura Económica.
- Bryman, A. (2004). The nature of Quantitative Research. En: *Social Research Methods* (pp. 62-68). London: Oxford University Press.
- Bunge, M. (1979). *La ciencia, su método y su filosofía*. Buenos Aires: Ediciones Siglo Veinte.
- De Vaus, D. A. (2001). *Research Design in Social Research*. London: Sage Publications.
- Grawitz, M. (1984). *Métodos y técnicas de las ciencias sociales*. Tomo I. México D.F.: Editia Mexicana.
- Koyré, A. (1977). *Estudios de historia del pensamiento científico*. Madrid: Siglo XXI.
- Piaget, J., y García, R. (1987). *Psicogénesis e historia de la ciencia*. México, D.F.: Siglo XXI editores.
- Popper, K.R. (1967). *La lógica de la investigación científica*. Madrid: Tecnos.
- Richards, S. (2005). *Filosofía y sociología de la ciencia*. México, D.F.: Siglo XXI editores.
- Twisk, J.W.R. (2004). *Applied longitudinal data analysis for Epidemiology*. Cambridge, C.B.: Cambridge University Press.
- Wallace, W. (1980). *La lógica de la ciencia en la Sociología*. Madrid: Alianza Universidad.

CAPÍTULO III

Cómo elaborar un proyecto de investigación

Rebeca Millán Guerrero

Introducción

El proyecto de investigación es la versión escrita del plan a través del cual se realiza un estudio. Para realizar un estudio de investigación, y saber si es factible dicha tarea, se debe contar con un proyecto de investigación. Frecuentemente se ignora el orden sobre cómo iniciar la investigación de manera científica, con los elementos básicos del método de la ciencia: la pregunta a investigar, el diseño, los individuos que tomarán parte en éste; la asignación, cubrir el tamaño de la muestra requerido en el tiempo fijado, el análisis, interpretación, conclusiones o extrapolación; problemas éticos, etcétera. El hecho de escribir los pensamientos, facilita que las ideas vagas se transformen en planes específicos. De esta manera, al planear el estudio de investigación y presentarlo ante un comité de expertos, nos permite solicitar consejo de colegas, adaptar el proyecto a los alcances y posibilidades de acuerdo con el sitio de trabajo para no invalidar los resultados y que nuestro esfuerzo sea inútil, y los resultados puedan ser integrados a la práctica clínica con éxito. Los estudios son útiles en tanto aporten inferencias válidas dentro de la muestra estudiada (validez interna) y que éstos puedan generalizar al resto de la población (validez externa).

Así pues, el objetivo del presente capítulo es informar sobre los pasos que deben seguirse para elaborar un proyecto integral de investigación que tenga los elementos necesarios para ser presentado ante el comité de expertos, quienes evaluarán su relevancia, pertinencia y factibilidad.

Estructura de un proyecto

A continuación se proponen los elementos a incluir cuando se elabora un proyecto de investigación:

Título

Éste debe ser corto, conciso y debe expresar, con el menor número de palabras, qué estudio se realizará, el universo de estudio, el diseño de investigación y las variables que se estudiarán (Hulley y Cummings, 1997).

Antecedentes

Consecuentemente debe realizarse una búsqueda exhaustiva de la literatura mundial, tanto internacional, nacional y local del tema que se investigará. Con este fin se recomienda que la extensión de dicho capítulo no sea mayor de dos cuartillas; para conseguir lo anterior, se resumirá lo relevante del estudio, incluir la historia y su relación con el momento actual en los ámbitos internacional y nacional (en este caso México). Este apartado tiene como objetivo que el investigador domine la bibliografía al respecto y con ello obtenga las herramientas necesarias para contestar cualquier duda en relación con el tema; asimismo evitará repetir estudios ya realizados con el propósito de que el suyo sea inédito. Finalmente este apartado explica el motivo de la investigación, ya sea por ser nuevo o porque no está ampliamente estudiado, etcétera (Friedman, Furberg y DeMets, 1985).

Planteamiento del problema

Este apartado se refiere a la pregunta que se investigará y es la parte más importante del proyecto. Las mejores preguntas surgen en la práctica o a partir de resultados de otras investigaciones. El investigador debe tener una pregunta (razón del estudio que iniciará), la cual será sencilla, específica y factible de realizar; en ella explicará el propósito u objetivo del estudio. Dicha pregunta deberá ser plausible de contestar, y deberá aportar datos útiles sobre los conocimientos que ya existan sobre el tema. Al escribirse, deberá expresarse entre signos de interrogación, por ejemplo: ¿Cuál es la prevalencia de diabetes *mellitus* en el hospital? ¿Es eficaz la aspirina en la fase aguda del infarto cerebral? (Hulley y Cummings, 1997; Friedman *et al.*, 1985; Rodríguez, Cabrera y Martínez-Cairo, 2001).

Objetivos

Siempre que se califica un proyecto, una de las preguntas que surgen es si están definidos adecuadamente los objetivos del estudio. Los objetivos de un proyecto de investigación resumen lo que se desea lograr con el estudio, es decir, deben estar relacionados con el planteamiento del problema. ¿Por qué deben formularse los objetivos de la investigación?: a) porque ayudan a enfocar el estudio; b) porque ayudan a evitar la recolección de datos innecesarios; y c) porque organizan el estudio. Los objetivos generales describen el propósito global que se espera obtener al realizar la investigación; por ello, se recomienda fragmentar en objetivos específicos, que si son formulados adecuadamente, facilitarán el desarrollo del proyecto. ¿Cómo se deben formular los objetivos? Se debe cuidar que los objetivos del estudio: a) cubran los diferentes aspectos del problema; b) estén claramente escritos y especificar exactamente lo que se hará, dónde y para qué; c) sean realistas; y d) usen verbos en infinitivo (demostrar, determinar, identificar, etcétera). Un objetivo es lo que el investigador desea demostrar, por ejemplo: que la aspirina es eficaz en la fase aguda del infarto cerebral (Riegelman y Hirsch, 1992: 49-51).

Hipótesis

Al definir los objetivos del estudio, es necesario formular una hipótesis específica. Una hipótesis es lo que piensa el investigador que encontrará, es una conjetura o argumento que trata de explicar determinados hechos. Generalmente el investigador ya tiene una idea del resultado esperado que desea confirmar y deberá probar. No se necesita hipótesis en estudios descriptivos que simplemente describen cómo se distribuyen las características de una población; se necesitan hipótesis en estudios que emplearán prueba de significación estadística para comparar los hallazgos entre los grupos. Una buena hipótesis debe basarse en una pregunta de investigación, sobre todo debe ser simple y específica. ¿Cuándo es simple?: cuando contiene una variable predictora o independiente y una variable de desenlace o dependiente, la cual se escribe como afirmación, ejemplo: la aspirina es eficaz en infarto cerebral (Rodríguez *et al.*, 2001; Hulley y Cummings, 1997).

Material y método

Este elemento contiene los siguientes apartados:

- *Diseño.* Existen diferentes tipos o diseños de estudios: Los descriptivos observacionales se emplean para describir enfermedades, incidencia en padecimientos, etcétera. Los descriptivos analíticos, como los de casos y controles, son útiles para conocer factores de riesgo sobre algunas enfermedades, causas de mortalidades, etcétera. Los estudios de cohortes sirven para conocer el pronóstico de enfermedades, cómo se desenvuelve una enfermedad en el tiempo, etcétera. Por último, los estudios de intervención, como los ensayos clínicos controlados, son de utilidad para conocer la eficacia terapéutica o demostrar causas con alguna maniobra que realice el investigador (Riegelman y Hirsch, 1992).
- *Universo de trabajo.* Éste identifica a los individuos que formarán parte del estudio; es decir, la elección de una

muestra de individuos que represente a la población objetivo o diana, con las características que puedan responder a la pregunta planteada. El grupo de individuos especificados en el protocolo, sólo puede ser un subconjunto de la población de interés, puesto que existen barreras prácticas que impiden estudiar al total de la población. En este apartado se describe en quiénes se realizará el estudio, por ejemplo, recién nacidos, niños, jóvenes, adultos, enfermos con cefalea, con infarto cerebral agudo, etcétera. Debe describirse el lugar del estudio, hospital o servicio donde se trabajará (Hulley y Cummings, 1997: 141-151).

- *Tamaño de muestra.* Planear el tamaño de la muestra es estimar un número adecuado de sujetos para un determinado diseño de estudio. Se debe calcular el tamaño de la muestra en la etapa del proyecto del estudio, cuando aún se pueden efectuar cambios (Hulley y Cummings, 1997: 141-151). Existen fórmulas que deducen cuántos enfermos bastan para hacer un análisis descriptivo e inferencial, sin tener que estudiar a toda la población existente y, de acuerdo al tipo de diseño, a la incidencia de la enfermedad, a la eficacia o fracaso de una terapéutica, se calcula el tamaño de muestra.
- *Criterios de selección (criterios de inclusión).* Definen la población que se estudiará. De acuerdo con la literatura se tomará qué características específicas deberán tener los sujetos; por ejemplo, características demográficas y clínicas: sexo, edad, grado de enfermedad, embarazadas, etcétera (Hulley y Cummings, 1997: 141-151). Estos criterios sugieren definir una población accesible, que no dificulte realizar el estudio, pero que incluya a una muestra útil, que no cause error en los resultados.
- *Criterios de no inclusión o exclusión.* También de acuerdo al conocimiento o antecedentes del caso, se deberá pensar qué pacientes no deben estudiarse. Debe tenerse

especial cuidado en no cometer errores de incluir o no incluir a enfermos útiles o inútiles (Hulley y Cummings, 1997: 141-151).

- *Criterios de eliminación.* Este apartado se refiere a aquellos pacientes que habían sido incluidos y que por alguna razón deben eliminarse del estudio, por ejemplo: que tengan datos incompletos, que no tengan su expediente, que ya no quieran formar parte del estudio, que se embaracen, que tengan reacción secundaria a los fármacos a prueba, que abandonen el tratamiento, etcétera.
- *Variables a estudiar.* Se le llama variable a lo que se está estudiando, lo que cambia o se modifica con nuestra observación o intervención, son las características de los individuos estudiados: se deben elegir las variables que representarán los fenómenos de interés. En un estudio descriptivo el investigador considera cada una de las variables por separado. En un estudio analítico, el investigador analiza la relación existente entre dos o más variables para realizar inferencia sobre causa-efecto. Hay dos tipos de variables: la independiente o predictora, que es la que no se modifica y hará cambiar las variables. Dependientes o de desenlace, que serán las que se modifican, por ejemplo, dolor, inflamación, etcétera. (Hulley y Cummings, 1997).
- *Procedimiento.* En este apartado se debe explicar con precisión cómo se efectuará el estudio: cómo se captarán los enfermos, cuándo llegan, quién los seleccionará, quién captará la información, quién vigilará que se efectúe un cuestionario, cómo se dará el medicamento, qué color tiene, quién lo administrará, quién medirá los cambios de las variables, quién medirá los efectos obtenidos, efectos secundarios, etcétera.

Análisis estadístico

El investigador debe planificar el modo de manejar y analizar los datos obtenidos a lo largo del estudio. En el mismo debe indicarse qué pruebas estadísticas se efectuarán para poder realizar estadística descriptiva o inferencial, por ejemplo: media, desviación estándar, *chi* cuadrado, *t* de *Student*, etcétera (Hulley y Cummings, 1997; Dawson y Trapp, 2001: 13-20).

Consideraciones éticas

De acuerdo con el proyecto de estudio y el diseño, habrá consideraciones éticas que deben tenerse en cuenta. Si es un estudio descriptivo y es una revisión de expedientes, es suficiente un consentimiento verbal. Si se trata de un estudio con información confidencial, es probable que se necesite consentimiento verbal o por escrito, pero si se trata de una maniobra terapéutica o una prueba especial con riesgo, debe haber una carta ética en donde se describa en qué consiste el estudio, qué medicamento es, sus características, beneficios o perjuicios, si se empleará un grupo placebo y el riesgo de formar parte de éste. Dicha carta deben firmarla el investigador, el enfermo y dos testigos.

Recursos materiales y humanos

En este punto se debe mencionar qué se necesitará para realizar el estudio; por ejemplo, si es una encuesta: papelería, mimeógrafo para los cuestionarios, personal médico, de enfermería, encuestadores, etcétera. Expedientes: Si es un fármaco, quién lo proporcionará, jeringas, radiografías (RX), tomografía axial computarizada (TAC), etcétera. Médicos familiares, no familiares, un especialista, etcétera.

Financiamiento

Asimismo debe mencionarse el costo de lo siguiente: medicamento, papelería, jeringas, fotocopias, servicio del mimeógrafo, etcétera. De esta manera es necesario que el investigador realice un cálculo de costo para planear el estudio y determinar factibilidad.

Cronograma de trabajo

Este paso es importante, ya que debe contemplarse el tiempo aproximado que llevará el estudio: por ejemplo: De septiembre a octubre de 2013 = elaboración del proyecto. Noviembre de 2013 = presentación del proyecto al comité de investigación. Diciembre de 2013 a marzo de 2014 = elaboración del estudio. De abril a mayo de 2014 = análisis de resultados. Y junio de 2014 = presentación de resultados y envío a publicación.

Productos entregables

Esto se refiere a que deben señalarse los productos para difusión de conocimientos, los cuales resultarán cuando se lleve a cabo el proyecto. Por ejemplo:

- Artículos científicos de divulgación nacional.
- Artículos científicos de divulgación internacional.
- Síntesis ejecutivas para tomadores de decisiones.
- Guías informativas para la población.
- Difusión de los resultados en foros nacionales de la especialidad, o de investigación; asimismo en foros internacionales, etcétera.

Grupo de trabajo

El grupo de trabajo que colaborará, el currículo y en qué tema son expertos y cómo colaborarán en el proyecto.

Bibliografía

La bibliografía que se consultó debe mencionarse en cada apartado del proyecto. En el manuscrito debe anotarse dentro de paréntesis el número y en el apartado final deberá señalarse la referencia del documento revisado (revista, libro, etcétera).

Conclusión

Un proyecto de investigación requiere que tomemos en cuenta diversos puntos antes de iniciar el trabajo; de igual manera, debemos ordenar nuestros pensamientos en tiempo y forma, para determinar su factibilidad.

Glosario

Ética de la investigación. El investigador está obligado a seguir determinadas normas de comportamiento relacionadas con el proyecto que planea realizar. El investigador debe proteger —a los participantes— contra riesgos, daños y amenazas que pudieran afrontar ellos y el equipo de investigación. También es importante que en todo estudio se observen los derechos de respeto por la dignidad humana, la igualdad, la autonomía individual y la libertad de expresión. Cuando nos encontramos en espacios privados o con comunicaciones privadas, los investigadores están obligados a obtener el consentimiento informado de los participantes, es decir, debe informárseles sobre los objetivos de la investigación y solicitar su aceptación explícita para participar en el proyecto.

Factibilidad de un proyecto. Se refiere a la disponibilidad de los recursos necesarios para llevar a cabo los objetivos o metas señalados.

Investigación científica. Es la aplicación del método científico para resolver problemas o tratar de explicar determinadas observaciones.

Proyecto de investigación. Es la versión escrita relativa al plan del estudio.

Autoevaluación

¿Qué es un proyecto de investigación?

¿Cuál es la relación entre el proyecto y el método científico?

¿Por qué un proyecto debe expresar una propuesta de investigación factible y apropiada para el contexto clínico que se quiera estudiar?

¿Cuáles son los elementos a incluir cuando se elabora un proyecto de investigación?

- ¿Qué son los antecedentes y cuál es su objetivo principal?
- ¿Por qué el planteamiento del problema es la parte más importante del proyecto?
- ¿Qué son los objetivos y por qué deben formularse con claridad?
- ¿Qué son las hipótesis y qué relación tienen con los objetivos?
- ¿Cuáles son los elementos que deben incluirse en la propuesta de material y métodos del proyecto?
- ¿Qué relevancia tiene en el proyecto la consideración de los aspectos éticos?
- ¿Qué aspectos éticos son los más importantes en un proyecto clínico?

Bibliografía

- Dawson, B., y Trapp, R.G. (2001). *Basic & Clinical Biostatistics*. New York: Lange Medical Books-McGraw-Hill.
- Friedman, L.M.; Furberg, C.D., y DeMets, D.L. (1985). *Fundamentals of Clinical Trials*. Littleton, MA: PSG Publishing Co, Inc.
- Hulley, S.B., y Cummings, S.R. (eds.) (1997). *Diseño de la Investigación Clínica*. Madrid: Harcourt Brace.
- Riegelman, R.K., y Hirsch, R.P. (1992). *Cómo estudiar un estudio y probar una prueba: lectura crítica de la literatura médica*. Washington DC: Organización Mundial de la Salud.
- Rodríguez, J.; Cabrera, H., y Martínez-Cairo, S. (2001). *Epidemiología Clínica. Pruebas diagnósticas*. México, D.F.: Arte y Cultura.

CAPÍTULO IV

Riesgo, causalidad y mediciones de riesgo

Benjamín Trujillo Hernández

Antes de iniciar este capítulo, se cree conveniente referir la siguiente terminología.

- **Causa.** Se define como la condición que sola o acompañada produce una enfermedad.
- **Criterios de causalidad.** En muchas ocasiones existe asociación estadística y causal entre factor de riesgo y enfermedad, sin embargo, para que un factor de riesgo se considere causa debe cumplir con los siguientes criterios: *Fuerza de asociación:* número de veces que puede enfermar si se expone a un factor-riesgo (riesgo relativo y razón de momios). *Efecto dosis-respuesta:* relación entre factor de riesgo y enfermedad. *Secuencia temporal:* el factor de riesgo precede a la enfermedad. *Coherencia externa:* que los resultados sean semejantes con otros. *Ausencia de sesgo o errores.* *Ausencia de explicaciones alternativas o hipótesis.* *Plausibilidad biológica.* *Resultados compatibles con lo descrito.* *Efecto de cesación o reversibilidad.* *Ausencia de enfermedad sin el factor de riesgo.* *Demostración experimental.*
- **Riesgo.** Probabilidad de desarrollar una enfermedad por estar expuestos a determinados factores.

- Factor de riesgo. Son las variables o conjunto de factores físicos, químicos, biológicos, psicológicos, y sociales, que preceden y se asocian a la enfermedad o incrementan la probabilidad de enfermar. Por ejemplo: *hipertensión arterial y nefropatía*.

Los humanos nos encontramos distribuidos en el mundo en diferentes sociedades o comunidades y, aunque parecidas, existen diferencias en los ámbitos genéticos, ambientales y sociales. Estas características hacen que la probabilidad de enfermar o de estar sano sea diferente en todas las comunidades. El conjunto de características genéticas, personales, ambientales y sociales que incrementan la posibilidad de enfermar se conocen como factores de riesgo.

El término *factor de riesgo* lo utilizó por primera vez Thomas Dawber en un estudio donde demostró que la cardiopatía isquémica se asociaba a la hipertensión arterial, hiperlipidemia y tabaquismo. A excepción de las enfermedades infecciosas, la mayoría de éstas son causadas por la interacción de dos o más factores de riesgo. Así por ejemplo se sabe que el factor genético es protector o causante de la enfermedad entre 25 a 30%, mientras que el ambiental y social es responsable 70 a 75%.

Probablemente, en un futuro mediano podamos modificar el factor genético, sin embargo, en la actualidad es responsable de una cuarta parte de las enfermedades. La buena noticia es que la modificación del ambiente y social es la más importante. Por ejemplo, en la *diabetes mellitus* tipo 2, la carga genética no se puede modificar, sin embargo, si se lleva una vida de constante actividad física, y mantenerse con un peso ideal, la probabilidad de padecer la enfermedad disminuye en buena medida.

Una de las características del factor de riesgo es que siempre es anterior a la enfermedad o evento mórbido. En diseños analíticos observacionales (estudios de cohorte y casos y controles) el

factor de riesgo es la variable independiente; debe mencionarse que no todos los factores de riesgo son causas, sin embargo, todas las causas son factores de riesgo.

Mediciones de riesgo

Determinación del Riesgo Relativo (RR)

Una de las mediciones más utilizadas para evaluar la asociación causal es la determinación del RR y OR (Odds Ratio). Es importante también determinar en los RR y OR el intervalo de confianza del 95% (IC de 95%) con sus límites inferior y superior. Para que un factor de riesgo sea considerado como tal, la determinación de su RR u OR, y su límite de confianza inferior debe ser <1 , mientras que un RR u OR y su límite de confianza superior deben ser <1 . Ejemplo, supongamos que en un estudio se encontró lo siguiente: la diabetes se asoció a cardiopatía con un RR de 2.5 (IC de 95% = 0.9 a 3.5); aparentemente, hay asociación causal, sin embargo, el límite inferior fue <1 , por lo tanto no hay significancia estadística, mientras que en otro estudio se encontró un RR de 4.2 (IC de 95% = 1.8 a 5.7) y en este caso sí hubo asociación porque el límite inferior fue >1 . La antítesis de lo anterior es el *factor protector*, es decir, un factor que protege o disminuye la probabilidad de enfermar; en este caso el RR y su límite de confianza superior deben ser <1 .

Los diseños analíticos son los encargados de determinar o evaluar el impacto o la probabilidad de que los factores de riesgo causen una enfermedad o evento morboso. Los diseños analíticos son dos:

- *Estudios de cohorte*. La medida de asociación es el riesgo relativo (RR).
- *Estudios de casos y controles*. La medida de asociación es la razón de momios u Odds Ratio (OR).

Estudios de cohorte

Éstos se refieren a un diseño que estudia dos grupos de individuos del mismo tamaño muestral y semejantes en sus carac-

terísticas, pero que difieren en a) individuos con el factor de riesgo, y b) individuos sin el factor de riesgo.

Posteriormente, durante meses o años se estudian hasta que aparece la enfermedad o mueran. Los estudios de cohorte son prospectivos o longitudinales ya que parten de la causa o factor de riesgo hacia la enfermedad (C---E). En los estudios de cohorte lo que interesa saber es cuántos pacientes enfermaron (se supone que los individuos expuestos a un factor de riesgo se van a enfermar más que lo no expuestos). Dicho de otra manera, el porcentaje de enfermos será mayor en los individuos con el factor de riesgo. Por ejemplo, para evaluar si el tabaquismo causa cáncer de pulmón, se estudiaron 120 individuos que fumaban y 120 que no fumaban; después 20 años de seguimiento se encontraron los siguientes datos (tabla 1):

Tabla 1
Tabla de 2 x 2 para estudios de cohorte

Tabaquismo	Cáncer	Sanos	Total
Positivo	a 90	b 30	120
Negativo	c 15	d 105	120
Total	105	135	240

Fuente: Elaboración propia.

Como se mencionó, la medida de asociación en este tipo de estudios es la determinación del riesgo relativo o RR, el cual se realiza de la siguiente manera:

- Primero se determina la incidencia de expuestos, que es el porcentaje de individuos con el factor de riesgo que enfermaron y se determina con la siguiente fórmula: $IE = a/a+b = 90/120 = 75\%$. Se entiende que 75% de los pacientes con tabaquismo desarrollaron cáncer.
- Segundo, se determina la incidencia de no expuestos, que es el porcentaje de individuos sin el factor de riesgo

que enfermaron y se determina con la siguiente fórmula: $I_o = c/c+d = 15/120 = 12\%$. Se entiende que 12% de los individuos que no fuman desarrollaron cáncer.

- Tercero, determinación del riesgo relativo: $RR = I_E/I_o = 75\%/12\% = 6.2$, es la división de los fumadores con cáncer o índice de expuestos (numerador) entre los no fumadores con cáncer o índice de no expuestos (denominador). En este caso interpretamos los resultados de la siguiente manera: a) que hay 6.2 fumadores con cáncer y 1 no fumador con cáncer; o b) que el fumar incrementa 6.2 veces la probabilidad de tener cáncer.

Riesgo atribuible (RA)

Es una medida que se utiliza para determinar la diferencia y la incidencia entre los expuestos y no expuestos. La fórmula es una resta entre índice de expuestos menos índice de no expuestos ($I_E - I_o$). En este caso $75\% - 12\% = 63\%$. Significa que el tabaquismo incrementa, en 63%, el riesgo de contraer cáncer.

Fracción atribuible (FA)

Es una medida que se utiliza para conocer el porcentaje de individuos que mejorarían o evitarían la enfermedad si se retirase el factor de riesgo. También puede interpretarse como el porcentaje de individuos que se enferman con el factor de riesgo. La misma consiste en la división entre $(I_E - I_o)/I_E \times 100$, o sea, se divide el RA entre índice de expuestos (RA/I_E) y el resultado se multiplica por 100. En este caso $63\%/75\% = 0.84$ y multiplicado por 100 = 84%. Ello significa que si los individuos dejaran de fumar evitarían o disminuirían la probabilidad, en 84%, de padecer cáncer.

Determinación de razón de momios u Odds Ratio (OR)

Estudio de casos y controles

Éste es un diseño en que se analizan dos grupos de individuos con las siguientes características:

- Individuos enfermos (casos).
- Individuos no enfermos (controles).

El estudio de casos y controles es un diseño en el que a los individuos se les practica una sola medición y son retrospectivos, ya que parten de la enfermedad a la causa o factor de riesgo (E----C). Por ejemplo, para determinar si el tabaquismo es un factor de riesgo con el que se contrae cáncer de pulmón, se tomaron 120 pacientes con cáncer (casos) y 120 pacientes sin cáncer (controles), y en ambos grupos se indagó el antecedente del hábito tabáquico. Los resultados fueron los siguientes (tabla 2):

Tabla 2

Tabla de 2 x 2 para calcular la razón de momios u *Odds Ratio*

Individuos que fumaban	Casos	Controles	Total
Positivo	a 90	b 30	120
Negativo	c 15	d 105	120
Total	105	135	240

Fuente: Elaboración propia.

Para determinar la asociación causal entre el tabaquismo y el cáncer, se determina la razón de momios u *Odds Ratio* (OR) de las dos siguientes formas:

$$A) OR = \frac{a / c}{b / d}$$

- (*a/c entre b/d*). Primero dividimos los enfermos que fumaban entre los enfermos que no fumaban ($90/15 = 6$), luego dividimos los sanos que fumaban entre los sanos que no fumaban ($30/105 = 0.28$) y posteriormente se divide la razón de los casos (a/c) entre la razón de los controles (b/d); en este caso el OR es $6/0.28 = 21$.

$$B) OR = \frac{a \times d}{b \times c}$$

- (*axd entre bxc*). Es la forma más sencilla, consiste en multiplicar $a \times d$ ($90 \times 105 = 9,450$), posteriormente multiplicar $b \times c$ ($15 \times 30 = 450$) y enseguida realizar una simple división; en este caso el OR es igual a $950/450 = 21$.

El resultado 21, se puede interpretar de las siguientes maneras: a) hay 21 casos que fuman por cada control que fuma, y b) el fumar incrementa la probabilidad de enfermar en 21 veces. Al respecto es importante mencionar que en los ejemplos anteriores los cálculos de los riesgos (RR y OR) fueron totalmente diferentes, no obstante que el número de pacientes estudiados fue el mismo en ambos estudios. La explicación es la siguiente: en los estudios de cohorte la prevalencia que se encuentra a través del seguimiento es la *real*, o sea, es muy parecida a lo que sucede en el resto de la población, ya que los investigadores parten de dos grupos de individuos (con y sin el factor de riesgo) y desconocen cuántos de ellos desarrollarán la enfermedad, mientras que en los estudios de casos y controles, los investigadores asignan con anticipación cuántos individuos serán casos y cuántos controles. En los ejemplos anteriores, se presentaron 135 enfermos de 240, es decir, la prevalencia de enfermedad fue de 56% ($135/240$), o sea, más de la mitad de la población padece cáncer de pulmón, lo cual es irreal, por eso la forma de calcular los riesgos son diferentes, sin embargo, como se observa el OR fue notablemente mayor; debido a esto se aconseja que en los estudios de casos y controles, preferentemente debe haber 1 caso y 2 o más controles.

Es importante saber que cuando calculemos el RR u OR, también deberemos determinar sus intervalos de confianza de 95% (IC de 95%). Como sabemos el IC de 95% tiene un límite superior y uno inferior. El IC de 95% puede calcularse en forma manual, sin embargo, es muy complicado, por lo que optamos por hacerlo a través de paquetes estadísticos. Asimismo debemos saber que para considerar un RR u OR como factor causal, el lí-

mite de confianza inferior debe ser >1 . Por ejemplo, se encontró que el OR del tabaquismo para cáncer de pulmón fue de 3.4 (IC de 95%, 2.2 a 4.6); en este caso el promedio del OR fue de 3.4 y el límite inferior de 2.2, luego entonces es significativo y sí existe asociación causal. Por otra parte, en ocasiones los factores de riesgo pueden ser protectores, esto es, su presencia disminuye la frecuencia de la enfermedad; en este caso el promedio del RR u OR y el límite de confianza superior debe ser <1 .

- Ejemplo. En un estudio previo se encontró que el RR del ejercicio para *diabetes mellitus* fue de 0.4 (IC de 95%, 0.2 a 0.6). En este caso se visualiza que el ejercicio protege para la *diabetes mellitus*.

Reducción absoluta de riesgo (RAR)

Es una medida que se utiliza en ensayos clínicos; es idéntico al riesgo atribuible (RA) y se realiza de la siguiente manera. Por ejemplo, en un ensayo clínico se compararon dos grupos de tratamientos y al final del estudio se encontró que en el grupo 1 la mejoría fue de 30%, mientras que en el grupo 2 fue de 60%; en este caso la reducción absoluta de riesgo es 30% (30%-60%) a favor del segundo grupo.

Número necesario a tratar (NNT)

Se utiliza en ensayos clínicos mediante la fórmula $100/\text{RAR}$. Por ejemplo, en un RAR de 30%, el NNT es $100/30 = 3.3$, o sea, se necesitan tratar 3 pacientes para obtener un éxito.

Pruebas diagnósticas

Antes de abordar este tema, es necesario definir los siguientes conceptos:

- *Validez*. Es el grado en que una prueba mide lo que se supone debe medir. La *sensibilidad* y la *especificidad*

son medidas de validez que forman parte de las pruebas diagnósticas.

- *Reproductividad*. Es la capacidad de una prueba para ofrecer los mismos resultados cuando se utiliza en circunstancias similares.
- *Seguridad*. Se refiere a la certeza o confianza que una prueba ofrece para detectar la presencia o ausencia una enfermedad. Los *valores predictivos positivo o negativo*, y *razones de probabilidad*, son *mediciones de seguridad*.

En el área médica existen diversas maneras de detectar la enfermedad, ya sea través de un síntoma, signo, dato de laboratorio, gabinete o biopsia. Generalmente la biopsia es la que presenta el porcentaje más alto para detectar una enfermedad. Sin embargo, en ocasiones la biopsia puede ser un procedimiento cruento, por lo que se busca un método alternativo que ofrezca casi la misma certeza para detectar una enfermedad y que no provoque incomodidad o riesgo al paciente. A la prueba que proporciona la mayor certeza para detectar correctamente la presencia o ausencia de una enfermedad se llama *prueba o regla de oro*, *prueba irrefutable*, *método de referencia* o simplemente *gold standar*.

Actualmente casi todas las enfermedades cuentan con su prueba de oro. El avance en la tecnología ha provocado que las pruebas de oro cambien y cada año muchas cambian. Por ejemplo, en los años setenta la prueba de oro para detectar úlcera péptica era la serie esofagogastroduodenal y en la actualidad es la endoscopia; para diagnosticar litiasis vesicular se usaba la colecistografía oral y ahora ultrasonido. La idea es que cada vez las pruebas sean más confiables y cómodas para el paciente. Pero, para llegar a esto es necesario utilizar nuevas pruebas y éstas deben compararse con las pruebas de oro. Para evaluar la efectividad de las pruebas nuevas, debemos determinar *la sensibilidad*, *especificidad*, *valor predictivo positivo*, *valor predictivo negativo* y *curva de eficacia diagnóstica* o *curva característica operativa del receptor*, o *curva ROC*.

- *Procedimiento.* Se seleccionan dos grupos de individuos: 1) individuos enfermos diagnosticados con la prueba de oro o *gold standar* y 2) individuos sanos de acuerdo con la prueba de oro. Generalmente se utilizan cien enfermos y cien sanos. Sin embargo, existen fórmulas para el cálculo que presenten tamaño de muestra para diseños de una prueba diagnóstica. Posteriormente, a los dos grupos se les aplica la *nueva prueba* diagnóstica y, después de ésta, obtendremos los siguientes cuatro grupos representados en la siguiente tabla (3).

Tabla 3

Tabla 2 x 2 utilizada para pruebas diagnósticas

Prueba nueva	Enfermo	Sanos	Total
Positiva	a 85 (VP)	b 40 (FP)	125
Negativa	c 15 (FN)	d 60 (VN)	75
Total	100	100	200

VP = Verdaderos positivos; FP = Falsos positivos; FN= Falsos negativos; VN = Verdaderos negativos.
Fuente: Elaboración propia.

Cuando el interés está centrado en saber cuántos sanos o enfermos fueron positivos o negativos a la prueba, realizamos los siguientes índices o porcentajes:

Sensibilidad o índice de resultados verdaderos positivos

Es el porcentaje de individuos enfermos que fueron positivos a la nueva prueba. De acuerdo con la tabla 3, la sensibilidad se realizaría así: $= VP/VP+FN = (85/100) \times 100 = 85\%$. O sea, 85% de los enfermos fueron positivos o verdaderos positivos (VP) y el 15% de los enfermos fueron negativos o falsos negativos (FN). Conclusión: La sensibilidad es la capacidad que posee la prueba para detectar enfermos.

Especificidad o índice de resultados verdaderos negativos

Es el porcentaje de individuos sanos que fueron negativos a la nueva prueba. De acuerdo con la tabla 3, la especificidad se realizaría así: $= VN/VN+FP = (60/100) \times 100 = 60\%$. O sea, 60% de los sanos fueron negativos o verdaderos negativos (VN) y 40% de los sanos fueron falsos positivos (FP). Conclusión: La especificidad es la capacidad que posee la prueba para detectar sanos.

Ahora, nos interesa saber o predecir la exactitud de una prueba, por lo que se utilizan los siguientes índices: Valor predictivo positivo (vpp) se define como el porcentaje de los individuos positivos a la nueva prueba que resultaron enfermos y, de acuerdo con la tabla 3 se determina así: $VP/VP+FP = (85/125) \times 100 = 68\%$. Conclusión: El vpp es la probabilidad de padecer la enfermedad si se obtiene un resultado positivo en la prueba.

Por su parte, el valor predictivo negativo (vpn) es el porcentaje de individuos negativos a la nueva prueba que resultaron sanos y, de acuerdo con la tabla 3, se determina así: $VN/VN+FP = (60/75) \times 100 = 80\%$. Conclusión: El vpp es la probabilidad de que un sujeto con un resultado negativo en la prueba esté realmente sano.

Sin embargo, el cálculo de los vpp y vpn presentan errores ya que no toman en cuenta la prevalencia de la población estudiada. Por ejemplo, si se realizara un diseño de una prueba diagnóstica y se seleccionaran 100 enfermos y 100 sanos, la frecuencia o prevalencia de la enfermedad sería de 50%. Ahora bien, si incrementamos el número de enfermos a 200 y continuamos con 100 sanos, entonces la prevalencia sería de 66.7%, los vpp cambian, mientras que la sensibilidad y especificidad permanecen estables.

Por lo anterior, queda claro que la prevalencia influye en los vpp o seguridad de una prueba. Debido a esto actualmente es necesario determinar las siguientes mediciones para determinar el grado de incertidumbre que presenta una prueba:

- *Cociente o Razón de probabilidad positiva (RP+)* = *sensibilidad / 1- especificidad* = $(a/a+c) / (b/b+d)$. Es el resultado de dividir el porcentaje de enfermos con la prueba + y el porcentaje de sanos con la prueba positiva (falsos positivos). De acuerdo con los resultados de la tabla 3 (= 85%/40% = 2.1), significa que hay 2 enfermos con prueba positiva y 1 sano con la prueba positiva.
- *Cociente o Razón de probabilidad negativa (RP-)* = *falsos negativos/ especificidad* = $(c/a+c) / (d/b+d)$. Resulta de dividir el porcentaje de enfermos con la prueba negativa o falsos negativos entre el porcentaje de sanos con la prueba negativa o verdaderos negativos. De acuerdo con los resultados de la tabla 3 (= 15%/60% = 0.25). Significa que hay 0.25 enfermos con prueba negativa y 1 sano con la prueba negativa.
- *Probabilidad preprueba*. Se refiere a la probabilidad de padecer una enfermedad y equivale a la prevalencia.
- *Razón de preprueba (RPre)* = . Es la división entre la probabilidad de que ocurra un fenómeno o prevalencia y la probabilidad de que no ocurra. O sea, se refiere al número de veces que la prueba sería positiva. Por ejemplo, si la probabilidad es de 80%, entonces la no probabilidad es de 20%, luego entonces la $RPre = 80/20 = 4$, significa que la prueba sería positiva en 4 y negativa en 1.
- *Cociente posprueba (CPos)*. Resulta de la multiplicación de la $RPre \times RP+$ en este caso $RPos = 0.4 \times 2.1 = 0.84$. = 0.456.
- *Probabilidad posterior (PP)*. $CPos / 1 + CPos$. Es la probabilidad que resulta de la división entre la posprueba (CPos) y $1 + CPos$; en este caso, la $PP = 0.456 / 1 + 0.456 = 0.313$.

Glosario

Causa. Se define como la condición que sola o acompañada produce una enfermedad.

Factor de riesgo. Son las variables o conjunto de condiciones físicas, químicas, biológicas, psicológicas, y sociales que preceden y se asocian con la enfermedad o incrementan la probabilidad de enfermar.

Mediciones del riesgo. Conjunto de mediciones utilizadas para evaluar la asociación causal entre variables.

Riesgo. Probabilidad de desarrollar una enfermedad por estar expuestos a determinados factores.

Validez. Es el grado en que una prueba mide lo que se supone debe medir.

Reproductividad. Es la capacidad de una prueba para ofrecer los mismos resultados cuando se utiliza en circunstancias similares.

Seguridad. Se refiere a la certeza o confianza que una prueba ofrece para detectar la presencia o ausencia de una enfermedad.

Sensibilidad. Capacidad de la prueba para detectar enfermos.

Especificidad. Capacidad de la prueba para detectar sanos.

Cociente o razón de probabilidad positiva. Es una medida de exactitud y es la posibilidad de que una prueba diagnóstica sea positiva en enfermos.

Cociente o razón de probabilidad negativa. Es una medida de exactitud y es la posibilidad de que la prueba diagnóstica sea negativa en sanos.

Autoevaluación

- ¿Qué es la causalidad y cuáles son sus principales modelos?
- ¿Cómo se define el riesgo de una enfermedad?
- ¿Qué es un factor de riesgo y qué criterios debe cumplir para ser una causa de enfermedad?
- ¿Cuáles son las medidas estadísticas que permiten evaluar la causalidad del riesgo en los estudios clínicos?
- ¿Qué son la sensibilidad y especificidad de una prueba diagnóstica? ¿Cuáles son sus aplicaciones?

Bibliografía

- Dawson, B., y Trapp, R. (2001). *Basic & Clinical Biostatistics*. New York: Lange Medical Books-McGraw-Hill.
- Downie, N.M., y Heath, R.W. (1986). *Métodos estadísticos aplicados*. México D.F.: Harla.
- Fleiss, J.L. (1981). *Statistical Methods for Rates and Proportions*. Nueva York: John Wiley & Sons.
- Longo D.; Kasper D.; Jameson L., et al. (2012). *Harrison's. Principles of Internal Medicine* (18ª ed.). New York: The McGraw-Hill Companies.
- Pita Fernández, S., y Pértegas Díaz, S. (2003). Pruebas diagnósticas: Sensibilidad y especificidad. En: *Cuadernos de Atención Primaria*, 10, pp. 120-124.
- Sackett, D.L.; Straus, S.E.; Richardson, W.S., et al. (2001). *Medicina basada en la evidencia*. Madrid: Harcourt.
- Sokal, R.R., y Rohlf, F.J. (1995). *Biometry*. New York: W.H. Freeman and Company.

CAPÍTULO V

Diseños de estudio

Rebeca Millán Guerrero

Introducción

El término *diseño* se deriva etimológicamente del término italiano *disegno*, que significa *dibujo*; es el proceso previo de configuración mental, *pre-figuración*, en la búsqueda de una solución en cualquier campo.

Por su parte, *diseño de investigación* es el conjunto de actividades coordinadas e interrelacionadas que deberán realizarse para responder la pregunta de la investigación. El diseño debe señalar todo lo que debe hacerse, de tal forma que cualquier investigador con conocimiento en el área pueda alcanzar los objetivos del estudio, responder las preguntas que se han planteado y asignar un valor de verdad a la hipótesis de la investigación. El diseño provee de todos los elementos necesarios para contrastar la hipótesis de la investigación.

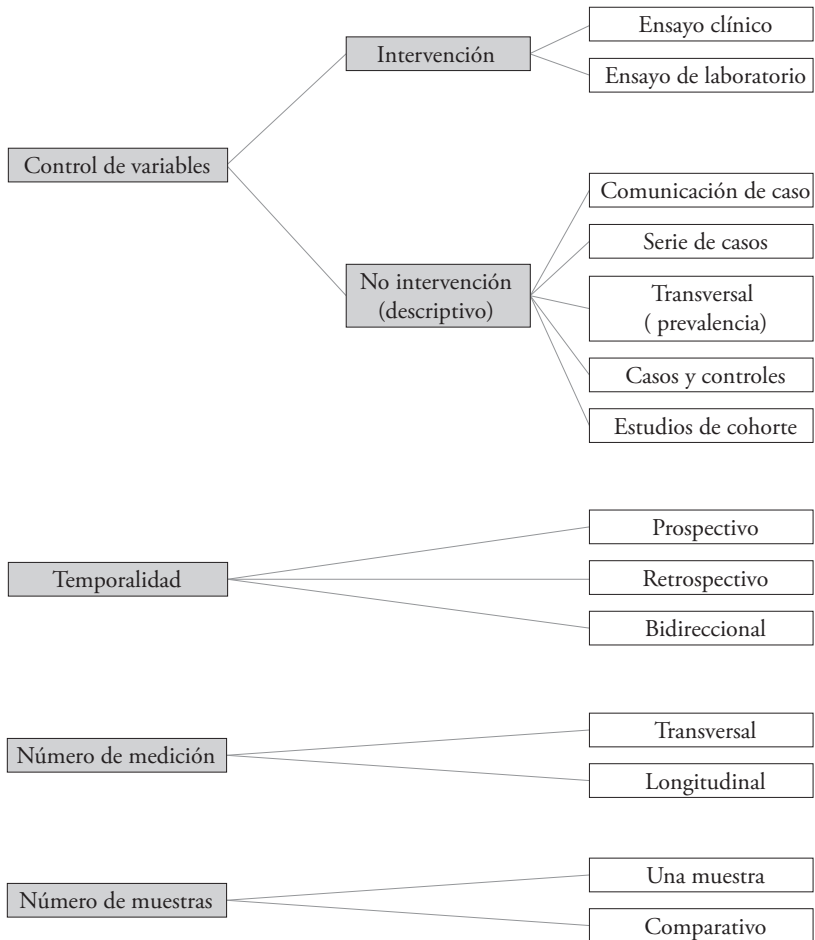
En la investigación metódica del conocimiento clínico, se diseñan estudios y se clasifican en función de la manera en que el investigador lleva el control de las variables de estudio y la dirección temporal que le dé; si el investigador elige mantenerse al margen y realizar una descripción de los hechos que ocurran en los individuos estudiados (estudio observacional), si su siguien-

te decisión es realizar mediciones es una sola ocasión (estudio transversal) o a lo largo de un periodo de tiempo (estudio longitudinal), asimismo si prefiere realizar una comparación y hacer un análisis de los factores potencialmente involucrados en el problema que se estudia o estudiar los efectos de una intervención sobre esos hechos (estudio de intervención [Thompson y Vega, 2001; Hulley *et al.*, 1997: 3-19]).

No existe un diseño preferible a los demás; para cada pregunta a investigar, se debe elegir el diseño más eficaz para obtener una respuesta satisfactoria (figura 3). Los errores en el diseño del estudio, son una razón muy habitual de donde se obtiene una respuesta equivocada a la pregunta que se investiga (Hulley *et al.*, 1997: 3-19).

El objetivo de este capítulo es describir cada uno de los diseños de investigación y dar ejemplos de cómo y cuándo emplear cada uno de ellos.

Figura 3
Diseños de estudio



Fuente: Tomado de Hulley *et al.*, 1997.

Tipos de diseños

Descriptivos

La característica principal de este tipo de estudios, es que no son manipuladas las variables de estudio. Este tipo de diseño es el indicado para estudiar un problema nuevo. El objetivo de estos estudios es explorar el terreno, describir la distribución de las en-

fermedades y características de una población; en cierta forma, se tiene un débil control sobre la población. En este diseño el investigador se concreta a observar las circunstancias en las que ocurren los eventos en forma natural. Se estudia una sola población y no se tiene hipótesis central, aunque puede servir para sugerir una hipótesis. El diseño descriptivo se utiliza para describir la historia natural o cuadro clínico en un grupo de pacientes, con una enfermedad en particular; también, para medir incidencia y prevalencia de enfermedades en una población.

Los tipos de estudios descriptivos son los siguientes:

- *Comunicación de un caso.* Describe la característica de un solo paciente y puede ser el reporte de una enfermedad rara o nueva. Estos estudios son muy valiosos y representan el vínculo entre la investigación clínica o de laboratorio; su desventaja es que son susceptibles a sesgos.
- *Comunicación de una serie de casos.* Es una encuesta de prevalencia realizada a un grupo de pacientes con una enfermedad determinada; se lleva a cabo en un momento único del tiempo. Este diseño tiene importancia en epidemiología y con frecuencia identifica el inicio de una epidemia. Su desventaja es que no se puede probar la presencia de asociaciones estadísticas válidas.
- *Estudios transversales.* Son también llamados estudios de prevalencia, porque nos permiten conocer la prevalencia de padecimientos o características de una población; con ellos no se determina causalidad, pero se conoce la magnitud del problema. El estímulo y la enfermedad se valoran en forma simultánea entre individuos de una población definida. Es como una fotografía instantánea de lo que ocurre en una comunidad. Suelen ser los primeros estudios en la búsqueda de los factores que se asocian con los problemas de salud y la distribución de éstos en la población. En estas investigaciones se define la

exposición y se estudia la frecuencia de la enfermedad. El límite de estos estudios es que se desconoce cómo ocurrió primero, ya que en el momento de la medición coinciden el estado basal, maniobra y resultado (Talavera y Rivas-Ruiz, 2012). Entre los diseños transversales se incluyen la encuesta transversal y el estudio de casos y controles:

La encuesta transversal

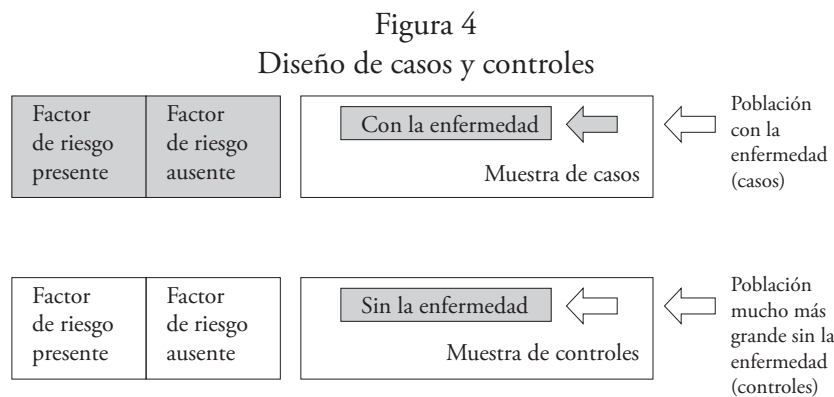
Éste es probablemente el diseño más utilizado en la investigación médica y se caracteriza por estudiar a una población en la cual los datos son recolectados a un mismo tiempo y en forma retrospectiva; es decir, al comenzar el análisis ya ocurrió el desenlace y la exposición a la maniobra; este estudio solamente lleva una medición o registro. En la encuesta transversal, en forma artificial se reconstruye la secuencia temporal de la enfermedad; la precisión de los datos depende del paciente y la medición de las variables puede ser errónea, ello se suma a que es difícil asociar un desenlace con una probable causa. Dentro de sus ventajas se encuentra la utilidad para describir variables y sus patrones de distribución, también para establecer la prevalencia o proporción de la población que presenta la enfermedad en un momento dado, además de que se realiza en corto tiempo, no hay riesgo de pérdidas en el seguimiento y su costo es menor.

