

Campo / Field	Contenido / Content
Tipo de contribución / Article type	Artículo original de investigación basado en análisis secundario de datos
Título en español	Automatización creativa con inteligencia artificial en publicidad digital: desempeño visual y efectos de la divulgación en dos conjuntos de datos abiertos
Title in English	Creative Automation with Artificial Intelligence in Digital Advertising: Visual Performance and Disclosure Effects across Two Open Datasets
Autorías / Authors	MARLON ELINOVICH TENECELA-CALDERON — UNIVERSIDAD DE GUAYAQUIL, Ecuador — ORCID: https://orcid.org/0009-0005-0559-7365
Correspondencia / Correspondence	MARLON ELINOVICH TENECELA-CALDERON — marlon.tenecela@asesoreseducativos-ec.com

Resumen

La inteligencia artificial generativa permite producir múltiples variantes publicitarias con rapidez, pero la utilidad de esa automatización depende tanto del desempeño visual como de la respuesta del público cuando se divulga la intervención de IA. Este estudio examinó ambas dimensiones mediante análisis secundarios independientes de dos conjuntos abiertos. En GenImageNet se analizaron 10.320 imágenes, 103.200 valoraciones humanas, siete modelos y ocho dominios; se estimaron descriptivos y ANOVA factoriales para calidad, realismo y estética. En un experimento de anuncios de Instagram se reprodujo la muestra analítica de 161 participantes —control, divulgación de IA en imagen y divulgación de IA en texto— y se compararon actitud hacia el anuncio, actitud hacia la marca, credibilidad y menor manipulación percibida, con corrección de Holm. El modelo generativo, el dominio y su interacción se asociaron con los tres desenlaces visuales (todos $p < 0,001$). El mayor efecto correspondió al modelo sobre realismo, η^2 cuadrado parcial = 0,127. Firefly 2 obtuvo la mayor media descriptiva de calidad (5,453; $n = 240$), mientras Realistic Vision obtuvo la mayor de realismo (5,411; $n = 2.400$). En el experimento, la divulgación redujo la actitud hacia el anuncio frente al control: diferencia = 0,555, IC 95 % [0,218; 0,892], p de Holm = 0,011. También hubo contrastes planificados en marca, credibilidad y menor manipulación percibida, pero solo el ómnibus de actitud hacia el anuncio permaneció significativo tras corregir cuatro pruebas. No se observaron diferencias entre divulgar IA en imagen o texto. La automatización creativa no ofrece un modelo universalmente superior y la transparencia puede producir costos de recepción. La selección de herramientas debe validarse por dominio y la divulgación debe evaluarse como parte del diseño comunicacional, sin renunciar a las obligaciones de transparencia.

Palabras clave: inteligencia artificial generativa; publicidad digital; creatividad computacional; imágenes sintéticas; divulgación de IA; actitud publicitaria

Abstract

Generative artificial intelligence enables rapid production of multiple advertising variants, yet the value of such automation depends on both visual performance and audience responses when AI involvement is disclosed. This study examined these dimensions through independent secondary analyses of two open datasets. GenImageNet comprised 10,320 images, 103,200 human ratings, seven models, and eight domains; descriptive statistics and factorial ANOVAs were estimated for quality, realism, and aesthetics. A separate Instagram advertising experiment was reproduced with an analytical sample of 161 participants assigned to control, image-AI disclosure, or text-AI disclosure conditions. Ad attitude, brand attitude, source credibility, and lower perceived manipulateness were compared using Holm correction. Generative model, image domain, and their interaction were associated with all three visual outcomes (all $p < .001$). The largest effect was the model effect on realism, partial η^2 = .127. Firefly 2 had the highest descriptive mean for quality (5.453; $n = 240$), whereas Realistic Vision had the highest

realism mean (5.411; $n = 2,400$). In the experiment, disclosure lowered ad attitude relative to control: difference = .555, 95% CI [.218, .892], Holm $p = .011$. Planned contrasts were also observed for brand attitude, credibility, and lower perceived manipulateness, although only the omnibus test for ad attitude remained significant after correcting four tests. Image- and text-specific AI disclosures did not differ. Creative automation therefore offers no universally superior model, and transparency may carry reception costs. Tool selection should be validated by domain, and disclosure should be treated as part of communication design without abandoning transparency obligations.

Keywords: generative artificial intelligence; digital advertising; computational creativity; synthetic images; AI disclosure; advertising attitude

1. Introducción

La automatización publicitaria ha evolucionado desde la compra programática y la optimización de audiencias hacia la producción automática de texto, imágenes y variaciones creativas. Los modelos generativos reducen el costo marginal de producir alternativas y amplían el espacio de experimentación, pero no convierten la creatividad en una propiedad uniforme ni garantizan que una pieza sea percibida como convincente. La literatura sobre IA generativa identifica oportunidades de productividad y colaboración humano-máquina junto con riesgos de error, homogeneización, propiedad intelectual y gobernanza (Dwivedi et al., 2023; Nah et al., 2023; Al-Busaidi et al., 2024). En marketing, el valor operativo depende de cómo se articulan generación, selección y evaluación, no solamente de la capacidad de emitir contenido (Guercini, 2023; Pagani & Wind, 2024).

