

Modelo matemático en primera aproximación para la relación entre los estadísticos de prueba de normalidad y el valor de p

A first-approximation mathematical model for the relationship between normality test statistics and the p-value

José Mendoza Rodríguez¹, Andrés Salazar Andrade²

¹Departamento de Matemática, Unidad Educativa Joaquín Lalama, Ambato, Ecuador, 180150;

²Departamento de Ciencias, Unidad Educativa Balandra-Cruz del Sur, Guayaquil, Ecuador, 090150;

andreuspd0909@hotmail.com

*Correspondencia: mendo-10@hotmail.com

Citación: Mendoza, J. & Salazar, A., (2026). Modelo matemático en primera aproximación para la relación entre los estadísticos de prueba de normalidad y el valor de p. *Novasinerгия*. 9(1). 138-159.

<https://doi.org/10.37135/ns.01.17.08>

Recibido: 20 junio 2025

Aceptado: 12 noviembre 2025

Publicado: 08 enero 2026

Novasinerгия

ISSN: 2631-2654

Resumen: Las pruebas de normalidad son consideradas como el punto de partida para la estadística paramétrica; sin embargo, es importante comprender el comportamiento de los estadísticos y el valor de p para tomar decisiones significativas. El objetivo de la investigación fue diseñar un modelo matemático que permita identificar patrones o tendencias existentes entre el estadístico y el valor de p. La metodología utilizada consistió en aplicar las pruebas de normalidad univariadas más comunes a variables como la precipitación, crudo y temperatura, mediante un análisis de submuestras según las características de cada prueba. Cada submuestra proporcionó el valor del estadístico y el valor de p, los cuales fueron procesados mediante modelos de regresión clásicos como primera aproximación: lineal, polinomial de grado 2 y 3, exponencial y logarítmica. La validación de los modelos se realizó mediante el coeficiente de determinación (R^2), con valores de R^2 más altos: en la prueba KS, para la variable crudo ($R^2 = 0,9784$); en la prueba AD, para la variable precipitación ($R^2 = 1$); en la prueba JB, para la variable temperatura ($R^2 = 0,8902$); en la prueba CvM, para la variable temperatura ($R^2 = 1$); en la prueba M.KS, para la variable temperatura ($R^2 = 0,936$); en la prueba P χ^2 , para la variable precipitación ($R^2 = 0,3798$); en la prueba SF, para la variable temperatura ($R^2 = 0,6019$); y en la prueba SW, para la variable crudo ($R^2 = 0,3712$). Realizar este tipo de análisis preliminar permite a los investigadores tomar decisiones más efectivas y óptimas al momento de comprobar el supuesto de la normalidad univariadas en un conjunto de datos.

Palabras clave: Estadístico de prueba, Modelo matemático, Normalidad, Regresión, valor de p.

Abstract: Normality tests are considered the starting point of parametric statistics; however, it is essential to understand the behavior of test statistics and p-values to make meaningful decisions. The objective of this study was to design a mathematical model capable of identifying existing patterns or trends between the test statistic and the p-value. The methodology consisted of applying the most common univariate normality tests to variables such as precipitation, crude oil, and temperature, using a subsample analysis tailored to each test. Each subsample provided both the test statistic and the p-value, which were processed using classical regression models as a first approximation: linear, second-, and third-degree polynomials, exponential, and logarithmic. Model validation was performed using the coefficient of determination (R^2), obtaining the highest R^2 values as follows: in the KS test, for the crude oil variable ($R^2 = 0.9784$); in the AD test, for the precipitation variable ($R^2 = 1$); in the JB test, for the temperature variable ($R^2 = 0.8902$); in the CvM test, for the temperature variable ($R^2 = 1$); in the M.KS test, for the temperature variable ($R^2 = 0.936$); in the P χ^2 test, for the precipitation variable ($R^2 = 0.3798$); in the SF test, for the temperature variable ($R^2 = 0.6019$); and in the SW test, for the crude oil variable ($R^2 = 0.3712$). Conducting this type of preliminary analysis enables researchers to make more effective and optimal decisions when assessing the assumption of univariate normality in a dataset.

Keywords: Test statistics, Mathematical model, Normality, Regression, p-value.



Copyright: 2026 derechos otorgados por los autores a Novasinerгия.

Este es un artículo de acceso abierto distribuido bajo los términos y condiciones de una licencia de Creative Commons Attribution (CC BY NC). (<http://creativecommons.org/licenses/by/4.0/>).

1. Introducción

Los datos pueden generar resultados incorrectos si no se realiza un análisis previo de los mismos. En muchos casos, es importante aplicar un filtrado y limpieza de los datos con la finalidad de que cumplan un formato específico y acorde a lo que se desea analizar, corrigiendo o eliminando errores e inconsistencias en la base de datos, el objetivo es obtener datos de calidad [1]. Es importante indicar que, si las técnicas de recolección de información son eficientes, el tiempo de la preparación de los datos disminuye considerablemente y resulta tomar decisiones más creíbles. En cambio, el uso de técnicas deficientes provocaría incertidumbre y resultados ambiguos [2].

La elección de un gráfico estadístico depende en gran medida de las variables de estudio, además, dicha representación debe adaptarse a las condiciones específicas y a la naturaleza de cada una. Es importante comprender el comportamiento del tipo de variable para que los datos a recolectar tengan mayor validez y puedan encajar en un gráfico cuya interpretación sea significativa, enriquecedora y de fácil comprensión [3]. La visualización se convierte en una herramienta indispensable para representar un conjunto de datos (data set), a través de los gráficos estadísticos se pueden generar patrones, relaciones entre variables, tendencias y posibles anomalías. La adaptabilidad es una de las características importantes de los gráficos, se ajustan a las necesidades de las variables, sus datos y el tipo de análisis a realizar, ya sea univariado, bivariado o multivariado [4].

La distribución normal también definida como una distribución gaussiana, tiene la finalidad de verificar que los datos se ajusten a una línea de tendencia basado en la media poblacional (μ) o muestral ($\bar{x}$) y a la desviación estándar poblacional (σ) o muestral (s) como medida de tendencia central y dispersión respectivamente. Tiene como característica principal la simetría de los datos donde existe una concentración en forma de campana de Gauss, lo que indica que la mayoría de los valores se agrupan cerca del centro de la distribución, reduciendo gradualmente hacia los extremos [5]. Si se refiere a una muestra, el criterio de simetría se convierte en una condición ideal, conociendo que los datos son recolectados con un error muy probable desde cómo se define la variable de estudio hasta su respectiva medición. En otras palabras, la existencia de una asimetría (skewness) en los datos tiene referencia a los sesgos o ligeros desvíos que se producen por los valores atípicos (outliers) o por la propia naturaleza de los datos [6].

La importancia de la distribución normal interfiere en la capacidad de modelar fenómenos en distintas áreas del conocimiento que son de base para la aplicación de técnicas de inferencia estadísticas que permiten predecir o estimar sobre los resultados basados en la normalidad. El tamaño de la muestra juega un papel crucial en la disminución del error muestral, mientras mayor sea el tamaño de la muestra, mayor precisión de las estimaciones y mejor representatividad con respecto a la población [7]. En muchos de los casos, los grandes volúmenes de datos al ser procesados representan un costo computacional significativo, por eso la importancia de trabajar con una muestra adecuada. En este contexto, seleccionar un muestreo probabilístico es la mejor alternativa para que los datos sean considerados en base a un criterio estadístico según las características propias de la población [8].

Las pruebas de normalidad están sujetas en comprobar si la hipótesis nula cumple con la normalidad de los datos en base a una muestra predefinida. Cada prueba de normalidad proporciona o genera un valor de p , que representa la probabilidad de obtener un resultado tan extremo basado en la suposición de que la hipótesis nula es verdadera [9]. El valor de p toma valores distintos cuando los tamaños muestrales crecen o decrecen, es decir, que están ligados por el tamaño de la muestra inicial. En muchos de los casos, la interpretación del valor de p debe ser un punto de inflexión para seguir profundizando en el análisis, el valor de p no debe ser la única vía para tomar decisiones en una investigación [10]. El valor de p debe ser respaldado por otros parámetros estadísticos que evidencien el comportamiento de una variable. Forzar en rechazar la hipótesis nula basado en un valor de p , puede ocasionar un error en seleccionar una prueba estadística para realizar inferencia sobre ciertas variables definidas en un estudio.

Al diseñar un modelo matemático en una primera aproximación, se pueden emplear modelos o expresiones matemáticas comunes o de mayor utilidad, como los modelos lineales, cuadráticas (polinomio de grado 2), cúbicas (polinomio de grado 3), exponenciales, logarítmicas, entre otros. Este tipo de modelos se utilizan para realizar ajustes de curva a los datos que provienen de una muestra. Además, son considerados como modelos iniciales o introductorios que orientan al investigador en la exploración de patrones o comportamientos de las variables.

La validación de un modelo es fundamental debido a que cada expresión matemática tiene un comportamiento distinto que depende de los coeficientes definidos en el modelo. El uso del coeficiente de determinación y la suma de los cuadrados de los errores son estadísticos que permiten estimar la calidad del ajuste del modelo a los datos.

