

CAPÍTULO

8

ANÁLISIS ESTADÍSTICO: SEGUNDA PARTE

INTRODUCCIÓN

Con este capítulo se complementa el 10 de Metodología de la investigación, 5ª edición, además de que se actualizó su contenido. Se presentan primero las hipótesis estadísticas, las puntuaciones z , cálculos inferenciales o estimaciones de parámetros y luego el cálculo del coeficiente alfa-Cronbach y el sustento del análisis de varianza unidireccional. Finalmente los análisis multivariados y algunas fórmulas, así como una secuencia de análisis en Minitab y otra en SPSS. Los autores asumimos que se revisó previamente el capítulo en cuestión.

HIPÓTESIS ESTADÍSTICAS

En el capítulo 10 se señaló que cada prueba estadística obedece a un tipo de hipótesis de investigación e hipótesis estadística distinta. Las hipótesis estadísticas son la transformación de las hipótesis de investigación, nulas y alternativas en símbolos

estadísticos. Algunas revistas académicas solicitan que se incluyan estas hipótesis y ayudan a conceptualizar ciertas pruebas revisadas en el capítulo 10 del texto impreso.

En ocasiones, el investigador traduce su hipótesis de investigación y nula (y cuando se formulan hipótesis alternativas, también éstas) en términos estadísticos. Básicamente hay tres tipos de hipótesis estadísticas, que corresponden a clasificaciones de las hipótesis de investigación y nula: 1) de estimación, 2) de correlación y 3) de diferencias de medias. A continuación mencionaremos y daremos ejemplos de cada una de ellas.

Hipótesis estadísticas de estimación

Corresponden a las que, al hablar de hipótesis de investigación, se les denomina hipótesis descriptivas de un dato que se pronostica. Sirven para evaluar la suposición de un investigador respecto del valor de alguna característica en una muestra de individuos, otros seres vivos, sucesos u objetos, y en una población. Se fundamentan en información previa. Supongamos que, basándose en ciertos datos, un investigador plantea la siguiente hipótesis: “el promedio mensual de casos de trastorno psiconeurótico caracterizados por reacción asténica, que serán atendidos en los hospitales de la ciudad de Linderbuck, resultará mayor a 20”. Además, desea transformar esta hipótesis de investigación en una hipótesis estadística. Lo primero que debe hacer es analizar cuál es la estadística a que su hipótesis hace referencia (en el ejemplo se trata de un promedio o media mensual de casos atendidos). El segundo paso consiste en encontrar cómo se simboliza esa estadística (promedio se simboliza como $\bar{X}$). El tercer paso consiste en traducir la hipótesis de investigación a una forma estadística:

Hi: $\bar{X} > 20$ (“el promedio mensual de casos atendidos será mayor a 20”).

La hipótesis estadística nula sería la negación de la hipótesis anterior:

Ho: $\bar{X} < 20$ (“el promedio mensual de casos atendidos será menor a 20”).

y la hipótesis alternativa podría ser:

Ha: $\bar{X} = 20$ (“el promedio mensual de casos... es igual a 20”).

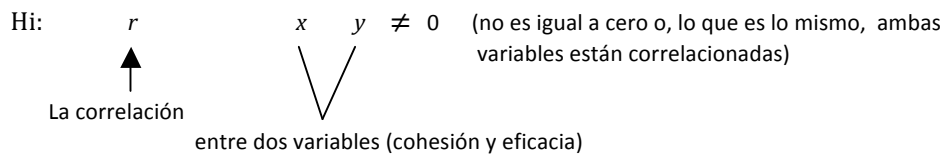
Después, el investigador comparará el promedio estimado por la hipótesis con el promedio actual de la muestra que seleccionó. La exactitud de su estimación se evalúa

con esta comparación. Y como señalan Black y Champion (1976), algunos investigadores consideran las hipótesis estadísticas de estimación como hipótesis de diferencia, pues en última instancia lo que se evalúa es la diferencia entre un valor planteado en la hipótesis y un valor observado en una muestra.

La estimación de estas hipótesis no se limita a promedios, ya que puede incluirse cualquier estadística: porcentajes, medianas, modas, etc. (Crosby et al., 2006).

Hipótesis estadísticas de correlación

Estas hipótesis tienen por objetivo traducir en términos estadísticos una correlación entre dos o más variables. El símbolo de una correlación entre dos variables es “r” (minúscula), y entre más de dos variables “R” (mayúscula). La siguiente hipótesis: “a mayor cohesión en un grupo, mayor eficacia en el logro de sus metas primarias”, se traduciría tal como se muestra en el esquema.



La hipótesis nula se traduciría:

Ho: $r_{xy} = 0$ (Las dos variables no están correlacionadas; su correlación es cero.)

Otro ejemplo:

Hi: $R_{xyz} \neq 0$ (La correlación entre las variables autonomía, variedad y motivación intrínseca no es igual a cero. Es decir, las tres variables “x”, “y”, “z” están asociadas).

Ho: $R_{xyz} = 0$ (No hay correlación)

Hipótesis estadísticas de la diferencia de medias u otros valores

En estas hipótesis se compara una estadística entre dos o más grupos. Supongamos que un investigador plantea la siguiente pregunta de estudio: ¿difieren los periódicos *Télex* y *Noticias* en cuanto al promedio de editoriales mensuales que dedicaron, durante el

último año, al tema del terrorismo internacional?¹ Su hipótesis de investigación podría ser: “existe una diferencia entre el promedio de editoriales mensuales que dedicó, durante el último año, al tema del terrorismo internacional el diario *Télex*, y el que dedicó el diario *Noticias*”. La estadística que se compara entre los grupos (editoriales de *Télex*, un grupo, y editoriales de *Noticias*, otro grupo) es el promedio mensual ($\bar{X}$). La hipótesis estadística se formularía así:

$$\begin{array}{c} \text{es diferente} \\ \downarrow \\ \text{Hi: } \bar{X}_1 \neq \bar{X}_2 \text{ (promedio del grupo 2: editoriales de } \textit{Noticias}) \\ \uparrow \\ \text{(promedio del grupo uno: editoriales de } \textit{Télex}) \end{array}$$

La hipótesis nula:

$$\text{Ho: } \bar{X}_1 = \bar{X}_2 \text{ (“No hay diferencia entre los promedios de los dos grupos de editoriales”).}$$

Con otra estadística (porcentaje) y tres grupos, se obtendrían hipótesis estadísticas como las siguientes:

$$\text{Hi: } \%_1 \neq \%_2 \neq \%_3 \text{ (“Los porcentajes de los tres grupos son distintos”).}$$

$$\text{Ho: } \%_1 = \%_2 = \%_3 \text{ (“No hay diferencias”).}$$

PUNTUACIONES z

Las **puntuaciones z** son transformaciones que se pueden hacer a los valores o las puntuaciones obtenidas, con el propósito de analizar su distancia respecto a la media, en unidades de desviación estándar. Una puntuación z nos indica la dirección y el grado en que un valor individual obtenido se

Puntuación z Medida que indica la dirección y el grado en que un valor individual se aleja de la media, en una escala de unidades de desviación estándar.

¹ Nombres completamente ficticios

aleja de la media, en una escala de unidades de desviación estándar. Como mencionan Nie et al. (1975), las puntuaciones z son el método más comúnmente utilizado para estandarizar la escala de una variable medida en un nivel por intervalos.

Su fórmula es:

$$z = \frac{X - \bar{X}}{s}$$

donde X es la puntuación o el valor a transformar, $\bar{X}$ es la media de la distribución y s la desviación estándar de ésta. El resultado z es la puntuación transformada en unidades de desviación estándar.

Supongamos que en una distribución de frecuencias obtuvimos una media de 60 y una desviación estándar de 10, y deseamos comparar una puntuación de “50” con el resto de la distribución. Entonces, transformamos esta puntuación o tal valor en una puntuación z . Tenemos que:

$$X = 50$$

$$\bar{X} = 60$$

$$s = 10$$

La puntuación z correspondiente a un valor de “50” es:

$$z = \frac{50 - 60}{10} = -1.00$$

Podemos decir que el valor “50” se localiza a una desviación estándar por debajo de la media de la distribución (el valor “30” está a tres desviaciones estándar por debajo de la media).

Estandarizar los valores permite comparar puntuaciones de dos distribuciones diferentes (la forma de medición es la misma, aunque se trata de distribuciones distintas). Por ejemplo, podemos contrastar una distribución obtenida en una

preprueba con otra obtenida en una posprueba (en un contexto experimental). Supongamos que se trata de un estímulo que incrementa la productividad. Un trabajador obtuvo en la preprueba una productividad de 130 (la media del grupo fue de 122.5 y la desviación estándar de 10). Y en la posprueba obtuvo 135 (la media del grupo fue de 140 y la desviación estándar de 9.8). ¿Mejoró la productividad del trabajador? En apariencia, la mejoría no es considerable. Sin transformar las dos calificaciones en puntuaciones z , no es posible asegurarlo porque los valores no pertenecen a la misma distribución. Entonces transformamos ambos valores a puntuaciones z , los pasamos a una escala común donde la comparación es válida. El valor de 130 en productividad en términos de unidades de desviación estándar es igual a:

$$z = \frac{130 - 122.5}{10.0} = 0.75$$

Y el valor de 135 corresponde a una puntuación z de:

$$z = \frac{135 - 140}{9.8} = -0.51$$

Como observamos, en términos absolutos 135 es una mejor puntuación que 130, pero no en términos relativos (en relación con sus respectivas distribuciones).

La distribución de puntuaciones z no cambia la forma de la distribución original, pero sí modifica las unidades originales a “unidades de desviación estándar” (Wright, 1979). La distribución de puntuaciones z tiene una media de 0 (cero) y una desviación estándar de 1 (uno). La figura 8.1 muestra la distribución de puntuaciones z .

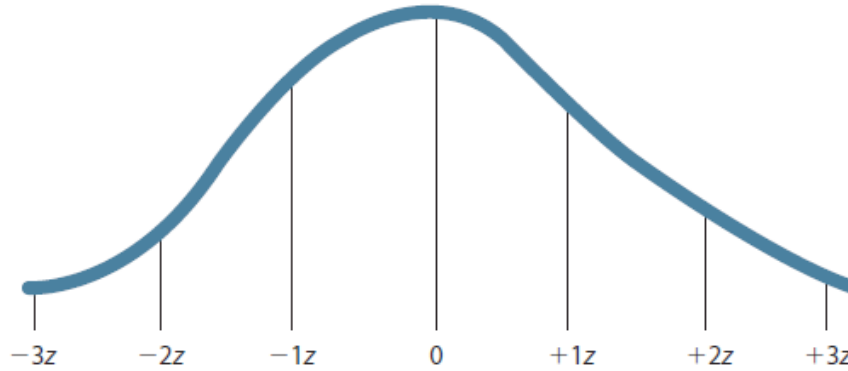


Figura 8.1 Distribución de puntuaciones z

Las puntuaciones z también sirven para comparar mediciones de distintas pruebas o escalas aplicadas a los mismos participantes (los valores obtenidos en cada escala se transforman en puntuaciones z y se comparan) (Delbaere et al., 2007). No debe olvidarse que los elementos de la fórmula específicamente son la media y la desviación estándar que corresponden al valor a transformar (de su misma distribución). También, las puntuaciones z sirven para analizar distancias entre puntuaciones de una misma distribución y áreas de la curva que abarcan tales distancias, o para sopesar el desempeño de un grupo de participantes en varias pruebas. Las puntuaciones z son un elemento descriptivo adicional que se agrega para analizar nuestros datos.

ESTIMACIÓN DE PARÁMETROS: UN EJEMPLO CON LA MEDIA POBLACIONAL

En el capítulo 10 del texto impreso se mencionó que una de las funciones de la estadística inferencial es la estimación de parámetros, pues bien, al calcular la media de nuestra muestra y si no estamos muy seguros de su representatividad podemos seguir un procedimiento para ver si nuestra hipótesis respecto a la media poblacional es aceptada o rechazada.² Lo anterior es para responder a la siguiente pregunta: ¿qué

² En este ejemplo se utiliza la media, tal vez el caso más conocido, pero puede ser cualquier otro parámetro poblacional.

hacemos para ver si nuestra hipótesis sobre la media poblacional es aceptada o rechazada? Pero antes de estudiar el procedimiento, es necesario hacer las siguientes consideraciones:

- a) La distribución muestral es una distribución normal de puntuaciones z , es decir, la base de la curva son unidades de desviación estándar.
- b) Las puntuaciones z son distancias que indican áreas bajo la distribución normal. En este caso, áreas de probabilidad.
- c) El área de riesgo es tomada como el área de rechazo de la hipótesis; por el contrario, el área de confianza, como el área de aceptación de la hipótesis.
- d) Se habla de una hipótesis acerca del parámetro (en este caso, media poblacional).

Si partimos de estas consideraciones, el procedimiento es:

1. Sobre bases firmes (revisión de la literatura e información disponible), establecer una hipótesis acerca del parámetro poblacional. Por ejemplo: el promedio de horas diarias que se exponen los niños de la ciudad de Valladolid a la televisión en fin de semana es de 3.0 horas.
2. Definir el nivel de significancia. Por ejemplo, .05.
3. Recolectar los datos en una muestra representativa. Vamos a suponer que obtuvimos una media de 2.9 horas y una desviación estándar de 1.2 horas; la muestra incluyó 312 niños.
4. Estimar la desviación estándar de la distribución muestral de la media utilizando la siguiente fórmula:

$$\bar{X} \quad S_{\bar{X}} = \frac{s}{\sqrt{n}}$$

Donde $S_{\bar{X}}$ es la desviación estándar de la distribución muestral de la media, s representa la desviación estándar de la muestra y n es el tamaño de la muestra.

En el ejemplo:

$$S\bar{X} = \frac{1.2}{\sqrt{312}}$$

$$S\bar{X} = \frac{1.2}{\sqrt{312}} = 0.0679$$

5. Transformar la media de la muestra en una puntuación z , en el contexto de la distribución muestral, con una variación de la fórmula ya conocida para obtener puntuaciones z :

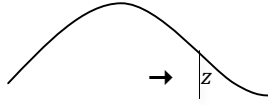
$$z = \frac{X - \bar{X}}{S\bar{X}}$$

donde X es la media de la muestra (recordemos que la distribución muestral es de medias y no de puntuaciones), $\bar{X}$ es la media hipotética de la distribución muestral (parámetro poblacional), $S\bar{X}$ es la desviación estándar de la distribución muestral de medias. Así, tenemos:

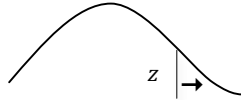
$$\alpha = \frac{1.984}{2.48}$$

6. En la *tabla de áreas bajo la curva normal* (ver apéndice, 4 tabla 1 de este CD), buscar aquella puntuación z que deje a 25% por encima de ella, que es 1.96. En la tabla del apéndice 4 se presenta la distribución de puntuaciones z , sólo la mitad, pues debemos recordar que es una distribución simétrica y se aplica igual para ambos lados de la media. Así se incluye en los textos de estadística. Se busca 2.5%, porque la tabla sólo abarca la mitad de la distribución y el riesgo que estamos afrontando es de 5% (2.5% del extremo de cada lado). La tabla mencionada contiene cuatro columnas: la primera indica puntuaciones z , la segunda expresa la distancia de la puntuación z a la media. La tercera, el área que está por de bajo de esa puntuación

desde el comienzo de la distribución, como se muestra en la gráfica:



Y la cuarta, del área que está por encima de esa puntuación:



Las áreas se expresan en proporciones. Lo que buscamos es una puntuación z que deje por encima un área de 0.0250 o 2.5% (la encontramos en la cuarta columna de la tabla); esta puntuación z es de 1.96. Siempre que nuestro nivel de significancia sea 0.05, tomamos la puntuación z de 1.96.

7. Comparar la media de la muestra transformada a puntuación z con el valor 1.96; si es menor, aceptar la hipótesis; si es mayor, rechazarla. Veamos el ejemplo:

Media de la muestra transformada a z
1.47

Nivel de significancia del 0.05
 ± 1.96

Decisión: Aceptar la hipótesis a un nivel de significancia de 0.05 (95% a favor y 5% de riesgo de cometer un error).

Si la media obtenida
al transformarse en z ,
Hubiera sido 3.25,

7.46 o un valor mayor $\rightarrow$ Rechazar la hipótesis

Por ejemplo :

Media de la muestra = 2.0

Desviación estándar de la muestra = 0.65

$$\begin{aligned}
 n &= 700 \\
 S\bar{x} &= 0.0246 \\
 z &= 40.65
 \end{aligned}$$

La media está situada a más de 40 desviaciones estándar de la media; se localiza en la zona crítica (más allá de 1.96 desviaciones estándar). Rechazar la hipótesis.

¿POR QUÉ ES IMPORTANTE OTRO CONCEPTO PARA LA ESTIMACIÓN DE PARÁMETROS: EL INTERVALO DE CONFIANZA?

Se ha hablado de la distribución muestral por lo que respecta a la prueba de hipótesis, pero otro procedimiento de la estadística inferencial es construir un intervalo donde se localiza un parámetro (Wiersma y Jurs, 2008). Por ejemplo, en lugar de pretender probar una hipótesis acerca de la media poblacional, puede buscarse un intervalo donde se ubique dicha media. Lo anterior requiere un nivel de confianza, al igual que en la prueba de hipótesis inferenciales. El nivel de confianza es al intervalo de confianza lo que el nivel de significancia es a la prueba de hipótesis. Es decir, se trata de una probabilidad definida de que un parámetro se va a ubicar en un determinado intervalo. Recordemos que los niveles de confianza –expresados en porcentajes– más comunes en la investigación son 0.95 y 0.99 (.05 y .01). Su sentido es similar al ya comentado, si es de 0.95; quiere decir que tenemos 95% en favor de que el parámetro se localice en el intervalo estimado, contra 5% de elegir un intervalo equivocado. El nivel de 0.99 señala 99% de probabilidad de seleccionar el intervalo adecuado. Tales niveles de confianza se expresan en unidades de desviación estándar. Una vez más se acude a la distribución muestral, concretamente a la tabla de áreas bajo la curva normal (apéndice 4, tabla 1 de este CD anexo o en STATS® “Áreas bajo la curva normal”)³, y se elige la puntuación z correspondiente al nivel de confianza seleccionado. Una vez hecho esto, se aplica la siguiente fórmula:

$$\text{Intervalo de confianza} = \text{estadígrafo} + \left(\begin{array}{c} \text{Puntuación } z \text{ que} \\ \text{expresa el nivel de} \\ \text{confianza elegido} \end{array} \right) \left(\begin{array}{c} \text{Desviación estándar de} \\ \text{la distribución muestral} \\ \text{correspondiente} \end{array} \right)$$

³ También STATS® contiene esta tabla.