Como desventaja, no existe la certeza de la relación temporal, aun cuando se establece asociación entre variables. No es útil para enfermedades poco frecuentes y resulta limitado para conocer la evolución natural de la enfermedad o pronóstico.

Estudio de casos y controles

Éste se trata de un estudio observacional, retrospectivo, en el cual los sujetos se seleccionan sobre la base de la presencia o no de la enfermedad o fenómeno a estudiar y luego se observa hacia atrás en el tiempo para detectar diferencias en las variables predictoras que pudiesen explicar por qué los casos desarrollan la enfermedad y los con-

troles no. Este diseño permite estudiar enfermedades con periodos de latencia largos; el investigador identifica a los individuos afectados y los observa retrospectivamente, sin tener que esperar a que se desarrolle la enfermedad. Los estudios de casos y controles parten de identificar pacientes con una enfermedad y se buscan los factores asociados a ella. En este caso se compara lo que acontece en dos grupos de personas: uno con la enfermedad, o fenómeno en estudio (casos), y el otro sin la enfermedad, o el fenómeno estudiado (controles). Después se determina la proporción de personas expuestas a uno o más factores que se sospecha están asociados con la enfermedad y determina el riesgo de enfermarse (Thomson y Vega, 2001). Al respecto véase la figura 4.



Fuente: Tomado de Hulley *et al.*, 1997: 84-95.

Los pasos para diseñar un estudio de casos y controles son los siguientes:

1. Seleccionar una muestra a partir de la población de individuos enfermos (casos).
2. Seleccionar una muestra a partir de una población de riesgo sin enfermedad (controles).
3. Medir las variables predictoras.

Lo que sí aporta este estudio (diseño de casos y controles) es información descriptiva sobre la características de los casos, todavía más importante, una estimación de la fuerza de asociación entre cada variable predictora. Esta estimación se presenta en forma de *odds ratio* (riesgo relativo indirecto o razón de *odds*, para lo cual utilizamos una tabla de 2 x 2 celdas (tabla 4):

Tabla 4
Análisis de casos y controles

Riesgo relativo indirecto	Casos	Controles
Expuestos	A	B
No expuestos	C	D

Riesgo = AD/CB .

Fuente: Elaboración propia.

Ejemplo: Al buscar asociación entre tabaquismo y cáncer pulmonar, integramos dos grupos, uno de éstos con el cáncer pulmonar y otro de similares características pero sin cáncer. Para ello interrogamos en ambos grupos cuántos elementos fumaron a lo largo de su vida. Los resultados se resumen en la tabla 5.

Tabla 5
Análisis de tabaquismo y cáncer pulmonar

Riesgo relativo indirecto	Caso (cáncer)	Controles (no cáncer)
Fumaron	16	4
No fumaron	4	16

Riesgo = $(16 \times 16) / (4 \times 4) = 256 / 16 = 16$.

Fuente: Elaboración propia.

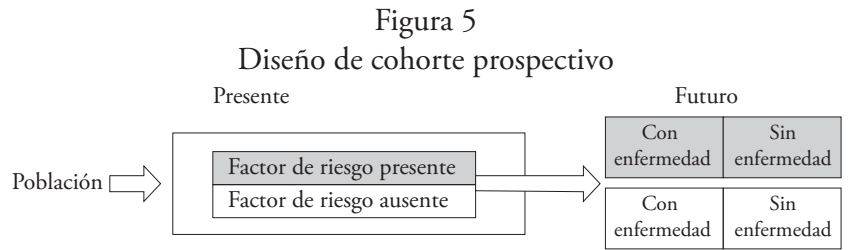
Con ello concluimos que una persona fumadora tiene 16 veces más posibilidades de desarrollar cáncer pulmonar que una persona no fumadora.

Las ventajas de estos estudios es que son baratos, es fácil identificar los casos, no es necesario esperar periodos largos para tener una respuesta y tienen mayor utilidad en enfermedades

poco frecuentes o con tiempo de latencia prolongado. Las limitaciones más importantes de mencionar, son que tanto el estímulo como la respuesta ya ocurrieron, no son útiles para conocer la relación temporal de los fenómenos, ya que los registros se basan en el pasado, por lo que la validación de la respuesta es difícil y no permiten comprobar la causalidad; además, son muy susceptibles a sesgos de selección, memoria y búsqueda. Sirven para analizar una o varias causas pero solamente un desenlace.

Estudios de cohorte

Este tipo de estudios implica efectuar un seguimiento de grupos, conformados por individuos, a lo largo del tiempo, y presenta dos propósitos principales: ser descriptivo y analítico (*cohorte* es un antiguo término romano que definía a un grupo de soldados que marchaban juntos hacia la batalla). Se trata de un estudio prospectivo, es decir, el investigador elige la muestra que todavía no presenta el desenlace de interés (Hulley *et al.*, 1997: 69-81). Véase la siguiente figura (5).



Fuente: Tomado de Hulley *et al.* (1997: 69-81).

Los pasos para realizar este tipo de estudio son los siguientes:

1. Seleccionar una muestra a partir de una población.
2. Medir las variables predictoras (factor de riesgo presente o ausente).
3. Efectuar un seguimiento de la cohorte.
4. Medir las variables de desenlace (enfermedad presente o ausente).

Normalmente los estudios descriptivos van seguidos de estudios analíticos en los que se observa asociación de variables, con el fin de descubrir relación causa-efecto y finalmente el paso final puede ser un estudio de experimento o intervención, para determinar los efectos de una intervención. Con frecuencia se presenta el ensayo aleatorizado como la solución óptima, pero todo dependerá de la pregunta del estudio.

Ensayos clínicos

Este diseño es un estudio de intervención, aporta evidencia científica sólida y debe tener tres características necesarias (Hulley *et al.*, 1997: 123-132):

- *Aleatorización.* Es un proceso mediante el cual se asignan al azar los sujetos en estudio, para evitar sesgos de selección en los grupos de estudio. Es decir, distribuye equitativamente las variables de confusión en los grupos experimental y testigo, y da validez a las pruebas estadísticas. Un proceso aleatorio es regido por el azar; es decir, no está determinado por el investigador. Dentro de los límites que impone la variación aleatoria, el investigador debe procurar que los grupos control y experimental sean similares al comienzo de la investigación y debe asegurarse de que los juicios personales no ejerzan influencia sobre el modo de asignación. Los diseños aleatorios comprenden el diseño paralelo y el cruzado. El diseño paralelo es aquel en el cual un grupo recibe una maniobra y se compara con otro grupo que no la recibe. La característica de este diseño es que los grupos en estudio reciben sus tratamientos respectivos en forma simultánea.
- *Cegamiento o enmascaramiento.* Disminuye el sesgo de quienes reciben la intervención, de quienes la administran y de quienes analizan la información. El cegamien-

to puede ser simple, cuando el investigador desconoce quién recibe la intervención. Lo recomendado es el cegamiento doble, es decir, cuando el investigador y el sujeto a estudiar desconocen el tipo de intervención recibida y puede existir el cegamiento triple cuando el investigador que analiza resultados, también está cegado.

- *Control.* Debe existir por lo menos un grupo control que no recibe la intervención, y quien recibe placebo ya sea con acción inerte, o la forma activa la cual es generalmente el tratamiento habitual que se indica para determinado problema.

Además, existen los estudios pseudoexperimentales: *no aleatorizado*, y *no cegado*. Las consecuencias de no realizar un ensayo clínico apropiado son pérdida de tiempo y dinero, además de que no se considera ético. Por ejemplo, si se desea determinar la eficacia de un nuevo agente quimioterapéutico en pacientes leucémicos, se requiere emplear como grupo testigo a otro grupo de leucémicos con el tratamiento quimioterapéutico más aceptado hasta ese momento; no sería ético tener un grupo control con placebo o sin tratamiento farmacológico.

Un ensayo clínico comienza por la formulación de una hipótesis. La asignación aleatoria permite balancear las diferencias entre los grupos a comparar. El enmascaramiento disminuye el sesgo de quienes reciben la intervención, quienes la administran y quienes procesan la información. El tamaño adecuado de la muestra disminuye el error aleatorio al tomarla. Las ventajas y desventajas de los ensayos clínicos controlados son las siguientes:

- *Ventajas.* El investigador diseña un protocolo de investigación en el que define mecanismos de control que operarán antes y durante el desarrollo de la fase experimental con el objeto de cautelar la seguridad del sujeto en quien se experimentará.

- *Desventajas:*

Complejidad. La posibilidad de manipular la variable independiente, determinar causalidad y experimentar en seres humanos, confiere a los ensayos clínicos un alto grado de complejidad.

Costo. La naturaleza de los estudios clínicos experimentales exige el uso de productos biológicos, farmacológicos o procedimientos terapéuticos y de control y monitoreo no exentos de costo.

Diseño cruzado

Éste permite que el sujeto en estudio sea su propio control. En este diseño cada sujeto recibe la intervención en forma secuencial. Esto es, a un grupo se le administra el tratamiento A y a otro el B, después se le mantiene en un periodo sin tratamiento (*lavado*) y el grupo que recibió al principio el tratamiento A recibe después el B y viceversa.

Los diseños cruzados tienen ventaja sobre los paralelos en relación con su eficiencia estadística ya que se requiere de menor número de sujetos y los costos son menores, pero existe un problema para este diseño, sobre todo cuando se emplean fármacos con larga acción terapéutica (meses), tiempo que deberemos esperar para el lavado y poder continuar con el estudio.

Los ensayos clínicos son una herramienta fundamental para evaluar la eficacia de los medicamentos y otros procedimientos terapéuticos. También optimizan el estudio de la cinética y algunos aspectos de la seguridad de los mismos. Cuando surge un fármaco con posible eficacia terapéutica se debe realizar la llamada *farmacología terapéutica* que presenta las siguientes fases:

- *Fase preclínica.* Se realiza en animales y deben ser tres especies diferentes, esta fase sirve para estudiar lo siguiente: Toxicidad estructural aguda: daño hepático, renal, de médula ósea.

Toxicidad fisiológica o funcional: sistema, cardiovascular, respiratorio, sistema nervioso central, etcétera.

Farmacocinética: valores de concentración en sangre, orina y tejidos, vida media y sus metabolitos.

- *Fase clínica I.* Es el primer paso en la investigación de un nuevo fármaco en seres humanos. Se realiza salvo excepciones, en individuos sanos y se obtienen datos preliminares sobre la tolerancia, farmacocinética y farmacodinamia.
- *Fase clínica II:* Estudio abierto, no cegado, en sujetos con la enfermedad y se realiza para observar lo siguiente: La evaluación de un nuevo fármaco en el ser humano. Se realiza en un grupo reducido de pacientes que padecen la enfermedad o la entidad clínica de interés, con criterios de selección, en general estrictos. Su principal objetivo es obtener información preliminar sobre la eficacia del fármaco, así como complementar los datos de seguridad obtenidos en la fase I. Los objetivos de la fase clínica II son los siguientes: 1) Demostrar si el nuevo fármaco es eficaz para una o más indicaciones clínicas. 2) Realizar una estimación inicial de la relación riesgo-beneficio. 3) Establecer la eficacia respecto a otros fármacos o métodos de tratamiento. 4) Estos estudios deben hacerse sólo si se cuenta con un protocolo bien establecido. El seguimiento de los protocolos es supervisado. Al final de la fase clínica II habrá una decisión clara de si el fármaco debe o no ser desarrollado como *agente terapéutico*.
- *Fase clínica III o ensayo clínico.* Éste es aleatorizado, cegado y controlado. Sus características son las siguientes: Tiene como objetivo evaluar la eficacia y seguridad de un fármaco en la indicación estudiada, por comparación con las alternativas disponibles. Estos estudios constituyen el soporte para la solicitud de registro y la autorización de comercialización de un fármaco.
- *Fase clínica IV.* Realiza estudios de fármacos existentes en el mercado. En esta fase se incluyen todos los aspectos

relativos a la investigación de fármacos una vez que ya han sido aprobados para su venta.

Después del ensayo clínico, el segundo diseño de investigación con mayor calidad de información es la cohorte; el clásico ejemplo de este último diseño lo tenemos en el famoso estudio *Framingham* (McKee *et al.*, 1971), que permitió conocer los factores de riesgo cardiovasculares varios años después de haberse iniciado y más tarde los factores de riesgo cerebrovasculares; este estudio implicó años de seguimiento, alto costo, tres generaciones de enfermos y tres de investigadores, pero el esfuerzo valió la pena.

Todos los diseños de investigación son valiosos y la elección dependerá del interés que persiga el investigador.

Conclusión

En este capítulo hemos descrito los diferentes tipos de diseños que pueden seguirse al realizar una investigación, sin embargo ninguno es mejor que otro, ya que debe elegirse el diseño que pueda responder a la pregunta planteada.

Glosario

Diseño de un estudio. Clasificación de los estudios en función de la manera en que el investigador lleva el control de las variables de estudio y la dirección temporal que le dé.

Diseños descriptivos. Tipo de estudios en los cuales no son manipuladas las variables de estudio. Este tipo de diseño es el indicado para estudiar un problema nuevo.

Estudios de cohorte. Este tipo de estudios implica efectuar, a lo largo del tiempo, un seguimiento de grupos conformados por individuos. Su propósito es descriptivo y analítico. Se trata de un estudio prospectivo, donde el investigador elige la muestra que todavía no presenta el desenlace de interés.

Estudio observacional. En este tipo de estudio no existe intervención en las variables.

Estudio transversal. Son también llamados estudios de prevalencia; nos permiten conocer la prevalencia de padecimientos o características de una población. El estímulo y la enfermedad se valoran en forma simultánea entre individuos de una población definida. Es como una fotografía instantánea de lo que ocurre en una comunidad.

Estudio longitudinal. Estudio de prevalencia en el cual las variables se miden varias veces a lo largo de un periodo.

Estudio de intervención. Estudio en el cual se intervienen las variables; son una herramienta fundamental para evaluar la eficacia de los medicamentos y otros procedimientos terapéuticos.

Autoevaluación

- ¿Qué es un diseño de investigación?
- ¿Qué aspectos determinan las características que debe contener un diseño de estudio?
- ¿Cómo clasifica los diseños de estudio propuestos por la autora?
- ¿Por qué la selección de un diseño de estudio debe basarse, sobre todo, en las preguntas de investigación que se quieran responder?

Bibliografía

- Hulley, S.B.; Gove, S., y Cummings, S.R. (1997). Concepción de la pregunta a investigar. En: S. B. Hulley, y S. R. Cummings (eds.), *Diseño de la Investigación Clínica* (pp. 3-19). Madrid: Harcourt Brace.
- Hulley, S.B.; Gove, S., y Cummings, S.R. (1997). Diseño de un nuevo estudio: I. Estudios de cohorte. En: S. B. Hulley, y S. R. Cummings (eds.), *Diseño de la Investigación Clínica* (pp. 69-81). Madrid: Harcourt Brace.
- Hulley, S.B.; Gove, S., y Cummings, S.R. (1997). Diseño de un nuevo estudio: II. Estudios transversales y estudios de casos y controles. En: S. B. Hulley, y S. R. Cummings (eds.), *Diseño de la Investigación Clínica* (pp. 84-95). Madrid: Harcourt Brace.
- Hulley, S.B.; Gove, S., y Cummings, S.R. (1997). Diseño de un nuevo estudio: IV. Experimentos. En: S. B. Hulley, y S. R. Cummings, (eds.), *Diseño de la Investigación Clínica* (pp. 123-132). Madrid: Harcourt Brace.
- McKee, P.A.; Castelli, W.P.; McNamara, P.M., y Kannel, W.B. (1971). The natural history of congestive heart failure: the Framingham study. En: *New England Journal of Medicine*, 285, pp. 1441-1446.
- Talavera, J.O. y Rivas-Ruiz, R. (2012). Investigación Clínica XII. Del juicio clínico a la encuesta transversal. En *Revista Médica del Instituto Mexicano del Seguro Social*, 50, pp. 41-44.
- Thompson, CH.OC., y Vega, F.L. (2001). Diseños de investigación en las ciencias biomédicas. En *Revista Mexicana de Pediatría*, 68, pp. 147-151.

CAPÍTULO VI

Variables y escalas de medición

Benjamín Trujillo Hernández

Introducción

La identificación de las variables es uno de los procesos claves para cualquier tipo de investigación y en innumerables ocasiones motivo de debate entre los estudiantes y asesores. Es frecuente observar que en proyectos de investigación avanzados, ya sean de posgrado o licenciatura, existen controversias entre los profesores que imparten los seminarios de investigación, y asesores, respecto a la identificación de las variables.

En nuestra experiencia, el conflicto principal es la identificación de las variables independientes y dependientes. Además, en la mayoría de las ocasiones hay un problema de conceptualización que no afecta la calidad del trabajo, sin embargo, puede llegar a ser motivo de discusiones encarnizadas entre los protagonistas, lo cual siempre afecta al estudiante. Creo que el problema radica en que existen muchos libros que hablan sobre el tema y clasifican a las variables en diferente manera.

Este apartado pretende contribuir a resolver este problema; con tal fin se sustenta en la revisión de diversos autores y se le agregó una pizca de experiencia; ambas nos han resultado eficaces para publicar manuscritos médicos en los ámbitos nacional e internacional.

Variable

Una variable es todo fenómeno físico, químico, biológico, psicológico, o social, que puede medirse y que varía entre los grupos. Esto es, todo lo que puede medirse es una variable.

Clasificación de las variables

Por su naturaleza

Se refiere a cómo podemos detectar o percibir a las variables o fenómenos. De acuerdo con lo anterior éstas pueden ser cualitativas y cuantitativas.

- *Cualitativa*. Se refiere a una cualidad o característica del fenómeno y para detectarla generalmente utilizamos los órganos de los sentidos, por lo tanto son subjetivos (ya que pueden variar de un observador a otro); ejemplo de estas variables son el color de ojos, etnia, la mayoría de síntomas y signos.
- *Cuantitativa*. Es un fenómeno que puede cuantificarse y generalmente necesitamos instrumentos para detectarla. Éstas se dividen en *continuas* y *discretas*. Las continuas las reportamos en fracciones o decimales, por ejemplo: hemoglobina 10.4 g/dL, peso 85.6 k, talla 1.67 m. Las discretas son aquellas reportadas en número enteros, por ejemplo: número de hijos, y leucocitos en sangre u orina (Dawson y Trapp, 2001; Downie y Heath, 1986; Fleiss, 1981).

Nivel de medición

El nivel de medición es útil cuando recolectamos información o la base de datos y sirve para identificar la característica de las variables estudiadas. De acuerdo con lo anterior las variables pueden ser de cuatro tipos:

- *Nominal*. Aquí anotamos una característica y se refiere a que le damos un nombre o cualidad que lo distingue. Esta característica es exclusiva y mutuamente excluyente,

o sea, no puedes pertenecer a más de una categoría, por ejemplo: el género (hombre o mujer), estado civil, enfermedad (sí o no), religión, tabaquismo (sí o no).

- *Ordinal*. En ocasiones el investigador desea darle un valor que puede ser subjetivo (rango) o al utilizar instrumentos validados (puntaje). El *ordinal de rango*¹ es frecuentemente utilizado en medicina, por ejemplo: grado de disnea (leve, moderada, severa), grado de ictericia, etcétera. El *ordinal de puntaje* utiliza valores numéricos *inventados* pero validados y aceptados por la comunidad médica, por ejemplo las siguientes escalas: APGAR, *Glasgow*, Coeficiente intelectual, APACHE, SOFA y todas las escalas que normalmente sirven para pronóstico. Una variable cuantitativa o de razón puede degradarse a nivel ordinal de puntaje cuando su comportamiento es anormal o tienen varianzas diferentes.
- *Intervalo*. Está casi en desuso, son valores que se clasifican en intervalos y no tienen cero absoluto, por ejemplo grupos de edades: 1-10 años, 11-20 años, 21-30 años, etcétera. Es importante saber que las variables de intervalo, en el análisis estadístico, se “degradan” a nivel nominal.
- *Razón*. Son valores numéricos que cuentan con cero absoluto y son determinados frecuentemente con instru-

¹ Debe tomar en cuenta que las variables ordinales de rango (en el análisis estadístico se *degradan* a variables nominales. Por ejemplo, se evaluó el grado de insuficiencia cardíaca en cuatro grados de acuerdo con la disnea (sin disnea, disnea con ejercicio leve, ejercicio moderado y disnea reposo); según lo anterior, y considerando que las variables son *ordinales*, el análisis estadístico que se cree conveniente utilizar son las siguientes pruebas *no paramétricas*: medianas, rango, U de Mann-Whitney, Wilcoxon, Kruskal-Wallis, etcétera. Sin embargo, se observa que en realidad hay cuatro categorías, o sea, son variables nominales y en este caso la estadística apropiada serían porcentajes y la *chi cuadrada* para comparar porcentajes entre los grupos, mientras que para las variables ordinales de puntaje (las cuales “aparentemente” son numéricas) idóneamente no se les debería realizar estadística paramétrica como promedios o desviación estándar, ya que el puntaje está apoyado en elementos subjetivos. Por ejemplo, la escala de Apgar está sustentada en las características observadas en el niño, esto es, que a pesar de llevar una calificación numérica del 1 al 10, ésta se sustenta en el grado de acuciosidad de cada médico y ésta es la explicación del porqué en muchas ocasiones los pediatras dan calificaciones diferentes al mismo niño. Debe recordarse que los médicos no podemos determinar en forma exacta los síntomas y signos, por eso se han inventado escalas que nos sirven para determinar el grado de enfermedad que tiene el paciente (Dawson y Trapp, 2001).

mentos como báscula, estadímetro, electrocardiograma, electroencefalograma, etcétera. Son ejemplos: peso, talla, presión arterial, parámetros bioquímicos, etcétera.²

Como sugerencia, cuando realice una investigación, preferentemente determine sus variables en nivel de razón. Por ejemplo, en un estudio se les preguntó a los pacientes si fumaban y cuántos cigarrillos fumaban en promedio; con estos datos podemos realizar los siguientes niveles de medición:

<i>Nominal</i>	Fuma o no fuma
<i>Ordinal de rango</i>	Tabaquismo leve (1-5 cigarrillos al día), moderado (10-20) y severo (>20)
<i>Razón</i>	Número de cigarrillos que fuma

Con esta información podemos determinar porcentajes (variables nominales y ordinales de rango), promedios, desviación estándar, coeficiente de variación, rango e intervalo de confianza (variable de razón). También podemos realizar pruebas estadísticas como *chi cuadrado*, *t de Student*, etcétera.

Sin embargo, en otro estudio sólo se les preguntó si fumaban o no. Luego entonces el nivel de medición que se puede realizar es nominal: fuma o no fuma. El análisis estadístico que se pue-

² Las variables de razón generalmente se analizan con estadística paramétrica. La medida de tendencia central es el promedio y las medidas de dispersión son: rango, intervalo, desviación estándar, varianza y coeficiente de variación. Las pruebas estadísticas para comparar promedios son *t Student* (pareada o no pareada), análisis de varianza (ANOVA) o ANOVA de medidas repetidas. Todos los programas estadísticos cuando comparan dos o más promedios, realizan pruebas de comparación de varianzas, para saber si hay diferencia significativa entre las mismas. En caso de que las varianzas sean diferentes o heterogéneas, la estadística que se debe utilizar es la *no paramétrica*. Por ejemplo, supongamos que se compararon los promedios y desviación estándar entre dos grupos de tratamiento: Grupo A) 89.5 ± 17.4 K y Grupo B) 87.4 ± 1.2 K. Desde ese punto de vista los promedios son casi iguales, sin embargo, el programa estadístico realiza la comparación de las varianzas (desviación estándar al cuadrado) con la prueba F (varianza mayor/varianza menor) y encuentra que son diferentes significativamente ($p < 0.05$). La prueba estadística para comparar dos promedios con varianzas iguales es la *t de Student*, sin embargo, como tiene varianzas diferentes, la prueba idónea es la prueba U de Mann-Whitney (prueba *no paramétrica* [Dawson y Trapp, 2001; Downie y Heath, 1986]).

de realizar es descriptivo como porcentajes o *chi cuadrado* en caso de que haya dos o más grupos.

Interrelación de las variables

Se refiere a la relación entre dos o más fenómenos o variables. Esto es, una afecta a la otra. De acuerdo con lo anterior podemos encontrar dos tipos de variables:

- *Variable independiente.* Es la variable precedente, factor de riesgo, causa o la que manipula el investigador.
- *Variable dependiente.* Normalmente es el fenómeno de interés, es el subsecuente, efecto o enfermedad.

Para identificar la relación de las variables se debe pensar en la relación genética de los padres y los hijos. Los padres siempre son primeros (variable independiente) y los hijos siempre son producto de los primeros (variable dependiente). Cuando se desea identificar a estas variables es muy importante saber el diseño metodológico y estadístico. Por ejemplo, en diseños observacionales o descriptivos el objetivo de estos es describir y no les interesa en forma explícita la relación causa-efecto, luego entonces en este tipo de estudio todas las variables *son dependientes*, mientras que en los estudios analíticos o experimentales sí hay variables independientes (causa o grupos de estudio) y dependiente (variable de interés).³

³ En estadística se llama variable independiente al número de grupos de estudio: si hay un grupo no hay variable independiente; si hay dos o más grupos, cada grupo es independiente del otro (grupo A, Grupo B y Grupo C) y las pruebas estadísticas que se deben utilizar son las *no pareadas* ya sean paramétricas (*t de Student* o ANOVA de una vía) o no paramétricas (*chi cuadrada*, U de Mann-Whitney o Kruskal-Wallis). Sin embargo, cuando un grupo (conjunto de individuos) se mide dos o más veces, se llaman grupos dependientes y las pruebas estadísticas son pareadas o dependientes que pueden ser paramétricas (*t Student* pareada o ANOVA de medidas repetidas) y no paramétricas (McNemar, Wilcoxon, Cochran, Friedman, etcétera [Sackett, Straus, Richardson, *et al.*, 2001; Sokal y Rohlf, 1995]).

Reversibilidad de las variables

Es un concepto poco usado y tiene relación con las variables dependiente e independiente y se refiere a la relación causa-efecto. Así, de acuerdo con lo anterior, se dice que las variables son irreversibles cuando A sólo causa B ($A \rightarrow B$) y reversibles cuando A causa B y B causa A ($A \leftrightarrow B$). Las variables irreversibles son las más frecuentes en el área de salud, por ejemplo, lo lógico es que la obesidad condicione o sea la causa de *diabetes mellitus* tipo 2, (obesidad $\rightarrow$ diabetes) o sea la obesidad es la variable independiente (A) y la *diabetes mellitus* tipo 2 la variable dependiente (B), y es poco probable que la relación sea reversible o sea que la diabetes provoque obesidad.

Los ejemplos de variables reversibles son más escasos, pero existen, por ejemplo, hipertensión arterial y nefropatía. En este caso es reversible porque ($A \leftrightarrow B$), esto es, tanto la hipertensión como la nefropatía pueden ser causa o efecto, o sea, cualquiera de las dos pueden ser variable independiente o dependiente. En este caso se busca una relación causal a través de un estudio analítico (estudios de cohorte o casos y controles) y el investigador debe decidir, apoyado en la secuencia lógica del tiempo, a cuál considerará como variable dependiente o independiente. Si el estudio es observacional y no tiene intención de demostrar causa-efecto, entonces las dos variables son dependientes.

Como se observa, conocer los conceptos anteriores nos permite plantear el análisis estadístico o diseño.

Nociones de estadística

La palabra *estadística* se utilizó por vez primera cuando el Estado o países quisieron saber cuántos habitantes existían para cobrar impuestos y para la milicia. En este libro abordaremos el tema con la menor cantidad de ecuaciones matemáticas. Apelamos a que actualmente existen los suficientes programas estadísticos que nos auxilian en el quehacer matemático, sin embargo, el

objetivo de este tema es que el lector comprenda y sea capaz de elegir el diseño o prueba estadística de acuerdo a su propósito.

En este texto definimos a la estadística como la herramienta que permite describir, analizar e interpretar un conjunto de datos o variables. Dentro de esta área son comunes los siguientes dos conceptos:

- *Deducción.* Conocimiento que se obtiene de lo general a lo particular, compuesta por dos premisas y una conclusión.
- *Inferencia.* Conocimiento que se obtiene de lo particular a lo general.

El médico como parte de su heurística (arte de investigar) utiliza el método hipotético-deductivo-inductivo. Esto es, obtiene conocimientos de los libros para aplicarlo a sus enfermos (deducción) pero también obtiene experiencias de algunos pacientes que aplica al resto de sus pacientes (inferencia). La inferencia en la investigación es sumamente importante y para esto es necesario que se tenga entre otros elementos un tamaño de muestra adecuado. De tal forma, si se cumple con los criterios indispensables, los conocimientos obtenidos en una muestra pueden ser aplicados al resto de la población (Dawson y Trapp, 2001; Downie y Heath, 1986; Fleiss, 1981).

Tipos de estadística:

- *Descriptiva.* Sólo describe, utiliza las funciones matemáticas básicas y gráficos. Ejemplos: porcentajes, media, mediana, moda, desviación estándar, varianza, intervalo de confianza, rango, centiles y coeficiente de variación.
- *Inferencial.* Utiliza fórmulas matemáticas o estadísticas, ejemplos: todas las pruebas estadísticas conocidas.

Conclusión

La identificación y características de las variables es altamente necesario para realizar un trabajo de investigación. La falta de esto, provoca que los resultados carezcan de validez interna y externa. De acuerdo con el tipo de variable de interés se clasifica en *paramétrica* (variables de razón) y *no paramétrica* (variables nominales y ordinales).

Glosario

Variable. Todo fenómeno que puede medirse y que varía entre los grupos.

Escalas de medición de las variables. Nominal, ordinal, intervalo y de razón.

Medidas no paramétricas. Mediciones sustentadas en observaciones subjetivas.

Medidas paramétricas. Mediciones sustentadas en observaciones objetivas.

Autoevaluación

¿Qué es una variable?

¿Qué es una escala de medición?

¿Cómo se clasifican las variables por su naturaleza? Señale ejemplos de variables cualitativas y cuantitativas.

¿Cómo se clasifican las variables por su escala de medición? ¿Qué pruebas estadísticas paramétricas y no paramétricas se pueden aplicar para contrastar hipótesis de nulidad en cada una de estas variables?

¿Cómo se clasifican las variables por su interrelación? ¿Qué se debe considerar para estar seguro de que una variable es independiente o dependiente?

¿Cómo se clasifican las variables por su reversibilidad? Dé ejemplos de variables irreversibles en el campo de la salud.

Bibliografía

- Dawson, B., y Trapp, R. (2001). *Basic & Clinical Biostatistics*. New York: Lange Medical Books-McGraw-Hill.
- Downie, N.M., y Heath, R.W. (1986). *Métodos estadísticos aplicados*. México D.F.: Harla.
- Fleiss, J.L. (1981). *Statistical methods for rates and proportions*. Nueva York: John Wiley & Sons.
- Sackett, D.L.; Straus, S.E.; Richardson, W.S., *et al.* (2001). *Medicina basada en la evidencia*. Madrid: Harcourt.
- Sokal, R.R., y Rohlf, F.J. (1995). *Biometry*. New York: W.H. Freeman and Company.

CAPÍTULO VII

Muestreo y tamaño de muestra

Benjamín Trujillo Hernández

Conceptos generales

A continuación se enlistan los conceptos que generalmente se utilizan en el muestreo:

Población o universo

Es un grupo de individuos que comparten una característica en común, ejemplos: mineros, trabajadores del área de la salud o región geográfica, etcétera. Debe mencionarse que el número de individuos pertenecientes al universo o población es variable; por ejemplo, si la variable fuera *mexicano*, el universo o población sería el país, si la variable fuera *mexicano del norte del país*, entonces los individuos de esta zona serían el universo.

Población de referencia

Algunos libros hacen mención de este concepto y se refiere a un subgrupo de individuos que comparten dos o más características y tiene relación con los criterios de inclusión. Por ejemplo, ser mexicano y ser originario del estado de Colima. Es importante aclarar que este concepto crea frecuentemente confusión y en muchas ocasiones no mejora la calidad de la investigación.

Muestra

Subconjunto de individuos que a través de una fórmula matemática (fórmula para calcular el tamaño de muestra) representan a la población o universo.

Muestreo

Actividad o herramienta utilizada para saber el porcentaje de la población que debe examinarse y representar fielmente a la población. Un mal muestreo puede provocar errores de muestreo o inferencia (los resultados de la muestra son erróneos y por desconocimiento de los investigadores se aplican al resto de la población).

Validez interna

Se refiere a que los resultados obtenidos en la muestra son válidos solo para el grupo de individuos o pacientes estudiados.

Validez externa o capacidad de generalización o inferencia

Se refiere a que los resultados obtenidos son válidos tanto para la muestra como para el resto de la población o universo.

Tipos de muestreo

Pueden dividirse en dos grandes grupos: muestreo probabilístico y no probabilístico.

Muestreo probabilístico

Son aquellos sustentados en el principio de equiprobabilidad, esto significa que todos los individuos de la población tienen la misma probabilidad de ser elegidos y formar parte de la muestra. Con este tipo de muestreo aseguramos que la muestra extraída sea representativa de la población y por ese motivo son los más recomendables. Los muestreos probabilísticos son los siguientes: aleatorio simple, sistemático, estratificado y por conglomerados:

Muestreo aleatorio simple

Es un procedimiento muy sencillo, pueden utilizarse tómbolas, tablas de números aleatorios, días de la semana, números pares o nones, programas de computación, etcétera. Por ejemplo, en el caso de que se quisiera elegir 60 pacientes, repartidos en dos grupos de 30 pacientes cada uno, podríamos utilizar una tómbola en donde se introducirían bolas negras y blancas, las cuales representarían a los dos grupos. Otra forma sería introducir los 60 números y que los números pares representaran a un grupo y los nones al otro. Los riesgos de este tipo de muestreo es que se pueden *cargar* los grupos y hacerse heterogéneos o de tamaño muestrales diferentes. Este fenómeno se puede presentar cuando utilizamos un muestreo con reemplazo. Por ejemplo, en una tómbola hay 2 bolas que representan a grupos de tratamiento (negra = grupo 1 y blanca = grupo 2); cada vez que llega un paciente debemos sacar una bola y, de acuerdo a su color, ingresarlo a un grupo. Una vez que se elige al paciente regresamos la bola a la urna, pero observamos que al término del estudio tenemos 45 pacientes del grupo 1 y 15 del grupo 2. Este fenómeno se observa frecuentemente en familias numerosas (10 o más hijos), en que se observa que la mayoría de los hijos son hombres o mujeres y en forma teórica 50% deberían ser mujeres u hombres. La forma de mejorar este fenómeno es meter en la urna los 60 números y que los nones representen a un grupo y los pares al otro; con este procedimiento aseguramos la homogeneidad de los grupos.

Muestreo aleatorio sistemático

Como su nombre lo dice, para este tipo de muestreo es necesario utilizar un sistema o regla. Por ejemplo, si quiere elegir 60 individuos de una población de 300, se podría realizar una división entre $300/60$ con el resultado de 5, lo que significa que los pacientes múltiplos de 5 (5, 10, 15, 20...) serían escogidos para ingresar al estudio. Otra forma sería ingresar en una tómbola

la 5 números, posteriormente sacar un número y a partir de este continuar recolectando pacientes; por ejemplo, si el número fuera el 2, entonces los pacientes que se invitaría a participar serían el 2, 7, 12, 17, etcétera. Lo importante es que haya orden o sistema.

El riesgo este tipo de muestreo es que aun siendo un muestreo sistemático el azar puede provocar que se cargue un grupo, por ejemplo, que al final del estudio existan 35 pacientes del grupo 1 y 25 del grupo 2. Una forma de disminuir la influencia del azar sería de la siguiente manera: supongamos que los pacientes que ingresarán al estudio son 2, 7, 12, 17, 22, 27, 32, hasta completar 60 pacientes; lo que debe hacerse es que los nones representen al grupo 1 y los pares al grupo 2 y con esto aseguramos la homogeneidad.

Muestreo aleatorio estratificado

Como su nombre lo menciona, en este tipo de muestreo se utilizan varios estratos o categorías o subgrupos. Por ejemplo, supongamos que queremos determinar la prevalencia de *diabetes mellitus* tipo 2 en México. Supongamos que el tamaño de muestra requerido fuera de 40 mil individuos y obviamente todos los estados deben aportar individuos que los representen. Sin embargo, los estados tienen poblaciones de diferentes tamaños y si queremos que sean representados fielmente, debemos primero saber el porcentaje de población de cada estado. Por ejemplo, el estado de Colima, México, tiene una población de 0.4% del país ($n = 40,000$), entonces este estado aportaría 160 individuos, mientras que el Estado de México tiene 20% de la población del país y su aportación sería de 8 mil individuos.

El objetivo del muestreo estratificado es que todas las categorías tengan representación adecuada en una muestra. Es importante aclarar que para la selección de los individuos de cada categoría o estrato se puede utilizar muestreo aleatorio o sistemático.

Muestreo aleatorio por conglomerados

En realidad la palabra conglomerar significa unir y, desde ese punto de vista, todas las poblaciones están unidas. Esto significa que los conglomerados son grupos de individuos; en los muestreos anteriores las unidades muestrales son cada uno de los individuos, mientras que en el muestreo por conglomerados la unidad muestral es un grupo de elementos de la población que forman una unidad. Estas unidades pueden ser hospitales, estados, municipios, obreros, etcétera.

Muestreos no probabilísticos

En ocasiones no es posible realizar un muestreo probabilístico (ya sea por el costo o porque no se tienen recursos humanos suficientes) y se recurre a métodos no probabilísticos. Sin embargo, es importante aclarar que los resultados de una investigación con este tipo de muestreo difícilmente pueden generalizarse. Los muestreos no probabilísticos más utilizados en investigación son los siguientes:

Muestreo por cuotas

Se le llama también *accidental*. Para realizar este tipo de muestreo, es necesario tener un conocimiento de los estratos de la población y elegir a los individuos más *representativos* o *adecuados*. En este sentido se asemeja al muestreo estratificado, sin embargo, carece de *aleatorización*. En este tipo de muestreo se le asigna a cada encuestador o investigador cuántos individuos o cuotas deben analizar en un periodo. Las cuotas son individuos que tienen semejanza entre ellos, por ejemplo grupos etarios, región geográfica, religión, etcétera. Este método se utiliza en encuestas de mercado, elecciones políticas o encuestas de opinión.

Muestreo intencional o de conveniencia

Como su nombre lo dice, el investigador realiza este muestreo para su provecho o conveniencia. Es muy utilizado por los partidos políticos y consisten en realizar encuestas preelectorales en zonas donde generalmente les ha favorecido el voto.

Bola de Nieve

Se llama así porque primero se localizan a algunos individuos, los cuales conducen a otros, y éstos a otros, y así hasta conseguir una muestra suficiente. Este tipo se emplea muy frecuentemente cuando se hacen estudios con poblaciones *marginales*, delincuentes, sectas, determinados tipos de enfermos, etcétera.

Muestreo discrecional o por autoridad

A criterio del investigador o por indicaciones de una autoridad, se decide el número de individuos que se estudiarán.

Tamaño de muestra

Idóneamente en todos los proyectos de investigación se debería determinar un tamaño de muestra (TM). Un TM adecuado asegura la posible generalización de los resultados de un estudio y denota calidad. Sin embargo, en muchas ocasiones por falta de tiempo o desconocimiento no se realiza el TM. En la actualidad un estudio sin un adecuado cálculo de TM difícilmente será publicado en revistas con rigor científico o con factor de impacto y a través del tiempo hemos observado que grandes ideas o tesis han sido desechadas por las revistas debido a que carecen de éste. Es por ello que en los cursos de epidemiología clínica se insiste en que los trabajos de investigación debe determinarse un TM. Este capítulo no pretende abordar en forma amplia el tema sobre el TM, para esto hay paquetes estadísticos y libros escritos por expertos.

Actualmente existen muchas fórmulas para los diferentes tipos de estudio, sin embargo, aquí abordaremos aquellas que consideramos son las básicas para casi todos los diseños metodológicos. Las fórmulas están divididas de acuerdo con dos criterios: *tipo de variable* y *número de grupos de estudio*. Las variables pueden ser de dos tipos:

- *Proporción*. Cuando la variable de interés es de naturaleza cualitativa y tiene nivel de medición nominal u ordinal.

- *Razón.* Cuando la variable es de naturaleza cuantitativa y tiene nivel de medición de intervalo/razón.

Por el número de grupos pueden ser de dos tipos:

- Cuando hay *un grupo* se llama *descriptivo no comparativo*.
- Cuando hay *dos o más grupos* se llama *descriptivo comparativo*.

De acuerdo a lo anterior dividimos las fórmulas para el TM en cuatro tipos:

Fórmula para estudios descriptivos no comparativos de proporción

Son utilizados en diseños metodológicos observacionales o descriptivos donde se estudiará un solo grupo y la variable de interés es nominal u ordinal. Generalmente los utilizamos para determinar la frecuencia o porcentaje de individuos con un fenómeno relacionado al proceso salud/enfermedad. Por ejemplo, encuestas poblacionales para determinar prevalencias o incidencias, estudios transversales, estudios longitudinales y estudio transversal-analítico. Para determinar el tamaño de muestra podemos utilizar dos fórmulas:

TM cuando no se conoce el tamaño de la población o infinita:

$$n = \frac{Z\alpha^2 (PQ)}{E^2}$$

Donde:

$Z\alpha = (1.96)^2 = 3.84 = 95\%$ de confianza

P = proporción esperada

Q = (1-p). Es el complemento de p, de tal forma que la suma de p + q debe ser igual a 100%.

E = Error o precisión. Generalmente utilizamos 10% del valor de p.

Por ejemplo, queremos saber la prevalencia de sobrepeso en adultos del estado de Jalisco. Estudios previos en México han reportado prevalencias de 65%. De acuerdo con lo anterior tenemos que:

$$P = 65\%$$

$$Q = 35\% \text{ (recordar que } P + Q = 100\%)$$

$$(PQ) = \text{Resultado de la multiplicación de } P \times Q (35 \times 65) = 2275$$

$$Z\alpha = (1.96)^2 = 3.84 = 95\% \text{ de confianza}$$

$$E = 6.5\% \text{ (recordar que utilizamos 10\% del valor de } P)$$

Al despejar tenemos que:

$$n = \frac{(3.84) \times (2275)}{(6.5)^2} = \frac{8,736}{42.25} = 206$$

De acuerdo con esta fórmula se necesitan 206 individuos, a cuya cantidad debe agregarse entre 10 y 15% por posibles pérdidas, lo que da un total de 235.

TM cuando se conoce el tamaño de la población o finita:

$$n = \frac{n z^2 (pq)}{e^2 (n-1) + z^2 (q)}$$

Por ejemplo, se quiere determinar la prevalencia de sobrepeso/obesidad en niños de edad escolar que pertenecen a una población de México:

$$P = \text{Prevalencia de } 22\%$$

$$Q = 1-p = 78\% (100\%-22\%)$$

$$E = \text{Precisión absoluta de } 2.2\% (10\% \text{ de } P).$$

$$Z = 1.96$$

$$N = 52,806$$

La muestra calculada con este estadístico es de 960 niños, a quienes debe sumarse 10 a 15% extras por posibles pérdidas, de tal forma que el TM final será de ~1050 sujetos.¹

Fórmula para estudios descriptivos comparativos de proporción

Esta fórmula se utiliza cuando el objetivo principal es comparar porcentajes entre los grupos. Hay muchos diseños metodológicos que utilizan comparaciones de porcentajes como por ejemplo los ensayos clínicos, estudios de casos y controles, estudios de cohorte y comparación entre dos grupos. De tal forma, de acuerdo al tipo de diseño metodológico se realizan modificaciones a la siguiente fórmula básica:

$$n = \frac{[Z\alpha \cdot \sqrt{2p(1-p)} + Z\beta \cdot \sqrt{p_1(1-p_1) + p_2(1-p_2)}]^2}{(p_1 - p_2)^2}$$

Donde:

$Z\alpha = 1.96$ (también se le conoce como *nivel de significancia*, *error alfa* y equivale a un intervalo de confianza de 95%).

$Z\beta = 1.65$ (también se llama *poder* y equivale a un intervalo de confianza de 90%).

P_1 = Porcentaje esperado en el grupo 1

P_2 = Porcentaje esperado en el grupo 2

$Q_1 (1-p_1)$ = Porcentaje o complemento del grupo 1

$Q_2 (1-p_2)$ = Porcentaje o complemento del grupo 2

¹ Muchos estudiantes e investigadores utilizan la primera fórmula porque con esta el TM es menor que con la segunda y se sustentan en el principio de que no hay información previa. El problema surge debido a que un TM inadecuado, no va permitir generalizaciones. Por lo anterior sugerimos que, preferentemente, utilicen la fórmula para poblaciones finitas. Otro fenómeno frecuente es utilizar una precisión o error de 5% independientemente del valor de P, lo que provoca que el tamaño de muestra calculado sea menor. Por tal motivo siempre sugerimos que se utilice como mínimo 10% del valor de prevalencia o frecuencia (P).

Por ejemplo, se ha encontrado que 85% de los pacientes con obesidad mejorarán con dieta hipocalórica y 55% mejorarán con ejercicio. De acuerdo con estos datos, ¿cuál es el TM requerido?

$$\begin{aligned} P1 &= 85\%, \text{ luego entonces } q1 = 15\% \\ &(\text{recordar que } q1 \text{ resulta de restar } p1 - 100\%), \\ P2 &= 55\%, \text{ luego entonces } q2 = 45\%. \\ Z\alpha &= 1.96 \\ Z\beta &= 1.65 \end{aligned}$$

De acuerdo con estos datos, el TM requerido será de 44.3 individuos por grupo y a cada grupo se le debe agregar entre 10% a 15% por pérdidas, de tal forma que el TM final será de ~50 por grupo.

Fórmula para estudios descriptivos de razón

Se utiliza en diseños descriptivos cuya variable de interés es de naturaleza cuantitativa, su nivel medición es de razón y hay un grupo. La fórmula es la siguiente:

$$n = \frac{(Z\alpha \ S^2)}{E^2}$$

$Z\alpha = 1.96$ (también se le conoce como *nivel de significancia, error alfa*, y equivale a un intervalo de confianza de 95%)

E = Error o precisión

S^2 = Varianza = Desviación estándar elevada al cuadrado.

Ejemplo: Estudios previos han reportado que el promedio de glucosa y su desviación estándar en la población son 105 ± 6.0 mg/dL. ¿Cuántos individuos necesitaríamos estudiar en un intervalo de confianza de 95% y un error o precisión de 3 mg/dL?

$$\begin{aligned} S^2 &= (6.0)^2 = 36 \\ Z\alpha &= (1.96)^2 = 3.84 \\ E &= (3)^2 = 9 \end{aligned}$$

De acuerdo con la fórmula tendríamos lo siguiente:

$$n = \frac{(3.84 \times 36)}{9} = \frac{138.2}{9} = 15.3$$

De acuerdo con estos datos necesitamos 15 individuos, sin embargo, debemos agregar 10% a 15% por pérdidas y el TM final debería ser ~18 sujetos.