2. Metodología

El tipo de investigación se direccionó al modelado matemático, ya que el objetivo principal fue diseñar una función o expresión matemática que relacione los estadísticos de prueba de normalidad univariada con sus respectivos valores de p [11]. Además, se consideró de tipo explicativo, dado que cada conjunto de datos seleccionados según las pruebas de normalidad, genera ciertos valores de p en función al test aplicado.

El estudio en lo referente al nivel de investigación fue exploratorio debido a que la intención era identificar patrones o relaciones en los datos como una primera aproximación [12]. Además, la finalidad de diseñar el modelo era encontrar la función o curva más adecuada que se ajuste a su comportamiento, según los criterios de cada prueba de normalidad univariada, ya sean estos direccionados en medir distancias o momentos, tomando como referencia una distribución normal. Por tal razón, la exploración de los datos fue indispensable para lograr un mejor ajuste y llevar el proceso de validación correspondiente.

La técnica utilizada fue la revisión documental mediante un análisis a profundidad de las instituciones gubernamentales del Ecuador que recopilan y centralizan información estadística. En particular, los datos se concentran en un Catálogo de Datos Abiertos donde se permitió acceder a la información sin necesidad de seguir un proceso de solicitud de acceso a la información pública. La finalidad de este catálogo es proporcionar criterios

técnicos y metodológicos a los usuarios para que puedan promover la investigación, el emprendimiento y la innovación de la sociedad. De esta forma, se obtuvo los datos para el desarrollo del estudio.

Para el diseño del modelo matemático en su primera aproximación se utilizaron diferentes ajustes de curva según el comportamiento de los datos. Para definir el mejor ajuste se tomó como referencia los valores del coeficiente de determinación (R^2) y la suma de los cuadrados de los errores SCE. R^2 oscila entre 0 a 1, su interpretación indica que, mientras más cercano sea el valor de R^2 a 1, existe un mejor ajuste del modelo a la distribución de los datos. En cambio, SCE indica que, si su valor es bajo, mejor es el ajuste o se considera más preciso el modelo a un conjunto de datos.



Figura 1. Esquema – relación R^2 y SCE

Tabla 1. Umbrales referenciales de R^2

R^2	Interpretación
$R^2 < 0,25$	Relación muy débil.
$0,25 \leq R^2 < 0,50$	Relación débil.
$0,50 \leq R^2 < 0,75$	Relación moderada.
$R^2 \geq 0,75$	Relación sustancial.

Nota. Adaptado de [13]

2.1. Población y muestra

Los datos considerados para el análisis estadístico fueron: la temperatura mínima absoluta, la producción de petróleo mensual operativo y la precipitación, los mismos que fueron descargados de la plataforma de datos abiertos del Ecuador (<https://www.datosabiertos.gob.ec/>) y proporcionados por el Instituto Nacional de Meteorología e Hidrología – INAMHI y la Empresa Pública de Hidrocarburos del Ecuador - EP PETROECUADOR.

En lo que concierne a la muestra, se empleó un muestreo no probabilístico por conveniencia, dado que los datos provienen de fuentes oficiales y confiables. Esta elección permite realizar un análisis estadístico válido y transparente, considerando que los datos son reales y cuya disponibilidad es de acceso libre. Además, utilizar un muestreo no probabilístico responde al propósito de caracterizar fenómenos observables a partir de datos recolectados, sin intervención ni manipulación de las variables.

Para la temperatura mínima absoluta, se consideró como muestra el primer semestre del año 2019 de las diferentes estaciones meteorológicas ubicadas en puntos estratégicos del territorio, el objetivo fue obtener mediciones distintas y representativas. En lo que corresponde a la producción de petróleo mensual operativo, se emplearon datos comprendidos entre los años 2022 y 2025. Para la variable precipitación, se analizaron datos registrados en los años 2018 y 2019. Las muestras específicas de cada variable se detallan en la tabla 2.

Tabla 2. Descripción de variables y muestra

Variable	Descripción	Unidad de medida	Muestra
Temperatura mínima absoluta	Datos sobre la temperatura máxima absoluta de las principales estaciones meteorológicas del INAMHI a nivel nacional.	Grado Celsius (°C)	418
Producción de petróleo mensual operativo	Datos sobre los volúmenes de barriles Promedio por mes (BPPM) por cada activo.	Barriles	535
Precipitación	Datos sobre la precipitación total mensual de las principales estaciones meteorológicas del INAMHI a nivel nacional.	Milímetros (mm)	669

Nota. Datos obtenidos de Datos Abiertos Ecuador: Producción de petróleo mensual operativo (<https://www.datosabiertos.gob.ec/dataset/produccion-mensual-petroecuador>), temperatura mínima absoluta (<https://www.datosabiertos.gob.ec/dataset/temperatura-minima-absoluta>) y precipitación (<https://www.datosabiertos.gob.ec/dataset/precipitacion-total-mensual>).

2.2. Tratamiento Estadístico

Sobre la metodología aplicada, se estructuró en base al enfoque, diseño y nivel de investigación. Se estableció un enfoque cuantitativo por la naturaleza que provienen los datos. Segundo, se utilizó un diseño no experimental debido a que no se manipuló la variable de estudio [14], mencionar que el objetivo fue analizar el comportamiento de la variable enfocado en los valores de las pruebas estadísticas preestablecidas con respecto a los valores de p . Finalmente, el nivel de investigación utilizado fue descriptivo debido a que el conjunto de datos analizados (muestra total n) se estableció en subconjuntos de datos (submuestras A_k) donde k es el valor incremental para cada iteración en base a un valor muestral inicial, n es el tamaño de la muestra para las pruebas Kolmogorov Smirnov (KS), Shapiro Wilk (SW), Anderson Darling (AD), Jarque – Bera (JB), Cramér-von Mises (CvM), Lilliefors (M.KS), Pearson Chi-Cuadrado ($P \chi^2$) y Shapiro-Francia (SF).

Para las pruebas KS, AD, JB, CvM, M.KS, $P \chi^2$ y SF, se utilizó $k > 50$ observaciones con incrementos de una observación adicional para cada submuestra y para la prueba SW $k \in \{50, 49, 48, \dots, 3\}$ con un proceso inverso donde existe decrementos en una observación para cada submuestra. Se fijó $k = 50$ para la primera submuestra según las características

propias de la prueba SW y las adaptaciones específicas en cuanto al tamaño muestral. Para las pruebas KS, SW, AD, JB, CvM, M.KS, $P\chi^2$ y SF se definieron las muestras $n_1 = 418$; $n_2 = 535$; $n_3 = 669$, según la variable de análisis.

Tabla 3. Valores mínimo y máximo para las iteraciones de cada submuestra

Prueba	Valor muestral inicial	Submuestra	Mínimo y máximo
KS AD JB CvM M.KS	$k = 51$	Incrementos de una observación hasta alcanzar el valor de n .	$\min(k) = 51, \max(k) = n$
$P\chi^2$ SF SW	$k = 50$	Decrementos de una observación hasta alcanzar el valor mínimo de 3.	$\min(k) = 3, \max(k) = 50$

Nota. Los valores mínimos y máximos para cada submuestra se establecieron siguiendo criterios de sensibilidad al tamaño muestral reportados en estudios comparativos recientes sobre pruebas de normalidad univariada. Los estadísticos de prueba y sus valores de p varían con el número de observaciones, lo que justifica los incrementos y decrementos aplicados en cada caso [15], [16] y [17].

Para aplicar las pruebas de normalidad univariada se adoptó un nivel de significancia de $\alpha = 0,05$, valor definido para conocer si las variables siguen una distribución normal y que es aceptado en la literatura estadística para determinar si existe o no diferencias significativas [18].

En las pruebas KS, AD, JB, CvM, M.KS, $P\chi^2$ y SF con n designada $\{x_1, x_2, x_3, \dots, x_n\}$, se definió las submuestras $A_k = \{x_1, x_2, x_3, \dots, x_k\}$, con $k \in \{51, 52, 53, \dots, n\}$. En el caso de la prueba SW se establecieron bloques sistemáticos (B_m) como se indica en la ecuación 2 con $k \in \{50, 49, 48, \dots, 3\}$. La finalidad de utilizar bloques fue cubrir la mayor cantidad de datos hasta acercarse al valor de n . La forma ampliada para las pruebas las pruebas KS, AD, JB, CvM, M.KS, $P\chi^2$ y SF se detallan a continuación:

$$M_k = \begin{pmatrix} x_1 & x_2 & x_3 & \cdots & x_{51} \\ x_1 & x_2 & x_3 & \cdots & x_{51} & x_{52} \\ x_1 & x_2 & x_3 & \cdots & x_{51} & x_{52} & x_{53} \\ \vdots & \vdots & \vdots & & \vdots & \vdots & \vdots & \ddots \\ x_1 & x_2 & x_3 & \cdots & x_{51} & x_{52} & x_{53} & \cdots & x_n \end{pmatrix} \quad (1)$$

Donde M_k es una matriz de orden $n-50$ fila y n columnas, las filas representa la cantidad de submuestras y las columnas indica el número de observaciones.

