



Educación matemática

ISSN: 0187-8298

ISSN: 2448-8089

Editorial Santillana

Inzunza Cazares, Santiago; Islas Anguiano, Eldegar
Análisis de una trayectoria de aprendizaje para desarrollar
razonamiento sobre muestras, variabilidad y distribuciones muestrales
Educación matemática, vol. 31, núm. 3, 2019, pp. 203-230
Editorial Santillana

DOI: <https://doi.org/10.24844/EM3103.08>

Disponible en: <https://www.redalyc.org/articulo.oa?id=40576153009>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

UAEH
redalyc.org

Sistema de Información Científica Redalyc
Red de Revistas Científicas de América Latina y el Caribe, España y Portugal
Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso
abierto

Análisis de una trayectoria de aprendizaje para desarrollar razonamiento sobre muestras, variabilidad y distribuciones muestrales

Analysis of a learning trajectory to develop reasoning about samples, variability and sampling distributions

Santiago Inzunza Cazares¹
Eldegar Islas Anguiano²

Resumen: En el presente artículo analizamos una trayectoria de aprendizaje mediada por el uso intensivo de tecnología, diseñada con el propósito de desarrollar razonamiento adecuado sobre muestras, variabilidad y distribuciones muestrales en estudiantes de ciencias sociales. Los resultados manifiestan que un enfoque informal basado en simulación del muestreo puede ayudar a los estudiantes a desarrollar un razonamiento correcto sobre las distribuciones muestrales y construir la base para la comprensión de la inferencia estadística. La idea de tamaño de muestra y su efecto en la variabilidad de una distribución muestral y precisión de una estimación de características poblacionales, ha sido comprendida en buena medida por los estudiantes, pero la confusión entre población, muestra y distribución muestral nos confirma la complejidad de estos conceptos. En cuanto a la interpretación de la información que contiene una distribución muestral, la mayor parte de los estudiantes han identificado correctamente la inusualidad de una muestra.

Fecha de recepción: 11 de febrero de 2019. **Fecha de aceptación:** 3 de julio de 2019.

¹ Universidad Autónoma de Sinaloa, sinzunza@uas.edu.mx, orcid.org/0000-0003-4014-6031.

² Universidad Autónoma de Sinaloa, eldegarislas@uas.edu.mx, orcid.org/0000-0002-8194-8922.

Palabras clave: *Análisis de datos, ambientes virtuales de aprendizaje, enseñanza universitaria, estadística.*

Abstract: In the present article we analyze a learning trajectory mediated by the intensive use of technology, designed with the purpose of developing adequate reasoning about samples, variability and sampling distributions in social science students. The results show that an informal approach based on sampling simulation can help students to develop correct reasoning about sampling distributions and build the basis for understanding statistical inference. The idea of sample size and its effect on the variability of a sampling distribution and precision of an estimation of population characteristics has been largely understood by the students, but the confusion between population, sample and sampling distribution confirms the complexity of these concepts. Regarding the interpretation of the information that contains a sampling distribution, most of the students have correctly identified the unusualness of a sample.

Keywords: *Data analysis, virtual learning environment, university teaching, statistics.*

1. EL PROBLEMA

El razonamiento sobre muestreo ha adquirido especial importancia para la educación estadística en los años recientes, como consecuencia del incremento en el uso de datos provenientes de muestras y experimentos aleatorizados en investigaciones de interés en las profesiones, la ciencia y la vida cotidiana. Un ejemplo muy recurrente de uso del muestreo son los resultados de encuestas de opinión que publican frecuentemente los medios de comunicación, en las cuales se reportan estimaciones sobre parámetros de una población (por ejemplo: proporciones y promedios). Otro caso representativo son los experimentos aleatorizados que se utilizan para determinar si existe diferencia entre dos o más tratamientos o condiciones, en campos tan diversos como la salud, nutrición, ingeniería, agronomía, entre otros.

El razonamiento inferencial es inductivo e implica ir más allá de los datos de una muestra o experimento para extraer conclusiones sobre un universo o población que no ha sido explorado en su totalidad (Inzunza, 2013); de esta

manera, y como consecuencia del azar y la variabilidad muestral, los resultados contienen inevitablemente márgenes de incertidumbre. Sin embargo, como señala Pfannkuch (2005), la integración de la estadística con la probabilidad presentó enormes desafíos en sus orígenes, lo que explica en parte las dificultades que entrañan para los estudiantes el aprendizaje y la aplicación de la inferencia estadística, incluso para muchos profesores e investigadores (Belia *et al.*, 2003).

Un concepto fundamental para comprender y desarrollar el razonamiento sobre la inferencia estadística, son las distribuciones muestrales. La distribución muestral de un estadístico (por ejemplo, la media o la proporción) representa el valor que puede tomar el estadístico en cada una de las muestras de un tamaño, dado que son posibles de extraer de una población; de tal forma, es importante que los estudiantes reconozcan que la muestra que están analizando, es sólo una del conjunto potencialmente infinito de muestras que podrían ser extraídas de una población o experimento, y que para hacer una inferencia, la distribución de todas las muestras para un estadístico en particular requiere ser conocida (Lipson, 2000).

De esta manera, la distribución muestral de un estadístico permite calcular y visualizar la posición que ocupa una muestra particular en la distribución, con lo cual se puede estimar su probabilidad de ocurrencia y generar así el precursor informal del p-valor, concepto que determina la significatividad del resultado de una prueba de hipótesis sobre un parámetro; también permite identificar la variabilidad muestral que conduce a la idea de error de muestreo y margen de error, componente fundamental de la estimación por intervalos de confianza (Garfield y Ben-Zvi, 2008).

Desde nuestra perspectiva, el enfoque que aún prevalece en la enseñanza de las distribuciones muestrales, el cual está basado en un enfoque abstracto y deductivo de la probabilidad (Inzunza, 2015), conduce a una distribución teórica de probabilidad en la que no fácilmente son visibles los conceptos que intervienen en las distribuciones muestrales, además este enfoque requiere de un lenguaje matemático que está fuera del alcance de muchos estudiantes universitarios, y mas aún de los estudiantes de bachillerato, nivel en el que en muchos currículos se contempla la enseñanza de estos tópicos.

En esta dirección, en los últimos años se han planteado diversas propuestas de enseñanza (por ejemplo: Cobb, 2007; Wild, Pfannkuch y Reagan, 2011; Batanero y Borovnick, 2016,) que se engloban en un enfoque denominado informal. Estas propuestas tienen como elemento integrador el uso de herramientas tecnológicas con amplio potencial de visualización, interactividad y

dinamismo con capacidad para simulación del muestreo y aleatorización, que permiten construir empíricamente una distribución muestral para una gran cantidad de muestras. Este enfoque reduce el nivel de abstracción de los conceptos involucrados y se relaciona más directamente con el proceso real de selección de muestras de una población que se utiliza en una investigación estadística.

En este contexto, nos hemos planteado como objetivo: diseñar un trayectoria de aprendizaje y un experimento de enseñanza para desarrollar el razonamiento sobre las distribuciones muestrales y su relación con la comprensión de la inferencia estadística con estudiantes universitarios con pocos antecedentes matemáticos, como es el caso de los estudiantes de ciencias sociales y humanidades. La trayectoria se caracteriza por la adopción de un enfoque informal mediado por el uso intensivo de tecnologías digitales para la simulación del muestreo y contextos de datos reales. Los resultados del presente artículo corresponden a la primera parte de la investigación relativa a las distribuciones muestrales. Pretendemos caracterizar el razonamiento que los estudiantes desarrollaron sobre el muestreo y las distribuciones muestrales, después de completar la trayectoria de aprendizaje, así como definir un modelo explicativo de las principales dificultades en el razonamiento y comprensión estocástica del muestreo.