Fórmula para estudios descriptivos comparativos de razón

Se pueden utilizar en estudios descriptivos y ensayos clínicos. Su fórmula es la siguiente:

$$n = \frac{2(z\alpha + z\beta)^2 \times s^2}{e^2}$$

n = Son los individuos necesarios en cada una de las muestras.

Z α = 1.96

Z β = Poder o potencia estadística o error tipo beta; normalmente se utilizan errores de 0.84 (80% de poder), 1.03 (85% de poder), 1.28 (90% de poder) o 1.64 (95% de poder).

E = Es la precisión o diferencia esperada entre dos promedios.

S² = Varianza

Ejemplo: Se desea utilizar un nuevo fármaco antihipertensivo y se considera que sería clínicamente eficaz si lograrse un descenso de la presión diastólica de 10 mm Hg respecto al tratamiento habitual. Por estudios previos sabemos que la desviación estándar de la presión diastólica de pacientes que reciben tratamiento habitual para la hipertensión arterial es de 13 mm Hg; se acepta un riesgo o intervalo de confianza de 95% y un poder o potencia estadística de 90%. Al despejar, y de acuerdo con la siguiente fórmula anterior:

$$n = \frac{2(1.96 + 1.28)^2 \times 13^2}{10^2} = \frac{2(10.49) \times 169}{100} = \frac{20.9 \times 169}{100} = \frac{3532}{100} = 35.4$$

De acuerdo con este resultado, se necesitan 35 pacientes por grupo, pero se debe sumar 10 a 15% por pérdidas, de tal forma que el TM final sería de ~40 pacientes por grupo.

Glosario

Universo. Es la totalidad de individuos.

Muestra. Subconjunto de individuos que representa al universo.

Muestreo. Técnica para recolectar individuos.

Muestreo aleatorio. Técnica de muestreo que implica el azar.

Muestreo no aleatorio. Técnica de muestreo con ausencia del azar.

Errores o sesgos de muestreo. Conjunto de acciones equivocadas utilizadas para recolectar individuos.

Autoevaluación

¿Qué es el muestreo y cuáles son los tipos de muestreo que existen?

¿Cuáles son las estrategias de muestreo probabilísticas? Señale ejemplos de cada una de éstas.

¿Por qué el muestreo probabilístico permite la generalización de los resultados a la población?

¿Cuáles son las estrategias de muestreo no probabilísticas? Señale ejemplos de cada una de éstas.

¿Qué es *cálculo del tamaño de la muestra*?

¿Cuál es la diferencia entre las variables de proporción y razón?

¿Cuáles son las cuatro fórmulas para el cálculo del tamaño de la muestra?

¿Qué fórmulas son más adecuadas para los diseños de estudio observacionales o descriptivos?

Bibliografía

- Abramson, J.H., y Abramson, Z.H. (1999). *Survey methods in Community Medicine*. Edinburgh, London, New York, Philadelphia, Sydney, Toronto: Churchill Livingstone.
- Canales, F.H., De Alvarado, E.L., y Pineda, E.B. (1994). Metodología de la investigación. Manual para el desarrollo de personal de salud. México, D.F.: Limusa.
- Dawson, B., y Trapp, R. (2001). *Basic & Clinical Biostatistics*. New York: Lange Medical Books-McGraw-Hill.
- Downie, N.M., y Heath, R.W. (1986). *Métodos estadísticos aplicados*. México: Harla.
- Fleiss, J.L. (1981). *Statistical methods for rates and proportions*. New York: John Wiley & Sons.
- Fuentelsaz Gallego, C. (2004). Cálculo del tamaño de la muestra. En: *Matronas Profesión*, 5, pp. 5-13.
- Hulley, S.P., y Cummings, S.R. (1993). *Diseño de la investigación clínica. Un enfoque epidemiológico*. Barcelona: Doyma.
- Lwanga, S.K., y Lemeshow, S. (1991). *Determinación del tamaño de las muestras en los estudios sanitarios: manual práctico*. Ginebra: Organización Mundial de la Salud.
- Mateu, E., y Casal, J. (2003). Tamaño de la muestra. En *Revista de Epidemiología y Medicina Preventiva*, 1, pp. 8-14.
- Silva, L.C. (1993). *Muestreo para la investigación en Ciencias de la Salud*. Madrid: Díaz de Santos.
- Sokal, R.R., y Rohlf, F.J. (1995). *Biometry*. New York: W.H. Freeman and Company.

CAPÍTULO VIII

Cómo elegir la prueba estadística

Benjamín Trujillo Hernández

Introducción

La elección de una prueba estadística es un tema que frecuentemente provoca dolores de cabeza a muchos estudiantes de pregrado o posgrado, debido a que actualmente existen un sinnúmero de pruebas estadísticas y al estudiante le cuesta trabajo elegir la prueba adecuada. Al revisar los artículos científicos publicados en revistas indizadas hemos notado que la mayoría de los investigadores coinciden en utilizar las mismas pruebas estadísticas. Son precisamente estas pruebas *comunes* o *básicas* las que mencionaremos en este capítulo y en caso de requerir otras pruebas le sugerimos al lector revise un libro de estadística que aborde más profundamente este tema. Consideramos que cuando menos 90% de las investigaciones o trabajos utilizarán, como mínimo, una de las pruebas que mencionaremos.

Elección de la prueba estadística adecuada

Para elegir la prueba estadística adecuada es necesario considerar tres elementos:

- Nivel de medición de la variable dependiente.
- Número de grupos de estudio que se compararán.
- Número de mediciones al mismo grupo de individuos.

Nivel de medición de la variable dependiente

De acuerdo con el nivel de medición las variables se dividen en: nominal, ordinal, intervalo y de razón.

Variables nominales

Son variables nominales cuando se refieren a una cualidad o característica que puede determinarse con los órganos de los sentidos, con preguntas y en ocasiones con exámenes de laboratorio y se utilizan con el fin de clasificar a un individuo en una categoría, por ejemplo, color de ojos (negros, cafés, verdes, azules), estado civil (soltero, casado, viudo, divorciado), estado de salud (enfermo o sano), tipo sanguíneo (O, A, B, AB) o síntomas de alguna enfermedad (con o sin síntomas). Es el nivel más básico de la medición y los estadísticos que permiten describir su distribución son los porcentajes o las frecuencias.

Variables ordinales

Como su mismo nombre lo dice es cuando utilizamos en la medición un orden de mayor a menor. Se subdividen de la siguiente manera:

- *Ordinales de rango.* Se refieren a ordenar las variables en grados, por ejemplo estado de gravedad (delicado, grave y muy grave), síntomas (disnea leve, moderada, severa), grado de edema (+, ++, +++). En realidad se trata de una variable nominal que tiene tres o más categorías, por lo que en este tipo de variables la estadística es la misma que para las variables nominales.
- *Ordinales de puntaje.* Son aquellas que utilizan un método de medición numérica, pero que están sustentados en observaciones o mediciones subjetivas, por ejemplo, coeficiente intelectual, escala de *Glasgow*, escala de Apgar, etcétera. Las medidas de tendencia central y dispersión que permiten describirlas son la mediana, los intervalos, el rango o los centiles.

Variables de razón

Son las variables que pueden ser cuantificadas en forma subjetiva o con aparatos de medición que calibran un cero absoluto. Por ejemplo, número de hijos, peso, talla, edad, glucosa, leucocitos, etcétera. Las medidas de tendencia central y dispersión que las describen son éstas: la media, la desviación estándar, la varianza, los intervalos o el rango.¹

Número de grupos que se compararán

De acuerdo con lo anterior sugerimos clasificar a los diseños de la siguiente manera: *univariado* (un grupo), *bivariado* (dos grupos) y *multivariado* (tres o más grupos). Debemos de aclarar que esta clasificación es una propuesta y creemos que ayudará en la mayoría de las veces a elegir la prueba estadística idónea. Para un mayor conocimiento sugerimos al lector consulte un tratado de estadística.

Número de mediciones

Se refiere al número de mediciones que se realiza a la muestra; éstas pueden ser a través de encuestas, evaluación clínica, de laboratorio o gabinete. De acuerdo con lo anterior los estudios se dividen en: transversales (una sola medición) y longitudinales (dos o más mediciones). Cuando a un conjunto de individuos se les mide dos o más veces la misma variable y se comparan entre ellos mismos se llama *pareo*. Ejemplos de estudios transversales: prevalencia, frecuencias o encuestas en una población. Ejemplos de estudios longitudinales: estudios de incidencia o estudio de una cohorte.

Una vez que se haya identificado el tipo de variable, número de grupos de estudio y número de mediciones, se podría escoger la(s) prueba(s) estadística(s) adecuada(s); lo único que debe hacerse es cruzar la(s) columna(s) con la(s) fila(s) correspondiente(s) y se selecciona la prueba estadística (véase tabla 6).

¹ Nota. La estadística *no paramétrica* se encarga de analizar las variables nominales y ordinales (de rango y de puntaje) o cualitativas, mientras que la estadística *paramétrica* analiza las variables de razón o cuantitativas.

Ejemplos: Se realizó un ensayo clínico para evaluar la eficacia de la nitazoxanida *versus* con mebendazol + quinfamida para el tratamiento de parasitosis intestinal. Se realizó un estudio coprológico a 677 niños con edades de 2 a 12 años; las variables de interés fueron presencia o ausencia de cualquier parásito (helminthos o protozoarios) y número de huevecillos. Después del coprológico se encontró que 275 niños presentaban parasitosis. Previo consentimiento firmado e informado, estos niños fueron asignados al azar y doble ciego a uno de los tratamientos siguientes: Grupo A (nitazoxanida 200 mg por 3 días) y Grupo B (quinfamida 100 mg al día + mebendazol 200 mg por 3 días). Un coprológico postratamiento fue realizado a los 14 días.

Como comentario se menciona lo siguiente. Primero se observa que son *dos tipos* de estudios, el primero es de tipo descriptivo y el segundo es un ensayo clínico. El primer estudio tiene *un grupo* de 677 niños a los cuales se les determina un coprológico para detectar la presencia o ausencia de parasitosis (nominal) y número de huevecillos (razón) De acuerdo con lo anterior, hay un grupo, una medición y dos variables de interés (nominal y razón). De acuerdo con la tabla 6 para este tipo de estudio podemos realizar estadística descriptiva como frecuencia o porcentaje de parasitosis (para variable nominal) y promedio de huevecillos (variable razón).

El segundo estudio es un ensayo clínico con *dos grupos* de tratamiento, *dos mediciones* o coprológicos en los que se determinó dos variables: nominal (presencia o ausencia de parasitosis) y de razón (número de huevecillos). La estadística que debe seguirse la tienen que elegir los autores, pero supongamos que éstos desean comparar la presencia o ausencia de parasitosis y el número de huevecillos antes y después del tratamiento en cada uno de los grupos. Esto es, cada grupo de individuos se va comparar con ellos mismos en dos ocasiones y, como habíamos mencionado, cuando el mismo grupo de individuo se mide dos o más veces y se compara se llama pareo.

Tabla 6
Algoritmo para seleccionar la prueba estadística

Grupos	Variable	Una medición	Dos mediciones (pareadas)	Tres o más mediciones (pareadas)
Uno	Nominal	Porcentajes	McNemar	Q Cochran
Uno	Ordinal de puntaje	Mediana	Wilcoxon	Friedman
Uno	Razón	Promedios	t de Student	ANOVA de medidas repetidas
Dos	Nominal	Chi cuadrado	McNemar	Q Cochran
Dos	Ordinal de puntaje	U de Mann-Whitney	Wilcoxon	Friedman
Dos	Razón	t de Student	t de Student pareada	ANOVA de medidas repetidas
Tres	Nominal	Chi cuadrado	McNemar	Q Cochran
Tres	Ordinal de puntaje	Kruskall-Wallis	Wilcoxon	Friedman
Tres	Razón	ANOVA de 1 vía	t de Student pareada	ANOVA de medidas repetidas

De acuerdo a lo anterior, vemos en la tabla 6 que cuando una variable nominal (ausencia o presencia de parasitosis) se mide y se compara dos veces en el mismo grupo de individuos (grupo A o grupo B) se debe utilizar la prueba de McNemar, mientras que la variable de razón (número de huevecillos) si se mide y compara dos veces en el mismo grupo de individuos, se debe analizar con la prueba *t de Student* pareada. De tal forma, debemos realizar dos pruebas de McNemar (grupo A y grupo B) y dos *t de Student* pareadas (grupo A y grupo B).

Ahora bien, si los autores desearan también realizar comparaciones entre los dos grupos (grupo A *versus* grupo B) antes y después del tratamiento entonces se deben realizar la prueba de *chi cuadrado* (variable nominal) y la prueba *t de Student* (variable de razón). O sea, se deben realizar dos chi cuadrados (antes y después del tratamiento) y dos *t de Student* (antes y después del tratamiento).

En conclusión, para este estudio se realizaría lo siguiente: estadística descriptiva como porcentajes, promedios y desviación estándar. La comparación intragrupal o dentro del grupo se determinaría con la pruebas pareadas de McNemar (presencia o ausencia de parasitosis) y *t de Student* pareada (número de huevecillos). La comparación intergrupala o entre los grupos se determinaría con las pruebas independientes de *chi cuadrado* (presencia o ausencia de parasitosis) y *t de Student* (número de huevecillos). Hay que recordar que la estadística es una herramienta y que cada investigador debe escoger el diseño estadístico de acuerdo a su pregunta de investigación. Sugerimos que no debe hacerse análisis estadístico por imposición de autoridad.

Otras pruebas estadísticas

Asociación causal

En estudios analíticos se debe determinar el grado de asociación causal a través del riesgo relativo (RR) o razón de productos cruzados (OR). Aquí se analiza un factor de riesgo y la enfermedad. Pero, si el objetivo es analizar la asociación de múltiples riesgos se utilizan las pruebas de Mantel-Hanzsel y regresión logística, entre otras.

Correlación

En ocasiones el objetivo es evaluar el efecto que tiene una variable (independiente) sobre otra (variable dependiente). Si las dos variables son de razón se utiliza la correlación de *r de Pearson*, mientras que si son de ordinal de puntaje se utiliza la correlación *rho de Spearman* y la correlación múltiple cuando son muchas variables.

Pruebas diagnósticas

En ocasiones el objetivo es evaluar una nueva prueba diagnóstica la cual compararemos con la prueba de oro o *gold estándar*. En este tipo de estudios se debe determinar la sensibilidad, especificidad, y los valores predictivos positivo y negativo.

Número necesario a tratar (100/RR)

Esta herramienta es cada vez más utilizada en ensayos clínicos y consiste en dividir 1 entre la diferencia de mejoría entre dos grupos de tratamiento o reducción de riesgo absoluto (RRA). Supongamos que se realizó un ensayo clínico en donde el objetivo fue determinar y comparar el porcentaje de pacientes hipertensos que tuvieron cifras de presión arterial normal después de haber sido asignados a dos tipos de tratamiento por dos meses. Al final del estudio el grupo A presentó mejoría en 80%, mientras que el grupo B fue de 40%, de acuerdo con esto, el grupo A fue mejor y hubo una reducción de riesgo absoluto o mejoría de 40% con respecto al grupo B. Luego entonces $100/40 = 2.5$. Significa que por cada 2.5 pacientes que se trate con el antihipertensivo A, habrá un paciente con cifras normales o que es necesario tratar 2.5 pacientes para que 1 presente cifras normales.

Glosario

Elección de una prueba estadística. Para elegir la prueba estadística adecuada es necesario considerar tres elementos: el nivel de medición de la variable dependiente, el número de grupos de estudio y el número de mediciones al mismo grupo de individuos.

Autoevaluación

¿Qué es una prueba estadística?

¿Cuáles son los elementos que se deben considerar para elegir una prueba estadística?

¿Cómo funciona el algoritmo para la selección de pruebas estadísticas que propone el autor?

¿Qué otras pruebas estadísticas describe el autor?

Bibliografía

- Dawson, B., y Trapp, R. (2001) *Basic & Clinical Biostatistics*. New York: Lange Medical Books-McGraw-Hill.
- Downie, N.M., y Heath, R.W. (1986). *Métodos estadísticos aplicados*. México: Harla.
- Fleiss, J.L. (1981). *Statistical methods for rates and proportions*. Nueva York: John Wiley & Sons.
- Sackett, D.L.; Straus, S.E.; Richardson, W.S., et al. (2001). *Medicina basada en la evidencia*. Madrid: Harcourt.
- Sokal, R.R., y Rohlf, F.J. (1995). *Biometry*. New York: W.H. Freeman and Company.

CAPÍTULO IX

Análisis de datos con el paquete estadístico SPSS

José Ramiro Caballero Hoyos

Expreso mi más grande reconocimiento al Lic. Freddy Carranza, estadígrafo boliviano, mi amigo y maestro, quien hace tres décadas me compartió, con generosidad, los procedimientos básicos para analizar datos con el SPSS. En agradecimiento, le dedico el presente capítulo como humilde fruto de la semilla de aprendizaje solidario.

José Ramiro Caballero Hoyos

Introducción

Las metodologías de investigación cuantitativa realizan observaciones empíricas que pretenden medir hechos sociales y, en ese afán, traducen lo observado en datos con códigos numéricos. El análisis de esos datos consiste en la aplicación sistemática de técnicas estadísticas orientadas a resumir su distribución, organizar sus relaciones e inferir los patrones y tendencias de su estructura.

A mediados del siglo xx, los análisis de datos se realizaban aún con herramientas de cálculo simples; pero, desde la década de los sesenta, el desarrollo de la tecnología y de la informática permitió una creciente incorporación de herramientas de cómputo a la actividad. Esta transición del análisis artesanal al computarizado fue revolucionario porque sacó a la actividad estadística de su oscura *biblioteca medieval*, para convertirla en una tarea de construcción colectiva.

Aparece el SPSS

En ese contexto de transición histórica del análisis de datos, apareció en el año 1968 el *Paquete Estadístico para las Ciencias Sociales* (más conocido bajo el acrónimo inglés spss del inglés *Statistical Package for the Social Sciences*). Uno de sus creadores fue Norman H. Nie, joven de 22 años de edad quien en 1967 era candidato doctoral de Ciencias Políticas en la Universidad de Standford.

Norman H. Nie impulsó el diseño y la producción del nuevo paquete estadístico porque estaba frustrado en su intento de usar programas complejos, procedentes de la Física y la Biología, para analizar los datos de su tesis. Tras pensar en los requerimientos técnicos y didácticos necesarios para el nuevo paquete, se asoció a dos compañeros expertos en informática (C. Hadlai Hull y Dale H. Bent), quienes le ayudaron a materializar el proyecto.

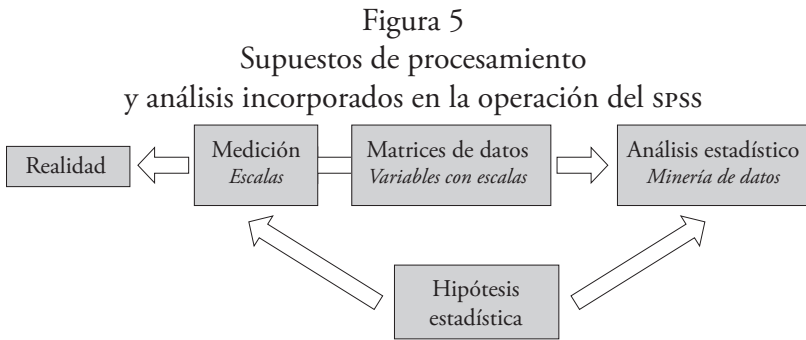
Originalmente, el paquete fue creado para ordenadores grandes y contaba con un manual de usuarios elaborado por Norman H. Nie y C. Hadlai Hall, en 1970. Por su mayor simplicidad de manejo y su potencial estadístico, el paquete se popularizó en universidades, centros de investigación e instituciones de gobierno de todo el mundo.

Durante el desarrollo del paquete (1969 y 1975), en su primera etapa, fue la Universidad de Chicago la que impulsó su producción y mercadeo. En la segunda etapa (1975 a 2009) el paquete pasó a manos privadas (*SPSS Inc.*) y su popularidad se hizo ubicua gracias al lanzamiento de la versión para microordenadores (1984) y la incorporación del ambiente de *Windows*. En su etapa más reciente (2009 en adelante), el paquete fue comprado por la empresa IBM y lleva el nombre de *IBM SPSS Statistics*. Esta empresa genera cada año una nueva versión que incorpora programas de análisis cada vez más complejos, con simplicidad didáctica y con el soporte de manuales multilingües.

Supuestos teóricos del análisis computarizado

En términos generales, el paquete SPSS es un sistema global de cómputo que permite realizar el análisis estadístico de datos numéricos clasificados en escalas de medición. Está diseñado bajo el ambiente del sistema operativo *Windows* que facilita el procesamiento (la captura de bases de datos y la transformación de variables) y el manejo de distintos niveles de análisis gráfico y numérico.

En su aplicación operativa, el SPSS incorpora los supuestos interrelacionados de procesamiento y análisis estadístico que se sintetizan en la figura 5.



Fuente: Elaboración propia.

Tales supuestos pueden resumirse de la siguiente manera:

1. La noción de la realidad como un fenómeno que se puede percibir sólo de manera indirecta, a través de observaciones basadas en sistemas de medición.
2. El sistema de medición que permite aterrizar los conceptos abstractos en registros observables codificados numéricamente. Los números que se asignan a los registros tienen propiedades numéricas y habilitan distintos niveles de análisis estadístico.
3. La construcción de matrices de datos y definiciones, donde, por una parte, se capturan las respuestas numéricas registradas en las observaciones bajo un formato

de casos (filas) y variables (columnas), y, por otra, se definen las características técnicas de las variables incluidas en la matriz (donde la definición principal es *escala de medición*).

4. Las matrices de datos y definiciones de variables habilitan el análisis estadístico de la información, en función de las escalas de medición de las variables.

En ese sentido, las variables categóricas o discretas se expresan en dos tipos de escalas (nominal y ordinal) que habilitan el cálculo de frecuencias, porcentajes, moda, mediana y rango, mientras que las variables de escala o continuas se expresan en escalas de intervalo y razón que habilitan el cálculo de algunas medidas de tendencia central y dispersión (tabla 7).

Tabla 7

Escalas de medición de las variables con sus propiedades numéricas y medidas de resumen que habilitan calcular

Tipos de variables	Escalas	Propiedad numérica	Medidas univariadas de resumen
<i>Catégoricas o discretas</i>	Nominal	Identifican atributos. Ejemplo: 1. Mujer 2. Hombre	Frecuencias, porcentajes, moda. Gráficas: diagramas de barras y sectores.
	Ordinal	Ordenan magnitudes. Ejemplo: 1. Menor 2. Mayor	Frecuencias, porcentajes, mediana, rango. Gráficas: diagramas de barras y sectores.
<i>Escala o continuas</i>	Intervalo	Distancias iguales y cero arbitrario (convencional). Ejemplo: grados de temperatura.	Tendencia central: media, mediana, moda. Dispersión: rango, desviación típica, varianza.
	Razón	Distancias iguales y cero absoluto (ausencia atributo) Ejemplo: precios de medicinas	Gráficas: histograma y normalidad.

Fuente: Elaboración propia.

Cuando el análisis se enfoca a buscar relaciones de variables, es central la orientación de las hipótesis estadísticas que se aplican bajo la lógica del contraste de independencia. Este contraste tiene como hipótesis nula a la independencia (relación nula entre variables) y como alterna a la dependencia (algún tipo de relación entre variables).

En dicha búsqueda, las combinaciones de pares de variables orientan el cálculo de determinadas pruebas estadísticas, en función de las escalas de medición que poseen. La tabla 8 mues-

tra que cuando se combinan variables nominales y ordinales será posible calcular la prueba *Chi-cuadrado*. Cuando se combinan variables de intervalo/razón con nominales u ordinales se podrán calcular las pruebas *t de Student* (es decir, cuando las categorías de la variable nominal son dos) o ANOVA (es decir, cuando las categorías de la variable nominal son tres o más), y, cuando se combinan variables de intervalo/razón entre sí será factible calcular la correlación *r de Pearson* y la *regresión lineal*.

Tabla 8
Combinación de pares de variables por sus escalas de medición y las pruebas estadísticas que habilitan calcular

Escalas x escalas	Nominal	Ordinal	Intervalo/razón
<i>Nominal</i>	Chi-cuadrado	Chi-cuadrado	<i>t de Student</i> (prueba <i>t</i>)
			ANOVA (prueba <i>F</i>)
<i>Ordinal</i>	Chi-cuadrado	Chi-cuadrado	<i>t de Student</i> (prueba <i>t</i>)
			ANOVA (prueba <i>F</i>)
<i>Intervalo/razón</i>	<i>t de Student</i> (prueba <i>t</i>)	<i>t de Student</i> (prueba <i>t</i>)	Correlación <i>r</i> de Pearson
	ANOVA (prueba <i>F</i>)	ANOVA (prueba <i>F</i>)	Regresión lineal

Fuente: Elaboración propia.

Cuando se relacionan más de dos variables, el uso de las herramientas informáticas ha orientado el análisis hacia el descubrimiento de relaciones significativas, patrones y tendencias, mediante técnicas especializadas de la estrategia llamada *minería de datos*. Esta estrategia persigue lograr el descubrimiento automático del conocimiento contenido en la información almacenada en las bases de datos, mediante dos tipos generales de técnicas estadísticas (Pérez, 2009: 19-32):

- *Explicativas o de dependencia*. Asumen, bajo criterios teóricos, la presencia de una o más variables dependientes como explicadas por otras variables independientes.

Cuando se generan modelos de análisis, se eligen las técnicas estadísticas inferenciales que son adecuadas a las escalas de medición de todas las variables relacionadas, entre las que se incluyen todos los tipos de regresión lineal, el análisis de varianza, el análisis discriminante, el análisis de series temporales y otras.

- *Descriptivas o de interdependencia.* No se originan en la teoría y asignan importancias equivalentes a todas las variables que se relacionan. Cuando se generan modelos de análisis, las técnicas seleccionadas buscarán inferir las relaciones estadísticas que se desprendan de la interconexión de las variables. Entre estas técnicas se encuentran el análisis de conglomerados, el escalamiento multidimensional, el escalamiento óptimo, el análisis de correspondencias, el análisis factorial de componentes principales, los árboles de decisión, las redes neuronales y otras.

Idea del capítulo

Provistos con los elementos de referencia previamente descritos (la idea referente al contexto de aparición del paquete SPSS en la historia del análisis estadístico, y los supuestos de procesamiento y análisis que incorpora el SPSS en su aplicación a los datos), ahora estamos en condiciones de proponer una introducción básica sobre el manejo del *SPSS versión 19* (Norusis, 2011) como una herramienta para el análisis de datos.

En esa introducción seguiremos un consejo clásico de los analistas (Lewis-Beck, 1995). Iniciaremos con los procedimientos más simples (exploración y descripción de variables individuales), para luego orientar, con los resultados obtenidos, el desarrollo de construcciones más complejas (contrastes de relación entre dos variables y entre múltiples variables). Al seguir el consejo asumiremos un enfoque exploratorio que plantea la necesidad de realizar, primero, un detallado conocimiento de las varia-

bles individuales, como requisito, cuyo propósito es lograr, después, un uso eficiente de éstas, las cuales servirán para construir relaciones de variables que aspiren a la generalización teórica.

Al realizar dicha introducción, tendremos como telón de fondo el constante ir y venir de los supuestos relativos al procesamiento y análisis de datos descritos previamente (véase figura 5), lo cual ilustrará el hecho de que el analista debe conocer y manejar esos supuestos, así como ser consciente de su influencia en la construcción de todo análisis; por ejemplo, cómo la escala de medición de las variables condiciona el tipo de estadísticos a calcular para su exploración y descripción; o bien cómo la escala de medición de dos variables lleva a decidir qué prueba será la adecuada para contrastar la hipótesis de su independencia.

Introducción al análisis estadístico con el SPSS

Con la idea de ilustrar los procedimientos de análisis, trabajaremos con bases de datos ya existentes. De una base seleccionaremos solamente ciertas variables de escala nominal y ordinal, para concentrarnos en describir algunos de los procedimientos simples y complejos calculables para las mismas. Luego seleccionaremos, de otra base de datos, algunas variables de intervalo y razón para enfocarnos en el cálculo de algunos de sus procedimientos simples y complejos. Al final, obtendremos dos resultados de análisis multivariados que serán expresados en tablas e interpretados.

Análisis de variables nominales y ordinales

En esta primera parte, se ilustrará el análisis de variables de escala nominal y ordinal tomadas de la base de datos que proporciona la Encuesta Nacional de Salud y Nutrición 2012 (ENSANUT 2012), sección *Adultos*, realizada en México por el Instituto Nacional de Salud Pública. Dicha encuesta consideró una muestra de 46,303 adultos, hombres y mujeres de 20 años o más, y fue representativa de 61.9% de la población estimada en México.

Se rescatarán de la ENSANUT 2012 las siguientes variables: el estrato de marginalidad y el estrato de urbanidad (dos variables estructurales que tienen teóricamente un peso importante como determinantes de la salud poblacional), el género, los grupos de edad, el uso del condón en la primera relación sexual, y el uso del condón en la última relación sexual y sus principales motivos (prevención del embarazo, infecciones de transmisión sexual [ITS] y virus de inmunodeficiencia humana [VIH]). La tabla 9 describe dichas variables con sus escalas de medición, los valores numéricos con que se codificaron y las etiquetas otorgadas a las categorías de los números.

Tabla 9
Escalas de medición, valores y etiquetas de las variables
de ejemplo de la ENSANUT 2012 (México)

Variables	Escala de medición	Valores y etiquetas
Estrato de marginalidad	Ordinal	1. Baja 2. Alta
Estrato de urbanidad	Nominal	1. Rural 2. Urbano 3. Metropolitano
Uso del condón en primera relación sexual	Nominal	0. No, 1. Sí
Uso del condón en última relación sexual	Nominal	0. No, 1. Sí
<i>Motivos usos del condón</i>		
Prevención embarazo	Nominal	0. No, 1. Sí
Prevención ITS	Nominal	0. No, 1. Sí
Prevención VIH	Nominal	0. No, 1. Sí
Sexo	Nominal	1. Hombre 2. Mujer
Grupos de edad	Ordinal	1. 20-29 2. 30-45 3. 46>

Fuente: Elaboración propia con variables tomadas de la ENSANUT 2012.

Iniciar el trabajo con el paquete SPSS requiere, primero, abrir la base de datos que contiene el registro numérico de las respuestas a las preguntas y las definiciones técnicas de las variables. La que usaremos en este trabajo es una base creada en el SPSS (su archivo tiene extensión “sav” y se llama “Adultosenfocada.sav”) que contiene sólo las variables que seleccionamos del amplio repertorio proporcionado por la ENSANUT 2012. Para abrir esta base se ingresa al paquete y se sigue el siguiente procedimiento: en la barra del menú que aparece en el SPSS se sigue la secuencia *Archivo/Abrir/Datos*. Se despliega una caja de diálogo que contiene la lista de los archivos disponibles y se selecciona el que interesa analizar: *Adultosenfocada.sav*. Finalmente, elige la opción *Abrir* para ejecutar el procedimiento (véanse figuras 6 y 7):

Figura 6
Secuencia a seguir en el menú del SPSS para ingresar a *Datos*

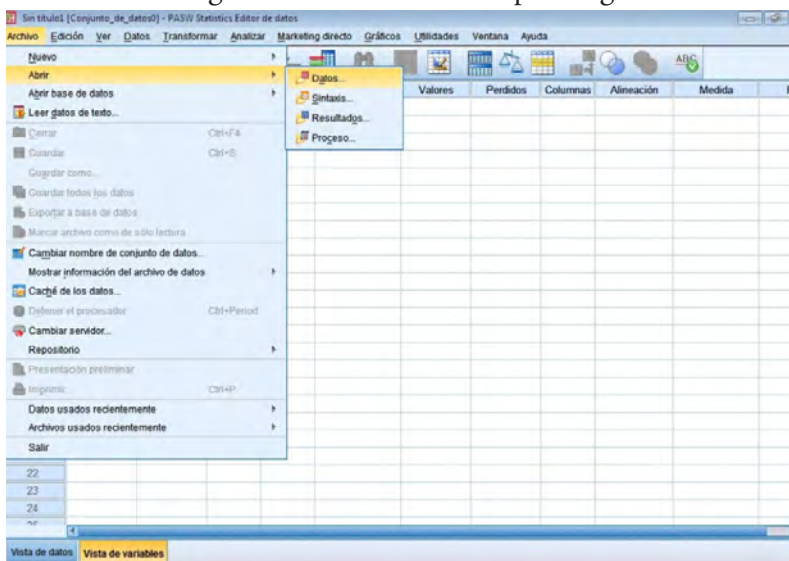
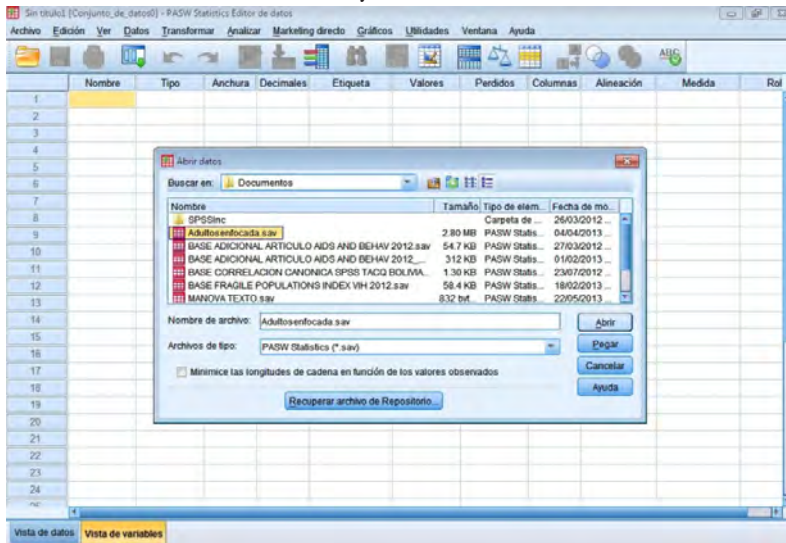


Figura 7
Caja de diálogo *Abrir/Datos*
donde se selecciona y abre el archivo de interés



Como resultado del procedimiento, se despliega en la pantalla una *matriz de datos* con los valores numéricos de las variables, a la cual se accede por la pestaña “Vista de datos” (figura 8) y una *matriz con las definiciones técnicas* de las variables, a la cual se ingresa por la pestaña “Vista de variables” (figura 9). Un detalle técnico esencial, en esta última matriz, es la definición de la escala de medición de las variables (véase la columna “Medida” donde se registran tres posibles tipos de escalas con sus correspondientes símbolos gráficos: nominales , ordinales  e intervalo/razón  (etiquetada como “escala”).

Figura 8
Matriz de datos con los valores numéricos de las variables¹

	folio	sexo	CONDONINSEX	CONDONULRESEX	CONDONVIH	CONDONITS	CONDONEMBA	P.P.R.R.	est_urb	est_marg	EDADREC
34	010053	1							3	1	3
35	010053	2							3	1	3
36	010053	2	1		0	1	0		3	1	1
37	010056	1							3	1	3
38	010057	1	1		0	1	0		3	1	1
39	010059	2	1		0	1	1		3	1	1
40	010061	1							3	1	3
41	010062	1	0		0	1	1		3	1	1
42	010063	2							3	1	3
43	010064	2							3	1	3
44	010065	1	0		0				3	1	2
45	010066	1							3	1	1
46	010067	1							3	1	3
47	010070	2							3	1	3
48	010071	1							3	1	3
49	010074	2							3	1	3
50	010076	2	0		0				3	1	3
51	010077	2							3	1	3
52	010077	1							3	1	2
53	010080	2	1		0	0	1		3	1	1
54	010081	2	0						3	1	2
55	010086	2	0		0				3	1	3
56	010090	2							3	1	3

Figura 9
Matriz de datos con las definiciones técnicas de las variables²

	Nombre	Tipo	Anchura	Decimales	Etiqueta	Valores	Pérdidos	Columnas	Alineación	Medida
1	folio	Cadena	10	0	Folio consecutivo	Ninguna	Ninguna	15	Izquierda	Nominal
2	sexo	Numérico	6	0	SEXO	(1, Hombre)	Ninguna	8	Derecha	Nominal
3	CONDONINSEX	Numérico	6	0	CONDONINSEX	(0, No)	Ninguna	8	Derecha	Nominal
4	CONDONULRESEX	Numérico	6	0	CONDONULRESEX	(0, No)	Ninguna	8	Derecha	Nominal
5	CONDONVIH	Numérico	6	0	CONDONVIH	(0, No)	Ninguna	8	Derecha	Nominal
6	CONDONITS	Numérico	6	0	CONDONITS	(0, No)	Ninguna	8	Derecha	Nominal
7	CONDONEMBAR	Numérico	6	0	CONDONEMBAR	(0, No)	Ninguna	8	Derecha	Nominal
8	PRUEBAVIH	Numérico	6	0	PRUEBA VIH ALGUNA VEZ	(1, Si)	Ninguna	8	Derecha	Nominal
9	PRUEBAVIH12	Numérico	6	0	PRUEBA VIH ULTIMO AÑO	(1, Si)	Ninguna	8	Derecha	Nominal
10	est_urb	Numérico	1	0	Estrato de Urbanidad	(1, Rural)	Ninguna	8	Derecha	Nominal
11	est_marg	Numérico	1	0	Estrato de Marginalidad	(1, Baja Ma)	Ninguna	8	Derecha	Ordinal
12	EDADREC	Numérico	8	0	Grupos de edad	(1, 20 A 29)	Ninguna	10	Derecha	Ordinal

¹ Se visualiza en el Editor de datos del SPSS, en la pestaña inferior de “Vista de datos”.

² Se visualiza en el Editor de datos del SPSS, en la pestaña inferior de “Vista de variables”.

Análisis univariado: tablas de frecuencia

Una vez que se tienen listas las matrices de datos, un primer paso de análisis consiste en explorar el contenido y distribución de las variables. Debido a que se trata de variables de escala nominal y ordinal, pueden calcularse las frecuencias y los porcentajes como medidas de resumen. Adicionalmente, pueden representarse sus respectivas gráficas (véase tabla 7).

En el caso de las variables de ejemplo, primero se podrían calcular las frecuencias y porcentajes de las variables estructurales (estrato de marginalidad, estrato de urbanidad) y de las demográficas (sexo y grupos de edad). Este análisis puede realizarse en el SPSS de la siguiente forma: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Estadísticos descriptivos/Frecuencias* (figura 10). Se despliega una caja de diálogo que contiene la lista de las variables disponibles. De esa lista se seleccionan las cuatro variables requeridas y se trasladan al espacio de *Variables*. En seguida se marca *Mostrar tablas de frecuencias* y se elige la opción *Aceptar* para ejecutar el procedimiento (figura 11).

Figura 10

Secuencia a seguir en el menú del SPSS para ingresar a *Frecuencias*

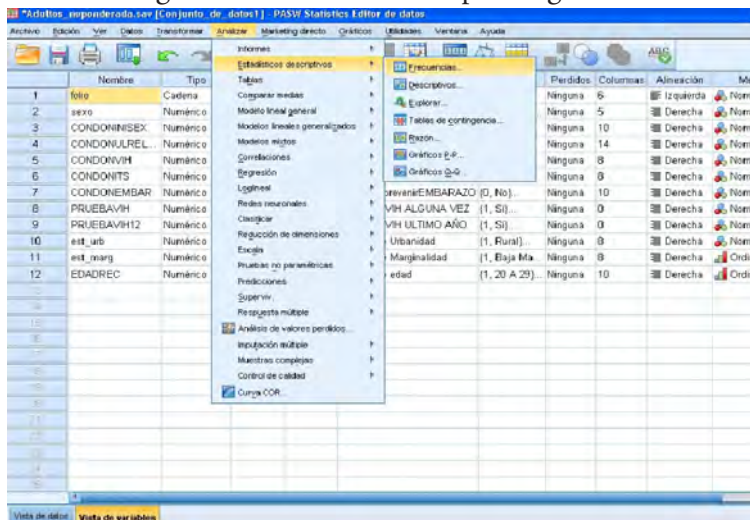
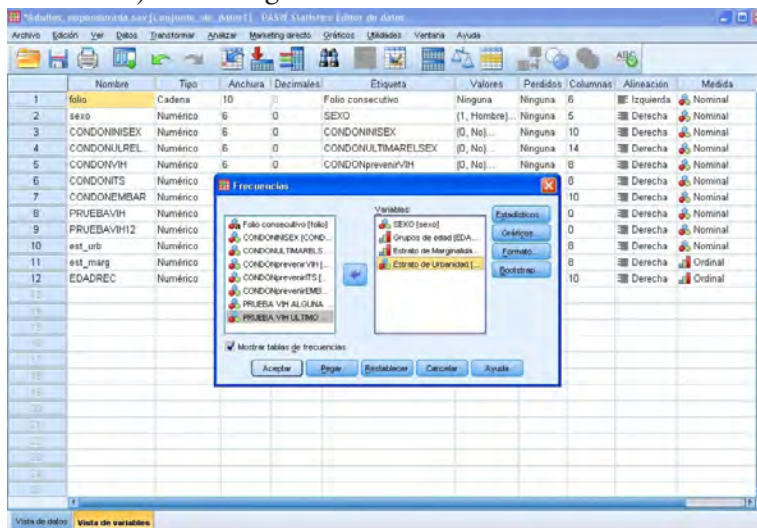


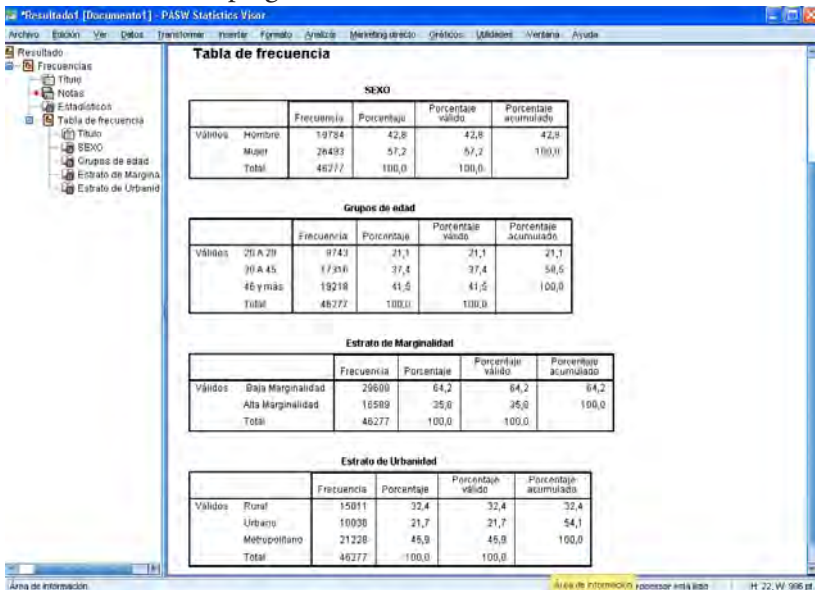
Figura 11
Caja de diálogo donde se abren las *Frecuencias*



Tal procedimiento genera tablas de frecuencias que se despliegan en el *Visor* de resultados (figura 12). En la parte izquierda se reporta un esquema de las variables analizadas, mientras que en la derecha aparecen las tablas que describen la frecuencia de casos en cada categoría de respuesta, los porcentajes y los totales.

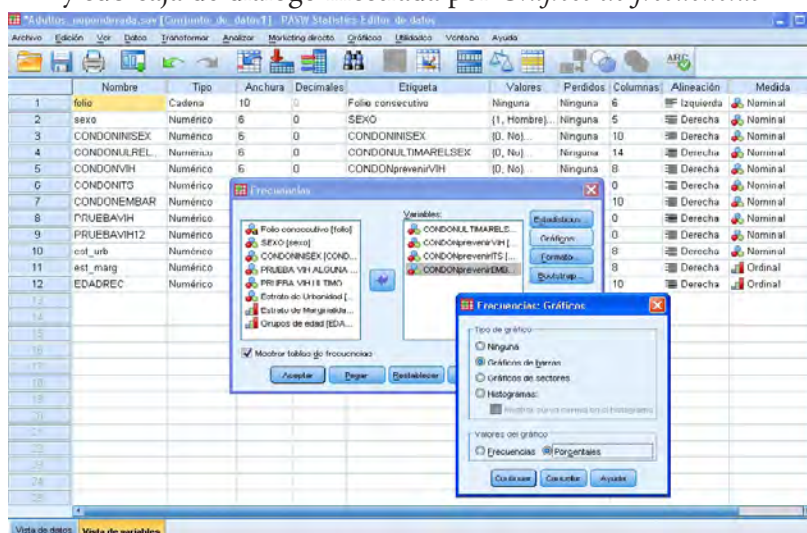
Después de leer las tablas se puede rescatar lo siguiente: a) en la variable sexo: la muestra estuvo constituida por 57.2 % de mujeres y 42.8% de hombres; b) en grupos de edad: 21.1% eran de 20 a 29 años, 37.4% de 30 a 45 años y 41.5% de 46 y más años (sirve también, en esta variable ordinal, el porcentaje acumulado que indica, por ejemplo, que hay 58.5% de casos de 20 a 45 años); c) en estratos de marginalidad: 64.2 % eran de baja marginalidad y 35.8% de alta marginalidad; y d) en estratos de urbanidad: 32.4% eran habitantes rurales, 21.7% urbanos (contextos de 2,500 a 100 mil habitantes) y 45.9% metropolitanos (ciudades con más de 100 mil habitantes).

Figura 12
Tablas de frecuencias de las variables demográficas y estructurales,
desplegadas en el *Visor* de resultados



En forma complementaria, también se pueden calcular las frecuencias y porcentajes para la variable “uso del condón en la última relación sexual” y las tres variables de “motivos de uso”, así como desplegar sus gráficos de barras. Lo anterior puede realizarse en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Estadísticos descriptivos/Frecuencias*. Se despliega una caja de diálogo que contiene la lista de las variables disponibles. De esa lista se seleccionan las variables requeridas y se trasladan al espacio de *Variables*. Luego se ingresa a *Gráficos* donde se marca *Gráficos de barras* en tipos de gráfico y *Porcentajes* en valores del gráfico. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 13).

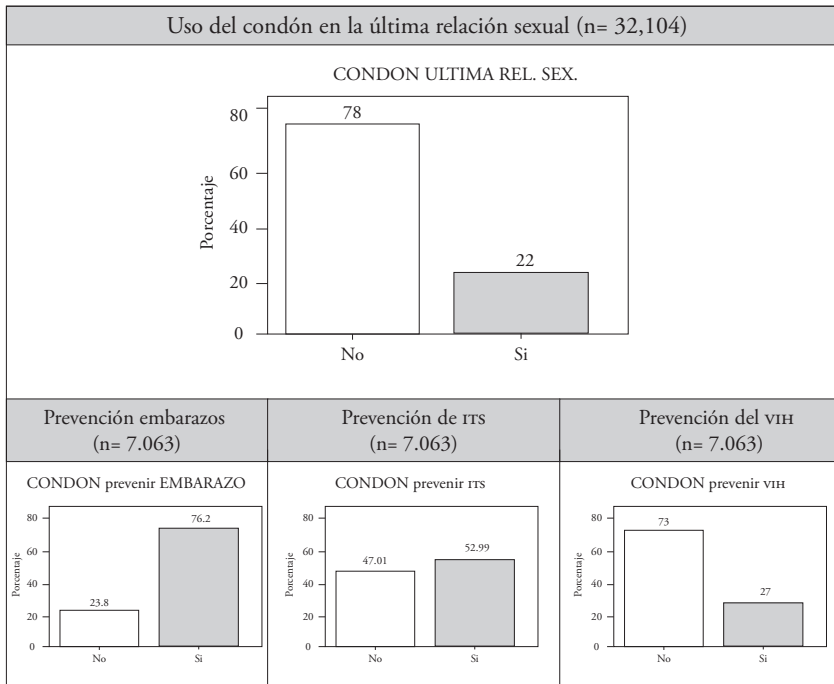
Figura 13
Caja de diálogo de *Frecuencias*
y sub caja de diálogo mostrada por *Gráficos de frecuencias*



Como efecto de tal procedimiento, las tablas de frecuencias y los gráficos se despliegan en el *Visor* de resultados. En la figura 14 se muestran editados, en un recuadro, solamente los gráficos de barras de las variables consideradas. El primer gráfico permite ver el porcentaje de casos con vida sexual activa ($n = 32,104$) que usaron el condón en su última relación sexual ($n = 7,063$ casos, 22%). En los gráficos restantes se desglosan los motivos de este último uso: 76.2% prevenir embarazos, 52.9% prevenir alguna ITS y 27% prevenir el VIH.