En la prueba SW, en lo referente a los bloques, cada bloque empieza con un total de 50 observaciones, luego se inicia el proceso de iteración hasta alcanzar el valor mínimo establecido.

$$B_m = \{x_{(m-1) \cdot 50+1}, x_{(m-1) \cdot 50+2}, \dots, x_{m/n(m \cdot 50, n)}\} \quad (2)$$

Donde $m = 1, 2, 3, \dots, \left\lceil \frac{n}{50} \right\rceil$ y $\min(m, 50, n)$ garantiza que no se exceda el tamaño muestral de cada variable de estudio.

Para cada bloque la ecuación 3 define su conjunto de observaciones que depende del número de bloque y sus iteraciones de manera decremental. La forma ampliada se representa en la ecuación 4.

$$A_{m,k} = \{x_{(m-1) \cdot 50+1}, x_{(m-1) \cdot 50+2}, x_{(m-1) \cdot 50+3}, \dots, x_{(m-1) \cdot 50+k}\} \quad (3)$$

$$M_{m,k} = \begin{pmatrix} x_1 & x_2 & x_3 & \cdots & x_{50} \\ x_1 & x_2 & x_3 & \cdots & x_{49} \\ x_1 & x_2 & x_3 & \cdots & x_{48} \\ \vdots & \vdots & \vdots & & \vdots \\ x_1 & x_2 & x_3 & \cdots & x_3 \end{pmatrix} \quad (4)$$

Donde cada fila de la matriz $M_{m,k}$ representa cada submuestra en base al bloque definido y las columnas indica el número de observaciones.

Cada una de las pruebas según la cantidad de submuestras y bloques analizados anteriormente, generan los estadísticos de prueba y los valores de p respectivos en cada iteración, por lo tanto, cada valor obtenido se representó en forma de conjunto de pares ordenados (λ). Cada submuestra fue evaluada de forma independiente, este procedimiento permitió analizar la variabilidad de los resultados según el tamaño muestral, identificando posibles fluctuaciones en el comportamiento del estadístico y el valor de p.

$$\lambda_{KS} = \{(D_k, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (5)$$

$$\lambda_{AD} = \{(A_k^2, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (6)$$

$$\lambda_{JB} = \{(JB_k, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (7)$$

$$\lambda_{CvM} = \{(W_k^2, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (8)$$

$$\lambda_{M.KS} = \{(D'_k, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (9)$$

$$\lambda_{P_{\chi^2}} = \{(P_k, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (10)$$

$$\lambda_{SF} = \{(W'_k, p_k) \mid k \in \{51, 52, \dots, n\}\} \quad (11)$$

$$\lambda_{SW} = \{(W_k, p_k) \mid k \in \{50, 49, \dots, 3\}\} \quad (12)$$

La información obtenida en la exploración de los datos, tanto en la visualización de gráficos estadístico como de los resultados derivados de la inferencia estadísticas mediante la aplicación de diferentes pruebas de normalidad univariada, fue procesada con el software RStudio (versión 2024.12.1) utilizando los paquetes y librerías: *readxl*, *ggplot2*, *nortest*, *tseries*, *GGally*, *dplyr*, *tidyr* y *systemfonts* [19].

3. Resultados

Las pruebas de normalidad con sus respectivos estadísticos están orientadas a medir distancias o momentos estadísticos, dependiendo del enfoque con el que se analice el conjunto de datos. Es importante explorar cada estadístico de forma individual para identificar patrones o tendencia en su comportamiento según las variables de estudio. Sin embargo, es adecuado, realizar una exploración por pares con el objetivo de evaluar la correlación y afinidad entre estadísticos.

Si se habla de correlación se debe tener en cuenta que su estadístico está acotado entre $[-1,1]$ donde -1 indica una correlación perfecta negativa y 1 una correlación perfecta positiva. Para la variable precipitación (ver Figura 2) mediante una matriz de correlación entre estadísticos, se obtuvo que las pruebas KS, AD, CvM, M.KS, $P \chi^2$ y SF tienen una alta correlación al compararse entre sí. El estadístico de la prueba JB tiene una correlación moderada en comparación con los demás estadísticos, esto sucede, debido a que la prueba JB está direccionado en medir el tercer y cuarto momentos estadístico (asimetría y curtosis) y las otras pruebas se centran en medir distancias entre la distribución empírica y la teórica. Por otro lado, en las pruebas KS y M.KS se obtuvo una correlación perfecta, esto se debe a que la prueba M.KS es una modificación o adaptación de la prueba KS, la diferencia está en el cálculo del valor de p.

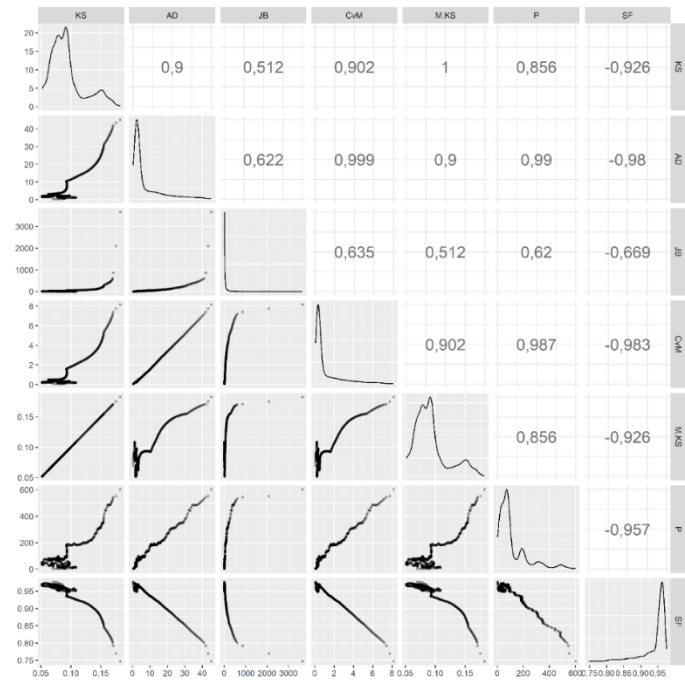


Figura 2. Matriz de correlación para la variable precipitación

Para la variable producción de petróleo mensual operativo (ver Figura 3) el comportamiento de los estadísticos difiere considerablemente. La prueba KS tiene una alta correlación con las pruebas M.KS (correlación perfecta, como se explicó previamente) y SF, sin embargo, la correlación es baja en los demás estadísticos. La prueba AD mantiene una alta correlación con las pruebas CvM y $P \chi^2$, mientras que con el resto la correlación es baja y moderada. La prueba JB presenta una correlación baja y moderada con respecto a los otros estadísticos

(evidente porque mide momentos). Por último, se observa una fuerte correlación entre las pruebas CvM y $P \chi^2$.

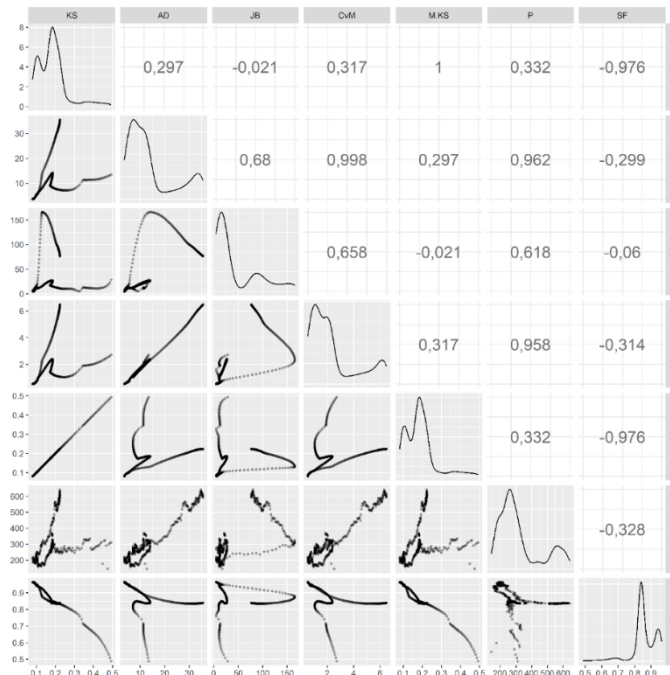


Figura 3. Matriz de correlación para la variable crudo

La variable temperatura mínima absoluta (ver Figura 4) la correlación entre estadísticos varía de moderada a alta, tanto positivas como negativas. Esto sugiere una consistencia o mayor concordancia en las pruebas de normalidad, evidenciando un mejor comportamiento general de los estadísticos al evaluar la distribución de esta variable.

