

Sesión 7 (C) — Los momentos no bastan: misma media y desviación típica, formas distintas

Marcos Bujosa

Descripción de la práctica

En la práctica A de esta misma sesión comprobamos que $\rho = \cos \theta$. Ahora damos un paso atrás y hacemos una pregunta más elemental, pero igual de importante: *¿basta con conocer la media y la desviación típica de una variable para saber cómo se distribuyen sus datos?* La respuesta, como veremos, es *no*: dos conjuntos de datos pueden compartir exactamente esos dos números y, sin embargo, tener distribuciones radicalmente distintas —simétrica, asimétrica, bimodal, con un valor atípico—. Esta práctica motiva por qué la estadística descriptiva necesita *más de dos herramientas*: además de media y desviación típica, también conviene analizar los histogramas, diagramas de caja, y estadísticos adicionales como la asimetría y la curtosis.

Narrativa: imaginamos cuatro empresas ficticias —TechNorte, Comercial Sur, Industrias Centro y Consultora Este— cuyo salario mensual (en euros) tiene, en las cuatro, la *misma* media y la *misma* desviación típica. ¿Tienen las cuatro plantillas “el mismo tipo” de estructura salarial?

Objetivo

1. Mostrar que estadísticos idénticos (o casi) no garantizan distribuciones parecidas.
2. Reforzar el uso de histograma y diagrama de caja como herramientas necesarias, no opcionales.
3. Introducir otros estadísticos descriptivos —asimetría, curtosis, mediana— que sí distinguen estas formas.

Nota técnica: en esta práctica los datos no son reales ni simulados al azar sin más: cada serie se genera con un tipo de distribución propia y luego se *reescala* para forzar una media y una desviación típica comunes. El procedimiento exacto se explica, con el detalle geométrico completo, en el Apéndice final. El tamaño muestral ($n = \$nobs$) es un parámetro (cuyo valor $\$nobs$ queda fijado inicialmente con `nulldata 40`): si al dibujar los histogramas la muestra pareciera escasa, basta con cambiar ese valor a 60 u 80 y volver a ejecutar el guión.

Actividad 1 - Generación de los datos

Generamos cuatro series de $n = 40$ observaciones con distribuciones muy distintas, pero forzando en las cuatro la misma media (μ_0) y la misma desviación típica (σ_0 , medida con la función `sd` de Gretl; véase la nota sobre el divisor en el Apéndice). El truco (reescalado de un vector en desviaciones) usa exactamente las propiedades de invarianza y homogeneidad demostradas en la lección 5; el detalle completo se explica en el Apéndice.

en línea de comandos:

```
nulldata 40
setobs 1 1 --cross-section
set seed 20250213

scalar n      = $nobs
```

Licencia: Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0).

```
scalar mu0 = 1800 # media común objetivo (euros)
scalar sigma0 = 250 # desviación típica común objetivo (euros)
```

Serie A: forma aproximadamente simétrica (TechNorte)

en línea de comandos:

```
series A_bruto = normal(0,1)
series A_z = (A_bruto - mean(A_bruto)) / sd(A_bruto)
series A = mu0 + sigma0 * A_z
setinfo A -d "Salario mensual, TechNorte (forma simétrica)" -n "TechNorte"
```

Serie B: asimetría positiva (Comercial Sur)

en línea de comandos:

```
series u1 = normal(0,1)
series u2 = normal(0,1)
series B_bruto = u1^2 + u2^2 # suma de cuadrados de normales: fuertemente asimétrica
series B_z = (B_bruto - mean(B_bruto)) / sd(B_bruto)
series B = mu0 + sigma0 * B_z
setinfo B -d "Salario mensual, Comercial Sur (asimetría positiva)" -n "Comercial Sur"
```

Serie C: bimodal (Industrias Centro)

en línea de comandos:

```
series grupo = (uniform(0,1) < 0.5) # dos "categorías" de empleados
series C_bruto = grupo*normal(-3,1) + (1-grupo)*normal(3,1)
series C_z = (C_bruto - mean(C_bruto)) / sd(C_bruto)
series C = mu0 + sigma0 * C_z
setinfo C -d "Salario mensual, Industrias Centro (bimodal: dos categorías)" -n "Industrias Centro"
```