Figura 14

Gráficos de barras que muestran las variables *uso del condón en la última relación sexual* y las tres variables de *motivos de uso*, desplegadas en el visor de resultados



En la descripción de las variables nominales, también puede identificarse si existen diferencias significativas entre la proporción de casos observados en las categorías y sus proporciones esperadas (por ejemplo, en la variable “sexo” se podría haber esperado una proporción similar de hombres y mujeres en la muestra: 50% de cada uno). La prueba que nos permitirá hacer tal identificación es la *binomial* y se puede ingresar a ésta en el SPSS de la siguiente forma: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Pruebas no paramétricas/Cuadros de diálogo antiguos/Binomial* (figura 15). Se despliega una caja de diálogo en la cual se elige la variable *Sexo* y se traslada al espacio *Lista Contrastar variables*. Luego, en definir dicotomía se

marca *Obtener de los datos* y en “proporción de prueba” se escribe el valor de la proporción esperada (0.50). Finalmente se elige *Aceptar* para correr el procedimiento (figura 16).

Figura 15
Secuencia a seguir en el menú del spss
para ingresar a la *Prueba binomial*

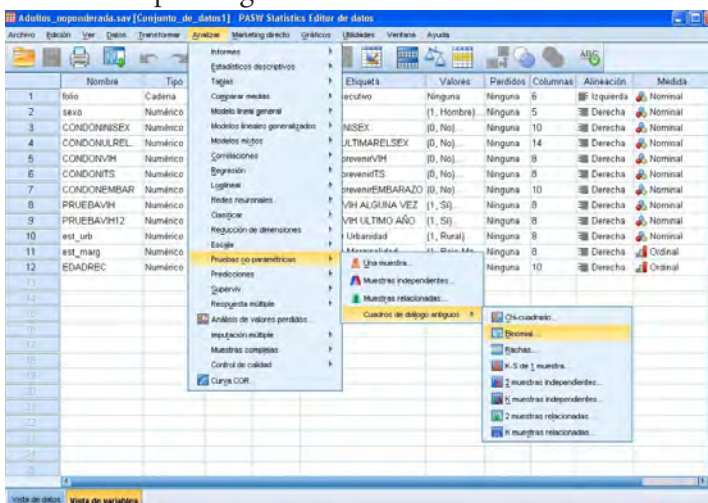
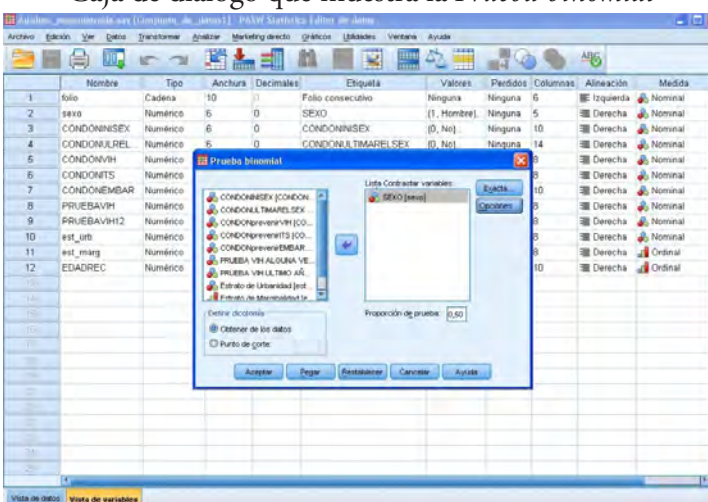


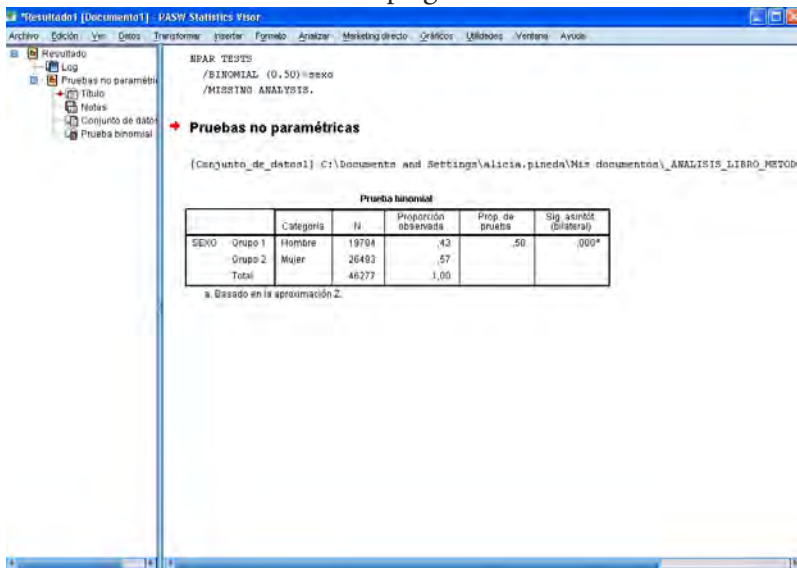
Figura 16
Caja de diálogo que muestra la *Prueba binomial*



En el *Visor de resultados* se despliega una tabla de la prueba binomial que muestra el número de casos por categorías, las proporciones observadas (0.43 en hombres y 0.57 en mujeres) y la proporción esperada en la prueba (0.50). En la última columna aparece el valor de significancia p (Sig. = 0.000) que permite identificar una diferencia significativa entre la proporción observada y la esperada de hombres y mujeres (figura 17).

Figura 17

Tabla de *Prueba binomial*, desplegada en el Visor de resultados



Análisis bivariado: tablas de contingencia

Una vez que se apreció el panorama respecto a la distribución de las variables individuales, se puede ahondar en el análisis con la búsqueda de relaciones entre pares de variables. La pregunta emergente sería: ¿existirá una relación significativa entre las dos variables de nuestro interés?

La relación se llevará a cabo mediante el cruce de valores de las variables dentro de una tabla de contingencia. Esta tabla incluye tanto las frecuencias y porcentajes individuales de cada variable (va-

lores marginales) como las condicionadas por la combinación de las variables (filas y columnas). Sus valores permiten sustentar interpretaciones generales sobre la existencia o no de patrones de variación conjunta entre variables.

En forma complementaria, los datos de la tabla de contingencia habilitan la posibilidad de calcular pruebas para contrastar hipótesis estadísticas. Estas pruebas sirven para establecer si existe o no relación entre las variables, en términos de significancia o probabilidad estadística. Es decir, permiten inferir si se acepta o rechaza la hipótesis nula y si la relación entre las variables se debe o no al azar.

Por ello, antes de analizar una relación de pares de variables, es mejor proponer una hipótesis de nulidad (las variables no tienen relación o son independientes) en contraste con una hipótesis alternativa (las variables tienen relación o son dependientes). Estas hipótesis se evalúan con el cómputo de las pruebas de significancia estadística.

En el caso de las variables de nuestro ejemplo, se podrían proponer las dos hipótesis estadísticas expresadas en la tabla 10, con el propósito de realizar su contrastación mediante pruebas de significancia.

Tabla 10

Hipótesis estadísticas para la relación de dos variables estructurales con la variable *uso del condón* en la última relación sexual

Hipótesis estadísticas
H_0 (Hipótesis nula): el estrato de marginalidad <i>no tiene</i> relación con el uso del condón en la última relación sexual (existe independencia entre las variables).
H_1 (Hipótesis alterna): el estrato de marginalidad <i>tiene</i> relación con el uso del condón en la última relación sexual (existe dependencia entre las variables).
H_0 (Hipótesis nula): el estrato de urbanidad <i>no tiene</i> relación con el uso del condón en la última relación sexual (existe independencia entre las variables).
H_1 (Hipótesis alterna): el estrato de urbanidad <i>tiene</i> relación con el uso del condón en la última relación sexual (existe dependencia entre las variables).

Fuente: Elaboración propia.

Bajo la orientación de tales hipótesis, primero, se realizará la descripción de la relación de valores de los pares de variables en tablas de contingencia. Para ello se seguirán los siguientes pasos en el SPSS: en la barra del menú del SPSS se sigue la secuencia *Analizar/Tablas de contingencia* (figura 18). Se despliega una caja de diálogo que contiene la lista de las variables disponibles. De esta lista se seleccionan, primero, las variables *Estrato de marginalidad* y *Estrato de urbanidad* y se trasladan al espacio *Fila*, en seguida se elige la variable *Uso del condón en la última relación sexual* y se mueve al espacio *Columnas*. Luego se marca *Mostrar los gráficos de barras agrupadas* y se ingresa a *Casillas*, donde se marca *Porcentajes de Fila y de Columna*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 19).

Figura 18
Secuencia a seguir en el menú del SPSS para ingresar
a *Tablas de contingencia*

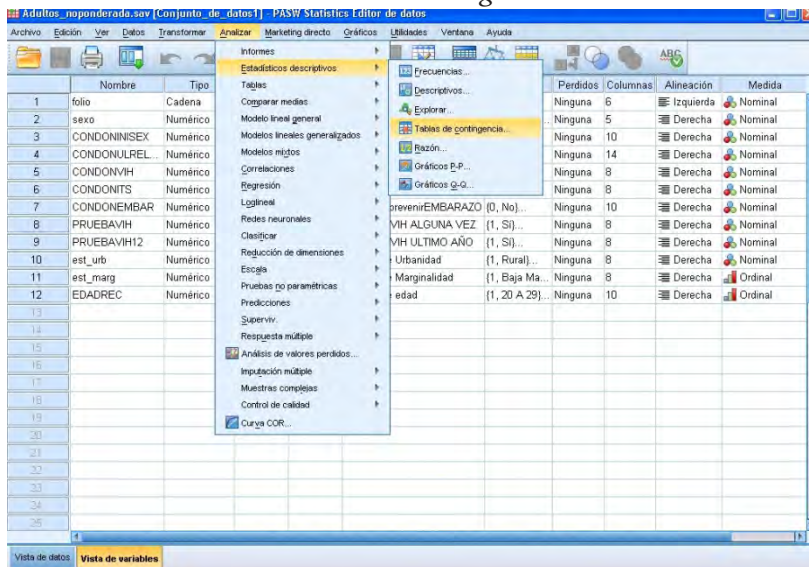
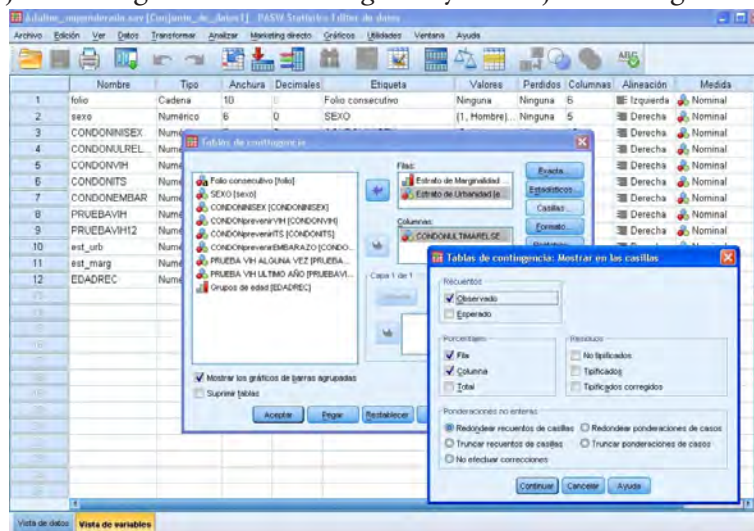


Figura 19
Caja de diálogo *Tablas de contingencia* y sub caja de diálogo *Casillas*



Dicho procedimiento conduce a que se desplieguen en el visor de resultados dos tablas de contingencia y dos gráficos de barras. En la tabla 11 aparece primero la tabla que resultó por el cruce de las variables “estrato de marginación” y “uso del condón en la última relación sexual”. La misma expresa las filas y columnas de valores cruzados de las variables y totales de cada fila y columna. Así, las filas de *recuento* expresan la frecuencia de casos en cada casilla (por ejemplo, 4,583 casos de baja marginalidad y 1,827 de alta marginalidad usaron el condón).

Las filas de % dentro de *Estrato de marginalidad* indican el porcentaje correspondiente a los casos en cada casilla para las categorías de marginalidad (por ejemplo, 24.8% de adultos de marginalidad baja y 17.23% de marginalidad usaron el condón). Finalmente, las filas de % dentro de “Condonultimarelsex” expresan el porcentaje correspondiente a los casos en cada casilla para las categorías *Uso del condón en la última relación sexual* (por ejemplo, 71.5% de casos que usaron el condón fueron de baja marginalidad y 28.5% de marginalidad alta).

En la figura 20 aparece el gráfico de barras correspondiente a las frecuencias ocurridas en la relación de variables que muestra la tabla de contingencia. El gráfico ilustra con claridad una mayor frecuencia de uso del condón en el estrato de marginación baja frente al de marginación alta.

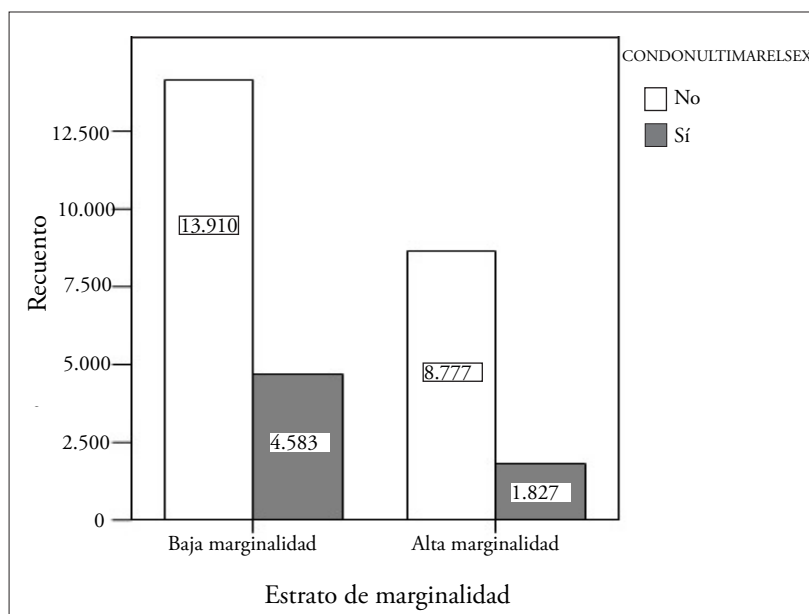
Tabla 11

Tabla de contingencia *estrato de marginalidad* de la variable *uso del condón en la última relación sexual*, desplegada en el visor de resultados*

Estrato de marginalidad		CONDONULTIMARELSEX		Total
		No	Sí	
Baja marginalidad	Recuento	13,910	4,583	18,493
	% dentro de estrato de marginalidad	75.2%	24.8%	100.0%
	% dentro de CONDONULTIMARELSEX	61.3%	71.5%	63.6%
Alta marginalidad	Recuento	8,777	1,827	10,604
	% dentro de estrato de marginalidad	82.8%	17.2%	100.0%
	% dentro de CONDONULTIMARELSEX	38.7%	28.5%	36.4%
Total	Recuento	22,687	6,410	29,097
	% dentro de estrato de marginalidad	78.0%	22.0%	100.0%
	% dentro de CONDONULTIMARELSEX	100.0%	100.0%	100.0%

* La estructura en que se presenta la tabla original, proporcionada por el spss, es distinta, sin embargo la información contenida es la misma que aquí se muestra. Idéntico procedimiento se siguió en las tablas restantes del presente capítulo.

Figura 20
Gráfico que ilustra las frecuencias de la relación
entre las variables mostradas en la tabla 11

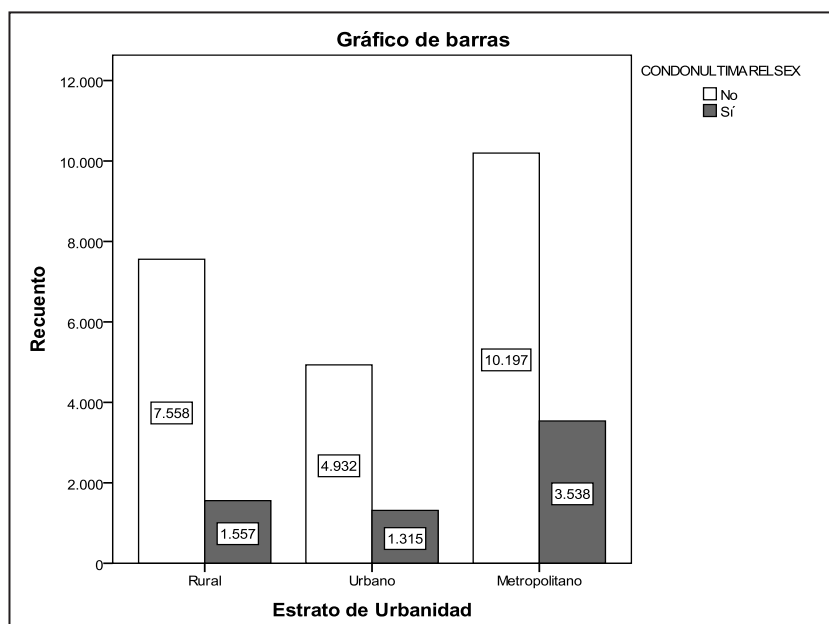


En la tabla 12 se muestran la tabla de contingencia y el gráfico de barras referente al cruce de las variables “estrato de urbanidad” y “uso del condón en la última relación sexual”. En la tabla se rescata que el porcentaje de uso es mayor en el estrato metropolitano y que el mismo decrece en los contextos urbano y rural. En el caso del gráfico (figura 21), se puede apreciar que si bien las barras *frecuencias de casos* muestran diferencias relevantes del uso en el estrato metropolitano frente a los otros, no son sensibles para expresar la diferencia de uso en los estratos urbano y rural. Ello debido a que el gráfico se limita a las frecuencias y no considera los porcentajes.

Tabla 12
 Tabla de contingencia *estrato de urbanidad* de la variable
uso del condón en la última relación sexual, desplegada
 en el visor de resultados

Estrato de urbanidad		CONDONULTIMARELSEX		Total
		No	Sí	
Rural	Recuento	7,558	1,557	9,115
	% dentro de estrato de urbanidad	82.9%	17.1%	100.0%
	% dentro de CONDONULTIMARELSEX	33.3%	24.3%	31.3%
Urbano	Recuento	4,932	1,315	6,247
	% dentro de estrato de urbanidad	78.9%	21.1%	100.0%
	% dentro de CONDONULTIMARELSEX	21.7%	20.5%	21.5%
Metropolitano	Recuento	10,197	3,538	13,735
	% dentro de estrato de urbanidad	74.2%	25.8%	100.0%
	% dentro de CONDONULTIMARELSEX	44.9%	55.2%	47.2%
Total	Recuento	22,687	6,410	29,097
	% dentro de estrato de urbanidad	78.0%	22.0%	100.0%
	% dentro de CONDONULTIMARELSEX	100.0%	100.0%	100.0%

Figura 21
Gráfico que ilustra las frecuencias de la relación
entre las variables mostradas en la tabla 12

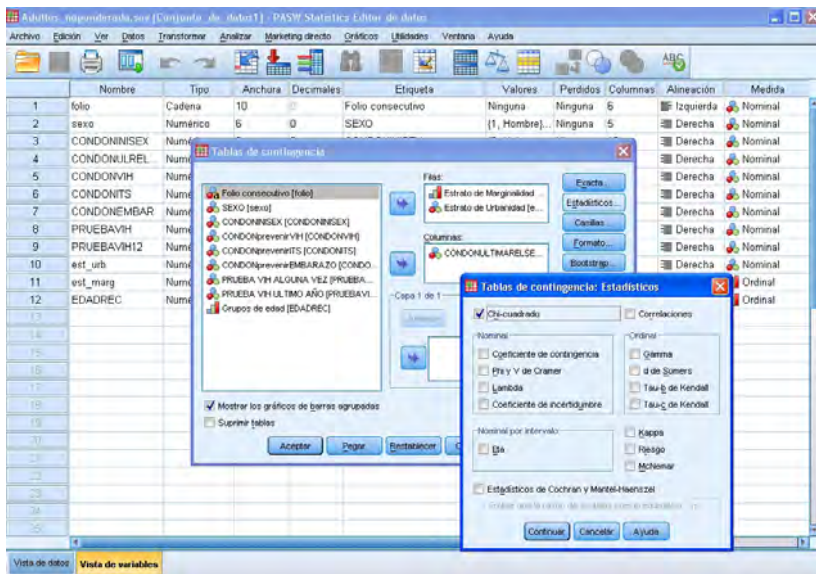


Hasta donde se han podido ver los resultados, los valores combinados de las variables analizadas parecen mostrar la existencia de relación entre ambos pares de variables: tanto los adultos del estrato de baja marginalidad como los del estrato metropolitano usaron el condón en mayor porcentaje que los demás. Este hallazgo es consistente con el supuesto teórico del mayor acceso a los métodos de prevención del riesgo sexual en donde existe menor porcentaje de marginalidad y en los contextos más urbanizados.

Sin embargo, para trascender el nivel de los datos descriptivos y de las relaciones aparentes, es necesario contrastar las hipótesis de independencia propuestas previamente (véase tabla 10), mediante el cálculo de pruebas de hipótesis adecuadas a las escalas de medición de las variables. En este caso, la prueba apropiada para las escalas nominal y ordinal de las variables a relacionar es el

Chi-cuadrado (véase tabla 8). Para generar esta prueba de hipótesis con el SPSS se sigue el siguiente camino: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Tablas de contingencia*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles y las ya seleccionadas. Allí se ingresa a *Estadísticos* donde se marca *Chi-cuadrado*. Finalmente se eligen Continuar y Aceptar para ejecutar el procedimiento (figura 22)

Figura 22
Caja de diálogo *Tablas de contingencia*
y sub caja de diálogo *Estadísticos*



Las tablas de *Chi-cuadrado* resultantes se despliegan en el visor de resultados (tablas 13 y 14). La prueba *Chi-cuadrado de Pearson* sobre “estrato de marginalidad” tiene un valor de 223,836, con 1 grado de libertad, donde su significancia (asintótica bilateral) es $P = 0.000$. Estos valores muestran que al ser la significancia de $p < 0.01$, se rechazaría la hipótesis nula (la no relación entre las variables) y se aceptaría la hipótesis alterna (la relación de dependencia entre las variables). Por su parte, los va-

lores de la prueba *Chi-cuadrado* correspondientes a “estrato de urbanidad” muestran también una significancia de $p < 0.01$ que permite rechazar la hipótesis nula y aceptar la relación de dependencia. Ambos resultados permiten corroborar la relación de dependencia entre los pares de variables que se pensaba existía con la sola inspección de las tablas de contingencia.

Tabla 13

Tabla de la prueba Chi-cuadrado referente a *uso del condón en la última relación sexual*, en el estrato *marginalidad*, desplegada en el visor de resultados

	Valor	gl	Significancia asintónica (bilateral)	Significancia exacta (bilateral)	Significancia exacta (unilateral)
Chi-cuadrado de Pearson	223,836 ^a	1	0.000		
Corrección por continuidad ^b	223,396	1	0.000		
Razón de verosimilitudes	230,001	1	0.000		
Estadístico exacto de Fisher				0.000	0.000
Asociación lineal por lineal	223,828	1	0.000		
N de casos válidos	29,097				

^a. 0 casillas (,0%) tienen una frecuencia esperada inferior a 5. La frecuencia mínima esperada es 2336,04.

^b. Calculado sólo para una tabla de 2x2.

Tabla 14

Tabla de la prueba Chi-cuadrado referente a *uso del condón en la última relación sexual*, en el estrato *urbanidad*, desplegada en el visor de resultados

	Valor	gl	Significancia asintónica (bilateral)
Chi-cuadrado de Pearson	224,629 ^a	2	0.000
Razón de verosimilitudes	248,526	2	0.000
Asociación lineal por lineal	244,223	1	0.000
N de casos válidos	29,097		

^a. 0 casillas (,0%) tienen una frecuencia esperada inferior a 5. La frecuencia mínima esperada es 1376,20.

Una vez que se establecieron las relaciones entre las variables estructurales y el “uso del condón en la última relación sexual” son estadísticamente significativas; sería oportuno indagar también si las variables estructurales tienen una asociación significativa con los motivos de uso del condón. Si se sigue el mismo procedimiento para generar tablas de contingencia y pruebas *Chi-cuadrado* en el spss (figura 23) se obtendrán los resultados que sintetizamos en la tabla 15. En la misma se puede ver, primero, que las variables estructurales se asocian con el “uso del condón para prevenir el VIH”, con diferencias significativas en favor del estrato de marginación baja (28.3% *versus* 24.2%; $p < 0.01$) y del estrato de urbanidad metropolitano (29.1% *versus* 27.3% y 22.2%; $p < 0.01$). También se puede rescatar la asociación significativa del “estrato de urbanidad” con el uso del condón para prevenir las ITS, con un porcentaje mayor en el estrato metropolitano (54% *versus* 50.3% y 52.0%, $p < 0.05$).

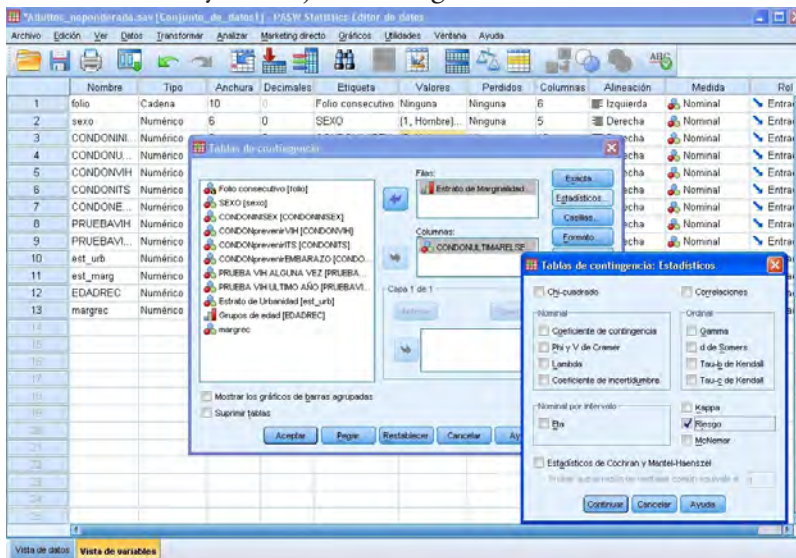
Tabla 15
Pruebas Chi-cuadrado de la relación entre variables estructurales
y motivos de uso del condón en la última relación sexual

Estratos	Motivos <i>Uso preventivo del condón</i>					
	Embarazo	χ^2 (g.l.) P ²	ITS	χ^2 (g.l.) P ²	VIH	χ^2 (g.l.) P ²
<i>Marginación</i>						
Baja (n = 3,015)	76.1	0.14 (1) 0.712	53.3	0.49 (1) 0.484	28.2	11.66 (1) 0.001**
Alta (n = 2,048)	76.5		53.2		24.2	
<i>Urbanidad</i>						
Rural	77.4	5.08 (2) 0.079	52.0	8.19 (2) 0.017*	22.2	29.49 (2) 0.000**
Urbano	77.4		50.3		27.3	
Metropolitano (n = 3,877)	75.1		54.4		29.1	

^a χ^2 = Chi-cuadrado; (g.l.) = grados de libertad; P = valor de significancia P
*Relación significativa al nivel P <0.005
**Relación significativa al nivel P <0.01

En el caso de las tablas de contingencia de 2 x 2 categorías, un cálculo complementario será *estimación del riesgo* para establecer el grado de asociación entre las variables. Para ello se calculará el índice de *Razón de ventajas* (conocido también como *Razón de momios* u *Odds Ratio*) y su intervalo de confianza de 95% (IC 95%); volviendo a las variables generales de nuestro ejemplo, se puede hacer este cálculo para la relación entre el “estrato de marginación” y el “uso del condón en la última relación sexual”. Para hacer el cálculo en el SPSS se realiza lo siguiente: en la barra del menú del SPSS se sigue la secuencia *Analizar/Tablas de contingencia*. Se seleccionan las variables requeridas y se insertan en los espacios de *Filas* y *Columnas*. Luego se ingresa a *Estadísticos* donde se marca *Riesgo*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 23).

Figura 23
Caja de diálogo que muestra las *Tablas de contingencia* y sub caja de diálogo *Estadísticos*



En el visor de resultados se despliega una tabla de contingencia referente a la relación de variables y otra sobre estimación de riesgo (tablas 16 y 17). La primera hace el recuento panorámico de las frecuencias de casos cruzados y los totales. En la segunda, los valores a leer e interpretar son los de la primera fila: el índice de razón de ventajas (*Razón de momios* u *Odds Ratio*) que tiene un valor de 1.583 y un intervalo de confianza de 95% (IC de 95%) que varía de 1.490 a 1.681. Estos valores señalan la presencia de 1.58 veces más probabilidad de usar el condón en los adultos de marginación baja (codificados como 1) comparados con los de marginación alta (codificados como 0), y que dicha probabilidad es mayor de 1.49 a 1.68 veces en 95% de los entrevistados. Se interpreta que al no pisar este intervalo el valor de 1.0 se corroboraría la existencia de la asociación.

Tabla 16
Tabla de contingencia de la relación marginación
y uso del condón en la última relación sexual

	Recuento		
	CONDONULTIMARELSEX		Total
	No	Sí	
Marginación alta	8,777	1,827	10,604
Marginación baja	13,910	4,583	18,493
Total	22,687	6,410	29,097

Tabla 17
Estimación de riesgo de uso del condón en la última relación sexual

	Valor	Intervalo de confianza al 95%	
		Inferior	Superior
Razón de las ventajas para marginalidad (Alta = 0 / Baja = 1)	1.583	1.490	1.681
Para la cohorte CONDONULTIMARELSEX = No	1.100	1.087	1.114
Para la cohorte CONDONULTIMARELSEX = Sí	0.695	0.662	0.730
N de casos válidos	29,097		

La estimación del riesgo también se puede realizar en estudios de casos y controles y de cohorte, como se observa en el ejemplo del recuadro 1:

Recuadro 1			
Estimación del riesgo en diseños de casos y controles y cohorte			
<i>Diseños de casos y controles</i> En los diseños de casos y controles, o retrospectivos, se forman dos grupos de personas a partir de alguna condición de enfermedad que interese indagar. Luego, se busca en la historia clínica de las mismas personas la presencia del algún factor que podría desencadenar la enfermedad. Por ejemplo, en los datos ficticios de la tabla siguiente se expresa la relación entre la presencia de enfermedad vascular y el consumo actual de tabaco como posible desencadenante de la enfermedad, en una muestra de 480 personas:			
Consumo actual de tabaco	Enfermedad vascular		Total
	Sí	No	
Sí	46	162	208
No	18	254	272
Total	64	416	480

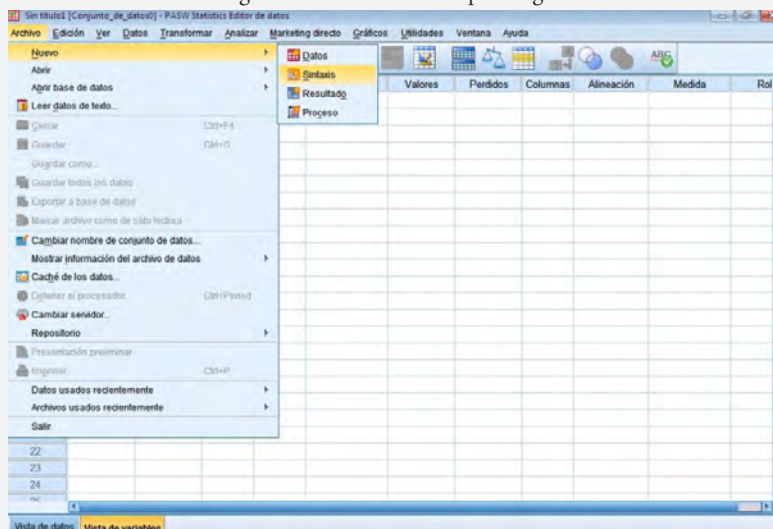
Continúa en la página 146

Viene de la página 145

La razón de *consumidores/no consumidores* de tabaco en el grupo con enfermedad vascular es de $46/18 = 2.555$ y en el de *sin enfermedad vascular* es de $162/254 = 0.638$. El índice de riesgo se obtiene al dividir ambas razones: $2.555/0.638 = 4.005$. Este índice estima el riesgo relativo y puede interpretarse como tal, pero también puede señalarse que entre las personas con enfermedad vascular es 4.005 más probable hallar consumidores actuales de tabaco que no consumidores.

En el paquete *SPSS* existe la posibilidad de generar tablas de contingencia tan sólo con los valores de las casillas, sin necesidad de contar con bases de datos completas, y calcular los índices de riesgo. Para generar la tabla de nuestro ejemplo se elige, en el menú *Archivo* del *SPSS*, la opción *Sintaxis* (figura A):

Figura A
Secuencia a seguir en le menú del *SPSS* para ingresar a *Sintaxis*



En seguida, aparecerá una pantalla en blanco donde se debe escribir con exactitud (en comandos y puntuación) el siguiente programa que indica cuáles serán las variables (consumo y vascular) y los valores de las casillas que contendrá la tabla de contingencia, así como los cálculos a realizar: frecuencias (*count*), valores esperados (*expected*) y el índice de riesgo (*risk*):

```
data list/consumo 1 vascular 3 wt 5-7.
weight by wt.
begin data.
end data.
```

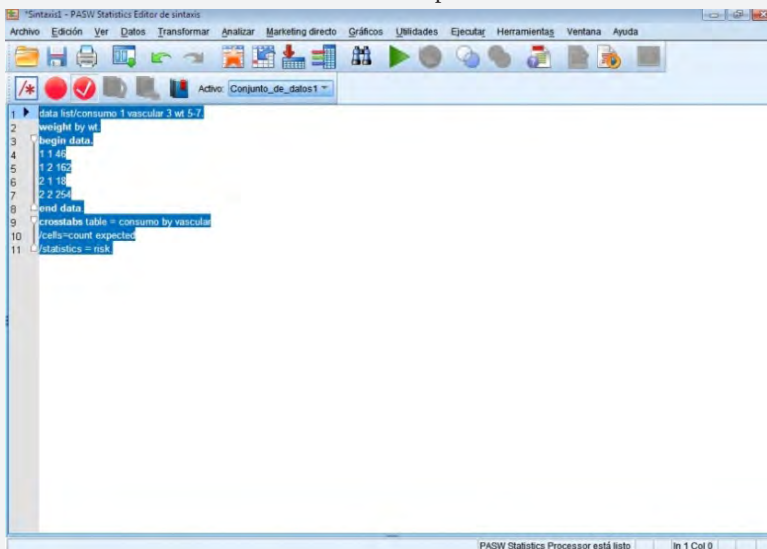
Continúa en la página 147

Viene de la página 146

```
crosstabs table = consumo by vascular
/cells=count expected
/statistics = risk.
```

Una vez escrito el programa en la pantalla de sintaxis, se seleccionará todo su texto y se hará clic en el botón  para ejecutarlo (figura B):

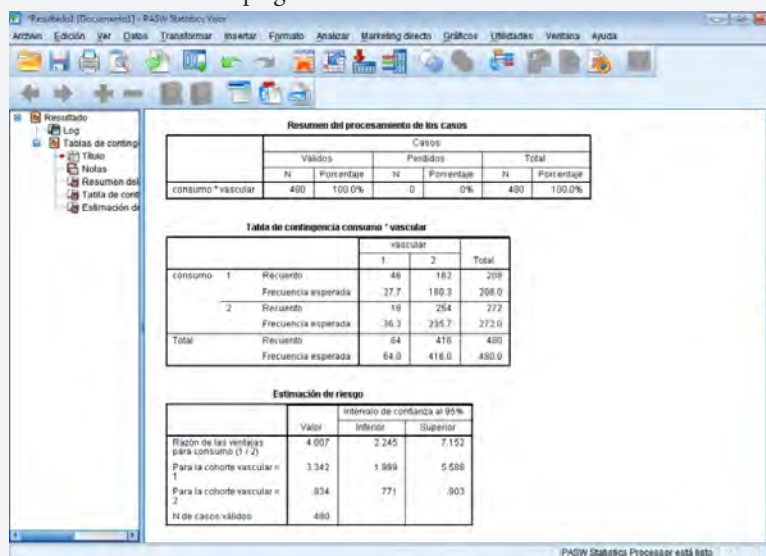
Figura B
Selección del texto en el SPSS y ejecución
mediante el botón  en la pantalla de *Sintaxis*



Como resultado de ejecutar el programa, se generan las tablas de *contingencia* y de *estimación de riesgo*. La primera coincide con la tabla del ejemplo (en frecuencias de las casillas) y la segunda expresa en su primera fila el valor de la razón de ventajas (4.007) y un IC de 95% (2.245 a 7.152) que indica la probabilidad de encontrar *consumo actual de tabaco/no consumo* entre las personas enfermas, con un intervalo de confianza aplicable a 95% de la muestra (figura C):

Continúa en la página 148

Figura C
Tablas de contingencia y de estimación de riesgo,
desplegadas en el visor de resultados



Diseños de cohorte

En los diseños de cohorte o prospectivos, se forman dos grupos de personas en función de la presencia o ausencia de un factor desencadenante (en nuestro ejemplo, el consumo actual de tabaco) y se les sigue en el tiempo para establecer en qué proporción de cada grupo ha ocurrido un determinado desenlace (en nuestro ejemplo, la presencia de enfermedad vascular).

Entre los consumidores actuales de tabaco, la proporción de casos con enfermedad vascular es de $46/208 = 0.221$, mientras que la de no consumidores es de $18/272 = 0.066$. Al dividir ambas proporciones se calcula el riesgo relativo (número de veces en que es más probable encontrar enfermedad vascular en consumidores actuales de tabaco que en no consumidores): $0.221/0.066 = 3.348$ veces más. El valor del índice calculado a mano se asemeja al que reporta la segunda fila de la tabla *estimación de riesgo* (figura C) para la cohorte vascular = 1 o con presencia de enfermedad: 3.342. El intervalo de confianza de 95% mayor a 1 expresa que la diferencia entre grupos es significativa y que posee un sentido de riesgo. El índice presentado en la tercera fila de la misma tabla corresponde a la cohorte vascular = 2 o sin presencia de enfermedad: 0.834. Su valor expresa el número de veces en que es más probable no encontrar enfermedad vascular en consumidores actuales de tabaco que en no consumidores. El intervalo de confianza de 95% menor a 1 expresa que la diferencia entre grupos es significativa y que posee un sentido de protección.

Como ya sabemos que existe relación estadísticamente significativa entre las variables estructurales y el “uso del condón en la última relación sexual”, es posible también ajustar la relación con una tercera variable, a fin de evaluar si su presencia podría generar algún efecto de confusión o sesgo en la relación encontrada.

Para ello, se incluirá en el análisis de contingencia a la variable “sexo” que, evaluada por separado con la prueba *Chi-cuadrado*, alcanzó una asociación significativa con el “uso del condón en la última relación sexual”: valor de $p < 0.01$ y Riesgo de 2.311 veces mayor de uso en los hombres que en las mujeres (IC de 95%: 2.185-2.446). Al respecto véanse las siguientes tablas (18 y 19).

Tabla 18
Pruebas de Chi-cuadrado

	Valor	gl	Significancia asintónica (bilateral)	Significancia exacta (bilateral)	Significancia exacta (unilateral)
Chi-cuadrado de Pearson	870,971 ^a	1	0.000		
Corrección por continuidad ^b	870,125	1	0.000		
Razón de verosimilitudes	861,421	1	0.000		
Estadístico exacto de Fisher				0.000	0.000
Asociación lineal por lineal	870,941	1	0.000		
N de casos válidos	29,097				

^a. 0 casillas (,0%) tienen una frecuencia esperada inferior a 5. La frecuencia mínima esperada es 2689,39.

^b. Calculado sólo para una tabla de 2x2.

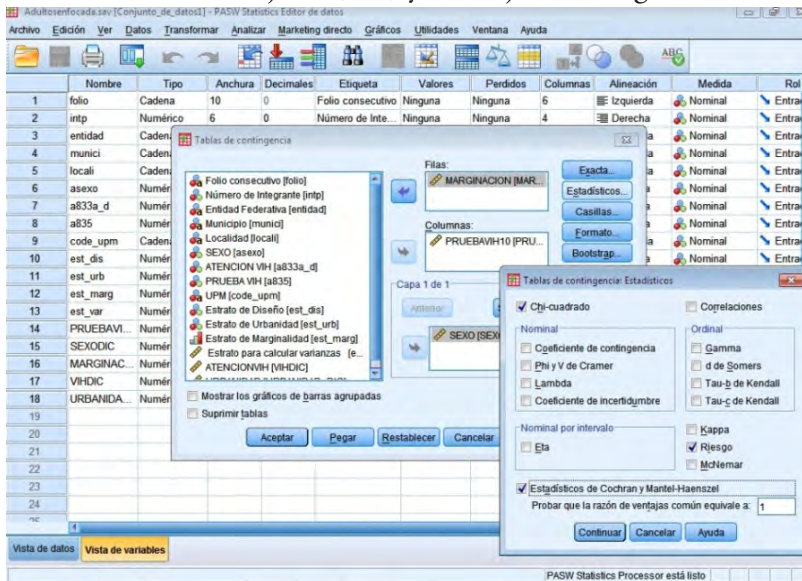
Tabla 19
Estimación de riesgo

	Valor	Intervalo de confianza al 95%	
		Inferior	Superior
Razón de las ventajas para sexo (Mujer = 0 / Hombre = 1)	2.311	2.185	2.446
Para la cohorte CONDONULTIMARELSEX = No	1.209	1.193	1.225
Para la cohorte CONDONULTIMARELSEX = Sí	0.523	0.501	0.546
N de casos válidos	29,097		

El ajuste con la tercera variable se realizará al generar tablas de contingencia individuales para cada estrato de la variable “sexo”, junto con sus pruebas de asociación *Chi-cuadrado* y *estimación de Riesgo* (valores de Razones de ventaja para cada estrato de “sexo” y de Razón de ventajas común de Mantel-Haenszel). Para realizar este análisis con el SPSS se procede de la siguiente forma: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Tablas de contingencia*. Se despliega una caja de diálogo que contiene la lista de las variables disponibles y las ya seleccionadas. Se elige la tercera variable en la lista de variables (*Sexo*) y se traslada al espacio *Capa*. Luego se ingresa a *Estadísticos* donde se marcan *Chi-cuadrado*, *Riesgo* y *Estadístico de Cochran y Mantel y Haenszel*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 24).

Figura 24

Caja de diálogo *Tablas de contingencia* con una variable de ajuste (*sexo*) y sub caja de diálogo *Estadísticos*



En el visor de resultados aparecen, primero, las tablas de contingencia para hombres y mujeres por separado. Las mismas permiten ver que, en ambos sexos, los porcentajes de “uso del condón en la última relación sexual” son mayores en el estrato de marginación baja que en el de marginación alta, con proporciones ligeramente más altas en las mujeres (tabla 20).

Tabla 20

Tabla de contingencia de *uso del condón en la última relación sexual y marginación* con una variable de ajuste (*sexo*), desplegada en el visor de resultados

Sexo	Estrato		CONDONULTIMARELSEX		Total
			No	Sí	
Mujeres	Marginación alta	Recuento	5,459	748	6,207
		% dentro de marginación	87.9%	12.1%	100.0%
		% dentro de CONDONULTIMARELSEX	38.4%	27.8%	36.8%
	Marginación baja	Recuento	8,739	1,943	10,682
		% dentro de marginación	81.8%	18.2%	100.0%
		% dentro de CONDONULTIMARELSEX	61.6%	72.2%	63.2%
	Total	Recuento	14,198	2,691	16,889
		% dentro de marginación	84.1%	15.9%	100.0%
		% dentro de CONDONULTIMARELSEX	100.0%	100.0%	100.0%
Hombres	Marginación alta	Recuento	3,318	1,079	4,397
		% dentro de marginación	75.5%	24.5%	100.0%
		% dentro de CONDONULTIMARELSEX	39.1%	29.0%	36.0%
	Marginación baja	Recuento	5,171	2,640	7,811
		% dentro de marginación	66.2%	33.8%	100.0%
		% dentro de CONDONULTIMARELSEX	60.9%	71.0%	64.0%
	Total	Recuento	8,489	3,719	12,208
		% dentro de marginación	69.5%	30.5%	100.0%
		% dentro de CONDONULTIMARELSEX	100.0%	100.0%	100.0%

Luego aparece una tabla con la prueba Chi-cuadrado que contiene los valores calculados para cada estrato de la variable *sexo* (tabla 21). En la misma se muestran relaciones estadísticamente significativas ($p < 0.05$ en ambos estratos) que confirman la relación entre las variables “estrato de marginación” y “uso del condón en la última relación sexual”, tanto en hombres como en mujeres.

Tabla 21

Tabla de Chi-cuadrado de *uso del condón en la última relación sexual* con una variable de ajuste (*sexo*) desplegada en el visor de resultados

Sexo		Valor	gl	Significancia asintónica (bilateral)	Significancia exacta (bilateral)	Significancia exacta (unilateral)
Mujeres	Chi-cuadrado de Pearson	110,441 ^a	1	0.000		
	Corrección por continuidad ^b	109,983	1	0.000		
	Razón de verosimilitudes	114,280	1	0.000		
	Estadístico exacto de Fisher				0.000	0.000
	Asociación lineal por lineal	110,435	1	0.000		
	N de casos válidos	16,889				
Hombres	Chi-cuadrado de Pearson	113,856 ^c	1	0.000		
	Corrección por continuidad ^b	113,419	1	0.000		
	Razón de verosimilitudes	116,115	1	0.000		
	Estadístico exacto de Fisher				0.000	0.000
	Asociación lineal por lineal	113,847	1	0.000		
	N de casos válidos	12,208				

^a. 0 casillas (,0%) tienen una frecuencia esperada inferior a 5. La frecuencia mínima esperada es 988,99.

^b. Calculado sólo para una tabla de 2x2.

^c. 0 casillas (,0%) tienen una frecuencia esperada inferior a 5. La frecuencia mínima esperada es 1339,49.

Finalmente, se despliegan tanto una tabla de estimación de riesgo con valores de *Razones de ventaja* para cada estrato, como una de *Razón de ventajas común* de Mantel-Haenszel (tablas 22 y 23). La primera tabla muestra que en ambos sexos el riesgo de “usar el condón en la última relación sexual” es mayor si pertenecen al “estrato de marginación baja” (1.570 veces más riesgo en los hombres y 1.623 veces más en las mujeres), mientras que la segunda presenta una *Estimación de razón de ventajas común* de 1.594 con un IC de 95%, de 1.499 a 1.695, lo cual indica que la variable “sexo” no genera un efecto de confusión en la relación evaluada, debido a que los valores de la *Razón de ventajas* de cada estrato y de la *Razón de ventajas común* no tuvieron cambios significativos al mantenerse mayores a 1.