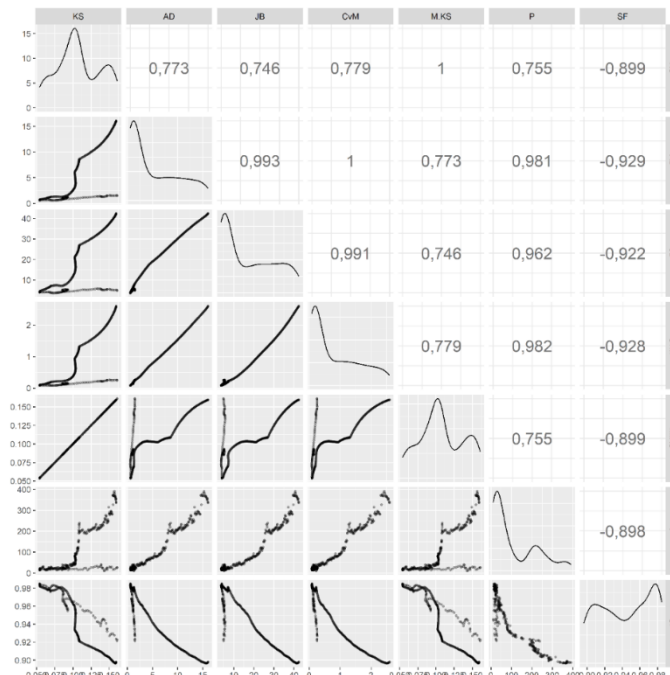


Figura 4. Matriz de correlación para la variable temperatura

Es importante señalar que, en las figuras 2, 3 y 4, la matriz de correlación presenta, en la diagonal principal, las funciones de densidad de cada estadístico según la variable

analizada. Debajo de la diagonal principal se encuentran los gráficos de dispersión entre pares, mientras que en la parte superior se muestran los coeficientes de correlación de Pearson correspondientes a cada par, indicando el grado de asociación entre estadísticos. Esta representación gráfica permite obtener una visión más clara sobre la relación de los estadísticos de las diferentes pruebas de normalidad univariada y así facilitar la interpretación de los resultados, identificación de patrones que se ajustan a un modelo determinado, corroboración de similitudes y diferencias entre pruebas que han sido modificadas o reajustadas y observación del comportamiento de los datos cuando se utilizan pruebas en medidas de distancia, en comparación con aquellas centradas en momentos estadísticos.

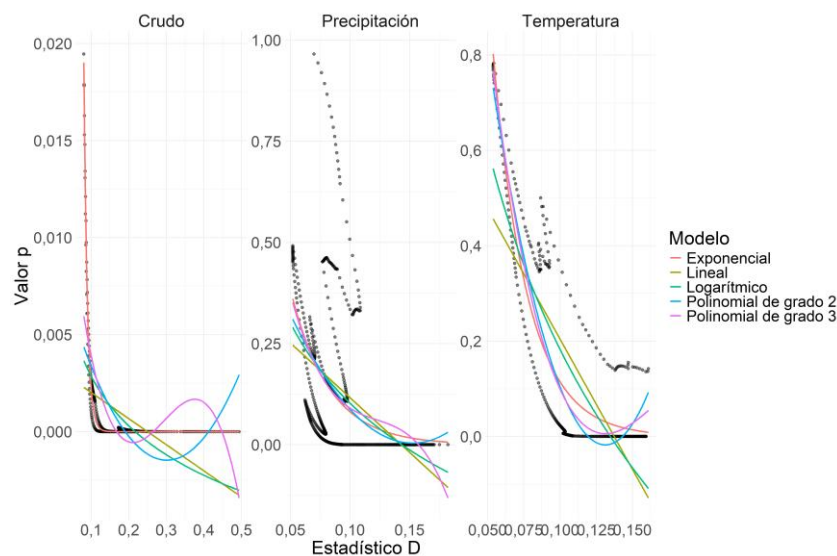


Figura 5. Ajuste de curva – valor de p y el estadístico D

En lo que corresponde a la prueba KS, las variables crudo y temperatura presentaron un mejor ajuste de curva en su primera aproximación mediante modelos exponencial (ver ecuación 13) y polinomial de grado 3 (ver ecuación 14) respectivamente. En el caso al modelo exponencial, se produce un decrecimiento debido al exponente negativo, esto significa que el valor de p disminuye exponencialmente a medida que el estadístico D aumenta, y viceversa. Por otro lado, en el modelo polinomial de grado 3, la interpretación de los coeficientes influye a medida que el estadístico D varía, el valor de p puede disminuir o incrementar dependiendo del equilibrio de los términos lineales, cuadrático y cúbico. Los modelos exponencial y polinomial de grado 3 están acotados con $D \in [0,0805;0,4935]$ y $D \in [0,0539;0,1610]$ respectivamente.

$$\hat{p} = 135,0617e^{-110,0911D} \quad (13)$$

$$\hat{p} = \left| 0,1621 - 3,1162D + 2,1967D^2 - 0,2921D^3 \right| \quad (14)$$

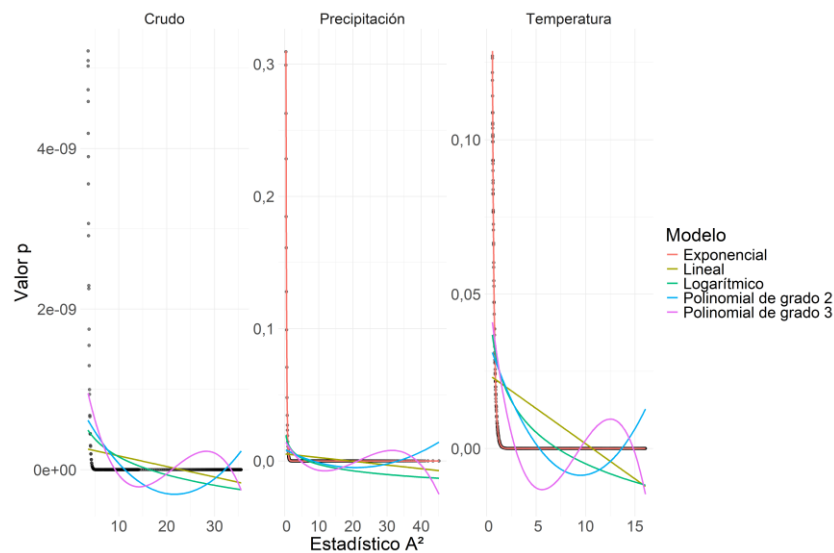


Figura 6. Ajuste de curva – valor de p y el estadístico A²

En lo que respecta a la prueba AD, las variables precipitación y temperatura presentaron un mejor ajuste de curva en su primera aproximación mediante modelos exponenciales (ver ecuaciones 15 y 16) en ambos casos. En los modelos exponenciales, se produce un decrecimiento debido al exponente negativo, esto significa que el valor de p disminuye exponencialmente a medida que el estadístico A^2 aumenta, y viceversa. Los modelos exponenciales están acotados con $A^2 \in [0, 4227; 45, 2029]$ y $A^2 \in [0, 5825; 16, 0537]$ respectivamente. Es importante indicar que, la notación matemática del estadístico es A^2 , el superíndice no indica que al estadístico se debe elevar al cuadrado.

$$\hat{p} = 3,4294e^{-5,6865(A^2)} \quad (15)$$

$$\hat{p} = 3,5297e^{-5,6845(A^2)} \quad (16)$$

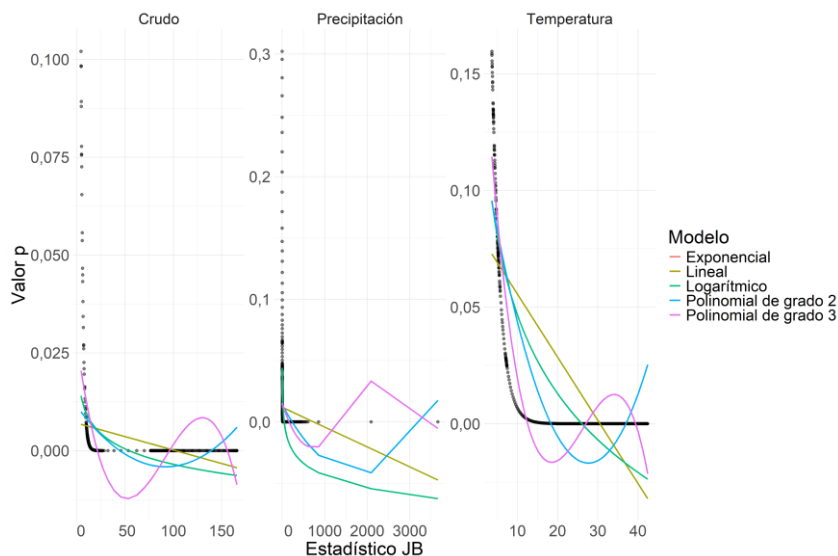


Figura 7. Ajuste de curva – valor de p y el estadístico JB

Con respecto a la prueba JB, las variables precipitación, producción de petróleo mensual operativo y temperatura mínima absoluta presentaron un mejor ajuste de curva en su

primera aproximación mediante modelos logarítmico (ver ecuación 17) y polinomial de grado 3 (ver ecuaciones 18 y 19) respectivamente. En el caso al modelo logarítmico, el coeficiente negativo que acompaña a la expresión logarítmica indica que, el valor de p disminuye a medida que el estadístico JB aumenta, y viceversa. Por otro lado, en el modelo polinomial de grado 3, la interpretación de los coeficientes influye a medida que el estadístico JB varía, el valor de p puede disminuir o incrementar dependiendo del equilibrio de los términos lineales, cuadrático y cúbico. Los modelos logarítmico y polinomial de grado 3 están acotados con $JB \in [2,394; 3661,102]$, $JB \in [4,565; 165,741]$ y $JB \in [3,670; 42,403]$ respectivamente.

