

Simulación de sistemas con variables aleatorias para la toma de decisiones estratégicas

María del Consuelo Argüelles Arellano

Instituto Tecnológico de Zacatecas,
México

mariac.aa@zacatecas.tecnm.mx

Resumen. La simulación de sistemas con variables aleatorias de distribución discreta o continua es una técnica muy útil en muchas situaciones en las que es necesario analizar el comportamiento de sistemas complejos en los que algunas variables tienen una distribución de probabilidad y se comportan de forma aleatoria. Esta técnica se utiliza para modelar la incertidumbre y analizar cómo afecta el comportamiento del sistema, evaluar diferentes escenarios y tomar decisiones basadas en datos, reducción de riesgos y costos asociados con la implementación de cambios en el sistema, identificar problemas y cuellos de botella, mejorar el rendimiento con resultados más eficientes y productivos en un sistema. La simulación de sistemas con variables aleatorias de distribución probabilística es una técnica muy útil para modelar sistemas complejos en los que algunas variables tienen una distribución discreta o continua y se comportan de forma aleatoria. Esta técnica permite a los usuarios evaluar diferentes escenarios, reducir riesgos, identificar problemas, mejorar el rendimiento y tomar decisiones basadas en datos.

Palabras clave. Simulación de sistemas, generación de variables aleatorias, ingeniería de sistemas, modelos de simulación.

Simulation of Systems with Random Variables for Making Strategic Decisions

Abstract. The simulation of systems with random variables of a discrete or continuous distribution is a very useful technique in many situations in which it is necessary to analyze the behavior of complex systems in which some variables have a probability distribution and behave randomly. This technique is used to model uncertainty and analyze how it affects the behavior of the system, evaluate different scenarios and make decisions based on data, reduce risks and costs associated with the implementation of system changes, identify problems and bottlenecks, improve performance with

more efficient and productive results in one system. The simulation of systems with random variables of probabilistic distribution is a very useful technique to model complex systems in which some variables have a discrete or continuous distribution and behave randomly. This technique allows users to evaluate different scenarios, reduce risks, identify problems, improve performance, and make data-driven decisions.

Keywords. Systems simulation, random variable generation, systems engineering, simulation models.

1. Introducción

La simulación, el modelo y el sistema son conceptos interrelacionados que se utilizan en diversas disciplinas para entender, analizar y predecir el comportamiento de sistemas complejos en el mundo real.

El concepto de simulación se refiere al proceso de imitar el comportamiento de un sistema en tiempo real utilizando un modelo computacional. La simulación se utiliza para predecir el comportamiento de un sistema en situaciones en las que no es posible realizar experimentos prácticos debido a limitaciones de tiempo, costo o éticas.

Por ejemplo, se pueden simular escenarios de terremotos para predecir el daño potencial en edificios, o se pueden simular sistemas de transporte para evaluar el impacto de cambios en la infraestructura. Las simulaciones pueden variar desde simples modelos matemáticos hasta modelos complejos que incorporan múltiples variables y factores ambientales.

El concepto de modelo se refiere a una representación simplificada y abstracta de un sistema real que se utiliza para comprender y

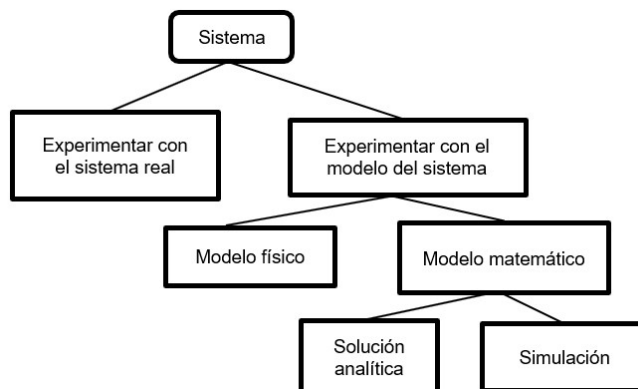


Fig. 1. Diagrama del análisis de un sistema

predecir su comportamiento. Los modelos pueden ser matemáticos, físicos o conceptuales, y se utilizan en una amplia variedad de disciplinas para representar sistemas complejos.

Por ejemplo, los modelos de árboles se utilizan en ecología para entender cómo los diferentes factores ambientales afectan el crecimiento de los árboles. Los modelos también se utilizan en la economía para predecir el comportamiento de los mercados financieros y en la ingeniería para diseñar sistemas de control y de producción.

El concepto de sistema se refiere a un conjunto de elementos interrelacionados que interactúan entre sí para lograr un objetivo común. Los sistemas pueden ser físicos, biológicos, sociales o tecnológicos, y se pueden estudiar a diferentes escalas, desde sistemas moleculares hasta sistemas planetarios.

Los sistemas se caracterizan por tener entradas, procesos internos y salidas. Por ejemplo, un sistema de producción de automóviles puede tener entradas como materias primas y energía, procesos internos como ensamblaje y pruebas, y salidas como automóviles completos.

Las relaciones entre la simulación, el modelo y el sistema son complejas y multifacéticas. La simulación es una herramienta que se utiliza para evaluar la precisión de un modelo al comparar las predicciones de la simulación con los datos reales del sistema.

Los modelos se utilizan para representar y simplificar los sistemas complejos, lo que permite una comprensión más profunda de su comportamiento. La simulación y el modelo

también pueden utilizarse juntos para probar diferentes escenarios y variables, lo que permite la evaluación de diferentes estrategias y decisiones.

En conclusión, la simulación, el modelo y el sistema son conceptos interrelacionados que se utilizan para entender y predecir el comportamiento de sistemas complejos en el mundo real.

La simulación es una herramienta valiosa para evaluar la precisión de un modelo, mientras que el modelo se utiliza para simplificar y representar el comportamiento de un sistema.

La comprensión de estas relaciones puede ayudar a los investigadores a desarrollar modelos y simulaciones más precisos y útiles para una amplia variedad de disciplinas.

Experimentar con un sistema real implica trabajar con un objeto o proceso físico en el mundo real, mientras que experimentar con un modelo de un sistema implica trabajar con una representación matemática o simulación computarizada del objeto o proceso. A continuación, se presentan algunas diferencias entre experimentar con un sistema real y experimentar con un modelo de un sistema:

Control: Al trabajar con un modelo de un sistema, es posible controlar y manipular las variables de entrada con mayor precisión y facilidad que en un sistema real. En un sistema real, las variables pueden ser más difíciles de medir y controlar debido a la complejidad del entorno y la naturaleza impredecible de algunas variables.

Costo: Experimentar con un modelo de un sistema suele ser más económico que

experimentar con un sistema real. En el caso de los sistemas físicos, pueden requerirse grandes cantidades de recursos, materiales y mano de obra para llevar a cabo experimentos, mientras que en los modelos se requieren recursos de computación y software.

Tiempo: El tiempo necesario para realizar experimentos con sistemas reales puede ser mucho mayor que el necesario para experimentar con modelos de sistemas. En el caso de los modelos, se pueden realizar simulaciones en paralelo y experimentar con diferentes configuraciones de variables de entrada en cuestión de minutos u horas, mientras que los sistemas reales pueden tardar días o incluso semanas en realizar una sola prueba.

Precisión: Los modelos de sistemas pueden ser más precisos que los sistemas reales en ciertos casos, ya que es posible ajustar las variables y los parámetros para obtener resultados más precisos. Además, los modelos pueden proporcionar resultados consistentes y repetibles que pueden ser difíciles de lograr en un sistema real debido a la variabilidad inherente.