En la fórmula, el estadígrafo es la estadística calculada en la muestra, la puntuación z es 1.96 con un nivel de 0.95 y de 2.58 con un nivel de 0.99, en tanto que el error estándar depende del estadígrafo en cuestión. Veámoslo con el ejemplo de la media en el caso de la exposición diaria a la televisión (en fin de semana) por parte de los niños de Valladolid:

Media = 2.9 horas

$s = 1.2$ horas

$\bar{Sx} = 0.0679$ (desviación estándar de la distribución muestral de la media).

Nivel de confianza = 0.95 ($z = 1.96$)

Intervalo de confianza = $2.9 \pm (1.96) (0.0679)$
= $2.9 \pm (0.133)$

Intervalo de confianza: la media poblacional está entre 2.767 y 3.033 horas, con 95% de probabilidades de no cometer error.

CÁLCULO DEL COEFICIENTE DE CONFIABILIDAD (FIABILIDAD) ALFA-CRONBACH

En los capítulos 9 y 10 se definió el coeficiente *alfa* de Cronbach (α), así como sus usos e interpretación. Los programas de análisis estadístico como SPSS, Minitab, SAS y otros, lo incluyen y calculan instantáneamente. Pero para quienes *no* tienen acceso a estos programas, presentamos la forma de obtenerlos. De acuerdo con Carmines y Zeller (1988, pp. 44 y 45), así como Corbetta (2003), existen tres procedimientos para determinar *el coeficiente “ α ” o alfa*:

1. *Sobre la base de la varianza de los ítems*, con la aplicación de la siguiente fórmula:

$$\Sigma s^2(Y_i)$$

En donde N representa el número de ítems de la escala, $\sum s^2(Y_i)$ es igual a la sumatoria de las varianzas de los ítems y " s^2x " equivale a la varianza de toda la escala.

2. Sobre la base de la matriz de correlación de los ítems, el procedimiento sería:

- a) Se aplica la escala.
- b) Se obtienen los resultados.
- c) Se calculan los coeficientes de correlación r de Pearson entre todos los ítems (todos contra todos de par en par).
- d) Se elabora la matriz de correlación con los coeficientes obtenidos. Por ejemplo:

	Ítems			
	1	2	3	4
1	—	0.451	0.399	0.585
2	ya fue calculado	—	0.489	0.501
3	ya fue calculado	ya fue calculado	—	0.541
4	ya fue calculado	ya fue calculado	ya fue calculado	—

Los coeficientes que se mencionan como "ya fue calculado", se ubican en la parte superior de las líneas horizontales (guiones). Es decir, cada coeficiente se incluye una sola vez y se excluyen los coeficientes que vinculan al ítem o puntuación consigo misma (1 con 1, 2 con 2, 3 con 3 y 4 con 4).

- e) Se calcula $\bar{r}$ (promedio de las correlaciones).

$$\bar{r} = \frac{\sum P}{NP} \quad (\sum P \text{ es la sumatoria de los valores de las correlaciones y } NP \text{ el número de correlaciones no repetidas o no excluidas}).$$

$$\bar{p} = \frac{0.451 + 0.399 + 0.585 + 0.489 + 0.541}{6}$$

$$\bar{p} = 0.494$$

f) Se aplica la fórmula siguiente:

$$\alpha = \frac{N\bar{p}}{1 + \bar{p}(N + 1)}$$

$$\alpha = \frac{N\bar{p}}{1 + \bar{p}(N + 1)}$$

Donde N es el número de ítems y $\bar{p}$ el promedio de las correlaciones entre ítems.

En el ejemplo:

$$\alpha = \frac{4(0.494)}{1 + 0.494(4 - 1)}$$

$$\alpha = \frac{1.984}{2.48}$$

$$\alpha = 0.798$$

$$\alpha = 0.80 \text{ (cerrando)}$$

Es un coeficiente aceptable y recordemos que todos los ítems de la escala deben estar medidos en intervalos o razón.

3. Mediante otra fórmula que se basa en la correlación promedio (Corbetta, 2003, p. 238).

Se usa la siguiente fórmula:

$$\alpha = \frac{nr}{1 + r(n-1)}$$

Donde n representa el número de ítems o elementos de la escala y r es su correlación promedio.

SUSTENTO DEL ANÁLISIS DE VARIANZA UNIDIRECCIONAL

En el capítulo 10 del texto impreso, se dijo que el ANOVA unidireccional produce un valor F , el cual se basa en una distribución muestral, conocida como *distribución F* , y tal valor compara las variaciones en las puntuaciones debidas a dos diferentes fuentes: variaciones entre los grupos que se contrastan y variaciones dentro de los grupos.

Si los grupos difieren realmente entre sí, sus puntuaciones variarán más de lo que puedan variar las puntuaciones entre los integrantes de un mismo grupo. Veámoslo con un ejemplo cotidiano. Si tenemos tres familias A , B y C . La familia A está integrada por Felipe, Angélica, Elena y José Luis. La familia B está compuesta por Chester, Pilar, Íñigo, Alonso y Carlos. Y la familia C está integrada por Rodrigo, Laura y Roberto. ¿Qué esperamos? Pues que los integrantes de una familia se parezcan más entre sí que a los miembros de otra familia. Esto se graficaría como en la figura 8.2.

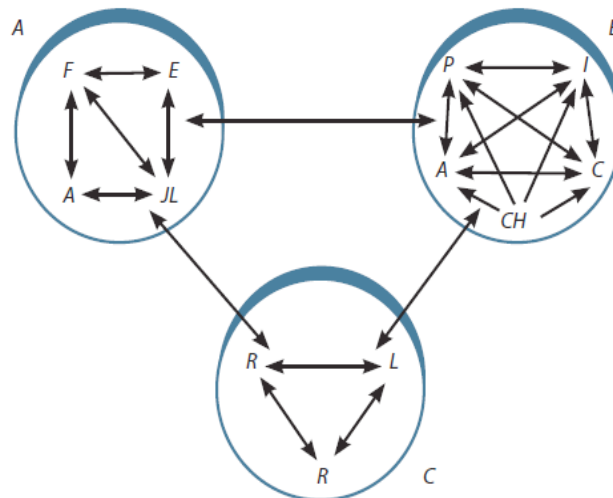


Figura 8.2 Ejemplo de las variaciones de los grupos comparados

Es decir, esperamos *homogeneidad* intrafamilias y *heterogeneidad* interfamilias. ¿Qué sucedería si los miembros de las familias se parecieran más a los integrantes de las otras familias que a los de la suya propia? Quiere decir que no hay diferencia entre los grupos (en el ejemplo, familias).

Esta misma lógica se aplica a la *razón F*, la cual nos indica si las diferencias entre los grupos son mayores que las diferencias intragrupos (dentro de éstos). Estas diferencias se miden en términos de varianza. La *varianza* es una medida de dispersión o variabilidad alrededor de la media y se calcula en términos de desviaciones elevadas al cuadrado. Recuerde que la *desviación estándar* es un promedio de desviaciones respecto a la media $X - \bar{X}$ y la *varianza es un promedio de desviaciones respecto a la media elevadas al cuadrado*. La varianza por eso se simboliza como S^2 y su *fórmula* es $\Sigma(X - \bar{X})^2 / n$. En consecuencia, la *razón F*, que es una razón de varianzas, se expresa así:

$$F = \frac{\text{Media cuadrática entre los grupos}}{\text{Media cuadrática dentro de los grupos}}$$

En donde la *media cuadrática* implica un promedio de varianzas elevadas al cuadrado. La media cuadrática entre los grupos se obtiene al calcular la media de las puntuaciones de todos los grupos (media total), después se obtiene la desviación de la media de cada grupo respecto a la media total y se eleva al cuadrado cada una de estas desviaciones, y luego se suman. Finalmente se sopesa el número de individuos en cada grupo y la *media cuadrática* se obtiene con base en los *grados de libertad intergrupales* (no se calcula con base en el número de puntuaciones). La *media cuadrática dentro de los grupos* se calcula obteniendo primero la desviación de cada puntuación respecto a la media de su grupo; posteriormente esta fuente de variación se suma y combina para obtener una media de la *varianza intragrupal* para todas las observaciones, tomando en cuenta los grados de libertad totales (Wright, 1979; Doncaster y Davey, 2007).

Las fórmulas de la media cuadrática son:

$$\text{Media cuadrática entre grupos} = \frac{\text{Suma de cuadrados entre grupos}}{\text{Grados de libertad entre grupos}}$$

Los grados de libertad entre grupos = $K + 1$ (donde K es el número de grupos).

$$\text{Media cuadrática dentro de los grupos} = \frac{\text{Suma de cuadrados intragrupos}}{\text{Grados de libertad intragrupos}}$$

Los grados de libertad intragrupos = $n - K$ (n es el tamaño de la muestra, la suma de los individuos de todos los grupos, y K recordemos que es el número de grupos).

Pues bien, cuando F resulta significativa, quiere decir que los grupos difieren “significativamente” entre sí. Es decir, se acepta la hipótesis de investigación y se rechaza la hipótesis nula.

Cuando se efectúa el *análisis de varianza* por medio de un programa computacional estadístico, se genera una tabla de resultados con los elementos de la tabla 8.1.

Tabla 8.1 Elementos para interpretar un análisis de varianza unidireccional realizado en SPSS

Fuente de variación (Source)	Sumas de cuadrados (Sums of squares)	Grados de libertad (Degrees of freedom)	Medias cuadráticas (Mean squares)	Razón F (F-ratio)	Significancia de F (F prob.)
Entre grupos (between groups)	SS entre +	gl entre	SS entre/gl entre	$\frac{MC \text{ entre}}{MC \text{ intra}}$	α
Intragrupos (within groups)	SS intra	gl intra	SS intra/gl intra		
Total	SS entre SS intra	gl entre + gl intra			
En Minitab se dan los siguientes elementos: Fuente GL (grados de libertad) SC (suma de cuadrados) MC (media cuadrática) F (valor) P (sig.)					

El valor *alfa* o probabilidad a elegir una vez más es 0.05 o 0.01. Si es menor que 0.05 es significativo en este nivel, y si es menor que 0.01 también es significativo en este nivel. Cuando el programa o paquete estadístico no incluye la significancia se acude a la tabla 3 del apéndice 4 (*tabla de la distribución F* o en STATS® a “Valores de “F” al nivel de confianza de .05 y .01”). Esta tabla contiene una lista de razones significativas (razones *F*) que debemos obtener para aceptar la hipótesis de investigación en los niveles de confianza de 0.05 y 0.01. Al igual que en caso de la razón *t*, el valor exacto de *F* que debemos obtener depende de sus grados de libertad asociados. Por lo tanto, la utilización de la tabla se inicia al buscar los dos valores *gl*, los *grados de libertad entre los grupos y los grados de libertad intragrupos*. Los grados de libertad entre grupos se indican en la parte superior de la página, mientras que los grados de libertad intragrupos se han colocado al lado izquierdo de la tabla. El cuerpo de la tabla de la distribución *F* presenta razones *F* significativas en los niveles de confianza de 0.05 y 0.01.

Si <i>F</i>	=	1.12
<i>gl</i> entre	=	2
<i>gl</i> intra	=	60

Este *valor F* se compara con el valor que aparece en la tabla de la distribución *F* que es 3.15 y como *el valor F* calculado es menor al de dicha tabla, rechazaríamos la hipótesis de investigación y aceptaríamos la hipótesis nula. Para que el *valor F* calculado sea significativo debe ser igual o mayor al de la tabla.

ANÁLISIS MULTIVARIADO

En el capítulo 10 del libro, cuando se analizaron los principales métodos estadísticos paramétricos, concretamente, después de revisar el ANOVA unidireccional, nos preguntábamos: ¿pero qué ocurre cuando tenemos diversas variables independientes y una dependiente, varias independientes y dependientes? Tal como observábamos en diagramas como el que se muestra en la figura 8.3.

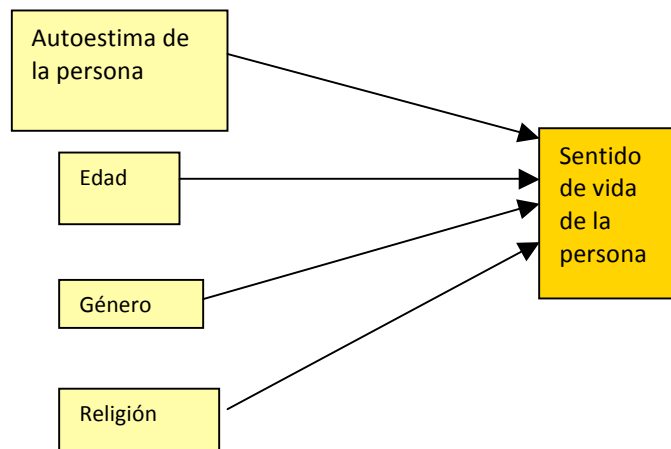


Figura 8.3 Ejemplo con cuatro variables independientes y una dependiente.

La respuesta era: entonces, requerimos de otros métodos estadísticos. Éstos son los que revisaremos a continuación y una vez más, sobre la base de que existen computadoras y programas como el SPSS, del mismo modo centrándonos en los elementos fundamentales de interpretación.

¿Qué son los métodos de análisis multivariado?

Los métodos de análisis multivariado son aquellos en que se analiza la relación entre diversas variables independientes y al menos una dependiente. Son métodos más complejos que requieren del uso de computadoras para efectuar los cálculos necesarios (normalmente se enseñan a nivel posgrado).

¿Qué es el análisis factorial de varianza?

ANOVA (análisis de varianza de k direcciones o varios factores)

Definición: Es una prueba estadística para evaluar el efecto de dos o más variables independientes sobre una variable dependiente.

Responde a esquemas como el que se muestra en la figura 8.4.

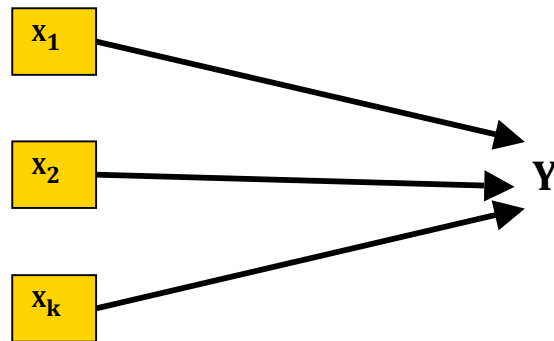


Figura 8.4 Esquema de un análisis factorial de varianza

Constituye una extensión del análisis de varianza unidireccional, solamente que incluye más de una variable independiente. Evalúa los efectos por separado de cada variable independiente y los efectos conjuntos de dos o más variables independientes.

Variables: dos o más variables independientes y una dependiente.

Nivel de medición de las variables: la variable dependiente (criterio) debe estar medida en un nivel por intervalos o razón, y las variables independientes (factores) pueden estar en cualquier nivel de medición, pero expresadas de manera categórica.

Interpretación y ejemplo

Hi: “la similitud en valores, la atracción física y el grado de realimentación positiva son variables que inciden en la satisfacción sobre la relación en parejas de novios”.

Contexto: Muestra de parejas de adultos jóvenes (23-29 años) de Santiago de Chile, pertenecientes a estratos económicos altos ($n = 400$).

El ANOVA efectuado mediante un paquete estadístico computacional como SPSS produce los siguientes elementos básicos:

- *Fuente de la variación (source of variation)*. Es el factor que origina variación en la dependiente. Si una fuente no origina variación en la dependiente, no tiene efectos.
- *Efectos principales (main effects)*. Es el efecto de cada variable independiente por separado; no está contaminado del efecto de otras variables independientes ni de error. Suele proporcionarse la suma de todos los efectos principales.
- *Interacciones de dos direcciones (2-way interactions)*. Representa el efecto conjunto de dos variables independientes, aislado de los demás posibles efectos de las variables independientes (individuales o en conjuntos). Suele proporcionarse la suma de los efectos de todas estas interacciones.
- *Interacciones de tres direcciones (3-way interactions)*. Constituye el efecto conjunto de tres variables independientes, aislado de otros efectos. Suele proporcionarse la suma de los efectos de todas estas interacciones.
- Puede haber efecto de *K*-direcciones, esto depende del número de variables independientes.

En nuestro ejemplo, tenemos los resultados que se muestran en la tabla 8.2.

Tabla 8.2 Ejemplo de resultados en el ANOVA

Fuente de variación (SOURCE OF VARIATION)	VARIABLE DEPENDIENTE: SATISFACCIÓN EN LA RELACIÓN				Significancia de F O P
	Suma de cuadrados (SUMS OF SQUARES)	Grados de libertad (DEGREES OF FREEDOM)	Medias cuadráticas (MEAN SQUARES)	Razón	
Efectos principales (main effects)	—	—	—	22.51	0.001**
Similitud	—	—	—	31.18	0.001**
Atracción	—	—	—	21.02	0.001**
Realimentación	—	—	—	11.84	0.004**
Interacción de dos direcciones (2-way interactions)	—	—	—	7.65	0.010*
Similitud	—	—	—	4.32	0.040*
Atracción	—	—	—	2.18	0.110
Realimentación	—	—	—	1.56	0.190
Interacción de tres direcciones (3-way interaction)	—	—	—	8.01	0.020*
Similitud	—	—	—	—	—

Atracción	—	—	—
Realimentación			
—Residual	—	—	—
—Total	—	—	—

Nota: A los estudiantes que se inician en el ANOVA normalmente les interesa saber si las razones “F” resultaron o no significativas; por tanto, sólo se incluyen estos valores. Por lo que es a ellos a quienes los autores recomiendan concentrarse en dichos valores y evitar confusiones. Desde luego, el investigador experimentado acostumbra estudiar todos los valores.

** Razón “F” significativa al nivel del 0.01 ($p < 0.01$)

* Razón “F” significativa al nivel del 0.05 ($p < 0.05$)

Como podemos ver en la tabla 8.2, la similitud, la atracción y la realimentación tienen un efecto significativo sobre la satisfacción en la relación. Respecto a los efectos de dos variables independientes conjuntas, sólo la similitud y la atracción tienen un efecto, hay un efecto conjunto de las tres variables independientes. La hipótesis de investigación se acepta y la nula se rechaza. Asimismo, se recuerda al lector que en el capítulo 5 del presente disco: diseños experimentales: segunda parte: series cronológicas, factoriales y cuasiexperimentos (en el apartado sobre diseños factoriales) se explica la noción de interacción entre variables independientes. Cabe agregar que el ANOVA es un método estadístico propio para los diseños experimentales factoriales.

¿Qué es el análisis de covarianza?