La imagen ocupa un lugar central en la publicidad digital. Su calidad percibida puede influir en atención, comprensión y evaluación, mientras el realismo y la estética responden a criterios relacionados pero no equivalentes. Una imagen visualmente atractiva puede resultar poco realista; otra puede parecer verosímil sin alcanzar la misma evaluación estética. Hartmann et al. (2025) plantearon que los sistemas generativos podían competir e incluso superar referencias humanas en tareas de marketing visual, y publicaron GenImageNet para facilitar comparaciones sistemáticas. El conjunto permite observar modelos y dominios en una escala que excede la evaluación anecdótica de unas pocas piezas.

Sin embargo, una clasificación global de modelos puede ser engañosa. Los generadores difieren en entrenamiento, arquitectura, filtros y tratamiento de instrucciones, y el rendimiento puede cambiar entre productos, alimentos, personas o anuncios. La interacción modelo $\times$ dominio es por ello sustantiva: si existe, la elección de herramienta debería depender del tipo de pieza. Además, una media superior no informa por sí sola sobre adecuación de marca, originalidad, legalidad, costo, tiempo de producción o eficacia comercial. Este estudio mantiene esos constructos fuera del alcance porque no están medidos por GenImageNet.

La segunda dimensión es la transparencia. Las divulgaciones publicitarias pueden activar conocimiento de persuasión y modificar la forma en que una audiencia procesa un mensaje (Friestad & Wright, 1994; Boerman et al., 2017; Rahmani, 2023). En el caso de contenido generado con IA, la etiqueta añade información sobre la agencia de producción. Esa información puede valorarse como honestidad, pero también funcionar como una señal que reduce autenticidad, esfuerzo percibido o afinidad humana. Estudios en arte, música, chatbots y anuncios han documentado formas de aversión algorítmica y penalización de la autoría artificial (Dietvorst et al., 2015; Luo et al., 2019; Bellaiche et al., 2023; Shank et al., 2023).

Las respuestas no tienen por qué ser uniformes. El marco de agencia de la máquina propone que las personas atribuyen funciones y responsabilidad de manera diferente cuando intervienen sistemas automáticos (Sundar, 2020). En publicidad, Campbell et al. (2022) destacaron que los contenidos manipulados y los anuncios generados por IA requieren analizar fuente, mensaje y contexto. Wortel et al. (2024) realizaron un experimento de tres condiciones con un anuncio ficticio de Instagram y diferenciaron divulgación relativa a la imagen, divulgación relativa al texto y ausencia de divulgación. Ese diseño permite separar el efecto de informar sobre IA del componente creativo al que se atribuye su uso.

La percepción de manipulación requiere especial cuidado. La escala de Campbell (1995) incluye formulaciones cuya codificación en el archivo reutilizado produce valores más altos para una evaluación más favorable, es decir, menor manipulación percibida. Ignorar la dirección podría invertir la interpretación. De igual modo, una correlación entre esa escala y las actitudes no demuestra mediación. El presente reanálisis se centra en comparaciones transparentes de medias, multiplicidad y tamaños de efecto, sin reconstruir el modelo de mediación causal del artículo fuente.

El problema de investigación se formula así: ¿cómo varía el desempeño visual de imágenes generadas entre modelos y dominios, y cómo cambia la recepción de un anuncio cuando se divulga que su imagen o su texto se creó con IA? El objetivo general fue caracterizar esas dos dimensiones complementarias de la automatización creativa mediante datos públicos. Los objetivos específicos fueron: describir calidad, realismo y estética por modelo; estimar efectos de modelo, dominio e interacción; comparar control frente a divulgación y divulgación de imagen frente a texto; evaluar la consistencia de las escalas; y derivar implicaciones responsables para selección y comunicación de herramientas.

No se fusionaron los conjuntos. El estudio visual observa imágenes y valoraciones agregadas; el estudio de divulgación observa participantes expuestos a un único anuncio ficticio. La síntesis final es conceptual y no implica que los modelos de GenImageNet hayan producido el estímulo de Instagram ni que el desempeño visual cause las actitudes observadas.

2. Metodología

2.1. Diseño general y fecha de corte

Se realizó un análisis secundario cuantitativo de dos estudios abiertos e independientes. La fecha efectiva y el corte bibliográfico fueron el 31 de julio de 2025. La búsqueda posterior se utilizó solo para verificar metadatos existentes antes de esa fecha. El protocolo se congeló después de auditorías preliminares de viabilidad; por ello se describe como plan analítico del reanálisis y no como prerregistro prospectivo.

2.2. Estudio A: GenImageNet

La fuente fue GenImageNet, depositada en OSF por Hartmann, Exner y Domdey y asociada con su artículo en *International Journal of Research in Marketing*. Se analizaron 10.320 imágenes generadas por DALL-E 3, Firefly 2, Imagen 2, Imagine, Midjourney v6, Realistic Vision y SDXL Turbo. El conjunto comprende ocho dominios de imagen. Uno aparece en la fuente con la etiqueta `ions_awards`, que se conserva en los datos y se presenta como Lions Awards con advertencia de codificación.

Cada imagen tuvo diez valoraciones, para 103.200 asignaciones y 1.238 evaluadores. La unidad analítica fue la imagen y los desenlaces fueron `qualityMean`, `realismMean` y `aestheticsScore`. Se comprobaron dimensiones, unicidad de nombres, asignaciones por imagen y datos faltantes. No se imputaron valores ni se excluyeron imágenes del análisis principal.