2. INVESTIGACIONES PREVIAS

Entre las primeras investigaciones sobre comprensión del muestreo, destaca la realizada por los psicólogos Kahneman y Tversky (1972), quienes reportan que cuando las personas evalúan la probabilidad de que una muestra ocurra, frecuentemente, basan sus juicios en la similitud de características que la muestra tiene con la población, aunque ésta sea pequeña. En una continuación de su trabajo, Kahneman y Tversky (1982) reportan que las personas, cuando hacen juicios a partir de muestras aleatorias son proclives a ver las muestras desde una perspectiva singular (muestras individuales) en lugar de una perspectiva de distribución (colección de muestras), se enfocan en las causas que produce un resultado particular y evalúan la probabilidad basándose en la información disponible del caso que tienen a la mano.

En relación con lo anterior, Saldanha y Thompson (2002) reportan dos concepciones de muestra y muestreo que surgieron en el contexto de un experimento de enseñanza con estudiantes de bachillerato. Una concepción, caracterizada como *aditiva*, donde se concibe a una muestra –no obstante la

repetición del muestreo– como una imagen aislada y una versión a pequeña escala de la población. La otra concepción, es caracterizada como *multiplicativa* e implica una red de imágenes interrelacionadas en la cual los resultados muestrales son vistos como un subconjunto de la población, que en la medida que el muestreo es repetido se ven como casos similares que se acumulan para formar una distribución. Desde esta perspectiva es posible determinar qué tan usual o inusual puede ser un resultado muestral, o establecer intervalos de valores muestrales entre los cuales podría estar un parámetro de la población. Saldanha y Thompson señalan que una concepción multiplicativa del muestreo es deseable para desarrollar un esquema adecuado sobre la inferencia estadística, sin embargo no es una actividad trivial, pues resultó ser una tarea compleja para la mayoría de los estudiantes que participaron en la investigación.

En el caso específico de distribuciones muestrales, un referente importante es el estudio realizado por delMas, Garfield y Chance (2004), el cual estuvo centrado en la comprensión de las propiedades de las distribuciones muestrales en relación con el teorema del límite central; para ello desarrollaron materiales de enseñanza y un programa de computadora denominado *Sampling Distributions*. La investigación se desarrolló en varias versiones, en cada una de las cuales se iban modificando las actividades o el software para tratar de mejorar la comprensión de los estudiantes. Para evaluación se diseñaron preguntas basadas en gráficas. En la segunda etapa se identificaron diferentes tipos de razonamiento en los sujetos de estudio:

- *Razonamiento correcto*. Los estudiantes seleccionaron los histogramas correctos de las distribuciones muestrales (para muestras pequeñas y muestras grandes) cuando los cuestionaron sobre la forma y la variabilidad muestral.
- *Buen razonamiento*. Los estudiantes eligieron un histograma que se parecía a la distribución normal para un tamaño de muestra grande y con menor variabilidad que el histograma para el tamaño de muestra pequeño, pero el histograma elegido para una muestra de tamaño pequeño se parecía más a la población que a la distribución normal.
- *Razonamiento de grande a pequeño*. Los estudiantes eligieron un histograma con menor variabilidad para el tamaño de muestra grande, pero eligieron un histograma con variabilidad similar a la población para , o bien, ambos histogramas tenían la forma de la población.
- *Razonamiento de pequeño a grande*. Los estudiantes eligieron un histograma con mayor variabilidad para el tamaño de muestra grande.

Una tercera versión no mejoró los resultados anteriores, debido a que las actividades y el software fueron más complejos y demandaban más atención de los estudiantes. La conclusión del estudio es que una presentación sencilla y clara no conduce necesariamente a una sólida comprensión de las distribuciones muestrales, además, mientras un software puede proporcionar los medios para experiencias ricas en el salón de clase, las simulaciones computacionales por sí solas no garantizan el cambio conceptual.

Otro referente importante es el estudio de Lipson (2000), quien analizó el papel que juegan las distribuciones muestrales en la comprensión de la inferencia estadística, apoyándose en actividades basadas en simulación computacional y analizando los esquemas que los estudiantes desarrollaron a lo largo de las actividades de enseñanza. Los sujetos de estudio fueron estudiantes de nivel de posgrado que tomaban un curso introductorio de estadística y sus antecedentes matemáticos eran pocos o prácticamente nulos. La evaluación de los esquemas de los estudiantes se llevó a cabo por medio de mapas conceptuales, comparando en cada una de las etapas del experimento de enseñanza los mapas contruidos por ellos con mapas expertos que contenían los diferentes conceptos y relaciones de un esquema adecuado.

La evaluación de los mapas de los estudiantes se hizo con base en el número de proposiciones clave identificadas para cada mapa experto, además al final del experimento se evaluó el éxito de los estudiantes en términos de comprensión procedimental y conceptual. Los estudiantes comprendieron varias de las proposiciones claves incluidas en el mapa experto, pero no lograron claridad en varias de las proposiciones fundamentales para la construcción de un esquema adecuado sobre distribuciones muestrales. Un análisis cualitativo de los mapas contruidos por los estudiantes llevó a la categorización de sus sujetos de estudio en tres grupos:

- Grupo 1: Los estudiantes mostraron evidencia del desarrollo del concepto de distribución muestral y lo integraron apropiadamente a sus esquemas para inferencia estadística.
- Grupo 2: Los estudiantes mostraron evidencia del desarrollo del concepto de distribución muestral pero no lo relacionaron a su esquema para inferencia.
- Grupo 3: Los estudiantes en ninguna etapa mostraron evidencia del desarrollo del concepto de distribución muestral.

Inzunza (2009) reporta una investigación sobre los significados que un grupo de estudiantes universitarios construyeron sobre distribuciones muestrales (medias y proporciones) en un ambiente de simulación computacional como el que proporciona el software Fathom. Los principales conceptos abordados en el estudio fueron la variabilidad muestral, el efecto del tamaño de muestra en el comportamiento de las distribuciones muestrales y en las probabilidades de resultados muestrales. Se reporta que el proceso de simulación no estuvo exento de dificultades y los estudiantes exhibieron algunas confusiones en algunas etapas del proceso, las cuales estuvieron relacionadas principalmente con el uso de las representaciones simbólicas del software. En cuanto a las confusiones entre conceptos importantes destacan: la confusión entre distribución de una muestra con una distribución muestral y la tendencia a ver a las muestras de manera aislada en lugar de verlas como distribuciones. Estas dificultades son persistentes, pues también han sido reportadas por Saldanha y Thompson (2002).

En el cuestionario posterior y las entrevistas se observó que algunos estudiantes tuvieron cierta evolución en los significados, aunque de manera diferenciada, llegando a obtener muy buenos resultados en algunos ítems, mientras que en otros, hubo un incremento apenas notable en las respuestas correctas y sus argumentaciones. Entre las ventajas observadas al resolver problemas mediante simulación se pueden mencionar las siguientes: Los resultados sugieren que es importante que antes de que los estudiantes se inicien en el estudio de la inferencia estadística, tengan experiencias de enseñanza donde exploren el concepto de distribuciones muestrales y sus propiedades, el teorema del límite central y el efecto del tamaño de muestra en la probabilidad de un determinado resultado muestral; un ambiente de estadística dinámica, apoyado en representaciones múltiples, parece ser apropiado para ello.