Definición: es un método estadístico que analiza la relación entre una variable dependiente y dos o más independientes, con el que se elimina o controla el efecto de al menos una de estas independientes. Similar al ANOVA, excepto que permite controlar la influencia de una variable independiente, la cual con frecuencia es una característica antecedente que puede variar entre los grupos (Mertens, 2005; Babbie, 2009) o influir los resultados y afectar la claridad de las interpretaciones.

Perspectivas o usos: Wildt y Ahtola (1978, pp. 8-9) destacan tres perspectivas para el análisis de covarianza:

- A. *Perspectiva experimental.* Se aplica a aquellas situaciones en que el interés del investigador se centra en las diferencias observadas en la variable dependiente, por medio de las categorías de la variable independiente (o variables independientes). Pero el experimentador asume que hay otras variables independientes cuantitativas que contaminan la relación y cuya influencia debe ser controlada (figura 8.5).

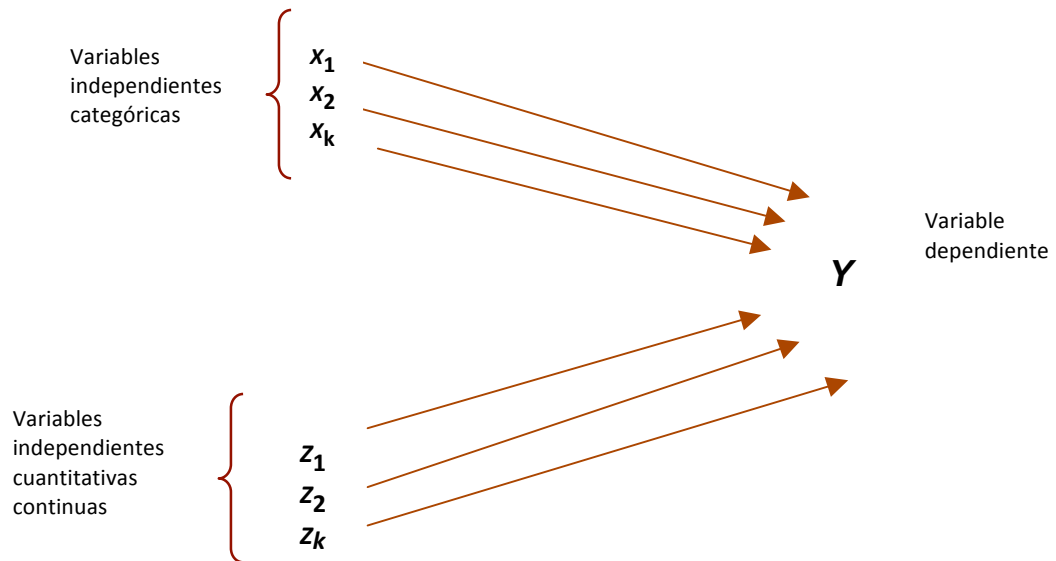


Figura 8.5 Ejemplo de variables independientes que afectan a una dependiente

Y el investigador únicamente se interesa por conocer la relación entre las variables independientes categóricas y la variable dependiente. Desea al mismo tiempo remover y controlar el efecto de las variables independientes cuantitativas no categóricas (continuas). Es decir, desea tener un esquema como el de la figura 8.6.

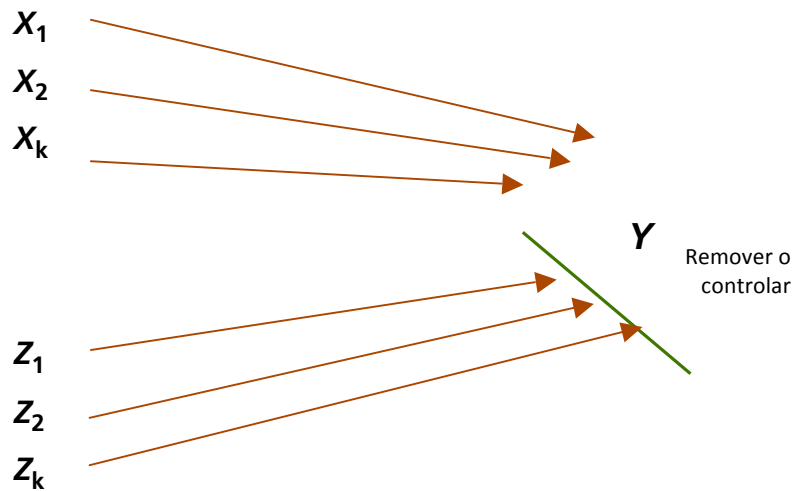


Figura 8.6 Ejemplo de control de variables independientes no categóricas.

El objetivo es “purificar la relación entre las independientes categóricas y la dependiente, mediante el control del efecto de las independientes no categóricas o continuas”.

Ejemplos de variables independientes categóricas serían: género (masculino, femenino), inteligencia (alta, media, baja), ingreso (menos de un salario mínimo, dos a cuatro salarios mínimos, cinco a 10 salarios mínimos, 11 o más salarios mínimos). Los niveles de medición nominal y ordinal son categóricos en sí mismos, mientras que los niveles de intervalos y razón deben transformarse en categorías más discretas. Estos últimos son en sí: cuantitativos, continuos y de categorías múltiples. Por ejemplo, el ingreso en su “estado natural” (pesos, dólares, euros, etc.) varía de la categoría cero hasta la categoría $(K)^k$, ya que puede haber millones de categorías.

Variable categórica — unas cuantas categorías o un rango medio.

Variable continua — muchas categorías (a veces una infinidad).

A dichas variables independientes cuantitativas continuas, cuya influencia se controla, se les denomina “covariables”. Una covariable se incluye en el análisis para remover su efecto sobre la variable dependiente, e incrementar el conocimiento de la relación entre las variables independientes categóricas de interés y la dependiente, lo cual aumenta la precisión del análisis (Doncaster y Davey, 2007).

En esta perspectiva, el *análisis de covarianza* puede ser concebido primero como un ajuste en la variable dependiente respecto a diferencias en la covariable o las covariables y, posteriormente, como una evaluación de la relación entre las variables independientes categóricas y los valores ajustados de la variable dependiente (Wildt y Ahtola, 1978). En términos de Creswell (2005): el procedimiento “ajusta” las puntuaciones en la dependiente para dar cuenta de la covarianza (por decirlo en términos sencillos: “hace equivalentes a los grupos en la(s) covariable(s)” y controla influencias potenciales que pueden afectar a la variable dependiente).

- B. *Perspectiva de interés por la covariable.* Esta perspectiva se ejemplifica con aquellas instancias en las cuales el interés principal se centra en analizar la relación entre la variable dependiente y la covariable (variable cuantitativa continua) o las covariables. Aquí el enfoque es distinto; la influencia que se remueve es la de las variables independientes categóricas. Primero se controla el efecto (en este caso “contaminante”) de estas variables y después se analiza el efecto “purificado” de las covariables.
- C. *Perspectiva de regresión.* En esta tercera perspectiva, tanto las variables independientes categóricas como las covariables resultan de interés para el investigador, quien puede desear examinar el efecto de cada variable independiente (covariables y no covariables, todas) y después ajustar o corregir los efectos de las demás variables independientes.

En cualquier caso, el *análisis de covarianza elimina influencias no deseadas sobre la variable dependiente*. Se puede utilizar en *contextos experimentales y no experimentales*. La mayoría de las veces la función del ANCOVA es “remover” la varianza compartida entre una o más covariables y la dependiente, de este modo, se valora en su justa dimensión la relación causal entre la(s) variable(s) independiente(s) de interés y la dependiente (Creswell, 2005). Veámoslo conceptualmente pero de forma gráfica con un ejemplo simple:

EJEMPLO

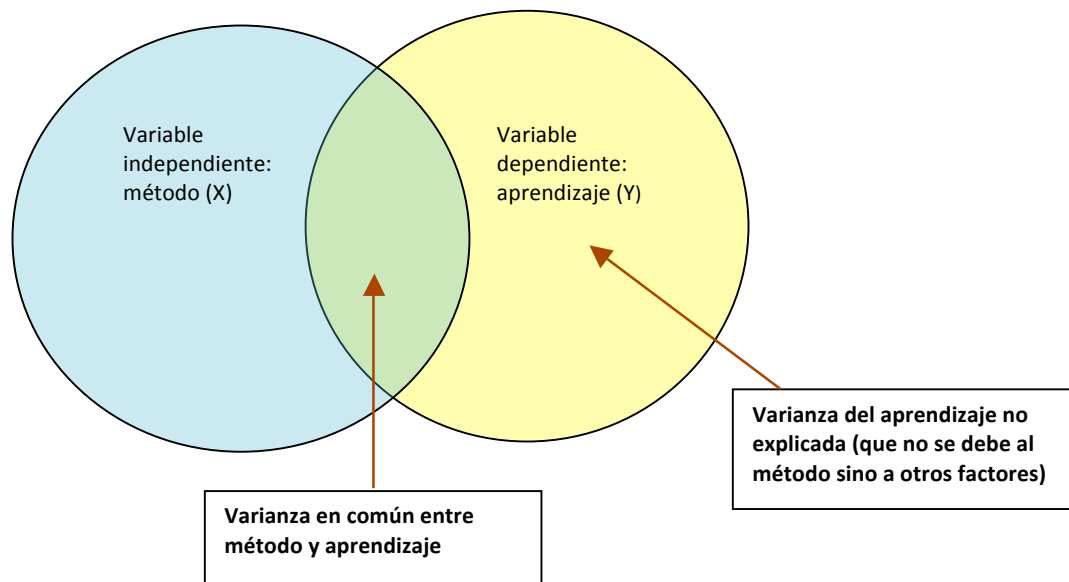
Estudio: Al investigador le interesa analizar el efecto en el aprendizaje de la computación, por medio de un nuevo método para su enseñanza a niños. La hipótesis es: “el nuevo método de enseñanza de la computación (MA-RH) provocará un mayor aprendizaje en los niños que un método tradicional”.

Entonces, implementa el siguiente experimento: a un grupo de infantes lo expone al nuevo método de enseñanza de computación (MA-RH); el otro grupo aprende con el método tradicional; finalmente, un tercer grupo, de control, *no* recibe ningún tipo de enseñanza en computación.

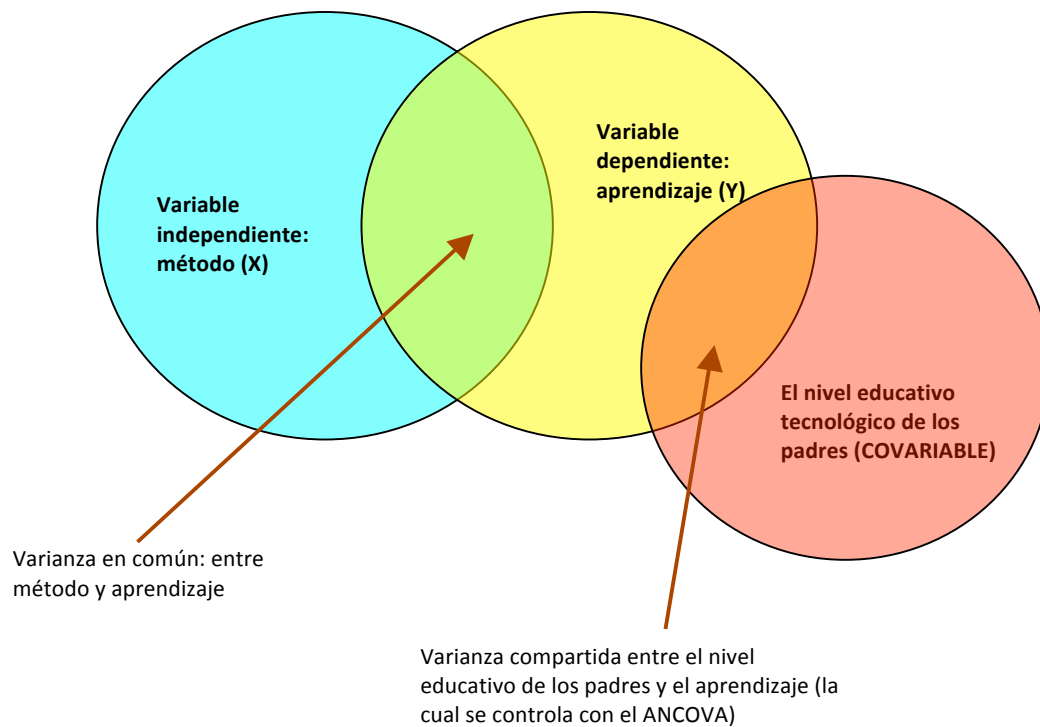
La variable independiente es el *tipo de método* con tres categorías o niveles (método nuevo, método tradicional y ausencia de método), la dependiente es el *aprendizaje en computación* (medida por una prueba estandarizada a nivel de intervalos). Se tiene un esquema como el de la figura 8.7.



El investigador sabe que el aprendizaje se puede deber a muchas razones, además del método. Es decir, el aprendizaje varía por diversos motivos, lo cual se representa en forma de conjuntos de la siguiente manera:



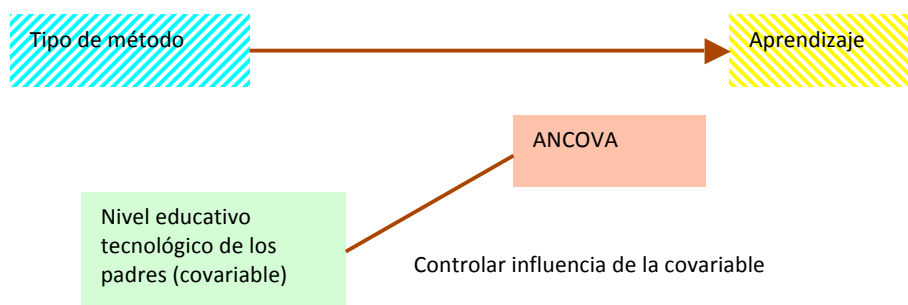
Con el experimento el investigador desea conocer la varianza en común entre método y aprendizaje (cuantificarla), la relación $X \rightarrow Y$ (pura). Si los niños son asignados al azar a los grupos del experimento y tiene grupos de tamaño aceptable, por el diseño mismo, remueve la influencia de las covariables que pudieran afectar. Pero si *no* es factible hacerlo y tiene un diseño cuasiexperimental (grupos intactos), debe remover tal influencia con el análisis de covarianza (eliminar al mínimo posible la varianza del aprendizaje no explicada), para evitar que las covariables impidan ver con claridad la relación $X \rightarrow Y$. Por ejemplo, el nivel educativo tecnológico de los padres puede influir (hace variar al aprendizaje) y este efecto debe ser controlado, al introducirlo como covariable.



Entre más covariables se controle, más se explicará sobre la varianza común entre el método y el aprendizaje.

Figura 8.7 Ejemplo del control de las covariables (con una covariable).

Lo que el investigador desea también se puede expresar gráficamente así:



Wildt y Ahtola (1978, p. 13) definen algunos usos del análisis de covarianza:

1. Incrementar la precisión en experimentos con asignación al azar.
2. Eliminar influencias extrañas o contaminantes que pueden resultar cuando las pruebas o los individuos no son asignados al azar a las diferentes condiciones experimentales (grupos de un experimento).
3. Eliminar efectos de variables que confundan o distorsionen la interpretación de resultados en estudios no experimentales.

Nivel de medición de las variables: la variable dependiente siempre está medida por intervalos o razón y las variables independientes pueden estar medidas en cualquier nivel. *Interpretación:* depende de cada caso específico, ya que el análisis de covarianza efectuado mediante un programa estadístico computacional, produce un cuadro de resultados muy parecido al del análisis de varianza. Los elementos más comunes pueden observarse en la tabla 8.3.

Tabla 8.3 Ejemplo de elementos comunes de un análisis de covarianza

Fuente de variación (Source of variation)	Sumas de cuadrados y productos cruzados (Sum of squares and cross products)	Sumas de cuadrados ajustadas (Adjusted sum of squares)	Grados de libertad (Degrees of freedom)	Medias cuadráticas	Razón F (F)	Significancia de F (Sig.)

La razón F es, igual que en el análisis de varianza, una razón de varianzas. El razonamiento estadístico es el mismo y F se interpreta igual, incluso se utiliza el mismo cuadro de la distribución F (tabla 3, apéndice 4 o en STATS® a “Valores de “ F ” al nivel de confianza de .05 y .01”). Solamente que las inferencias y conclusiones se hacen al considerar que las medias de la variable dependiente, a través de las categorías de las variables independientes, se han ajustado, de este modo eliminan el efecto de la

covariable o covariables.

EJEMPLO

Diseño de investigación que utiliza el análisis de covarianza

Hi: “los trabajadores que reciban retroalimentación verbal sobre el desempeño de parte de su supervisor mantendrán un nivel mayor de productividad que los trabajadores que reciban retroalimentación sobre el desempeño por escrito, y más aún que los trabajadores que no reciban ningún tipo de retroalimentación”.

$$\begin{array}{ccccc} \text{Hi: } \bar{X}_1 & > & \bar{X}_2 & > & \bar{X}_3 \\ \text{(verbal)} & & \text{(por escrito)} & & \text{(ausencia)} \end{array}$$

El investigador plantea un diseño experimental para intentar probar su hipótesis. Sin embargo, no puede asignar aleatoriamente a los trabajadores a los tres grupos del experimento. El diseño sería con grupos intactos (cuasiexperimental) y se esquematizaría así:

G_1	X_1	$\bar{X}_1$
G_2	X_2	$\bar{X}_2$
G_3	—	$\bar{X}_3$

Asimismo, el investigador presupone que hay un factor que puede contaminar los resultados (actuar como fuente de invalidación interna): la motivación. Diferencias iniciales en motivación pueden invalidar el estudio. Como la asignación al azar está

ausente, no se sabe si los resultados se ven influidos por dicho factor. Entonces, el experimentador decide eliminar o controlar el efecto de la motivación sobre la productividad para conocer los efectos de la variable independiente: tipo de retroalimentación. La motivación se convierte en covariable. El esquema es el que se muestra en la figura 8.8.

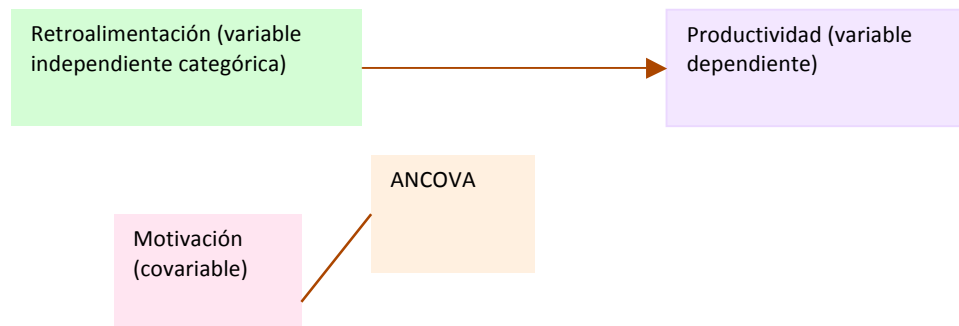


Figura 8.8 Ejemplo donde la motivación es covariable

Cabe destacar que, para introducir una *covariable* en el análisis, de preferencia debe medirse antes del inicio del experimento.