Serie D: con un valor atípico (Consultora Este)

en línea de comandos:

```
series D_bruto = normal(0,0.3) # plantilla muy homogénea...
D_bruto[1] = 15 # ...salvo un directivo con un salario disparado
series D_z = (D_bruto - mean(D_bruto)) / sd(D_bruto)
series D = mu0 + sigma0 * D_z
setinfo D -d "Salario mensual, Consultora Este (con un valor atípico)" -n "Consultora Este"
```

Actividad 2 - Estadísticos descriptivos clásicos

en línea de comandos:

```
summary A B C D --simple
```

	Media	Mediana	D. T.	Mín	Máx
A	1800	1787	250,0	1222	2296
B	1800	1726	250,0	1610	3055
C	1800	1922	250,0	1405	2126
D	1800	1753	250,0	1713	3324

A la vista solo de estos números (media y desviación típica), ¿diría que las cuatro empresas tienen una estructura salarial parecida? Anote su respuesta antes de continuar.

Actividad 3 - Histogramas

- GUI: seleccione cada variable → clic derecho → Gráficos → Distribución de frecuencias.
o bien teclee en línea de comandos:

```
freq A
freq B
freq C
freq D
```

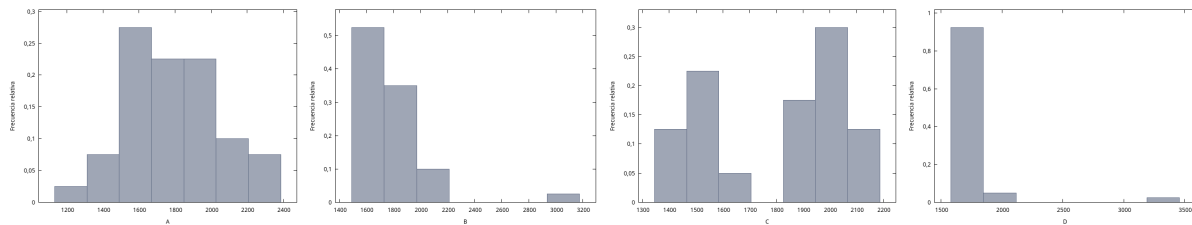


Figura 1: Histogramas.

¿Cambia su respuesta a la pregunta de la Actividad 2 al ver estos cuatro histogramas?

Actividad 4 - Diagramas de caja

en línea de comandos:

```
boxplot A B C D
```

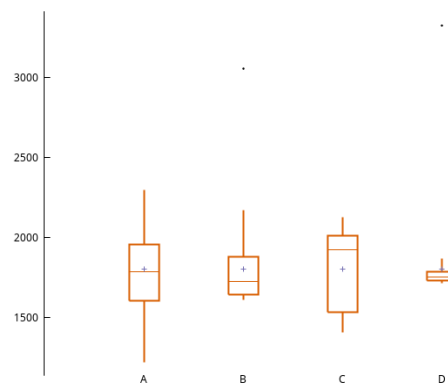


Figura 2: Diagramas de caja.

Compare las cuatro cajas: mediana, rango intercuartílico, longitud de los bigotes y presencia de puntos marcados como atípicos. ¿Qué revela el diagrama de caja que el histograma no mostraba con tanta claridad? ¿Y al revés?