Tabla 22
Tabla *estimación de riesgo* desplegada en el visor de resultados

Sexo		Valor	Intervalo de confianza al 95%	
			Inferior	Superior
Mujeres	Razón de las ventajas para marginalidad (Alta = 0 / Baja = 1)	1.623	1.482	1.777
	Para la cohorte CONDONULTIMARELSEX = No	1.075	1.061	1.089
	Para la cohorte CONDONULTIMARELSEX = Sí	0.663	0.613	0.716
	N de casos válidos	16,889		
Hombres	Razón de las ventajas para marginalidad (Alta = 0 / Baja = 1)	1.570	1.445	1.706
	Para la cohorte CONDONULTIMARELSEX = No	1.140	1.114	1.167
	Para la cohorte CONDONULTIMARELSEX = Sí	0.726	0.683	0.771
	N de casos válidos	12,208		

Tabla 23

Estimación de la razón de las ventajas común de Mantel-Haenszel*

Estimación			1.549
In(estimación)			0.446
Error típ. de In(estimación)			0.031
Significancia asintótica (bilateral)			0.000
Intervalo de confianza asintótica al 95%	Razón de ventajas común	Límite inferior	1.499
		Límite superior	1.695
	In(Razón de ventajas común)	Límite inferior	0.405
		Límite superior	0.528

*La estimación de la razón de las ventajas común de Mantel-Haenszel se distribuye de manera asintóticamente normal bajo el supuesto de razón de las ventajas común igual a 1,000. Lo mismo ocurre con el log natural de la estimación.

Cabe agregar que los efectos de confusión se pueden detectar cuando los valores de la *Razón de ventajas común* y de los estratos se vuelven contradictorios, debido al ajuste. De acuerdo con la forma de contradicción de los valores, los tipos de confusión podrían ser por relación espuria, por relación enmascarada o por relación invertida o paradójica, como se aprecia en la siguiente tabla (24).

Tabla 24

Tipos de efectos de confusión según la contradicción en los valores de las *Razones de ventajas*

Tipos de confusión	Valores de las Razones de ventajas (RM)
Relación espuria	La RM común muestra una asociación significativa ($RM > 1$), mientras que las RM de los estratos muestran relaciones no significativas ($RM = 1$).
Relación enmascarada	La RM común muestra una relación no significativa ($RM = 1$), mientras que las RM de los estratos muestran asociaciones significativas ($RM > 1$).
Relación invertida o paradójica	Las RM de los estratos muestran asociaciones significativas ($RM > 1$), mientras que la RM común muestra una asociación significativa en sentido inverso ($RM < 1$).

Análisis multivariado: regresión logística binaria

A estas alturas es posible realizar un análisis más complejo para ir más allá de las relaciones de pares de variables expresadas en las tablas de contingencia. La técnica de la regresión logística binaria nos puede ayudar a realizar un análisis donde se evalúe el efecto de distintas variables independientes sobre una variable dependiente dicotómica. Esta evaluación abarcará dos aspectos: 1) el ajuste del efecto de cada variable con el de las otras variables independientes, a fin de controlar su potencial efecto de confusión; y 2) la estimación del efecto independiente de cada variable como factor predictor de la variable dependiente dentro de un modelo que sea estadísticamente adecuado.

En la regresión logística binaria, se incluye una variable dependiente binaria (con categorías 1 = sí y 0 = no) y dos o más variables independientes que tengan escalas de medición dicotómicas (categorías 1 y 0) o de intervalo/razón.

El primer paso orientado a construir un modelo de regresión logística consiste en identificar las variables a incluir en el mismo, para lo cual se siguen dos criterios generales de inclusión: 1) la elección de aquellas variables que tuvieron asociación significativa ($p < 0.05$) en el análisis previo de pares de variables (tablas de contingencia). Este criterio es laxo y abierto a la posibilidad de incluir variables con valores cercanos a la significancia estadística ($p = 0.10$ o $p = 0.15$); y 2) la inclusión de variables sin asociación significativa pero con relevancia teórica para la explicación de la variable dependiente.

En nuestro análisis de tablas de contingencia encontramos relaciones significativas de las variables estructurales (estratos de marginalidad y urbanidad) y del sexo con la que será nuestra variable dependiente (uso del condón en la última relación sexual [véanse Figuras 14 y 21]). Por lo que, si seguimos el primer criterio de inclusión, tendríamos ya identificadas a tres variables independientes para construir nuestro modelo de regresión logística.

Dado el repertorio de variables con las que contamos, nos interesaría también incluir otras dos variables independientes a nuestro modelo: el “uso del condón en la primera relación sexual” y los “grupos de edad”. Sin contar con los cálculos de significancia, asumiremos, por los antecedentes reportados en la literatura especializada, que estas variables son relevantes en la explicación de nuestra variable dependiente. De tal manera, tendríamos ya un conjunto de variables que nos permiten proponer la posibilidad de construir el modelo que se expresa en la siguiente tabla (25).

Tabla 25

Variables independientes y variable dependiente del modelo de regresión logística binaria que se intentará construir

Variables independientes	Variable dependiente
$\bar{x}_1$, Estrato de marginalidad	Y. Uso del condón en la última relación sexual
$\bar{x}_2$, Estrato de urbanidad	
$\bar{x}_3$, Uso del condón en la primera relación sexual	
$\bar{x}_4$, Sexo	
$\bar{x}_5$, Grupos de edad	

Debido a que nuestras variables de ejemplo son de escala nominal y ordinal, será necesario recodificar sus categorías a dicotómicas para habilitar su incorporación al modelo de regresión que se intenta construir. Tal transformación las convertirá en las llamadas “variables ficticias” (conocidas como variables *dummy*) que permitirán la comparación de categorías en el cálculo de la regresión (1 = variable calculada *versus* 0 = variable de referencia o comparación).

La referida recodificación se hará, en las variables de nuestro modelo, según el detalle de la tabla 26. Ésta muestra que en las variables nominales y ordinales con dos categorías, la reco-

dificación a dicotómica (1 y 0) es suficiente, mientras que en el caso de las variables ordinales con tres categorías (estrato de urbanidad y grupos de edad), la recodificación requiere crear una serie de dicotomías (00, 10, 01) para habilitar la comparación (10 y 01 = variables calculadas *versus* 00 variable de referencia).

Tabla 26
Variables de ejemplo recodificadas como variables ficticias

Variables	Escala	Valores y etiquetas	Recodificación a variables ficticias
Estrato de marginalidad	Ordinal	1. Baja, 2. Alta	1. Baja, 2. Alta
Estrato de urbanidad	Nominal	1. Rural, 2. Urbano, 3. Metropolitano	00. Rural, 10. Urbana, 01. Metropolitana
Sexo	Nominal	1. Hombres, 2. Mujeres	1. Hombres, 0. Mujeres
Grupos de edad	Ordinal	1. 20-29; 2. 30-45; 3. 46>	10. 20-29; 01. 30-45; 00. 46>
Uso del condón en el inicio sexual	Nominal	1. Sí, 2. No	1. Sí, 0. No

Fuente: Elaboración propia.

La transformación a variables ficticias se puede realizar en el paquete SPSS como un elemento preliminar de preparación de las variables que se incorporarán en el modelo de regresión logística binaria. Para realizarla, se procederá de la siguiente manera: en la barra del menú del SPSS se elige la secuencia *Analizar/Regresión/Logística binaria* (figura 25). Se despliega una caja de diálogo que contiene la lista de las variables disponibles. Se elige, primero, la variable dependiente y se traslada al espacio *Dependientes*. Luego se elige la variable independiente a transformar (*estrato urbanidad*) y se mueve al espacio *Covariables categóricas*. En la opción “cambiar de contraste” de la parte inferior se selecciona la *Categoría de referencia* “primero” (esto porque se quiere

que la primera categoría que es “rural” sea la comparativa con valor 0) y se ingresa a *Cambiar* (figura 26). De la misma manera se eligen las otras variables y se trasladan a *Covariables categóricas*. Allí se selecciona para cada una de éstas su *Categoría de referencia* correspondiente (figura 27). Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento.

Figura 25
Secuencia a seguir en el menú del spss
para ingresar a *Logística binaria*

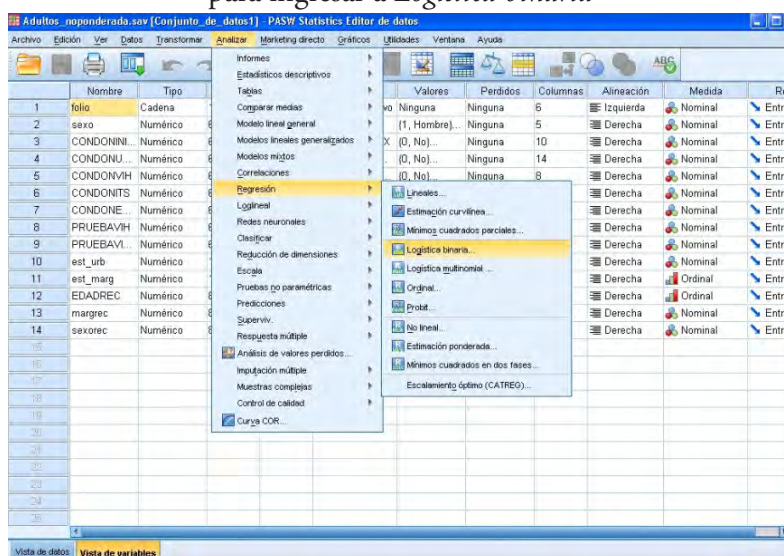


Figura 26
Caja de diálogo Logística binaria
y sub caja de diálogo *Definir variables categóricas*

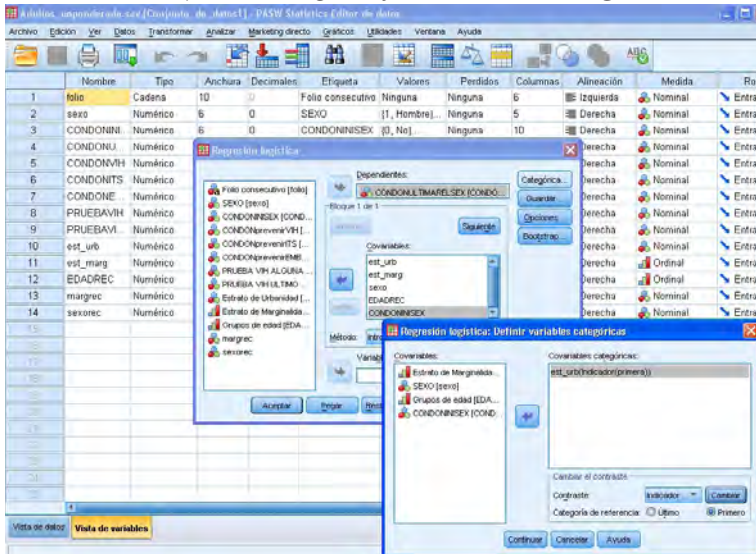
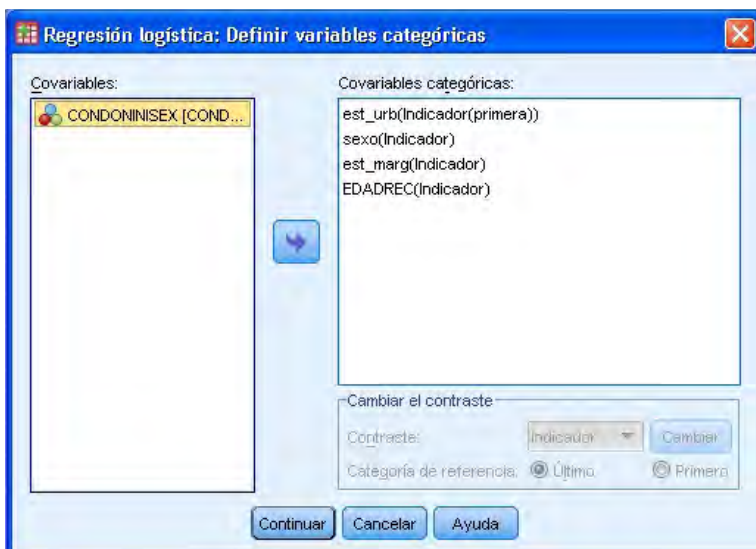


Figura 27
Sub caja de diálogo *Definir variables categóricas*
con las variables del modelo transformadas a ficticias



Tras ejecutar el procedimiento, en el visor del SPSS se desplegarán varias tablas que refieren el análisis de regresión logística. La única que nos interesará inspeccionar, por ahora, es la de *codificaciones de variables categóricas* que reporta los valores dicotómicos asignados a las variables (tabla 27). Estos valores coinciden con los previstos en la tabla 26 y asignan el código 1 a las categorías comparativas y el 0 a las de referencia.

Tabla 27
Tabla de codificaciones
de variables categóricas desplegada en el visor de resultados

		Frecuencia	Codificación de parámetros	
			(1)	(2)
Grupos de edad	20 a 29	8,774	1.000	0.000
	30 a 45	16,916	0.000	1.000
	46 y más	3,407	0.000	0.000
Estrato de urbanidad	Rural	9,115	0.000	0.000
	Urbano	6,247	1.000	0.000
	Metropolitano	13,735	0.000	1.000
Sexo	Hombre	12,208	1.000	
	Mujer	16,889	0.000	
Estrato de marginalidad	Baja	18,493	1.000	
	Alta	10,604	0.000	

Luego de proponer el modelo de regresión logística binaria y de transformar sus variables nominales y ordinales en dicotómicas o ficticias, se iniciará el segundo paso que consiste en la estimación y contraste de los parámetros o coeficientes de regresión logística (β). El propósito de este paso será detectar tanto el efecto lineal como la fuerza del efecto de cada variable independiente sobre la probabilidad de que la variable dependiente tome un determinado valor (py).

En ese sentido, se contrastarán las siguientes hipótesis estadísticas del efecto lineal y de la fuerza del efecto (tabla 28): a) en el contraste del efecto lineal, la hipótesis nula (H_0) es que todos los coeficientes de regresión logística (β) serán igual a cero (no existe relación lineal), mientras que la hipótesis alterna (H_1) expresa que por lo menos algún coeficiente de regresión será distinto de cero y tendrá una relación significativa con la PY; y b) en el contraste de la fuerza del efecto, la H_0 es que el valor del *Riesgo estimado* (*Razón de ventajas* que simbolizaremos como RM) será igual a 1 (no hay asociación significativa), y H_1 es que el valor del *Riesgo estimado* será diferente de 1 (asociación significativa).

Tabla 28

Hipótesis estadísticas relativas a la estimación de coeficientes de regresión logística del modelo a construir

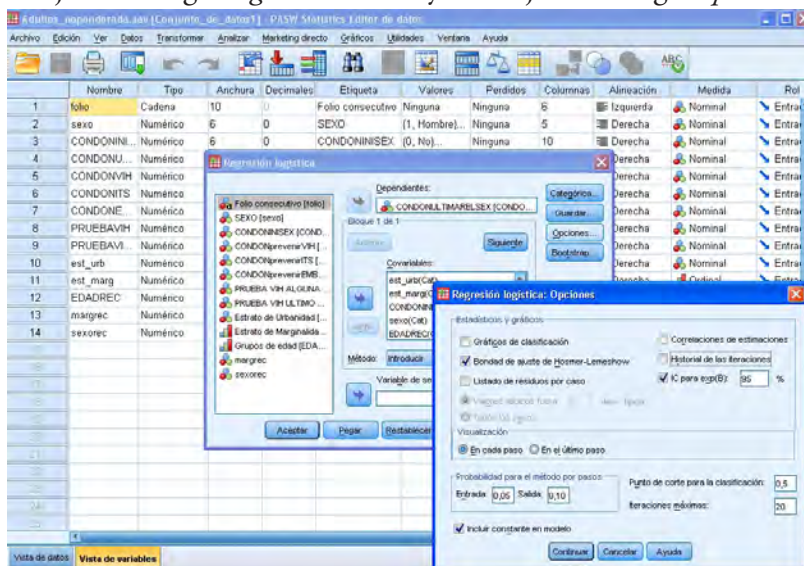
Variables independientes	Variable dependiente	Hipótesis estadísticas complementarias
$\bar{x}_1$. Estrato de marginalidad β_1	PY. Uso del condón en la última relación sexual	Efecto lineal: H_0 . $\beta_1 = \beta_2 = \beta_k = 0$ H_1 . Algún $\beta_{1,5} \neq 0$
$\bar{x}_2$. Estrato de urbanidad β_2		
$\bar{x}_3$. Uso del condón en la primera relación sexual β_3		Fuerza del efecto: H_0 . Razón de ventajas (RM) = 1 H_1 . Razón de ventajas (RM) $\neq 1$
$\bar{x}_4$. Sexo β_4		
$\bar{x}_5$. Grupos de edad β_5		

Para contrastar la hipótesis del efecto lineal se calcularán los coeficientes de regresión (β) y la prueba de significancia p del estadístico *Wald*, mientras que el contraste de la fuerza del efecto se realizará mediante la estimación de las *Razones de ventajas* (RM) y sus correspondientes intervalos de confianza de 95% (IC de 95%). Tales cálculos se pueden efectuar en el SPSS de la siguiente manera: en la barra del menú del SPSS se sigue la secuencia *Analizar/Regresión/Logística binaria*. Se despliega una caja de diálogo que contiene la lista de las variables disponibles. Se elige,

primero, la variable dependiente y se traslada al espacio *Dependientes*. Luego se eligen todas las variables independientes y se mueven al espacio *Covariables*. En la opción “Método” se elige *Introducir*. Se ingresa a *Opciones* donde se seleccionan *IC para exp(B)* 95 e *Incluir constante para el modelo*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 28).

Figura 28

Caja de diálogo *Logística binaria* y sub caja de diálogo *Opciones*



En el visor de resultados aparecerán todas las tablas del análisis de regresión logística. Para esta segunda fase, solamente debemos concentrarnos en la tabla *Variables en la ecuación*, la cual despliega para cada variable independiente, los siguientes valores: coeficiente de regresión logística (β), estadístico *Wald* y su significancia estadística (Sig.), razón de ventajas RM [Exp(b)] y límites inferior y superior de su intervalo de confianza del 95% [IC de 95% para EXP(B)]. Al respecto véase la tabla 29.

Tabla 29
Tabla de codificaciones de variables categóricas y tabla de la ecuación de regresión logística, desplegadas en el visor de resultados

	B	E.T.	Wald	gl	Sig.	Exp(B)	IC de 95% para EXP(B)	
							Inferior	Superior
Paso 1 ^a			24,743	2	0.000			
est_rural								
est_urbano(1)	0.092	0.049	3,523	1	0.061	1.096	0.996	1.207
est_metropolit(2)	0.205	0.043	23,091	1	0.000	1.227	1.129	1.334
est_marg_baja(1)	0.188	0.039	23,450	1	0.000	1.207	1.118	1.302
CONDONINISEX	1.842	0.033	3096,975	1	0.000	6.307	5.911	6.729
Sexo_hombres(1)	0.582	0.032	330,559	1	0.000	1.790	1.681	1.906
EDAD_46>			253,619	2	0.000			
EDAD_20-29(1)	0.703	0.063	125,879	1	0.000	2.021	1.787	2.285
EDAD30-45(2)	0.213	0.060	12,367	1	0.000	1.237	1.099	1.392
Constante	-2,883	0.066	1888,012	1	0.000	0.056		
^a . Variable(s) introducida(s) en el paso 1: est_urb, est_marg, CONDONINISEX, sexo, EDADREC.								

En la lectura de la tabla, el primer aspecto que se rescata es el de los coeficientes de regresión (β) de las variables: a) todos tienen signo positivo, lo cual las postula como predictores de una mayor probabilidad del uso del condón en la última relación sexual; y b) los tres coeficientes más elevados corresponden a: “uso del condón en la primera relación sexual” (1.842), edad de 20 a 29 años (0.703) y hombres (0.582).

El segundo aspecto a rescatar son los valores de significancia estadística p (Sig.) de la prueba Wald. En casi todos los casos, excepto en la categoría “urbano”, las significancias estadísticas fueron de $p < 0.01$ y permiten aceptar la hipótesis de la relación lineal.

Finalmente, también se rescatan los valores de las Razones de ventajas ajustadas (RM) y de sus IC, del 95%, que contrastarán la hipótesis de la fuerza del efecto de las variables independientes sobre la dependiente. En casi todos los casos, excepto en la categoría urbano, las RM son mayores a 1, lo cual lleva a interpretar que las variables *estrato metropolitano*, *baja marginación*, *uso del condón en la primera relación sexual*, *sexo masculino* y *edades menores* son predictores independientes significativos de una mayor probabilidad de “usar el condón en la última relación sexual”. En estas variables, el rango de los IC de 95% es mayor a 1 y otorga a las variables un sentido de riesgo o favorecedor de una mayor probabilidad de uso del condón.

En correspondencia con los valores de los coeficientes de regresión (β), los valores más altos de RM los tuvieron las variables “uso del condón en la primera relación sexual” (6.307), edad de 20 a 29 años (2.021) y hombres (1.790).

Un factor que se evaluó en esta etapa fue la presencia de efectos de confusión e interacción entre las variables independientes. Se encontró que no existían, como se aprecia en el recuadro 2:

Recuadro 2

Efectos de confusión e interacción en el modelo

En esta segunda fase, la inclusión de distintas variables en el modelo permitió ajustar el efecto de cada variable por el de las otras y no hubo mayores modificaciones en los valores de las *Razones de ventajas* (RM). Esto quiere decir que no hubo *efectos de confusión* entre las variables independientes, en el sentido propuesto en la tabla 24 (página 156).

Otro aspecto que se consideró fue la posibilidad de incluir en el modelo al *efecto de interacción* entre algunas variables independientes. Pensamos en la posibilidad de que las variables estructurales generarían dicho efecto por tener en común una raíz socioeconómica. Por ello se hizo un análisis preliminar de interacción entre las variables “estrato de urbanidad” y “estrato de marginalidad” para decidir si se incluía en el modelo.

Este análisis se realizó en el spss de la siguiente manera: en la barra del menú del spss se sigue la secuencia *Analizar/Regresión/Logística binaria*. Se despliega una caja de diálogo que contiene la lista de las variables disponibles. Se elige, primero, la variable dependiente y se traslada al espacio *Dependientes*. Luego se eligen las variables independientes de interés y se mueven al espacio *Covariables*. Adicionalmente, se selecciona la primera variable de interés en el espacio de variables disponibles y se oprime la tecla *Ctrl*. Sin soltar la tecla *Ctrl* se selecciona la segunda variable de interés. Se activará el botón de interacción  (figura D) y habrá que cliquearlo para que las dos variables seleccionadas ingresen juntas al espacio *Covariables*. En la opción “Método” se elige *Introducir*. Se ingresa a *Opciones* donde se seleccionan *IC para EXP(B)* 95 e *Incluir constante para el modelo*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura E).

Continúa en la página 168

Viene de la página 167

Figura D
Caja de diálogo *Logística binaria* donde se eligen las variables que se incluyen en el análisis de *interacción*

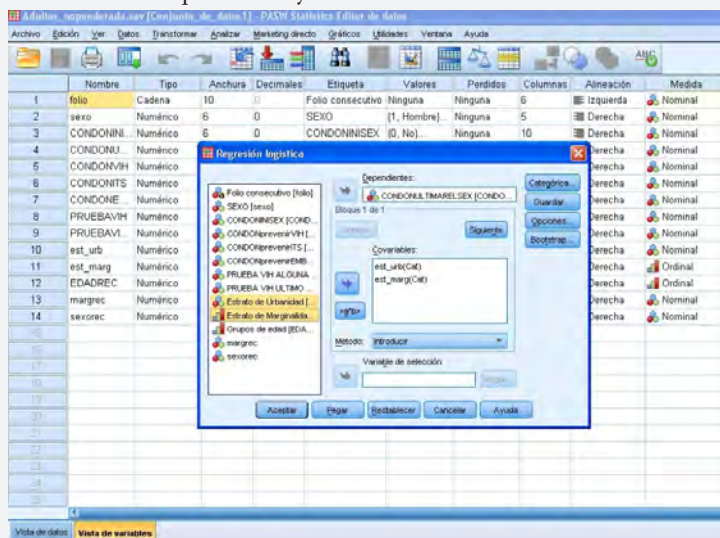
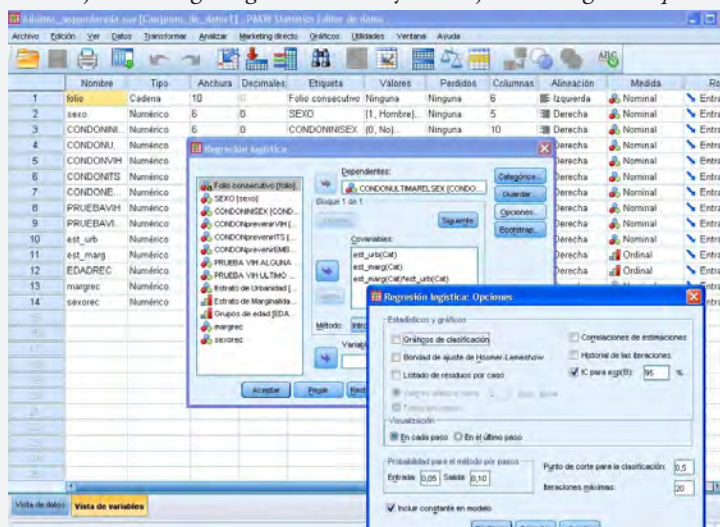


Figura E
Caja de diálogo *Logística binaria* y sub caja de diálogo de *Opciones*



Continúa en la página 169

Viene de la página 168

Tal procedimiento generó las distintas tablas del análisis de regresión que se desplegaron en el visor de resultados del SPSS. Nos interesó solamente leer e interpretar la *tabla de variables en la ecuación* (véase la tabla A de este recuadro) donde se expresaron los valores calculados de las variables estructurales individuales y en interacción (multiplicadas). Notamos que los valores de la significancia p (*Sig.*) de la prueba *Wald* para las interacciones eran mayores a 0.05 y ello nos permitió aceptar la hipótesis nula de la interacción (es decir que no había efecto de interacción entre las variables estructurales):

Tabla A
Tabla de variables en la ecuación con efectos
de interacción desplegada en el visor de resultados

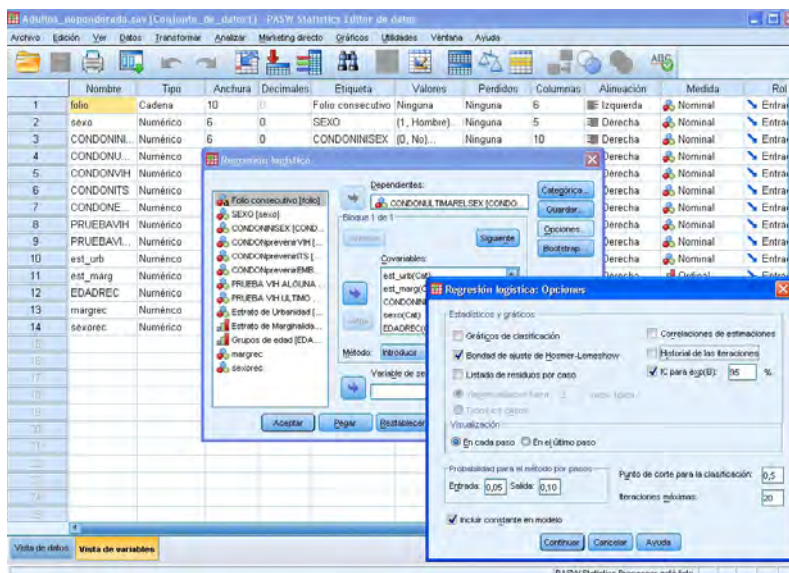
		B	E.T.	Wald	gl	Sig.	Exp(B)	IC de 95% para EXP(B)	
								Inferior	Superior
Paso 1a	est_urb			49.018	2	0.000			
	est_urb(1)	0.020	0.074	.074	1	0.786	1.020	0.883	1.178
	est_urb(2)	0.406	0.060	46.389	1	0.000	1.501	1.335	1.687
	est_marg(1)	0.302	0.059	26.343	1	0.000	1.352	1.205	1.517
	est_marg * est_urb			4.616	2	0.099			
	est_marg(1) by est_urb(1)	0.148	0.095	2.446	1	0.118	1.159	0.963	1.395
	est_marg(1) by est_urb(2)	-0.046	0.079	0.344	1	0.558	0.955	0.817	1.115
	Constante	-1.667	0.034	2382.080	1	0.000	0.187		

En consecuencia, al no haber efectos de interacción, no se consideró su inclusión como variable independiente en el modelo general.

Finalmente, el tercer paso en la construcción de nuestro modelo de regresión logística binaria consistirá en evaluar su significancia estadística al estimar la proporción de variabilidad de su explicación y de su bondad de ajuste como modelo predictor. El primer cálculo se hará mediante el rango de determinación R^2 de Cox y Snell, y R^2 de Nagelkerke; y el segundo con la prueba Hosmer-Lemeshow.

Para generar tal evaluación en el SPSS se realiza lo siguiente: en la barra del menú del SPSS se sigue la secuencia: *Analizar/Regresión/Logística binaria*. Se despliega una caja de diálogo que contiene la lista de las variables disponibles. Se elige, primero, la variable dependiente y se traslada al espacio *Dependientes*. Luego se eligen las variables independientes de interés y se mueven al espacio *Covariables*. Se ingresa a *Opciones* donde se selecciona *Bondad de ajuste de Hosmer-Lemeshow*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 29).

Figura 29
Caja de diálogo *Logística binaria* y sub caja de diálogo de *Opciones*



En el visor de resultados del spss se despliegan todas las tablas del análisis de regresión logística, pero en esta fase solamente nos interesarán las de *resumen del modelo* y la prueba de Hosmer y Lemeshow (tablas 30 y 31).

Tabla 30
Resumen del modelo, desplegado en el visor de resultados

Paso	-2 log de la verosimilitud	R cuadrado de Cox y Snell	R cuadrado de Nagelkerke
1	25151,910 ^a	0.173	0.266

^a. La estimación ha finalizado en el número de iteración³ 5 porque las estimaciones de los parámetros han cambiado en menos de ,001.

Tabla 31
Prueba de Hosmer y Lemeshow, desplegada en el visor de resultados

Paso	Chi-cuadrado	gl	Sig.
1	11,670	8	0.167

La primera tabla (30) expresa los valores que posee el rango de la proporción de la variabilidad respecto a la variable dependiente explicada por el modelo. El límite inferior del rango (R^2 de Cox y Snell) es de 0.173 y el superior (R^2 de Nagelkerke) de 0.266, lo cual significa que el modelo lograría explicar hasta 26.6% de la variabilidad sobre el “uso del condón en la última relación sexual”. En forma complementaria, se asumiría que la proporción restante podría ser explicada por la influencia de otras variables no consideradas en el estudio.

Por su parte, la tabla (31) de la prueba de Hosmer y Lemeshow incluye los valores de la prueba Chi-cuadrado (11.670), grados de libertad (gl = 8) y significancia estadística P (Sig. = 0.167). Su interpretación sería que el modelo tiene una bondad de ajuste estadísticamente significativa, debido a que tiene un

³ Iteración es el número de veces que se repite un proceso.

valor de Chi-cuadrado relativamente pequeño (menor a 12) y un valor de significancia alto ($p > 0.05$). En este caso la hipótesis nula (H_0) es que el valor de la significancia sea $p < 0.05$ y la alterna (H_1) que el valor sea $p > 0.05$ (como indicador de ajuste adecuado del modelo).

Síntesis

Nuestro proceso de análisis de las variables nominales y ordinales siguió etapas interrelacionadas que nos permitieron trabajar los datos desde su escrutinio más simple hasta una construcción más compleja. Nos permitió explorar el contenido de las variables individuales, contrastar la relación de pares de variables mediante tablas de contingencia y pruebas de significancia, y, sobre esa base, generar un modelo de regresión logística binario.

La tabla 32 sintetiza el modelo construido y detalles descriptivos que lo contextualizan, de manera que se pueda realizar la comunicación de los resultados ante la comunidad científica con cierta precisión y rigor estadístico. Como se aprecia, la tabla es el resultado de todo el proceso de análisis desarrollado y su síntesis destaca los siguientes elementos: 1) el porcentaje de uso del condón en la última relación sexual; 2) *las Razones de ventajas* (RM) univariadas (no ajustadas por las otras variables) y las multivariadas (ajustadas por las otras variables) y sus IC, del 95%; 3) los valores de significancia p de la prueba Wald; y 4) los valores de la prueba Hosmer y Lemeshow de bondad de ajuste del modelo.

Tabla 32

Modelo de regresión logística de factores asociados al uso del condón en la última relación sexual, en población adulta**

Factores	Uso del condón en última relación sexual		
	Sí (%)	RM univariado (IC 95%)	RM multivariado (IC 95%)
<i>Marginación</i>			
Baja (n = 18,493)	24.8	1.6 (1.49-1.68)*	1.2 (1.12-1.30)*
Alta (n = 10,604)	17.2	Referente	Referente
<i>Urbanidad</i>			
Metropolitana (n = 13,735)	25.8	1.7 (1.58-1.80)*	1.2 (1.13-1.33)*
Urbana (n = 6,247)	21.1	1.3 (1.19-1.40)*	1.1 (0.99-1.21)
Rural (9,115)	17.1	Referente	Referente
<i>Uso del condón en primera relación sexual</i>			
Sí (n = 8,222)	50.3	8.3 (7.78-8.80)*	6.3 (5.91-6.73)*
No (n = 20,869)	10.9	Referente	Referente
<i>Sexo</i>			
Hombres (n = 12,208)	30.5	2.2 (2.18-2.45)*	1.8 (1.68-1.91)*
Mujeres (16,889)	15.9	Referente	Referente
<i>Grupos de edad</i>			
20-29 (n = 8,774)	34.8	3.9 (3.46-4.34)*	2.0 (1.79-2.28)*
30-45 (16,916)	17.8	1.6 (1.44-1.80)*	1.2 (1.10-1.39)*
46 > (3,407)	11.8	Referente	Referente
<i>Constante</i>			0.06

Bondad de ajuste del modelo Hosmer & Lemeshow: Chi-cuadrado = 11.7, grados de libertad = 8, P = 0.167

*Prueba Wald: P < 0.01

RM = Razón de ventajas (Razón de Momios u Odds Ratio)

**México, 2012. N = 29,097.

Fuente: Elaboración propia.

La lectura de la anterior tabla (32) nos permite ver que todas las variables son predictores significativos de la variable dependiente como factores de riesgo, y que algunas tienen un

mayor efecto independiente que otras. La prueba Hosmer y Lemeshow indica que el modelo construido tiene una bondad de ajuste adecuada.

El procedimiento seguido para lograr la construcción del modelo de regresión logística nos permitió compenetrarnos a fondo con los datos y aprovechar parte de su riqueza en función de las hipótesis que nos interesaron. Sin embargo, el análisis no agotó toda la potencialidad de los datos y se podría profundizar en la búsqueda de nuevos hallazgos. Ésta es la esencia del enfoque de la minería de datos.

En ese sentido, se piensa que las relaciones halladas en el análisis de regresión logística se podrían trabajar más a fondo para identificar relaciones que pertenecen sólo a subgrupos específicos y que facilitarían la construcción de reglas para la predicción de la variable dependiente. Ello nos conduciría a desarrollar un análisis de Árboles de decisión que en el paquete SPSS se encuentra en *Analizar/Clasificar/Árbol*. Nosotros no realizaremos este nuevo análisis y nos detendremos aquí porque nuestro interés era solamente ilustrar el proceso de análisis de variables nominales y ordinales hasta este nivel.

Análisis de variables de intervalo/razón

En esta segunda parte, se ejemplificará el análisis de variables de escalas de intervalo/razón con información tomada de diferentes bases de datos elaboradas por algunas instituciones internacionales para evaluar el desarrollo humano y el bienestar social en el mundo.

Cabe mencionar que la medición que realizan dichas instituciones surgió a mediados del siglo XX bajo el nombre “movimiento de indicadores sociales”, con el propósito de generar información que pudiera influir en el diseño de mejores políticas públicas en los países de menor desarrollo. En la actualidad, este movimiento ha ganado protagonismo como fuente que impulsa la agenda de promoción de políticas en los distintos organismos

internacionales, mediante un amplio repertorio de indicadores sociales aplicados a múltiples áreas del desarrollo humano (Neuman, 1997: 168-169).

Desde el enfoque de la epidemiología social, se considerarán algunos indicadores que podrían constituirse en determinantes estructurales de la prevalencia del VIH en personas de 15 a 49 años y de la tasa de incidencia de tuberculosis, en el mundo. Los indicadores elegidos son los siguientes: el nivel de ingreso de los países reflejado en su Producto Interno Bruto (PIB), el índice GINI de desigualdad social, el índice compuesto de fragilidad de los países, el índice de felicidad planetaria y los coeficientes de dispersión étnica y lingüística.

La tabla 33 describe las variables consideradas, sus definiciones, sus códigos y etiquetas, y las fuentes institucionales de los datos.

Tabla 33
Variables por escalas de medición, definiciones, y valores y etiquetas

Variables y escalas	Definiciones	Valores y etiquetas
<i>Ordinal</i>		
PIB de países en 2010 ^a	Niveles de ingreso de los países	1. Mediano-bajo, 2. Mediano, 3. Mediano-alto, 4. Alto
<i>Intervalo/razón</i>		
Índice de desigualdad social ^a	Índice de la inequidad en la distribución del ingreso	Índice GINI de 0 a 100
Índice de fragilidad de los países ^b	Grado de fragilidad económica, política, desarrollo humano, demográfica, ambiental y seguridad	Índice compuesto de 0 a 10
Índice de felicidad planetaria ^c	Percepción de bienestar, esperanza de vida y ecología	Porcentaje
Coeficiente de dispersión étnica ^d	Grado de heterogeneidad étnica	Coeficiente de 0 a 1
Coeficiente de dispersión de lenguas ^e	Grado de heterogeneidad de uso de lenguas	Coeficiente de 0 a 1
Prevalencia de VIH ^e	% de personas de 15 a 49 años infectados con VIH	Porcentaje
Tasa de incidencia de tuberculosos ^e	Cantidad estimada de nuevos casos de TB pulmonar, frotis positivo y extrapulmonar	Casos x 100,000 personas

^aFuente: Banco Mundial, 2010.

^bFuente: Carleton University Ottawa, 2010.

^cFuente: Centre for Well-being at New Economics Foundation, 2010.

^dFuente: Alesina *et al.* Fractionalization. Journal of Economic Growth, 8, 2003: 155-194.

^eFuente: Banco Mundial, 2011.

Los indicadores tomados de bases de datos específicas de las fuentes consultadas, fueron trasladados a una base de datos común que es la que nos servirá para ejemplificar el proceso de análisis. Esta base de datos se generó en el paquete SPSS 19 mediante la construcción de dos matrices: 1) una matriz de datos donde se capturaron los valores numéricos de las variables (figura 30), y 2) una matriz de definiciones técnicas de las variables (figura 31).

Figura 30

Matriz de datos con los valores numéricos de las variables*

	PAIS	REGION	HPIINDEX	R1_FRA_GILTY	TB2011	VIH2011	FRACETNIA	FRACLEN	FRACRE	GUA	LIGION	GINI	PIB_2010REC	PID2010_DIC
1	Afghanistan	0	36.8	6.69	189	.1	.77	.61	27				1	1
2	Albania	0	54.1	4.70	13		.22	.04	47	34.5			1	1
3	Algeria	0	52.2	5.37	90		.34	.44	01	35.3			2	1
4	Andorra	0		2.94	6		.71	.68	23					
5	Angola	2	33.2	6.35	310	2.1	.79	.79	63				2	1
6	Antigua and Barbuda	0		4.45	7		.16	.11	68				3	2
7	Argentina	1	54.1	4.09	26	4	.26	.06	22	45.8			3	2
8	Armenia	0	46.0	4.99	55	2	.13	.13	46	30.9			1	1
9	Aruba	0		4.35	6			.39	41					
10	Australia	0	42.0	3.16	6	2	.09	.33	82	30.5			4	2
11	Austria	3	47.1	3.09	4	4	.11	.15	41	26.0			4	2
12	Azerbaijan	0	40.9	5.49	113	1	.20	.21	49	33.7			2	1
13	Bahamas	0		3.88	13	2.8	.42	.19	68				3	2
14	Bahrain	0	26.6	4.58	18		.50	.43	55				3	2
15	Bangladesh	0	56.3	5.79	225	1	.05	.09	21	33.2			1	1
16	Barbados	0		3.35		9	.14	.09	69				3	2
17	Belarus	0	37.4	4.50	70	4	.32	.47	61	27.2			2	1
18	Belgium	3	37.1	3.38	8	3	.56	.54	21	28.0			4	2
19	Belize	1	59.3	4.66	40	2.3	.70	.63	58				2	1
20	Benin	0	31.1	5.68	70	1.2	.79	.79	55	36.5			1	1
21	Bhutan	0		5.32	192	3	.61	.61	38				1	1
22	Bolivia	1	43.6	5.13	131	3	.74	.22	21	58.2			1	1

*Se visualiza en el Editor de datos del SPSS, en la pestaña inferior *Vista de datos*.

Figura 31
Matriz de datos con las definiciones técnicas de las variables*

	Nombre	Tipo	Anchura	Decim.	Etiqueta	Valores	Perd.	Column.	Alineación	Medida	Rol
1	PAIS	Cadena	35	0		Ninguna	Ninguna	15	Izqui...	Nominal	Entrada
2	REGION	Numérico	8	0	REGIONES DEL MUNDO	(0, RESTO ...	Ninguna	5	Dere...	Nominal	Entrada
3	HPINDEX	Numérico	7	1	IND_FELICIDAD	Ninguna	Ninguna	7	Dere...	Escala	Entrada
4	R1_FRAGILITY	Numérico	18	2	IND_FRAGILIDAD	Ninguna	Ninguna	5	Dere...	Escala	Entrada
5	TB2011	Numérico	8	0	INCIDENCIA_TB	Ninguna	Ninguna	5	Dere...	Escala	Entrada
6	VIH2011	Numérico	8	1	PREVALENCIA_VIH	Ninguna	Ninguna	5	Dere...	Escala	Entrada
7	FRACETNIA	Numérico	8	2	DISPER_ETNICA	Ninguna	Ninguna	9	Dere...	Escala	Entrada
8	FRACLENGUA	Numérico	8	2	DISPER_LENGUAS	Ninguna	Ninguna	6	Dere...	Escala	Entrada
9	FRACRELIGION	Numérico	8	2	DISPER_RELIGION	Ninguna	Ninguna	5	Dere...	Escala	Entrada
10	GINI	Numérico	2	1	GINI_DESIGUALDAD	Ninguna	Ninguna	5	Dere...	Escala	Entrada
11	PIB_2010REC	Numérico	8	0	PIB_2010_4NIVELES	(1, Mediano...	Ninguna	9	Dere...	Ordinal	Entrada
12	PID2010_DIC	Numérico	8	0	PIB_2010_2NIVELES	(1, Mediano...	Ninguna	11	Dere...	Ordinal	Entrada
13											
14											
15											
16											
17											
18											
19											
20											
21											
22											
23											
24											
25											

*Se visualiza en el Editor de datos del SPSS, en la pestaña inferior *Vista de datos*.

Análisis univariado: medidas de tendencia central y de dispersión

Luego de crear las matrices de datos, un primer paso de análisis consiste en explorar el contenido y distribución de las variables. Debido a que se trata de variables de escala *intervalo/razón*, se pueden calcular sus medidas de tendencia central y de dispersión (véase tabla 7). Las definiciones de tales medidas se presentan en la tabla 34.

Tabla 34
Definiciones de las medidas de tendencia central y de dispersión

<i>Medidas de tendencia central: ¿dónde tienden a concentrarse los valores de una distribución?</i>
• Media aritmética ($\bar{x}$) es el promedio aritmético de un conjunto de observaciones en una distribución.
• Mediana (Md) es el valor medio en un conjunto de valores ordenados dentro de una distribución.
• Moda (Mo) es la observación que ocurre con más frecuencia en una distribución.
<i>Medidas de dispersión: ¿cómo y cuánto varían los valores de una distribución?</i>
• Rango: cuantifica la diferencia entre los valores mínimo y máximo de una distribución.
• Mínimo: el valor mínimo observado en una distribución.
• Máximo: el valor máximo observado en una distribución.
• Varianza: cuantifica el grado de variación entre el conjunto de valores de una distribución.
• Desviación típica (desviación estándar): es la raíz cuadrada del valor calculado de la varianza.

Fuente: Hardy, M. (2004: 35-64).

En el caso de las variables de ejemplo, el cálculo de dichas medidas se puede realizar en el SPSS de la siguiente manera: en la barra del menú del SPSS se sigue la secuencia *Analizar/Estadísticos descriptivos/Frecuencias* (figura 32). Se despliega una caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables estructurales y las dependientes y se les traslada al espacio *Variables*. Luego se ingresa a *Estadísticos* donde se marcan, como opciones de tendencia central, la media, la mediana y la moda; y, como opciones de dispersión, desviación típica, varianza, rango máximo y mínimo. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 33).

Figura 32
Secuencia de opciones a seleccionar
en el menú del spss para ingresar a *Frecuencias*

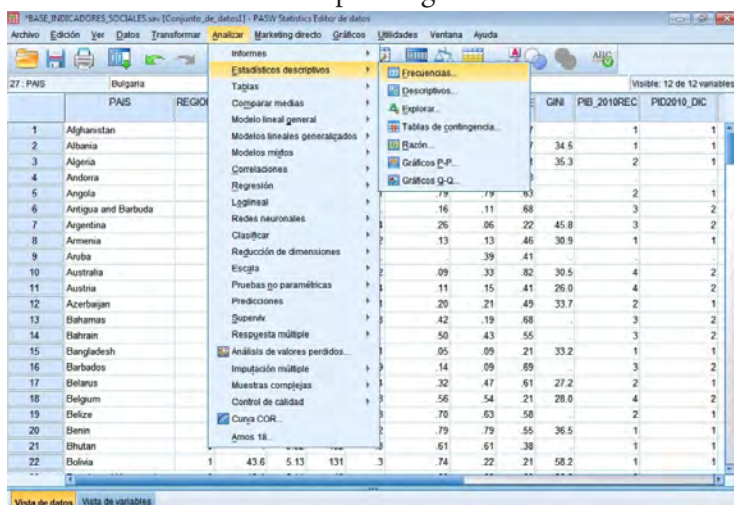
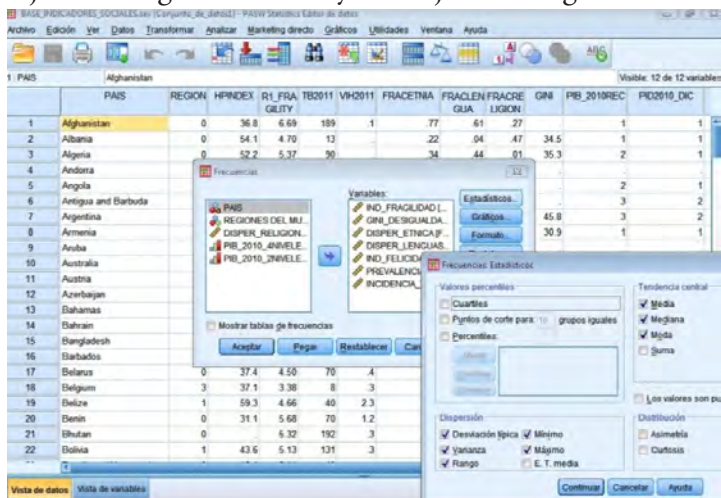


Figura 33
Caja de diálogo *Frecuencias* y sub caja de diálogo *Estadísticos*



Tal procedimiento genera una tabla de estadísticos que se despliega en el visor de resultados (tabla 35). La tabla describe los valores de tendencia central y de dispersión de las variables consideradas en el análisis.