$$\hat{p} = 0,0567 - 0,0145 \log(JB) \quad (17)$$

$$\hat{p} = \left| 0,004 - 0,0704JB + 0,0701JB^2 - 0,1207JB^3 \right| \quad (18)$$

$$\hat{p} = \left| 0,0352 - 0,6677JB + 0,4426JB^2 - 0,3241JB^3 \right| \quad (19)$$

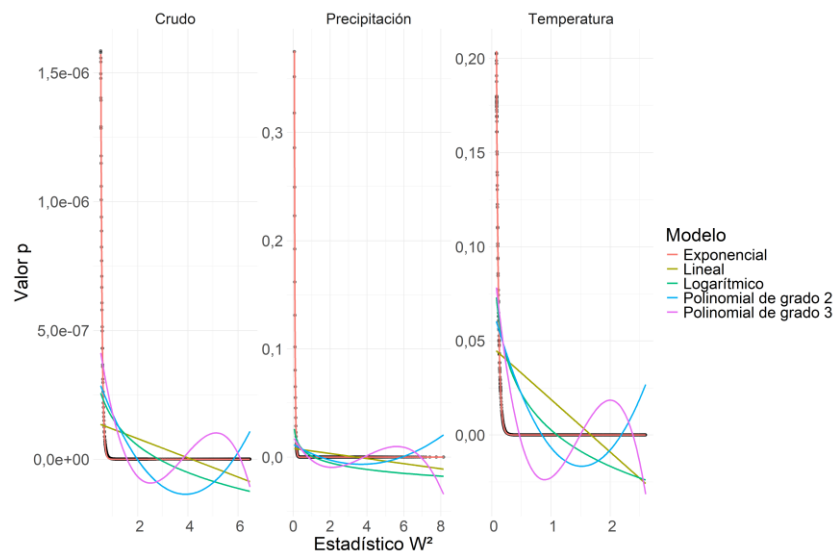


Figura 8. Ajuste de curva – valor de p y el estadístico W^2

En lo que corresponde a la prueba CvM, las variables precipitación, crudo y temperatura presentaron un mejor ajuste de curva en su primera aproximación mediante modelos exponenciales (ver ecuaciones 20 a 22) en todos casos. En los modelos exponenciales, se produce un decrecimiento debido al exponente negativo, esto significa que el valor de p disminuye exponencialmente a medida que el estadístico W^2 aumenta, y viceversa. Los modelos exponenciales están acotados con $W^2 \in [0,0597; 8,1505]$, $W^2 \in [0,5255; 6,4766]$ y $W^2 \in [0,0804; 2,5972]$ respectivamente. Es importante indicar que, la notación matemática del estadístico es W^2 , el superíndice no indica que al estadístico se debe elevar al cuadrado.

$$\hat{p} = 2,3317e^{-30,5783(W^2)} \quad (20)$$

$$\hat{p} = 0,0485e^{-19,6693(W^2)} \quad (21)$$

$$\hat{p} = 2,4986e^{-31,1508(W^2)} \quad (22)$$

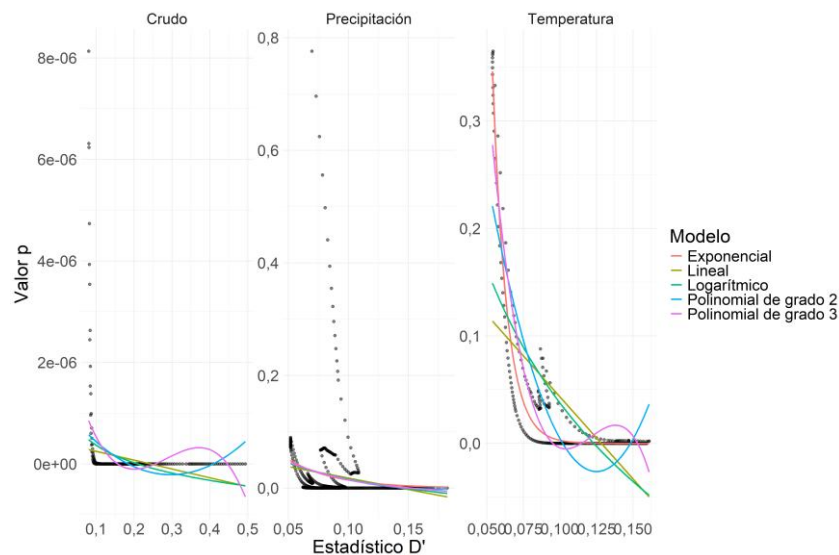


Figura 9. Ajuste de curva – valor de p y el estadístico D'

En lo que respecta a la prueba M.KS, la variable temperatura presentó un mejor ajuste de curva en su primera aproximación mediante un modelo exponencial (ver ecuación 23). El modelo produce un decrecimiento debido al exponente negativo, esto significa que el valor de p disminuye exponencialmente a medida que el estadístico D' aumenta, y viceversa. El modelo exponencial está acotado con $D' \in [0,0539;0,1610]$.

$$\hat{p} = 93,6107e^{-103,8542(D')} \quad (23)$$

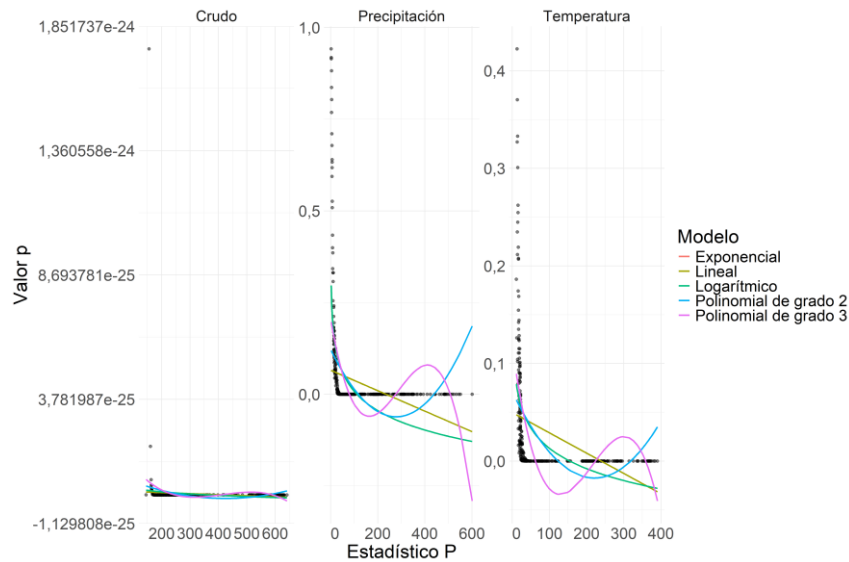


Figura 10. Ajuste de curva – valor de p y el estadístico P

Con respecto a la prueba $P \chi^2$, las variables precipitación y temperatura presentaron un mejor ajuste de curva en su primera aproximación mediante modelos logarítmico (ver ecuación 24) y polinomial de grado 3 (ver ecuación 25) respectivamente. En el caso al modelo logarítmico, el coeficiente negativo que acompaña a la expresión logarítmica indica que, el valor de p disminuye a medida que el estadístico P aumenta, y viceversa. Por otro lado, en el modelo polinomial de grado 3, la interpretación de los coeficientes influye a medida que el estadístico P varía, el valor de p puede disminuir o incrementar dependiendo

del equilibrio de los términos lineales, cuadrático y cúbico. Los modelos logarítmico y polinomial de grado 3 están acotados con $P \in [2, 296; 601, 453]$ y $P \in [11, 28; 390, 59]$ respectivamente.

$$\hat{p} = 0,3598 - 0,0763 \log(P) \quad (24)$$

$$\hat{p} = \left| 0,0271 - 0,4353P + 0,3416P^2 - 0,4059P^3 \right| \quad (25)$$

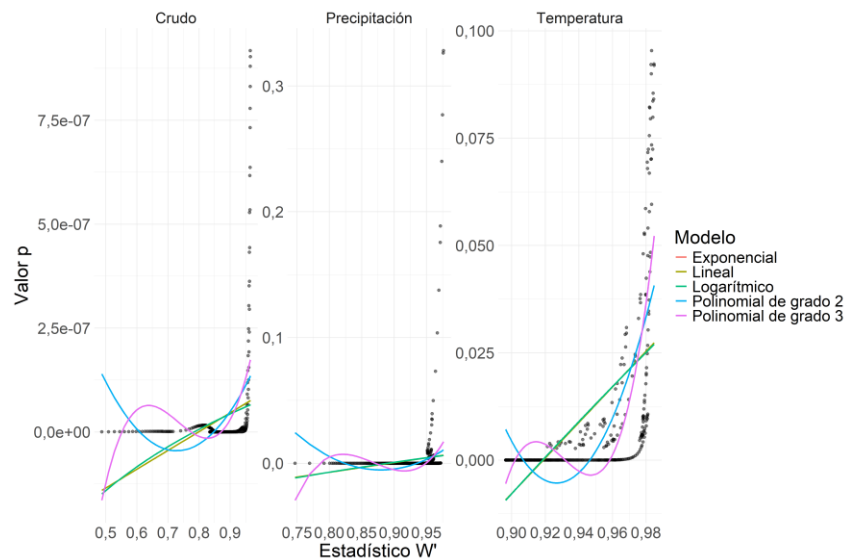


Figura 11. Ajuste de curva – valor de p y el estadístico W'