Respuesta

El diagrama de caja pone en primer plano, de forma directamente comparable entre las cuatro empresas y sobre una única escala vertical, varios resúmenes que el histograma deja implícitos: la *mediana* (que en Consultora Este se separa claramente de la media, delatando el efecto del directivo con salario disparado, algo que un histograma con pocos intervalos podría disimular en una barra casi invisible); el *rango intercuartílico*, una medida de dispersión robusta que —a diferencia de la desviación típica, idéntica por construcción en las cuatro series— sí puede diferir entre empresas si los datos están distribuidos de forma distinta dentro de ese rango central; y, sobre todo, el *marcado explícito de valores atípicos* como puntos aislados fuera de los bigotes: en Consultora Este, un único punto extremo salta a la vista de inmediato, mientras que en un histograma ese mismo dato podría pasar desapercibido como una barra solitaria de altura mínima en el extremo del eje. El boxplot es, en definitiva, el instrumento idóneo para comparar *varias* distribuciones de un vistazo sobre los mismos ejes.

Sin embargo, el diagrama de caja tiene un punto ciego precisamente donde el histograma es insustituible: la *forma global* de la distribución. El boxplot de Industrias Centro (cuya distribución es bimodal, con dos grupos de empleados bien separados y prácticamente ningún salario intermedio) puede tener un aspecto perfectamente ordinario —una caja, una mediana, unos bigotes— sin insinuar en absoluto la existencia de dos picos separados; los cinco números que resume un boxplot (mínimo, Q_1 , mediana, Q_3 , máximo, más atípicos) simplemente no tienen capacidad para codificar cuántas modas tiene una distribución. Solo el histograma, al mostrar la densidad de puntos en cada tramo del rango, revela sin ambigüedad los dos grupos separados por un hueco central. La lección general es que ambas herramientas son complementarias: el boxplot compara resúmenes de posición y dispersión entre grupos; el histograma muestra la forma que esos resúmenes, por sí solos, no pueden capturar.

Actividad 5 - Más allá de la media y la desviación típica

Calculemos, para las cuatro series, algunos estadísticos adicionales: la *mediana* (`median`), la *asimetría* (`skewness`) y la *curtosis* (`kurtosis`, curtosis *excedente* respecto a la normal, que vale 0). Son estadísticos vistos en la asignatura de estadística del cuatrimestre anterior.

en línea de comandos:

```
printf "%-15s %10s %10s %10s %10s\n", "Empresa", "Media", "Mediana", "SD", "Asimetría", "Curtosis"
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "TechNorte", mean(A), median(A), sd(A), skewness(A), kurtosis(A)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Comercial Sur", mean(B), median(B), sd(B), skewness(B), kurtosis(B)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Ind. Centro", mean(C), median(C), sd(C), skewness(C), kurtosis(C)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Consultora Este", mean(D), median(D), sd(D), skewness(D), kurtosis(D)
```

Empresa	Media	Mediana	SD	Asimetría	Curtosis
TechNorte	1800,00	1786,56	250,00	0,1190	-0,2722
Comercial Sur	1800,00	1726,48	250,00	3,3545	14,0124
Ind. Centro	1800,00	1922,45	250,00	-0,3539	-1,5767
Consultora Este	1800,00	1753,34	250,00	5,8738	33,3403

Observe que la asimetría y la curtosis *sí* distinguen numéricamente lo que media y sd no distinguían. Pero incluso con esta tabla ampliada, ningún número delata con claridad la *bimodalidad* de Industrias Centro: solo el gráfico la muestra sin ambigüedad.

Conviene subrayar la idea general: *asimetría*, *curtosis* y *mediana* mejoran claramente la descripción, pero siguen siendo *resúmenes escalares* de toda la distribución. Cada uno comprime los datos en un único número y, por tanto, inevitablemente pierde información sobre la forma completa. Por eso pueden distinguir bastante bien una cola larga o un atípico, pero no sustituyen a un gráfico cuando lo que queremos detectar es una estructura más global —por ejemplo, dos grupos bien separados.

Actividad 6 - Síntesis

La estadística descriptiva es un *conjunto de herramientas complementarias* —medidas numéricas y gráficos— y no un único número: cada una responde a una pregunta distinta sobre los datos. Media y desviación típica resumen *posición* y *dispersión*; asimetría y curtosis dicen algo sobre la *forma*; pero solo el histograma y el diagrama de caja permiten “ver” patrones —como una bimodalidad— que ningún estadístico habitual resume bien en un solo número.