Limitaciones: Los modelos de sistemas pueden ser limitados por la calidad y la precisión de los datos de entrada y por la complejidad del sistema en sí. Además, los modelos no siempre pueden capturar todas las interacciones y complejidades de un sistema real, lo que puede llevar a resultados inexactos. En el caso de los sistemas reales, puede haber limitaciones en cuanto a lo que es posible medir y controlar, lo que puede afectar la precisión y la validez de los resultados.

En resumen, experimentar con un modelo de un sistema puede ser más económico, rápido y controlable que experimentar con un sistema real, aunque los modelos pueden tener limitaciones y no siempre pueden capturar todas las complejidades de un sistema real.

Por otro lado, experimentar con sistemas reales puede proporcionar resultados más precisos y representativos del mundo real, aunque puede ser más costoso y requerir más tiempo.

Experimentar con un modelo físico y un modelo matemático son dos formas de realizar experimentos en ingeniería y ciencias físicas, y aunque ambas tienen como objetivo estudiar sistemas, hay algunas diferencias importantes entre ellas:

Naturaleza del modelo: Un modelo físico es una representación tangible de un sistema, que se construye utilizando materiales físicos y herramientas de fabricación. Por otro lado, un modelo matemático es una representación abstracta de un sistema que se construye utilizando ecuaciones y fórmulas matemáticas.

Flexibilidad: Un modelo matemático es más flexible que un modelo físico, ya que se pueden realizar cambios rápidamente y modificar la representación matemática según se necesite. En cambio, los modelos físicos pueden requerir más tiempo y recursos para hacer cambios en la construcción y en la configuración del sistema.

Costo: Los modelos físicos pueden ser más costosos de construir y mantener que los modelos matemáticos, ya que requieren materiales, herramientas y espacio físico para su construcción y pruebas. En cambio, los modelos matemáticos pueden ser más económicos de construir y pueden ser ejecutados en cualquier computadora.

Precisión: Los modelos físicos pueden proporcionar mediciones precisas de un sistema real, ya que utilizan sensores físicos y herramientas de medición para recopilar datos. Por otro lado, los modelos matemáticos pueden proporcionar una alta precisión si se conocen con exactitud los parámetros de entrada del modelo, aunque pueden estar sujetos a errores de aproximación.

Alcance: Los modelos matemáticos son más fáciles de escalar que los modelos físicos. Por ejemplo, un modelo matemático de un sistema hidráulico se puede utilizar para predecir el comportamiento de sistemas a gran escala, mientras que un modelo físico tendría que ser construido específicamente para la escala deseada.

En resumen, tanto los modelos físicos como los modelos matemáticos tienen sus ventajas y desventajas, y la elección del modelo dependerá de los objetivos del experimento y del sistema que se esté estudiando.

Los modelos físicos son más adecuados cuando se requiere una alta precisión en la medición, y los modelos matemáticos son más adecuados cuando se requiere una mayor flexibilidad en la configuración del sistema y una mayor capacidad de escalar el modelo.

Una solución analítica y una simulación son dos herramientas importantes en la ingeniería y las ciencias físicas. A continuación, se presentan algunas diferencias entre una solución analítica y una simulación:

Naturaleza de la solución: Una solución analítica es una solución matemática cerrada que se puede obtener a partir de ecuaciones y fórmulas matemáticas. Por otro lado, una simulación es una representación numérica del sistema que se obtiene mediante el uso de herramientas de simulación computacional.

Precisión: Una solución analítica puede ser más precisa que una simulación, ya que se puede obtener una solución exacta si se conocen todas las variables y parámetros del sistema. Sin embargo, en la práctica, esto no siempre es posible, y las soluciones analíticas pueden estar sujetas a errores de aproximación y suposiciones simplificadoras. Las simulaciones, por otro lado, pueden proporcionar resultados más precisos si se conocen con precisión los parámetros de entrada y se han validado las suposiciones utilizadas en el modelo.

Flexibilidad: Una solución analítica puede ser menos flexible que una simulación, ya que las ecuaciones y las fórmulas matemáticas pueden ser complicadas y difíciles de modificar. En cambio, una simulación puede ser más flexible, ya que los parámetros del modelo se pueden modificar y ajustar para ver cómo afectan al comportamiento del sistema.

Escalabilidad: Las soluciones analíticas pueden ser más difíciles de escalar que las simulaciones, ya que las ecuaciones matemáticas pueden volverse más complejas a medida que se aumenta la complejidad del sistema. Las simulaciones pueden ser más fáciles de escalar, ya que se pueden ejecutar en diferentes configuraciones de hardware y software para manejar sistemas más grandes y complejos.

Costo: En general, una solución analítica puede ser más económica que una simulación, ya que se requiere menos recursos computacionales y menos tiempo de desarrollo. Sin embargo, esto depende del problema que se esté resolviendo y de la complejidad de las ecuaciones y los modelos matemáticos involucrados.

En resumen, tanto las soluciones analíticas como las simulaciones tienen ventajas y

desventajas, y la elección de una herramienta dependerá de los objetivos del problema y de la naturaleza del sistema que se esté estudiando.

Las soluciones analíticas son útiles cuando se puede obtener una solución exacta o se necesitan resultados rápidos y simples. Las simulaciones son útiles cuando se necesita una mayor flexibilidad y precisión en la modelización del sistema, especialmente cuando la complejidad del sistema es alta y la solución analítica no es factible.

2. Diseño de un modelo de simulación

Para el estudio de cualquier sistema, vemos que el diseño de un modelo de simulación puede ser un proceso complejo, que requiere de conocimientos en distintas áreas, como la modelación matemática, la programación y la ingeniería de sistemas. A continuación, se presentan algunas etapas clave en el diseño de un modelo de simulación:

Identificación del sistema: Se debe definir claramente qué sistema se quiere modelar y cuáles son los objetivos de la simulación. Se debe considerar cuál es el alcance de la simulación y qué aspectos del sistema se quieren analizar.

Recopilación de datos: Se deben recopilar datos relevantes sobre el sistema, como tiempos de procesamiento, capacidades de los equipos, demandas de los clientes, entre otros. Estos datos pueden ser obtenidos a través de mediciones en el sistema real o a través de información de fuentes secundarias.

Diseño del modelo: Se debe definir un modelo matemático del sistema, que permita representar los procesos y comportamientos del sistema en cuestión. Esto implica la selección de las variables relevantes, las relaciones entre ellas y las reglas que gobiernan su comportamiento.

Validación del modelo: Una vez que se ha definido el modelo, se debe verificar que es adecuado para representar el sistema real. Esto se puede hacer comparando los resultados de la simulación con los datos del sistema real, o a través de la validación cruzada con otros modelos similares.

Implementación del modelo: El modelo debe ser implementado en un software de simulación, que permita su ejecución y la obtención de resultados. En esta etapa se deben definir los parámetros y las condiciones iniciales del modelo.

Ejecución de la simulación: Se deben ejecutar múltiples simulaciones, variando los parámetros de entrada del modelo y analizando los resultados obtenidos. Se pueden utilizar técnicas de análisis estadístico para obtener conclusiones más robustas.