Se calcularon *n*, media, desviación estándar, error estándar e intervalo del 95 % por modelo y por dominio. Para cada desenlace se ajustó un ANOVA factorial con modelo, dominio e interacción. El término residual tuvo 10.264 grados de libertad. Se informaron *F* y *eta* cuadrado parcial. Dada la magnitud de la muestra, la interpretación no se basó únicamente en significación estadística.

El diseño no constituye un ensayo publicitario de conversión. Las valoraciones describen propiedades percibidas de imágenes bajo las condiciones del estudio fuente; no miden clics, ventas, costo, duración del trabajo, cumplimiento de marca, originalidad jurídica ni desempeño frente a creativos humanos encargados para una campaña específica.

2.3. Estudio B: divulgación en Instagram

La segunda fuente fue el experimento de Wortel, Vanwesenbeeck y Tomas (2024), depositado en OSF. El diseño original fue entre sujetos y utilizó un anuncio de una marca ficticia de mochilas. Los participantes fueron asignados a control, divulgación de que la imagen fue creada con IA o divulgación de que el texto fue creado con IA.

El archivo contenía 231 filas. Se aplicaron los filtros documentados de uso de Instagram, atención y comprobación de manipulación, con una muestra final de 161: 55 en control, 55 en divulgación de imagen y 51 en divulgación de texto. El reanálisis no reclutó personas ni modificó la asignación.

Los desenlaces fueron actitud hacia el anuncio, actitud hacia la marca, credibilidad de la fuente y la puntuación derivada de los ítems de intención manipulativa. Debido a la dirección de codificación, esta última se denomina menor manipulación percibida: un valor alto representa una evaluación más favorable. La aversión a la IA se utilizó en análisis exploratorios.

La consistencia interna se estimó mediante alfa de Cronbach. Se calcularon medias, desviaciones estándar e intervalos del 95 % por condición. Para cada desenlace se ajustó un ANOVA de una vía. Se definieron dos contrastes planificados: control frente al promedio de las dos divulgaciones, e imagen frente a texto. Se reportaron diferencia, intervalo del 95 %, diferencia estandarizada y p ajustado. Holm corrigió por separado cuatro pruebas ómnibus y ocho contrastes.

Como exploración se estimó el bloque de interacción condición $\times$ aversión a la IA en modelos lineales y se corrigieron cuatro pruebas mediante Holm. También se calcularon correlaciones entre menor manipulación percibida y los tres desenlaces favorables. Esas correlaciones son descriptivas y no se presentan como mediación.

2.4. Privacidad, ética y licencias

El archivo público original del experimento incluía direcciones IP y trazas técnicas. Esos campos no eran necesarios para la pregunta científica. La fuente se preservó inmutable, pero antes del análisis se creó una extracción desidentificada que excluyó direcciones IP, fechas individuales, tiempos de clic, navegación y residuales. Los derivados, el cuaderno y los manuscritos contienen únicamente variables necesarias y resultados agregados.

El estudio no intentó reidentificar participantes. El experimento original informó aprobación ética. Para este análisis secundario no fue necesaria una nueva aprobación ética. GenImageNet se reutilizó respetando la licencia indicada en su depósito; no se redistribuyeron las imágenes ni los datos brutos, y solo se informaron agregados. Las huellas SHA-256 y el código se conservaron en el expediente.

3. Resultados

3.1. Integridad de los conjuntos

GenImageNet incluyó 10.320 nombres de archivo únicos, 103.200 asignaciones y exactamente diez valoraciones por imagen. No se encontraron duplicados ni ausencias en los campos analíticos. El experimento conservó 161 participantes después de reproducir los filtros y no tuvo datos faltantes en los cuatro desenlaces.

3.2. Diferencias visuales por modelo

Las medias de calidad se situaron entre 5,126 para SDXL Turbo y 5,453 para Firefly 2. Realistic Vision y Midjourney v6 tuvieron valores próximos, 5,350 y 5,349. En realismo, Realistic Vision alcanzó 5,411, seguido de Firefly 2 con 5,299 y Midjourney v6 con 5,231. En estética, Firefly 2 obtuvo 5,573, Realistic Vision 5,477 e Imagine 5,462. Firefly 2 aportó 240 imágenes, frente a 2.400 en varios modelos; su posición descriptiva debe interpretarse con esa diferencia de soporte.

Tabla 1. Desempeño visual por modelo generativo

Modelo	n	Calidad, M	Realismo, M	Estética, M
DALL-E 3	2.400	5,221	4,819	5,356
Firefly 2	240	5,453	5,299	5,573
Imagen 2	240	5,308	5,143	5,453
Imagine	240	5,336	5,161	5,462
Midjourney v6	2.400	5,349	5,231	5,266
Realistic Vision	2.400	5,350	5,411	5,477
SDXL Turbo	2.400	5,126	5,048	5,294