Cuando dispongamos de la recta de ajuste y del R^2 (lección 8), veremos que el mismo fenómeno puede ocurrir con *dos* variables a la vez y su recta de ajuste: es el célebre *cuarteto de Anscombe*, que reservamos para más adelante.

Apéndice - Cómo se han construido estos datos

Esta construcción no es un truco de programación ajeno al curso: es una aplicación directa de las propiedades demostradas en la lección 5 para la media y la desviación típica.

Dado cualquier vector $\mathbf{x}$ (con la forma que sea) con media $\mu_{\mathbf{x}}$ y desviación típica $\sigma_{\mathbf{x}} > 0$, definimos su versión *estandarizada*:

$$\mathbf{z} = \frac{\mathbf{x} - \mu_{\mathbf{x}}\mathbf{1}}{\sigma_{\mathbf{x}}}.$$

Por construcción, $\mathbf{z}$ tiene media 0 y desviación típica 1 (es el vector en desviaciones de $\mathbf{x}$, reescalado por su propia norma estadística). Esto es lo que hicimos en cada serie con **A_z**, **B_z**, **C_z**, **D_z**. Nótese que estandarizar *no cambia la forma* de la distribución: solo la traslada y la reescala; la asimetría o la bimodalidad de $\mathbf{x}$ se conservan intactas en $\mathbf{z}$.

Nota técnica sobre Gretl y el divisor. En esta práctica usamos **sd** tanto para estandarizar como para verificar el resultado final. Por eso, *sea cual sea el convenio interno exacto de Gretl para esa función*,¹ la operación

$$\mathbf{z} = \frac{\mathbf{x} - \mu_{\mathbf{x}}\mathbf{1}}{sd(\mathbf{x})}$$

produce una serie con desviación típica 1 *según ese mismo convenio*, y al multiplicar después por **sigma0** obtenemos una serie con desviación típica **sigma0** *según ese mismo convenio*. Es decir: el truco fuerza exactamente la desviación típica objetivo en el *mismo sentido operativo* en que Gretl la mide. Esto no contradice la observación de la práctica B: allí comparábamos *dos convenios distintos* para la varianza (divisor n frente a $n - 1$); aquí, en cambio, usamos el *mismo* convenio al normalizar y al reconstruir.

A continuación, reescalamos $\mathbf{z}$ a los valores objetivo μ_0 y σ_0 :

$$\mathbf{y} = \mu_0\mathbf{1} + \sigma_0\mathbf{z}.$$

Por la *linealidad* de la media (lección 4: sumar una constante desplaza la media exactamente esa constante) y, para la sd, por su *invarianza* ante traslaciones y *homogeneidad* ante escala (lección 5), la media de $\mathbf{y}$ es exactamente μ_0 ; por las mismas propiedades aplicadas a la desviación típica, la **sd** de $\mathbf{y}$ es exactamente σ_0 —con independencia de la forma de $\mathbf{x}$ de partida—. Esto es lo que hace la línea **mu0 + sigma0*A_z** (y análogas): fuerza los dos primeros momentos sin tocar la forma.

Verifiquemos numéricamente que las cuatro series comparten, en efecto, media y sd (salvo el redondeo de coma flotante):

en línea de comandos:

```
printf "\nVerificación: medias y sd deberían coincidir (salvo redondeo)\n"
printf "Media   -> A: %.6f  B: %.6f  C: %.6f  D: %.6f\n", mean(A), mean(B), mean(C), mean(D)
printf "SD      -> A: %.6f  B: %.6f  C: %.6f  D: %.6f\n", sd(A),  sd(B),  sd(C),  sd(D)
printf "(objetivo: mu0 = %.2f ,  sigma0 = %.2f)\n", mu0, sigma0
```

¹en concreto, **sd** divide por $n - 1$, como se comprueba en la práctica B de esta misma sesión; pero el argumento que sigue no depende de ello.