Tabla 35
Estadísticos de tendencia central
y dispersión desplegados en el visor de resultados

	IND_ FRAGILIDAD	GINI_ DESIGUALDAD	DISPER_ ÉTICA	DISPER_ LENGUAS	IND_ FELICIDAD	PREVALENCIA_ VIH	INCIDENCIA_ TB
N	Válidos	132	185	180	151	146	190
	Perdidos	61	8	13	42	47	3
Media	4.8419	40.255	0.4397	0.3906	42.244	1.907	129.23
Mediana	4.9650	39.100	0.4373	0.3631	42.000	0.400	60.00
Moda	4.57 ^a	26.0 ^a	.26 ^a	.02 ^a	46.0	0.1	21
Desviación típica	1.04752	9.9564	0.25859	0.28117	9.1166	4.2934	184.327
Varianza	1.097	99.130	0.067	0.079	83.113	18.433	33976.371
Rango	4.23	47.7	0.93	0.92	41.4	25.9	1316
Mínimo	2.56	23.0	0.00	0.00	22.6	0.1	1
Máximo	6.79	70.7	0.93	0.92	64.0	26.0	1317

^a.Existen varias modas. Se mostrará el menor de los valores.

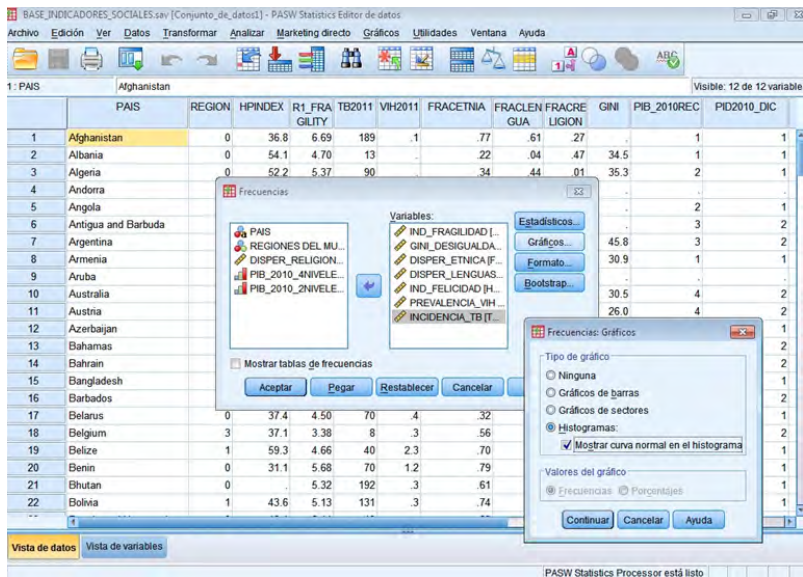
La primera interpretación que se puede hacer de los estadísticos calculados tiene que ver con la magnitud de sus valores. Por ejemplo, la media de la variable *prevalencia de VIH* es de 1.9, mientras que la mediana es mucho menor (0.4). Esto se debe a que hay una diferencia de 25.9% (rango) entre la prevalencia mínima (0.1) y la máxima (26.0) que eleva el valor de la media, por ser ésta una medida sensible a ser arrastrada hacia los valores mayores. Esta distorsión no existe en la mediana (0.1) que expresa el valor medio exacto de la distribución por no tener la sensibilidad de la media. A la vez, el valor de la moda (0.1) expresa cuál es la prevalencia más frecuente. Se podría hacer una lectura similar con las otras variables.

La segunda interpretación se orientará a iniciar la evaluación de la forma que presenta la distribución de las variables. La primera pregunta a responder será: ¿hasta qué punto es simétrica la distribución de los valores de la media, la mediana y la moda? Un indicio de simetría en la distribución es que los tres valores sean iguales o muy semejantes, condición que es más o menos cumplida sólo por las variables *índice de fragilidad* (4.84, 4.96 y 4.57) e *índice de felicidad* (42.24, 42.00 y 46.00). Al respecto véase la tabla anterior (35).

Un supuesto de la distribución simétrica de las variables es que pueden representarse de manera precisa por una curva normal con forma de campana (campana de Gauss). Para ver si esto se cumple, se podría generar un gráfico de histograma con curva normal para las variables *índice de fragilidad* e *índice de felicidad*. Ello se realiza en el SPSS de la siguiente forma: en la barra del menú del SPSS se sigue la secuencia: *Analizar/Frecuencias*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables requeridas y se trasladan al espacio *Variables*. Luego se ingresa a *Gráficos* donde se marcan *Histograma* y *Mostrar curva normal en el histograma*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 34).

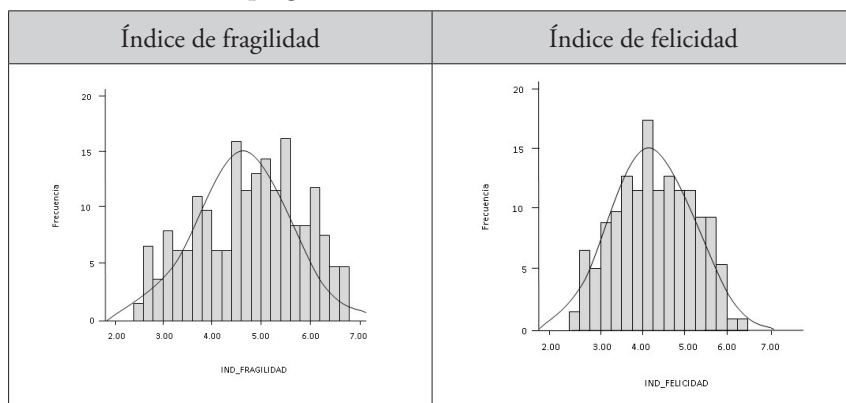
Figura 34

Caja de diálogo del comando de cálculo *Frecuencias* y sub caja de diálogo *Gráficos* donde se generan histogramas con curva normal



Dicho procedimiento genera los histogramas de las variables en el visor de resultados (figura 35). En éstos, el trazo de la curva normal permite apreciar que la distribución de las variables tiene una forma muy cercana a la de una campana (sobre todo en la variable índice de felicidad), lo cual refuerza la hipótesis de la distribución normal en las variables consideradas.

Figura 35
Gráficos de histograma con curva normal
desplegados en el visor de resultados



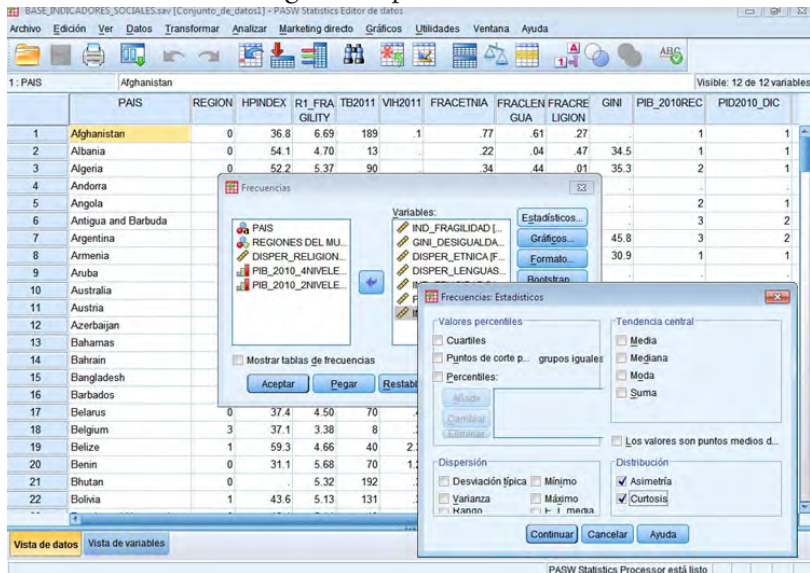
Fuente: Elaboración propia.

Una segunda pregunta a responder en la evaluación de la forma es la siguiente: ¿hasta qué punto los valores de la distribución se desvían de la curva normal? La pregunta se responderá mediante el cálculo de los estadísticos de la asimetría (que evalúa si los valores de una variable tienden a concentrarse en el lado derecho o izquierdo de la curva normal) y de la curtosis (que evalúa si una variable tiene muchos o pocos casos en el medio de la curva normal).

Estos cálculos se realizan en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia: *Analizar/Frecuencias*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables de interés y se trasladan al espacio *Variables*. Luego se ingresa a *Estadísticos* donde se marcan las dos opciones de distribución: *Asimetría* y *Curtosis*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 36).

Figura 36

Caja de diálogo del comando *Frecuencias* y sub caja *Estadísticos* donde se eligen las opciones en *Distribución*



En el visor de resultados se genera una tabla con los estadísticos de asimetría y curtosis para todas las variables consideradas. En términos generales, cuando los valores de ambos estadísticos son de ± 0.5 tienden a expresar la presencia de una distribución simétrica o normal, esto más o menos acontece con las variables índice de fragilidad, índice de felicidad y *Gini_desigualdad* (tabla 36).

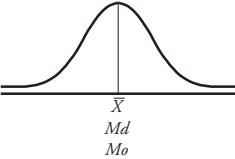
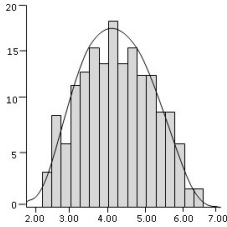
Tabla 36
Estadísticos de *asimetría* y *curtosis*
desplegados en el visor de resultados

N	IND_ FRAGILIDAD	GINI_ DESIGUALDAD	DISPER_ ÉTNICA	DISPER_ LENGUAS	IND_ FELICIDAD	PREVALENCIA_ VIH	INCIDENCIA_ TB
Válidos	192	132	185	180	151	146	190
	Perdidos	61	8	13	42	47	3
Asimetría	-0.306	0.581	-0.013	0.256	0.063	3.851	2.912
Error típico de asimetría	0.175	0.211	0.179	0.181	0.197	0.201	0.176
Curtosis	-0.733	-0.010	-1.257	-1.244	-0.673	15.742	11.756
Error típico de curtosis	0.349	0.419	0.355	0.360	0.392	0.399	0.351

En conjunto, la evaluación de todos los estadísticos de forma nos permite apreciar los tipos de distribución de las variables, en función de su curva normal y de las características de sus valores obtenidos. Así, en nuestras variables de ejemplo, el *índice de felicidad* tiende a una distribución normal con una curva simétrica (platicúrtica), el índice de fragilidad tiende a una distribución normal con una curva asimétrica negativamente sesgada (platicúrtica) y la *prevalencia de VIH* tiene una distribución no normal con una curva asimétrica positivamente sesgada (leptocúrtica). Al respecto véase la tabla 37.

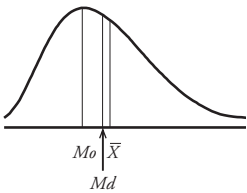
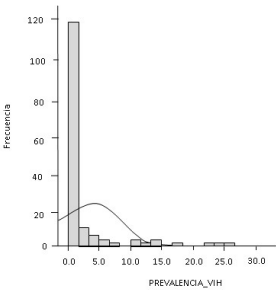
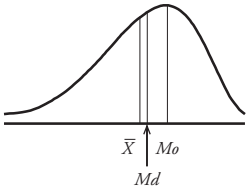
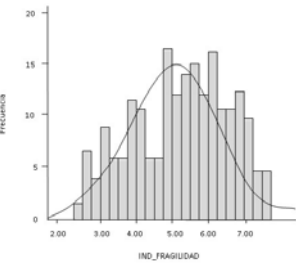
Tabla 37

Distribución de la curva según las características de los estadísticos calculados en tres variables de ejemplo

Tipos de distribución	Características	Variables de ejemplo
<i>Normal o simétrica</i> 	• La media ($\bar{x}$), la mediana (Md) y la moda (Mo) tienden a tener un mismo valor. • La forma de la curva de representación es el de una campana simétrica. • <i>Asimetría</i>: Cero o cercana a cero. • <i>Curtosis</i>: Valor negativo menor que 0. • <i>Forma de la curva</i>: platicúrtica.	ÍNDICE DE FELICIDAD $\bar{x} = 42.2$ Md = 42.0 Mo = 46.0 Asimetría = 0.063 Curtosis = -0.673 

Continúa en la página 188

Viene de la página 187

Tipos de distribución	Características	Variables de ejemplo
<i>Asimétrica positivamente sesgada</i> 	• La media ($\bar{x}$), por su sensibilidad a las magnitudes de los valores, es arrastrada al extremo derecho de la distribución. Tiene un valor mayor a la mediana (Md) y la moda (Mo). • La curva muestra que la mayoría de los valores están a la izquierda de la media ($\bar{x}$). • <i>Asimetría:</i> Valor positivo mayor que 0. • <i>Curtosis:</i> Valor positivo mayor que 0. • <i>Forma de la curva:</i> leptocúrtica.	PREVALENCIA DE VIH $\bar{x} = 1.9$ $Md = 0.4$ $Mo = 0.1$ Asimetría = + 3.8 Curtosis = + 15.742 
<i>Asimétrica negativamente sesgada</i> 	• La media ($\bar{x}$), por su sensibilidad a las magnitudes de los valores, es arrastrada al extremo izquierdo de la distribución. Su valor tiende a ser menor a la mediana (Md) y la moda (Mo). • La curva muestra que la mayoría de los valores están a la derecha de la media ($\bar{x}$). • <i>Asimetría:</i> Valor negativo menor que 0. • <i>Curtosis:</i> Valor negativo menor que 0. • <i>Forma de la curva:</i> platocúrtica.	ÍNDICE DE FRAGILIDAD $\bar{x} = 4.8$ $Md = 5.0$ $Mo = 4.6$ Asimetría = -0.306 Curtosis = -0.733 

Fuente: Elaboración propia.

Si bien la evaluación realizada hasta el momento nos ha permitido explorar si la distribución de las variables se acerca o no a la normalidad, se trató sólo de una aproximación tentativa. Ahora podríamos corroborar la distribución normal de las varia-

bles mediante el cálculo de una prueba específica para contrastar la hipótesis de la normalidad: la prueba Kolmogorov-Smirnov. Esto se realiza en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Explorar* (figura 37). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables requeridas y se trasladan al espacio *Lista de dependientes*. Luego se ingresa a *Gráficos* donde se marca *Gráficos con pruebas de normalidad*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 38).

Figura 37

Secuencia a seguir en el menú del SPSS para ingresar a *Explorar*

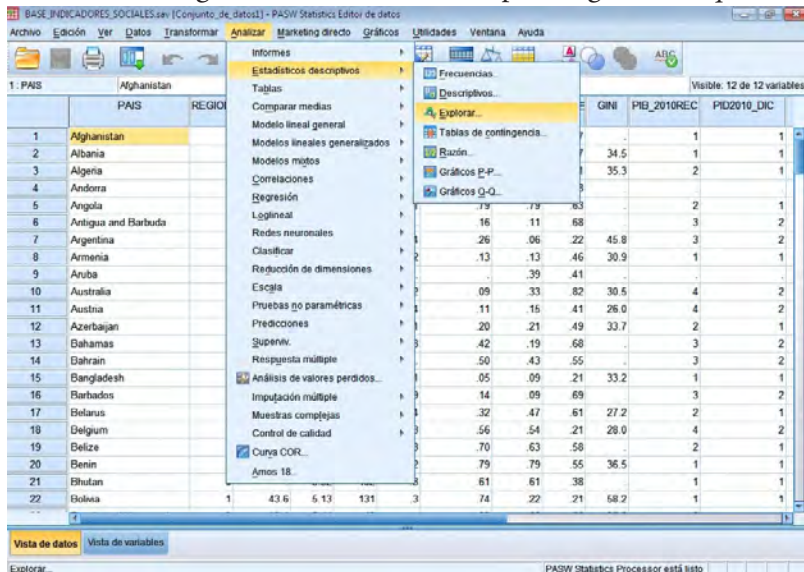
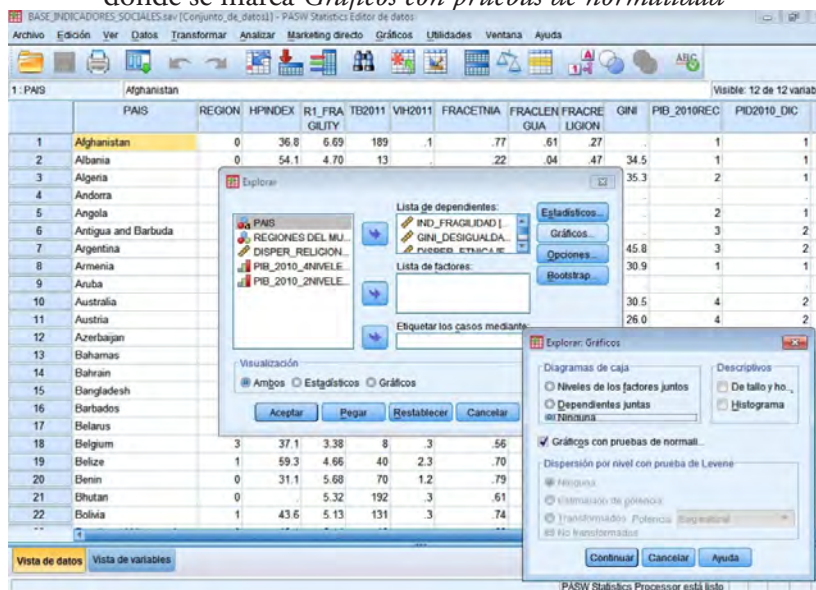


Figura 38

Caja de diálogo del comando *Explorar* y sub caja *Gráficos* donde se marca *Gráficos con pruebas de normalidad*



En el visor de resultados se genera, en primera instancia, una tabla con los estadísticos de la prueba Kolmogorov-Smirnov y sus valores de significancia p (Sig.) para todas las variables consideradas en el análisis. En el contraste que hace esta prueba, la hipótesis nula (H_0) es la distribución normal y la hipótesis alterna (H_1) la distribución no normal, razón por la cual una significancia de $p < 0.05$ nos llevaría a rechazar la normalidad y a aceptar la hipótesis alterna. En tal sentido, las dos variables de nuestro ejemplo que se aceptarían como distribuidas normalmente serían el índice de felicidad ($p = 0.200$) y el índice de desigualdad *de Gini* ($p = 0.169$), mientras que las demás tendrían una distribución no normal por tener valores de $p < 0.05$ (tabla 38). Cabe agregar que la tabla también reporta la prueba Shapiro-Wilk que realiza el mismo contraste de normalidad, pero ajustado a muestras pequeñas (menos de 50 casos). En la interpretación de nuestros datos, solamente consideraremos los

valores de Kolmogorov-Smirnov por contar con una muestra más grande.

Tabla 38
Tabla relativa a *pruebas de normalidad*
desplegada en el visor de resultados

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Estadístico	gl	Sig.	Estadístico	gl	Sig.
IND_FRAGILIDAD	0.103	110	0.006	0.958	1	0.002
GINI_DESIGUALDAD	0.075	110	0.169	0.965	110	0.005
DISPER_ÉTNICA	0.091	110	0.026	0.954	110	0.001
DISPER_LENGUAS	0.126	110	0.000	0.918	110	0.000
IND_FELICIDAD	0.052	110	0.200*	0.992	110	0.779
PREVALENCIA_VIH	0.371	110	0.000	0.445	110	0.000
INCIDENCIA_TB	0.239	110	0.000	0.685	110	0.000

^a. Corrección de la significación de Lilliefors.

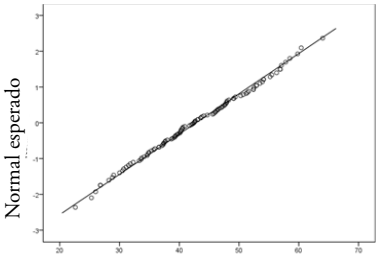
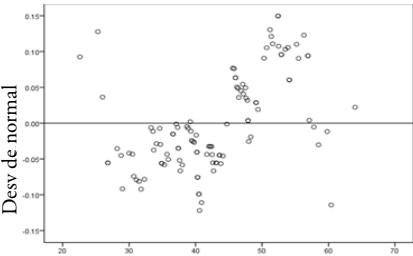
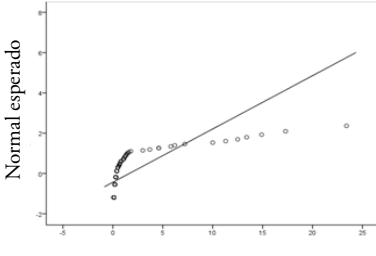
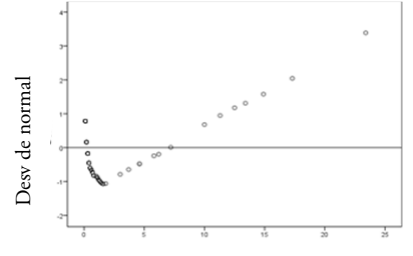
* Éste es un límite inferior de la significación verdadera.

En segunda instancia, se generan en el visor de resultados los gráficos Q-Q normales y Q-Q normales sin tendencia que representan en diagramas complementarios la tendencia a la normalidad de las variables. Es necesario mencionar que en el gráfico Q-Q normal, cada valor observado (eje x) es comparado con la puntuación típica que se esperaría para ese valor en una distribución normal (eje y), por lo cual una distribución normal se representa como una mancha de puntos ajustados a la línea diagonal del diagrama, mientras que en el gráfico Q-Q normal sin tendencia se muestran las distancias verticales existentes entre cada punto y la línea diagonal del gráfico Q-Q normal, motivo por el cual una distribución normal debería expresarse como una mancha de puntos dispersos en torno a la línea horizontal del diagrama.

En la tabla 39 se muestran los gráficos correspondientes a las variables índice de felicidad y *prevalencia de VIH*. La primera tuvo una distribución normal, según la prueba de Kolmogorov-

Smirnov, por lo cual su representación en el gráfico Q-Q normal muestra que sus valores se ajustan a la línea diagonal, mientras que en el gráfico Q-Q sin tendencias sus valores aparecen dispersos en torno a la línea recta horizontal. En el caso de la segunda variable (prevalencia de VIH), el cálculo mostró que no tiene una distribución normal, por lo cual en el gráfico Q-Q normal sus valores no se ajustan a la línea diagonal, y en el gráfico Q-Q sin tendencia sus valores muestran una forma no aleatoria (o sin dispersión) en torno a la línea recta horizontal.

Tabla 39
Gráficos de normalidad Q-Q normales y Q-Q normales sin tendencias, desplegados en el visor de resultados

Gráficos Q-Q normales	Gráficos Q-Q normales sin tendencias
Variable índice de felicidad	
 Normal esperado Valor observado	 Desv de normal Valor observado
Variable <i>prevalencia de VIH</i>	
 Normal esperado Valor observado	 Desv de normal Valor observado

Fuente: Elaboración propia.

Como se ha podido ver, el análisis univariado nos permitió describir los valores de tendencia central y dispersión de las variables, y también explorar la forma de su distribución. Esta última tarea es relevante para el desarrollo de la segunda fase de análisis que se presentará a continuación, debido a que la selección de pruebas estadísticas para contrastar hipótesis de diferencias de medias y de relaciones de asociación está condicionada a que las variables tengan o no una distribución normal.

Análisis bivariado: diferencias de medias y relaciones de asociación

Respecto a la segunda fase de análisis, se exploró la presencia de relaciones significativas entre pares de variables con escala de intervalo/razón. La exploración puede realizarse en dos sentidos diferentes: a) la búsqueda de patrones de diferencias de medias, en el afán de contestar a la siguiente pregunta: ¿las diferencias entre las distribuciones de dos variables son estadísticamente significativas?; y b) la búsqueda de patrones de asociación en el intento de responder la pregunta: ¿hay una asociación significativa entre las distribuciones de valores de dos variables? En seguida, tales formas de exploración se desarrollarán por separado.

Búsqueda de patrones de diferencias de medias

En el caso de la exploración de patrones de diferencia, la selección de pruebas varía de acuerdo con el número de muestras (o categorías) de la variable comparativa:

- Si la variable comparativa tiene dos muestras, se elige la prueba *t* de Student cuando la variable independiente tiene una distribución normal; o bien, se opta por la prueba *U de Mann-Whitney* cuando dicha variable no tiene una distribución normal.
- Si la variable tiene tres o más muestras se escoge la prueba *análisis de varianza* (ANOVA) cuando la variable independiente tiene una distribución normal, o se opta

por la prueba *H de Kruskal-Wallis* cuando tal variable no tiene una distribución normal.

A continuación se ejemplificarán por separado los análisis de relaciones de pares de variables, a partir de pruebas de diferencia de medias para variables comparativas que incluyen dos y tres, o más muestras. Para generar tales análisis, se plantearán previamente hipótesis estadísticas concretas para su posterior contrastación con el uso de las pruebas que se seleccionen.

Diferencia de medias para variables comparativas de dos muestras

Para ejemplificar el análisis de diferencias de medias, respecto a una variable comparativa de dos muestras, se seleccionarán algunas variables de nuestro repertorio disponible (véase tabla 33). En nuestro afán de continuar la exploración de las desigualdades basadas en el efecto de variables estructurales, nos será útil elegir como variable comparativa al índice Gini de desigualdad (dividido en dos muestras o categorías: mayor desigualdad con un valor de ≥ 40.3 ; y menor desigualdad con un valor de < 40.3) y como variables independientes a la prevalencia de VIH, la incidencia de Tb y el índice de felicidad.

En la siguiente tabla (40) se proponen las hipótesis estadísticas de diferencias de medias para las variables independientes y se seleccionan las pruebas de significancia adecuadas a su tipo de distribución (normal o no normal), mismo que fue diagnosticado en la fase anterior.

Tabla 40

Hipótesis estadísticas de diferencia de medias
y selección de pruebas de significancia adecuadas
al tipo de distribución que presentan las variables

Variable de comparación de dos muestras	Hipótesis estadísticas	Pruebas de significancia
Índice <i>Gini de desigualdad</i> con dos muestras o categorías: 1 = mayor desigualdad (valor ≥ 40.3); y 0 = menor desigualdad (valor < 40.3)*	H_0 . Diferencias no significativas de medias de índice de felicidad. H_1 . Diferencias significativas de medias de índice de felicidad.	t de Student (paramétrica), debido a que la variable índice de felicidad tiene distribución normal.
	H_0 . Diferencias no significativas de medias de <i>prevalencia VIH</i> . H_1 . Diferencias significativas de medias de <i>prevalencia VIH</i> .	U de Mann-Whitney (no paramétrica), debido a que la variable <i>prevalencia de VIH</i> no tiene distribución normal.
	H_0 . Diferencias no significativas de medias de <i>incidencia Tb</i> . H_1 . Diferencias significativas de medias de <i>incidencia Tb</i> .	U de Mann-Whitney (no paramétrica), debido a que la variable <i>incidencia de Tb</i> no tiene distribución normal.

*El punto de corte de 40.3% es el valor de la media y fue elegido arbitrariamente como frontera ordinal.

Luego de planteadas las hipótesis, se procede a su contrastación. En primer lugar se generará el análisis de diferencia de medias con la prueba paramétrica t de Student, a fin de establecer si existen diferencias significativas entre los valores del índice de felicidad, según grados de desigualdad del índice de Gini. Este análisis se realiza en el SPSS de la siguiente forma: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/ Comparar medias/Prueba t para muestras independientes* (figura 39). Se despliega la caja de diálogo que contiene la lista de las

variables disponibles. Se selecciona la variable independiente y se traslada al espacio *Variables para contrastar*. Luego se elige la variable de comparación y se mueve también al espacio *Variables para contrastar*. Al final se selecciona la variable de comparación y se traslada al espacio *Variable de agrupación*. Se ingresa a *Definir grupos* y aparece una sub caja de diálogo donde se elige *Punto de corte* y se anota el valor 40.3. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 40).

Figura 39
Secuencia a seguir en el menú del SPSS para ingresar
a *Prueba t para muestras independientes*

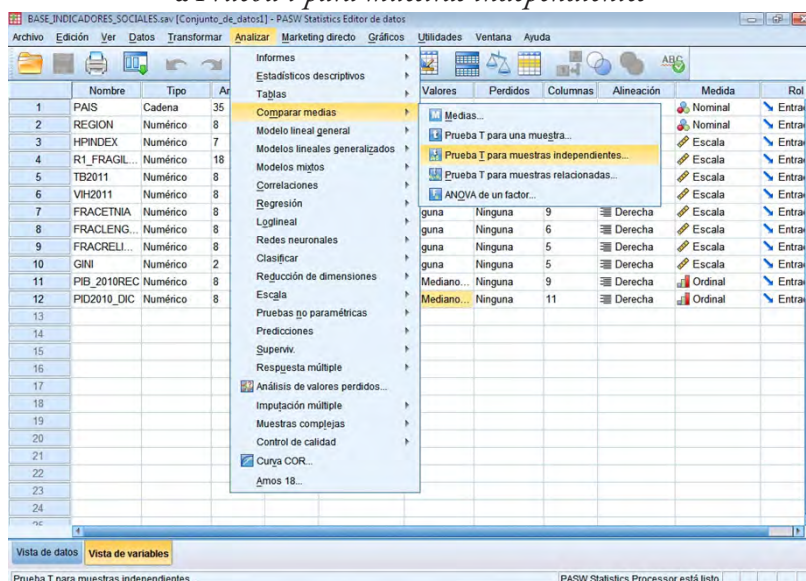
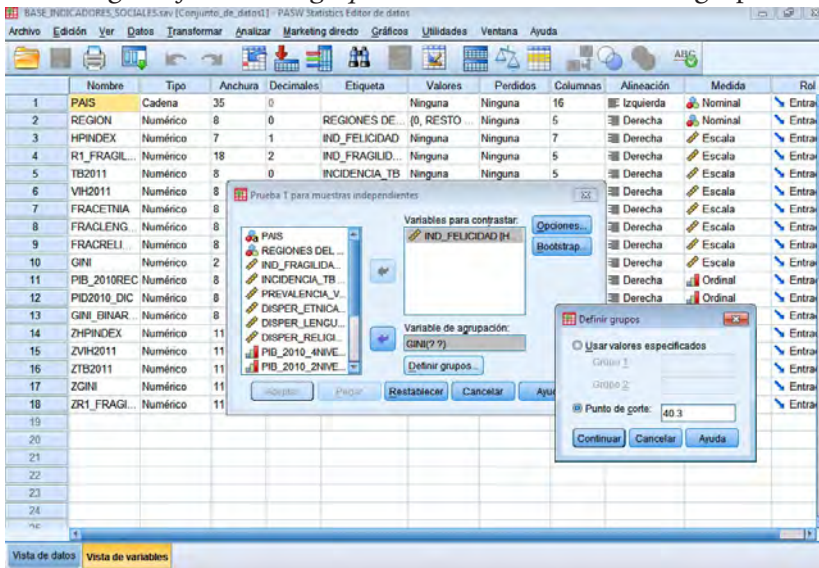


Figura 40

Caja de diálogo de la *Prueba t para muestras independientes* y sub caja de diálogo *Definición de grupos* relativa a la variable de agrupación



En el visor de resultados se despliegan las tablas *Estadísticos de grupo* y *Prueba de muestras independientes* (tabla 41 y 42). En la tabla *Estadísticos de grupo* se aprecia que la media del índice de felicidad es ligeramente mayor en la categoría de mayor desigualdad de Gini que en la de menor desigualdad (43.42 *versus* 42.76). En tanto, en la tabla de *Prueba de muestras independientes* se deben interpretar los valores de significancia *p* de la *prueba de Levene* y de la *prueba T*:

- Por una parte, la prueba de Levene evalúa si la variable independiente cumple el requisito de la homogeneidad de varianzas (homocedasticidad) en su distribución. La hipótesis nula (H_0) es la homogeneidad y la alterna (H_1) es la heterogeneidad. El valor de significancia de nuestra variable de ejemplo (índice de felicidad) es de $p = 0.002$ (es decir $p < 0.05$), por lo cual se rechaza la

hipótesis nula y se acepta la hipótesis alterna (es decir la heterogeneidad de varianzas en su distribución).

- Por otra, el hallazgo de la prueba de Levene conduce a interpretar el valor de significancia de la prueba t correspondiente a la fila *No se han asumido varianzas iguales* (que se aplica a una variable con distribución heterogénea de varianzas). En esa fila, el valor es de $p = 0.689$, lo cual quiere decir que no hay diferencias significativas entre las medias del índice de felicidad pertenecientes a los niveles de mayor y menor desigualdad. En términos descriptivos, el resultado es comprensible porque la diferencia de medias es de sólo 0.6624.

Tabla 41
Tabla estadísticos de grupo, desplegada en el visor de resultados

IND_FELICIDAD	GINI_DESIGUALDAD	N	Media	Desviación típica	Error típico de la media
	> = 40.3	54	43.424	10.1869	1.3863
	< = 40.3	73	42.762	7.6863	0.8996

Tabla 42
Prueba T de muestras independientes desplegada en el visor de resultados

IND_ FELICIDAD		Prueba de Levene para la igualdad de varianzas		Prueba T para la igualdad de medias						
		F	Sig.	t	gl	Sig. (bilateral)	Diferencia de medias	Error típico de la diferencia	95% intervalo de confianza para la diferencia	
									Inferior	Superior
Se han asumido varianzas iguales		9.604	0.002	0.418	125	0.677	0.6624	1.5855	-2.4755	3.8004
				0.401	94.680	0.689	0.6624	1.6526	-2.6185	3.9434
No se han asumido varianzas iguales										

Para continuar con el contraste de las hipótesis de diferencias de medias, se generará ahora un análisis adecuado a las variables que no cuentan con una distribución normal (prevalencia de VIH e incidencia de Tb), mediante el cálculo de la prueba U de Mann Whitney. Esta prueba se basa en un procedimiento de ordenamiento de posiciones de los valores de las muestras que permitirá establecer el número de veces en que alguna de ellas obtiene una posición mayor que la otra y si la diferencia observada es estadísticamente significativa.

Para correr la prueba U de Mann Whitney, en el SPSS, primero se debe recodificar la variable índice Gini de desigualdad como una variable ordinal de dos categorías: 0 = menor desigualdad y 1 = mayor desigualdad. Se procede de la siguiente manera:

En la barra del menú que aparece en el SPSS se sigue la secuencia *Transformar/Recodificar en distintas variables* (figura 41). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se selecciona la variable que se quiere recodificar y se traslada al espacio *Variable numérica -> Variable resultado*. Luego se escribe el nuevo nombre para la variable resultado (en nuestro ejemplo GINI_BINARIO) y se oprime el botón *Cambiar*. Se ingresa a *Valores antiguos y nuevos* y aparece una sub caja de diálogo donde se recodificarán los valores de la siguiente forma: 1) se marca la opción *Valor* y se pone como rango más bajo al valor 40.2. Luego se marca la opción *Valor* en el espacio *Valor nuevo* y se pone el número “0”. Al final, se elige *Añadir* para que el nuevo código aparezca automáticamente en el espacio *Antiguo-Nuevo* como: *Lowest thru 40.2 → 0*; 2) se marca la opción *Valor* y se pone 40.3 al valor más alto. Luego se marca la opción *Valor* en el espacio *Valor nuevo* y se escribe el número “1”. Al final, se elige *Añadir* para que el nuevo código aparezca automáticamente en el espacio *Antiguo-Nuevo* como: *40.3 thru highest → 1*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 42).

Figura 41
Secuencia a seguir en el menú del SPSS
para ingresar a *Recodificar en distintas variables*

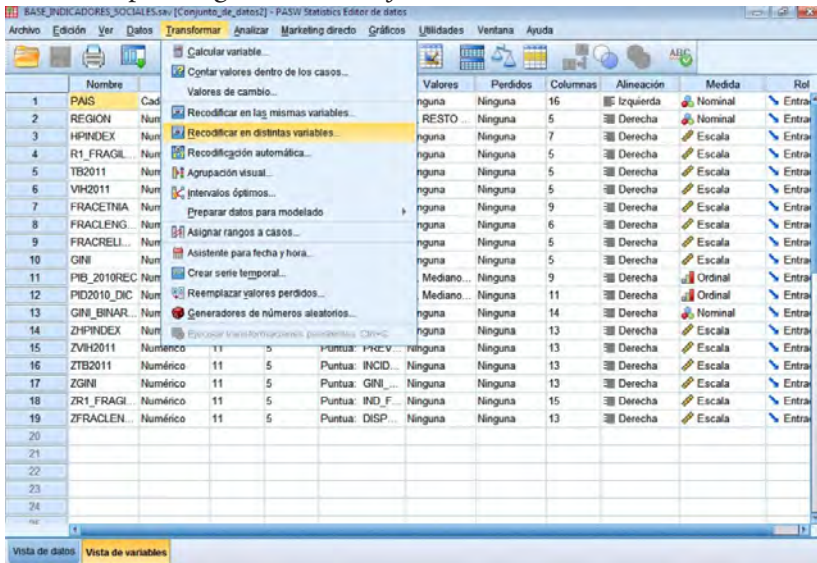
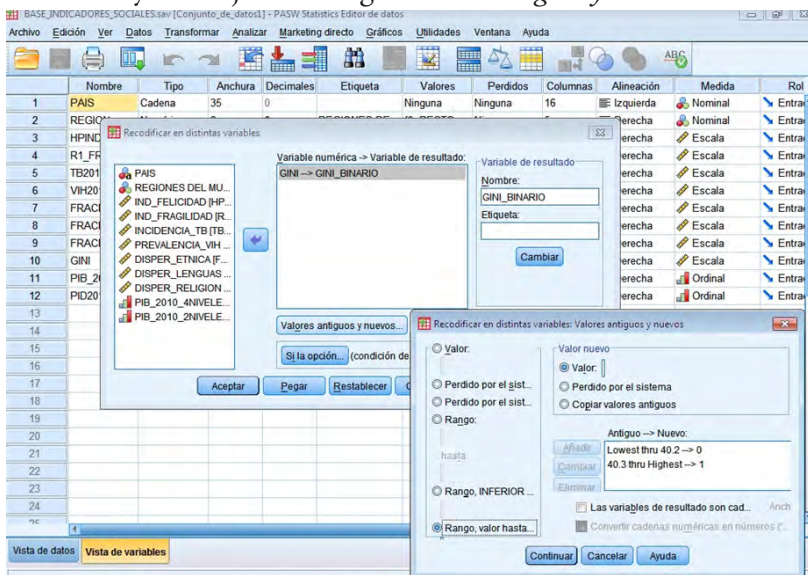


Figura 42
Caja de diálogo *Recodificar en distintas variables*
y sub caja de diálogo *Valores antiguos y nuevos*



Con la variable de comparación habilitada como ordinal, bajo el nombre de GINI_BINARIO, ahora se podrá correr la prueba U de Mann-Whitney en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Pruebas no paramétricas/Cuadros de diálogo antiguos/2 Muestras independientes* (figura 43) Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se selecciona la variable independiente y se traslada al espacio *Lista contrastar variables*. Luego se elige la variable de comparación (GINI_BINARIO) y se mueve al espacio *Variable de agrupación*. Se ingresa a *Definir* y aparece una sub caja de diálogo donde se pondrán los códigos de las categorías de la variable de agrupación: 0 en la opción “grupo 1” y 1 en “grupo 2”. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 44).

Figura 43

Secuencia a seguir en el menú del SPSS para ingresar a *Pruebas no paramétricas* y elegir *Dos Muestras independientes*

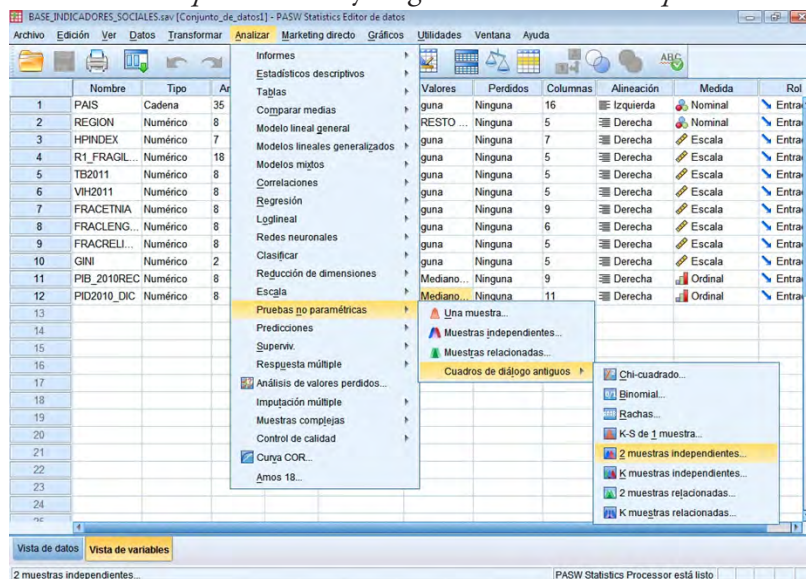
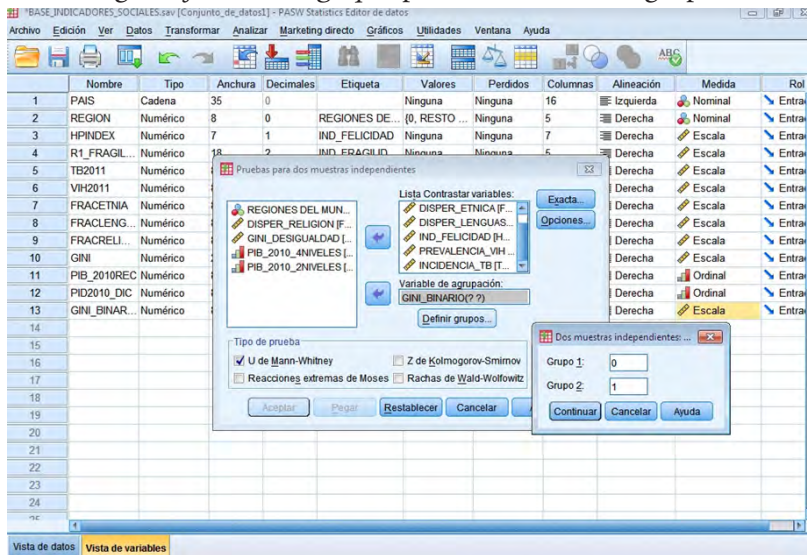


Figura 44

Caja de diálogo *Prueba para dos muestras independientes* y sub caja de diálogo *Definición de grupos* para la variable de agrupación



Tal procedimiento permitirá que, en el visor de resultados del SPSS, se generen dos tablas (tabla 43 y 44). En la primera, *tabla de rangos*, se expresan los valores promedio de la ordenación de posiciones de cada variable, dentro de las categorías de la variable comparativa. Se puede ver, por ejemplo, que en la categoría de mayor desigualdad de GINI (código 1), los valores de la *prevalencia de VIH* y de la *incidencia de Tb* tienen rangos promedio más altos que en la de menor desigualdad (75.46 *versus* 47.58 para la primera; y 83.5 *versus* 53.45 para la segunda). Al no tener rangos promedio similares, se podría pensar que las diferencias serán estadísticamente significativas, aunque hay que corroborarlo con el resultado de la prueba de significancia de U de Mann-Whitney que aparece en la tabla 44.

Tabla 43

Tabla *diferencia de rangos*, desplegada en el visor de resultados

	GINI_ BINARIO	N	Rango promedio	Suma de rangos
PREVALENCIA_VIH	0	66	47.58	3140.50
	1	53	75.46	3999.50
	Total	119		
INCIDENCIA_TB	0	74	53.45	3955.50
	1	58		4822.50
	Total	132		

En la tabla (44) sobre *estadísticos de contraste*, se expresa el valor de estadístico U de Mann-Whitney y el de la significancia p (Sig. asintótica bilateral) calculada para un valor estandarizado z que corrige sesgos de empate de rangos. Se puede ver que los valores de p de las variables de interés son menores a 0.01, lo cual permite, en ambos casos, rechazar las hipótesis de nulidad y aceptar las de diferencias significativas entre categorías de desigualdad social.

Tabla 44

Estadísticos de contraste de la prueba U de Mann-Whitney, desplegada en el visor de resultados

	PREVALENCIA_VIH	INCIDENCIA_TB
U de Mann-Whitney	929.500	1180.500
W de Wilcoxon	3140.500	3955.500
Z	-4.417	-4.427
Sig. asintótica (bilateral)	0.000	0.000

Diferencia de medias para variables que incluyen tres o más muestras independientes

Para ejemplificar el análisis de diferencia de medias para una variable comparativa que incluye tres o más muestras o cate-

gorías, seleccionaremos otras variables de nuestro repertorio disponible (véase tabla 33). Si elegimos como variable comparativa al “Nivel de ingreso de los países” (dividido en cuatro categorías ordinales: 1 = ingreso mediano bajo, 2 = ingreso mediano, 3 = ingreso mediano alto y 4 = ingreso alto) y como variables independientes a la prevalencia de VIH, la incidencia de Tb y el índice de felicidad, podremos explorar las diferencias de medias de estas últimas variables, en función de los distintos niveles de ingreso de los países.

En la tabla 45 se proponen las hipótesis estadísticas de diferencias de medias para las variables independientes y se seleccionan las pruebas de significancia adecuadas a su tipo de distribución (normal o no normal), diagnosticado en la fase anterior.

Tabla 45

Hipótesis estadísticas de diferencia de medias y selección de pruebas de significancia adecuadas al tipo de distribución de las variables

Variable de agrupación	Hipótesis estadísticas	Pruebas de significancia
Nivel de ingreso de los países con cuatro categorías: 1 = mediano bajo, 2 = mediano, 3 = mediano alto, 4 = alto	H_0 . Diferencias no significativas de medias de índice de <i>felicidad</i> . H_1 . Diferencias significativas de medias de índice de <i>felicidad</i> .	ANOVA, debido a que la variable índice de <i>felicidad</i> tiene distribución normal.
	H_0 . Diferencias no significativas de medias de <i>prevalencia</i> VIH. H_1 . Diferencias significativas de medias de <i>prevalencia</i> VIH.	H de Kruskal-Wallis, debido a que la variable <i>prevalencia de VIH</i> no tiene distribución normal.
	H_0 . Diferencias no significativas de medias de <i>incidencia</i> Tb. H_1 . Diferencias significativas de medias de <i>incidencia</i> Tb.	H de Kruskal-Wallis, debido a que la variable <i>incidencia de Tb</i> no tiene distribución normal.

Fuente: Elaboración propia.

Tras plantear las hipótesis estadísticas, se procederá a su contrastación. En primer lugar se generará el análisis de diferencia de medias con la prueba paramétrica ANOVA, a fin de establecer si existen diferencias significativas entre los valores del índice de felicidad, según el nivel de ingresos que presentan los países. Para correr la ANOVA en el SPSS se sigue el siguiente camino:

En la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Comparar medias/ANOVA de un factor* (figura 45). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se selecciona la variable “índice de felicidad” y se traslada al espacio *Lista de dependientes*. Luego se elige la variable comparativa “Nivel de ingresos” y se mueve al espacio *Factor*. Se ingresa a *Opciones* y aparece una sub caja de diálogo donde se seleccionan *Descriptivos* y *Prueba de homogeneidad de varianzas*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 46).