En un estudio posterior y utilizando el software Geogebra (Inzunza, 2015), analizó el razonamiento e imágenes que estudiantes universitarios de ciencias sociales construyeron sobre conceptos de población, muestra, variabilidad muestral, efecto del tamaño de muestra y error de muestreo. Los resultados muestran que algunos estudiantes lograron identificar la influencia del tamaño de muestra en la variabilidad muestral, el error muestral y en forma de las distribuciones muestrales, y sobre todo concebir al muestreo como un proceso aleatorio y repetible de una población, que aunque produce resultados que varían, éstos son en la mayoría de los casos cercanos a la media poblacional. Sin embargo, mostraron confusión al momento de identificar los datos de una muestra y una población.

Alvarado, *et al* (2013), en un estudio sobre distribuciones muestrales y sus implicaciones en los intervalos de confianza advierte sobre la complejidad del tema al analizar las respuestas de estudiantes universitarios a dos problemas abiertos y un cuestionario después de varias sesiones de enseñanza. Una proporción importante de los estudiantes comprendieron la distribución de la media muestral en poblaciones normales con sus respectivos parámetros. El procedimiento algebraico de estandarización no presentó mayores dificultades en los estudiantes. Se destaca además la apropiación de las propiedades de la media y varianza de la distribución de la media muestral en poblaciones normales. No obstante, hubo dificultades para determinar los errores de estimación en poblaciones normales. Respecto a la comprensión del teorema central del límite, se observaron dificultades, en la aplicación de propiedades de la media y varianza del estimador y en establecer sus argumentaciones.

Begué, Batanero y Gea (2018), analizan la comprensión intuitiva de la relación entre la proporción en una población y el valor esperado de la proporción muestral y de la variabilidad de dicha proporción en función del tamaño de la muestra – conceptos muy ligados a las distribuciones muestrales– con estudiantes españoles de segundo y cuarto curso de la educación secundaria obligatoria. Los resultados indican una buena percepción del valor esperado de la proporción en las situaciones de muestreo propuestas, pero menos de la cuarta parte de los estudiantes parecen percibir correctamente el efecto del tamaño de la muestra sobre la variabilidad muestral. La mayoría de los estudiantes, en especial los de cuarto curso, alcanzan el nivel proporcional de razonamiento sobre el muestreo, pero pocos llegan al nivel distribucional de razonamiento sobre muestreo.

3. MARCO CONCEPTUAL

El muestreo y las distribuciones muestrales

El análisis conceptual de las distribuciones muestrales desde una perspectiva informal ubica al muestreo y la variabilidad como conceptos centrales que se deben tener en cuenta en el diseño de tareas de aprendizaje. Dierdorff, *et al*. (2016, p. 39) identifican como parte de la red conceptual que subyace al muestreo, las ideas de tamaño de muestra, proceso aleatorio, distribución, intervalo intuitivo de variabilidad y relación entre muestra y población. En este sentido, los estudiantes requieren:

- Comprender que el incremento en el tamaño de muestra conduce a estimaciones más precisas de la probabilidad y de características poblacionales.
- Comprender que las mediciones repetidas de un mismo fenómeno, es un proceso aleatorio que conduce a resultados distintos y, que las inferencias están influenciadas por la muestra seleccionada.
- Comprender que la extracción de muchas muestras permite formar una colección o distribución que puede ser descrita por medidas de centro y variabilidad.
- Tener una idea razonable de variabilidad alrededor de un valor esperado en un proceso aleatorio, tal como un intervalo de confianza.
- Comprender que como consecuencia de la variabilidad puede proporcionar una imagen distorsionada de la población.

En el mismo análisis de conceptos que rodean al muestreo, Harradine, Batanero y Rossman (2011), señalan que para su comprensión se requiere tener claridad sobre tres tipos de distribuciones que intervienen en la construcción conceptual de una distribución muestral:

1. La *distribución de probabilidad* que modela los valores de una variable aleatoria y depende de algún parámetro de la población.
2. La *distribución los valores de una variable en una muestra aleatoria*, sobre la cual se calculan estadísticos (por ejemplo: media y desviación estándar) que pueden ser usados para inferir sobre parámetros de la población.
3. La *distribución muestral del estadístico*, que modela todos los valores que puede tomar un estadístico en el conjunto de las muestras posibles de un tamaño dado que pueden ser seleccionadas de una población.

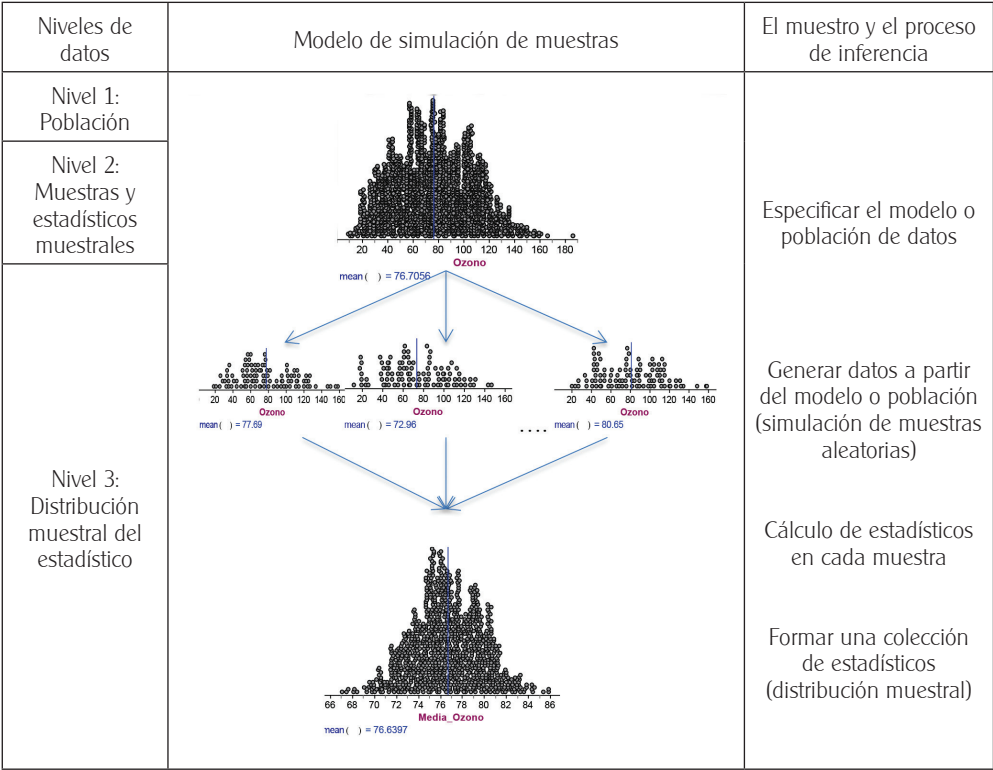
En concordancia con lo anterior, en el proceso hipotético de selección repetida de muestras de una población que caracteriza a un enfoque informal a las distribuciones muestrales, se pueden identificar tres niveles o ámbitos que guían el proceso de simulación, ya sea física o con computadora:

1. Definir el modelo o distribución de la población.
2. Seleccionar una muestra de la población y calcular el estadístico de interés en forma repetida para muchas muestras de la población.