Validación de los resultados: Los resultados de la simulación deben ser validados en base a los objetivos establecidos previamente. Se debe analizar si el modelo es adecuado para realizar predicciones y si las conclusiones obtenidas son coherentes con el sistema real.

Comunicación de los resultados: Finalmente, se deben comunicar los resultados de la simulación, a través de informes técnicos, gráficos o presentaciones. Es importante que los resultados sean comprensibles para los distintos actores involucrados en el sistema, como gerentes, operarios o especialistas técnicos.

Un modelo de simulación se puede clasificar en discreto o continuo. En el primer caso es cuando se tienen variables de tipo discreto, en donde se manejan números enteros como por ejemplo líneas de espera. En el segundo caso es cuando se tienen las variables de tipo continuo, en donde se manejan números reales como estudio de fluidos, intercambio de calor.

Adicionalmente un modelo de simulación se puede clasificar en determinista cuando en el modelo se utilizan variables cuyos valores no se ven afectados por variaciones aleatorias y se conocen con exactitud. Por ejemplo, el modelo de inventarios conocido como lote económico.

En caso contrario se clasifica como aleatorio o estocástico, que es cuando en el modelo se utilizan variables cuyos valores sufren modificaciones aleatorias con respecto a un valor promedio; dichas variaciones pueden ser manejadas mediante distribuciones de probabilidad.

Un buen número de estos modelos se pueden encontrar en la teoría de líneas de espera. En este caso es necesario crear un modelo de simulación que imite el comportamiento de las variables.

3. Simulación con variables aleatorias

Las variables aleatorias de entrada son aquellas que representan las incertidumbres en un modelo de simulación. Estas variables pueden incluir factores como la variabilidad en los tiempos de llegada de clientes, la variabilidad en la duración de los procesos, la variabilidad en las tasas de fallas de los equipos, entre otros.

La importancia de las variables aleatorias de entrada radica en que su inclusión adecuada en el modelo es esencial para que la simulación se aproxime a la realidad. Si se ignoran o se subestiman las incertidumbres en el modelo, se corre el riesgo de generar resultados poco realistas y poco confiables.

Por ejemplo, si se está simulando el tiempo de espera de los clientes en una tienda, es importante considerar la variabilidad en los tiempos de llegada de los clientes, la variabilidad en el tiempo que cada cliente tarda en hacer sus compras, y la variabilidad en el tiempo que los cajeros tardan en procesar las transacciones.

Si se ignora la variabilidad en alguna de estas variables, los resultados simulados podrían ser muy diferentes a los observados en la realidad.

En resumen, para que un modelo de simulación se aproxime a la realidad, es crucial identificar y modelar adecuadamente las variables aleatorias de entrada que representan las incertidumbres en el sistema o proceso que se está simulando.

Cómo determinar el tipo de distribución que tienen los datos de entrada de un sistema

Para determinar el tipo de distribución que tienen un conjunto de datos, se pueden utilizar diversas técnicas estadísticas y gráficas. Las más comunes son: histograma, gráfico de probabilidad normal, gráfico de caja y bigotes, prueba de bondad de ajuste y análisis de momentos. La prueba de bondad de ajuste es más utilizada y consiste en determinar si los datos se ajustan a una distribución teórica. Si el valor p de la prueba es mayor que el nivel de significancia elegido (generalmente 0.05), se puede concluir que los datos se ajustan a la distribución teórica.

Prueba de Bondad de Ajuste

La prueba de bondad de ajuste es una técnica estadística utilizada para evaluar si una muestra de datos proviene de una distribución de probabilidad específica o no. En otras palabras, se utiliza para determinar si una muestra de datos se ajusta adecuadamente a una distribución teórica conocida.

El objetivo de la prueba de bondad de ajuste es determinar si la muestra de datos se ajusta a la distribución teórica de interés. Si la muestra de datos no se ajusta a la distribución teórica, esto puede indicar que la distribución teórica no es la más adecuada para modelar los datos o que los datos son insuficientes para hacer una conclusión definitiva.

Existen diferentes métodos para realizar una prueba de bondad de ajuste, que varían en su complejidad y en los supuestos subyacentes. Algunos de los métodos más comunes son la prueba de chi-cuadrado, la prueba de Kolmogórov-Smirnov, la prueba de Anderson-Darling y la prueba de Shapiro-Wilk.

Es importante destacar que la prueba de bondad de ajuste no garantiza que la distribución teórica sea la más adecuada para modelar los datos, pero puede proporcionar evidencia para apoyar o refutar la hipótesis de que la muestra de datos proviene de una distribución teórica específica. Además, la prueba de bondad de ajuste se utiliza a menudo en la validación de modelos estadísticos y de simulación, donde es necesario comprobar si los datos simulados se ajustan adecuadamente a la distribución teórica asumida.

Prueba de Anderson-Darling

La prueba de Anderson-Darling es una prueba de bondad de ajuste que se utiliza para determinar si un conjunto de datos proviene de una distribución teórica específica, como la distribución normal, exponencial o uniforme, entre otras. Es una prueba más poderosa que la prueba de Kolmogórov-Smirnov para muestras grandes y puede detectar diferencias más sutiles entre la distribución teórica y los datos observados.

La prueba de Anderson-Darling se basa en el cálculo de una estadística de prueba que mide la

discrepancia entre la función de distribución empírica de los datos y la función de distribución acumulada de la distribución teórica. Cuanto mayor sea el valor de la estadística de prueba, mayor será la evidencia en contra de la hipótesis nula de que los datos provienen de la distribución teórica.

El resultado de la prueba de Anderson-Darling es un valor de estadística de prueba y un valor p asociado que indica la probabilidad de que los datos provengan de la distribución teórica. Si el valor p es menor que un nivel de significancia previamente elegido (por ejemplo, 0.05), se rechaza la hipótesis nula y se concluye que los datos no siguen la distribución teórica.

Si el valor p es mayor que el nivel de significancia, no se puede rechazar la hipótesis nula y se concluye que los datos siguen la distribución teórica.

Es importante tener en cuenta que la prueba de Anderson-Darling requiere que los datos sean independientes e idénticamente distribuidos (IID) y que la distribución teórica se especifique de antemano. Además, la prueba puede ser sensible a la elección de los parámetros de la distribución teórica y a la presencia de valores atípicos en los datos.

Por lo tanto, se recomienda utilizar la prueba de Anderson-Darling en combinación con otras técnicas para evaluar la distribución de los datos y tener en cuenta el contexto y conocimiento del problema para interpretar los resultados. El procedimiento general de la prueba es:

1. Obtener n datos de la variable aleatoria a analizar.
2. Calcular la media y la varianza de los datos.
3. Ordenar los datos en forma ascendente Y_i para $i=1, 2, \dots, n$.
4. Ordenar los datos en forma descendente Y_{n+1-i} para $i=1, 2, \dots, N$.
5. Establecer de manera explícita la hipótesis nula, al proponer una distribución de probabilidad.
6. Calcular la probabilidad esperada acumulada para cada número Y_i $PEA(Y_i)$ y la probabilidad esperada acumulada para cada número Y_{n+1-i} $PEA(Y_{n+1-i})$ a partir de la función de probabilidad propuesta.