Figura 45

Secuencia a seguir en el menú del spss para ingresar a *ANOVA de un factor*

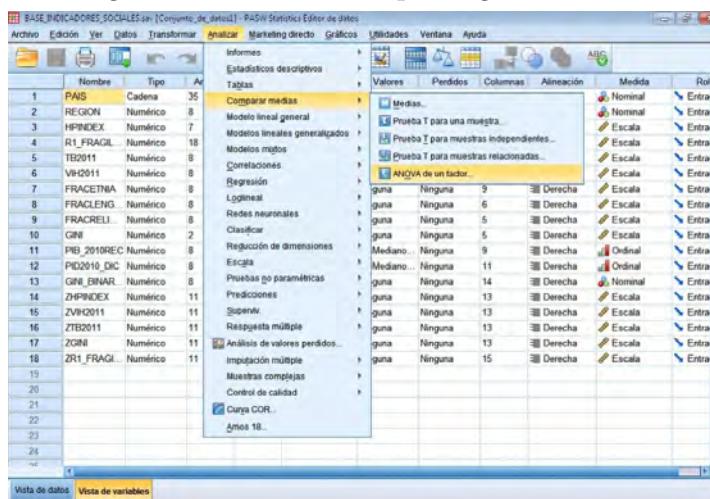
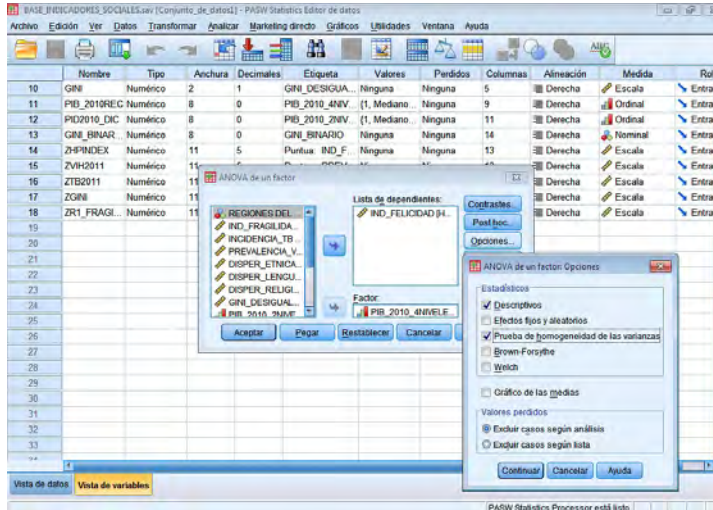


Figura 46

Caja de diálogo *ANOVA de un factor* y sub caja de diálogo de *Opciones*

En el visor de resultados del SPSS se despliegan tres tablas (véanse tablas 46, 47 y 48). La primera, *tabla de descriptivos*, expresa los valores de las medias y desviaciones típicas de la variable índice de felicidad para las categorías *nivel de ingresos*. Se puede ver que las diferencias de los valores no son marcadas (el más elevado es 46.584 para el ingreso mediano y el más bajo 40.671 para el ingreso alto, donde la media grupal es de 42.183).

La segunda tabla, *prueba de homogeneidad de varianzas* (tabla 47), presenta el resultado del estadístico de Levene que, como vimos anteriormente, al tocar la *t* de Student, evalúa si la variable independiente cumple el requisito de la homogeneidad de varianzas (homocedasticidad) en su distribución. En esta evaluación se asume como hipótesis nula (H_0) a la homogeneidad y como hipótesis alterna (H_1) a la heterogeneidad, por lo que al ser el valor de $P = 0.245$ (es decir $P > 0.05$) se rechaza la hipótesis alterna y se acepta la nula (la homogeneidad).

Finalmente, la tercera tabla (48) muestra el valor de la ANOVA; ésta establece que la diferencia de las medias no es significativa ($p = 0.116$), por lo cual se acepta la hipótesis de nulidad.

Tabla 46
Tabla descriptivos de diferencia de medias, desplegada en el visor de resultados

IND_FELICIDAD	N	Media	Desviación típica	Error típico	Intervalo de confianza para la media al 95%		Mínimo	Máximo
					Límite inferior	Límite superior		
Mediano-bajo	64	41.105	9.6368	1.2046	38.697	43.512	24.7	60.4
Mediano	19	46.584	9.0269	2.0709	42.233	50.935	28.3	59.8
Mediano-alto	42	42.590	8.9969	1.3883	39.787	45.394	22.6	64.0
Alto	21	40.671	7.4378	1.6231	37.286	44.057	25.2	51.4
Total	146	42.183	9.1877	0.7604	40.680	43.686	22.6	64.0

Tabla 47

Tabla *prueba de homogeneidad de varianzas*,
desplegada en el visor de resultados

IND_FELICIDAD			
Estadístico de Levene	gl	gl2	Sig.
14.01	3	142	0.145

Tabla 48

Tabla de la ANOVA, desplegada en el visor de resultados

IND_FELICIDAD					
	Suma de cuadrados	gl	Media cuadrática	F	Sig.
Inter-grupos	497.414	3	165.805	2.005	0.116
Intra-grupos	11742.593	142	82.694		
Total	12240.007	145			

Cabe señalar que al no haber encontrado la ANOVA diferencias significativas entre grupos, ya no será necesario realizar el contraste de diferencias entre todas las categorías, mediante la selección de alguna prueba de significancia POST HOC. En caso contrario, si la ANOVA hubiera detectado diferencias significativas, lo que hubiera correspondido hacer es la selección de una prueba POST HOC para variables independientes que cumplen el requisito de la homogeneidad de varianzas (esto debido a que el estadístico de Levene tuvo un valor de $p > 0.05$).

Para continuar con el contraste relativo a las hipótesis de diferencias de medias, se realizará ahora un análisis adecuado a las variables que no cuentan con una distribución normal (la prevalencia de VIH y la incidencia de Tb), mediante el cálculo de la prueba no paramétrica H de Kruskal-Wallis. Esta prueba se basa en un procedimiento donde se realiza el ordenamiento de posiciones de los valores de tres o más muestras que intenta establecer el número de veces en que alguna de ellas obtiene una

posición mayor que las otras y si las diferencias observadas son estadísticamente significativas.

El cálculo de H de Kruskal-Wallis se realiza en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia: *Analizar/Pruebas no paramétricas/Cuadros de diálogo antiguos/K muestras independientes* (figura 47). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables de interés y se trasladan al espacio *Lista contrastar variables*. Luego se elige la variable comparativa “Nivel de ingresos” y se mueve al espacio *Variable de agrupación*. Se ingresa a *Definir rango* y aparece una sub caja de diálogo donde se indican los valores del rango: se pone 1 en “mínimo” y 4 en “máximo”. Se elige *Continuar* y luego en “Tipo de prueba” se elige “H de Kruskal-Wallis”. Finalmente se selecciona *Aceptar* para ejecutar el procedimiento (figura 48).

Figura 47
Secuencia a seguir en el menú del SPSS
para ingresar a *K muestras independientes*

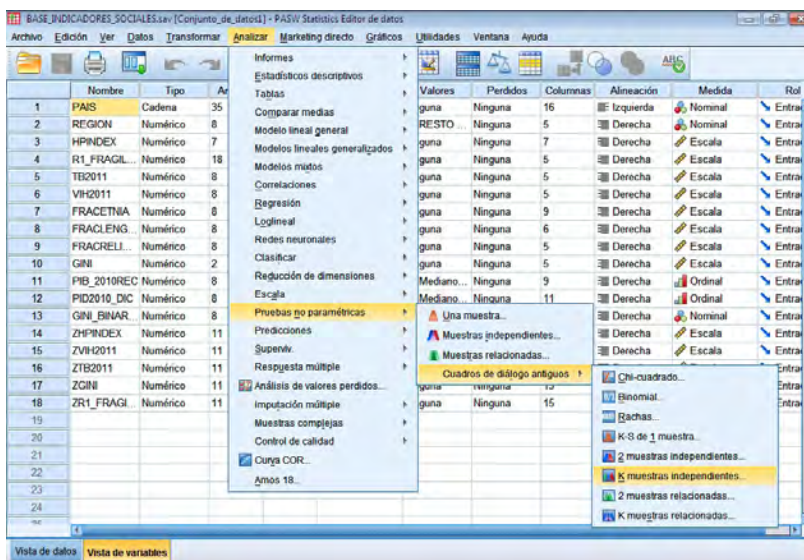
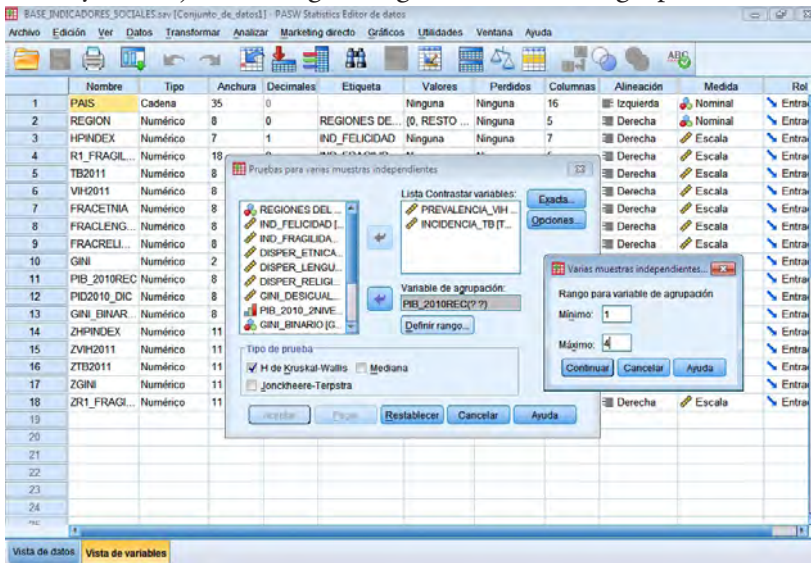


Figura 48

Caja de diálogo Prueba para varias muestras independientes
y sub caja de diálogo Rango de la variable agrupación



En el visor de resultados del SPSS se reportan dos tablas (49 y 50). La primera de ellas, *tabla de rangos*, expresa los valores de rango promedio de la ordenación de posiciones de cada variable, en las categorías *nivel de ingresos*. Se puede ver, por una parte, que las prevalencias de VIH y las incidencias de Tb son más elevadas conforme los niveles de ingreso son más bajos, y, por otra, que las diferencias de rango promedio son amplias y que podrían ser estadísticamente significativas.

En la segunda tabla, *estadísticos de contraste*, presenta como estadístico de contraste a la prueba Chi-cuadrado y sus valores de significancia *p* (Sig. asintótica). Se puede ver que para las dos variables consideradas, los valores de *p* fueron menores a 0.01 y que permiten rechazar la hipótesis nula y aceptar la hipótesis alterna de diferencias significativas entre rangos promedios por niveles de ingresos de cada país.

Tabla 49
Tabla de la *prueba Kruskal-Wallis* (rangos),
desplegada en el visor de resultados

	PIB_2010_4NIVELES	N	Rango promedio
PREVALENCIA_VIH	Mediano-bajo	67	84.20
	Mediano	14	72.89
	Mediano-alto	42	59.71
	Alto	18	46.72
	Total	141	
INCIDENCIA_TB	Mediano-bajo	82	124.88
	Mediano	20	88.20
	Mediano-alto	55	63.22
	Alto	22	28.59
	Total	179	

Tabla50
Tabla *estadísticos de contraste*^{a,b} *Chi-cuadrado*,
desplegada en el visor de resultados

	PREVALENCIA_VIH	INCIDENCIA_TB
Chi-cuadrado	16.832	82.772
gl	3	3
Sig. asintótica	0.001	0.000

^a. Prueba de Kruskal-Wallis.

^b. Variable de agrupación: PIB_2010_4NIVELES.

Cuando las diferencias de rangos promedio son estadísticamente significativas, es posible realizar, en forma adicional, comparaciones entre categorías para apreciar cuáles diferencias específicas son significativas. Ello se realiza mediante pruebas POST HOC de comparaciones múltiples que no asumen varianzas iguales. Para correr dichas pruebas en el SPSS se realiza lo siguiente:

En la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Comparar medias/ANOVA de un factor* (figura 49). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan primero las variables “Prevalencia de VIH” e “Incidencia de Tb” y se trasladan al espacio *Lista de dependientes*. Luego se selecciona la variable comparativa *Nivel de ingresos* y se mueve al espacio *Factor*. Se ingresa a POST HOC y aparece una sub caja de diálogo donde se selecciona *T2 de Tamhane* como prueba de contraste que no asume varianzas iguales. Luego se elige *Continuar*. Se ingresa a *Opciones* y se marca *Gráfico de medias*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 50).

Figura 49
Caja de diálogo *ANOVA de un factor*
y sub caja de diálogo de pruebas POST HOC

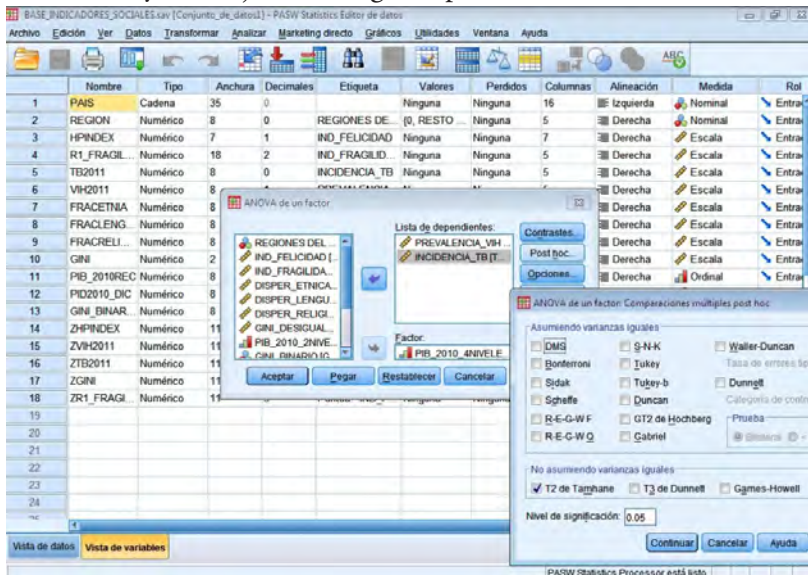
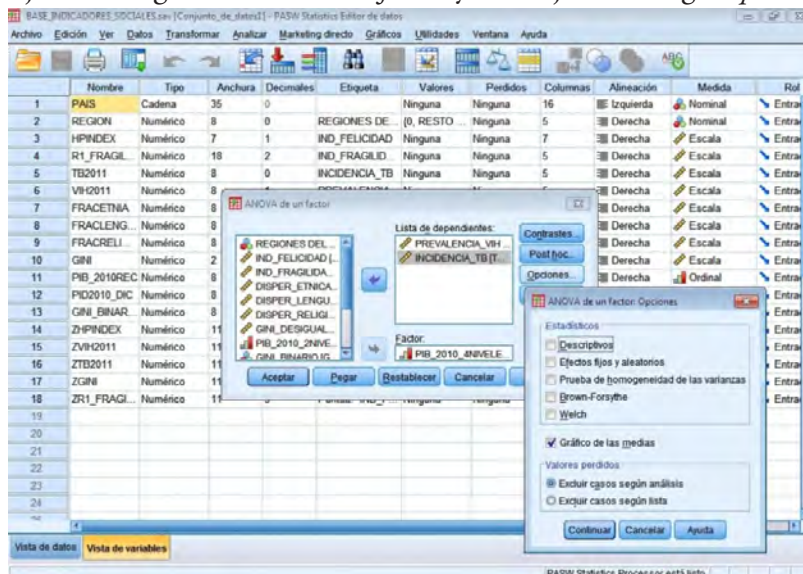


Figura 50

Caja de diálogo de *ANOVA de un factor* y sub caja de diálogo *Opciones*



En el visor de resultados del SPSS se despliega, en primera instancia, la tabla *comparaciones múltiples de Tamhane*, en la cual se cotejan las diferencias de las medias de todas las categorías de las variables seleccionadas y se expresan sus valores de significancia (Sig.), al respecto véase tabla 51. En el caso de la variable *prevalencia de VIH* existe una diferencia estadísticamente significativa entre la media de los países de ingreso “mediano-bajo” y de ingreso “alto” (diferencia de +2.4 a favor de los primeros, $p = 0.001$), mientras que en la variable *incidencia de Tb* se aprecian diferencias significativas entre la media de los países de ingreso “mediano-bajo”, “mediano-alto” y “alto” (diferencias de 136.6 y 191.1 a favor de los países de menor ingreso frente a los otros, $p < 0.01$ en los dos casos).

En segunda instancia, aparecen en el visor de resultados los *gráficos de diferencia de medias* que ilustran las marcadas diferencias en las prevalencias de VIH e incidencias de Tb según los niveles de ingresos de los países. Los gráficos permiten apreciar la amplia diferencia entre los países de ingresos más bajos y más altos (tabla 52).

Tabla 51
Tabla comparaciones múltiples de Tamhane, desplegada en el visor de resultados

Variable dependiente	(i) PIB_2010_4NIVELES	(j) PIB_2010_4NIVELES	Diferencia de medias (i-j)	Error típico	Sig.	Intervalo de confianza para la media al 95%	
						Límite inferior	Límite superior
PREVALENCIA_VIH	Mediano-bajo	Mediano	.9874	1.1039	0.943	-2.159	4.134
		Mediano-alto	1.0612	0.9117	0.818	-1.389	3.511
		Alto	2.4009*	0.6049	0.001	0.761	4.041
	Mediano	Mediano-bajo	-.9874	1.1039	0.943	-4.134	2.159
		Mediano-alto	0.738	1.1491	1.000	-3.175	3.323
		Alto	1.4135	0.9247	0.624	-1.448	4.275
	Mediano-alto	Mediano-bajo	-1.0612	0.9117	0.818	-3.511	1.389
		Mediano	-.0738	1.1491	1.000	-3.323	3.175
		Alto	1.3397	0.6838	0.296	-0.550	3.229
	Alto	Mediano-bajo	-2.4009*	0.6049	0.001	-4.041	-0.761
		Mediano	-1.4135	0.9247	0.624	-4.275	1.448
		Mediano-alto	-1.3397	0.6838	0.296	-3.229	0.550

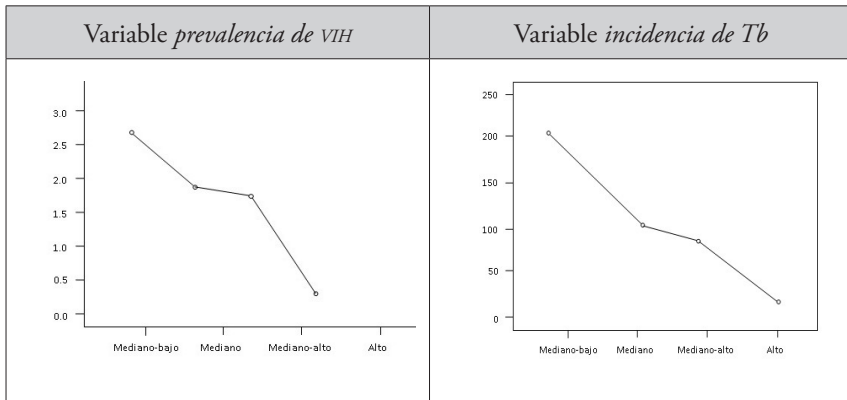
Continúa en la página 216

Viene de la página 215

Variable dependiente	(i) PIB_2010_4NIVELES	(j) PIB_2010_4NIVELES	Diferencia de medias (i-j)	Error típico	Sig.	Intervalo de confianza para la media al 95%	
						Límite inferior	Límite superior
INCIDENCIA_TB	Mediano-bajo	Mediano	102.780	42.057	0.113	-14.56	220.12
		Mediano-alto	136.315*	30.234	0.000	55.56	217.07
		Alto	191.091*	22.199	0.000	131.30	250.88
	Mediano	Mediano-bajo	-102.780	42.057	0.113	-220.12	14.56
		Mediano-alto	33.535	41.533	0.964	-82.76	149.83
		Alto	88.311	36.107	0.136	-17.38	194.00
	Mediano-alto	Mediano-bajo	-136.315*	30.234	0.000	-217.07	-55.56
		Mediano	-33.535	41.533	0.964	-149.83	82.76
		Alto	54.775	21.189	0.072	-2.96	112.52
	Alto	Mediano-bajo	-191.091*	22.199	0.000	-250.88	-131.30
		Mediano	-88.311	36.107	0.136	-194.00	17.38
		Mediano-alto	-54.775	21.189	0.072	-112.52	2.96

*La diferencia de medias es significativa en el nivel 0.05.

Tabla 52
Gráficos referentes a *diferencia de medias*
desplegados en el visor de resultados



Relación de pares de variables mediante pruebas de asociación

Hasta este momento, la exploración de relaciones significativas entre pares de variables de intervalo/razón se enfocó a la búsqueda de patrones de diferencias de medias para variables comparativas de dos y más categorías.

Ahora se enfatizará en la búsqueda de patrones de relación mediante pruebas que contrasten si existe una asociación significativa entre cómo se distribuyen los valores de las variables. Si al explorar los patrones de asociación las variables de interés tienen una distribución normal, habitualmente se seleccionan pruebas paramétricas como la *Correlación r de Pearson* y la *Regresión lineal*. Si no tienen una distribución normal se recurre a las pruebas no paramétricas como la *correlación ρ de Spearman*. En general, dichas pruebas estiman la magnitud, dirección y significancia de la asociación entre dos variables.

Para generar pruebas de asociación, es necesario proponer hipótesis estadísticas sobre las relaciones de variables. En la tabla 53 se hace una proposición que incluye los tipos de variables (independiente y dependiente), la magnitud esperada de la co-

relación y la selección de pruebas de significancia adecuadas a la distribución de las variables (normal o no normal).

Tabla 53
Hipótesis estadísticas relativas a la asociación
de pares de variables y selección de pruebas
de significancia adecuadas a la distribución de las variables

Variables independientes	Hipótesis estadísticas	Pruebas de significancia
Índice de felicidad	H ₀ . Baja correlación y asociación no significativa con la variable dependiente <i>prevalencia de VIH</i> . H ₁ . Alta correlación y asociación significativa con la variable dependiente <i>prevalencia de VIH</i> .	Correlación <i>r de Pearson</i> , debido a que las variables independientes tienen distribución normal
Índice Gini de desigualdad	H ₀ . Baja correlación y asociación no significativa con la variable dependiente <i>incidencia de Tb</i> . H ₁ . Alta correlación y asociación significativa con la variable dependiente <i>incidencia de Tb</i> .	
Índice de fragilidad	H ₀ . Baja correlación y asociación no significativa con la variable dependiente <i>prevalencia de VIH</i> . H ₁ . Alta correlación y asociación significativa con la variable dependiente <i>prevalencia de VIH</i> .	<i>rho de Spearman</i> , debido a que las variables independientes no tienen distribución normal
Dispersión étnica	H ₀ . Baja correlación y asociación no significativa con la variable dependiente <i>incidencia de Tb</i> . H ₁ . Alta correlación y asociación significativa con la variable dependiente <i>incidencia de Tb</i> .	
Dispersión de lenguas	H ₀ . Baja correlación y asociación no significativa con la variable dependiente <i>incidencia de Tb</i> . H ₁ . Alta correlación y asociación significativa con la variable dependiente <i>incidencia de Tb</i> .	

Una vez que las hipótesis fueron propuestas, se genera primero el análisis de *correlación r de Pearson* para las variables independientes que tienen una distribución normal. El análisis se realiza en el SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Correlaciones/Bivariadas* (figura 51). Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables de interés y se trasladan al espacio *Variables*. Luego se eligen las opciones *Pearson*, y *Unilateral* (debido a que se asume una relación de dependencia) y *Marcar las correlaciones significativas*. Finalmente se elige *Aceptar* para ejecutar el procedimiento (figura 52).

Figura 51
Secuencia a seguir en el menú del SPSS
para ingresar a *Correlaciones bivariadas*

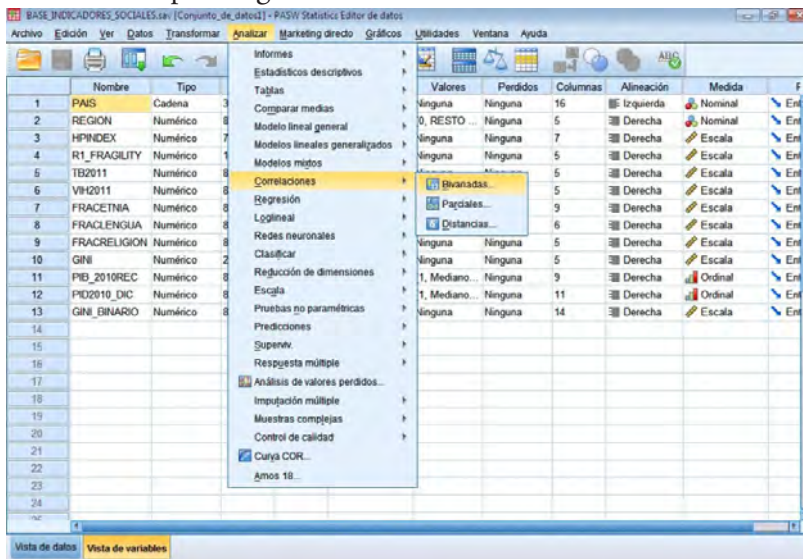
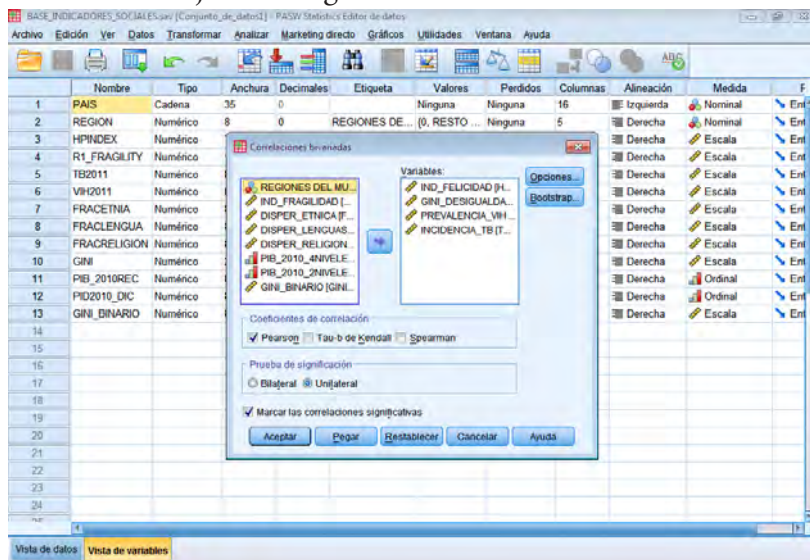


Figura 52
Caja de diálogo Correlaciones bivariadas



En el visor de resultados se despliega una tabla de correlaciones de todas las variables seleccionadas (tabla 54). Nosotros leeremos solamente los cruces de variables que corresponden a las hipótesis propuestas (celdas sombreadas):

- En primer lugar, vemos que el índice de felicidad tiene un coeficiente de correlación de r de Pearson = -0.365 con la *prevalencia de VIH* y que ambas variables tienen una asociación estadísticamente significativa (Sig. unilateral de $p = 0.000$). Dos asteriscos al lado del coeficiente hacen notar que la relación es significativa. El valor del coeficiente indica una correlación baja y el signo negativo una relación inversa (a mayor felicidad, una menor prevalencia de VIH).
- En el caso de la relación del índice de felicidad con la incidencia de Tb, el coeficiente también indica una correlación baja e inversa y estadísticamente significativa (r de Pearson = -0.302, $p < 0.01$).

- El índice Gini de desigualdad tiene correlaciones regulares o modestas de signo positivo con la *prevalencia de VIH* (r de Pearson = 0.484) y la *incidencia de Tb* (r de Pearson = 0.497). Ambas correlaciones son estadísticamente significativas en el nivel $p < 0.001$.

Tales resultados conducen a rechazar las hipótesis nulas de no asociación y, en consecuencia, a aceptar las hipótesis de asociación significativa entre las variables analizadas.

Tabla 54
Tabla correlaciones *r* de Pearson, desplegada en el visor de resultados

	IND_FELICIDAD	GINI_DESIGUALDAD	PREVALENCIA_VIH	INCIDENCIA_TB
IND_FELICIDAD	Correlación de Pearson	1	0.022	-0.302**
	Sig. (unilateral)		0.402	0.000
	N	151	127	151
GINI_DESIGUALDAD	Correlación de Pearson	0.022	1	0.484**
	Sig. (unilateral)	0.402		0.000
	N	127	132	132
PREVALENCIA_VIH	Correlación de Pearson	-0.365	0.484	1
	Sig. (unilateral)	0.000	0.000	0.000
	N	129	119	145
INCIDENCIA_TB	Correlación de Pearson	-0.302	0.497	0.757
	Sig. (unilateral)	0.000	0.000	1
	N	151	132	190

**La correlación es significativa en el nivel 0,01 (unilateral).

En el recuadro 3 se presentan algunos criterios generales para interpretar los coeficientes de correlación:

Recuadro 3

Interpretación de los coeficientes de correlación

El valor y el signo de un coeficiente r de Pearson permiten interpretar en forma conjunta el peso y la dirección de una correlación lineal (tabla B). Por una parte, el peso se interpreta de acuerdo con el valor que adopta el coeficiente: a) un valor de 0 indica la ausencia de correlación y un valor de 1 una correlación perfecta; b) si los valores se acercan a 0 expresan una correlación débil y, si se aproximan a 1, una correlación fuerte; y c) los valores intermedios señalan una correlación regular o modesta.

Por otra parte, la dirección se puede ver en el signo del coeficiente: a) cuando el signo es positivo (+) la correlación tiene un sentido positivo; y b) cuando es negativo (-) tiene un sentido negativo o inverso.

Tabla B

Interpretación del peso y la dirección de los coeficientes de correlación

Perfecta					Perfecta
negativa		Ninguna			positiva
-1.00		0			1.00
Fuerte	←————→ Débil		Débil ←————→	Fuerte	

Fuente: Bryman & Cramer (2009: 217).

De esa manera, la interpretación del coeficiente siempre abarca el peso y la dirección de la correlación. Por ejemplo, si el valor es de $r = -0.12$, se interpreta que la correlación es débil y de sentido negativo o inverso. Si el valor es de $r = 0.85$ se interpreta que hay un correlación fuerte y positiva.

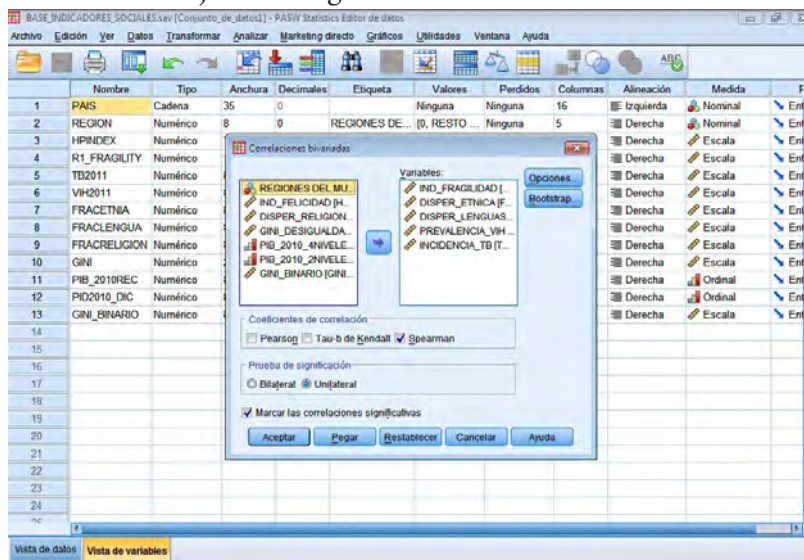
El anterior esquema de interpretación es muy general y no entra en detalles sobre cuáles serían los puntos de corte adecuados para expresar los distintos niveles de peso. Algunos autores sugieren categorizar los niveles de correlación en: ninguna, muy baja, baja, regular o modesta, alta y muy alta, pero lo hacen en un contexto donde no hay un consenso sobre cuáles serán los intervalos que abarcarán las categorías.

A continuación, se realizará el análisis de correlación *Rho de Spearman* para las variables independientes con distribución no normal de nuestro ejemplo: Índice de fragilidad, dispersión

étnica y dispersión de lenguas. Este análisis se procesa en el SPSS de la siguiente forma:

En la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Correlaciones/Bivariadas*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles. Se seleccionan las variables de interés y se trasladan al espacio *Variables*. Luego se eligen las opciones *Spearman*, *Unilateral* (debido a que se asume una relación de dependencia) y *Marcar las correlaciones significativas*. Finalmente se elige *Aceptar* para ejecutar el procedimiento (figura 53).

Figura 53
Caja de diálogo *Correlaciones bivariadas*



En el visor de resultados se despliega una tabla de correlaciones de las variables seleccionadas (tabla 55). Los pares de variables que contrastan las hipótesis propuestas son los de las celdas sombreadas. De la lectura de la tabla se puede destacar lo siguiente:

- El índice de fragilidad tuvo una correlación regular o modesta con la prevalencia de VIH ($\rho = 0.419$) y una

correlación alta con la incidencia de Tb ($\rho = 0.693$). Su asociación con ambas variables fue estadísticamente significativa en el nivel de $p < 0.01$.

- La dispersión étnica tuvo correlaciones regulares o modestas con la prevalencia de VIH ($r = 0.526$) y la incidencia de Tb ($r = 0.418$). Su asociación con ambas variables fue estadísticamente significativa en el nivel de $p < 0.01$.
- La dispersión de lenguas tuvo correlaciones regulares o modestas con la prevalencia de VIH ($r = 0.493$) y la incidencia de Tb ($r = 0.432$). Su asociación con ambas variables fue estadísticamente significativa en el nivel de $p < 0.01$.

Tales resultados nos permiten rechazar las hipótesis nulas de no asociación y, por ello, aceptar las hipótesis de asociación significativa entre las variables analizadas.

Tabla 55
Tabla correlaciones Rho de Spearman desplegada en el visor de resultados

Rho de Spearman	IND_FRAGILIDAD	IND_ FRAGILIDAD	DISPER_ ÉTNICA	DISPER_ LENGUAS	PREVALENCIA_ VIH	INCIDENCIA_ TB
	IND_FRAGILIDAD	Coeficiente de correlación				
		Sig. (unilateral)				
		N				
	DISPER_ÉTNICA	Coeficiente de correlación				
		Sig. (unilateral)				
		N				
	DISPER_LENGUAS	Coeficiente de correlación				
		Sig. (unilateral)				
		N				
	PREVALENCIA_VIH	Coeficiente de correlación				
		Sig. (unilateral)				
		N				
	INCIDENCIA_TB	Coeficiente de correlación				
		Sig. (unilateral)				
		N				

**La correlación es significativa en el nivel 0,01 (unilateral).

Debido a que todas las variables correlacionadas tuvieron asociaciones significativas, se cuenta con una base importante para generar y contrastar un modelo de regresión múltiple que sea predictor de una variable dependiente e incluya el efecto ajustado de todas las variables independientes. Esta tarea se realizará en el siguiente apartado.

Análisis multivariado: modelo de regresión múltiple

A continuación, se ejemplificará la aplicación del análisis de regresión lineal múltiple a nuestros datos de intervalo/razón. Dicho análisis busca construir un modelo en el cual distintas variables independientes ($\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$) expliquen el comportamiento de una variable dependiente (y).

El primer paso en la construcción del modelo de regresión es identificar qué variables se incluirán en el mismo. Un primer criterio de inclusión sería elegir aquellas variables con correlación significativa ($p < 0.05$) en su asociación con la variable dependiente. Este criterio es flexible y abre la posibilidad de incluir variables con valores cercanos a la significancia estadística ($p = 0.10$ o $p = 0.15$). Un segundo criterio sería incluir variables sin correlación significativa pero con relevancia teórica en la explicación de la variable dependiente.

En nuestro ejemplo, dimos este primer paso al seguir el primer criterio, puesto que en la fase anterior pudimos identificar correlaciones significativas, en el nivel de $p < 0.01$ entre nuestras variables de interés (véanse tablas 54 y 55). Tal identificación nos permite proponer la posibilidad de construir dos modelos con sus respectivas variables independientes y dependientes (tabla 56).

Tabla 56
Variables independientes y variable dependiente de los modelos de regresión lineal múltiple que se intentarán construir

Variables independientes	Variable dependiente
<i>Modelo 1:</i> $\bar{x}_1$. Índice de felicidad $\bar{x}_2$. Índice Gini de desigualdad $\bar{x}_3$. Índice de fragilidad $\bar{x}_4$. Dispersión étnica $\bar{x}_5$. Dispersión de lenguas	Y. Incidencia de Tb
<i>Modelo 2:</i> $\bar{x}_1$. Índice de felicidad $\bar{x}_2$. Índice Gini de desigualdad $\bar{x}_3$. Índice de fragilidad $\bar{x}_4$. Dispersión étnica $\bar{x}_5$. Dispersión de lenguas	Y. Prevalencia de VIH

Fuente: Elaboración propia.

Luego de proponer los modelos de regresión que deseamos construir, el segundo paso consistirá en la estimación de los mismos mediante dos procedimientos complementarios (tabla 57):

- *La estimación y contraste de los parámetros o coeficientes de regresión beta (β), con el propósito de detectar el efecto de cada variable explicativa sobre la variable dependiente (y). En este procedimiento se contrastan las siguientes hipótesis estadísticas: la hipótesis nula (H_0) es que todos los coeficientes de regresión (β) son igual a cero (no existe relación lineal), mientras que la hipótesis alterna (H_1) expresa que por lo menos algún coeficiente de regresión será distinto de cero y tendrá una relación significativa con y (tabla 57).*
- *La estimación de la variable dependiente (y) mediante la descomposición de dos elementos de su variabilidad: a) el coeficiente de ajuste de la recta de regresión a la va-*

riable Y (coeficiente de determinación R^2) que expresa el valor de la fuerza de la asociación lineal del modelo (varía de 0 a 1, donde 0 es la ausencia y 1 la asociación perfecta). Su valor, multiplicado por 100, indica el porcentaje de la variabilidad de Y explicada por el modelo; y b) el contraste de las hipótesis estadísticas de regresión mediante la prueba F de *Snedecor* del Análisis de Varianza (ANOVA). Esta prueba contrasta las hipótesis nulas y alternas de los coeficientes de determinación, para establecer si se acepta o rechaza el modelo explicativo (tabla 57).

Tabla 57
Hipótesis estadísticas de la estimación
de coeficientes de regresión de los modelos a construir

Variables independientes	Variable dependiente	Hipótesis estadísticas complementarias
Modelo 1: $\bar{x}_1$. Índice de felicidad β_1 $\bar{x}_2$. Índice Gini de desigualdad β_2 $\bar{x}_3$. Índice de fragilidad β_3 $\bar{x}_4$. Dispersión étnica β_4 $\bar{x}_5$. Dispersión de lenguas β_5	Y. Incidencia de Tb	H_0 . $\beta_1 = \beta_2 = \beta_k = 0$ H_1 algún $\beta_{1,5} \neq 0$ y H_0 . $R^2 = 0$ H_1 . $R^2 \neq 0$
Modelo 2: $\bar{x}_1$. Índice de felicidad β_1 $\bar{x}_2$. Índice Gini de desigualdad β_2 $\bar{x}_3$. Índice de fragilidad β_3 $\bar{x}_4$. Dispersión étnica β_4 $\bar{x}_5$. Dispersión de lenguas β_5	Y. Prevalencia de VIH	H_0 . $\beta_1 = \beta_2 = \beta_k = 0$ H_1 algún $\beta_{1,5} \neq 0$ y H_0 . $R^2 = 0$ H_1 . $R^2 \neq 0$

Fuente: Elaboración propia.

Para correr esta fase de análisis en el SPSS, respecto al modelo donde la variable dependiente es la incidencia de Tb, se realiza lo siguiente: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Regresión/Lineales* (figura 54). Se despliega

la caja de diálogo que contiene la lista de las variables disponibles. Se selecciona, primero, la variable dependiente “Incidencia de Tb” y se traslada al espacio *Dependientes*. Luego se seleccionan las variables independientes y se mueven al espacio *Independientes*. En el espacio *Método* se elige *Pasos sucesivos* (método que incluye primero las variables con mayor correlación y luego las de menor correlación). Se ingresa a *Estadísticos* y se despliega una sub caja de diálogo donde se seleccionan *Estimaciones*, *Intervalos de confianza del 95%*, *Ajuste del modelo*, *Cambio en R cuadrado*, *Descriptivos* y *Correlaciones parciales y semiparciales*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 55).

Figura 54
Secuencia a seguir en el menú del spss
para ingresar a *Regresiones Lineales*

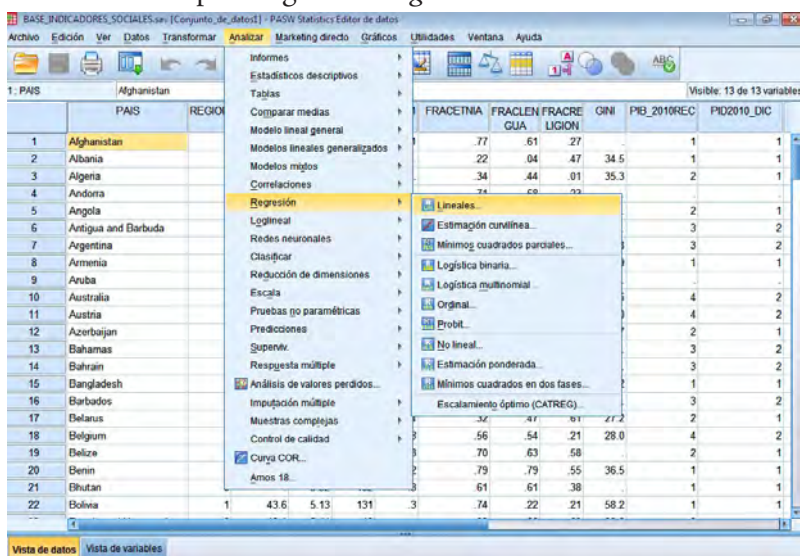
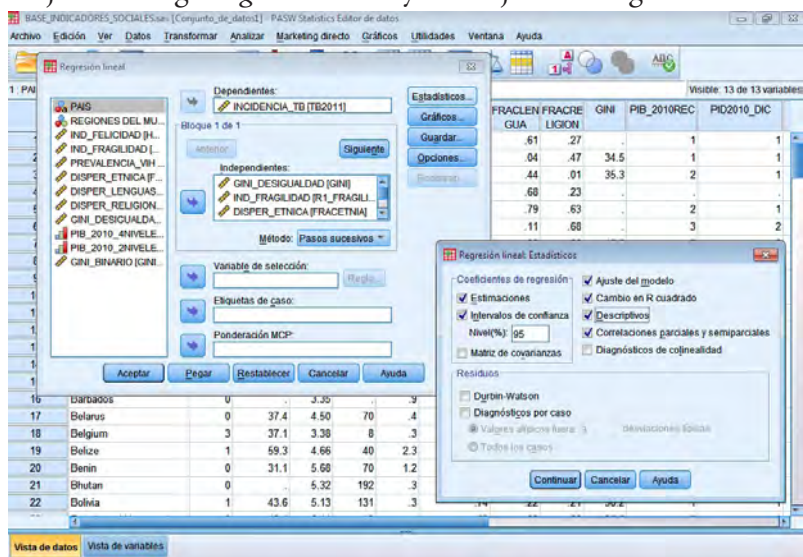


Figura 55

Caja de diálogo *Regresión lineal* y sub caja de diálogo *Estadísticos*

En el visor de resultados se despliegan inicialmente tres tablas: *estadísticos descriptivos*, *resumen del modelo* y *prueba ANOVA* para el modelo global (tablas 58, 59 y 60).

Tabla 58
Tabla estadísticos descriptivos

	Media	Desviación típica	N
INCIDENCIA_TB	114.91	162.548	122
GINI_DESIGUALDAD	39.740	9.7254	122
IND_FRAGILIDAD	4.6904	1.04393	122
DISPER_ÉTNICA	0.4273	0.24833	122
DISPER_LINGUAS	0.3854	0.28293	122
IND_FELICIDAD	42.993	8.8475	122

Tabla 59
Resumen del modelo

Modelo	R	R cuadrado	R cuadrado corregida	Error típico de la estimación	Estadísticos de cambio				
					Cambio en R cuadrado	Cambio en F	gl1	gl2	Sig. Cambio en F
1	0.526	0.276	0.270	138.861	0.276	45.802	1	120	0.000
2	0.650	0.423	0.413	124.533	0.147	30.201	1	119	0.000
3	0.677	0.459	0.445	121.077	0.036	7.891	1	118	0.006
4	0.694	0.481	0.463	119.081	0.022	4.989	1	117	0.027

A. Variables predictoras: (Constante), GINI_DESIGUALDAD.

B. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas.

C. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas, ind_felicidad.

D. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas, ind_felicidad, ind_fragilidad.

Tabla 60
Prueba de ANOVA^E

Modelo		Suma de cuadrados	gl	Media cuadrática	F	Sig.
1	Regresión	883171.701	1	883333171.701	45.802	0.000
	Residual	2313888.289	120	19282.402		
	Total	3197059.990	121			
2	Regresión	1351541.841	2	675770.921	43.574	0.000
	Residual	1845518.149	119	15508.556		
	Total	3197059.990	121			
3	Regresión	1467218.803	3	489072.934	33.362	0.000
	Residual	1729841.187	118	14659.671		
	Total	3197059.990	121			
4	Regresión	1537958.868	4	384489.717	27.114	0.000
	Residual	1659101.122	117	14180.351		
	Total	3197059.990	121			

A. Variables predictoras: (Constante), GINI_DESIGUALDAD.

B. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas.

C. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas, ind_felicidad.

D. Variables predictoras: (Constante), gini_desigualdad, disper_lenguas, ind_felicidad, ind_fragilidad.

E. Variables predictoras: (Constante), INCIDENCIA_TB.

De esta forma, la lectura de cada tabla nos permite rescatar los siguientes aspectos:

- La tabla *estadísticos* presenta los estadísticos descriptivos con los valores de las medias aritméticas y las desviaciones típicas de cada una de las variables incluidas en el modelo.
- En la tabla de *resumen del modelo* destacan los siguientes elementos: a) el modelo fue generado en cuatro pasos sucesivos, donde la variable más significativa incluida en el primer paso fue GINI_DESIGUALDAD. En la parte inferior de la tabla se señalan las otras variables incluidas en los siguientes pasos. Se puede apreciar que se eliminó del modelo la variable DISPER_ÉTNICA por no ser significativa al ajustarse con las otras variables; b) en la fila del cuarto paso se generan los valores adecuados al modelo final y allí se encuentran los coeficientes de determinación ($R^2 = 0.481$ y R^2 corregido = 0.463). En este caso se asumirá el R^2 corregido porque ajusta el sesgo del efecto de las variables sobre el coeficiente. Su valor indica una asociación lineal regular en el modelo, con 46% de las variaciones de la “Incidencia de Tb” explicadas por las variables independientes.
- Finalmente, en la tabla de ANOVA el valor de la F de Snedecor para el cuarto paso es de 27.114 y la probabilidad de obtener dicho valor en caso de que algún Beta (β) fuera 0 es de $p = 0.000$. Estos valores permiten rechazar la hipótesis nula y aceptar la hipótesis alterna de asociación lineal entre la variable dependiente (incidencia de Tb) y las variables independientes incluidas al modelo de regresión en el cuarto paso.

En forma adicional, se despliega en el visor de resultados una cuarta tabla que corresponde al modelo final de coeficientes de regresión beta (β) ajustados entre sí. Incluye los valores individuales de los errores típicos, de la t de Student y su significancia, y de los coeficientes de correlación (tabla 61).

Tabla 61
 Tabla que refiere el modelo de coeficientes de regresión beta (β) desplegada en el visor de resultados*

Modelo	Coeficientes no estandarizados		Sig.	t	Intervalo de confianza de 95.0% para B		Correlaciones		
	B	Error típico			Límite inferior	Límite superior	Orden cero	Parcial	Semiparcial
1	(Constante)	-234.197	53.094	-4.411	0.000	-339.319	-129.076		
	GINI_DESIGUALDAD	8.785	1.298	6.768	0.000	6.215	11.355	0.526	0.526
2	(Constante)	-184.870	48.454	-3.815	0.000	-280.814	-88.927		
	GINI_DESIGUALDAD	7.573	1.185	6.392	0.000	5.227	9.919	0.526	0.445
	Puntua:	62.931	11.451	5.496	0.000	40.256	85.605	0.474	0.383
	DISPER_Lenguas								
3	(Constante)	-31.596	72.087	-0.438	0.662	-174.348	111.157		
	GINI_DESIGUALDAD	7.937	1.159	6.847	0.000	5.642	10.233	0.526	0.533
	Puntua:	47.275	12.451	3.797	0.000	22.620	71.930	0.474	0.330
	DISPER_Lenguas								
4	IND_FELICIDAD	-3.909	1.391	-2.809	0.006	-6.664	-1.153	-0.332	-0.250
	(Constante)	23.412	75.055	0.312	0.756	-125.229	172.054		
	GINI_DESIGUALDAD	7.020	1.212	5.793	0.000	4.620	9.420	0.526	0.386
	DISPER_Lenguas	32.465	13.925	2.331	0.021	4.887	60.043	0.474	0.211
	IND_FELICIDAD	-4.245	1.377	-3.083	0.003	-6.972	-1.519	-0.332	-0.274
	IND_FRAGILIDAD	30.217	13.529	2.234	0.027	3.424	57.010	0.477	0.202

*Variable dependiente: INCIDENCIA_TB.