En lo referente a la prueba SF, la variable temperatura mínima absoluta presentó un mejor ajuste de curva en su primera aproximación mediante un modelo polinomial de grado 3 (ver ecuación 26). En el modelo polinomial, la interpretación de los coeficientes influye a medida que el estadístico W' varía, el valor de p puede disminuir o incrementar dependiendo del equilibrio de los términos lineales, cuadrático y cúbico. El modelo polinomial de grado 3 está acotado con $w' \in [0,8965; 0,9848]$.

$$\hat{p} = \left| -362,8148 + 1170,145(W') - 1257,5917(W')^2 + 450,3867(W')^3 \right| \quad (26)$$

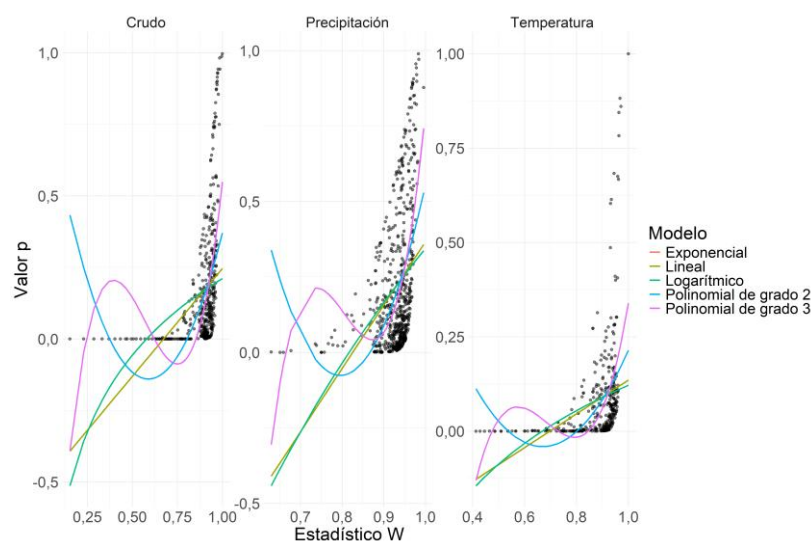


Figura 12. Ajuste de curva – valor de p y el estadístico W

En lo que respecta a la prueba SW, las variables precipitación, producción de petróleo mensual operativo y temperatura mínima absoluta presentaron un mejor ajuste de curva en su primera aproximación mediante un modelo polinomial de grado 3 (ver ecuaciones 27 a 29). En el modelo polinomial, la interpretación de los coeficientes influye a medida que el estadístico W varía, el valor de p puede disminuir o incrementar dependiendo del equilibrio de los términos lineales, cuadrático y cúbico. Los modelos polinomiales de grado 3 están acotados con $w \in [0,6298;0,9959]$, $w \in [0,1530;1,0000]$ y $w \in [0,4128;1,0000]$ respectivamente.

$$\hat{p} = \left| -72,4629 + 274,1903(w) - 342,7814(w)^2 + 141,8534(w)^3 \right| \quad (27)$$

$$\hat{p} = \left| -1,6939 + 11,6496(w) - 22,4849(w)^2 + 13,0778(w)^3 \right| \quad (28)$$

$$\hat{p} = \left| -4,5067 + 21,0047(w) - 31,5604(w)^2 + 15,4017(w)^3 \right| \quad (29)$$

Tabla 4. Comparación de modelos

Prueba	Variable	Modelo	R ²	SCE
KS	Precipitación	Lineal	0,1634	17,9472
		Polinomial de grado 2	0,1856	17,4715
		Polinomial de grado 3	0,1953	17,2635
		Exponencial	0,1966	17,2356
		Logarítmico	0,1834	17,5191
	Crudo	Lineal	0,1632	0,0027
		Polinomial de grado 2	0,3798	0,002
		Polinomial de grado 3	0,5081	0,0016
		Exponencial	0,9784*	1e ⁻⁴
		Logarítmico	0,3018	0,0022
	Temperatura	Lineal	0,5326	8,5229
		Polinomial de grado 2	0,7972	3,6974
		Polinomial de grado 3	0,8019*	3,612
		Exponencial	0,7857	3,908
		Logarítmico	0,6438	6,4957
AD	Precipitación	Lineal	0,0098	0,3926
		Polinomial de grado 2	0,0245	0,3868
		Polinomial de grado 3	0,0523	0,3758
		Exponencial	1*	0
		Logarítmico	0,074	0,3672
	Crudo	Lineal	0,0394	0
		Polinomial de grado 2	0,1394	0
		Polinomial de grado 3	0,2316	0
		Exponencial	NA	NA
		Logarítmico	0,1045	0
	Temperatura	Lineal	0,1697	0,2354
		Polinomial de grado 2	0,2793	0,2039
		Polinomial de grado 3	0,4002	0,1697
		Exponencial	0,9999*	0
		Logarítmico	0,3442	0,1855
JB	Precipitación	Lineal	0,0072	0,8001
		Polinomial grado 2	0,0224	0,7879
		Polinomial grado 3	0,0322	0,78
		Exponencial	NC	NC
		Logarítmico	0,2003*	0,6445
	Crudo	Lineal	0,0568	0,0824
		Polinomial grado 2	0,113	0,0775
		Polinomial grado 3	0,2797*	0,0629
		Exponencial	NC	NC
		Logarítmico	0,1732	0,0722
		Lineal	0,5315	0,3931

Prueba	Variable	Modelo	R ²	SCE
CvM	Temperatura	Polinomial grado 2	0,765	0,1972
		Polinomial grado 3	0,8902*	0,0921
		Exponencial	NC	NC
		Logarítmico	0,7436	0,2151
	Precipitación	Lineal	0,0126	0,6455
		Polinomial de grado 2	0,028	0,6354
		Polinomial de grado 3	0,0535	0,6188
		Exponencial	1*	0
		Logarítmico	0,0858	0,5977
	Crudo	Lineal	0,0701	0
		Polinomial de grado 2	0,2161	0
		Polinomial de grado 3	0,3361	0
		Exponencial	0,9998*	0
		Logarítmico	0,1867	0
	Temperatura	Lineal	0,1975	0,722
		Polinomial de grado 2	0,3303	0,6025
		Polinomial de grado 3	0,4664	0,4801
		Exponencial	1*	0
		Logarítmico	0,4172	0,5243
M.KS	Precipitación	Lineal	0,0257	3,0186
		Polinomial de grado 2	0,0274	3,0133
		Polinomial de grado 3	0,0275	3,0132
		Exponencial	0,0272	3,0139
		Logarítmico	0,0275	3,0131
	Crudo	Lineal	0,0408	0
		Polinomial de grado 2	0,1053	0
		Polinomial de grado 3	0,162	0
		Exponencial	NC	NC
		Logarítmico	0,0834	0
	Temperatura	Lineal	0,3773	1,2508
		Polinomial de grado 2	0,7419	0,5184
		Polinomial de grado 3	0,8585	0,2842
		Exponencial	0,936*	0,1286
		Logarítmico	0,5055	0,9933
P χ^2	Precipitación	Lineal	0,0637	8,9564
		Polinomial de grado 2	0,1686	7,953
		Polinomial de grado 3	0,3068	6,6308
		Exponencial	NC	NC
		Logarítmico	0,3798*	5,9327
	Crudo	Lineal	0,0049	0
		Polinomial de grado 2	0,0191	0
		Polinomial de grado 3	0,0324	0
		Exponencial	NC	NC
		Logarítmico	0,0092	0
	Temperatura	Lineal	0,1395	1,1687
		Polinomial de grado 2	0,2254	1,052
		Polinomial de grado 3	0,3467*	0,8873
		Exponencial	NC	NC
		Logarítmico	0,2829	0,9739
SF	Precipitación	Lineal	0,0106	0,4465
		Polinomial de grado 2	0,0246	0,4402
		Polinomial de grado 3	0,0435	0,4316
		Exponencial	NC	NC
		Logarítmico	0,01	0,4468
		Lineal	0,0885	0
		Polinomial de grado 2	0,193	0

Prueba	Variable	Modelo	R^2	SCE
SF	Crudo	Polinomial de grado 3	0,2595	0
		Exponencial	NC	NC
		Logarítmico	0,0702	0
	Temperatura	Lineal	0,3336	0,108
		Polinomial de grado 2	0,5072	0,0799
		Polinomial de grado 3	0,6019*	0,0645
		Exponencial	NC	NC
		Logarítmico	0,328	0,1089
	Precipitación	Lineal	0,15	28,6284
		Polinomial de grado 2	0,2425	25,514
		Polinomial de grado 3	0,3081*	23,3045
		Exponencial	NA	NA
		Logarítmico	0,1348	29,1406
SW	Crudo	Lineal	0,1431	20,3126
		Polinomial de grado 2	0,2623	17,4858
		Polinomial de grado 3	0,3712*	14,9038
		Exponencial	NC	NC
		Logarítmico	0,0945	21,4644
	Temperatura	Lineal	0,1133	6,5641
		Polinomial de grado 2	0,1857	6,0284
		Polinomial de grado 3	0,2442*	5,5952
		Exponencial	NC	NC
		Logarítmico	0,0944	6,704

Nota. El símbolo * indica el modelo con mejor ajuste. NC indica que el modelo no converge

4. Discusión

Para las variables precipitación, crudo y temperatura consideradas en el análisis, se aplicaron las pruebas de normalidad univariadas KS, AD, JB CvM, M.KS, P χ^2 y SF en base a incrementos de una observación con muestra inicial de 50 hasta completar la muestra total correspondiente a cada variable. Cada submuestra generó un par ordenado compuesto por el estadístico de cada prueba y el valor de p.