Desempeño visual por modelo generativo

Medias por imagen; barras en escala de 1 a 7

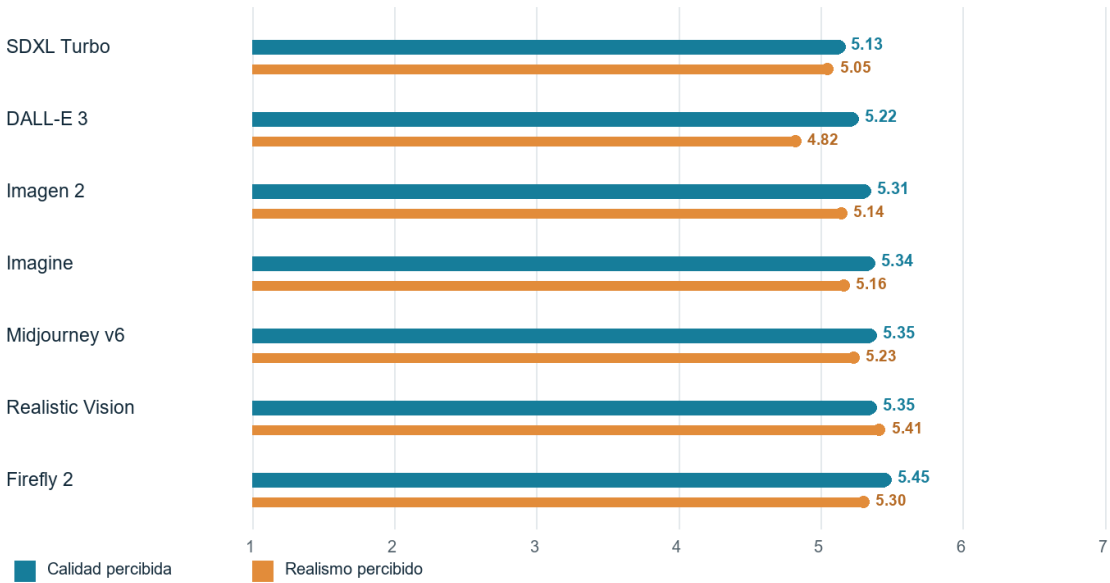


Figura 1. Medias e intervalos de confianza del desempeño visual por modelo.

3.3. Modelo, dominio e interacción

En calidad, los efectos de modelo, dominio e interacción fueron estadísticamente distinguibles. Modelo: $F(6, 10264) = 80,001$, $p < 0,001$, eta cuadrado parcial = 0,0447; dominio: $F(7, 10264) = 72,549$, $p < 0,001$, eta cuadrado parcial = 0,0471; interacción: $F(42, 10264) = 3,731$, $p < 0,001$, eta cuadrado parcial = 0,0150.

En realismo, el modelo explicó una proporción mayor: $F(6, 10264) = 249,001$, $p < 0,001$, eta cuadrado parcial = 0,1271. El dominio alcanzó $F(7, 10264) = 107,139$, eta cuadrado parcial = 0,0681, y la interacción $F(42, 10264) = 4,859$, eta cuadrado parcial = 0,0195.

Para estética, modelo $F(6, 10264) = 39,214$, eta cuadrado parcial = 0,0224; dominio $F(7, 10264) = 107,992$, eta cuadrado parcial = 0,0686; e interacción $F(42, 10264) = 3,532$, eta cuadrado parcial = 0,0142; todos $p < 0,001$. La existencia de interacciones en los tres desenlaces indica que las diferencias entre modelos no son constantes entre dominios.

Tabla 2. ANOVA factorial de los desenlaces visuales

Desenlace	Efecto	F	gl	p	Eta cuadrado parcial
Calidad	Modelo	80,001	6; 10264	<0,001	0,0447
Calidad	Dominio	72,549	7; 10264	<0,001	0,0471
Calidad	Modelo × dominio	3,731	42; 10264	<0,001	0,0150
Realismo	Modelo	249,001	6; 10264	<0,001	0,1271
Realismo	Dominio	107,139	7; 10264	<0,001	0,0681
Realismo	Modelo × dominio	4,859	42; 10264	<0,001	0,0195
Estética	Modelo	39,214	6; 10264	<0,001	0,0224
Estética	Dominio	107,992	7; 10264	<0,001	0,0686
Estética	Modelo × dominio	3,532	42; 10264	<0,001	0,0142

3.4. Recepción de la divulgación

El control obtuvo medias más favorables en los cuatro desenlaces. La actitud hacia el anuncio fue 4,689 en control, 4,118 con divulgación de imagen y 4,150 con divulgación de texto. La actitud hacia la marca fue 4,724, 4,160 y 4,153; la credibilidad, 4,234, 3,718 y 3,824; y menor manipulación percibida, 5,077, 4,702 y 4,681.

Tabla 3. Resultados por condición experimental

Desenlace	Control M (DE)	IA imagen M (DE)	IA texto M (DE)	F(2,158)	p de Holm	Eta cuadrado
Actitud hacia el anuncio	4,689 (0,992)	4,118 (1,027)	4,150 (1,060)	5,316	0,023	0,063
Actitud hacia la marca	4,724 (1,090)	4,160 (1,330)	4,153 (1,256)	3,855	0,057	0,047
Credibilidad	4,234 (0,795)	3,718 (1,115)	3,824 (1,068)	4,056	0,057	0,049
Menor manipulación percibida	5,077 (0,605)	4,702 (1,011)	4,681 (0,786)	4,014	0,057	0,048

Solo la prueba ómnibus de actitud hacia el anuncio permaneció por debajo de 0,05 después de corregir las cuatro pruebas. En los contrastes planificados, control superó al promedio de divulgación para actitud hacia el anuncio: diferencia = 0,555, IC 95 % [0,218; 0,892], p de Holm = 0,011, diferencia estandarizada = 0,541. También se observaron contrastes ajustados en actitud hacia la marca, diferencia = 0,567, p = 0,037; credibilidad, diferencia = 0,463, p = 0,037; y menor manipulación percibida, diferencia = 0,385, p = 0,037. Estas pruebas responden a comparaciones a priori, pero la discrepancia con los ómnibus corregidos aconseja describir las tres últimas como evidencia limitada.