Verificación: medias y sd deberían coincidir (salvo redondeo)
 Media -> A: 1800,000000 B: 1800,000000 C: 1800,000000 D: 1800,000000
 SD -> A: 250,000000 B: 250,000000 C: 250,000000 D: 250,000000
 (objetivo: $\mu_0 = 1800,00$, $\sigma_0 = 250,00$)

Verifiquemos que la estandarización logra que las cuatro series tengan media cero y sd igual a uno (salvo el redondeo de coma flotante):

en línea de comandos:

```
printf "\nChequeo de la estandarización usada en Gretl\n"
printf "sd(A_z) = %.6f sd(B_z) = %.6f sd(C_z) = %.6f sd(D_z) = %.6f\n", sd(A_z), sd(B_z), sd(C_z), sd(D_z)
printf "mean(A_z) = %.6f mean(B_z) = %.6f mean(C_z) = %.6f mean(D_z) = %.6f\n", mean(A_z), mean(B_z), mean(C_z),
↪ mean(D_z)
```

```
Chequeo de la estandarización usada en Gretl
sd(A_z) = 1,000000 sd(B_z) = 1,000000 sd(C_z) = 1,000000 sd(D_z) = 1,000000
mean(A_z) = 0,000000 mean(B_z) = -0,000000 mean(C_z) = 0,000000 mean(D_z) = 0,000000
```

Lo único que *no* controla este procedimiento es la forma: la asimetría de $\mathbf{x}$ (o su bimodalidad, o su atípico) pasa intacta a $\mathbf{y}$, porque tanto la traslación como el reescalado por una constante positiva son transformaciones que *no alteran* la forma relativa de los datos —la misma idea, aplicada aquí de forma constructiva en lugar de solo verificatoria, que empleamos en el Apéndice de la práctica A para mostrar la invarianza de ρ ante el divisor—.

Preguntas de interpretación para la clase

1. ¿Qué estadístico (de los calculados en la Actividad 5) delata mejor la asimetría de Comercial Sur?
2. ¿Por qué ni la asimetría ni la curtosis bastan para detectar, con claridad, la bimodalidad de Industrias Centro?
3. Si un analista solo mirara la tabla de **summary** (media y sd) y nunca los gráficos, ¿qué conclusión errónea podría sacar sobre las cuatro empresas?
4. En Consultora Este, ¿qué estadístico —media o mediana— describe mejor “el salario típico” de un empleado cualquiera de esa empresa? ¿Por qué difieren tanto entre sí?
5. El procedimiento de generación (estandarizar y reescalar) fuerza la misma media y la misma sd en las cuatro series. ¿Por qué este mismo procedimiento *no puede* forzar, además, la misma asimetría o la misma forma?

Para profundizar:

- Tukey, J. W. (1977). *Exploratory Data Analysis*. Addison-Wesley. (Origen del diagrama de caja y del énfasis en la visualización de datos.)
- Anscombe, F. J. (1973). “Graphs in Statistical Analysis”. *The American Statistician*, 27(1), 17-21. (Mención breve: el mismo fenómeno con dos variables y su recta de ajuste, que veremos más adelante.)
- Cottrell, A. y Lucchetti, R. (2023). *Gretl User’s Guide*. Funciones **skewness**, **kurtosis**, **median**.
- Véase la lección 5 para las propiedades de invarianza y homogeneidad usadas en el Apéndice.