El *análisis de covarianza* “quita” a la variabilidad de la dependiente lo que se debe a la *covariable*. *Ajusta la varianza de la variable dependiente en las categorías de la independiente*, al basarse en la covariable. En el ejemplo, ajusta la varianza de la productividad debida a la motivación, en las categorías experimentales (tratamientos o grupos). El ajuste se realiza sobre la base de la correlación entre la covariable y la dependiente. Esto se muestra esquemáticamente en la tabla 8.4.

Una vez realizado el *análisis de covarianza*, se evalúa si *F* es o no significativa. *Cuando F resulta significativa se acepta la hipótesis de investigación*.

Si el resultado fuera:

$$G_1 = 35$$

$$G_2 = 36$$

La correlación entre la calificación en motivación y las puntuaciones en productividad es la base para el ajuste.

$$G_3 = 38$$

$$Gl \text{ entre} = K - 1 = 3 - 1 = 2$$

$$Gl \text{ intra} = N - K = 107$$

$$F = 1.70$$

Comparamos con el valor de la tabla respectiva: en el nivel de 0.05 es igual a 3.07, y nuestra razón F a 1.70 es menor a este valor. Por tanto, rechazamos la hipótesis de investigación y aceptamos la hipótesis nula. Esto se contrasta y profundiza con las medias ajustadas de los grupos que proporcione el análisis de covarianza (no las medias obtenidas en el experimento por cada grupo, sino las ajustadas con base en la covariable).

Recordemos que SPSS y Minitab nos proporcionan automáticamente la significancia de F .

Tabla 8.4 Ejemplo de un diseño de investigación que utiliza el análisis de covarianza como herramienta para ajustar diferencias en motivación entre los grupos

	Covariable Calificación en motivación	Variable independiente Tipo de realimentación	Variable dependiente Puntuaciones en productividad ajustadas, tomando en cuenta la covariable
G_1	0	X_1	0
G_2	0	X_2	0
G_3	0	—	0

¿Qué es la regresión múltiple?

Es un método para analizar el efecto de dos o más variables independientes sobre una dependiente. Asimismo, constituye una extensión de la regresión lineal sólo que con mayor número de variables independientes. Es decir, sirve para predecir el valor de una variable dependiente, cuando se conoce el valor y la influencia de las variables

independientes incluidas en el análisis. Si queremos conocer el efecto que ejercen las variables: *a)* satisfacción sobre los ingresos percibidos, *b)* antigüedad en la empresa, *c)* motivación intrínseca en el trabajo y *d)* percepción del crecimiento y desarrollo personal en el trabajo; sobre la variable “permanencia en la empresa” (duración o estancia), el modelo de regresión múltiple es el adecuado para aplicarlo a los datos obtenidos. Otro ejemplo sería el siguiente:

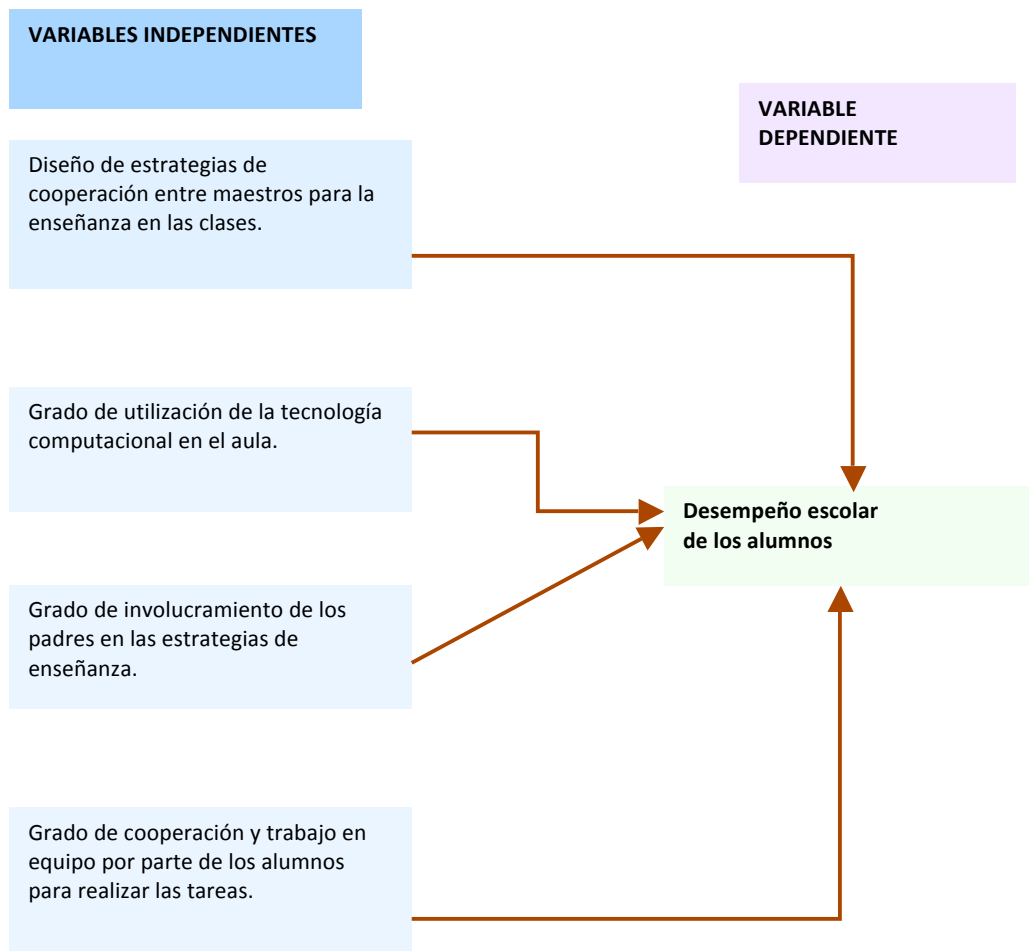


Figura 8.9 Esquema de un modelo con una variable dependiente y varias independientes, donde se conoce el efecto de cada una de éstas

Es decir, el modelo de regresión múltiple nos indica:

- La relación entre cada variable independiente y la única dependiente (cómo cambios en la independiente se vinculan con cambios en la dependiente).

- La relación entre todas las variables independientes (en conjunto) y la dependiente (cómo cambios en las independientes se vinculan con cambios en la dependiente).
- La predicción de la dependiente a partir de las independientes.
- La correlación entre las variables independientes (colinealidad).

Las variables independientes se denominan “predictoras”⁴ y anteceden temporalmente a la variable dependiente o criterio.

La información básica que proporciona la regresión múltiple es el coeficiente de correlación múltiple (R) y la ecuación de regresión.

Coeficiente de correlación múltiple (R). Señala la correlación entre la variable dependiente y todas las variables independientes tomadas en conjunto.

El coeficiente puede variar de cero a uno; cuanto más alto sea su valor, las variables independientes estarán más correlacionadas con la variable dependiente y explicarán en mayor medida sus fluctuaciones (varianza); en consecuencia, son factores más efectivos para predecir el comportamiento de esta última.

En el capítulo 10 del libro, se comentó el coeficiente de correlación de Pearson y se mencionó que cuando el coeficiente r se eleva al cuadrado (r^2), se obtiene el coeficiente de determinación y el resultado indica la *varianza de factores comunes*, esto es, el porcentaje de la variación de una variable debido a la variación de la otra y viceversa (o cuánto explica o determina una variable la variación de la otra). Pues bien, algo similar ocurre con el coeficiente de correlación múltiple, solamente que tenemos más variables a considerar. Cuando el coeficiente R se eleva al cuadrado (R^2), se produce el llamado *coeficiente de determinación o correlación parcial*, que nos señala la varianza explicada de la variable dependiente por todas las independientes (dicho de otra forma, el porcentaje de variación en la dependiente es debido a las independientes consideradas).

Veámoslo gráficamente en la figura 8.10 con dos independientes y una dependiente, a fin de que resulte menos complejo de entender.

⁴ Término anglosajón.

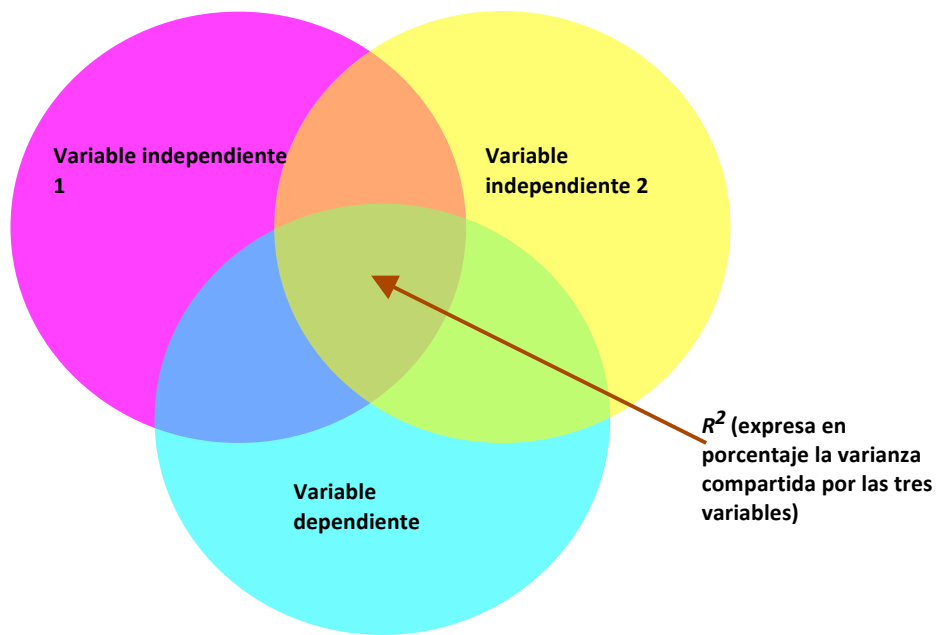


Figura 8.10 Esquema de un coeficiente de determinación o correlación parcial

Este coeficiente (R^2) resulta útil también para determinar la cantidad de varianza que una variable interviniente explica tanto de la variable independiente como de la dependiente, de este modo, se puede remover la varianza compartida de la interviniente con la variable independiente o la dependiente (Creswell, 2005; Sengupta, 2009), que es algo similar a lo que se efectúa con el análisis de covarianza.

Tal sería el caso de una relación del siguiente tipo:

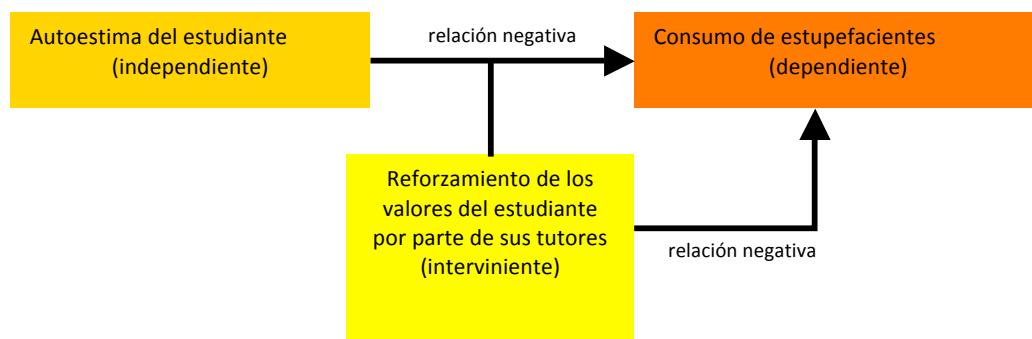


Figura 8.11 Ejemplo del coeficiente de determinación (correlación parcial)

Si resumimos lo visto en el capítulo 10 del libro sobre correlación y regresión lineal y lo expuesto hasta aquí, tenemos los coeficientes que se resumen en la tabla 8.5.

Tabla 8.5 Resumen de coeficientes de correlación bivariada y múltiple

Coeficiente	Símbolo	Información producida
Pearson	r	Grado de asociación entre dos variables (oscila entre 0 y 1).
Coeficiente de determinación	r^2	Varianza de factores comunes (porcentaje de la variación de una variable debido a la variación de la otra variable y viceversa). Oscila entre 0 y 100%.
Múltiple	R	Correlación entre la variable dependiente y todas las variables independientes tomadas en conjunto. Oscila entre 0 y 1.
Determinación (correlación parcial)	R^2	Varianza explicada de la dependiente por todas las independientes. Oscila entre 0 y 100%.

Otra información relevante producida por el análisis de regresión múltiple son los valores “beta” (β o b) que indican el peso o la influencia que tiene cada variable independiente sobre la dependiente, al controlar la varianza de todas las demás independientes. Cada peso *beta* es un coeficiente que señala la magnitud de la predicción de una variable independiente para una variable dependiente (criterio), después de remover los efectos de todas las demás independientes. Los valores *beta* están estandarizados, es decir, no importa que cada variable predictora esté medida en una escala diferente (como ocurría con las puntuaciones z) y se interpretan como el coeficiente de Pearson, de -1.00 a $+1.00$ (Creswell, 2005).

También el análisis proporciona coeficientes de correlación bivariados entre la dependiente y cada independiente (Shaw, 2009).

Para predecir valores de la variable dependiente se aplica la ecuación de regresión múltiple:

$$y = a + b_1X_1 + b_2X_2 + b_3X_3 + \dots b_kX_k$$

Donde a es una constante de regresión para el conjunto de puntuaciones obtenidas, $b_1, b_2, b_3, \dots, b_k$ son los pesos “beta” de las variables independientes. Mientras que X_1, X_2, X_3 y $\dots, X_k$ son valores de las variables independientes que fija el investigador para hacer la predicción.

La variable dependiente debe estar medida en un nivel por intervalos o de razón. Las independientes, en cualquier nivel de medición (el modelo estandariza mediciones). Cuando se utilizan variables categóricas (nominales u ordinales, como género, grupo étnico, nivel jerárquico, etc.) se transforman en variables “dummy” y se introducen al modelo como predictores. Los códigos *dummy* son series de números asignados para indicar la pertenencia a un grupo en cualquier categoría exhaustiva y mutuamente excluyente.

De acuerdo con Mertens (2005), la cantidad de varianza que cada independiente aporta para la variable dependiente puede tener cambios con diferentes órdenes de entrada de las variables independientes. Al respecto no hay reglas, se usa la lógica del investigador o criterios como los siguientes:

- Ingresar las variables de acuerdo con la fuerza de su correlación con la variable dependiente, de la más alta a la más baja.
- Seguir el orden en que se han introducido en estudios previos.
- Proceder de acuerdo con la teoría.
- Orden cronológico (tiempo en que se introducen las variables en un experimento o al medirse, si es que su medición fue por etapas, de la primera a la última).

Los resultados *más relevantes* que produce SPSS sobre la regresión múltiple se muestran en las tablas 8.6, 8.7 y en la figura 8.12, que corresponden a un estudio para predecir el clima laboral (criterio o dependiente) sobre la base de las siguientes variables independientes (Hernández Sampieri, 2005):

- Normalización (formalización de políticas en documentos).
- Avance del proceso de calidad en el departamento (un nuevo esquema de trabajo introducido en el 2004).
- Innovación departamental.

- Identificación del trabajo individual en los resultados generales de la organización.
- Comunicación (percepción del grado en que la información relevante de la empresa les es transmitida a los empleados de su departamento).
- Desempeño (índice de productividad del empleado).
- Motivación general.
- Antigüedad en la empresa (en meses).
- Satisfacción general en el trabajo.
- Liderazgo (percepción del superior inmediato como líder).
- Cultura (arraigo de la cultura organizacional definida por la dirección de la empresa).
- Pago (salario).

Primero. Se presentan las variables introducidas en el modelo de regresión:

Tabla 8.6 Variables introducidas en el ejemplo de regresión múltiple

Variables introducidas /eliminadas			
Modelo	Variables introducidas (a)	Variables eliminadas	Método (b)
	Pago Innovación Antigüedad Motivación Normalización Proceso de calidad Cultura Identificación Desempeño Liderazgo Satisfacción Comunicación		Introducir
a. Todas las variables solicitadas introducidas b. Variable dependiente: clima			

Segundo. Se presentan resultados de varianzas (ANOVA), los cuales omitimos, y los coeficientes *beta* y estadísticas de colinealidad (tabla 8.7).

Tabla 8.7 Ejemplo de resultados básicos de la regresión múltiple

Coeficientes ^a									
Modelo	Coeficientes no estandarizados		Coeficiente estandari- zados						
	B	Error tip.	Beta		Orden 0	Parcial	Semiparcial	Tolerancia	FIV
1 (Constante)	1.02E010	.000		Pesos beta detallado izquierda ← Correlaciones detallado derecha →					
Normalización	.083	.000	.106		.783	1.000	0.63	.352	1.045
Proceso de calidad	.083	.000	.135		.831	1.000	.0009	.305	1.378
Innovación	.083	.000	.060		.747	1.000	.059	.429	1.334
Identificación	.083	.000	.103		.812	1.000	.056	.197	1.355
Comunicación	.083	.000	.104		.858	1.000	.052	.253	1.080
Desempeño	.083	.000	.084		.005	1.000	.052	.201	1.320
Motivación	.083	.000	.078		.785	1.000	.044	.320	3.042
Antigüedad	.083	.000	.100		.713	1.000	.070	.488	1.048
Satisfacción	.083	.000	.105		.854	1.000	.054	.260	1.049
Liderazgo	.083	.000	.135		.855	1.000	.055	.274	1.047
Cultura	.083	.000	.043		.827	1.000	.055	.357	1.805
Pago	.083	.000	.126		.718	1.000	.092	.534	1.872
a. Variable dependiente: clima									

Tercero. Se muestran los valores estadísticos sobre los residuos (residuales).

Cuarto. Se realiza una regresión lineal de cada variable independiente sobre la dependiente, se tendrán tantas gráficas como variables predictoras). Mostramos el ejemplo de la variable cultura en la figura 8.12.