Ningún contraste imagen frente a texto fue significativo; todos los valores p de Holm fueron 1,000. El contenido específico de la divulgación no alteró detectablemente la respuesta en este estímulo.

Recepción publicitaria según divulgación de IA

Medias e intervalos de confianza del 95%; escala de 1 a 7

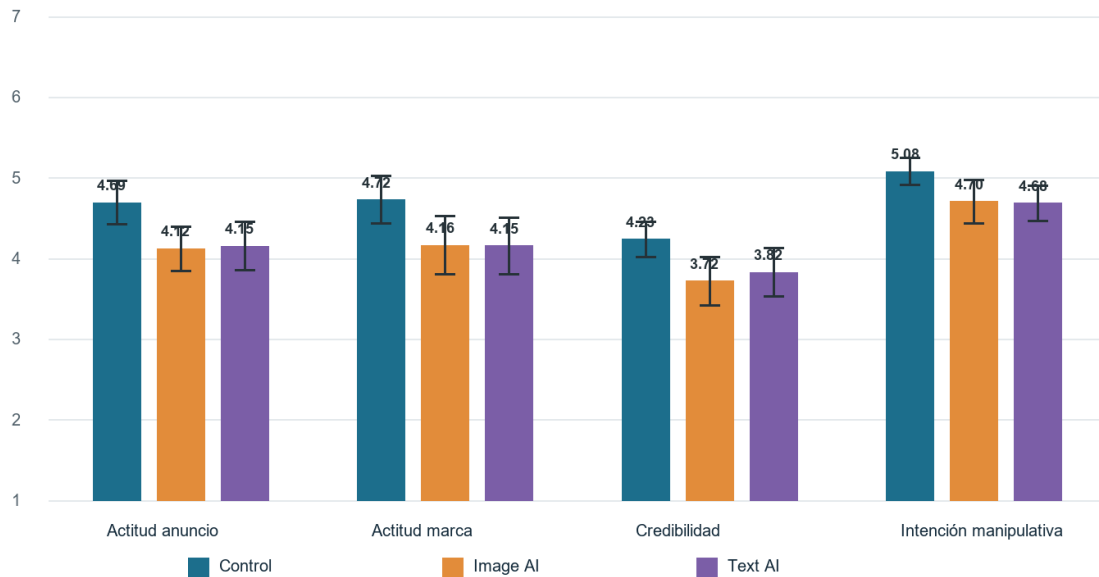


Figura 2. Medias e intervalos de confianza de recepción por condición de divulgación.

3.5. Fiabilidad y exploraciones

La consistencia interna fue $\alpha = 0,881$ para actitud hacia el anuncio, $0,940$ para actitud hacia la marca, $0,895$ para credibilidad, $0,805$ para menor manipulación percibida y $0,862$ para aversión a la IA.

Las interacciones condición $\times$ aversión no fueron robustas después de Holm. Para actitud hacia el anuncio, $F(2, 155) = 4,506$, p ajustado = $0,0501$; para marca, $F = 3,936$, $p = 0,0645$; credibilidad y menor manipulación tuvieron $p = 1,000$. El primer resultado es próximo, pero superior, al umbral definido y no se clasifica como significativo.

Menor manipulación percibida se correlacionó con actitud hacia el anuncio, $r = 0,491$; marca, $r = 0,513$; y credibilidad, $r = 0,632$; todos p de Holm $< 0,001$. La orientación positiva respalda la etiqueta aplicada a la puntuación, pero no establece una cadena causal.

4. Discusión

Los dos estudios revelan una tensión central de la automatización creativa. La generación visual puede alcanzar evaluaciones favorables a gran escala, pero su desempeño varía entre modelos y dominios; además, informar al público sobre la intervención de IA puede reducir la respuesta hacia un anuncio. La eficiencia productiva no elimina la necesidad de selección humana, validación contextual y diseño transparente de la comunicación.

En GenImageNet, ningún modelo fue superior de forma uniforme en todos los criterios. Firefly 2 encabezó descriptivamente calidad y estética, pero tuvo un tamaño menor; Realistic Vision encabezó realismo y se mantuvo entre los primeros en las otras dimensiones. Midjourney v6 tuvo calidad similar a Realistic Vision, pero una media estética menor que varios competidores. Estas diferencias confirman que “mejor imagen” no es un constructo único. Una campaña que necesita verosimilitud de producto podría priorizar criterios distintos de otra orientada a fantasía, impacto o estilo.

