



IPSOS VIEWS

EL PODER DE LOS TEST DE PRODUCTOS CON DATOS SINTÉTICOS

Humanizar la IA,
segunda parte

Colin Ho, Ph.D
Dr. Nikolai Reynolds



En Ipsos, defendemos la combinación única de Inteligencia Humana (IH) e Inteligencia Artificial (IA) para impulsar la innovación y ofrecer a nuestros clientes hallazgos con impacto y centrados en el ser humano.

Nuestra Inteligencia Humana proviene de la experiencia en ingeniería de *prompts*, *data science* y nuestros conjuntos de datos únicos y de alta calidad – que incorporan creatividad, curiosidad, ética y rigor en nuestras soluciones de IA, impulsadas por nuestra plataforma Ipsos Facto GenAI. Nuestros clientes se benefician de perspectivas que son más seguras, rápidas y basadas en el contexto humano.

#IpsosHiAi



En Ipsos, creemos que los datos sintéticos abren nuevas posibilidades para la investigación de mercados, especialmente en el campo de los test de productos.



Los datos sintéticos están a punto de cambiar el mundo. Desde acelerar el desarrollo de fármacos en el sector de la salud, simular transacciones fraudulentas en los servicios financieros y fomentar las pruebas de vehículos autónomos en el mercado automovilístico, ya está demostrando su valor en diversos contextos empresariales.

En Ipsos, creemos que los datos sintéticos abren nuevas posibilidades para la investigación de mercados, especialmente en el campo de los test de productos. Sin embargo, muchas empresas siguen teniendo dudas sobre la calidad de los datos sintéticos o sobre cómo evaluarlos. Este documento pretende cubrir estas brechas.

En su calidad de mayor y más importante asesor mundial en test de productos, Ipsos ha estado a la vanguardia del uso de

tecnologías de vanguardia para acelerar la innovación y el crecimiento de empresas de todo el mundo. Las siguientes páginas presentan las percepciones de Ipsos sobre los test de productos con datos sintéticos, proporcionando a los lectores:

- Recomendaciones para generar y evaluar conjuntos de datos sintéticos de alta calidad
- Aplicaciones específicas en test de productos para bienes de consumo y servicios

Hoy en día, encontramos diferentes tipos de datos sintéticos en la industria de la investigación de mercados, cada uno con sus fortalezas y debilidades. En este artículo, nos centramos en el aumento de datos, es decir, en la mejora de los conjuntos de datos con datos sintéticos.

Tipos de enfoques para utilizar datos sintéticos

Aumento de datos

Potenciar los conjuntos de datos para crear una muestra más completa con datos sintéticos, manteniendo la integridad estadística

Imputación de datos y fusión

Rellenar los datos que faltan con la información existente

Agentes de IA Gen y persona bots

Asistentes digitales personalizados que imitan a los segmentos de consumidores y ofrecen respuestas sintetizadas

Datos sintéticos completos

Utilización de muestras totalmente artificiales compuestas por encuestados sintéticos



Comenzamos con una visión general de lo que se necesita para generar datos sintéticos de alta calidad y cómo evaluar su calidad. Dado que los efectos y matices de los datos sintéticos sólo pueden comprenderse realmente cuando se consideran en ámbitos de aplicación específicos, investigamos cómo pueden aplicarse los datos sintéticos a la

prueba de productos, concretamente. En este artículo sólo se aborda la generación de datos numéricos sintéticos, el formato más utilizado por los investigadores cuantitativos. No cubre la aplicación de imágenes sintéticas, vídeo, imputación de datos o personas sintéticas, todos los cuales también entran dentro del amplio paraguas de los datos sintéticos¹.

Como se explica en el documento de Ipsos Views, **Synthetic Data: From Hype to Reality**¹, los datos sintéticos son datos artificiales generados a partir de un modelo entrenado para imitar las propiedades estadísticas y los patrones de los datos del mundo real.





Si una IA no ha sido entrenada con datos reales que sean relevantes para su negocio, no será capaz de generar datos sintéticos que compartan las mismas propiedades que los datos reales. Así de sencillo.



Generación y evaluación de datos sintéticos

Utilizamos los datos para tomar mejores decisiones empresariales en el mundo real, y los datos sintéticos pueden emplearse de muchas maneras para apoyar la toma de decisiones. Por ello, aunque los datos sintéticos no se correspondan con hechos o personas reales, deben imitar las propiedades y patrones estadísticos de los datos del mundo real. Esto plantea dos cuestiones fundamentales:

- 01 ¿Qué se necesita para generar datos sintéticos que se asemejen a los reales?
- 02 ¿Cómo se puede evaluar la semejanza de los datos sintéticos con los datos reales?