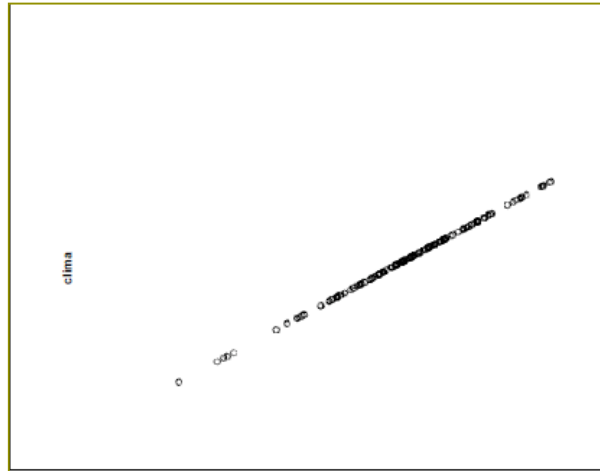


Figura 8.12 Gráfica de una variable en el ejemplo de regresión múltiple

Finalmente, mostramos un ejemplo de interpretación del coeficiente R^2 de otro de los estudios de Hernández Sampieri (2005), en el cual la variable dependiente es el clima organizacional total (medido por la escala de la Universidad de Celaya o ECOUNI) y las independientes son: moral, apoyo de la dirección, innovación, percepción de la empresa, comunicación, percepción del desempeño, motivación intrínseca, autonomía, satisfacción general, liderazgo, visión y recompensas o retribución. Las variables independientes fueron evaluadas a través de diferentes mediciones no incluidas en la ECOUNI⁵).

⁵ Mediciones “clásicas” de origen distinto a las de la ECOUNI. Por ejemplo, para la variable *visión* se usó la escala de Anderson y West (1998), para satisfacción el Job Satisfaction Survey, versión en español (Spector, 1997), para motivación la escala de Wang y Guthrie (2004), etcétera.

EJEMPLO

Para cierto laboratorio químico-farmacéutico (empresa), la R^2 fue de 0.989 (0.988 corregida). Todos los pesos *beta* tuvieron una significancia menor al 0.01 (excepto recompensas, la cual fue de 0.175). Una vez más, esta variable parece *no* ser “predictora” del clima organizacional. La tendencia resultante es tan contundente que poco puede comentarse al respecto, tal como lo muestra el diagrama respectivo. En el caso de la institución educativa, el coeficiente R^2 fue de prácticamente 0.80. Pero no tuvieron pesos *beta* significativos: desempeño, motivación intrínseca y liderazgo. Autonomía se encuentra en la frontera de la significancia (0.078). Por lo tanto, estas cuatro dimensiones *no* pueden considerarse —en la muestra— predictoras de la escala de clima organizacional total. La tendencia se presenta en el diagrama correspondiente.

Quinto. Se presenta el diagrama de dispersión de la regresión de todo el modelo que explica a la variable dependiente (figura 8.13)

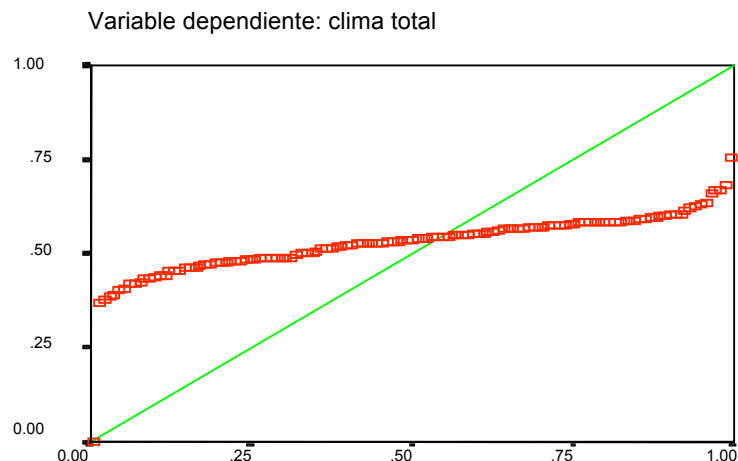


Figura 8.13 Diagrama de dispersión del clima organizacional: laboratorio

¿Qué es el análisis lineal de patrones o path analysis?

Es una técnica estadística multivariada para representar interrelaciones entre variables a partir de regresiones, así como analizar la magnitud de la influencia de unas variables

sobre otras, influencia directa e indirecta. Se trata de un modelo causal. Supongamos que tenemos el diagrama que se observa en la figura 8.14 y deseamos probarlo.

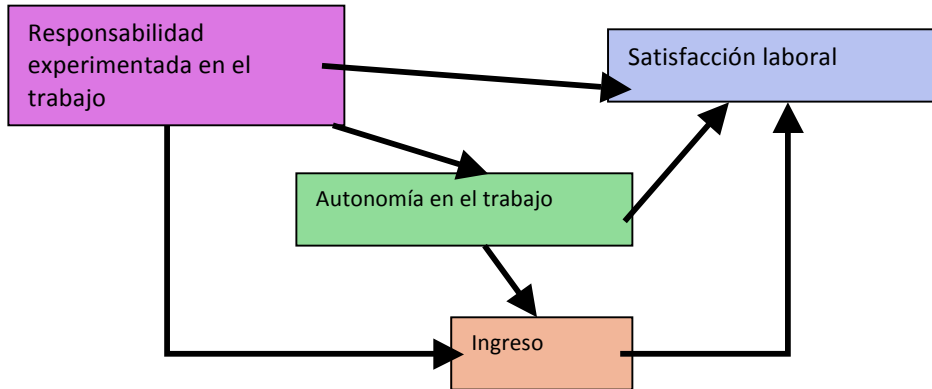


Figura 8.14 Un esquema propicio para el análisis de patrones o vías

El *análisis path* constituye un método para someterlo a prueba y una extensión de la regresión múltiple (Webley y Lea, 1997). La información principal que proporciona son los *coeficientes path*, los cuales representan la fuerza de las relaciones entre las variables (son coeficientes de regresión estandarizados como los pesos *beta*).

También cuantifica efectos. En el modelo puede haber variables independientes, intervinientes y dependientes; incluso una variable puede ser dependiente en una parte del modelo e independiente en otra (en la figura 8.14, ingreso es dependiente de autonomía y responsabilidad, pero también es independiente en relación con la satisfacción laboral). La ecuación del análisis es la siguiente⁶ para cada secuencia causal:

$$\text{variable dependiente} = b_1 (\text{variable independiente 1}) + b_2 (\text{variable independiente 2}) + b_k (\text{variable independiente } k).$$

En el ejemplo tenemos tres secuencias causales, por lo tanto, las correspondientes ecuaciones serían:

⁶ Adaptado de Webley y Lea (1997).

$$\text{Satisfacción} = b_{11} \text{ responsabilidad} + b_{12} \text{ autonomía} + b_{13} \text{ ingreso} + e_1$$

$$\text{Ingreso} = b_{21} \text{ responsabilidad} + b_{22} \text{ autonomía} + e_2$$

$$\text{Autonomía} = b_{31} \text{ responsabilidad} + e_3$$

Al igual que Webley y Lea (1997), se utiliza una notación diferente para los coeficientes de Bryman and Cramer, con la finalidad de clarificar que b_{11} en la primera ecuación es diferente de b_{21} en la segunda; e_1 , e_2 y e_3 representan el error o los términos de la varianza no explicada.

En la figura 8. 15 se muestra un ejemplo para ilustrar este tipo de análisis.⁷ Cuanto más se acerque un coeficiente “path” a cero menor efecto tendrá (recordemos que son equivalentes en su interpretación al coeficiente de Pearson).

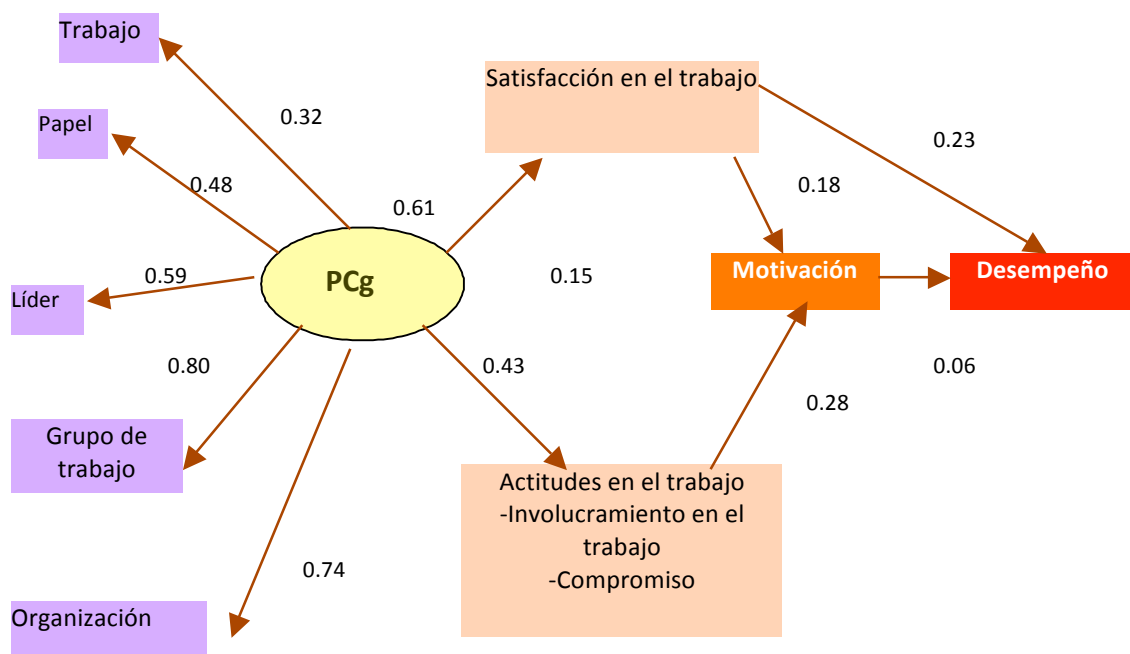


Figura 8.15 Ejemplo de un análisis “path” con el modelo de dos etapas del clima organizacional

⁷ Extraído de Parker *et al.* (2003), quienes buscaron por medio del análisis probar el modelo de dos etapas del clima organizacional, el cual ya se comentó en el texto.

Al visualizar el análisis completo nos damos cuenta de la mayor o menor influencia de unas variables sobre otras.

Para profundizar en la técnica, se sugiere a Loehlin (2009) y Webley y Lea (1997).

¿Qué es el análisis de factores?

Es un método estadístico multivariado que se utiliza para determinar el número y la naturaleza de un grupo de constructos subyacentes en un conjunto de mediciones. Un *constructo* es un atributo para explicar un fenómeno (Wiersma y Jurs, 2008). En este análisis se generan “variables artificiales” (denominadas factores) que representan constructos. Los *factores* se obtienen de las variables originales y deben ser interpretados de acuerdo con éstas. Como mencionan Naghi (1984) y Loehlin (2009), es una técnica para explicar un fenómeno complejo en función de unas cuantas variables. Es muy útil para la validez de constructo. Las variables deben estar medidas en un nivel por intervalos o razón.

Veamos varios ejemplos del empleo de esta técnica.⁸

Primer ejemplo: relación vendedor-comprador⁹

El estudio pretendió analizar los factores que determinan la relación entre los vendedores y los compradores industriales de las tres principales ciudades de México (Distrito Federal, Guadalajara y Monterrey). Se midieron variables entre las que destacan: coordinación (coord.), conflicto (confl.), frecuencia de la relación comprador-vendedor (frec.), reciprocidad económica en la relación (RF_2), reciprocidad en el manejo de consideraciones administrativas (RF_1) e importancia de la relación (monto de las operaciones) (impor.). Los resultados se muestran en la tabla 8.8.

⁸ En los ejemplos, no se incluyen todos los análisis, con el fin de no complicar su entendimiento.

⁹ Paniagua (1988) con la colaboración de los autores.

Tabla 8.8 Ejemplo de algunos resultados en un análisis de factores

MATRIZ DE PATRÓN FACTORIAL ELEGIDA SUBMUESTRA “COMPRAS”									
VARIABLE O COMPONENTE ESCALADO	ÍTEM	COMU- NALI- DAD	<i>FI</i>	<i>FII</i>	<i>FIII</i>	<i>FIV</i>	<i>FV</i>	<i>FVI</i>	FACTOR EN QUE CARGA
Coord.	23	.66997	.84392	-.00895	-.11828	-.03405	.06502	-.27645	<i>FI</i>
	24	.46446	.71642	-.05609	-.01958	.07106	.00043	-.07127	
	25	.59113	.67853	-.11175	.02581	-.09507	-.02857	.14978	
	26	.63988	.74737	.04695	.13472	-.04837	.07117	.02225	
Confl.	47	.55724	-.05110	.62553	-.20945	-.05248	-.27963	-.06387	<i>FII</i>
	48	.59072	.06230	.65163	.17884	-.10916	.31061	.04245	
	49	.35530	-.05665	.55503	-.11163	-.12946	-.07981	-.17370	
	50	.54716	-.07908	.61007	.08675	.20431	-.24058	.14142	
Frec.	15	.46081	-.14748	-.06276	.63660	-.12476	-.06299	.13488	<i>FIII</i>
	16	.44245	.03606	.08501	.62797	.00401	.11692	-.03968	
	18	.50307	-.05359	-.02787	.71169	.03786	-.10795	-.06775	
	19	.71311	.06664	.02881	.84213	.07161	-.16806	-.11058	
RF2	42	.46781	-.01380	-.18030	.09416	.63133	.17716	.06309	<i>FIV</i>
	43	.50097	.10175	-.07970	-.16207	.64202	-.02808	-.18530	
RFI	40	.72202	.01579	-.03548	.04181	-.18914	.77312	.14292	<i>FV</i>
	41	.48405	.15684	-.18489	-.06425	.01958	.58187	.19379	

Impor.	53	.31524	.02822	.02945	-.10069	.08605	.01579	.55431	
	54	.44550	-.04376	.08383	.01731	-.18396	.13956	.58137	<i>FVI</i>
	58	.53236	.26836	-.05219	.10026	-.11741	-.02893	.55080	
EIGENVALUE		5.36433	2.53081	2.47621	1.55248	1.23464	1.06932		
% VARIANZA EXPLICADA DELTA =.00		37.7%	17.8%	17.4%	10.9%	8.7%	7.5%	T = 100%	
<i>FI</i> = Coordinación(explica 37.7% de la varianza) <i>FII</i> = Conflicto(explica 17.8% de la varianza) <i>FIII</i> = Frecuencia(explica 17.4% de la varianza) Y así sucesivamente.									

Eigenvalue: representa la cantidad de varianza con que contribuye cada factor.

Observe que debajo de las columnas *FI* a *FVI* aparecen coeficientes que corresponden a los ítems de una escala. Si estos coeficientes son medios o elevados, se dice que los ítems *cargan* o forman parte del factor correspondiente. Por ejemplo, los ítems 23, 24, 25 y 26 cargan en el primer factor (obtienen valores de 0.84392, 0.71642, 0.67853 y 0.74737, respectivamente) y *no* pesan o cargan en otros factores (tienen valores bajos). Así, descubrimos una estructura de seis factores (*F*) en 19 ítems. Los *factores* reciben un nombre para saber qué constructos se encuentran subyacentes (el cual debe reflejar al factor y generalmente se extrae de la teoría). El análisis de factores también proporciona la varianza explicada y puede diagramarse en forma gráfica en las coordenadas *X* y *Y*.

Segundo ejemplo: escala del clima organizacional¹⁰

Para la validación del instrumento sobre el clima organizacional se consideraron varias muestras independientes. Entre éstas, un laboratorio químico farmacéutico y una institución educativa. El primero de 500 trabajadores, dos subunidades o centros de trabajo, con la inclusión de una planta y oficinas; 19 áreas funcionales y una antigüedad de más de 76 años. Se trata de una organización de alta tecnología y parte de un grupo corporativo internacional. El tamaño de muestra final fue de 421 casos válidos (*n*), 216 hombres y 186 mujeres (19 personas no especificaron). De los cuales 90% tienen 18 a 40 años (63% menores a 33); mientras que solamente 2% fue de nivel gerencial o mayor.

El instrumento pretendió medir las siguientes dimensiones: *moral, apoyo de la dirección, innovación, percepción de la empresa-identidad-identificación, comunicación, percepción del desempeño, motivación intrínseca, autonomía, satisfacción general, liderazgo, visión y recompensas o retribución*. Recordemos que lo que se evalúa son las percepciones sobre tales atributos organizacionales. Constó de 73 ítems con frases de la escala tipo Likert y 23 preguntas (cuyas respuestas se escalaron también tipo Likert), 96 reactivos en total.

Primero, se realizó un análisis sin rotar los factores (solución inicial o simple, sin consideraciones específicas). Los resultados completos de la rotación inicial se muestran en el apéndice respectivo, que produce 19 factores, pero debe hacerse notar que prácticamente los ítems “cargan” en un *solo* factor. El resto de los factores *no* muestra reactivos con “cargas” o “pesos” significativos. Posteriormente, los factores fueron “rotados” para visualizar posibles cambios (con los métodos *varimax, equamax* y *quartimax*). Los resultados prácticamente no variaron.

Esta solución de un único factor significativo y 18 factores sin pesos importantes, nos indica que la escala del clima organizacional es un constructo básicamente homogéneo, al menos en esta muestra. La solución factorial se presenta parcialmente en la tabla 8.9 (no se incluye toda, pues como ya se comentó exclusivamente un factor fue significativo, los restantes factores muestran pesos para los ítems como el segundo, tercero, cuarto y quinto factor).

¹⁰ Este ejemplo fue tomado de una aplicación del instrumento de la Universidad de Celaya para medir el clima laboral (Hernández Sampieri, 2005).