Los efectos del dominio fueron comparables o superiores a los del modelo para calidad y estética. La interacción, aunque de magnitud menor, fue consistente. Operativamente, esto favorece evaluaciones por caso de uso: un equipo no debería seleccionar un generador con base en una tabla global y asumir que conservará su ventaja en alimentos,

personas, productos o piezas publicitarias. Una estrategia responsable sería mantener conjuntos de prueba por dominio, criterios explícitos y revisión humana antes de desplegar producción masiva.

El mayor tamaño de efecto correspondió al modelo sobre realismo. Esto resulta relevante porque la verosimilitud puede afectar expectativas, especialmente cuando la imagen representa productos o situaciones que el consumidor interpretará como evidencia. Sin embargo, más realismo no equivale a más verdad. Los sistemas sintéticos pueden producir escenas plausibles de hechos inexistentes. Campbell et al. (2022) e Ienca (2023) advierten que la manipulación con IA depende de cómo las capacidades técnicas interactúan con intención, contexto y vulnerabilidad. Por ello, la calidad técnica debe acompañarse de controles de exactitud y procedencia.

El experimento de Instagram muestra un costo de divulgación especialmente claro en actitud hacia el anuncio. El tamaño estandarizado del contraste control-divulgación fue moderado. El resultado coincide con trabajos donde etiquetar intervención algorítmica activa respuestas negativas o preferencias por producción humana (Luo et al., 2019; Bellaiche et al., 2023; Shank et al., 2023). No obstante, no demuestra que ocultar la IA sea una estrategia recomendable. La transparencia puede ser una obligación ética o regulatoria y permite a la audiencia interpretar la procedencia. La pregunta práctica no es si se debe engañar para preservar una actitud favorable, sino cómo diseñar una divulgación comprensible y cómo crear piezas cuyo valor resista esa información.

Las comparaciones de actitud de marca, credibilidad y menor manipulación percibida requieren cautela. Los contrastes planificados sobrevivieron su familia de ocho pruebas, pero los correspondientes omnibus quedaron en p ajustado = 0,057. La evidencia apunta en la misma dirección, aunque no sostiene una conclusión fuerte para cada variable bajo la corrección más conservadora. Presentar únicamente los valores nominales habría exagerado certeza; informar ambas familias permite una lectura proporcionada.

La ausencia de diferencias entre “imagen creada con IA” y “texto creado con IA” sugiere que, para este anuncio, la señal general de intervención artificial pesó más que el componente nombrado. Wortel et al. (2024) llegaron a una conclusión compatible mediante su modelo original. Esto no implica equivalencia universal: otras formulaciones, tamaños, ubicaciones o productos podrían producir respuestas diferentes. La literatura sobre divulgación patrocinada muestra que formato, momento y alfabetización influyen en la activación del conocimiento de persuasión (Boerman et al., 2017; Beckert et al., 2021).

La puntuación de menor manipulación percibida fue más alta en control, un patrón que puede parecer contraintuitivo si se asume que divulgar siempre aumenta honestidad percibida. La dirección de la escala es crucial. Además, su asociación positiva con actitud y credibilidad indica que quienes evaluaron el anuncio como menos manipulativo también lo evaluaron más favorablemente. No se puede concluir que esta percepción medie causalmente el efecto: desenlace y escala se midieron en la misma sesión, y este reanálisis no estimó una intervención sobre el mediador.

La aversión a la IA no mostró moderación robusta después de multiplicidad. El valor ajustado de 0,0501 para actitud hacia el anuncio ilustra por qué no conviene convertir un límite numérico en una afirmación dicotómica. El resultado justifica replicación con mayor potencia y una muestra más diversa, pero no respalda segmentar campañas como si el efecto estuviera establecido. La escala AIAS-4 mostró buena consistencia interna, de modo que la falta de robustez no parece explicarse simplemente por fiabilidad deficiente.

La integración conceptual de ambos estudios ofrece cuatro implicaciones. Primera, la automatización debe gestionarse como un portafolio de modelos y no como una sustitución única. Segunda, los criterios de calidad tienen que corresponder al objetivo de la pieza y al dominio. Tercera, la divulgación forma parte de la experiencia creativa: contenido, etiqueta y contexto se perciben conjuntamente. Cuarta, la evaluación interna no puede detenerse en métricas visuales; debe incluir autenticidad, comprensión, credibilidad y riesgo de manipulación.

Estas implicaciones son relevantes para organizaciones pequeñas, pero los datos no permiten estimar ahorros, adopción ni resultados de microempresas ecuatorianas. Sería improcedente convertir la reducción potencial del costo marginal en una promesa empírica. La evidencia solo muestra diferencias perceptuales y visuales en dos fuentes internacionales. La transferencia a otra población o mercado requiere pruebas locales.

La gobernanza de activos también importa. Los modelos pueden crear variaciones a gran velocidad, lo que aumenta el número de piezas que deben revisarse. Al-Busaidi et al. (2024) muestran que propiedad intelectual y conocimiento organizacional requieren límites claros. Laine et al. (2025) proponen abordar la ética generativa mediante principios establecidos y nuevos. En publicidad, una bitácora mínima debería registrar modelo, versión, instrucciones, materiales de referencia, revisión humana, derechos y decisión de divulgación.