Distribución	Estadístico de prueba ajustado	Valores críticos α			
		0.1	.05	.025	.001
Parámetros conocidos $n \geq 5$	A_n^2	1.933	2.492	3.070	3.857
Normal	$A_n^2 \left(1 + \frac{4}{n} - \frac{25}{n^2} \right)$	0.632	0.751	0.870	1.029
Exponencial	$A_n^2 \left(1 + \frac{3}{5n} \right)$	1.070	1.326	1.587	1.943
De Weibull	$A_n^2 \left(1 + \frac{1}{5\sqrt{n}} \right)$	0.637	0.757	0.877	1.038
Log-logística	$A_n^2 \left(1 + \frac{1}{4\sqrt{n}} \right)$	0.563	0.660	0.769	0.906

Fig. 2. Estadísticos de prueba y valores críticos para la prueba de Anderson-Darling

ν grados de libertad	$D_{\alpha=0.1}$	$D_{\alpha=0.05}$	$D_{\alpha=0.01}$
1	0.950	0.975	0.995
2	0.776	0.842	0.929
3	0.642	0.708	0.828
4	0.564	0.624	0.733
5	0.510	0.565	0.669
6	0.470	0.521	0.618
7	0.438	0.486	0.577
8	0.411	0.457	0.543
9	0.388	0.432	0.514
10	0.368	0.410	0.490
11	0.352	0.391	0.468
12	0.338	0.375	0.450
13	0.325	0.361	0.433
14	0.314	0.349	0.418
15	0.304	0.338	0.404
16	0.295	0.328	0.392
17	0.286	0.318	0.381
18	0.278	0.309	0.371
19	0.272	0.301	0.363
20	0.264	0.294	0.356
25	0.250	0.270	0.320
30	0.220	0.240	0.290
35	0.210	0.230	0.270
40	0.193	0.215	0.258
45	0.182	0.203	0.243
Para valores mayores a 35	$\frac{1.22}{\sqrt{n}}$	$\frac{1.36}{\sqrt{n}}$	$\frac{1.63}{\sqrt{n}}$

Fig. 3. Estadísticos de prueba y valores críticos para la prueba de Kolmogórov-Smirnov

Distribución	Generador	Parámetros
Uniforme U_i	$U_i = a + (b - a)r_i$	a = Límite inferior de la distribución uniforme. b = Límite superior de la distribución uniforme. r_i = Número aleatorio con distribución uniforme entre 0 y 1.
Triangular T_i	$T_i = \begin{cases} a + \sqrt{(b-a)(c-a)}r_i & \text{si } r_i \leq \frac{(c-a)}{(b-a)} \\ b - \sqrt{(b-a)(b-c)(1-r_i)} & \text{si } r_i > \frac{(c-a)}{(b-a)} \end{cases}$	a = Límite inferior de la distribución triangular. c = Moda de la distribución triangular. b = Límite superior de la distribución triangular.
Triangular T_i	$T_i = \begin{cases} a + (c-a)\sqrt{r_i} & \text{si } r_i \leq \frac{(c-a)}{(b-a)} \\ b - [(b-c)\sqrt{1-r_i}] & \text{si } r_i > \frac{(c-a)}{(b-a)} \end{cases}$	a = Límite inferior de la distribución triangular. c = Moda de la distribución triangular. b = Límite superior de la distribución triangular.
χ^2	$\chi_i^2 = \sum_{j=1}^n Z_j^2$	Z_j = Números aleatorios con distribución normal estándar. n = Grados de libertad.
Erlang ER_i	$ER_i = -\frac{1}{k\lambda} \ln \prod_{i=1}^k (1-r_i)$	$1/\lambda$ = Valor esperado. k = Parámetro de forma.
Normal N_i	$N_i = \left[\left(\sqrt{-2 \ln(1-r_i)} \right) \cos(2\pi r_i) \right] \sigma + \mu$ $N_i = \left[\left(\sqrt{-2 \ln(1-r_i)} \right) \sin(2\pi r_i) \right] \sigma + \mu$	μ = Media de la distribución normal. σ = Desviación estándar de la distribución normal.
Normal N_i	$N_i = \left(\sum_{i=1}^{12} r_i - 6 \right) \sigma + \mu$	μ = Media de la distribución normal. σ = Desviación estándar de la distribución normal.

Fig. 4. Generadores de variables aleatorias (parte 1)

- Calcular el estadístico de prueba $A_n^2 = [n + \frac{1}{n} \sum_{i=1}^n (2i-1) [\ln(\text{PEA}(Y_i)) + \ln(\text{PEA}(Y_{n+1-i}))]]$.
- Ajustar el estadístico de prueba de acuerdo con la distribución de probabilidad propuesta
- Definir el nivel de significancia de la prueba α y determinar su valor crítico α_n según tabla anexa.
- Comparar el estadístico de prueba con el valor crítico. Si el estadístico de prueba es menor

que el valor crítico no se puede rechazar la hipótesis nula.

Prueba de Kolmogórov-Smirnov

La prueba de Kolmogórov-Smirnov es una técnica estadística utilizada para evaluar si una muestra de datos proviene de una distribución de probabilidad específica. En otras palabras, se utiliza para determinar si una muestra de datos se

Exponencial E_i	$E_i = -\frac{1}{\lambda} \ln(1-r_i)$	$1/\lambda$ = Media de la distribución exponencial.
Weibull W_i	$W_i = \gamma + \beta \sqrt[\alpha]{-\ln(1-r_i)}$	β = Parámetro de escala. α = Parámetro de forma. γ = Parámetro de localización.
Gamma G_i	$G_i = -\frac{1}{k\lambda} \ln \prod_{j=1}^k (1-r_j)$	$1/\lambda$ = Valor esperado. k = Parámetro de forma.
Lognormal LN_i	$LN_i = e^{N_i}$ donde: $N_i = \left(\sum_{j=1}^k r_j - G \right) \cdot \left(\ln \left(1 + \frac{\sigma^2}{\mu^2} \right) \right)^{1/2} + \left(\ln \frac{\mu^2}{\sqrt{\mu^2 + \sigma^2}} \right)$	μ = Valor esperado. σ^2 = Varianza.
Bernoulli BE_i	$BE_i = \begin{cases} 0 & \text{si } r_i \in (0, 1-p) \\ 1 & \text{si } r_i \in (1-p, 1) \end{cases}$	p = Probabilidad de ocurrencia del evento $x = 1$. $1-p$ = Probabilidad de ocurrencia del evento $x = 0$.
Binomial B_i	$B_i = \sum_{j=1}^N BE_j$	BE_j = Números aleatorios con distribución de Bernoulli. N = Número del evento máximo de la distribución binomial. p = Probabilidad de éxito de la distribución binomial que se involucra al generar los Bernoulli.
Poisson P_i	<i>Inicialización.</i> Hacer $N = 0$, $T = 1$ y generar un aleatorio r_i <i>Paso 1.</i> Calcular $T' = (T)(r_i)$ <i>Paso 2.</i> Si la $T' \geq e^{-\lambda}$, entonces hacer $N = N + 1$, $T = T'$, calcular otro r_i y regresar al paso 1. Si $T' < e^{-\lambda}$, entonces la variable generada esta dada por: $P_i = N$. Para generar la siguiente variable de Poisson, regresar a la fase de inicialización.	λ = Media de la distribución de Poisson. N = Contador. T = Contador.

Fig. 5. Generadores de variables aleatorias (parte 2)

ajusta adecuadamente a una distribución teórica conocida.

La prueba de Kolmogórov-Smirnov se basa en la comparación de la función de distribución empírica (FDE) de los datos con la función de distribución teórica (FDT) correspondiente. La

FDE es una estimación de la FDT a partir de los datos observados, mientras que la FDT se conoce a priori.