**Tabla 8.9 Cargas de factores en el ejemplo de la escala
para medir el clima organizacional (laboratorio químico-farmacéutico)**

Pregunta	Factor o componente				
Frases Likert	Factor 1	Factor 2	Factor 3	Factor 4	Factor 5
F1	.352	.286	-.276	-.160	.365
F2	.508	.382	-7.527E-03	-1.853E-02	.245
F3	.511	.211	.304	.153	.153
F4	.555	.359	-1.285E-02	-4.903E-02	.247
F5	.631	.325	-.120	-.137	4.398E-02
F6	.586	.312	-.121	-.287	-4.812E-02
F7	.615	-.224	.162	-.262	-6.974E-02
F8	.595	-.165	.125	-.330	4.410E-02
F9	.609	-.272	.325	-.296	-5.500E-03
F10	.655	-.235	.294	-.293	-2.404E-02
F11	.659	8.963E-02	.140	3.780E-02	-.167
F12	.589	.152	-.161	-5.420E-02	-.107
F13	.591	-.217	.189	.231	5.625E-02
F14	.636	-.198	.113	.237	3.174E-02
F15	.675	-.217	5.034E-02	-7.586E-02	4.885E-02
F16	.646	-.166	.243	-.215	3.701E-02
F17	.651	.164	.213	-7.052E-02	-8.041E-03
F18	.534	.328	.269	.276	-7.664E-02
F19	.690	-3.630E-02	-9.095E-05	-6.007E-02	7.306E-02
F20	.590	-9.375E-02	-6.703E-02	.359	3.371E-02
F21	.727	-.150	-.404	5.516E-02	-8.518E-03
F22	.765	-.213	-.389	2.251E-02	-5.801E-03
F23	.649	-.211	2.260E-02	.141	3.218E-02
F24	.656	.335	-8.049E-02	-1.521E-02	.211
F25	.534	-9.697E-03	.342	7.291E-02	-.135
F26NEG	-2.383E-02	3.496E-02	.124	.187	.280

F27	.592	-.257	4.450E-02	.410	-3.095E-02
F28	.593	.231	.216	.384	-.123
F29	.398	.103	-8.613E-02	.326	-.170
F30	.677	-8.654E-02	-.223	-5.095E-02	-3.149E-02
F31NEG	.236	.210	.114	.102	.333
F32	.673	.317	5.273E-02	-3.608E-02	.204
F33	.657	-.276	.226	-.277	-8.926E-02
F34	.604	.397	3.055E-02	-1.101E-02	1.358E-02
F35	.547	.417	3.127E-02	-2.232E-04	-3.890E-02
F36	.669	-.256	-9.381E-02	.296	4.097E-04
F37NEG	.163	-.144	-.254	.161	.367
F38NEG	.555	-.176	-.255	1.392E-02	.226
F39	.701	.312	-9.353E-02	-.209	.184
F40	.643	.412	-.144	-.149	.130
F41	.730	-.269	.235	-.210	2.546E-02
F42NOUNI	.518	-.336	.161	-.167	.255
F43	.229	6.211E-02	-3.422E-03	-2.360E-02	9.347E-02
F44NEG	.246	-.223	-.105	.263	.292
F45NEG	8.139E-02	-.207	-.170	7.145E-02	.404
F46	.642	-.141	-.339	3.685E-02	-.175
F47	.764	-.155	-.338	-5.616E-02	2.326E-02
F48	.612	-.186	-.359	-.192	-8.310E-02
F49	.720	-.148	-.339	-.105	-.117
F50NOUNI	.505	-.339	.191	-9.964E-02	.260
F51	.676	.389	-2.925E-02	-5.744E-02	.226
F52NEG	.376	-.164	-6.835E-02	.239	.363
F53NEG	.156	-.214	.187	.336	.244
F54	.542	.128	.117	6.809E-02	-.115
F55	.509	.344	.233	.333	-.101
F56	.467	1.753E-02	-.273	.343	-.132
F57	.528	.393	5.363E-02	.321	-6.305E-02

F58	.617	7.204E-02	-.184	7.046E-02	-.256
F59	.737	-.114	-.448	-8.039E-04	-.138
F60	.584	.181	1.196E-03	.341	-3.219E-02
F61	.395	-2.439E-02	-1.207E-02	.330	-3.374E-02
F62NEG	.424	-.241	-.308	6.336E-02	.213
F63	.684	-.138	.132	-1.317E-02	3.409E-02
F64	.565	-.142	.235	2.183E-02	-3.918E-02
F65	.540	-6.075E-02	.154	.189	-6.796E-02
F66	.746	-.171	-.341	-6.588E-02	-6.952E-02
F67	.742	4.266E-02	-.249	-9.802E-02	-2.049E-02
F68	.469	-.167	.267	.180	-.137
F69	.400	6.046E-02	9.498E-02	.430	-.135
F70NOUNI	.488	-.284	.297	-.115	.221
F71NOUN N	4.747E-02	-.153	.129	.113	.324
F72NOUNI	.272	-.132	7.687E-02	-8.865E-02	-.108
F73NOUNI	.555	-.205	.275	-6.823E-02	6.611E-03
Preguntas P1	.588	4.995E-02	.345	4.504E-02	7.539E-02
P2	.482	.396	-.230	-.122	.336
P3	.576	.397	2.461E-02	-.124	.272
P4	.738	-.208	.202	-.266	-7.068E-02
P5	.729	-.236	.180	-.239	3.344E-02
P6	.651	-5.214E-02	.152	.169	9.607E-02
P7	.687	-.203	.334	-1.692E-02	-3.240E-03
P8	.714	5.413E-03	.424	4.213E-03	-5.577E-02
P9	.700	.375	9.235E-03	-9.991E-02	.100
P10	.485	.403	4.916E-03	-4.660E-02	2.116E-03
P11	.564	.275	9.553E-02	-3.650E-02	-.135
P12	.708	-.123	.161	-.104	6.737E-02
P13	.668	-.143	-.133	-9.723E-02	.102

P14	.653	-.367	5.745E-02	.311	7.085E-02
P15	.767	-.225	-.405	6.120E-02	-.112
P16	.759	-.188	-.431	3.359E-02	-.118
P17	.792	-.196	-.408	-2.678E-03	-.106
P18	.727	5.139E-02	-3.487E-03	-1.964E-02	-.170
P19	.718	.117	-1.238E-02	-2.477E-02	-.237
P20	.702	6.121E-02	2.808E-03	-4.372E-02	-.282
P21	.653	.246	8.018E-02	-.169	-.265
P22	.661	.284	4.514E-02	-.102	-.224
P23	.660	.117	2.673E-02	-5.360E-02	-.267

En tablas o cuadros como el anterior resulta muy fácil “perderse en un mundo de cifras”. Por ello es necesario concentrarse en lo relevante:

- La tabla es una matriz de correlaciones de cada ítem con los factores (estructura completa de ítems o reactivos).
- Las cargas factoriales (valores en las celdas) son una especie de coeficientes de correlación y en consecuencia se interpretan como tales.
- El análisis de factores también proporciona la varianza explicada por cada factor de la estructura, así como otros valores que omitimos deliberadamente para no complicar la interpretación al estudiante de nivel licenciatura.
- El factor 1 (sombreado) es el único que realmente emergió de la administración de la escala del clima organizacional¹¹ (véanse los valores para cada ítem del factor 1 comparados con los valores de los demás factores en los mismos ítems).

Un ejemplo (pregunta 17):

¹¹ Da cuenta de 39% de la varianza total (los tres primeros factores generan 60%). Al factor se le denominó: Proceso de juicio común para evaluar las percepciones del entorno laboral.

Ítem	Factor 1	Demás factores			
P17	.792	-.196	-.408	-2.678E-03	-.106

The diagram illustrates the factor loadings for item P17. A box labeled 'carga factorial alta' (high factor loading) has an arrow pointing to the value .792 in the Factor 1 column. A bracket groups the loadings for other factors (-.196, -.408, -2.678E-03, -.106), with an arrow pointing to a box labeled 'cargas factoriales bajas' (low factor loadings).

- Los ítems con cargas bajas en todos los factores deben desecharse de la escala, no miden realmente lo que nos interesa: los componentes o dimensiones del clima organizacional (afectan la validez del instrumento). Por ejemplo, el ítem 53 (frase):

Ítem	Factor 1	Demás factores			
F53NEG	.156	-.214	.187	.336	.244

The diagram illustrates the factor loadings for item F53NEG. A box labeled 'carga factorial baja' (low factor loading) has an arrow pointing to the value .156 in the Factor 1 column. A bracket groups the loadings for other factors (-.214, .187, .336, .244), with an arrow pointing to a box labeled 'cargas factoriales bajas' (low factor loadings).

- En este caso, la conclusión principal del resultado sería:

El hecho de que el análisis de factores haya revelado un único factor significativo en la muestra, nos lleva a la conclusión provisional de que el *clima organizacional* es un constructo “molar”, en el cual se “funden” distintas percepciones sobre aspectos centrales del ambiente de trabajo.

Por lo anterior, hemos de decir que los resultados respaldan la noción de Parker et al. (2003) respecto a que, detrás de las dimensiones laborales del clima, se encuentra presente un “proceso de juicio común”, el cual se refleja en las distintas mediciones de la percepción del entorno de trabajo. Es un “proceso subyacente” que se expresa de diversas maneras.

Asimismo, el análisis de factores y la matriz de correlación entre dimensiones apoyan el modelo de dos niveles postulado principalmente por L.A. James, L.R. James y C.P. Parker.

La confiabilidad *alfa* para la escala fue 0.9747, que aumenta a 0.98 si se eliminan los ítems que no cargan en algún factor.

Tercer ejemplo: validación de un instrumento para medir el espíritu empresarial en estudiantes¹²

Objetivo del estudio: validar un instrumento que evalúa el espíritu empresarial en una muestra de estudiantes mexicanos.

Instrumento desarrollado en 2003 por Leslie Borjas Parra, Universidad Metropolitana de Venezuela.

Confiabilidad total (*alfa*): 0.925.

A partir del análisis de factores por componentes principales, se puede observar una carga mayor a 0.50 de 24 ítems hacia el factor 1 (F_1), lo que proporciona confianza respecto de que el instrumento realmente mide lo que pretende (validez de constructo).

Tabla 8.10 Ejemplo de análisis de factores para la validación de un instrumento (espíritu empresarial en estudiantes)

PREGUNTA	F_1	F_2	F_3	F_4	F_5	F_6	F_7	F_8
P9_CREATIVIDAD	0.60	-0.40	0.26	0.19	0.02	-0.22	-0.03	0.12
P8_GRUPO	0.53	-0.23	-0.24	0.09	0.24	0.01	-0.39	-0.31
P7_CREATIVIDAD	0.63	-0.06	0.27	-0.18	0.29	-0.15	-0.17	-0.22
P6_RIESGOS	0.63	-0.43	0.02	-0.14	-0.04	-0.19	-0.19	-0.14
P5_RIESGOS	0.54	-0.44	-0.15	-0.11	0.06	0.17	-0.13	0.10
P4_EVENTOS	0.57	0.00	0.31	-0.17	0.02	0.27	-0.07	0.39
P32_GRUPO	0.59	-0.43	-0.24	0.07	-0.12	0.01	-0.12	0.24
P31_EVENTOS	0.48	0.46	0.30	-0.25	0.00	0.04	0.26	-0.04
P30_HONESTIDAD	0.54	0.05	0.31	0.35	-0.38	0.02	0.17	0.26
P3_CREATIVIDAD	0.55	-0.15	0.21	-0.09	-0.49	0.02	0.08	-0.06
P29_EVENTOS	0.48	0.57	0.25	0.35	0.05	-0.09	-0.08	-0.06

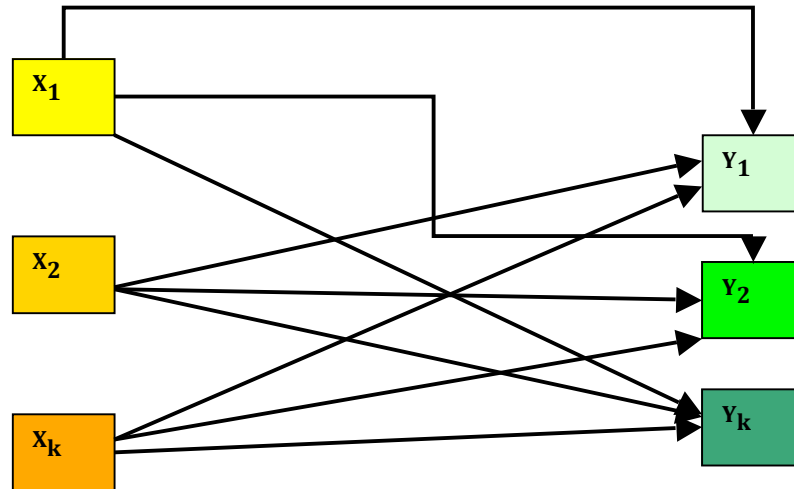
¹² Este ejemplo fue proporcionado por Moisés Mendoza Escobar (2006).

P28_AUTODETERMINACI	0.52	0.29	-0.01	0.46	0.01	-0.22	-0.03	0.23
P27_SOCIAL	0.39	0.32	-0.02	0.61	0.21	0.12	0.00	-0.23
P26_AUTODETERMINACI	0.64	0.06	0.38	0.16	0.06	-0.28	0.08	0.06
P25_AUTODETERMINACI	0.68	0.08	0.28	-0.04	-0.22	-0.08	-0.12	-0.10
P24_GRUPO	0.60	0.21	-0.12	-0.33	-0.01	-0.23	0.13	-0.15
P23_CAMBIO	0.59	-0.16	0.06	-0.26	0.25	-0.31	0.04	0.32
P22_HONESTIDAD	0.60	-0.05	-0.43	0.04	-0.26	0.08	0.26	-0.20
P21_HONESTIDAD	0.41	0.08	-0.17	-0.07	0.46	-0.26	0.47	-0.15
P20_SOCIAL	0.53	0.20	-0.38	0.19	-0.26	-0.22	-0.19	-0.03
P2_CREATIVIDAD	0.64	0.04	-0.26	0.27	0.24	0.20	-0.01	0.07
P19_SOCIAL	0.49	0.40	-0.37	-0.32	0.06	-0.05	-0.03	0.37
P18_CAMBIO	0.70	0.07	-0.17	-0.17	-0.05	-0.11	0.07	-0.10
P17_SOCIAL	0.64	0.42	-0.21	-0.25	-0.13	-0.02	0.02	-0.07
P16_GRUPO	0.68	0.19	-0.27	-0.17	-0.26	0.16	0.00	-0.06
P15_CAMBIO	0.56	-0.15	-0.40	0.19	-0.15	0.31	0.20	-0.03
P14_RIESGOS	0.58	-0.21	0.36	-0.23	0.04	0.39	0.08	-0.13
P13_RIESGOS	0.57	0.17	0.28	-0.04	0.21	0.35	0.02	-0.06
P12_CAMBIO	0.72	-0.42	0.09	0.05	-0.01	-0.08	-0.23	-0.12
P11_EVENTOS	0.36	0.59	0.14	-0.12	0.14	0.32	-0.35	0.03
P10_AUTODETERMINACI	0.31	-0.48	0.31	0.14	0.04	0.12	0.37	-0.12
P1_HONESTIDAD	0.45	-0.26	-0.33	0.11	0.41	0.17	0.15	0.25

Para quien desee compenetrarse con esta técnica recomendamos consultar Sengupta (2009), Brown (2006), Harman (1976), Gorsuch (1983), Nie et al. (1975), Kim y Mueller (1978a y 1978b), así como Hunter (1980). Del mismo modo, para aplicarlos se sugiere revisar a Cooper y Curtis (1976), y en español a Padua (2004) y otras referencias más sobre el paquete SPSS. Aunque es requisito conocer el programa estadístico computacional.

¿Qué es el análisis multivariado de varianza (MANOVA)?

Es un modelo para analizar la relación entre una o más *variables independientes* y dos o más *variables dependientes*. Es decir, es útil para *estructuras causales* del tipo:



La técnica posee varios usos, entre los que destacan:

- Evaluar diferencias entre grupos a través de múltiples variables dependientes (medidas por intervalos o razón). La(s) variable(s) independiente(s) es(son) categórica(s) (no métricas). Tiene el poder de evaluar no solamente las diferencias totales, sino diferencias entre las combinaciones de las dependientes.

En este sentido representa una extensión del análisis de varianza (ANOVA) para cubrir casos donde hay más de una variable dependiente y/o cuando las variables dependientes simplemente no pueden ser combinadas. En otras palabras, reconoce si los cambios en la(s) variable(s) independiente(s) tienen un efecto significativo en las dependientes. Señala qué grupos difieren en una variable o en el conjunto de variables dependientes (Steyn y Ellis, 2009).

- Identificar las interacciones entre las variables independientes y la asociación entre las dependientes.

Las tres clases principales del MANOVA son:

1. Hotelling's T. Es parecida a la prueba t (dos grupos) pero con más dependientes: una variable independiente dicotómica y varias dependientes.
2. MANOVA unidireccional. Análogo al ANOVA de una sola vía, pero con más dependientes: una variable independiente multicategórica y varias dependientes.
3. MANOVA factorial. Similar al ANOVA factorial, solamente que con dos o más dependientes: varias independientes categóricas y varias dependientes.

Los modelos del MANOVA tienen en común que forman combinaciones lineales de las dependientes que discriminan mejor entre los grupos en un experimento o una situación no experimental. Es una prueba de significancia de las diferencias en los grupos en un espacio multidimensional donde cada dimensión está definida por combinaciones lineales del conjunto de variables dependientes.

Una pregunta que suele hacer el estudiante al revisar el MANOVA es ¿por qué no hacemos ANOVAS separados, uno para cada dependiente? La respuesta: las dependientes están correlacionadas muy frecuentemente, por lo cual los resultados de varios ANOVA pueden ser redundantes y difíciles de integrar. He aquí una síntesis de la explicación de Wiersma y Jurs (2008) sobre este tipo de análisis:

Al incluir dos o más variables dependientes simultáneamente no se consideran las diferencias entre las medias en cada variable, sino las diferencias en variables canónicas. El interés no sólo es saber si los grupos definidos por las variables independientes difieren en las *variables canónicas*, sino conocer la naturaleza de éstas. Una *variable canónica* es una variable artificial generada a partir de los datos. Representa *constructos* y se compone de variables reales, las cuales deben ser descritas en términos de variables dependientes. Lo anterior se efectúa por medio de las cargas de los coeficientes de correlación entre una variable dependiente y una variable canónica. Si una carga entre la *variable canónica* y la dependiente es positiva y elevada, significa que altos valores en la dependiente se asocian con altos valores en la canónica. Por ejemplo, si una variable dependiente consiste en puntuaciones a una prueba sobre innovación, y dichas puntuaciones se correlacionan en forma considerable con una

variable canónica, inferimos que la variable canónica representa un constructo que involucra esencialmente a la innovación.

En los cálculos que se hacen en el MANOVA, se generan variables canónicas hasta que se encuentra que *no* hay una diferencia estadística significativa entre las categorías o los grupos de las variables independientes; o bien, hasta que se agotan los grados de libertad de las variables independientes (lo que ocurra primero). El número de variables canónicas no puede exceder el número de variables dependientes, pero es común que el número de dependientes sea mayor que el de *variables canónicas* estadísticamente significativas o los grados de libertad.

La hipótesis general de investigación en el MANOVA postula que las medias de los grupos o las categorías de la(s) variable(s) independiente(s) difieren entre sí en las variables canónicas. La hipótesis nula postula que dichas medias serán iguales. Se calculan diversas estadísticas para evaluar ambas hipótesis, entre las que destacan: F (total, toma en cuenta el modelo completo), la prueba *Hotelling's T-Square*, T^2 (cuando hay dos grupos formados por las variables independientes), *Wilks' lambda*, U (cuando hay más de dos grupos formados por las variables independientes), y Pillai-Bartlett (cuando hay coeficientes canónicos);¹³ y si resultan significativas en un nivel de confianza, se acepta la hipótesis de investigación de diferencia de medias. Esto indica que hay, por lo menos, una variable canónica significativa (pero puede haber varias). Si diversas variables canónicas son significativas, esto muestra que se presentan diferencias en las variables canónicas en cuestión, entre los grupos o categorías de las independientes.