La automatización creativa, por tanto, se asemeja menos a un botón que produce creatividad y más a un sistema sociotécnico de generación, curaduría, prueba y rendición de cuentas. El aporte de la IA está en ampliar alternativas; la responsabilidad de seleccionar, verificar y comunicar permanece en la organización y sus autores humanos.

5. Conclusiones

El desempeño de imágenes generadas con IA varió significativamente por modelo, dominio y combinación de ambos. Realistic Vision obtuvo el mayor realismo promedio, mientras Firefly 2 encabezó calidad y estética en un subconjunto menor. Estos resultados no sustentan un ganador universal: la selección debe validarse en el dominio y criterio relevante.

En el experimento de Instagram, divulgar la intervención de IA redujo de forma robusta la actitud hacia el anuncio frente al control. No se detectaron diferencias entre divulgar IA en la imagen o en el texto. Las señales para actitud de marca, credibilidad y menor manipulación fueron concordantes, pero más sensibles a la familia de corrección empleada.

La eficiencia de generación y la aceptación del público son dimensiones distintas. Una estrategia de automatización creativa responsable debe combinar pruebas por dominio, curaduría humana, control de derechos y exactitud, y divulgaciones claras evaluadas con usuarios. La transparencia no debe abandonarse por su posible costo actitudinal; ese costo debe anticiparse y tratarse mediante mejor diseño y evidencia.

6. Limitaciones

Los dos conjuntos se analizaron de manera independiente y no permiten vincular el modelo que produjo una imagen con la respuesta a una divulgación. GenImageNet utiliza tareas, modelos y versiones específicas; las plataformas cambian rápidamente y sus posiciones relativas pueden variar. Los tamaños por modelo fueron desiguales y los análisis no cubren costo, tiempo, originalidad, coherencia de marca, accesibilidad o conversión.

Las valoraciones visuales están anidadas en imágenes y evaluadores; el reanálisis utilizó medias por imagen y ANOVA factorial, por lo que no reconstruyó toda la dependencia individual. Las pruebas con gran n pueden detectar diferencias pequeñas. Los tamaños de efecto y el dominio son más informativos que el valor p aislado.

El experimento tuvo 161 participantes, muestreo principalmente universitario, un producto, una marca ficticia, una plataforma y una forma de divulgación. No representa automáticamente otras edades, países, marcas, productos, videos o mercados. La multiplicidad redujo la solidez de algunos resultados y los análisis de moderación fueron exploratorios.

El archivo público original contenía direcciones IP. Aunque los derivados fueron desidentificados y verificados, su presencia en la fuente constituye una limitación de gobernanza del dato. No se redistribuye el archivo ni identificadores. El presente reanálisis no requirió una nueva aprobación ética.

Agradecimientos

Se reconoce a Jochen Hartmann, Yannick Exner y Sebastian Domdey, y a Caroline Wortel, Ini Vanwesenbeeck y Frédéric Tomas, por publicar los datos y materiales que hicieron posible este análisis. Este reconocimiento no implica participación en el reanálisis ni aprobación de sus conclusiones.

Declaraciones

Campo / Field	Contenido / Content
Contribuciones CRediT	Conceptualización; metodología; software; validación; análisis formal; investigación; curación de datos; visualización; redacción del borrador original; redacción, revisión y edición; administración del proyecto.
Financiamiento	La investigación no recibió financiación externa específica.
Conflictos de interés	La persona autora declara no tener conflictos de interés.
Aprobación ética	Los estudios fuente informaron sus condiciones éticas. El presente reanálisis utilizó datos secundarios públicos desidentificados, no reclutó participantes, no recopiló datos nuevos y no requirió una nueva aprobación ética.
Consentimiento	No aplicable al presente análisis secundario; no se reclutaron participantes ni se recopilaron datos nuevos.
Disponibilidad de datos	OSF: https://osf.io/8ctjy/ y https://osf.io/8fj pz/ . Código y derivados desidentificados en el expediente.
Uso de inteligencia artificial	Codex de OpenAI apoyó la programación reproducible, la organización del manuscrito, la traducción académica y los controles editoriales. La persona autora revisó el contenido y asume la responsabilidad íntegra por la exactitud, originalidad, interpretación y versión final del artículo.

Referencias

- Al-Busaidi, A. S., Raman, R., Hughes, L., Albashrawi, M., Malik, T., Dwivedi, Y. K., et al. (2024). Redefining boundaries in innovation and knowledge domains: Investigating the impact of generative artificial intelligence on copyright and intellectual property rights. *Journal of Innovation & Knowledge*, 9, 100630. <https://doi.org/10.1016/j.jik.2024.100630>
- Arango, L., Singaraju, S. P., & Niininen, O. (2023). Consumer responses to AI-generated charitable giving ads. *Journal of Advertising*, 52(4), 486–503. <https://doi.org/10.1080/00913367.2023.2183285>
- Beckert, J., Koch, T., Viererbl, B., & Schulz-Knappe, C. (2021). The disclosure paradox: How persuasion knowledge mediates disclosure effects in sponsored media content. *International Journal of Advertising*, 40(7), 1160–1186. <https://doi.org/10.1080/02650487.2020.1859171>
- Bellaiche, L., Shahi, R., Turpin, M. H., Ragnhildstveit, A., Sprockett, S., Barr, N., Christensen, A., & Seli, P. (2023). Humans versus AI: Whether and why we prefer human-created compared to AI-created artwork. *Cognitive Research: Principles and Implications*, 8, 42. <https://doi.org/10.1186/s41235-023-00499-6>
- Boerman, S. C., Willemsen, L. M., & Van der Aa, E. P. (2017). “This post is sponsored”: Effects of sponsorship disclosure on persuasion knowledge and electronic word of mouth in the context of Facebook. *Journal of Interactive Marketing*, 38, 82–92. <https://doi.org/10.1016/j.intmar.2016.12.002>