- 3. Registrar los valores de los estadísticos calculados para generar una colección y representarlos gráficamente para construir la distribución muestral.

Saldanha y Thompson (2002) identifican los pasos anteriores como un proceso de tres niveles que conforman un esquema centrado en imágenes repetidas del proceso de muestreo de la población. Los autores reportan que cuando los estudiantes pueden visualizar el proceso de simulación a través del este esquema, pueden lograr una comprensión del proceso y lógica de la inferencia. Un modelo visual de dicho esquema es proporcionado por Lane-Getaz (2006) se muestra en la figura 1.

Figura 1. Modelo visual del proceso de generación empírica de una distribución muestral proporcionado por Lane-Getaz (2006) adaptado al software Fathom.



El primer nivel lo representa la población, el segundo nivel las muestras y los estadísticos calculados en cada una de ellas, y el tercer nivel es representado por la distribución muestral del estadístico. En el último nivel, una muestra cualquiera puede ser comparada con el resto de la distribución para determinar si es una muestra inusual o es una muestra que tiene mucha probabilidad de ocurrir, entre otro tipo de información útil para una inferencia.

Concepción estocástica del muestreo

En el contexto de la probabilidad, Liu y Thompson (2007, p. 122) caracterizan un concepción estocástica de la siguiente manera: “una persona con una concepción estocástica o aleatoria de un evento concibe un resultado observado como una expresión de un proceso subyacente repetible, el cual después de una gran cantidad de repeticiones producirá una distribución estable de resultados”. Pfannkuch *et al.* (2012) extienden dicha concepción al muestreo y la inferencia estadística. Una concepción estocástica del muestreo significa:

1. Concebir el muestreo como un proceso aleatorio. Esto es, seleccionar una muestra de la población, registrar el dato de cada elemento de la muestra y calcular el estadístico en cuestión para estimar el correspondiente parámetro de la población.
2. Imaginar muestras de un mismo tamaño tomadas repetidamente y registrar el valor del estadístico en cada una de ellas.
3. Comprender que este proceso producirá una colección de resultados que serán en su mayoría diferentes del parámetro poblacional que deseamos estimar (la distribución muestral).
4. Comprender que debido al proceso de selección aleatoria hay variabilidad en los resultados, pero en una gran cantidad de repeticiones la distribución de los resultados llegará a ser estable y centrada en el verdadero valor del parámetro.

La tecnología y la simulación computacional del muestreo

La interactividad y la multiplicidad de representaciones visuales dinámicas de las que disponen muchas herramientas de software que existen en la actualidad,

así como el poder de simulación para extraer una gran cantidad de muestras de una población casi de forma simultánea, pueden ayudar, siempre que se combinen con una estrategia pedagógica adecuada, a que los estudiantes accedan a las grandes ideas de la inferencia estadística. De acuerdo con Wild, Pfannkuch y Reagan (2011), la tecnología proporciona excitantes posibilidades para cambiar el panorama de la educación estadística en las escuelas, de forma que podrían hacerla irreconocible.

La simulación es una actividad mediante la cual se pueden extraer conclusiones acerca del comportamiento de un sistema dado, estudiando el comportamiento de un modelo cuyas relaciones de causa y efecto son las mismas (o similares) a las del sistema original (Gottfried, 1984). En nuestro caso, nos centraremos en un tipo particular de simulación, como es la simulación aleatoria, en la cual el sistema consiste de situaciones que contienen elementos de incertidumbre (situaciones aleatorias). Estamos interesados en situaciones aleatorias que pueden ser expresadas a través de un modelo matemático, que puede ser codificado en lenguaje de computadora para ser experimentado por medio de ella; en tal caso, estaremos hablando de una simulación computacional. Este tipo de simulación (computacional y aleatoria) se fundamenta en el comando *random* que genera números en forma aleatoria de acuerdo con una distribución de probabilidad previamente establecida. En este sentido concebimos a la simulación de acuerdo con la metáfora de la caja de herramientas señalada por Sánchez y Yáñez (2003), donde para construir modelos relacionados con situaciones aleatorias se ligan situaciones y problemas con comandos de software a través de modelos.

4. MÉTODO

Los sujetos de estudio y el escenario de investigación

Los sujetos de estudio fueron 16 estudiantes de la carrera de estudios internacionales (19-20 años) de la Universidad Autónoma de Sinaloa que estaban tomando el curso de probabilidad y estadística. Sus antecedentes en el tema consistían de estadística descriptiva y probabilidad básica. El tema de distribuciones muestrales es parte del programa de estudios y decidimos abordarlo desde una perspectiva informal, apoyándonos en el software de análisis de datos Fathom. El estudio se llevó a cabo en un aula en la que los estudiantes

individualmente tuvieron acceso a una computadora con acceso a internet en la que se instaló el software para responder las actividades planteadas.

Las actividades e instrumentos de recolección

La trayectoria de aprendizaje esta formada por cuatro actividades y fue diseñada con el propósito que los estudiantes desarrollen un razonamiento adecuado y comprensión estocástica de los conceptos involucrados en las distribuciones muestrales. Cada actividad iniciaba con la lectura de algún artículo, video o resultados de una encuesta para introducir a los estudiantes a la temática en cuestión. Los instrumentos de recolección fueron hojas de trabajo para cada actividad con instrucciones precisas y con espacio para sus argumentaciones, los archivos del software fueron recolectados en cada actividad y un cuestionario de tres ítems con varios incisos aplicado al final de la trayectoria.

En todas las actividades se consideraron poblaciones con parámetros conocidos y se seleccionaron muestras de tamaño creciente repetidamente para formar una distribución empírica de al menos 1000 muestras (ver cuadro 1), siguiendo el modelo de simulación descrito en la figura 1. En las hojas de trabajo se planteaban preguntas orientadoras para que los estudiantes centraran su atención en los conceptos de interés. Por ejemplo en la actividad 1 las preguntas fueron: ¿qué sucede con los estadísticos (media de ozono y proporción de días con buena calidad del aire) cuando se selecciona una y otra muestra?, ¿qué efecto tiene incrementar el tamaño de la muestra en la variabilidad de los estadísticos?, ¿qué efecto tiene incrementar el tamaño de la muestra en la estimación de los parámetros y en el error de muestreo?

Cuadro 1: Actividad, contexto, variables y tamaños de muestra

Actividad	Contexto y población	Variable, parámetros y estadísticos	Tamaños de muestra
1	Contaminación en la ciudad de México. (Población con datos de una variable).	Variable: ozono Media de ozono. Proporción de días con buena calidad de aire	50, 100, 200
2	Encuesta de opinión sobre muro entre EU y México (La población se modela en la computadora).	Variable: opinión sobre el muro entre EU y México. Proporción de personas a favor del muro.	100, 500, 1000
3	Encuesta de opinión sobre muro entre EU y México. (La población se modela en la computadora).	Variable: opinión sobre el muro entre EU y México. Proporción de personas a favor del muro.	500, 800, 1000
4	Contaminación en la ciudad de México. (Población con datos de una variable).	Variable: dióxido de azufre. Media de dióxido de azufre	500, 1000

La primera y cuarta actividad tenían como contexto los niveles máximos diarios de contaminación en la ciudad de México (2012-2016), los cuales constituyen la población objetivo con parámetros medibles: *media diaria de ozono*, *media diaria de dióxido de azufre* y *proporción de días con buena calidad del aire* (ver tabla 1). La segunda y tercera actividad tenían como contexto una encuesta realizada por Gallup sobre la opinión de los ciudadanos de EU acerca de la construcción del muro en la frontera con México. Se requería en este caso construir un modelo simbólico con el software, teniendo como información la *proporción de ciudadanos que están a favor del muro*, derivada de una pregunta de la encuesta.