Los paquetes estadísticos que contiene el MANOVA suelen posicionar a los grupos de las variables independientes por puntuaciones discriminantes; éstas se calculan con una función discriminante, que es una ecuación de regresión para un compuesto de variables dependientes. A cada grupo se le asigna una puntuación discriminante en cada variable canónica. Las puntuaciones discriminantes de una variable independiente pueden ser cero o tener un valor positivo o negativo. Una puntuación discriminante positiva y elevada para un grupo indica que éste se coloca por encima de los demás en la respectiva variable canónica. Y deben considerarse las cargas, las cuales son positivas o

¹³ Además, para comparaciones específicas están las pruebas *post hoc* del ANOVA, ya revisadas en el capítulo 10 del libro.

negativas. Las puntuaciones discriminantes se usan para interpretar las separaciones de los grupos en las variables canónicas, en tanto que las cargas se usan para evaluar y ligar los resultados de las variables dependientes (Wiersma y Jurs, 2008). Un ejemplo de las cargas de los coeficientes de correlación entre las variables dependientes y las variables canónicas se muestra en la tabla 8.11, mientras que un ejemplo de las *puntuaciones discriminantes* viene en la tabla 8.12.

Tabla 8.11 Cargas de los coeficientes de correlación entre las variables dependientes y las variables canónicas

VARIABLE DEPENDIENTE	VARIABLES CANÓNICAS		
	I (MOTIVACIÓN INTRÍNSECA)	II (ATRIBUCIÓN DE CAUSALIDAD EXTERNA)	III (DESEMPEÑO LABORAL)
— Motivación intrínseca (escala intrínseca del inventario de características del trabajo)	0.90	0.05	0.07
— Atribuciones internas	0.86	-0.07	-0.09
— Sentimientos de éxito en el trabajo	0.70	0.80	0.00
— Atribuciones externas	0.03	0.61	0.12
— Productividad	-0.12	0.07	0.74
— Eficiencia	0.18	-0.04	0.48
— Calidad	0.19	0.13	0.57

Tabla 8.12 Puntuaciones discriminantes con cuatro grupos en tres variables canónicas

GRUPO	VARIABLES CANÓNICAS		
	I	II	III
Ejecutivos	1.97	0.95	-1.69
Secretarias	-0.19	1.18	1.25
Empleados	1.40	-1.01	-0.49
Obreros	-3.18	-1.12	0.93

Como observamos en la tabla 8.11, se obtuvieron tres *constructos subyacentes* en las puntuaciones recolectadas de la muestra: motivación intrínseca, atribución de causalidad externa y desempeño laboral. Vemos en la tabla 8.12 que los grupos (niveles en la empresa) están separados en las tres variables canónicas (los grupos difieren), particularmente en la primera variable canónica (motivación intrínseca) y los obreros ocupan la posición más baja. Las variables dependientes enmarcadas en un recuadro en la primera variable canónica se cargan en ella (tabla 8.11); en consecuencia, los ejecutivos tienen las puntuaciones más altas en motivación intrínseca medida por la escala mencionada, en atribuciones internas y en sentimientos de éxito en el trabajo. Así se interpretan todas las variables canónicas y dependientes.

En el MANOVA se incluyen *razones F* y *análisis de varianza*. Algunos paquetes estadísticos agregan una prueba denominada *correlación canónica*, que es muy similar al MANOVA. Ésta es la máxima correlación que llega a obtenerse entre los conjuntos de puntuaciones y las relaciones entre las variables independientes, entre las variables dependientes y entre los conjuntos de ambas (dependientes e independientes) (Kerlinger, 1979). Las variables en el MANOVA y la correlación canónica asumen que las variables dependientes están medidas en un nivel de intervalos o razón. Tal correlación se interpreta como otras; pero el contexto de interpretación varía de acuerdo con el número de variables involucradas.

Distancias euclideanas

Es una medida para analizar “distancias” entre constructos (agrupamiento: cercanía o lejanía). En el programa SPSS se pueden obtener básicamente dos valores opuestos: distancia euclidiana o de disimilitud y distancia de proximidad o similitud. Entre mayor sea el valor de una distancia euclidiana, más “alejadas” se encuentran las variables; y entre más alto sea el valor de proximidad, “más cercanas y relacionadas” están las variables (este último valor se proporciona en unidades de correlación de Pearson).

Por ejemplo, Hernández Sampieri (2009) midió el clima organizacional en función del Modelo de los Valores en Competencia (MVC) de Kim S. Cameron y Robert E. Quinn. Consideró cuatro dimensiones del clima organizacional en relación con la cultura organizacional (denominados cuadrantes en el MVC) que a su vez contienen diversas variables:

- 1) Cuadrante de organización familiar o clan (relaciones humanas), constituido por las siguientes variables: bienestar de los empleados, autonomía, comunicación, énfasis en el entrenamiento, integración y soporte del superior inmediato.
- 2) Cuadrante de organización jerárquica (proceso interno): formalización y tradición.
- 3) Cuadrante de organización adhocrática (sistemas): flexibilidad/innovación, enfoque externo y reflexividad.
- 4) Cuadrante que corresponde a la organización de mercado (metas racionales): claridad de metas organizacionales, esfuerzo, eficiencia, calidad, presión para producir y realimentación del desempeño.

Con la finalidad de analizar la proximidad de tales dimensiones o cuadrantes realizó un análisis de disimilitudes y otro de similitudes, los cuales se presentan en las tablas 8.13 Y 8.14.

Tabla 8.13 Matriz de disimilitudes (distancias euclidianas).

	Relaciones humanas	Proceso interno	Sistemas	Metas
Relaciones humanas	.000	24.149	10.710	13.097
Proceso interno	24.149	.000	23.727	17.413
Sistemas	10.710	23.727	.000	11.819
Metas	13.097	17.413	11.819	.000

Tabla 8.14 Matriz de proximidad o similitudes (en unidades de correlación de Pearson).

Correlaciones entre vectores				
	Relaciones humanas	Proceso interno	Sistemas	Metas
Relaciones humanas	1.000	-.164	.796	.666
Proceso interno	-.164	1.000	-.133	.062
Sistemas	.796	-.133	1.000	.739
Metas	.666	.062	.739	1.000

En las tablas podemos observar que “relaciones humanas”, “sistemas” y “metas” son dimensiones o cuadrantes “ceranos y asociados”. “Proceso interno” es una dimensión “lejana” al resto de los cuadrantes.

¿Hay otros métodos multivariados?

En la actualidad, hay muchos métodos multivariados de análisis, los cuales se desarrollaron con la evolución de la computadora. Los investigadores disponen del *análisis discriminante*, cuando las variables independientes se miden por intervalos o razón, y la dependiente es categórica. Tal análisis sirve para predecir la pertenencia de un caso a una de las categorías de la variable dependiente, sobre la base de varias independientes (dos o más). Se utiliza una ecuación de regresión llamada *función*

discriminante. Por ejemplo, si queremos predecir el voto obtenido por dos partidos contendientes (variable dependiente nominal con dos categorías) sobre la base de cuatro variables independientes, aplicaremos el análisis discriminante, para resolver una ecuación de regresión; así se obtienen las predicciones individuales. En el ejemplo, hay dos categorías (votar por *A* o votar por *B*); por lo tanto, los valores a predecir son 0 y 1 (*A* y *B*, respectivamente). Si el sujeto obtiene una puntuación más cercana a cero, se predice que pertenece al grupo que votará por *A*; si logra una puntuación más cercana a 1, se predice que pertenece al grupo que votará por *B*. Además, se consigue una medida del grado de discriminación del modelo.

Se cuenta también con el análisis de conglomerados o *clusters* (técnica para agrupar los casos o elementos de una muestra en grupos con base en una o más variables), el escalamiento multidimensional (para diseñar escalas que midan a los sujetos en diversas variables y los ubiquen simultáneamente en los ejes de las distintas variables, así como para conocer la estructura de las variables entre sí), el análisis de series cronológicas o de tiempo (para analizar la evolución de los casos en una o más variables a través del tiempo y predecir el comportamiento de las variables o sucesos) y la elaboración de mapas multidimensionales (donde establecemos distancias entre casos, al basarnos en mediciones múltiples de varias dimensiones o variables), para los cuales se requieren bases sólidas en materia de estadística y en matemáticas avanzadas. Sugerimos a Loehlin (2009), Sengupta (2009), Shaw (2009), Ferrán (2001) y Lehamn et al. (2005) para una revisión de tales pruebas y modelos. También se publica una revista fundamental en la materia: *Multivariate Behavioral Research*.

FÓRMULAS Y PROCEDIMIENTOS ESTADÍSTICOS

En este documento simplemente se incluyen las fórmulas y los procedimientos de cálculo de ciertas estadísticas que anteriormente se encontraban en el capítulo 10 del texto impreso, aunque insistimos que ya solamente se usan para ilustrar la definición de

ciertas estadísticas, porque hoy en día los cálculos y las estimaciones se efectúan mediante programas computarizados de análisis estadístico, como SPSS.

Fórmula y cálculo de la media o promedio

La fórmula de la media es:

$$\bar{X} = \frac{\bar{X}_1 + \bar{X}_2 + \bar{X}_3 + \bar{X}_k}{N}$$

Por ejemplo, si tuviéramos las siguientes puntuaciones:

8 7 6 4 3 2 6 9 8

La media sería igual a:

$$\bar{X} = \frac{8+7+6+4+3+2+6+9+8}{9} = 5.888$$

La fórmula simplificada de la media es:

$$\bar{X} = \frac{\sum X}{N}$$

El símbolo “ Σ ” indica que debe efectuarse una sumatoria, X es el símbolo de una puntuación y N es el número total de casos o puntuaciones. En nuestro ejemplo:

$$\bar{X} = \frac{53}{9} = 5.888$$

La media sí es sensible a valores extremos. Si tuviéramos las siguientes puntuaciones:

8 7 6 4 3 2 6 9 20

La media sería:

$$\bar{X} = \frac{65}{9} = 7.22$$

Fórmula de la desviación estándar

La fórmula de la desviación estándar es:

$$s = \sqrt{\frac{\sum (X - \bar{X})^2}{N}}$$

Esto es, la desviación de cada puntuación respecto a la media se eleva al cuadrado, se suman todas las desviaciones cuadradas, se divide entre el número total de puntuaciones, y a esta división se le saca raíz cuadrada.

Cálculo de la prueba *t* con fórmulas y tablas

El valor *t* se obtiene en muestras grandes mediante la fórmula:

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

En donde $\bar{X}_1$ es la media del primer grupo, $\bar{X}_2$ la media del segundo grupo, s_1^2 representa la desviación estándar del primero elevada al cuadrado, n_1 es el tamaño del primer grupo, s_2^2 simboliza la desviación estándar del segundo grupo elevada al cuadrado y n_2 es el tamaño del segundo grupo. En realidad, el denominador es el *error estándar de la distribución muestral de la diferencia entre medias*.

Para saber si el valor *t* es significativo, se aplica la fórmula y se calculan los grados de libertad.

Recordemos que los grados de libertad se calculan así:

$$gl = (n_1 + n_2) - 2$$

Una vez calculados el valor t y los grados de libertad, se elige *el nivel de significancia* y se compara el valor obtenido contra el valor que le correspondería, en la tabla 2 del apéndice 4 de este CD anexo (*tabla de la distribución t de Student*) o en STATS® “Distribución “ t ” de Student. Si el valor calculado es igual o mayor al que aparece en la tabla, se acepta la hipótesis de investigación. Pero si es menor, se acepta la hipótesis nula.

En la tabla se busca el valor con el cual vamos a comparar el que hemos calculado, con base en el nivel de confianza elegido (0.05 o 0.01) y los grados de libertad. La tabla contiene los niveles de confianza como columnas y los grados de libertad como renglones. Los niveles de confianza adquieren el significado del que se ha hablado (el 0.05 implica 95% de que los grupos en realidad difieran significativamente entre sí y 5% de posibilidad de error).

Cuanto mayor sea el valor t calculado respecto al valor de la tabla y menor sea la posibilidad de error, mayor será la certeza en los resultados.

EJEMPLO

Hi: “los varones le atribuyen mayor importancia al atractivo físico en sus relaciones heterosexuales que las mujeres”.

Ho: “los varones no le atribuyen más importancia al atractivo físico en sus relaciones heterosexuales que las mujeres”.

La variable atractivo físico fue medida a través de una prueba estandarizada y el nivel de medición es por intervalos. La escala varía de 0 a 18.

La hipótesis se somete a prueba con los estudiantes de clase media de dos universidades de la ciudad de Monterrey.

n_1 (hombres) = 128

n_2 (mujeres) = 119

Resultados:

$$\overline{X}_1 \text{ (hombres)} = 15$$

$$\overline{X}_2 \text{ (mujeres)} = 12$$

$$s_1 \text{ (hombres)} = 4$$

$$s_2 \text{ (mujeres)} = 3$$

$$t = \frac{15 - 12}{\sqrt{\frac{(4)^2}{128} + \frac{(3)^2}{119}}}$$

$$t = 6.698$$

$$Gl = (128 + 119) - 2$$

$$Gl = 245$$

Al acudir a la tabla de la distribución t de *Student* (tabla 2, apéndice 4, incluido en este CD), buscamos los grados de libertad correspondientes y elegimos en la columna de gl , el renglón α , que se selecciona siempre que se tienen más de 200 grados de libertad. La tabla contiene los siguientes valores:

gl	0.05	0.01
α (mayor de 200)	1.645	2.326

Nuestro valor calculado de t es 6.698 y resulta superior al valor de la tabla en un nivel de confianza de 0.05 ($6.698 > 1.645$). Entonces, la conclusión es que aceptamos la hipótesis de investigación y rechazamos la nula. Incluso, el valor t calculado es superior en un nivel de confianza del 0.01 ($6.698 > 2.326$).

Comentario: Efectivamente, en el contexto de la investigación (del ejemplo anterior), los varones le atribuyen más importancia al atractivo físico en sus relaciones de pareja que las mujeres.

Si tuviéramos 60 grados de libertad y un valor t igual a 1.27, al comparar este valor con los de la tabla obtendríamos:

gl	0.05	0.01
60	1.6707	2.390

El valor t calculado es menor a los valores de la tabla. Se *rechaza la hipótesis de investigación y se acepta la hipótesis nula*.

CHI CUADRADA MEDIANTE STATS® Y EJEMPLO DE TABLA DE CONTINGENCIA PRODUCIDA POR SPSS CON OTROS ELEMENTOS

Para calcular la chi cuadrada en STATS®, primero se elige la opción correspondiente en el menú principal.

Después, en Cuadrícula colocamos los datos que se nos requiere (número de categorías de cada variable, las de la variable que representa a las filas y las de la variable que representa las columnas).

Supongamos que tenemos las dos variables del ejemplo del capítulo 10, *Intención del voto* (candidata A Guadalupe Torres y candidata B Liz Almanza) y *género* (masculino y femenino).

Tenemos dos categorías en cada variable, por lo tanto, la Cuadrícula se llenaría así:

Filas	2
Columnas	2

Damos clic en Aceptar. STATS® abre la tabla y nosotros colocamos las frecuencias observadas o contadas:

40	58
32	130

En el siguiente paso, STATS® nos pide que seleccionemos si los valores esperados están capturados o no. Tenemos dos opciones: marcar que no (los valores esperados no están calculados) o calcularlos. Si elegimos la primera, simplemente damos clic en Calcular y se nos proporcionan los siguientes resultados: chi-Cuadrada y grados de libertad, pero no la significancia. Tenemos que acudir a la tabla de chi-cuadrada (apéndice 4, tabla 4 de este CD o a la de STATS –Valores de α^2 a los niveles de confianza de .05 y .01–). Vemos nuestros grados de libertad y elegimos el nivel. Si es igual o superior al valor de la tabla, implica que chi cuadrada fue significativa dependiendo del nivel seleccionado (de nuevo, .05 o .01).

Si nuestra opción es: calcular los valores esperados están calculados, los obtenemos siguiendo este procedimiento:

La frecuencia esperada de cada celda, casilla o recuadro, se calcula mediante la siguiente *fórmula aplicada a la tabla de frecuencias observadas*.

$$f_e = \frac{(\text{Total o marginal de renglón})(\text{total o marginal de columna})}{n}$$

En donde n es el número total de frecuencias observadas.

Para la primera celda (candidata A y género masculino), la frecuencia esperada sería:

$$f_e = \frac{(98)(72)}{260} = 27.13$$

Veamos de dónde salieron los números:

	Género			Total
		Masculino	Femenino	
Intención del voto				98
Total		72		260

Para este ejemplo, la tabla de frecuencias esperadas sería la tabla 8.15.

Tabla 8.15 Cuadro de frecuencias esperadas para el ejemplo

	Género		Total
		Masculino	Femenino
Intención del voto	Candidata A Guadalupe Torres	27.1	70.9
	Candidata B Liz Almanza	44.9	117.1
Total		72.0	188.0

Se colocan los datos en el STATS® y damos clic en calcular, obteniendo el valor de chi-cuadrada, los grados de libertad y la probabilidad de que los valores observados y esperados sean distintos en porcentaje (si es mayor a 95% es significativo en este nivel, si es mayor a 99% resulta más significativa; recordemos que STATS® proporciona la significancia en términos del porcentaje en nuestro favor).

Si se hiciera manualmente, aplicaríamos la fórmula *de chi cuadrada*:

$$\chi^2 = \sum \frac{(O - E)^2}{E}$$

donde Σ significa sumatoria.

O es la frecuencia observada en cada celda.

E es la frecuencia esperada en cada celda.

Es decir, se calcula para cada celda la diferencia entre la frecuencia observada y la esperada; esta diferencia se eleva al cuadrado y se divide entre la frecuencia esperada. Finalmente, se suman tales resultados y la sumatoria es el valor de χ^2 obtenida.

Obtenemos el valor de chi-cuadrada y los grados de libertad:

$$gl = (r - 1) (c - 1)$$

En donde r es el número de renglones del cuadro de contingencia y c el número de columnas; en nuestro caso: 1.

Y acudimos con los grados de libertad que corresponden a la tabla 4 del apéndice 4 de este CD (Distribución de *chi* cuadrada) o STATS® en “Valores de α^2 a los niveles de confianza de .05 y .01”, eligiendo nuestro nivel de confianza (.05 o .01). Si nuestro valor calculado de α^2 es igual o superior al de la tabla, decimos que las variables están relacionadas (α^2 fue significativa).