Antes de que una IA pueda generar datos sintéticos que reflejen los datos del mundo real, **necesita ser entrenada con datos del mundo real**. Como se explica en el primer artículo de Ipsos de la serie Humanizar la IA, las IA no son más que algoritmos; no tienen inteligencia propia hasta que se las entrena. Es mediante el aprendizaje a partir de datos de entrenamiento como las IA adquieren la inteligencia que asociamos a ellas. Este es el punto más importante que hay que recordar de este artículo:

si una IA no ha sido entrenada con datos reales que sean relevantes para su negocio, no será capaz de generar datos sintéticos que compartan las mismas propiedades que los datos reales. Así de sencillo.

El proceso de evaluación también es sencillo. Los datos numéricos sintéticos deben, como mínimo, **reflejar los datos del mundo real en medidas estadísticas comunes**, como medias, distribuciones de datos, varianzas y relaciones entre variables (por ejemplo, correlaciones). Una comparación directa entre los datos sintéticos y los humanos en estas métricas comunes nos dará una idea de hasta qué punto un conjunto de datos sintéticos se aproxima a los datos humanos. Cuanto más se aproximen los datos sintéticos a los humanos, menos riesgo asumiremos al utilizarlos, pero siempre existe cierto riesgo porque los datos sintéticos nunca pueden imitar perfectamente a los datos reales en todos los aspectos. Por tanto, sólo debemos utilizar datos sintéticos cuando estemos dispuestos a aceptar cierto riesgo.

Generación de datos sintéticos mediante LLM

Los métodos de generación de datos sintéticos pueden dividirse en dos categorías: Los LLM (Large Language Models) y los no-LLM, que se diferencian por su naturaleza textual y numérica, respectivamente. Los LLM comerciales o públicos, que se entrenan previamente con amplios conjuntos de datos como páginas web, libros en línea y publicaciones en redes sociales, pueden generar datos sintéticos de calidad en las áreas temáticas incluidas en su entrenamiento.

Sin embargo, los LLM comerciales tienen limitaciones a la hora de producir datos sintéticos realistas (véase la Figura 1)^{2,3}. Primero, sus datos de entrenamiento tienen una cobertura limitada: muchos temas son demasiado mundanos o privados para encontrarlos en Internet. En segundo lugar, los LLM suelen estar sesgados hacia los países occidentales de habla inglesa, debido

al predominio de este tipo de datos en sus conjuntos de entrenamiento. Los estudios han demostrado, por ejemplo, que los valores culturales generados por los LLM se alinean más con la anglosfera y la Europa protestante que con los de otros países⁴. En tercer lugar, la información puede quedar obsoleta rápidamente.

Por lo tanto, para generar datos sintéticos de alta calidad utilizando LLM, es crucial entrenarlos con datos reales actualizados y específicos de cada país que sean relevantes para el tema de interés. Este proceso requiere acceso a datos recientes, pertinentes y especializados, conocimientos estadísticos y de ciencia de datos, y una inversión considerable de tiempo y esfuerzo para garantizar que los datos sintéticos reflejen con precisión las propiedades y patrones estadísticos del mundo real⁵.

Figura 1: Comparaciones entre un conjunto de datos totalmente humanos y otro de humanos-sintéticos aumentados



Fuente:
Ipsos

Generación de datos sintéticos con enfoques no-LLM

Mucho antes de que los LLM acapararan la atención, los científicos de datos han utilizado algoritmos de aprendizaje profundo (Deep Learning, DL)⁶ para generar datos numéricos sintéticos⁷. Los algoritmos de DL, incluidos los tipos utilizados en los LLM, son herramientas potentes para la generación de datos sintéticos, y cada uno posee ventajas únicas.

Los LLM son especialmente eficaces para generar datos de texto similares a los humanos. Pueden proporcionar datos textuales detallados y contextualmente ricos, lo que los hace muy valiosos en aplicaciones como la creación de contenidos, la traducción de idiomas y los chatbots⁸.

La DL no-LLM ofrece ventajas significativas para generar datos numéricos sintéticos. Los

algoritmos de DL son especialmente eficaces en la producción de datos numéricos sintéticos que reflejan fielmente las propiedades estadísticas de los conjuntos de datos del mundo real. Mientras que los LLM preentrenados como ChatGPT están diseñados para tareas de lenguaje natural, los modelos de DL pueden entrenarse específicamente para la síntesis de datos numéricos, permitiendo la personalización para aplicaciones específicas del dominio o relevantes para el mercado. En Ipsos, contamos con un sólido historial de aprovechamiento de técnicas de DL para la síntesis de datos. En la siguiente sección, detallamos los resultados de aplicar modelos de DL para generar datos sintéticos para test de productos, analizados producto por producto.

Por qué en test de productos

Fuera de la investigación de mercados, muchas aplicaciones de datos sintéticos se centran en el anonimato (por ejemplo, la confidencialidad de los datos médicos anónimos). En la investigación de mercados, el principal beneficio que buscan muchas empresas es el ahorro de tiempo y dinero que supone no tener que recopilar datos del mundo real.