Tabla 1: Niveles máximos de contaminación diarios en zona centro de la Ciudad de México (2012-2016).

Contaminantes 2012-2016 Centro DF

	Fecha	Ozono	Dioxido_Azufre	Dioxido_Nitrogeno	Carbono	Particulas	Calidad
1793	11/27/2016	101	9	34	12	102	Mala
1794	11/28/2016	37	6	34	23	102	Buena
1795	11/29/2016	47	4	39	30	81	Buena
1796	11/30/2016	73	4	41	18	83	Regular
1797	01/12/2016	67	4	31	18	97	Regular

Análisis de los resultados

Los resultados que se muestran en este artículo provienen de un cuestionario de cuatro ítems que fue aplicado al final de la trayectoria de aprendizaje. Además de reportar la frecuencia de respuestas correctas e incorrectas en cada ítem, analizamos y clasificamos las argumentaciones de acuerdo su nivel de complejidad en un modelo jerárquico que hemos construido para las distribuciones muestrales basado en el modelo SOLO (Biggs y Collis, 1982).

1. Preestructural

Los estudiantes tienen ideas muy inconexas o superficiales sobre los conceptos que intervienen en una distribución muestral (población, muestra, centro, forma, variabilidad muestral, tamaño de muestra).

2. Uniestructural

Los estudiantes relacionan correctamente alguna de las propiedades de la distribución muestral como centro, variabilidad y forma.

3. Multiestructural

Los estudiantes relacionan correctamente varias propiedades de la distribución muestral, como centro.

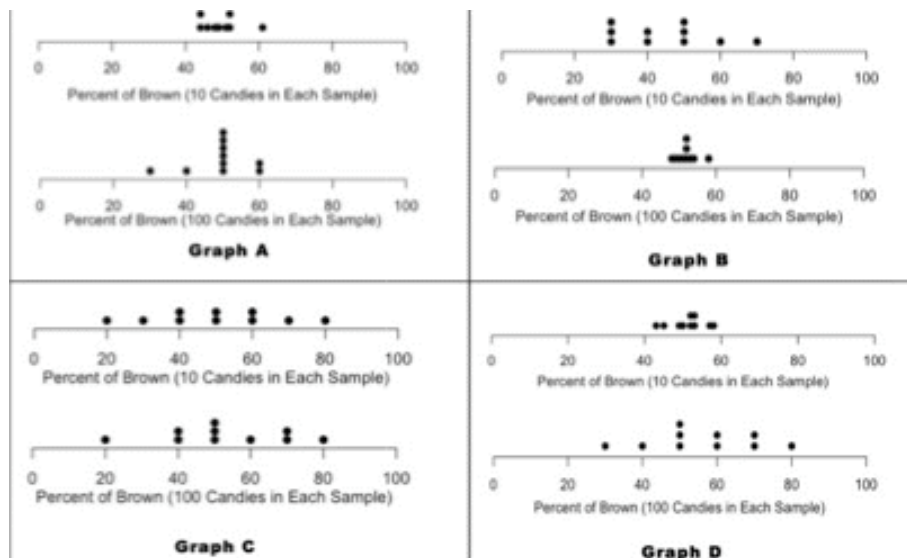
4. Relacional

Los estudiantes comprenden varias propiedades de las distribuciones muestrales y las relacionan correctamente.

4. RESULTADOS Y DISCUSIÓN

ÍTEM 1

Imagina que tienes un barril que contiene miles de dulces de diferentes colores. Se sabe que el fabricante produce 50% de dulces color café. Cada uno de 10 estudiantes selecciona una muestra aleatoria de 10 dulces y registra el porcentaje de dulces color café que aparece en cada muestra. Otros 10 estudiantes seleccionan una muestra aleatoria de 100 dulces y registran el porcentaje de dulces color café en cada muestra. ¿Cuál de las siguientes pares de gráficas representa las distribuciones mas plausibles para el porcentaje de dulces color café obtenidos en las muestra de cada grupo de 10 estudiantes?



El ítem fue tomado de la prueba AIRS (Assessment Inferential Reasoning in Statistics) desarrollada por Park (2012) y tiene el propósito de evaluar la comprensión sobre el efecto que el tamaño de muestra tiene en la variabilidad de la distribución muestral de un estadístico (en este caso de la proporción) y establecer relación entre los resultados muestrales con la población de la que provienen, aspectos relevantes que han mostrado ser difíciles para muchos

estudiantes, como lo reportan Chance, delMas y Garfield (2004) e Inzunza, 2009). La variable aleatoria involucrada es la proporción de dulces color café. El parámetro de la población es conocido, $P=0.50$. Un total de 14 estudiantes eligieron la gráfica correcta (gráfica B), no obstante algunos no argumentan adecuadamente su elección por lo que son clasificados con un razonamiento de nivel preestructural. El análisis de respuestas nos ha permitido clasificarlas de la siguiente manera:

Nivel	Cantidad
Preestructural	4
Uniestructural	3
Multiestructural	9
Total	16

Ejemplos de respuestas en cada nivel:

- Nivel preestructural:

En él se ubican los estudiantes que eligieron la gráfica incorrecta y otros que eligieron la correcta pero sus argumentos no involucran los conceptos y relaciones adecuadas.

Porque es la gráfica más cercana a la realidad (Kasandra).

Porque me parece que es la que da la información solicitada (Alina).

- Nivel uniestructural:

Se ubican los estudiantes que atendieron alguna propiedad de las distribuciones muestrales, como el centro o la variabilidad.

Porque la tendencia que muestran los gráficos tiende a agrupar los datos en la media, la información muestra que el 50% de los dulces producidos deberán ser de color café, la distribución muestral del gráfico B afirma esa condición (Nicole).

- Nivel multiestructural:

En este nivel se ubican los estudiantes que atendieron más de una propiedad de las distribuciones muestrales.

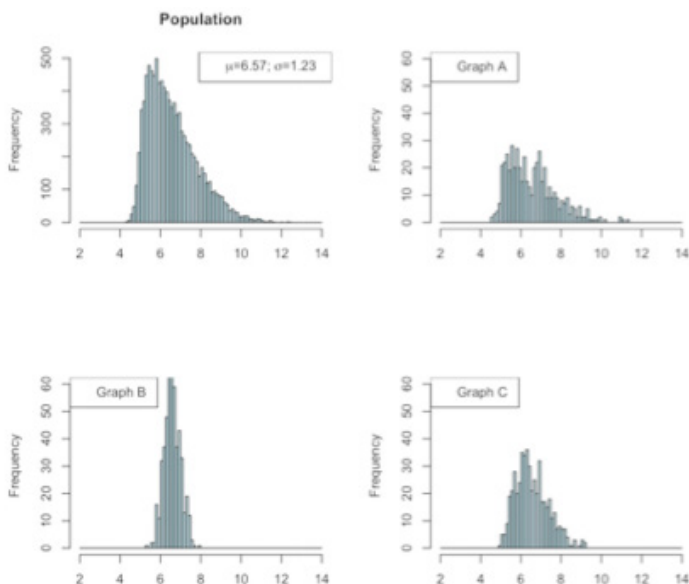
Porque en la gráfica B la distribución muestral da estimaciones entre 30 y 70 dulces cafés en cada uno, lo que es válido si las muestra son de 10 dulces con una

proporción de 50% de café, en la parte donde las muestras son de 100, las estimaciones de la distribución muestral se concentran casi todas en los 50 (Gustavo). Porque la primera muestra es muy pequeña, por lo que se entiende que su gráfica sea más dispersa, pero a al seleccionar una muestra más grande, los valores tienden a concentrarse en el centro, y vemos que en la gráfica B esto es lo que sucede. Además de que las proporciones de dulces café que se seleccionaron rondan en el 50% (María Alejandra).