- Campbell, C., Plangger, K., Sands, S., & Kietzmann, J. (2022). Preparing for an era of deepfakes and AI-generated ads: A framework for understanding responses to manipulated advertising. *Journal of Advertising*, 51(1), 22–38. <https://doi.org/10.1080/00913367.2021.1909515>
- Campbell, M. C. (1995). When attention-getting advertising tactics elicit consumer inferences of manipulative intent: The importance of balancing benefits and investments. *Journal of Consumer Psychology*, 4(3), 225–254. https://doi.org/10.1207/s15327663jcp0403_02
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E., Jeyaraj, A., Kar, A. K., et al. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>
- Grassini, S. (2023). Development and validation of the AI attitude scale (AIAS-4): A brief measure of general attitude toward artificial intelligence. *Frontiers in Psychology*, 14, 1191628. <https://doi.org/10.3389/fpsyg.2023.1191628>
- Guercini, S. (2023). Marketing automation and the scope of marketers’ heuristics. *Management Decision*, 61(13), 295–320. <https://doi.org/10.1108/MD-07-2022-0909>
- Hartmann, J., Exner, Y., & Domdey, S. (2025). The power of generative marketing: Can generative AI create superhuman visual marketing content? *International Journal of Research in Marketing*, 42(1), 13–31. <https://doi.org/10.1016/j.ijresmar.2024.09.002>
- Ienca, M. (2023). On artificial intelligence and manipulation. *Topoi*, 42(3), 833–842. <https://doi.org/10.1007/s11245-023-09940-3>
- Laine, J., Minkkinen, M., & Mäntymäki, M. (2025). Understanding the ethics of generative AI: Established and new ethical principles. *Communications of the Association for Information Systems*, 56, 1–29. <https://doi.org/10.17705/1CAIS.05601>
- Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Machines versus humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Marketing Science*, 38(6), 937–947. <https://doi.org/10.1287/mksc.2019.1192>
- Nah, F. F.-H., Zheng, R., Cai, J., Siau, K., & Chen, L. (2023). Generative AI and ChatGPT: Applications, challenges, and AI-human collaboration. *Journal of Information Technology Case and Application Research*, 25(3), 277–304. <https://doi.org/10.1080/15228053.2023.2233814>
- Pagani, M., & Wind, Y. (2024). Unlocking marketing creativity using artificial intelligence. *Journal of Interactive Marketing*, 59(4), 335–347. <https://doi.org/10.1177/10949968241265855>
- Rahmani, V. (2023). Persuasion knowledge framework: Toward a comprehensive model of consumers’ persuasion knowledge. *AMS Review*, 13, 12–33. <https://doi.org/10.1007/s13162-023-00254-6>
- Shank, D. B., Stefanik, C., Stuhlsatz, C., Kacirek, K., & Belfi, A. M. (2023). AI composer bias: Listeners like music less when they think it was composed by an AI. *Journal of Experimental Psychology: Applied*, 29(3), 676–692. <https://doi.org/10.1037/xap0000447>
- Sundar, S. S. (2020). Rise of machine agency: A framework for studying the psychology of human–AI interaction. *Journal of Computer-Mediated Communication*, 25(1), 74–88. <https://doi.org/10.1093/jcmc/zmz026>
- Wang, X., & Wu, Y. C. (2024). Balancing innovation and regulation in the age of generative artificial intelligence. *Journal of Information Policy*, 14, 1–27. <https://doi.org/10.5325/jinfopoli.14.2024.0012>

Wortel, C., Vanwesenbeeck, I., & Tomas, F. (2024). Made with artificial intelligence: The effect of artificial intelligence disclosures in Instagram advertisements on consumer attitudes. *Emerging Media*, 2(3), 547–570. <https://doi.org/10.1177/27523543241292096>

Fuentes de datos

Hartmann, J., Exner, Y., & Domdey, S. (2024). *GenImageNet* [Data set]. OSF. <https://osf.io/8ctjy/>

Wortel, C., Vanwesenbeeck, I., & Tomas, F. (2024). *Made with artificial intelligence: Data, analyses and supplemental materials* [Data set]. OSF. <https://osf.io/8fj pz/>

Anexos

No se incluyen anexos en el manuscrito editorial.

Información editorial

Campo / Field	Contenido / Content
Recepción / Received	15 de julio de 2025 / July 15, 2025
Revisión / Revised	14 de agosto de 2025 / August 14, 2025
Aceptación / Accepted	18 de septiembre de 2025 / September 18, 2025
Publicación / Published	30 de diciembre de 2025 / December 30, 2025
Volumen, número y año / Volume, issue and year	Vol. 1, Núm. 2, julio-diciembre de 2025

Licencia prevista: Creative Commons Atribución 4.0 Internacional (CC BY 4.0), salvo condiciones aplicables al material de terceros.