Los valores a interpretar en la tabla son los de la fila del paso 4 referente a la inclusión de variables (fila que nosotros sombrearemos). Allí aparecen los coeficientes β de cuatro variables independientes: índice Gini de desigualdad ($\beta_1 = 7.020$), dispersión de lenguas ($\beta_2 = 32.465$), índice de felicidad ($\beta_3 = -4.245$) e índice de fragilidad ($\beta_4 = 30.217$). Los cuatro coeficientes tienen significancias de t de Student menores a 0.05, por lo cual contribuyen significativamente a la explicación de la variable dependiente (incidencia de Tb). Tres de tales coeficientes tienen signo positivo y el incremento en cada una de sus unidades de medida se asocia al incremento de algunas unidades de la variable dependiente, mientras que la variable “índice de felicidad tiene signo negativo y el aumento en cada una de sus unidades se asocia al decremento de algunas unidades de la incidencia de Tb. Por ejemplo, en esta última variable, por cada unidad de incremento en el índice de felicidad, la incidencia de Tb disminuirá en -4.245 casos por 100.000.

Tres aspectos adicionales al leer la tabla son: a) cuando las medidas de las variables contempladas en el modelo son heterogéneas e interesa compararlas, la interpretación puede basarse en los coeficientes beta (β) estandarizados con valores de 0 a 1; b) los intervalos de confianza de 95% confirman la significancia si sus límites inferior y superior no tienen signos diferentes y, también, el sentido de la relación lineal por el signo de sus límites (ver por ejemplo que en la variable “índice de felicidad” ambos límites tienen signo negativo); y c) los valores de las correlaciones de orden cero son bivariadas y las de correlación parcial y semiparcial son ajustadas por el resto de las variables. Estas últimas actúan en el orden de inclusión de variables por el *método de pasos sucesivos*.

Del modelo al análisis de residuos

Luego de llevar a cabo la estimación del modelo de regresión, se procede a la tercera fase de análisis que consiste en el estudio

detallado de los residuos (E). Estos se definen como la diferencia entre los valores observados y estimados por el modelo para la variable dependiente.

En términos generales, el estudio de los residuos permite diagnosticar si el modelo construido cumple o no con las siguientes asunciones del análisis de regresión: a) la relación lineal entre las variables independientes y la dependiente; b) la normalidad de la distribución de los residuos; c) la igualdad de varianzas de los residuos (homocedasticidad); y d) la relevancia de las variables independientes (que no existan autocorrelación ni colinealidad en los residuos).

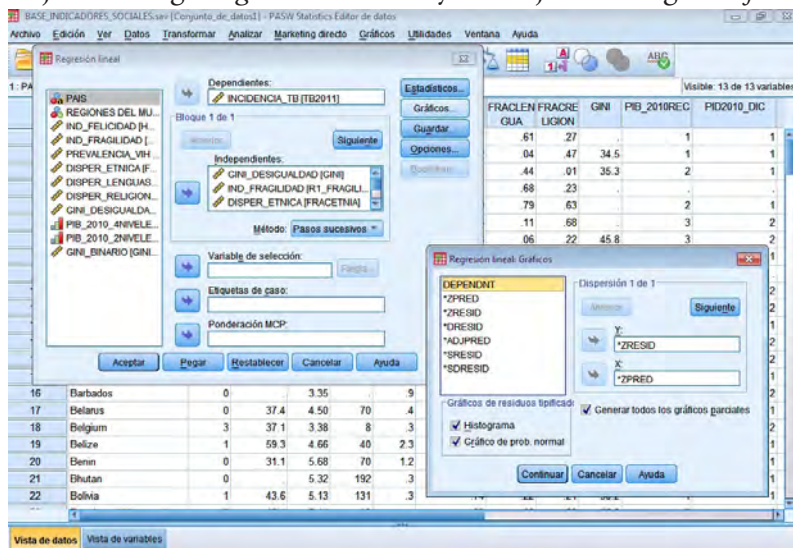
A continuación se realizará el estudio de los residuos de nuestro modelo, en dos momentos: a) en primer lugar, se diagnosticarán las asunciones de linealidad, normalidad y homocedasticidad, mediante procedimientos gráficos; y b) en segundo lugar, se diagnosticarán las asunciones de autocorrelación y colinealidad en los residuos.

El primer momento de análisis (diagnóstico de las asunciones de linealidad, normalidad y homocedasticidad) se procesa con el paquete SPSS de la siguiente manera:

En la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Regresión/Lineales*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles y las previamente seleccionadas como *Dependiente* e *Independientes*. Se ingresa a *Gráficos* y aparece una sub caja de diálogo donde se elige, primero, la variable *ZRESID (residuo tipificado) y se traslada al espacio de Y. Después, se elige la variable *ZPRED (valor pronosticado tipificado) y se mueve al espacio de X. Luego se marcan tres opciones de “gráficos de residuos tipificados”: *Histograma*, *Gráfico de Prob. Normal* Y *Generar todos los gráficos parciales*. Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 56).

Figura 56

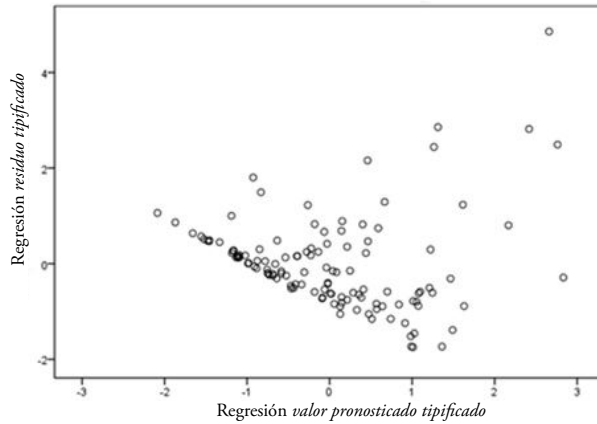
Caja de diálogo *Regresiones lineales* y sub caja de diálogo *Gráficos*



En el visor de resultados se despliega, primero, el gráfico global de dispersión de los residuos tipificados frente a los valores pronosticados tipificados. Este gráfico nos permite contrastar las hipótesis de linealidad y homocedasticidad de la regresión al apreciar la forma de la distribución de sus puntos en el cuadrante. El gráfico obtenido muestra una nube de puntos con una forma aleatoria, dispersa, sin una tendencia clara, lo cual nos lleva a aceptar las asunciones de linealidad y homocedasticidad en nuestro modelo (figura 57).

Figura 57

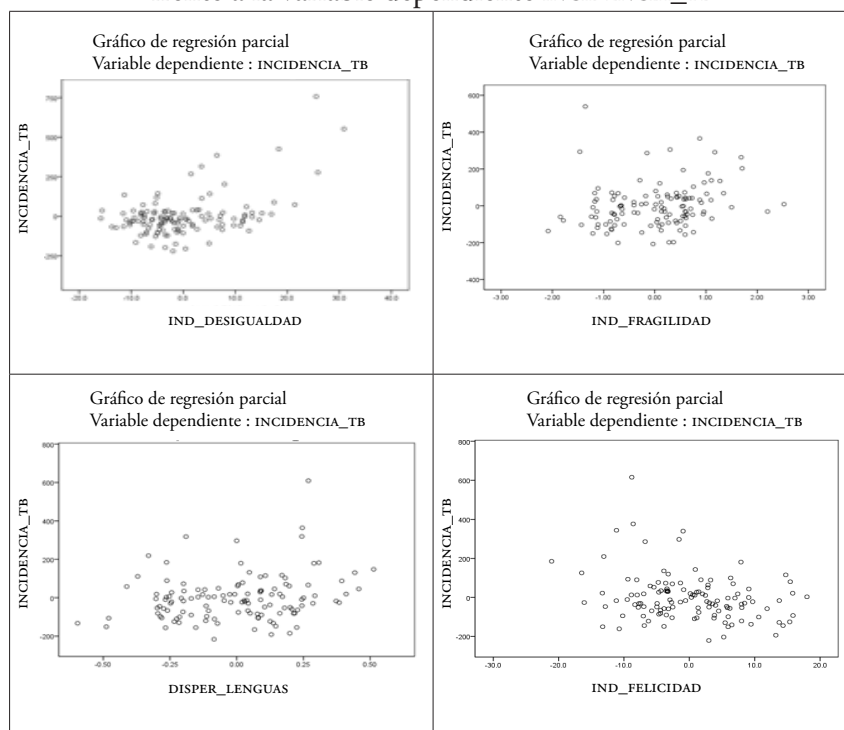
Gráfico de dispersión donde se muestran los residuos tipificados frente a los valores pronosticados tipificados*



*Variable dependiente: INCIDENCIA_TB.

En forma complementaria se despliegan los gráficos de dispersión parciales de cada una de las variables independientes del modelo. Al tener todas las nubes de puntos una forma aleatoria, permiten aceptar la hipótesis de la linealidad de cada variable frente a la variable dependiente (figura 58).

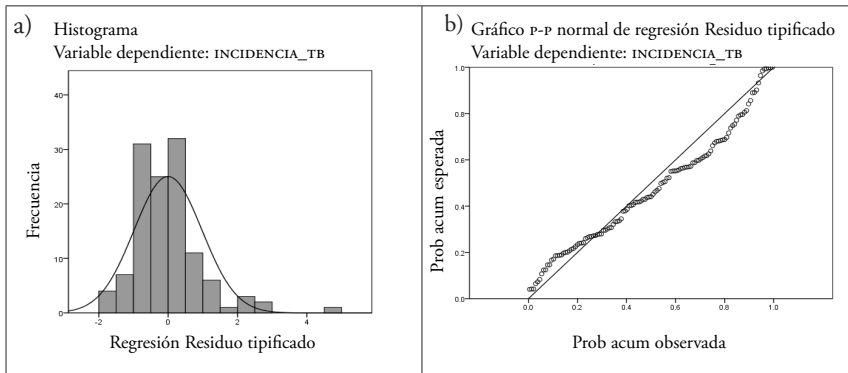
Figura 58
Gráficos de residuos parciales de cada variable del modelo
frente a la variable dependiente *INCIDENCIA_TB*



Fuente: Elaboración propia.

Finalmente, en el visor de resultados se despliegan un histograma y un gráfico probabilístico normal, ambos nos permitirán contrastar la hipótesis de normalidad de los residuos. Por una parte, la forma de la distribución del histograma se acerca a la de una campana y, por otra, los valores del gráfico P-P probabilístico normal tienden a agruparse alrededor de la línea diagonal de regresión. Ambas formas gráficas nos permiten aceptar la hipótesis de la normalidad (figura 59).

Figura 59
Gráficos histograma (a) y P-P probabilístico normal
de regresión (residuo tipificado [b])*



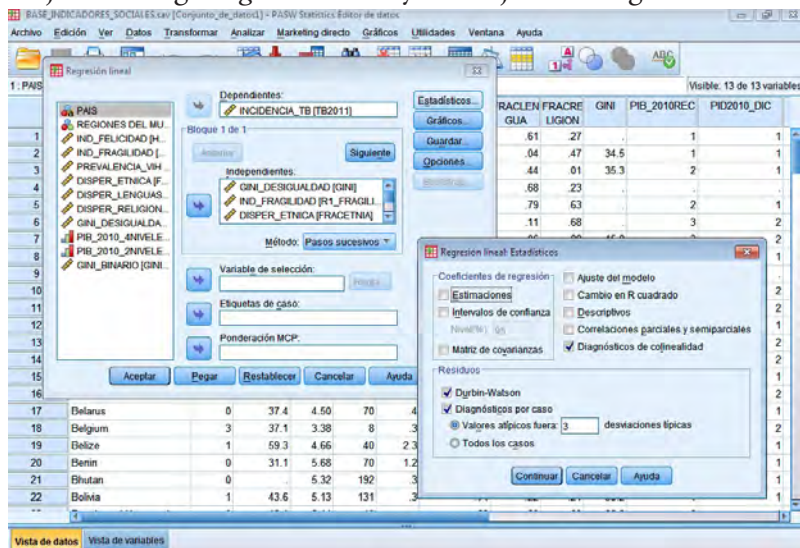
*Variable dependiente: INCIDENCIA_TB.

Fuente: Elaboración propia.

El segundo momento del análisis de residuos implica diagnosticar las asunciones de autocorrelación y colinealidad, a fin de asegurar que las variables incluidas en el modelo sean relevantes. El diagnóstico lo podemos realizar con el paquete SPSS de la siguiente manera: en la barra del menú que aparece en el SPSS se sigue la secuencia *Analizar/Regresión/Lineales*. Se despliega la caja de diálogo que contiene la lista de las variables disponibles y las previamente seleccionadas como *Dependiente* e *Independientes*. Se ingresa a *Estadísticos* y aparece una sub caja de diálogo donde se eligen *Diagnósticos de colinealidad*, *Durbin-Watson* y *Diagnósticos por caso* (con valores atípicos fuera de 3 desviaciones típicas). Finalmente se eligen *Continuar* y *Aceptar* para ejecutar el procedimiento (figura 60).

Figura 60

Caja de diálogo *Regresión lineal* y sub caja de diálogo *Estadísticos*



En el visor de resultados se despliegan primero las tablas *Resumen del modelo* y *Diagnósticos por caso* (tablas 62 y 63). Sus datos nos permiten interpretar lo siguiente:

1. La tabla *resumen del modelo* contiene el valor del estadístico Durbin-Watson que nos permite contrastar la hipótesis de autocorrelación. Su valor calculado es de 1.766 dentro de un rango de 0 a 4. Al acercarse a 2 se acepta la hipótesis de no autocorrelación (si los valores se hubieran acercado a 0 o 4 se hubiera aceptado la hipótesis de la autocorrelación).
2. La tabla *diagnósticos por caso* indica que el caso 160 es un valor atípico (se ubica a más de 3 veces la desviación típica por elevadísima tasa de incidencia de Tb) que podría generar errores y no ajustarse al modelo. Ante esta situación, correspondería estudiar la posibilidad de eliminarlo o no del análisis para obtener un mejor modelo.

Tabla 62
Resumen del modelo*

Modelo	Durbin-Watson
4	1.766

*Variables predictoras: (Constante), gini_desigualdad, disper_lenguas, ind_felicidad, ind_fragilidad. Variable dependiente: incidencia_tb.

Tabla 63
Diagnósticos por caso*

Número de casos	Residuo típico	INCIDENCIA_TB	Valor pronosticado	Residual
160	4.856	993	414.74	578.258

*Variable dependiente: INCIDENCIA_TB.

También se despliegan otras dos tablas que permiten contrastar la hipótesis de colinealidad: la tabla de los estadísticos *Tolerancia* y *factor de inflación de la varianza* (FIV [tabla 64]) y *Diagnósticos de colinealidad* (tabla 65). La primera muestra los valores de los estadísticos *tolerancia* y FIV para cada variable del modelo. Estos valores se interpretan en forma recíproca para decidir si se acepta o no la hipótesis de la colinealidad, como se muestra en la tabla 65. Dado que en nuestro modelo, los valores de *tolerancia* fueron altos (cerca de 1) y, a la vez, los de FIV bajos (menores a 2), decidimos aceptar la hipótesis nula o la no colinealidad.

Tabla 64
Coeficientes*

Modelo		Estadísticos de colinealidad	
		Tolerancia	FIV
1	GINI_DESIGUALDAD	1.000	1.000
2	GINI_DESIGUALDAD	0.965	1.036
	DISPER LENGUAS	0.953	1.036
3	GINI_DESIGUALDAD	0.772	1.049
	DISPER LENGUAS	0.799	1.295
	IND_FELICIDAD	0.799	1.251
4	GINI_DESIGUALDAD	0.844	1.185
	DISPER LENGUAS	0.597	1.675
	IND_FELICIDAD	0.790	1.266
	IND_FRAGILIDAD	0.645	1.551

*Variable dependiente: INCIDENCIA_TB.

Tabla 65

Valores recíprocos de los estadísticos de *tolerancia* y FIV en el contraste de las hipótesis estadísticas de colinealidad (datos ficticios)

Hipótesis	Tolerancia (en función de 1)	FIV (en función de 2)
H ₀ . No colinealidad	Valores altos (ceranos a 1)	Valores bajos (menores a 2)
	0.58	1.13
	0.67	1.72
	0.81	1.65
H ₁ . Colinealidad	Valores bajos (ceranos a 0)	Valores altos (mayores a 2)
	0.12	18.12
	0.13	9.26
	0.36	4.13

Fuente: Elaboración propia.

En segundo lugar se despliega una tabla de *diagnósticos de colinealidad* (tabla 66). Esta tabla establece cuáles son las variables del modelo que presentan colinealidad. Describe la relación entre el índice de condición y las proporciones de la varianza de las variables incluidas en el modelo como indicador de colinealidad. Así, cuando existen en una misma fila dos variables con proporciones de varianza ≥ 0.6 y un índice de condición con valor elevado, es señal de colinealidad entre las dos variables. En nuestro ejemplo —parte sombreada para el paso cuatro— no encontramos que dos variables tengan proporciones de varianza ≥ 0.6 en la misma fila ni tampoco índices de condición elevados, por lo que aceptamos la hipótesis nula o la no colinealidad entre variables.

Tabla 66
Diagnósticos de colinealidad*

Modelo	Dimensión	Autovalores	Índice de condición	Proporciones de la varianza				
				(Constante)	GINI_DESI-GUALDAD	DISPER LENGUAS	IND_FELI-CIDAD	IND_FRAGI-LIDAD
1	1	1.972	1.000	0.01	0.01			
	2	0.028	8.326	0.99	0.99			
2	1	2.728	1.000	0.01	0.01	0.04		
	2	0.243	3.350	0.04	0.03	0.96		
	3	0.028	9.798	0.96	0.96	0.00		
3	1	3.638	1.000	0.00	0.00	0.02	0.00	
	2	0.307	3.440	0.00	0.00	0.67	0.02	
	3	0.041	9.367	0.02	0.88	0.12	0.20	
	4	0.013	16.665	0.98	0.11	0.20	0.78	
4	1	4.613	1.000	0.00	0.00	0.01	0.00	0.00
	2	0.308	3.873	0.00	0.00	0.52	0.02	0.00
	3	0.042	10.522	0.02	0.71	0.12	0.22	0.01
	4	0.025	13.584	0.00	0.25	0.30	0.12	0.89
	5	0.012	19.286	0.98	0.03	0.05	0.64	0.11

*Variable dependiente: INCIDENCIA_TB.

En síntesis, las tres fases del análisis de regresión lineal nos permitieron construir un modelo de variables independientes predictoras de la “incidencia de Tb” que cumple con todas las asunciones técnicas.

En el caso de la variable “prevalencia del VIH”, se construyó el modelo de regresión que se presenta en la parte inferior de la tabla 67. Para ello se siguieron las mismas tres fases del análisis ilustradas en esta sección (procedimiento no mostrado). En dicho modelo, el índice Gini de desigualdad y el índice de felicidad fueron factores predictores independientes de la “prevalencia del VIH”, y el índice Gini como factor de riesgo de una mayor probabilidad y el índice de felicidad como factor protector. El modelo fue estadísticamente significativo (ANOVA $P < 0.01$) y explicó 38% de la variabilidad de la variable dependiente.

Tabla 67
Modelos de regresión lineal sobre variables predictoras
de la incidencia de Tb y la prevalencia del VIH, en el mundo *

Modelos y variables predictoras ^a	Coefficiente β [IC de 95%]	p	R ² corregido	Cambio en R ²
<i>Modelo 1: Predictores de la incidencia de Tb^b</i>				
Índice Gini de desigualdad	7.02 [4.62 a 9.42]	<0.01	0.46	0.02 (p<0.05)
Dispersión de lenguas	32.46 [4.89 a 60.04]	<0.05		
Índice de felicidad	-4.24 [-6.97 a -1.52]	<0.05		
Índice de fragilidad	30.22 [3.42 a 57.01]	<0.05		
Constante	23.41 [-125.23 a 172.05]	0.76		
<i>Modelo 2: Predictores de prevalencia de VIH^c</i>				
Índice Gini de desigualdad	0.19 [0.13 a 0.25]	<0.01	0.38	0.16 (p<0.05)
Índice de felicidad	-0.17 [-0.23 a -0.11]	<0.01		
Constante	1.44 [-2.09 a 4.97]	0.42		

*En el año 2010.

^a Variables incluidas en los modelos iniciales: índice Gini de desigualdad, índice de felicidad, índice de fragilidad, dispersión étnica y dispersión de lenguas.

^b Variables significativas en el paso 4: F (1, 117) = 27.11; p <0.01; Durbin-Watson = 1.77.

^c Variables significativas en el paso 2: F (1, 107) = 34.60; p <0.01; Durbin-Watson = 1.76.

Síntesis

El proceso de análisis de las variables *intervalo/razón* siguió etapas interrelacionadas que permitieron trabajar los datos desde su escrutinio más simple hasta una construcción de relaciones más compleja. Permitió explorar el contenido de las variables individuales, contrastar la relación de pares de variables mediante pruebas de diferencias de medias y de asociación, y, sobre esa base, generar dos modelos de regresión lineal.

La tabla 67 sintetiza los modelos construidos y detalles descriptivos que los contextualizan, de manera que se pueda realizar la comunicación de los resultados ante la comunidad científica con cierta precisión y rigor estadístico. Como se aprecia, la tabla es el resultado de todo el proceso de análisis desarrollado y su síntesis destaca los siguientes elementos: 1) los coeficientes de regresión β y los IC (de 95%) de cada variable; 2) el valor de significancia (p) de la prueba t de Student de cada variable independiente; 3) el coeficiente de determinación R^2 corregido; 4) el valor del cambio de R^2 en el paso en que se generó el modelo; 5) el valor de la significancia de ANOVA para el modelo; y 6) el valor de la prueba de contraste de colinealidad Durbin-Watson.

La lectura de la referida tabla (67) muestra que hay dos variables que fueron más consistentes que las otras en la predicción de las dos variables dependientes de los modelos: el índice Gini de desigualdad social y el índice de felicidad (en sentido inverso). Aunque la tuberculosis y el VIH son dos fenómenos muy relacionados entre sí, los hallazgos indican también que en la tuberculosis las variables de la pobreza (índice de fragilidad y dispersión de lenguas) son factores que refuerzan el riesgo en forma independiente. Los valores de ANOVA señalan que los modelos son significativos y la R^2 corregidas que los modelos explican 46% y 38% de la varianza respecto a las variables dependientes.

El procedimiento seguido para lograr un resultado complejo, permitió compenetrarse a fondo con los datos y aprove-

char parte de su riqueza en función de las hipótesis de nuestro interés. Sin embargo, el análisis no agotó toda la potencialidad de los datos y podría considerarse profundizar en la búsqueda de nuevos elementos de las relaciones (ésta es la esencia del enfoque de la minería de datos). Por ejemplo, se podría generar un Análisis de Camino (*Path Analysis*) para representar un modelo de efectos causales, a partir de los coeficientes de correlación parcial estandarizados que se calcularon en la regresión lineal (Tacq, 1997: 140-182).

Se trata de un análisis que puede ser generado desde la secuencia *Regresión/Lineal* del paquete SPSS, o recurriendo al paquete AMOS de IBM SPSS que se especializa en construir *Modelos de Ecuaciones Estructurales*. Nosotros no realizaremos este análisis y lo postergaremos como una posibilidad a seguir, debido a que nuestro interés se limitó a ilustrar el proceso de análisis de variables de intervalo/razón hasta llegar al modelo de regresión.

Palabras finales

Queremos hacer una reflexión general sobre el proceso seguido a lo largo de esta breve introducción al análisis de datos con el paquete estadístico SPSS.

El proceso de análisis propuesto tiene claramente un enfoque exploratorio donde se escarba en detalle la información simple para sustentar la posterior construcción de bloques sólidos de relaciones complejas. En ese intento, el analista adopta y hace interactuar dos actitudes: una de apego fiel a la norma estadística y al cumplimiento de los supuestos del procesamiento y análisis (véase figura 5), y otra de cierta apertura a la novedad en los patrones estadísticos y a la creatividad en la conducción de los procedimientos.

Con ello, el trabajo de procesamiento y análisis se convierte en una actividad de búsqueda constante de respuestas a preguntas e hipótesis estadísticas y de renovación de ideas para profundizar los hallazgos con nuevos procedimientos. En esa búsqueda, el paquete SPSS es una herramienta heurística que habilita para transformar

la información de las bases de datos originales y con ello favorecer el desarrollo de diversas estrategias de visualización y cálculo que permitan construir interpretaciones y nuevas hipótesis.

La propuesta de esta breve introducción al análisis se ciñó a indagar las relaciones de las bases de datos que sirvieron de ejemplo y buscó ilustrar las estrategias de análisis adecuadas a las variables disponibles. En ese afán, se limitó a ejemplificar análisis de datos transversales y, por ello, no abarcó estrategias para estudios longitudinales ni de diseños experimentales.

El autor considera que no hay una receta universal para todos los análisis posibles y que sería difícil reunir en una introducción todo el repertorio de estrategias que normalmente son recopiladas por sabios tratadistas. Esto es natural porque cada base de datos es, en sí misma, un universo complejo susceptible de ser explorado en función de las hipótesis que se quieran contrastar, mediante estrategias específicas de análisis.

Está claro que si el analista tiene escasos conocimientos de estadística y una débil conciencia de los supuestos de procesamiento y análisis que se aplican en el paquete de cómputo, difícilmente podrá aprovechar el potencial del mismo. Sin embargo, todo se aprende gradualmente y un aprendizaje histórico que tenemos es que el SPSS ha ayudado a varias generaciones de estudiantes y científicos a acercarse al mundo del análisis con procedimientos simples, lo cual generó un gran impacto en la proliferación de la investigación.

El autor de esta introducción, espera que tanto la propuesta del capítulo, como los ejemplos de análisis, motiven a los lectores a realizar la actividad de procesamiento y análisis con el SPSS. Una sugerencia es tomar algunas variables de interés que ya existan en una base de datos y seguir cada uno de los pasos propuestos como una guía. Así, observar qué resulta de ello: ojalá que un mayor interés por analizar datos con la mente puesta en plantear problemas y contrastar hipótesis estadísticas interesantes.

Glosario

Análisis de datos. Consiste en la aplicación sistemática de técnicas estadísticas orientadas a resumir su distribución, organizar sus relaciones y evaluar inductivamente los patrones y tendencias de su estructura.

Archivo de datos. Archivo electrónico que contiene los datos con códigos numéricos.

Caja de diálogo del SPSS. Es un cuadro rectangular que presenta al usuario el repertorio de variables disponibles y de opciones a seguir para correr un procedimiento de análisis estadístico.

Datos. Observaciones traducidas a valores con códigos numéricos.

Editor de datos del SPSS. Ventana donde despliega la matriz de definiciones que contienen la especificación de variables, sus nombres y tipos de escala.

Hipótesis. Una proposición sobre la naturaleza de una relación esperada entre dos o más variables.

Medición. Es un proceso deductivo que implica tomar un concepto o idea y desarrollar una medida para su observación empírica basada en escalas numéricas con principios matemáticos. Las escalas más usadas son las propuestas por Stevens: nominal, ordinal e intervalo/razón.

Minería de datos. Estrategia que persigue lograr el descubrimiento automático del conocimiento contenido en la información almacenada en las bases de datos, mediante técnicas explicativas o de dependencia y técnicas descriptivas o de interdependencia.

Pruebas no paramétricas. Pruebas estadísticas de significancia las cuales no requieren que las variables de análisis cumplan las asunciones de normalidad y de homogeneidad de varianzas.

Pruebas paramétricas. Pruebas estadísticas de significancia que asumen una distribución normal en las variables de análisis.

Realidad. Un fenómeno que se puede percibir sólo de manera indirecta, a través de observaciones basadas en sistemas de medición.

Sistema Windows. Plataforma de cómputo o sistema operativo que despliega gráficamente un menú de operaciones disponibles.

Visor de datos del SPSS. Ventana donde se despliega la matriz de datos que contiene los nombres y valores de las variables (columnas) correspondientes a cada caso (filas).

Visor de resultados del SPSS. Ventana donde se despliegan los resultados de un procedimiento de análisis.

Autoevaluación

¿Qué es el análisis de datos?

¿Cuáles son los principales supuestos que presenta el análisis de datos computarizado?

¿Qué es el paquete estadístico SPSS y cómo contribuyó a que la disciplina del análisis de datos tenga un mayor desarrollo y popularidad en el mundo?

¿Cuáles son los procedimientos de cómputo básicos para realizar análisis exploratorios?

¿Cuáles son los procedimientos de cómputo básicos para realizar análisis univariados, según tipos de variables (nominal, ordinal, intervalo/razón)?

¿Cómo se procesan las pruebas de asociación *Chi-cuadrado* y *correlación*, y qué tipos de variables requieren?

¿Cómo se procesan las pruebas de diferencia de medias T y ANOVA, y qué tipos de variables requieren?

¿Qué criterios se siguen para desarrollar análisis multivariados cuando las variables son nominales y ordinales?

¿Qué criterios se siguen para desarrollar análisis multivariados cuando las variables son de intervalo/razón?

¿Qué es el enfoque de la minería de datos y cómo se aplica en el análisis computarizado?

Bibliografía

- Bryman, A., y Cramer, D. (2009). *Quantitative Data Analysis with SPSS 14, 15 & 16. A guide for Social Scientists*. London: Routledge.
- Hardy, M., y Bryman, A. (eds.) (2004). *Handbook of Data Analysis*. London: Sage Publications.
- Hardy, M. (2004). Summarizing distributions. En: M. Hardy y A. Bryman (eds.), *Handbook of Data Analysis* (pp. 35-64). London: Sage Publications.
- Lewis-Beck, M.S. (1995). *Data Analysis. An Introduction*. Thousand Oaks, CA: Sage Publications.
- Neuman, W.L. (1997). *Social Research Methods. Qualitative and Quantitative approaches*. Boston, MA: Allyn and Bacon.
- Norusis, M.J. (2011). *IMB SPSS Statistics 19 Guide to Data Analysis*. Upper Saddle River, NJ: Prentice Hall Press.
- Pérez, C. (2009). *Técnicas de análisis de datos con SPSS 15*. Madrid: Pearson/Prentice Hall.
- Tacq, J. (1997). *Multivariate Analysis Techniques in Social Science Research*. London: Sage Publications.

CAPÍTULO X

Medicina basada en evidencias aplicada a la neurología

Rebeca Millán Guerrero

La medicina basada en evidencias

La Medicina Basada en Evidencias (MBE) es la integración de las mejores evidencias de la investigación, con la experiencia clínica y los valores del paciente; éstos se explican a continuación (Sackett, *et al.*, 2000; Candelise, 2007):

- *Mejor evidencia de la investigación.* Se entiende como la investigación clínica centrada en el paciente, con exactitud y precisión de las pruebas diagnósticas (incluyendo exploración clínica), marcadores pronósticos, eficacia, seguridad de terapéutica, rehabilitación y prevención. Las nuevas evidencias de la investigación invalidan lo aceptado previamente y son remplazados por otros nuevos, más potentes, más exactos, más eficaces y más seguros.
- *Experiencia clínica.* Se entiende por ésta la capacidad de utilizar nuestras habilidades clínicas y experiencia del pasado, para identificar rápidamente el estado de salud y diagnóstico específico de cada paciente, sus riesgos-beneficios de posibles acciones, así como valores personales y expectativas.

- *Valores del paciente.* Se entiende como las preferencias, preocupaciones y expectativas del paciente, las que se deben integrar a la decisión médica.

Cuando estos tres elementos se integran, los clínicos y los pacientes forman una alianza diagnóstica y terapéutica que optimiza los resultados clínicos y la calidad de vida.

Estas ideas no son nuevas, han estado presentes durante mucho tiempo en la medicina china antigua, en Francia, etcétera. El concepto de MBE fue desarrollado por un grupo de internistas y epidemiólogos clínicos, liderados por Gordon Guyatt, de la Escuela de Medicina de la Universidad McMaster de Canadá (*Evidence-Based Medicine Working Group*, 1992). El mismo grupo que años atrás incluyó el término *epidemiología clínica*. El propósito de crear la MBE fue para que toda acción médica de diagnóstico, pronóstico y terapéutica, tenga evidencia cuantitativa sólida, basada en la mejor investigación epidemiológica clínica.

En palabras de David Sackett (1996: 71-72), “la MBE es la utilización consciente, explícita y juiciosa de la mejor evidencia clínica disponible para tomar decisiones sobre el cuidado de los pacientes individuales”. En esencia, la MBE pretende aportar más ciencia al arte de la medicina, siendo su objetivo disponer de la mejor información científica disponible (la evidencia) para aplicarla a la práctica clínica (Guerra-Romero, 1996).

La utilidad de la MBE en la neurología

En la práctica diaria, el neurólogo ocupado tiene que tomar decisiones acerca de intervenciones terapéuticas, como por ejemplo: timectomía o no para *miastenia gravis*, valproato o carbamazepina en fase temprana de epilepsia parcial, levodopa o agonistas dopaminérgicos en fase temprana de Parkinson, penicilina para meningitis por meningococo, etcétera. Es posible que, para un experto, esas dudas sean fáciles de resolver. Sin embargo, para un neurólogo joven, o cuando han surgido nuevos tratamientos,

es necesario que, experto o no, se vea en la necesidad de revisar la literatura y buscar la mejor terapéutica demostrada con evidencia (Sackett, *et al.*, 2000; Candelise, 2007).

¿Cómo practicamos realmente la MBE?

- *Paso 1:* La necesidad de información (sobre prevención, diagnóstico, pronóstico, terapia, causalidad, etcétera) conduce a generar una pregunta.
- *Paso 2:* Buscar la mejor evidencia para contestar esa pregunta. Con frecuencia surge dificultad para la búsqueda de esta respuesta, porque en ocasiones falta información y se tiene dificultad, que paulatinamente ha sido resuelta con la ayuda de satélites y del archivo Cochrane que cuenta con las revisiones más completas y actuales.
- *Paso 3:* Evaluar en forma crítica la validez de la evidencia encontrada en impacto y aplicabilidad, es decir: ¿la podemos aplicar en la práctica clínica?
- *Paso 4:* Integrar los datos encontrados para nuestra especialidad clínica.
- *Paso 5:* Evaluar los resultados de los pasos (Sackett, *et al.*, 2000; Candelise, 2007).

¿Se pueden medir los resultados?

La medición del resultado depende de la pregunta clínica y se debe tener cuidado al interpretarlo. ¿Será lo mismo significancia estadística y significancia clínica? Habrá resultados que muestren una significancia estadística (IC de 95% para un valor de $p < 0.05$), pero no tienen significancia clínica, porque mientras que la significancia estadística está directamente relacionada con el tamaño de la muestra, y puede ser aumentada al aumentar el número de sujetos estudiados, la significancia clínica no. Como ejemplo tenemos a *donepezil* (*Aricept*) para disminuir la pro-

gresión de demencia: la evidencia en un metanálisis de Cochrane muestra mejoría estadística significativa de 2.9 puntos (ADAS-cog escala) a las 24 semanas con 10 mg de donepezil; éstos clínicamente son cambios menores a 5% de la escala: ¿tendrá significancia clínica? (Candelise, 2007).

Entonces, ¿cómo se deben interpretar los resultados? Al revisar la literatura, el médico debe tener cuidado al estudiar el diseño y resultados, ver sesgos del estudio, sesgos de selección, si se realizó aleatorización, y ver el grupo control. Asimismo si hubo cegamiento adecuado, empleo de placebo inactivo o activo, etcétera. Es decir, tener en cuenta conocimientos de metodología.

Para cada patología, el estudio sistemático debe adecuarse a la pregunta. Por ejemplo, en migraña (Candelise, 2007), si queremos buscar lo siguiente:

- Prevalencia: debemos buscar en la literatura los estudios epidemiológicos relacionados en determinado país, o estudios multicéntricos; además, que cumplan con un adecuado tamaño de muestra, selección de pacientes, etcétera.
- Método diagnóstico: ¿existe alguna prueba de oro para esta enfermedad?, ¿cuál es el mejor método diagnóstico hasta el momento?
- Terapéutica del evento agudo: buscar estudios con la mejor evidencia en ensayos clínicos; asimismo buscar guías de la *International Headache Society*, a pesar de que existen críticas a las guías y algoritmos, porque para algunos médicos, limitan la acción clínica (Lambert, 2006). ¿Qué se recomienda para tratar el evento agudo?, ¿cuáles son los mejores abortivos?
- Terapéutica en profilaxis: búsqueda semejante para el caso del punto previo. ¿Cuál es el mejor fármaco y cuanto tiempo lo debo emplear?
- Genética en migraña, etcétera.

Diferentes escalas de gradación para medir la MBE

En función del rigor científico que se emplee en el diseño de los estudios, pueden construirse escalas de clasificación jerárquica de la evidencia, a partir de las cuales pueden establecerse recomendaciones respecto a la adopción de un determinado procedimiento médico o intervención sanitaria (Jovel y Navarro-Rubio, 1995; Guyatt *et al.*, 1995). Aunque hay diferentes escalas de gradación para medir la calidad de la evidencia científica, todas ellas son muy similares entre sí, algunas dirigidas a medicina general, otras a enfermedades neurológicas, relacionadas con prevalencia, con método diagnóstico, la mejor terapéutica, etcétera:

Se han creado diferentes escalas para medir la evidencia y están relacionadas con el tipo de diseño, con recomendaciones, etcétera. En la siguiente escala (tabla 68), el nivel de evidencia se califica del I al III, donde I es el mejor nivel, representado por el diseño de ensayo clínico.

Tabla 68
Jerarquía de los estudios por el tipo de diseño (USPSTF)

Nivel de evidencia	Tipo de estudio
I	Al menos un ensayo clínico controlado y aleatorizado, diseñado de forma apropiada.
II-1	Ensayos clínicos controlados, bien diseñados, pero no aleatorizados.
II-2	Estudios de cohortes o de casos-controles, bien diseñados, preferentemente multicéntricos.
II-3	Múltiples series comparadas en el tiempo con o sin intervención, y resultados sorprendentes en experiencias no controladas.
III	Opiniones basadas en experiencias clínicas, estudios descriptivos, observaciones clínicas o informes de comités de expertos.

Fuente: Harris *et al.* (2001: 21-35).

La tabla 69 muestra la calificación de recomendación para emplear algún método diagnóstico o terapéutico, la calificación es ordinal y se representa por letras, designando a la letra A como una buena evidencia.

Tabla 69
Establecimiento de las recomendaciones (USPSTF)

Calidad de la evidencia	Beneficio neto sustancial	Beneficio neto moderado	Beneficio neto pequeño	Beneficio neto nulo o negativo
Buena	A	B	C	D
Moderada	B	B	C	D
Mala	E	E	E	E

Fuente: Harris *et al.* (2001: 21-35).

Finalmente, se presenta la propuesta del *Centre for Evidence-Based Medicine* (CEBM) de Oxford, cuya escala tiene en cuenta no sólo las intervenciones terapéuticas y preventivas, sino también aquellas ligadas al diagnóstico, pronóstico, factores de riesgo y evaluación económica (tabla 70):

Tabla 70
Niveles de evidencia^a

Nivel de evidencia	Tipo de estudio
1a	Revisión sistemática de ensayos clínicos aleatorizados, con homogeneidad.
1b	Ensayo clínico aleatorizado con intervalo de confianza estrecho.
1c	Práctica clínica (“todos o ninguno”) (*).
2a	Revisión sistemática de estudios de cohortes, con homogeneidad.
2b	Estudio de cohortes o ensayo clínico aleatorizado de baja calidad (**).
2c	<i>Outcomes research</i> (***), estudios ecológicos.
3a	Revisión sistemática de estudios caso-control, con homogeneidad.
3b	Estudio caso-control.
4	Serie de casos o estudios de cohortes y caso-control de baja calidad (****).
5	Opinión de expertos sin valorización crítica explícita, o basados en la fisiología, <i>bench research</i> o <i>first principles</i> (*****).

Fuente: Saha *et al.*, 2001; Harbour y Miller, 2001: 334-336.

^aEn esta calificación, de nuevo, se toma en cuenta el diseño del estudio y se señalan las siguientes recomendaciones: Se debe añadir un signo menos (-) para indicar que el nivel de evidencia no es concluyente: 1) Ensayo clínico aleatorizado con intervalo de confianza amplio y no estadísticamente significativo; 2) Revisión sistemática con heterogeneidad estadísticamente significativa.

*Cuando todos los pacientes mueren antes de que un determinado tratamiento esté disponible, y con él algunos pacientes sobreviven, o bien cuando algunos pacientes morían antes de su disponibilidad, y con él no muere ninguno.

**Por ejemplo, con seguimiento inferior a 80%.

***El término *outcomes research* hace referencia a estudios de cohortes de pacientes con el mismo diagnóstico en los que se relacionan los eventos que suceden con las medidas terapéuticas que reciben.

****Estudio de cohorte: sin clara definición de los grupos comparados y/o sin medición objetiva de las exposiciones y eventos (preferentemente ciega) y/o sin identificar o controlar adecuadamente variables de confusión conocidas y/o sin seguimiento completo y suficientemente prolongado. Estudio caso-control: sin clara definición de los grupos comparados y/o sin medición objetiva de las exposiciones y eventos (preferentemente ciega) y/o sin identificar o controlar adecuadamente variables de confusión conocidas.

*****El término *first principles* hace referencia a la adopción de determinada práctica.

Conclusión

La Medicina Basada en Evidencias (MBE) es la integración de las mejores evidencias de la investigación, donde se toman en cuenta el método científico, sin olvidar la experiencia laboral, que permite al profesional considerar los años de trabajo.

Glosario

Medicina Basada en Evidencias. Es la integración de las mejores evidencias de la investigación, con la experiencia clínica y los valores del paciente.

Autoevaluación

- ¿Qué es la Medicina Basada en Evidencias (MBE)?
- ¿Quiénes desarrollaron la MBE y con qué finalidad?
- ¿Cuáles son los pasos para desarrollar la MBE?
- ¿Por qué se plantea que es necesario distinguir entre significancia clínica y significancia estadística cuando se interpretan las evidencias de los estudios?

Bibliografía

- Candelise, L. (ed.) (2007). *Evidence-based Neurology: Management of Neurological Disorders*. Oxford: Blackwell Publishing.
- Hughes, R., Liberati, A., y Uitdehaag, B.M.J. (2007). Evidence-based medicine: its contributions in the way we search, appraise and apply scientific information to patient care. En: L. Candelise (ed.), *Evidence-based Neurology: Management of Neurological Disorders* (pp. 2-10). Oxford: Blackwell Publishing.
- Evidence-Based Medicine Working Group (1992). Evidence-based medicine. A new approach to teaching the practice of medicine. En: *Journal of American Medical Association*, 268, pp. 2420-2425.
- Guerra-Romero, L. (1996). La medicina basada en la evidencia: un intento de acercar la ciencia al arte de la práctica clínica. En: *Medicina Clínica*, 107, pp. 377-382.
- Guyatt, G.H.; Sackett, D.L.; Sinclair, J.C., *et al.* (1995). Users' Guides to the Medical Literature: IX. A method for grading health care recommendations. En: *Journal of American Medical Association*, 274, pp. 1800-1804.

- Harbour, R. y Miller, J. (eds.) (2001). Scottish Intercollegiate Guidelines Network Grading Review Group. En: *British Medical Journal*, 323, pp. 334-336.
- Harris, R.P.; Helfand, M.; Woolf, S.H., *et al.* (2001). Current methods of the U.S. Preventive Services Task Force: a review of the process. En: *American Journal of Preventive Medicine*, 20(3S), pp. 21-35.
- Jovell, A.J., y Navarro-Rubio, M.D. (1995). Evaluación de la evidencia científica. En: *Medicina Clínica*, 105, pp. 740-743.
- Lambert, H. (2006). Accounting for EBM: Notions of evidence in medicine. En: *Social Science & Medicine*, 62, pp. 2633-2645.
- Sackett, D.L.; Rosenberg, W.M.C.; Gary, J.A.M., *et al.* (1996). Evidence based medicine: what is it and what it isn't. En: *British Medical Journal*, 312, 71-72.
- Sackett, D.L.; Straus, S.E.; Richardson, W.S.; Rosenberg, W., y Hynes, R.B. (2000). *Medicina Basada en la Evidencia. Cómo practicar y enseñar la MBE*. Barcelona: Churchill-Livingstone.
- Saha, S.; Hoerger, T.J.; Pignone, M.P., *et al.* (2001). The art and science of incorporating cost effectiveness into evidence-based recommendations for clinical preventive services. En: *American Journal of Preventive Medicine*, 20(3S), pp. 36-43.

Introducción a la epidemiología clínica y estadística, de Rebeca Millán Guerrero, Benjamín Trujillo Hernández y José Ramiro Caballero Hoyos, fue editado en la Dirección General de Publicaciones de la Universidad de Colima, avenida Universidad 333, Colima, Colima, México, www.ucol.mx. La impresión se terminó en mayo de 2015 con un tiraje de 500 ejemplares sobre papel bond ahuesado de 90 gramos para interiores y sulfatada de 12 puntos para la portada. En la composición tipográfica se utilizó la familia Adobe Garamond Pro. El tamaño del libro es de 22.5 cm de alto por 16 cm de ancho. Programa Editorial: Alberto Vega Aguayo. Gestión Administrativa: Inés Sandoval Venegas. Diseño: José Luis Ramírez Moreno. Corrección y cuidado de la edición: José Augusto Estrella.

El presente libro es una herramienta didáctica para el aprendizaje de la epidemiología clínica. Con un lenguaje sencillo y sin tecnicismos, los autores exponen distintos elementos teóricos y empíricos que ayudarán al lector a incorporar el método científico y los principios de la estadística en el ejercicio de su práctica clínica. En tal sentido, desarrollan en el texto la relación entre epidemiología clínica y estadística, la aplicación del método científico, la elaboración del protocolo de investigación, la elección de diseños de investigación clínica, la relación entre causalidad y riesgo, la medición de los factores de riesgo, el diseño de la muestra, la selección de pruebas de hipótesis, el análisis de datos con el paquete SPSS y la concepción de la medicina basada en evidencias.

Rebeca Millán Guerrero

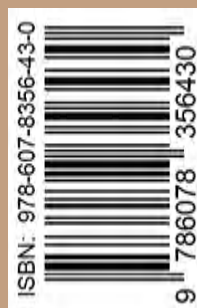
Mexicana. Doctora en ciencias médicas. Nivel II del Sistema Nacional de Investigadores (SNI). Investigadora Titular A del Instituto Mexicano del Seguro Social (IMSS). Profesora B en la Universidad de Colima. Línea de investigación: neurociencias.

Benjamín Trujillo Hernández

Mexicano. Doctor en ciencias médicas. Nivel II del SNI. Investigador Titular A del IMSS. Profesor B de la Universidad de Colima. Línea de investigación principal: epidemiología clínica y molecular de enfermedades crónico-degenerativas.

José Ramiro Caballero Hoyos

Boliviano. Doctor en sociología. Nivel II del SNI. Investigador Asociado D del IMSS. Profesor de la Universidad de Colima. Líneas de investigación: comunicación y riesgo epidemiológico.



UNIVERSIDAD DE COLIMA