Como puede observarse, la relación entre el tamaño de muestra con el centro y la variabilidad de una distribución muestral, parece ser comprendida por la mayor parte del estudiantado. Quienes clasificaron e identificaron en un razonamiento multiestructural, que el centro de una distribución muestral debe ser igual o cercano al parámetro poblacional y, que a mayor tamaño de muestra hay menor variabilidad en la distribución muestral, estableciendo así una relación entre población y resultados muestrales. Quienes mostraron un razonamiento uniestructural señalaron alguna de las dos propiedades anteriores, como fue el caso de Nicole que relaciona el parámetro de la población con el centro de las distribuciones muestrales para decidir al respuesta correcta. Es importante resaltar el nivel de intuición en el uso de lenguaje estadístico que mostraron algunos estudiantes en sus argumentaciones, al señalar aspectos como “concentrarse en el centro” o “ceranos al punto medio” en lugar de referirse a la variabilidad como tal.

ÍTEM 2

Se presentan cuatro gráficas a continuación. La primera, es una distribución de una población de calificaciones de un examen. La calificación media es de 6.57 y la desviación estándar es 1.23. Observa cuidadosamente y selecciona una gráfica apropiada para cada una de las siguientes dos preguntas. Este ítem fue tomado de la investigación de Chance, delMas y Garfield (2004), y tiene el propósito de evaluar si los estudiantes distinguen una muestra de una población y de una distribución muestral. Esta relación ha mostrado ser uno de los aspectos mas complejos del muestreo y las distribuciones muestrales, como se reporta por Inzunza (2015) y Saldanha y Thompson (2002). Harradine, Batanero y Rossman (2011), quienes señalan la importancia de tener claridad sobre el ámbito y la información que contienen estos tres tipos de distribuciones.



a) ¿Cuál gráfica (A, B o C) piensas que representa una sola muestra aleatoria de 500 valores desde esta población?

Un total de 13 estudiantes eligieron la gráfica A (correcta). Sin embargo en sus argumentaciones se observan inconsistencias y falta de claridad sobre los conceptos involucrados, por lo que la mayoría fueron clasificados con razonamientos de nivel preestructural.

Nivel	Cantidad
Preestructural	10
Uniestructural	3
Multiestructural	3
Total	16

Ejemplos de respuestas en cada nivel:

- Nivel preestructural

La gráfica C es la correcta en mi opinión pues cuenta con una forma de campana lo que significa que los datos están correctamente distribuidos (Kassandra).

La gráfica C representa un tamaño de muestra proporcional al de la población, concentrándose en una media de 6.57, sin embargo el tamaño de la muestra es pequeño y para ser más preciso debería contar con una muestra aleatoria mayor a la de 500 para ser más precisa (Salma).

- Nivel uniestructural

La selección de la gráfica A, es debido a que de las tres opciones, es la que tiene el mismo sesgo (hacia la derecha) como la gráfica de población (Hascibe)

Ya que al tener solamente una muestra, esto hace que la gráfica sea mas dispersa (Janeth)

- Nivel multiestructural

Tomo la gráfica A debido a que, al igual que la gráfica de la población, comprende la misma cantidad de intervalos (4 a 12) y también está sesgada hacia la derecha (Gustavo).

En resumen, la mayor parte de los estudiantes (13 de 16) tuvieron la intuición que la gráfica A representa una muestra de la población, pero sus argumentaciones carecen de los elementos conceptuales y el lenguaje estadístico que permita les permita justificar adecuadamente su elección, lo que explica la complejidad para distinguir y argumentar la diferencia entre las tres distribuciones.

b) ¿Cuál gráfica (A, B o C) piensas que representa una distribución de 500 medias muestrales de tamaño 9.

Este inciso fue mucho más complicado que el anterior, ya que sólo 4 estudiantes eligieron la gráfica correcta (grafica B), otros 10 estudiantes eligieron la gráfica C y 2 más la gráfica A.

Nivel	Cantidad
Preestructural	12
Uniestructural	4
Multiestructural	0
Total	16

Ejemplos de respuestas en cada nivel:

- Nivel Preestructural

Porque la gráfica C tiene una distribución normal, en forma de campana, lo que es común en las distribuciones muestrales. También, porque en dicha distribución, las estimaciones se concentran principalmente cerca del 6.57, que es la media de la población. Gustavo

La distribución de 500 medias de muestras de tamaño 9, se ven reflejadas en la gráfica C, ya que al ser la representación de medias muestrales, la distribución es normal. Hascibe

- Nivel Uniestructural

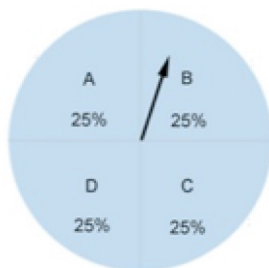
Porque la gráfica B concentra todos las medias de las muestras aleatorias de 9, y las medias se encuentran cerca de la calificación media (6.57), por lo que esta gráfica representa totalmente la distribución de las medias muestrales (Marcela).

Gustavo y Hascibe proporcionan respuestas típicas de quienes eligieron la gráfica C (incorrecta), mencionan la normalidad de la distribución muestral como un argumento en el que se basaron para elegirla, lo cual es correcto, pero en realidad la gráfica C tiene un sesgo muy similar al de la población que no debería presentar para un tamaño de muestra igual a 9.

Conjeturamos que a los estudiantes les pareció que la gráfica C tenía distribución normal y no así la gráfica B que era la correcta, debido a la forma prototípica que usualmente se tiene de la distribución normal como distribución cercana a la forma estandarizada. En un estudio realizado por Inzunza (2006) reporta este tipo de confusiones, donde los estudiantes no consideran distribución normal a una distribución angosta y alta, pero si a una distribución más dispersa similar a la forma estándar $N(0,1)$. En suma los estudiantes tienen una idea correcta de la normalidad de una distribución para muestras de tamaño 9, pero la idea de una distribución normal con poca altura y dispersa les hizo que eligieran una opción incorrecta. Por su parte Marcela que contesta correctamente, se basa en la cercanía de la medias con la media poblacional, pero además hace referencia a medias dejando claro que una distribución muestral es una colección de medias.