Como primera aproximación, los modelos evaluados fueron: lineal, polinomio de grado 2 y 3, exponencial y logarítmica. La validación de cada modelo se realizó considerando el coeficiente de determinación (R^2), estableciendo un valor mínimo referencial (umbral mínimo) de 0,2 para aceptar la expresión matemática obtenida. Los modelos con un R^2 inferior a 0,2 no fueron considerados. El valor de R^2 está acotado entre [0,1], donde 0 indica ausencia de relación entre las variables y 1 representa una relación perfecta o ajuste de curva ideal (ver umbrales de R^2 en la tabla 1).

El R^2 ajustado resulta de mucha utilidad cuando se incorpora variables predictoras (variables de entrada o independientes), las cuales no formaron parte en la presente investigación. Además, en el análisis de cada variable según las pruebas de normalidad aplicadas, los valores de R^2 ajustado fueron inferiores a los de R^2 . Con respecto a la SCE, se utilizó únicamente para verificar el comportamiento siguiendo el esquema relacional entre R^2 y SCE, como se muestra en la figura 1.

Para la prueba KS, las variables crudo y temperatura presentaron un mejor ajuste a los modelos exponencial ($R^2 = 0,9784$) y polinomial de grado 3 ($R^2 = 0,8019$) respectivamente, al analizar la relación entre el estadístico D y el valor de p , dichos resultados son considerados como una relación sustancial. Por otro lado, en la prueba AD, las variables precipitación y temperatura se ajustaron de mejor manera al modelo exponencial ($R^2 = 1$ y $R^2 = 0,9999$ respectivamente) en la relación entre el estadístico A^2 y el valor de p , siendo los resultados una relación sustancial.

En el caso de la prueba JB, las variables precipitación, crudo y temperatura presentaron un mejor ajuste a los modelos logarítmicos ($R^2 = 0,2003$), polinomial de grado 3 en las dos últimas variables ($R^2 = 0,2797$ y $R^2 = 0,8902$), al relacionar el estadístico JB y el valor de p . Estos resultados se interpretan como relaciones muy débil y sustancial según los valores obtenidos. En lo referente a la prueba CvM, las variables precipitación, crudo y temperatura se ajustaron adecuadamente al modelo exponencial en todos los casos ($R^2 = 1$, $R^2 = 0,9998$ y $R^2 = 1$) respectivamente. Esto significa que el estadístico W^2 y el valor de p tienen una relación sustancial.

Con respecto a la prueba M.KS, la variable temperatura tuvo un mejor ajuste al modelo exponencial ($R^2 = 0,936$), lo que indica que la relación entre el estadístico D' y el valor de p es sustancial. Por otro lado, para la prueba $P \chi^2$, las variables precipitación y temperatura mostraron un mejor ajuste a los modelos logarítmico ($R^2 = 0,3798$) y polinomial de grado 3 ($R^2 = 0,3467$) respectivamente. Estos resultados indican que la relación entre el estadístico P y el valor de p es débil. Finalmente, para la prueba SF, la variable temperatura presentó un ajuste adecuado al modelo polinomial de grado 3 ($R^2 = 0,6019$), lo que sugiere que la relación entre el estadístico W' y el valor de p es moderada.

Finalmente, la prueba SW, las variables precipitación, crudo y temperatura (analizado por bloques) mostraron un mejor ajuste al modelo polinomial de grado 3 en todos los casos con $R^2 = 0,3081$, $R^2 = 0,3712$ y $R^2 = 0,2442$. Estos valores obtenidos de R^2 indican que la relación entre el estadístico W y el valor de p es débil y muy débil.

Diseñar un modelo matemático que relacione las pruebas de normalidad univariadas con sus respectivos estadísticos y los valores de p constituye el primer paso para profundizar en el criterio del nivel de significancia. El valor de p puede generar controversias en la toma de decisiones, por ende, es importante tomar una posición metodológica rigurosa para que el proceso de validar una hipótesis tenga sustento estadístico. Tener un valor de p menor que 0,05 no garantiza que las variables de estudio estén relacionadas, más bien, depende en gran medida del análisis previo de los datos que el investigador realice, a fin de que el estudio sea más significativo. Lo mencionado coincide con lo señalado por [20], donde se indica que realizar un estudio comparativo de las pruebas de normalidad, variando el tamaño muestral, resulta de gran utilidad para argumentar que el mismo estadístico de prueba puede producir valores de p muy distintos según el contexto y las variables analizadas.

Al momento de utilizar las pruebas de normalidad univariadas es importante conocer la forma de aplicación. Como se mencionó en apartados anteriores, algunas pruebas se basan en medir distancias de cada observación con respecto a la media aritmética, donde el orden

de los datos es importante. En cambio, pruebas como la de Jarque-Bera está direccionada en evaluar momentos estadísticos (asimetría y curtosis) sin importar el ordenamiento de los datos.

La prueba de Jarque-Bera está disponible en el paquete "tseries", en la librería library(tseries) de R donde viene incorporado por defecto el coeficiente de asimetría de Fisher clásico, por esta razón, el estadístico JB genera valores excesivamente altos, afectando el comportamiento de las variables sobre todo cuando existen valores atípicos. Por esta razón, se propone reajustar el cálculo del estadístico JB mediante el uso del coeficiente de asimetría de Fisher ajustado, lo cual permitiría corregir los sesgos existentes en un conjunto de datos y podría mejorar la aproximación a la distribución normal.

Para reforzar el análisis de los estadísticos de las pruebas de normalidad univariada, sería factible incorporar otro tipo de pruebas actualizadas, con el objetivo de comparar la utilidad de los diferentes momentos estadísticos y evaluar el comportamiento del supuesto de normalidad, así como su aplicabilidad en diferentes tipos de variables. Se recomienda ampliar los tamaños muestrales mediante el incremento progresivo de submuestras para verificar el comportamiento de los estadísticos de las pruebas de normalidad y los valores de p , esto permitiría observar patrones o tendencias que contribuyan en obtener resultados más eficientes, consistentes y robustos [21].

En estadística paramétrica, la verificación del supuesto de normalidad se ha convertido en el punto de partida en los trabajos de investigación, especialmente cuando se emplean variables continuas. Por ende, se recomienda diseñar un modelo matemático preliminar que se ajuste adecuadamente a un conjunto amplio de datos, con el propósito de obtener valores de p más consistente y significativo. Según [22], se evidencia que el estadístico de prueba puede arrojar valores de p distintos a medida que el tamaño muestral aumenta o disminuye, lo que refuerza la necesidad de aplicar el submuestreo para una interpretación precisa del valor de p .

En lo que concierne a los modelos más sofisticados y flexibles que se pueden aplicar para aumentar R^2 y tener un mejor ajuste de curva, se recomienda aplicar los modelos de regresión avanzada como: regresión spline (splines cúbicos, B-splines o P-splines) para dividir el conjunto de datos en tramos y realizar un ajuste flexible con continuidad y suavidad [23]; regresión local (LOESS o LOWESS) cuando el comportamiento de los datos son no paramétricos y necesita realizar un ajuste local para cada dato, permitiendo encontrar patrones complejos [24]; regresión polinomial con regularización (ridge o lasso) para evitar sobreajuste en polinomios de grados superiores [25]; además se propone aplicar modelos aditivos generalizados (GAM, por sus siglas en inglés) para modelar relaciones no lineales entre las variables de manera flexible [26].

5. Conclusiones

Se utilizaron distintas pruebas de normalidad univariadas con el objetivo de verificar el comportamiento de cada uno de los estadísticos y los respectivos valores de p . Este análisis permitió encontrar patrones consistentes que se ajustaron a los modelos de regresión clásicos para determinar la relación entre las variables.