Dado que la obtención de datos en línea es cada año más rápida y menos costosa, hay que considerar detenidamente si el ahorro de tiempo y dinero compensa la disminución de la precisión que conllevan los datos sintéticos. Por ejemplo, utilizando [Ipsos.Digital](#), la plataforma ágil de pruebas de Ipsos, un investigador en Estados Unidos puede realizar una encuesta a 300 encuestados por unos 2.000 dólares y obtener resultados en unas 24 horas. Hipotéticamente, si el investigador aprovecha los datos sintéticos para responder a las preguntas que los 300

humanos reales habrían proporcionado, las respuestas con datos sintéticos pueden tardar 12 horas en generarse, costar 1.500 dólares y ofrecer una precisión menor en comparación con los datos humanos reales. En esta situación, puede tener más sentido recopilar datos reales en lugar de generar datos sintéticos, ya que el ahorro de costos y tiempo no parece compensar la disminución de la precisión. ¿Por qué no gastarse los 500 dólares adicionales y esperar 12 horas más para obtener los datos reales?

Debido al *trade-off* inherente a los datos sintéticos, queríamos probarlos en situaciones en las que el costo de la investigación típicamente es más elevado. Los test de productos encajan perfectamente en este caso, debido a los muchos gastos que conllevan:



Fabricación: Fabricación o compra de productos y prototipos, o enmascaramiento de los mismos para pruebas de producto



Envíos y devoluciones: Costos de entrega de productos, y devolución o destrucción de envases vacíos



Muestreo: Costos de contratación, sobre todo cuando las empresas tienen que ser selectivas con su base de usuarios

Debido a los costos de fabricación, envío y muestreo, cualquier reducción del número de participantes en una prueba de producto puede suponer un ahorro significativo. Esto no quiere decir que los datos sintéticos no puedan utilizarse en otras áreas de la investigación de mercado, simplemente significa que el riesgo puede superar a los beneficios si el ahorro de costos es mínimo.

Seguimos necesitando humanos

La experiencia del producto es inherentemente humana. La IA por sí sola no puede captar los cinco sentidos, las emociones, las expectativas o el impacto del contexto que los humanos experimentan con los productos. Por tanto, nuestro objetivo al aplicar los datos sintéticos a los test de productos no era sustituir por completo

la aportación humana, sino aumentarla. Nuestro desafío consistía en establecer el número mínimo de encuestados humanos necesarios para probar productos junto con datos sintéticos, a fin de garantizar resultados viables. Para lograrlo, los equipos de innovación de Ipsos llevaron a cabo dos líneas de investigación:

Línea de investigación 1

Determinar el menor número de seres humanos necesario para aproximar los resultados de los test de productos a partir de muestras más grandes (por ejemplo, 200-300 seres humanos) *sin datos sintéticos*

En la **línea de investigación 1**, aprovechamos algunos de los datos de la base de datos de test de productos de Ipsos, considerando 40.000 encuestados y 185 productos seleccionados de bienes de consumo envasados (CPG) probados en todo el mundo⁹, para evaluar con qué número mínimo de muestras humanas obtenemos una correlación suficientemente alta con los resultados de los test de productos de

Línea de investigación 2

Validar que una pequeña muestra humana, cuando se aumenta con datos sintéticos, produce los mismos resultados que una muestra totalmente humanas

muestras más grandes. Determinamos que cuando el producto con mejores resultados difiere del producto con peores resultados en al menos un 8%, una muestra de 50 encuestados humanos es suficiente para reproducir las clasificaciones de resultados de los mejores y peores productos (coeficiente de correlación $r = 0,8$).



La experiencia del producto es inherentemente humana. La IA por sí sola no puede captar los cinco sentidos, las emociones, las expectativas o el impacto del contexto que los seres humanos experimentan con los productos.



El reducido número de participantes necesario para replicar los resultados de muestras más grandes se debe probablemente a que la varianza en los datos de los test de productos está influida principalmente por las diferencias en la tecnología del producto (por ejemplo, la cantidad de azúcar utilizada). **En consecuencia, la varianza observada en los test de productos suele ser menor que en otros ámbitos de investigación, como la evaluación de las actitudes de los consumidores o la percepción de las marcas, por ejemplo.**

En un escenario ideal, las empresas actuarían sobre la base de estos resultados y realizarían test de productos con 50 participantes, especialmente durante las primeras fases de las pruebas, cuando el riesgo es menor.

Sin embargo, la mayoría de las empresas no proceden así por dos razones:

01 Una muestra de 50 no permite a las empresas obtener información sobre los subgrupos. A veces, las empresas necesitan analizar segmentos específicos dentro de la población de consumidores.

02 Un tamaño de muestra de 50 da lugar a una baja potencia estadística para detectar diferencias. Esto es problemático, porque las empresas suelen confiar en normas de actuación basadas en pruebas estadísticas. En la práctica, utilizar una muestra de 50 puede impedir que las empresas sigan adelante con los resultados de la investigación. En términos estadísticos, las muestras pequeñas aumentan el riesgo de errores de Tipo II, es decir, la posibilidad de no detectar una diferencia real cuando existe.