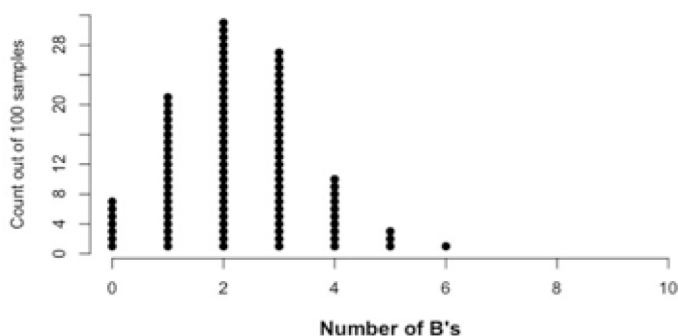
ÍTEM 3

Un estadístico realiza una prueba estadística para examinar el equilibrio de una ruleta (ver figura) usando un programa de cómputo de simulación. La simulación genera aleatoriamente alguna de las cuatro letras de la ruleta (A, B, C y D) con probabilidad de $\frac{1}{4}$. Se obtienen 100 muestras de 10 letras cada una y se registra el número de B's en cada muestra.



Este es un ítem muy completo que fue tomado de la prueba AIRS (Park, 2012) y se relaciona con el tercer nivel del modelo visual para la simulación de muestras (ver figura 1), el cual contempla la interpretación de información que contiene una distribución muestral, tal como la inusualidad de una muestra medida a través de su probabilidad, relación de la información muestral con características de la población y el efecto del tamaño muestral en la precisión de una estimación.

- La siguiente gráfica (distribución muestral) representa el número de B's para cada una de las 100 muestras. ¿Qué piensas sobre los resultados observados de 5 Bs de los 10 giros de la ruleta?



- a) 5 Bs no son inusuales porque 5 o menos Bs aparecieron en más de 90 muestras de 100.
- b) 5 Bs no son inusuales porque 5 o más B aparecieron en 4 muestras de 100.
- c) 5 Bs son inusuales porque 5 Bs aparecieron en sólo 3 muestras de 100.
- d) 5 Bs son inusuales por 5 o más Bs aparecieron en sólo 4 muestras de 100.
- e) No hay información suficiente para saber decidir si 5 Bs son inusuales o no.

La respuesta correcta está representada por el inciso c. Un total de 9 estudiantes eligieron la opción correcta, 3 estudiantes la opción d que es muy cercana a la opción correcta, pero que contempla un intervalo (5Bs o más) en lugar de un valor puntual (5 Bs), y 4 más consideraron que no hay suficiente información para tomar la decisión. En resumen podemos considerar que hubo un buen desempeño de los estudiantes al interpretar la información de la distribución muestral.

- ¿Qué puedes concluir acerca del equilibrio de la ruleta?
 - a) La ruleta es más probable que esté equilibrada porque 2 y 3 B's ocurren en la mayoría de la simulación.
 - b) La ruleta es más probable que esté equilibrada porque 5 o menos B's no son inusuales en la simulación.
 - c) La ruleta es más probable que esté equilibrada porque 5 o más B's son raras en la simulación.
 - d) La ruleta es más probable que esté equilibrada porque la distribución en la simulación parece sesgada.
 - e) No sabemos si la ruleta esté equilibrada porque el tamaño 10 de la muestra es pequeño

La proporción esperada de B's para una ruleta equilibrada que se hace girar 10 veces es de 2.5, por lo que la respuesta correcta la proporciona el inciso a. Sin embargo, sólo 3 estudiantes seleccionaron la respuesta correcta, 8 eligieron el inciso e, que señala que la muestra es muy pequeña para determinar si la ruleta está equilibrada. Nuestra conjetura es que estos estudiantes tienen una concepción aislada del muestreo, contrario a la concepción distribucional donde las muestras se ven como parte de una colección con características similares que pueden ser contabilizadas y expresadas como frecuencias; de haber razonando de esta manera, los estudiantes hubieran sido conscientes que la gráfica de la distribución tiene 100 muestras y que en los resultados más frecuentes fueron

2 y 3 B's, por lo que hay evidencia que orienta hacia el equilibrio de la ruleta. Estas concepciones han sido reportadas por Saldanha y Thompson (2002) y de nuevo se hacen presentes en esta investigación.

- Digamos que el estadístico hace otra simulación, pero esta vez cada muestra consta de 20 giros. Se calcula la proporción de B's en cada muestra (el número de B's dividido por 20). ¿Qué se puede esperar de la distribución de la proporción de B's obtenida en 100 muestras de 20 giros cada una comparada con la distribución de la proporción de B's obtenidas de 100 muestras de 10 giros cada una?
 - a) La distribución de la proporción de B's para 100 muestras de 20 giros cada una, será más ancha porque tiene el doble de giros en cada muestra.
 - b) La distribución de la proporción de B's para 100 repeticiones de 20 giros cada una, será más angosta porque tienes más información de cada muestra.
 - c) Ambas distribuciones tendrán la misma amplitud porque la probabilidad de obtener cada letra es la misma para 10 giros que para 20 giros.

Este ítem tiene el propósito de relacionar el tamaño de muestra con la variabilidad de una distribución muestral. Los estudiantes (11 de 16) responden correctamente al elegir el inciso b, con lo que reafirman resultados de ítems anteriores, pero 4 razonan que la variabilidad será mayor porque es mayor el tamaño de muestra, lo cual es incorrecto; esta forma de razonamiento fue clasificada como de *razonamiento de pequeño a grande*, en el estudio de Chance, delMas y Garfield (2004), cuando los estudiantes eligieron un histograma con mayor variabilidad para el tamaño de muestra grande.

- ¿Cuáles de los siguientes resultados, 5 B's en 10 giros o 10 B's en 20 giros, ofrece mayor evidencia de que la ruleta no está equilibrada? ¿Por qué?
 - a) 10 B's de 20 giros, porque mayores muestras tienen menor variabilidad, por lo que es menos probable obtener un resultado inusual con una ruleta equilibrada.
 - b) 5 B's de 10 giros, porque pequeñas muestras tienen mayor variabilidad, por lo que es más probable obtener un resultado inusual con una ruleta equilibrada.
 - c) Ambos resultados proveen de la misma evidencia porque hay la misma proporción de B's ($\frac{1}{2}$) en cada una de las dos muestras

Este ítem tiene el propósito de evaluar si los estudiantes comprenden que una muestra de mayor tamaño proporciona mayor precisión para estimar una característica de la población que una muestra pequeña, en este caso el equilibrio de la ruleta, expresado por una proporción de B's. Un total de 14 estudiantes eligieron el inciso a que es el correcto, con lo cual muestran evidencia que comprenden la relación entre tamaño de muestra y la probabilidad de obtener un resultado inusual en una distribución muestral.

CONCLUSIONES

Los resultados obtenidos son diferenciados, pues mientras en algunos conceptos los estudiantes razonaron adecuadamente en otros tuvieron dificultades. Es importante hacer notar que las dificultades identificadas son coincidentes con resultados de otros estudios, lo cual nos precisa sobre la dificultad de los conceptos y nos permite hacer una aproximación a un modelo explicativo que pueda orientar a continuar la investigación en el tema.

Los resultados muestran que la trayectoria de aprendizaje ha ayudado a muchos estudiantes a desarrollar un esquema de razonamiento correcto sobre los diversos conceptos que se involucran en las distribuciones muestrales, sin embargo, la falta de lenguaje estadístico observada en las argumentaciones hizo que algunos estudiantes fueran clasificados en un razonamiento de preestructural a uniestructural en los diversos ítems que se plantearon; esta problemática en las argumentaciones también ha sido reportada por Alvarado, *et al.* (2013).