Con la aplicación de los modelos clásicos como primera aproximación se logró un avance significativo en el comportamiento de los estadísticos de las pruebas de normalidad univariadas (KS, AD, JB CvM, M.KS, P χ^2 , SF y SW), así como los valores de p obtenidos en cada iteración al incrementar o decrementar el tamaño de la muestra.

Los modelos que mostraron un mejor ajuste de curva a las variables precipitación, producción de petróleo mensual operativo y temperatura mínima absoluta fueron los exponenciales y polinomial de grado 3, presentando los valores de R^2 más altos: en la prueba KS, para la variable crudo ($R^2 = 0,9784$); en la prueba AD, para la variable precipitación ($R^2 = 1$); en la prueba JB, para la variable temperatura ($R^2 = 0,8902$); en la prueba CvM, para la variable temperatura ($R^2 = 1$); en la prueba M.KS, para la variable temperatura ($R^2 = 0,936$); en la prueba P χ^2 , para la variable precipitación ($R^2 = 0,3798$); en la prueba SF, para la variable temperatura ($R^2 = 0,6019$); y en la prueba SW, para la variable crudo ($R^2 = 0,3712$).

Contribuciones de los autores

Conceptualización, J.M. y A.S.; metodología, J.M. y A.S.; software, J.M.; validación, J.M.; análisis formal, J.M.; investigación, J.M.; recursos, J.M.; curación de datos, J.M.; redacción—preparación del borrador original, J.M. y A.S.; redacción—revisión y edición, J.M. y A.S.; visualización, J.M. Todos los autores han leído y aprobado la versión publicada del manuscrito.

Conflicto de Interés

Los autores declaran que no existen conflictos de interés de naturaleza alguna con la presente investigación.

Declaración sobre el uso de Inteligencia Artificial Generativa

En la preparación de este artículo, se utilizó ChatGPT para corrección gramatical y la revisión de fragmentos de código de programación.

Referencias

- [1] A. Ferencek y M. K. Borštnar, "Data quality assessment in product failure prediction models," *Journal of Decision Systems*, vol. 29, n. 1, pp. 79–86, jun. 2020, doi: 10.1080/12460125.2020.1776927.
- [2] R. H. Taplin, "Quantifying data quality after removing respondents who fail data quality checks" *Current Issues in Tourism*, vol. 28, n. 16, pp. 2570-2581, jul. 2024, doi: 10.1080/13683500.2024.2378611.
- [3] A. Inglis, A. Parnell, y C. B. Hurley, "Visualizing variable importance and variable interaction effects in machine learning models," *Journal of Computational and Graphical Statistics*, vol. 31, n. 3, pp. 766–778, ene. 2022, doi: 10.1080/10618600.2021.2007935.
- [4] L. M. Hudiburgh y D. Garbinsky, "Data visualization: Bringing data to life in an introductory statistics course," *Journal of Statistics Education*, vol. 28, n. 3, pp. 262–279, ago. 2020, doi: 10.1080/10691898.2020.1796399.

- [5] C. Avram y M. Mărușteri, "Normality assessment, few paradigms and use cases," *Revista Romana de Medicina de Laborator*, vol. 30, n. 3, pp. 251–260, jul. 2022, doi: 10.2478/rrlm-2022-0030.
- [6] M. Kumagai, Y. Ando, A. Tanaka, K. Tsuda, Y. Katsura, y K. Kurosaki, "Effects of data bias on machine-learning-based material discovery using experimental property data," *Science and Technology of Advanced Materials: Methods*, vol. 2, n. 1, pp. 302–309, ago. 2022, doi: 10.1080/27660400.2022.2109447.
- [7] S. Pawel y L. Held, "Closed-form power and sample size calculations for Bayes factors", *The American Statistician*, vol. 79, n. 3, pp. 330–344, abr. 2025, doi: 10.1080/00031305.2025.2467919.
- [8] M. Javed y M. Irfan, "A simulation study: New optimal estimators for population mean by using dual auxiliary information in stratified random sampling," *Journal of Taibah University for Science*, vol. 14, n. 1, pp. 557–568, ene. 2020, doi: 10.1080/16583655.2020.1752004.
- [9] Y. Pawitan, "Defending the p-value," *arXiv*, 2020, doi: 10.48550/arXiv.2009.02099.
- [10] L. P. M. Diaz-Ballve, "El valor p: un concepto estadístico-metodológico omnipresente en la investigación biomédica. ¿Lo interpretamos correctamente?," *Argentinian Journal of Respiratory & Physical Therapy*, vol. 2, n. 1, feb. 2020, doi: 10.58172/ajrpt.v2i1.91.
- [11] P. C. Austin, I. Eekhout, y S. van Buuren, "Evaluating the median p-value method for assessing the statistical significance of tests when using multiple imputation," *Journal of Applied Statistics*, vol. 52, n. 6, pp. 1161–1176, oct. 2024, doi: 10.1080/02664763.2024.2418473.
- [12] A. Ebbelind, "A functional view on language: a methodology for mathematics education to study shifts in prospective teachers' discursive patterns," *International Journal of Mathematical Education in Science and Technology*, vol. 54, n. 8, pp. 1731–1745, may. 2023, doi: 10.1080/0020739X.2023.2204506.
- [13] J. F. Hair, C. M. Ringle, y M. Sarstedt, "PLS-SEM: Indeed a silver bullet," *Journal of Marketing Theory and Practice*, vol. 19, n. 2, pp. 139–152, abr. 2011, doi: 10.2753/MTP1069-6679190202.
- [14] E. E. Pinto Aragón, A. R. Villa Navas, y H. A. Pinto Aragón, "Estrés académico en estudiantes de la Universidad de La Guajira, Colombia," *REV CIENC SOC-VENEZ*, vol. 28, no. 5, pp. 87–99, may. 2022. Disponible en: <https://dialnet.unirioja.es/servlet/articulo?codigo=8471675>.
- [15] V. Glinskiy, Y. Ismayilova, S. Khrushchev, A. Logachov, O. Logachova, L. Serga, A. Yambartev y K. Zaykov, "Modifications to the Jarque–Bera test," *Mathematics*, vol. 12, n. 16, p. 2523, ago. 2024, doi: 10.3390/math12162523.
- [16] B. O. Emmanuel, N. T. Maureen, y N. Wonu, "Detection of non-normality in data sets and comparison between different normality tests," *Asian Journal of Probability and Statistics*, vol. 5, n. 4, pp. 1–20, ene. 2020, doi: 10.9734/ajpas/2019/v5i430149.
- [17] A. Kamath, S. Poojari, y K. Varsha, "Assessing the robustness of normality tests under varying skewness and kurtosis: a practical checklist for public health researchers," *BMC Medical Research Methodology*, vol. 25, n. 1, p. 206, sep. 2025, doi: 10.1186/s12874-025-02641-y.
- [18] S. Tenny y I. Abdelgawad, "Statistical Significance," *StatPearls [Internet]*, Treasure Island (FL): StatPearls Publishing, Updated Nov. 23, 2023. Disponible en: <https://www.ncbi.nlm.nih.gov/books/NBK459346/>
- [19] R Core Team, "R: A Language and Environment for Statistical Computing." [r-project.org](https://www.R-project.org/). <https://www.R-project.org/>
- [20] S. S. Uyanto, "An extensive comparisons of 50 univariate goodness-of-fit tests for normality," *Austrian Journal of Statistics*, vol. 51, n. 3, pp. 45–97, ago. 2022, doi: 10.17713/ajs.v51i3.1279.
- [21] S. Korkmaz y Y. Demir, "Investigation of some univariate normality tests in terms of type-I errors and test power," *Journal of Scientific Reports-A*, n. 52, pp. 376–395, mar. 2023, doi: 10.59313/jsr-a.1222979.

- [22] S. Demir, "Comparison of normality tests in terms of sample sizes under different skewness and kurtosis coefficients," *International Journal of Assessment Tools in Education*, vol. 9, n. 2, pp. 397–409, jun. 2022, doi: 10.21449/ijate.1101295.
- [23] J. Gauthier, Q. V. Wu, y T. A. Gooley, "Cubic splines to model relationships between continuous variables and outcomes: a guide for clinicians," *Bone Marrow Transplant.*, vol. 55, n. 4, pp. 675–680, oct. 2019, doi: 10.1038/s41409-019-0679-x.
- [24] Y. Xie, Z. Jing, H. Pan, X. Xu, y Q. Fang, "Redefining the high variable genes by optimized LOESS regression with positive ratio," *BMC Bioinformatics*, vol. 26, n. 104, abr. 2025, doi: 10.1186/s12859-025-06112-5.
- [25] S. K. Safi, M. Alsheryani, M. Alrashdi, R. Suleiman, D. Awwad, y Z. N. Abdalla, "Optimizing linear regression models with lasso and ridge regression: A study on UAE financial behavior during COVID-19," *Migration Letters*, vol. 20, n. 6, pp. 139–153, sep. 2023, doi: 10.59670/ml.v20i6.3468.
- [26] S. N. Wood, "Inference and computation with generalized additive models and their extensions," *TEST*, vol. 29, n. 2, pp. 307–339, abr. 2020, doi: 10.1007/s11749-020-00711-5.