Por lo tanto, en lugar de reclutar a 200 humanos para probar un producto, por ejemplo, podríamos reclutar a 50 humanos para probar un producto, generar 150 encuestados sintéticos a partir de los 50 datos humanos sin replicar ni volver a muestrear a los 50 humanos y, a continuación, combinar los encuestados humanos y sintéticos para probar con una muestra mixta de 200. Consideremos los 50 humanos como una muestra «semilla» para entrenar a la IA, permitiéndole generar datos sintéticos que imiten con precisión las respuestas humanas a los productos.

Aquí es donde entra en juego la **línea de investigación 2**, para validar si las pequeñas

muestras humanas, aumentadas con datos sintéticos, seguirían produciendo los mismos resultados que las muestras totalmente humanas. Para garantizar la generalización, en nuestras pruebas piloto con datos sintéticos cubrimos un conjunto diverso de países y categorías. También experimentamos con muestras de semillas más grandes (por ejemplo, 75, 100), pero para abreviar, sólo compartiremos los resultados de nuestras pruebas con 50 humanos.

A continuación, los datos de los 50 humanos se utilizaron para entrenar un algoritmo de DL que generara datos sintéticos.

No utilizamos un LLM estándar porque los datos públicos con los que se han entrenado previamente los LLM no incluyen las experiencias multisensoriales de las personas con productos de la categoría en cuestión (por ejemplo, prototipos, nuevas fórmulas) y no proporcionan datos numéricos sólidos sobre los encuestados. Además, no consideramos la ponderación de los datos como una alternativa a la DL, porque la ponderación no ayuda a crear tamaños de muestra adicionales para subgrupos y a menudo no se acepta en los test de productos.



[Ipsos] validó los resultados comparando las conclusiones de un conjunto de datos exclusivamente humanos con las de un conjunto de datos ampliado con datos sintéticos.



En general, los datos sintéticos funcionan

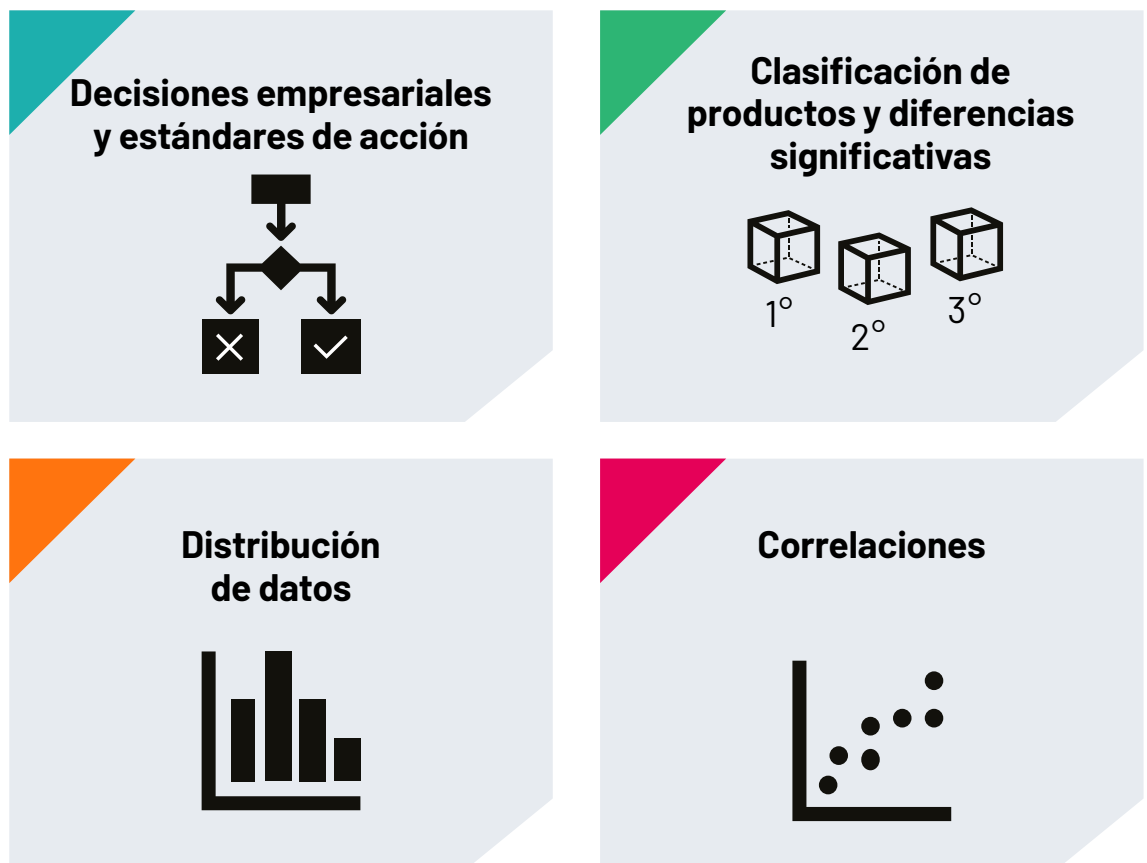
En nuestros experimentos, validamos los resultados comparando las conclusiones de un conjunto de datos exclusivamente humanos con las de un conjunto de datos ampliado con datos sintéticos (véase la Figura 1)¹⁰. Como recordatorio, en nuestro enfoque no nos limitamos a replicar o copiar y pegar los datos existentes de los encuestados.