La idea fundamental de tamaño de muestra y su efecto en la variabilidad de una distribución muestral y precisión de una estimación de características poblacionales, la cual tiene presencia en los tres niveles del esquema de simulación del muestreo proporcionado por Lane-Getaz (2006), ha sido comprendida en buena medida por los estudiantes, con algunas excepciones, según el contexto del ítem de evaluación.

En cuanto a la interpretación de la información que contiene una distribución muestral, la mayor parte de los estudiantes han identificado correctamente la inusualidad de una muestra –concepto clave en las pruebas de hipótesis–, pero tuvieron dificultades para ver las muestras como colección de casos similares que les permite calcular una frecuencia para decidir sobre el equilibrio de una

ruleta, relacionando los resultados muestrales con el valor esperado ligado al parámetro poblacional.

La confusión entre población, muestra y distribución muestral nos confirma la complejidad de estos conceptos, la cual ha sido reportada en la literatura por diversos autores (Saldahna y Thompson, 2002; Inzunza, 2006). En este ámbito muchos estudiantes lograron identificar la muestra entre población y distribuciones muestrales, pero fallaron al identificar la distribución muestral para una tamaño de muestra dado, interfiriendo en ello la imagen prototípica que se tiene una distribución muestral que se contrapone con una distribución alta y angosta como la que se mostraba en el ítem de evaluación.

Consideramos que los resultados nos permiten elaborar un rediseño a la trayectoria de aprendizaje en la búsqueda de mejorar el razonamiento donde se observaron las principales dificultades, para realizar una experimentación con una mayor cantidad de estudiantes, y superar así la principal limitación de este estudio.

REFERENCIAS

- Alvarado, H., Galindo, M. y Retamal, L. (2013). Comprensión de la distribución muestral mediante configuraciones didácticas y su implicación en la inferencia estadística. *Revista Enseñanza de las Ciencias* 31(2), 75-91.
- Batanero, C. y Borovnick, M. (2016). *Statistics and Probability in High School*. Rotterdam Netherlands: Sense Publishers
- Begué, N., Batanero, C. y Gea, M. (2018). Comprensión del valor esperado y variabilidad de la proporción muestral por estudiantes de educación secundaria obligatoria. *Revista Enseñanza de las Ciencias* 36(2), 63-79.
- Belia, S., Fidler, F., Williams, J., y Cumming, G. (2005). Researchers Misunderstand Confidence Intervals and Standard Error Bars. *Psychological Methods* 10(4), 389-396.
- Biggs, J. B. y Collis, K. F. (1982). *Evaluating the Quality of Learning: The Solo Taxonomy*. New York: Academic Press.
- Chance, B., delMas, R. y Garfield, J. (2004). Reasoning about Sampling Distributions. En D. Ben-Zvi y J. Garfield (Eds.), *The Challenge of Developing Statistical Literacy, Reasoning and Thinking* (pp. 295-323). Netherlands: Kluwer Academic Publishers.
- Cobb, G. W. (2007). The introductory statistics course: A Ptolemaic curriculum? *Technology Innovations in Statistics Education* 1(1).

- Dierdorff, A., Bakker, A., Maanen, J. y Eijkelhof, H. (2016). Supporting Students to Develop Concepts Underlying Sampling and to Shuttle Between Contextual and Statistical Spheres. En D. Ben-Zvi y K. Makar (Eds.), *The teaching and learning statistics: International perspectives* (pp. 37-48). Suiza: Springer.
- Garfield J. y Ben-Zvi, D. (2008). *Developing Students' Statistical Reasoning: Connecting Research and Teaching Practice*. Nueva York, N. Y.: Springer.
- Gottfried, B. (1984). *Elements of Stochastic Process Simulation*. Nueva Jersey, N. J.: Prentice Hall.
- Harradine, A., Batanero, C. y Rossman, A. (2011). Students and teachers' knowledge of sampling and inference. En C. Batanero, G. Burrill y C. Reading (Eds.), *Teaching Statistics in School Mathematics Challenges for Teaching and Teacher Education: A joint ICME/IASE Study* (pp. 235-246). Nueva York, N. Y. Springer Science + Business Media
- Inzunza, S. (2006). Significados que estudiantes universitarios atribuyen a las distribuciones muestrales en un ambiente de simulación computacional y estadística dinámica. Tesis doctoral. Departamento de Matemática Educativa. Cinvestav-IPN, CDMX, México.
- Inzunza, S. (2009). Construcción de significados sobre distribuciones muestrales y conceptos previos a la inferencia en un ambiente de simulación computacional. *Educación Matemática* 21(1), 119-150.
- Inzunza, S. (2013). Un acercamiento informal a la inferencia estadística mediante un ambiente computacional con estudiantes de bachillerato. *Revista electrónica de la Asociación Mexicana de Investigadores del uso de Tecnología en Educación Matemática* 1(1), 60-75.
- Inzunza, S. (2015). Dificultades en el desarrollo de una concepción estocástica de las distribuciones muestrales utilizando un ambiente computacional. En J. M. Contreras, C. Batanero, J. D. Godino, G. R. Cañadas, P. Arteaga, E. Molina, M. M. Gea y M.M. López (Eds.), *Didáctica de la Estadística, Probabilidad y Combinatoria* 2, (pp. 197-205). Granada, España: SEIEM.
- Kahneman D. y Tversky, A. (1972). Subjective Probability: A Judgment of Representativeness. *Cognitive Psychology* 3, 430-454.
- Kahneman, D., Slovic, P. y Tversky, A. (1982). *Judgment under uncertainty: Heuristics and biases*. Cambridge, R. U.: Cambridge University Press
- Lane-Getaz, S. J. (2006). What is statistical thinking, and how is it developed? En G. F. Burrill (Eds.), *Thinking and reasoning about data and chance: Sixty-eighth NCTM Yearbook* (pp 273-289). Reston, VA: National Council of Teachers of Mathematics.

- Lipson, K. (2000). *The Role of the Sampling Distribution in Developing Understanding of Statistical Inference*. Tesis doctoral no publicada. Universidad Tecnológica de Swinburne, Australia.
- Liu, Y. y Thompson, P. W. (2007). Teachers' understandings of probability. *Cognition and Instruction* 25(2), 113-160.
- Park, J. (2012). Developing and Validating an Instrument to Measure College Students' Inferential Reasoning in Statistics: An Argument-Based Approach to Validation. Tesis doctoral no publicada. Universidad de Minnesota.
- Pfannkuch, M. (2005). Probability and statistical inference: How can teachers enable learners to make the connection? En G. A. Jones (Ed.), *Exploring probability in school: Challenges for teaching and learning* (pp. 267-294). Nueva York, N. Y.: Springer.
- Saldanha, L. A. y Thompson, P. W. (2002). Conceptions of sample and their relationship to statistical inference. *Educational Studies in Mathematics* 51(3), 257-270.
- Sánchez, E. y Yáñez, G. (2003). Computational Simulation and Conditional Probability Problem Solving. En *Proceedings of the International Association of Statistical Education. Satellite Conference on Statistics Education: Statistics Education and the Internet*. Berlin [CD-ROM], Voorburg Netherlands. International Statistics Institute.
- Wild, Ch., Pfannkuch, M. y Reagan, M. (2011). Towards more accessible conceptions of statistical inference. *Journal of the Royal Statistical Society Series A* 174 (2), 247-295.

SANTIAGO INZUNZA CAZARES

Dirección: Circuito de las Amapas 785-51, Residencial Amapas
Culiacán de Rosales Sinaloa, CP. 80060

Teléfono: 667-1611055