La prueba de Kolmogórov-Smirnov calcula la distancia máxima (D) entre la FDE y la FDT, y se utiliza para determinar si los datos se ajustan

adecuadamente a la distribución teórica. Si la distancia máxima es pequeña, entonces los datos se ajustan bien a la distribución teórica. Por el contrario, si la distancia máxima es grande, entonces los datos no se ajustan bien a la distribución teórica.

La prueba de Kolmogórov-Smirnov se utiliza comúnmente en la estadística para evaluar la bondad de ajuste de una distribución teórica a los datos observados. También se utiliza en la validación de modelos estadísticos y de simulación para determinar si los datos simulados se ajustan adecuadamente a la distribución teórica asumida.

Es importante tener en cuenta que la prueba de Kolmogórov-Smirnov asume que los datos son independientes e idénticamente distribuidos (IID) y que la distribución teórica es continua. Si los datos no cumplen con estos supuestos, entonces la prueba puede no ser adecuada y se deben considerar otras pruebas de bondad de ajuste. El procedimiento general de la prueba es:

1. Obtener al menos 30 datos de la variable aleatoria a analizar.
2. Calcular la media y la varianza de los datos
3. Crear un histograma de $m = \sqrt{n}$ intervalos, y obtener la frecuencia observada en cada intervalo O_i .
4. Calcular la probabilidad observada en cada intervalo $PO_i = O_i/n$, esto es, dividir la frecuencia observada O_i entre el número total de datos n .
5. Acumular las probabilidades PO_i para obtener la probabilidad observada hasta el i -ésimo intervalo, POA_i .
6. Establecer de manera explícita la hipótesis nula, para esto se propone una distribución de probabilidad que se ajuste a la forma del histograma.
7. Calcular la probabilidad acumulada para cada intervalo, PEA_i , a partir de la función de probabilidad propuesta.
8. Calcular el estadístico de prueba $c = \max |PEA_i - POA_i|$ para $i=1,2,3,\dots,k,\dots,m$.
9. Definir el nivel de significancia de la prueba α , y determinar el valor crítico de la prueba $D_{\alpha,n}$, como muestra la gráfica.
10. Comparar el estadístico de prueba con el valor crítico. Si el estadístico de prueba es menor

que el valor crítico no se puede rechazar la hipótesis nula.

Prueba de Kolmogórov-Smirnov vs Anderson Darling

Tanto la prueba de Kolmogórov-Smirnov como la prueba de Anderson-Darling son pruebas de bondad de ajuste utilizadas para evaluar si una muestra de datos proviene de una distribución de probabilidad específica.

La principal diferencia entre estas dos pruebas es que la prueba de Anderson-Darling es más sensible a las desviaciones en la cola de la distribución, mientras que la prueba de Kolmogórov-Smirnov es más sensible a las desviaciones en el centro de la distribución.

La prueba de Anderson-Darling utiliza una medida de distancia basada en la función de distribución empírica (FDE), al igual que la prueba de Kolmogórov-Smirnov, pero utiliza pesos diferentes para cada punto de datos, lo que hace que sea más sensible a las desviaciones en la cola de la distribución. La prueba de Kolmogórov-Smirnov, por otro lado, se basa en la distancia máxima entre la FDE y la función de distribución teórica (FDT) y es más adecuada para detectar desviaciones en el centro de la distribución. En resumen, la elección entre la prueba de Kolmogórov-Smirnov y la prueba de Anderson-Darling dependerá de la forma de la distribución de los datos y de la ubicación de las desviaciones en relación con la distribución teórica.

Es recomendable realizar ambas pruebas para tener una evaluación más completa de la bondad de ajuste de la distribución teórica a los datos observados.

Aplicaciones para hacer las pruebas de bondad de ajuste de un conjunto de datos

Stat Fit es una herramienta de ProModel, un software de simulación empresarial utilizado para modelar, analizar y optimizar sistemas complejos. Stat Fit es una función dentro de ProModel que se utiliza para ajustar modelos de distribución de probabilidad a los datos históricos o a los datos de entrada de un modelo de simulación.

Con Stat Fit, los usuarios pueden seleccionar una distribución de probabilidad de una amplia

variedad de opciones, incluyendo distribuciones continuas y discretas, y luego ajustar los parámetros de la distribución a los datos observados.

La herramienta utiliza varios métodos de ajuste, como el método de máxima verosimilitud y el método de los momentos, para encontrar los parámetros óptimos de la distribución.

Una vez que se ha ajustado la distribución, los usuarios pueden utilizarla para generar datos aleatorios que se ajusten a la distribución ajustada.

Estos datos se pueden utilizar para realizar análisis de sensibilidad y simulaciones en ProModel, lo que permite a los usuarios evaluar diferentes escenarios y tomar decisiones informadas. La herramienta Stat Fit de ProModel es una herramienta útil para ajustar modelos de distribución de probabilidad a los datos y generar datos aleatorios que se ajusten a la distribución.

Esto permite a los usuarios modelar sistemas complejos y tomar decisiones informadas basadas en los resultados de la simulación. Esta herramienta utiliza la Prueba de Kolmogórov-Smirnov y la prueba de Anderson Darling.

Generadores de variables aleatorias

Una vez que se ha determinado el tipo de distribución que tienen los datos de entrada, es necesario generar datos aleatorios de entrada con la distribución que tienen los datos, como se presenta en Fig.4 y Fig. 5.

Generadores de variables aleatorias muestra una tabla con los generadores más comunes.

4. Análisis de resultados

La validación y la verificación son procesos críticos en el desarrollo de modelos de simulación porque garantizan que el modelo se ajuste a su propósito y sea confiable para su uso.

Estos procesos también ayudan a identificar errores y deficiencias en el modelo, lo que permite corregirlos y mejorar la precisión y la calidad de los resultados.

La validación se refiere al proceso de asegurarse de que el modelo simula de manera precisa el comportamiento del sistema real. Esto

se logra mediante la comparación de los resultados del modelo con los datos reales, los resultados de otros modelos similares y los resultados esperados de la teoría. La validación es importante porque un modelo no validado puede generar resultados inexactos y llevar a decisiones equivocadas.

La verificación, por otro lado, se refiere al proceso de asegurarse de que el modelo se haya construido correctamente y que sea confiable para su uso. Esto se logra mediante la revisión y la prueba del modelo para detectar y corregir errores, y garantizar que el modelo esté diseñado de manera coherente con los requisitos y especificaciones del usuario. La verificación es importante porque un modelo no verificado puede contener errores y ser inadecuado para su propósito.

La validación y la verificación son procesos críticos para garantizar la precisión y la calidad de los resultados de un modelo de simulación. La validación asegura que el modelo simule de manera precisa el comportamiento del sistema real, mientras que la verificación asegura que el modelo se haya construido correctamente y sea confiable para su uso. Un modelo validado y verificado puede ayudar a tomar decisiones informadas y lograr resultados positivos.

Algunos de los aspectos que es necesario tomar en cuenta en la validación son: recopilar datos reales, definir los criterios de validación para comparar con los datos reales como error absoluto medio o error cuadrático o correlación, ejecutar la simulación y comparar los datos simulados con los datos reales, realizar un análisis de sensibilidad para determinar si al realizar cambios en el modelo los resultados se comportan como en la realidad y finalmente ajustar el modelo de simulación en caso de que sea necesario.