En general, comprobamos que los dos conjuntos de datos eran notablemente similares, en términos de:

- Los resultados relativos de los productos (por ejemplo, clasificaciones, significación estadística)
- La distribución de los datos (por ejemplo, la distribución de las respuestas de los encuestados entre las opciones de respuesta de cada pregunta)
- Las relaciones entre las variables de los datos (por ejemplo, las correlaciones entre el agrado general y los atributos de los productos)

Y lo que es más importante, los dos conjuntos de datos mostraron diferencias en las varianzas, pero condujeron a la misma decisión comercial en todos los conjuntos de datos que probamos (véase la figura 2).

Figura 2: Criterios utilizados para determinar la exactitud de la réplica en parte humana y en parte sintética de los resultados del conjunto de datos totalmente humanos



Fuente:
Ipsos

Una ventaja clave de los test de productos con datos sintéticos es la posibilidad de aumentar los datos para poblaciones de difícil acceso. Una vez aumentados, las diferencias que antes no eran estadísticamente significativas pueden llegar a serlo debido al aumento del tamaño de la muestra. En una de nuestras pruebas, por ejemplo, generamos respuestas sintéticas para aumentar el número de personas que utilizaban una determinada marca de producto. En nuestra muestra totalmente humana,

teníamos unos 100 usuarios de una marca por producto. En el conjunto de datos humanos, había algunas diferencias entre los tres productos probados, pero las diferencias no alcanzaban significación estadística. Al añadir 100 usuarios sintéticos de la marca por producto, las diferencias entre productos pasaron a ser estadísticamente significativas debido al aumento del tamaño de la muestra (véase la figura 3)

Figura 3: Usuarios humanos de marcas aumentados con usuarios sintéticos de marcas



Fuente:
Ipsos



Cuando buscamos opiniones sobre productos de un grupo objetivo específico, debemos establecer cuotas de reclutamiento que se ajusten a los objetivos empresariales



Se requiere cierta cautela

Hemos presentado una imagen prometedora de los datos sintéticos. Como ya se ha mencionado, los datos sintéticos por sí solos conllevan una desventaja inherente, ya que nunca podrán igualar por completo la precisión de los datos del mundo real. La capacidad de un conjunto de datos en parte humanos y en parte sintéticos para reproducir los resultados de un conjunto de datos totalmente humanos depende de lo siguiente:

- **La representatividad de los datos humanos utilizados para entrenar la IA.** En nuestras



pruebas piloto iniciales, cuando no controlamos a los 50 participantes humanos para alinearlos con la estructura de la muestra original de 200 personas y generamos datos sintéticos a partir de los 50, observamos diferencias en los resultados. Sin dicho control, en el 20-30% de estas muestras, los datos sintéticos generados no captaban el rendimiento del producto reflejado en los 200 humanos originales. La conclusión es que los datos sintéticos pueden fallar si la muestra semilla humana no es representativa. Sin embargo, al comparar la estructura de

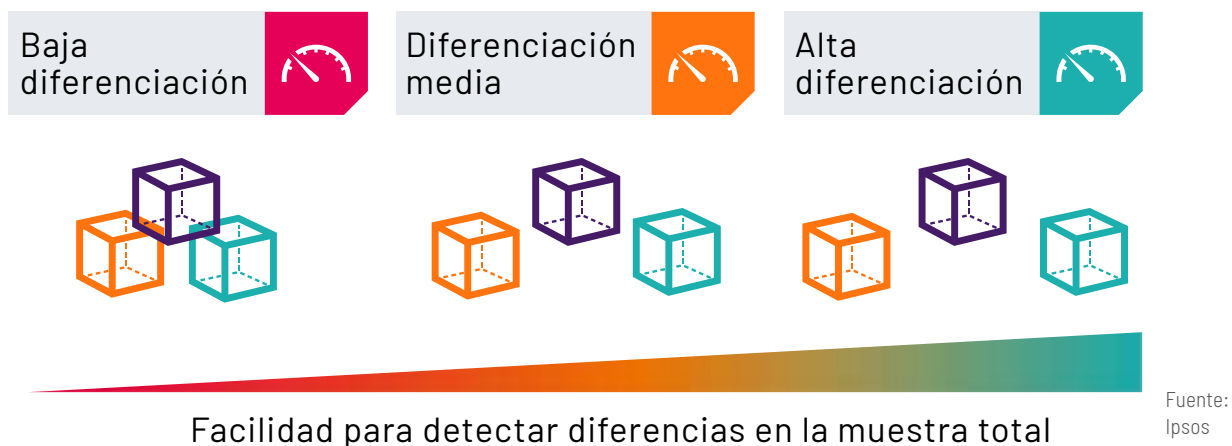
los 50 participantes humanos con la de la muestra original de 200 personas, la decisión empresarial habría sido la misma utilizando una muestra combinada de 50 humanos y 150 sintéticos que utilizando 200 humanos. En resumen, cuando buscamos opiniones sobre productos de un grupo objetivo específico, debemos establecer cuotas de reclutamiento que se ajusten a los objetivos empresariales. Al igual que hicimos en el pasado con muestras de mayor tamaño, esto garantiza que estamos muestreando el grupo objetivo adecuado.