Código completo de la práctica

```
# -----
# Copyright (C) 2026 Marcos Bujosa
# Licencia: GNU General Public License v3.0 o posterior
# Este código es software libre y puede ser redistribuido y/o modificado bajo los términos de la GPL.
# Ver el archivo LICENSE del repositorio para más detalles.
# -----

# Los dos primeros comandos son necesarios para que Gretl guarde los resultados de la práctica en el
↳ directorio de trabajo
# al ejecutar lo siguiente desde un terminal (use los nombres y ruta que correspondan)
#
#   DIRECTORIO="Nombre_Directorio_trabajo" gretlcli -b ruta/nombre_fichero_de_la_practica.inp
#
# Si esto no le funciona en su sistema, comente las siguientes dos líneas y sitúese en el directorio de
↳ trabajo de gretl
# que corresponda (configure dicho directorio de trabajo desde la ventana principal de Gretl).

string directory = getenv("DIRECTORIO")
set workdir "@directory"

nulldata 40
setobs 1 1 --cross-section
set seed 20250213

scalar n      = $nobs
scalar mu0    = 1800 # media común objetivo (euros)
scalar sigma0 = 250  # desviación típica común objetivo (euros)

series A_bruto = normal(0,1)
series A_z     = (A_bruto - mean(A_bruto)) / sd(A_bruto)
series A       = mu0 + sigma0 * A_z
setinfo A -d "Salario mensual, TechNorte (forma simétrica)" -n "TechNorte"

series u1 = normal(0,1)
series u2 = normal(0,1)
series B_bruto = u1^2 + u2^2 # suma de cuadrados de normales: fuertemente asimétrica
series B_z     = (B_bruto - mean(B_bruto)) / sd(B_bruto)
series B       = mu0 + sigma0 * B_z
setinfo B -d "Salario mensual, Comercial Sur (asimetría positiva)" -n "Comercial Sur"

series grupo = (uniform(0,1) < 0.5) # dos "categorías" de empleados
series C_bruto = grupo*normal(-3,1) + (1-grupo)*normal(3,1)
series C_z     = (C_bruto - mean(C_bruto)) / sd(C_bruto)
series C       = mu0 + sigma0 * C_z
setinfo C -d "Salario mensual, Industrias Centro (bimodal: dos categorías)" -n "Industrias Centro"

series D_bruto = normal(0,0.3) # plantilla muy homogénea...
D_bruto[1] = 15 # ...salvo un directivo con un salario disparado
series D_z     = (D_bruto - mean(D_bruto)) / sd(D_bruto)
series D       = mu0 + sigma0 * D_z
setinfo D -d "Salario mensual, Consultora Este (con un valor atípico)" -n "Consultora Este"

outfile --quiet estadisticos.txt
summary A B C D --simple
end outfile

freq A --plot="histA.png"
freq B --plot="histB.png"
freq C --plot="histC.png"
freq D --plot="histD.png"

boxplot A B C D --output="boxplotABCD.png"

outfile --quiet tablaestadisticosadicionales.txt
```

```

printf "%-15s %10s %10s %10s %10s %10s\n", "Empresa", "Media", "Mediana", "SD", "Asimetria", "Curtosis"
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "TechNorte", mean(A), median(A), sd(A),
↳ skewness(A), kurtosis(A)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Comercial Sur", mean(B), median(B), sd(B),
↳ skewness(B), kurtosis(B)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Ind. Centro", mean(C), median(C), sd(C),
↳ skewness(C), kurtosis(C)
printf "%-15s %10.2f %10.2f %10.2f %10.4f %10.4f\n", "Consultora Este", mean(D), median(D), sd(D),
↳ skewness(D), kurtosis(D)
end outfile

outfile --quiet verificacionmomentos.txt
printf "\nVerificación: medias y sd deberían coincidir (salvo redondeo)\n"
printf "Media    -> A: %.6f  B: %.6f  C: %.6f  D: %.6f\n", mean(A), mean(B), mean(C), mean(D)
printf "SD      -> A: %.6f  B: %.6f  C: %.6f  D: %.6f\n", sd(A), sd(B), sd(C), sd(D)
printf "(objetivo: mu0 = %.2f , sigma0 = %.2f)\n", mu0, sigma0
end outfile

outfile --quiet estandarizacioninterna.txt
printf "\nChequeo de la estandarización usada en Gretl\n"
printf "sd(A_z) = %.6f  sd(B_z) = %.6f  sd(C_z) = %.6f  sd(D_z) = %.6f\n", sd(A_z), sd(B_z),
↳ sd(C_z), sd(D_z)
printf "mean(A_z) = %.6f  mean(B_z) = %.6f  mean(C_z) = %.6f  mean(D_z) = %.6f\n", mean(A_z),
↳ mean(B_z), mean(C_z), mean(D_z)
end outfile