Análisis de sensibilidad

El análisis de sensibilidad ayuda a identificar qué variables o parámetros tienen el mayor impacto en los resultados de la simulación, lo que permite enfocar los esfuerzos en la mejora de la precisión de esos parámetros. También puede ayudar a identificar las incertidumbres y las limitaciones en el modelo de simulación.

Existen diferentes técnicas de análisis de sensibilidad, como el análisis de varianza (ANOVA), el análisis de sensibilidad univariante y multivariante, y la simulación de Monte Carlo. En general, estas técnicas implican la ejecución de la simulación con diferentes valores de los parámetros de entrada y la comparación de los resultados para evaluar cómo cambian los resultados a medida que cambian los parámetros de entrada.

El análisis de sensibilidad puede ser útil para tomar decisiones informadas en áreas como la planificación, el diseño y la toma de decisiones. Por ejemplo, en un modelo de simulación de gestión de inventario, el análisis de sensibilidad puede ayudar a determinar cuál es el nivel óptimo de inventario que minimiza los costos y maximiza los beneficios, y cuál es la sensibilidad de los resultados a los cambios en los costos de los materiales o los plazos de entrega.

Los resultados de una simulación dependen de la distribución de probabilidad de las variables aleatorias. Los intervalos de confianza en una simulación nos permiten evaluar la precisión y la incertidumbre asociadas con nuestras estimaciones y predicciones, lo que puede ser útil para tomar decisiones informadas basadas en los resultados de la simulación.

El intervalo de confianza cuando la variable aleatoria sigue una distribución normal es calculado con la siguiente formula:

$$IC = \left[\bar{x} - \frac{s}{\sqrt{r}} (t_{\alpha/2, r-1}), \bar{x} + \frac{s}{\sqrt{r}} (t_{\alpha/2, r-1}) \right]. \quad (1)$$

En el caso de que la variable aleatoria siga otro tipo de distribución el intervalo de confianza es relativamente más amplio y se calcula como:

$$IC = \left[\bar{x} - \frac{s}{\sqrt{r\alpha}}, \bar{x} + \frac{s}{\sqrt{r\alpha}} \right], \quad (2)$$

en donde:

r = número de réplicas,

α = nivel de rechazo,

$$\bar{x} = \frac{1}{r} \sum_{i=1}^r x_i,$$

$$s = \left[\frac{1}{r-1} \sum_{i=1}^r (x_i - \bar{x})^2 \right]^{1/2}.$$

Para que el resultado de una variable aleatoria llegue al estado estable en una simulación no terminal o de estado estable, es decir, que los

resultados de la simulación se estabilizan y la longitud de las réplicas ya no tiene un impacto significativo en los resultados.

En una simulación, la longitud de las réplicas se refiere a la duración de cada ejecución independiente de la simulación. Una réplica se ejecuta para generar un conjunto de datos de salida que se utiliza para analizar el comportamiento del sistema simulado.

La longitud de las réplicas debe ser determinada con cuidado, ya que puede tener un impacto significativo en los resultados de la simulación. Si las réplicas son demasiado cortas, puede que no se capturen adecuadamente los efectos estocásticos o aleatorios del sistema y los resultados pueden ser poco confiables.

Si las réplicas son demasiado largas, la simulación puede requerir un tiempo de computación excesivo y puede ser difícil de manejar en términos de almacenamiento de datos y análisis.

La longitud óptima de las réplicas depende del sistema que se está simulando, de los objetivos de la simulación y de la precisión deseada en los resultados. En general, se recomienda que la longitud de las réplicas sea lo suficientemente larga como para capturar los efectos estocásticos importantes del sistema, pero lo suficientemente corta como para no requerir un tiempo de computación excesivo.

Una forma común de determinar la longitud óptima de las réplicas es realizar un análisis de sensibilidad para evaluar cómo varían los resultados de la simulación en función de la longitud de las réplicas.

Esto permite determinar el punto en el que los resultados de la simulación se estabilizan y la longitud de las réplicas ya no tiene un impacto significativo en los resultados.

En el caso de que los datos tengan una distribución normal, el tamaño de la corrida de la simulación se calcula con la siguiente formula:

$$n = \left[\frac{\sigma Z_{\alpha/2}}{\epsilon} \right]^2. \quad (3)$$

Cuando la suposición de normalidad en los datos no existe, la longitud de la réplica se calcula así:

$$n = \frac{1}{\alpha} \left[\frac{\sigma}{\epsilon} \right]^2. \quad (4)$$

5. Caso práctico

En una empresa zacatecana de manufactura se realizó un estudio de líneas de espera en el área de inspección. Se tomó el tiempo entre llegadas de las piezas y el tiempo que transcurría durante la inspección de cada una. Considerando una muestra de 50 piezas se tienen los siguientes datos:

Luego de la prueba de bondad de ajuste tenemos los siguientes resultados para el tiempo entre llegadas. Es importante señalar que los datos se redondearon a 4 dígitos después del punto decimal para mostrarlos, sin embargo, en la realidad no se muestran así en ambos casos.

La prueba de bondad de ajuste para el tiempo durante la inspección: Generamos el modelo de simulación, en donde el tiempo entre llegadas tiene una distribución exponencial con una media de 4.52 minutos y el tiempo durante la inspección con una distribución normal con una media de 3.99 minutos y una desviación estándar de 0.617 minutos.

Aunque tiene un 100% la distribución Lognormal, no haremos con la distribución normal con un 99.6% de aceptación.

Aquí se muestran solamente las primeras 10 piezas, pero continua hasta las corridas necesarias según la formula.

Graficamos la columna con la variable de interés, en este caso los minutos que está la pieza en el sistema. Como observamos el valor de $p < 0.05$, por lo anterior, los datos no siguen una distribución normal.

Utilizamos entonces la fórmula de abajo con una desviación estándar de 4.958 y un error de 2 minutos (definido por analista) y una alfa de 0.05.

$$n = (1/0.05) * (4.958/2)^2 = 122.91$$
 corridas para llegar a la estabilidad, sin embargo, en la prueba piloto se realizó con 500 corridas, por lo que se sugiere dejarlo así.

Posteriormente se realizan las réplicas para calcular el intervalo de confianza, se sugiere manejar al menos 30 réplicas para tener un mejor desempeño.

Para efecto de la presentación se hizo un corte. Se calcula el promedio de todos los promedios de las 30 réplicas.

Tiempo entre llegadas (min/pieza)				
3.4304	5.2266	0.8756	3.7216	4.9293
3.8017	3.0991	2.0299	6.2574	9.5237
3.9118	2.9633	0.2235	1.6228	0.0179
1.4472	0.6613	2.6925	0.5423	2.9871
3.3982	0.8759	0.6608	4.3546	5.5448
5.8023	3.7024	7.1132	0.1460	8.2216
0.7991	1.9219	9.2331	0.7438	12.4336
6.6452	8.3873	7.6080	11.6309	8.5193
23.4671	0.8839	10.0814	0.6456	0.1455
0.5215	11.4354	0.2187	4.9872	6.9904

Fig. 6. Tiempo entre llegadas de piezas (min/pieza)

Tiempo de inspección (min/pieza)				
3.7976	4.2998	3.8570	4.8499	3.9797
3.6466	4.0194	3.8323	3.3240	4.5743
3.3788	4.6158	4.2479	3.8906	4.3924
4.2785	3.1710	4.3600	4.1699	3.8693
4.3745	5.6806	4.0262	3.7346	4.2834
4.2897	3.9904	3.7298	4.6117	4.4796
4.7278	4.4618	4.9857	3.9811	4.2737
3.9050	4.1210	4.4449	4.2525	4.0797
4.7730	3.1981	3.3293	4.1189	3.7828
5.5990	4.1802	4.2097	4.3854	4.3049

Fig. 7. Tiempo de inspección de cada pieza (min/pieza)

El valor de $p < 0.05$ por lo que los datos no siguen una distribución normal y calculamos el intervalo de confianza con la fórmula de la Fig. 20.