- **Cómo se diferencian los productos en primer lugar.** Cuando hay diferencias significativas entre productos en el conjunto de datos humanos, la IA puede detectarlas fácilmente y reproducirlas en los datos sintéticos. Cuando las diferencias entre productos no son estadísticamente significativas, la IA, al igual que en la muestra humana, tampoco puede detectarlas (véase la figura 4). En las numerosas pruebas piloto que realizamos en varios países y categorías, identificamos una ventaja adicional del uso de muestras aumentadas en situaciones de baja discriminación entre productos. Al



analizar subgrupos dividiendo la muestra original en subconjuntos más pequeños, como usuarios actuales de marcas y no usuarios de marcas, descubrimos que el uso de datos sintéticos generados por IA para aumentar los tamaños de los subgrupos más pequeños mejoraba la detección de diferencias en el rendimiento de los productos, en función de las preferencias del grupo objetivo. Esto fue evidente a nivel de subgrupo, donde los datos sintéticos ayudaron a aumentar la potencia estadística y a descubrir preferencias que podrían no haberse detectado solo con la muestra original.

Figura 4: La exactitud de los datos sintéticos depende de las diferencias reales entre productos

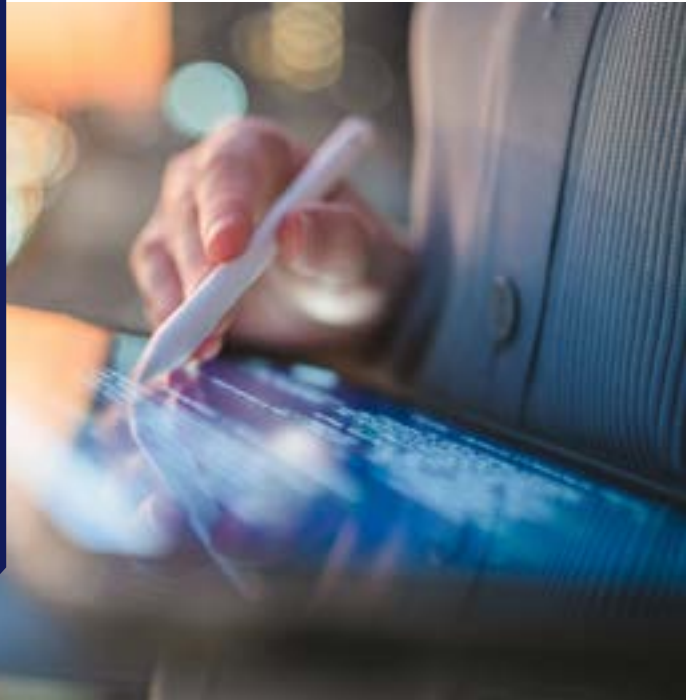


Como nota final sobre la validación, para evitar que el lector se vaya con la impresión de que puede empezar a hacer todos los estudios de mercado con sólo 50 humanos, debemos dejar claro que las conclusiones aquí presentadas son específicas de los test de productos. Las condiciones que hicieron posible utilizar sólo 50 seres humanos en la prueba de productos pueden no darse en otros ámbitos de investigación. Como asesor líder mundial en test de productos, tomamos

algunas enseñanzas de la base de datos de test de productos de Ipsos. Por no mencionar que los test de productos tienen características de datos únicas que pueden ayudar a garantizar la calidad de los datos sintéticos (por ejemplo, las características del producto que influyen en los sentidos humanos). Los datos sintéticos deben validarse siempre, caso por caso, para la aplicación específica deseada (por ejemplo, segmentación).



A menudo se considera que los datos son la parte vital de las empresas, ya que permiten tomar decisiones más inteligentes que las ayudan a crecer y prosperar



La promesa de los datos sintéticos: del revuelo a la realidad

A menudo se considera que los datos son la parte vital de las empresas, ya que permiten tomar decisiones más inteligentes que las ayudan a crecer y prosperar. La promesa de poder generar datos sintéticos a voluntad y a gran escala es, por tanto, muy atractiva. Sin embargo, la opinión pública sobre los datos sintéticos está bastante polarizada. Las empresas que ofrecen servicios de datos sintéticos pueden decir «esto es todo lo que necesita, ¡no se necesitan humanos! Los investigadores más precavidos prefieren esperar y ver qué pasa, y por el momento dudan en utilizar datos sintéticos. Bajo estas opiniones subyace una tendencia a categorizar el mundo de forma binaria: ¿Buenos o malos? ¿Datos sintéticos o del mundo real? Hemos dado un primer paso para aportar claridad sobre este tema, para demostrar que los datos sintéticos no son diferentes de otras herramientas de investigación. Los datos sintéticos tienen sus puntos fuertes y débiles y son más adecuados para determinadas situaciones.