```

Enlace al guión: [S07-Prct-C-momentos-no-bastan.inp](#)

Respuestas

1. ¿Qué estadístico delata mejor la asimetría de Comercial Sur? La *asimetría* (**skewness**), por definición: es precisamente el estadístico diseñado para capturar numéricamente el desequilibrio entre la cola izquierda y la derecha de una distribución. Un valor claramente positivo y alejado de 0 confirma la cola larga hacia la derecha que se ve en el histograma. La curtosis, en cambio, mide el “apuntamiento” o el peso de las colas en general, no la dirección del desequilibrio, así que no es el indicador más directo aquí.
2. ¿Por qué asimetría y curtosis no bastan para detectar la bimodalidad de Industrias Centro? Porque ambos son, igual que la media y la varianza, *momentos* (medidas que resumen toda la distribución en un único número mediante una suma). Un resumen numérico de este tipo puede ser prácticamente idéntico para una distribución con dos picos simétricos y separados y para una distribución con un único pico centrado: si los dos grupos están colocados simétricamente en torno a la media global, la asimetría puede salir cercana a 0 (parece “simétrica”), y la curtosis simplemente no fue diseñada para distinguir “un pico ancho” de “dos picos estrechos”. La bimodalidad es una característica de la *forma completa* de la distribución, no capturable en un solo escalar: por eso hace falta el histograma, que muestra la densidad en cada zona del rango, no solo un resumen agregado.
3. Si un analista solo mirara **summary** sin gráficos, ¿qué conclusión errónea podría sacar? Podría concluir que las cuatro empresas tienen “el mismo tipo” de plantilla salarial —una estructura homogénea con un salario típico de 1800€ y una dispersión de 250€ alrededor de esa cifra—, cuando en realidad una tiene una distribución razonablemente simétrica, otra está sesgada hacia salarios altos poco frecuentes, otra tiene en realidad *dos grupos* de empleados con salarios muy distintos entre sí (y ningún empleado con salario “intermedio”), y la última tiene una plantilla muy homogénea salvo un único directivo con un salario disparado. Decisiones de política salarial basadas solo en la media y la sd (p. ej., “subir un 5% a todos”) tendrían sentido muy distinto según el caso.
4. En Consultora Este, ¿media o mediana describe mejor “el salario típico”? ¿Por qué difieren? La *mediana* describe mucho mejor el salario típico: al ser un estadístico de *orden* (el valor central cuando se ordenan los datos), apenas se ve afectada por un único valor extremo. La media, en cambio, es una *suma* dividida por n , y el salario disparado del directivo “tira” de ella hacia arriba, de modo que la media puede terminar por encima de lo que gana la inmensa mayoría de la plantilla. La diferencia entre ambas (media > mediana, aquí de forma marcada) es en sí misma un indicio —sin necesidad de mirar el histograma— de que hay al menos un valor atípico influyente hacia arriba.
5. ¿Por qué el procedimiento de generación no puede forzar también la misma asimetría o forma? Porque el procedimiento (estandarizar y reescalar: $\mathbf{y} = \mu_0 \mathbf{1} + \sigma_0 \frac{\mathbf{x} - \mu_{\mathbf{x}} \mathbf{1}}{\sigma_{\mathbf{x}}}$) solo aplica dos tipos de transformación —una traslación (sumar una constante) y un cambio de escala (multiplicar por una constante positiva)—. Estas dos operaciones, por construcción, no alteran la forma *relativa* de los datos: si una distribución tiene una cola larga a la derecha, trasladarla o estirla/encogerla no la vuelve simétrica ni le añade un segundo pico. Para cambiar la forma (asimetría, número de modas) haría falta una transformación distinta —no lineal, o que combinara varias fuentes de variación de manera distinta—, precisamente lo que se hizo *antes* de estandarizar, al elegir cada distribución “en bruto” (normal, suma de cuadrados de normales, mezcla de dos normales, etc.).