Límite superior = $12.55 + (3.12/RAIZ(30*0.05/2)) = 16.1526$.

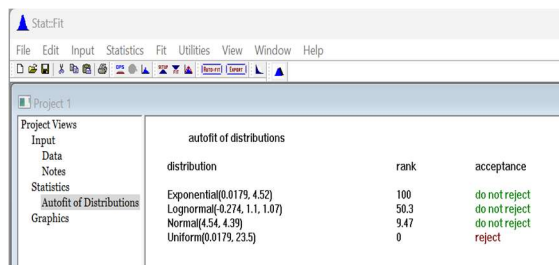
Límite inferior = $12.55 - (3.12/RAIZ(30*0.05/2)) = 8.9473$.

La simulación nos concluye que el promedio de una pieza en el sistema de inspección (línea de espera + tiempo de inspección) es de 12.55 minutos con una desviación estándar de 3.12 minutos. El intervalo de confianza con un 95% de seguridad está entre 8.9473 y 16.1526 minutos.

6. Conclusiones

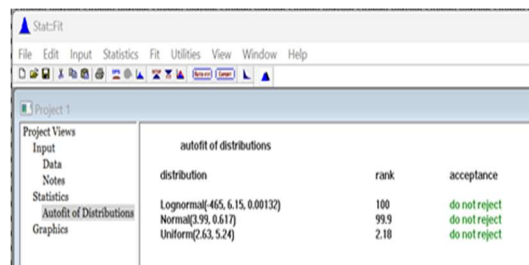
Las conclusiones que se pueden obtener después de una simulación dependen del propósito y del alcance de la simulación en sí.

Algunas de las conclusiones generales que se pueden obtener son las siguientes: Comportamiento del sistema: La simulación permite analizar el comportamiento del sistema en diferentes condiciones y escenarios.



distribution	rank	acceptance
Exponential(0.0179, 4.52)	100	do not reject
Lognormal(0.274, 1.1, 1.07)	50.3	do not reject
Normal(4.54, 4.39)	9.47	do not reject
Uniform(0.0179, 23.5)	0	reject

Fig. 8. Comparativo de las distribuciones analizadas. Prueba de bondad de ajuste con el software Promodel, módulo stat-fit. No se rechaza la hipótesis de que los datos tengan una distribución exponencial con un 100%



distribution	rank	acceptance
Lognormal(465, 6.15, 0.00132)	100	do not reject
Normal(1.99, 0.617)	99.9	do not reject
Uniform(2.63, 5.24)	2.18	do not reject

Fig. 11. Comparativo de las distribuciones analizadas. Prueba de bondad de ajuste con el software Promodel, módulo stat-fit. No se rechaza la hipótesis de que los datos tengan una distribución lognormal con un 100%, ni tampoco se rechaza la hipótesis de que los datos tengan una distribución normal con un 99.9%

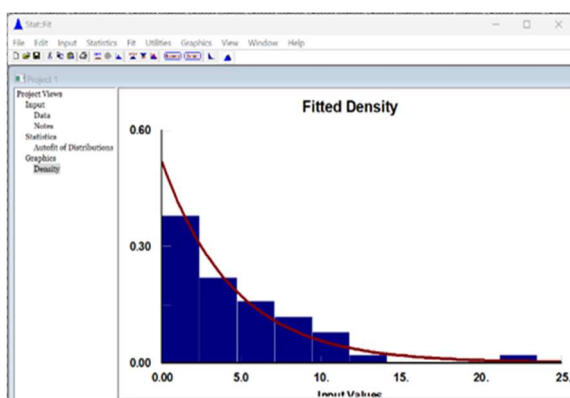


Fig. 9. Histograma de los datos comparada con la distribución exponencial

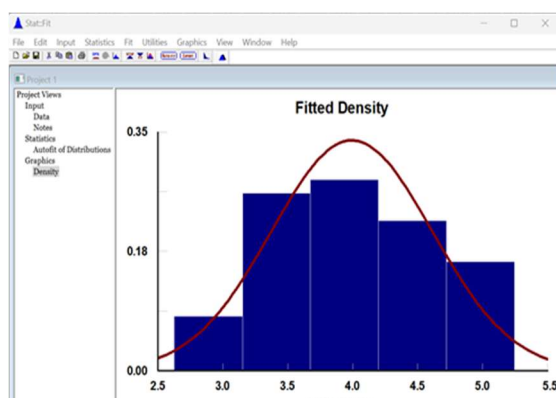
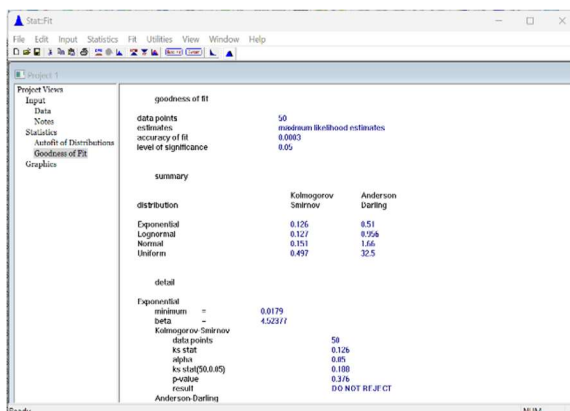
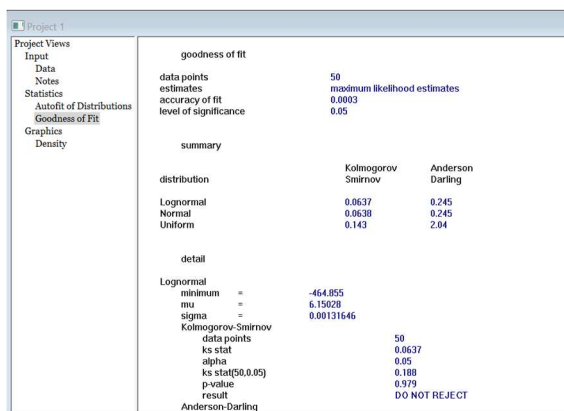


Fig. 12. Histograma de los datos comparada con la distribución normal



distribution	Kolmogorov Smirnov	Anderson Darling
Exponential	0.126	0.51
Lognormal	0.127	0.956
Normal	0.151	1.66
Uniform	0.497	32.5

Fig. 10. Resultado de las pruebas de bondad de ajuste de los 50 datos con un nivel de significancia de 0.05, muestra ambas Kolmogórov-Smirnov y Anderson-Darling



distribution	Kolmogorov Smirnov	Anderson Darling
Lognormal	0.0637	0.245
Normal	0.0638	0.245
Uniform	0.143	2.04