En nuestros pilotos, demostramos que la IA podía generar datos sintéticos para imitar

los del mundo real, pero antes necesita datos humanos de calidad para entrenarse. Por tanto, la respuesta no es datos sintéticos o del mundo real. Necesitamos ambos. La precisión de los datos sintéticos no es buena o mala; más bien, «depende». Si las diferencias de producto son pequeñas entre humanos, tenemos que buscar subgrupos. Si los datos humanos utilizados para entrenar la IA no son representativos del grupo objetivo o relevantes para la empresa, la precisión de los datos sintéticos se verá comprometida. Si queremos utilizar datos sintéticos, debemos aceptar que pueden no funcionar en algunas condiciones. **Como investigadores, nuestra responsabilidad es garantizar que utilizamos datos sintéticos sólo cuando es apropiado: en condiciones que maximicen el éxito.** Aumentar los datos sintéticos ofrece varias ventajas en comparación con el uso de muestras más pequeñas, como la posibilidad de realizar análisis de subgrupos, conservar la potencia estadística y llevar a cabo análisis más complejos.

En general, no creemos que los datos sintéticos deban sustituir por completo a los humanos, al menos no en los test de productos. En la película de 1997 «Good Will Hunting», el fallecido actor Robin Williams interpretaba a un profesor que tutelaba a un joven genio, interpretado por Matt Damon. El prodigio posee una gran cantidad de conocimientos, gracias a su capacidad sobrehumana para absorber información de los libros. En una de las escenas de la película, el profesor aconseja al prodigio sobre la distinción entre el conocimiento de los libros y la experiencia del mundo real. El profesor le dice: «Pero apuesto a que no puedes decirme cómo huele la Capilla Sixtina. Nunca has estado allí y has contemplado ese hermoso techo». El verdadero conocimiento proviene de la vida, no de libros, fotos, vídeos o cualquier otra representación del mundo real.

Como el prodigio de la película, una IA puede recibir todos los conocimientos del mundo que existan, pero nunca podrá experimentar el mundo como un ser humano. Hay algo único y hermoso en ser humano y poder sentir el calor del sol en la cara, disfrutar de la melodía o el ritmo de la música, o la capacidad de contemplar una hermosa puesta de sol: experiencias que la IA nunca podrá reproducir del todo, por muy avanzada que llegue a ser. La forma en que los seres humanos reaccionamos ante los productos, o ante la vida en general, no se capta únicamente en el cerebro como conocimiento factual o semántico, sino que nuestro cuerpo y nuestras experiencias sensoriales también desempeñan un papel importante.



El verdadero conocimiento proviene de la vida, no de libros, fotos, videos o cualquier otra representación del mundo.



Puntos claves

01

Los datos sintéticos nunca serán humanos.

La IA por sí sola nunca podrá hacer eco de nuestras experiencias de producto, que combinan los cinco sentidos, las emociones, las expectativas y el contexto. Por tanto, nuestro objetivo es aumentar la información humana con datos sintéticos, no sustituirla análisis detallados de subgrupos.