Cabe señalar que en SPSS es más rápido, y en segundo lugar con STATS® y la opción “los valores esperados no están calculados” (porque solamente vemos en la tabla con nuestros grados de libertad si igualamos o superamos el valor).

En el capítulo 10 del libro, señalamos que SPSS produce un resumen de los casos válidos y perdidos para cada variable y una tabla de contingencia sencilla, o bien una tabla más compleja con diversos resultados por celda. Tal es el caso de la tabla 8.16, con el ejemplo de la *intención de voto* por Guadalupe o Liz y el *género* (el programa genera aún más, pero hemos reproducido solamente los que son más útiles para el lector en la interpretación de la tabla).

Frecuencia (*count*): número de casos o frecuencias observadas en cada celda.

Frecuencia esperada (*expected count*): número de casos en cada celda que se esperarían, si no hubiera relación entre variables.

Porcentajes dentro de cada celda que representan las frecuencias observadas (en relación con una variable, con la otra y con el total). Por ejemplo, la primera celda (intención por candidato A y género masculino = 40), representa 40.8% dentro de la intención del voto del candidato A (marginal horizontal o de su fila = 98), implica 55.6% de los hombres (marginal de su columna = 72), y 15.4% respecto al total (muestra = 260).¹⁴

¹⁴ Los valores son redondeados por el programa.

Tabla 8.16 Ejemplo de la tabla de contingencia producida por SPSS (resumida)

Intención de voto por candidato * Género del votante					
			Género del votante		Total
			Masculino	Femenino	
Intención de voto por candidato	Candidato A Guadalupe Torres	Frecuencia	40	58	98
		Frecuencia esperada	27.1	70.9	98
		% dentro de intención de voto	40.8%	59.2%	100.0%
		% dentro del género del votante	55.6%	30.9%	37.7%
		% del total	15.4%	22.3%	37.7%
	Candidato B Liz Almanza	Frecuencia	32	130	162
		Frecuencia esperada	44.9	117.1	162
		% dentro de intención de voto	19.8%	80.2%	100.0%
		% dentro del género del votante	44.4%	69.1%	62.3%
		% del total	12.3%	50.0%	62.3%
Total	Frecuencia		72	188	260
	Frecuencia esperada		72.0	188.0	260
	% dentro de intención de voto		27.7%	72.3%	100.0%
	% dentro del género del votante		100.0%	100.0%	100.0%
	% del total		27.7%	72.3%	100.0%

En este caso, el valor de chi cuadrada es significativo al nivel del .01, es decir, existe relación entre las variables *género* e *intención de voto* por los diferentes candidatos, se acepta la hipótesis de investigación (Liz Almanza gana, pero sobre todo por el voto femenino).

Con la tabla se hacen algunas conclusiones, como las que se presentaron en el apartado: ¿Qué otra aplicación tienen las tablas de contingencia?

EJEMPLOS DE SECUENCIAS DE ANÁLISIS EN MINITAB Y SPSS

Para completar este capítulo del CD y el 10 del texto impreso, incluimos dos ejemplos de secuencias de análisis de los datos.

Secuencia en Minitab

Esta secuencia corresponde al ejemplo desarrollado en el texto impreso sobre **la televisión y el niño**. La secuencia se presenta en la figura 8.16.

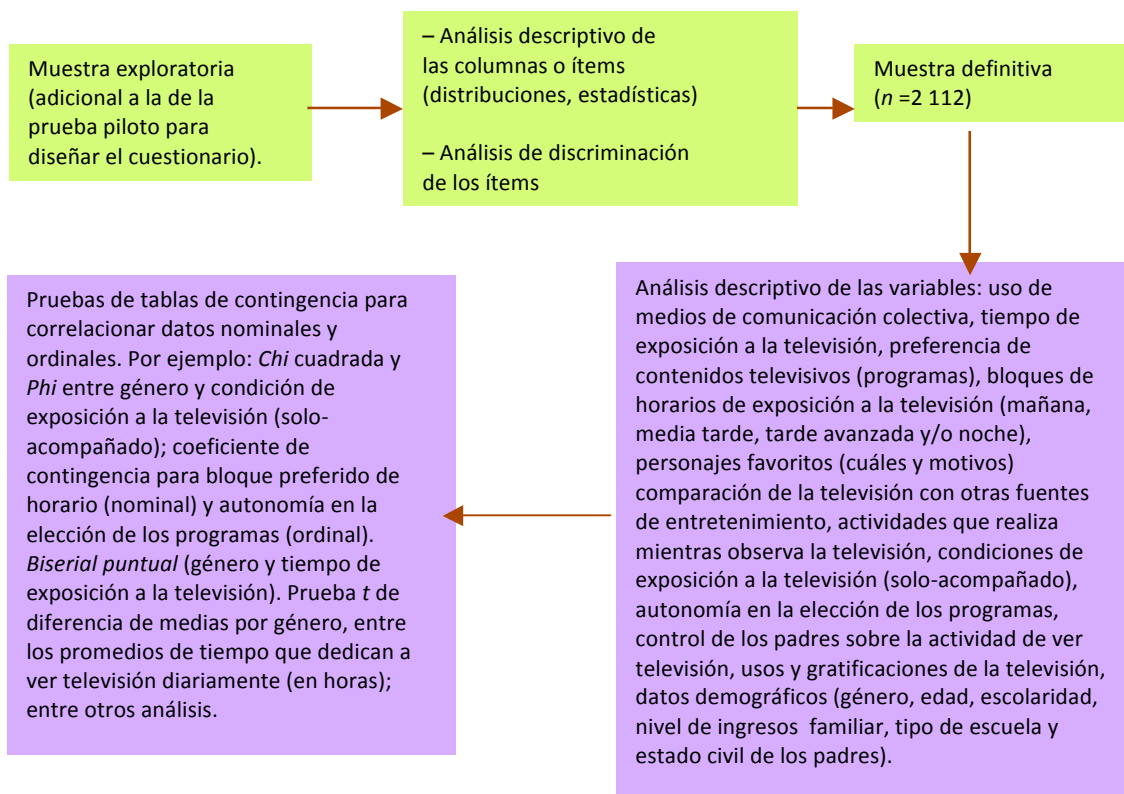


Figura 8.16 Secuencia de análisis con Minitab

Secuencia en SPSS

La secuencia corresponde al estudio respecto al **clima organizacional** citado en el libro impreso (Hernández Sampieri, 2005).

Objetivo central:

- Validar un instrumento para medir el clima organizacional en el ámbito laboral mexicano.

Pregunta de investigación:

- ¿Cuáles son las variables o dimensiones más pertinentes que deben incluirse para medir el clima laboral de una manera válida y confiable?

Comentarios sobre los análisis:

Los análisis del estudio se llevaron a cabo de manera independiente para cada muestra, debido a la naturaleza contingente del clima organizacional. Cada empresa tiene su propia historia, procesos de producción, estructura, orientación, filosofía y otros factores situacionales que la hacen única. Sin embargo, la validación en diferentes muestras va consolidando y robusteciendo al instrumento de medición. A continuación se presenta la secuencia de análisis y algunos resultados obtenidos en una de las muestras (una institución de educación superior).

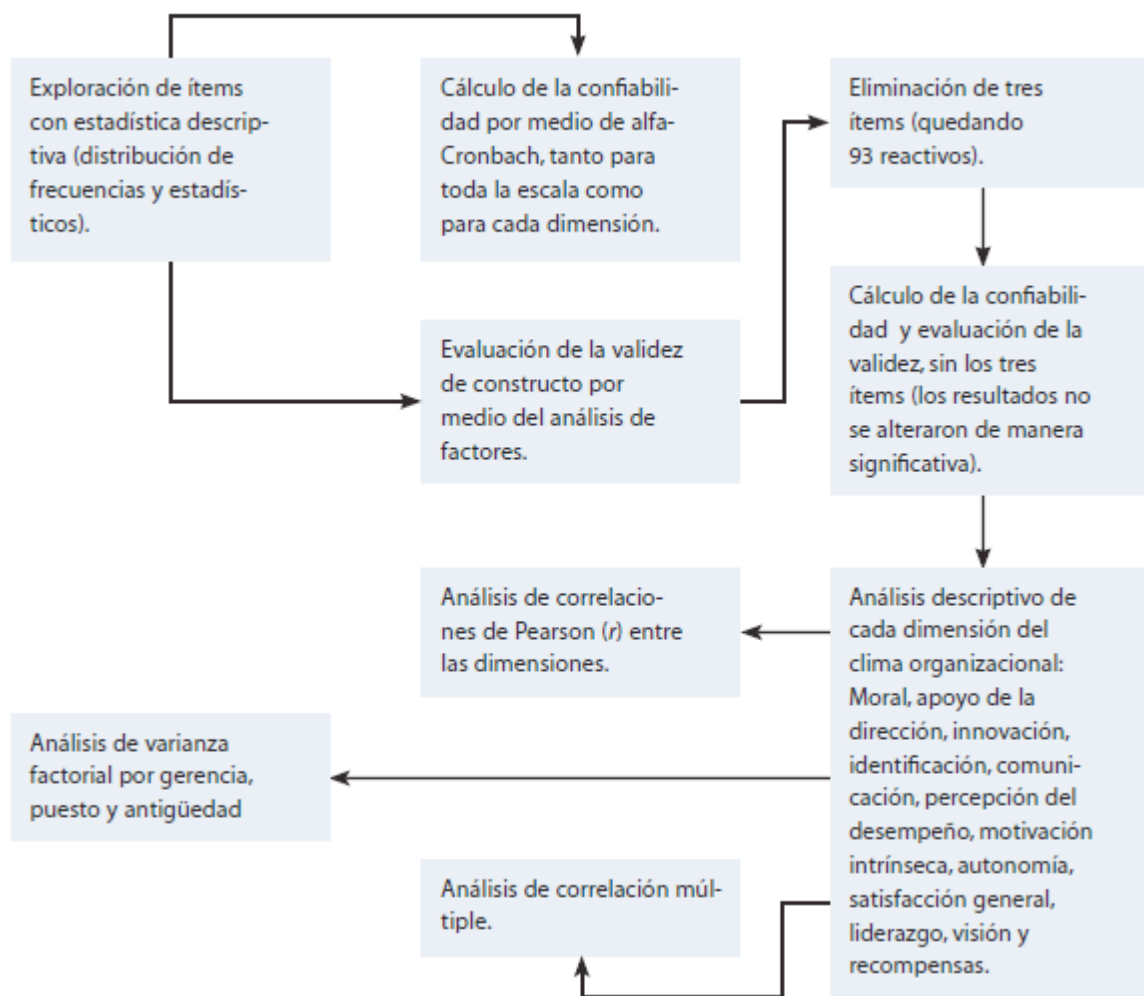


Figura 8.17 Secuencia de análisis con el SPSS.

La confiabilidad (*alfa*) del instrumento fue de 0.98 (muy elevada, $n = 163$). El análisis de factores reveló una dimensión única y después de efectuarlo, se eliminaron tres ítems. Los coeficientes *alfa* para las dimensiones, fueron los que se muestran en la tabla 8.17.

Estadística descriptiva en el nivel de toda la escala: en la tabla 8.18, se presentan las principales estadísticas de los resultados a toda la escala del clima organizacional. El promedio (3.7) y la mediana (3.8) son bastante favorables e indican que esta organización posee un clima organizacional positivo. La respectiva normalización, se presenta en la gráfica de la figura 8.18.

Las correlaciones entre los componentes del clima organizacional para esta segunda muestra, se presentan en la tabla 8.19 (matriz de correlaciones de Pearson).

La mayoría se encuentra entre rangos que oscilan entre 0.55 y 0.69, es decir, correlaciones medias y considerables. Destacan las correlaciones entre comunicación y dirección (0.800), satisfacción y liderazgo (0.772), satisfacción y visión (0.721), identificación y motivación intrínseca (0.716); así como, moral y liderazgo (0.698). Llama la atención la correlación entre autonomía y motivación intrínseca (0.613), que respalda la tradicional vinculación entre ambos conceptos, propuesta por los modelos sobre características del trabajo. Identificación y liderazgo es de las pocas correlaciones bajas (0.425).

Tabla 8.17 Los coeficientes *alfa* para las dimensiones del clima organizacional

Dimensión	<i>n</i>	<i>Alfa</i>
Moral	176	0.898
Dirección	176	0.928
Innovación	177	0.788
Identificación	175	0.823
Comunicación	179	0.827
Percepción del desempeño	174	0.709
Motivación intrínseca	178	0.792
Autonomía	179	0.827
Satisfacción	179	0.865
Liderazgo	177	0.947
Visión	176	0.902

Tabla 8.18 Estadística descriptiva de la muestra para la escala completa del clima organizacional

Estadística	Valor
Media	3.7
Mediana	3.8
Moda	3.96
Desviación estándar	0.62
Varianza	0.38
Asimetría	-0.442
Curtosis	-0.116
Mínimo	1.84
Máximo	4.92

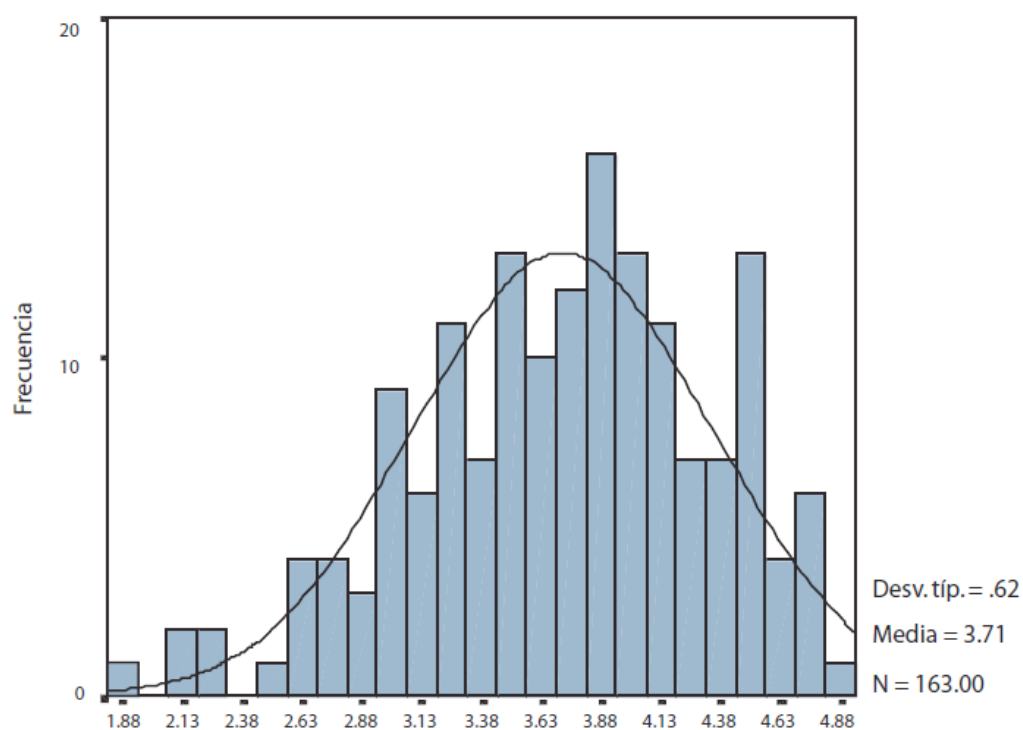


Figura 8.18 Histograma de toda la escala del clima organizacional en una institución educativa

Tabla 8.19 Matriz de correlaciones de Pearson entre variables del clima organizacional

		MOR	DIR	INN	IDENT	COM	DES	MOT	AUT	SAT	LID	VIS
MOR	Correlación de	1	0.510*	0.684*	0.466*	0.583*	0.673*	0.588*	0.491*	0.681*	0.698*	0.692*
	Sig. (bilateral)	.	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
	N	176	174	174	173	175	171	176	176	175	176	173
DIR	Correlación de	0.510*	1	0.604*	0.601*	0.800*	0.588*	0.632*	0.649*	0.603*	0.589*	0.644*
	Sig. (bilateral)	0.000	.	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
	N	174	176	173	172	175	170	175	176	175	175	174
INN	Correlación de	0.684*	0.604*	1	0.544*	0.555*	0.612*	0.653*	0.505*	0.636*	0.640*	0.661*
	Sig. (bilateral)	0.000	0.000	.	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
	N	174	173	177	173	177	172	176	176	176	174	173
IDEN	Correlación de	0.466*	0.601*	0.544*	1	0.606*	0.556*	0.716*	0.539*	0.616*	0.425*	0.625*
	Sig. (bilateral)	0.000	0.000	0.000	.	0.000	0.000	0.000	0.000	0.000	0.000	0.000
	N	173	172	173	175	174	170	175	174	174	173	171
COM	Correlación de	0.583*	0.800*	0.555*	0.606*	1	0.579*	0.550*	0.684*	0.659*	0.637*	0.647*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	.	0.000	0.000	0.000	0.000	0.000	0.000
	N	175	175	177	174	179	173	177	178	178	176	175
DES	Correlación de	0.673*	0.588*	0.612*	0.556*	0.579*	1	0.682*	0.466*	0.579*	0.470*	0.672*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	.	0.000	0.000	0.000	0.000	0.000
	N	171	170	172	170	173	174	173	173	174	171	171
MOT	Correlación de	0.588*	0.632*	0.653*	0.716*	0.550*	0.682*	1	0.613*	0.586*	0.532*	0.642*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	0.000	.	0.000	0.000	0.000	0.000
	N	176	175	176	175	177	173	178	177	177	176	174
AUT	Correlación de	0.491*	0.649*	0.505*	0.539*	0.684*	0.466*	0.613*	1	0.616*	0.611*	0.585*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	0.000	0.000	.	0.000	0.000	0.000
	N	176	176	176	174	178	173	177	179	178	177	176
SAT	Correlación de	0.681*	0.603*	0.636*	0.616*	0.659*	0.579*	0.586*	0.616*	1	0.772*	0.721*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	.	0.000	0.000
	N	175	175	176	174	178	174	177	178	179	176	175
LID	Correlación de	0.698*	0.589*	0.640*	0.425*	0.637*	0.470*	0.532*	0.611*	0.772*	1	0.675*
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	.	0.000
	N	176	175	174	173	176	171	176	177	176	177	174
VIS	Correlación de	0.692*	0.644*	0.661*	0.625*	0.647*	0.672*	0.642*	0.585*	0.721*	0.675*	1
	Sig. (bilateral)	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	.
	N	173	174	173	171	175	171	174	176	175	174	176

*La correlación es significativa al nivel 0.01 (bilateral).