Fig. 13. Resultado de las pruebas de bondad de ajuste de los 50 datos con un nivel de significancia de 0.05, muestra ambas Kolmogórov-Smirnov y Anderson-Darling

							estadísticos	
							promedio	12.906
							desv estándar	8.8270
							máximo	48.6687
							mínimo	2.4774
a. modelo de simulación para 500 piezas con gráfica de estabilidad								
Piezas	número aleatorio	Tiempo entre llegadas (min/pieza)	Minutos en que llega la pieza	Minutos en que inicia la inspección	número aleatorio	Tiempo de inspección (min/pieza)	Minutos en que finaliza la inspección	Minutos que esta la pieza en el sistema
1	0.7400	6.0904	6.0904	6.0904	0.9542	5.0302	11.1216	5.0312
2	0.5904	4.0341	10.1245	11.1236	0.3501	3.7524	14.8759	4.7495
3	0.8713	9.2674	19.3919	19.3919	0.9680	5.1418	24.5337	5.1418
4	0.4803	2.9582	22.3501	24.5337	0.3900	3.0176	28.3512	6.0012
5	0.9361	12.4382	34.7833	34.7833	0.6878	4.2621	39.0754	4.2621
6	0.5251	3.3369	38.1492	39.0754	0.3724	3.7082	42.8645	4.7553
7	0.2722	1.4360	39.5852	42.8645	0.8885	4.6806	47.5451	7.9599
8	0.0997	0.4746	40.0599	47.5451	0.2930	3.6540	51.1991	11.1392
9	0.5407	3.3564	43.5762	51.1991	0.8149	4.6408	55.7421	12.1618
10	0.6450	19.4487	57.5714	57.5714	0.6687	3.6415	61.9544	23.6455

Probability Plot of Minutos que esta la pieza en el

Normal - 95% CI

Percent

Minutos que esta la pieza en el

Mean	9.220
St Dev	4.958
N	500
AD	12.954
P-Value	0.005

Donde:
 s = Desviación Estándar
 r = Número de réplicas.
 α = Valor alfa (Nivel de rechazo)
 ϵ = Error Permitido
 n = Tamaño de cada corrida de simulación

	estadísticos	estadísticos	estadísticos	estadísticos	estadísticos	estadísticos	estadísticos
promedio	14.8853	13.1037	13.0606	14.1141	17.0804	9.2130	9.6575
desv estándar	13.7177	8.7150	9.0058	10.5733	11.3936	5.3496	6.7711
máximo	70.3840	43.7902	45.7914	53.6358	52.9841	32.8437	33.2559
mínimo	2.9122	2.8608	2.9331	2.4866	2.9345	2.2756	2.9398
	Replica 1	Replica 2	Replica 3	Replica 4	Replica 5	Replica 29	Replica 30
							promedio
	3.6883	4.0165	3.5723	3.8394	4.0638	4.5169	3.7168
	3.9296	6.0184	4.4366	4.9786	5.0049	3.3963	4.7268
	4.7646	5.5178	4.4961	6.0456	4.5670	4.2126	5.9973
	4.6589	6.0215	4.6583	5.5842	4.5669	5.1043	4.4365
	5.2890	6.9444	5.2941	5.2942	4.4785	5.7743	5.4560
	5.6725	7.8862	4.9006	5.1862	4.3629	5.9628	5.6126
	6.0373	9.0056	5.2732	5.3438	4.4482	5.6688	5.8701
	6.6630	9.6404	5.5399	5.2568	4.2774	5.4788	5.5513
	7.4161	9.0013	5.8921	5.3536	4.6033	5.3155	5.4614

En resumen, las conclusiones que se pueden obtener después de una simulación son importantes para entender el comportamiento del sistema, identificar puntos críticos, evaluar alternativas y tomar decisiones informadas basadas en evidencia.

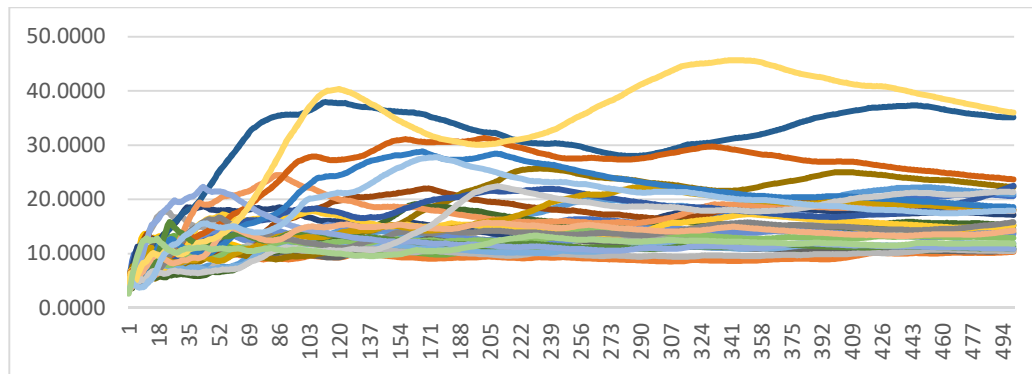


Fig. 18. Gráfica de probabilidad de la variable de interés "minutos que esta la pieza en el sistema" de las 30 réplicas

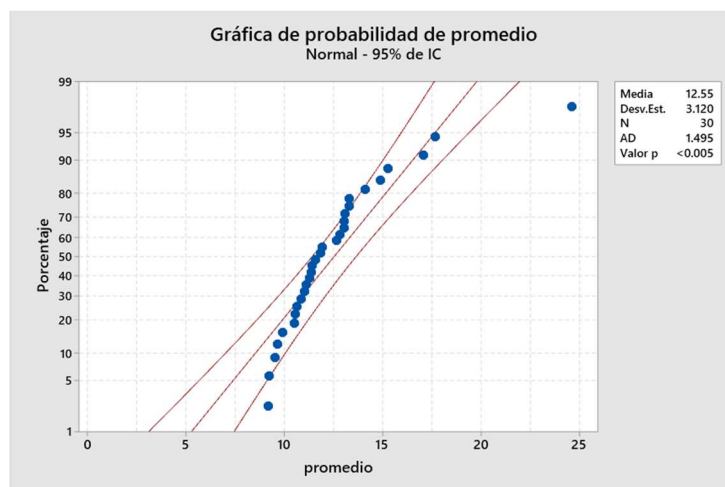


Fig. 19. Gráfica de probabilidad los 30 promedios de las réplicas referentes a la variable de interés "minutos que la pieza está en el sistema"

SIMULACIONES TERMINALES

Intervalos de confianza bajo el supuesto de normalidad:

$$IC = \bar{x} \pm \frac{s}{\sqrt{r}} (t_{\alpha/2, r-1})$$

Intervalos de confianza bajo el supuesto de NO normalidad:

$$IC = \bar{x} \pm \frac{s}{\sqrt{\frac{r\alpha}{2}}}$$

Fig. 20. Intervalo de confianza después de las réplicas.

Referencias

1. **Law, A. (2015).** Simulation, modeling and analysis. New York, 5ed, McGraw-Hill Education.
2. **García-Dunna, E., García-Reyes, H., Cárdenas-Barrón, L. E. (2013).** Simulación y

análisis de sistemas con ProModel. México, 2ed, Pearson Educación.

*Article received on 21/02/2023; accepted on 17/05/2023.
Corresponding author is María del Consuelo Argüelles Arellano.*