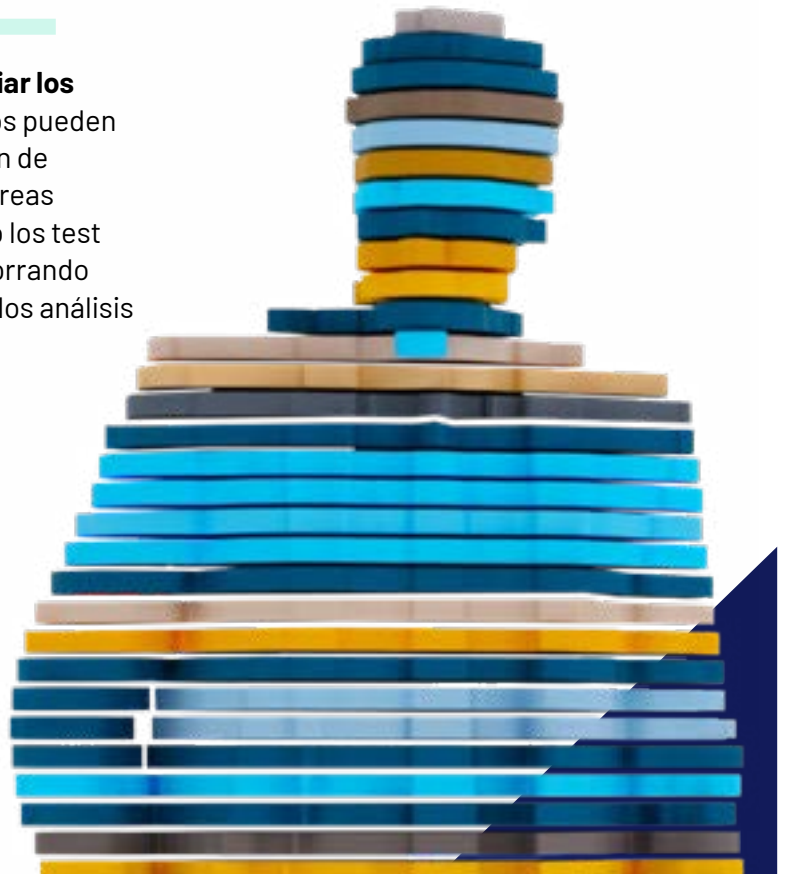
02

La precisión depende de los datos de

entrenamiento. El valor de los datos sintéticos no es binario (bueno o malo); la precisión de los datos sintéticos depende de muchos factores, entre ellos las diferencias en los datos que intentamos replicar y la representatividad de los datos del mundo real con los que estamos entrenando a una IA para que aprenda. El uso de datos sintéticos debe ser estratégico, teniendo en cuenta los riesgos y beneficios asociados.

03

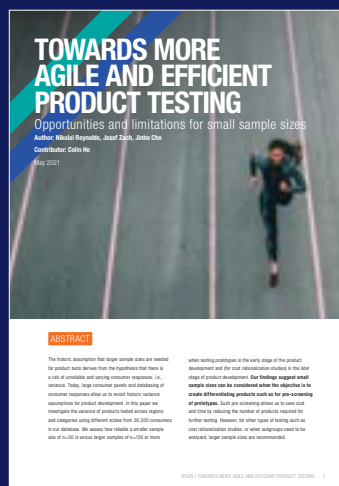
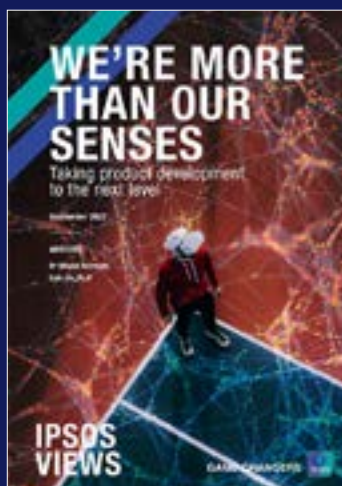
Cuando son precisos, pueden potenciar los test de productos. Los datos sintéticos pueden aumentar la agilidad de la investigación de mercado, por lo que son ideales para áreas que requieren muchos recursos, como los test de productos, reduciendo costos y ahorrando tiempo, con ventajas adicionales para los análisis detallados de subgrupos.



Notas

- 1 Guidi, M., Hubert, B., Sava, C., & Timpone, R. (2024). Synthetic Data: From Hype to Reality – a guide to responsible adoption. Ipsos POV
- 2 Illic, Maya, Bangia, Ajay, Legg Jim (2024). Conversations with AI Part V. Is there depth and empathy with AI twins? Ipsos Views
- 3 Moore Chris, Stronge Cameron, Bhudiya, Manjula (2024). Judgment Day: The Machines Have Arrived – But how good are they at answering choice experiments? Sawtooth Conference.
- 4 Yan Tao, Olga Viberg, Ryan S. Baker and René F. Kizilcec (2024). Cultural bias and cultural alignment of large language models. PNAS Nexus, Vol. 3, No. 9
- 5 Lisa P. Argyle, Ethan C. Busby, Nancy Fulda, Joshua Gubler, Christopher Rytting, and David Wingate (2022). Out of One, Many: Using Language Models to Simulate Human Samples. arXiv.
- 6 AI-based Deep Learning is a way for computers to learn by analyzing large amounts of data and finding patterns, much like how humans learn from experience. It uses neural networks inspired by the human brain to recognize information and make decisions without being explicitly programmed.
- 7 Ian J. Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, Yoshua Bengio (2014). Generative Adversarial Networks. arXiv.
- 8 Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). Language Models are Unsupervised Multitask Learners.
- 9 Reynolds, N., Zach, J., Cho, J, Ho, C. (2021). Towards more agile and efficient product testing. Opportunities and limitations for smaller sample sizes, Ipsos POV.
- 10 We also compared the results from the all-synthetic data to the all-human data; the findings are like those reported in this paper

Para leer más



ABRIL, 2025

EL PODER DE LOS TEST DE PRODUCTOS CON DATOS SINTÉTICOS

Humanizar la IA,
segunda parte

AUTORES

Colin Ho, Ph.D

Chief Research Officer,
Innovation and Market Strategy
& Understanding, Ipsos

Dr. Nikolai Reynolds

Global Head of Product Testing,
Ipsos

Los artículos de **IPSOS
VIEWS** son producidos por el
Ipsos Knowledge Centre.

www.ipsos.com

@Ipsos

