

Divulgación, credibilidad y vinculación parasocial ante contenido generado por inteligencia artificial (influencers virtuales y deepfakes): Una revisión sistemática rápida de sus efectos comunicativos y psicológicos en las audiencias

Disclosure, Credibility, and Parasocial Bonding with AI-Generated Content (Virtual Influencers and Deepfakes): A Rapid Systematic Review of Communicative and Psychological Effects on Audiences

Divulgazione, credibilità e legame parasociale di fronte a contenuti generati dall'intelligenza artificiale (influencer virtuali e deepfake): una revisione sistematica rapida dei loro effetti comunicativi e psicologici sul pubblico

Luis Arturo Rosero Constante ^I

lroseroc@unemi.edu.ec

<https://orcid.org/0000-0002-5590-7481>

Edison David Andrade-Sánchez ^{II}

eandrades@unemi.edu.ec

<https://orcid.org/0000-0001-8385-8034>

Kevin Acosta-Barreno ^{III}

jacostab@unemi.edu.ec

<https://orcid.org/0000-0002-7062-5773>

Iván Leonardo Pincay Aguilar ^{IV}

ipincaya@unemi.edu.ec

<https://orcid.org/0000-0002-9093-7838>

Correspondencia: lroseroc@unemi.edu.ec

Artículo de Investigación

*** Recibido:** 26 de julio de 2026 ***Aceptado:** 28 de agosto de 2026 *** Publicado:** 22 de septiembre de 2026

- I. Universidad Estatal de Milagro (UNEMI), Milagro, Ecuador.
- II. Universidad Estatal de Milagro (UNEMI), Milagro, Ecuador.
- III. Universidad Estatal de Milagro (UNEMI), Milagro, Ecuador.
- IV. Universidad Estatal de Milagro (UNEMI), Milagro, Ecuador.

Resumen

La proliferación de personas sintéticas generadas por IA —influencers virtuales y deepfakes— genera efectos comunicativos (credibilidad, confianza, divulgación) y psicológicos (bienestar, autoestima, ansiedad) aún no integrados en una síntesis única. El objetivo de este estudio fue sintetizar la evidencia empírica sobre cómo la divulgación, la credibilidad de la fuente y la interacción parasocial ante contenido generado por IA se relacionan con efectos comunicativos y psicológicos en las audiencias. Se realizó una revisión sistemática rápida (rapid review) siguiendo PRISMA 2020, con búsquedas en Scopus, SciELO y Redalyc (2016-2026); tras cribar 553 registros únicos, 257 estudios resultaron potencialmente elegibles, y se evaluó a texto completo una muestra representativa (n=49), estratificada por diseño y eje temático, cuya calidad metodológica se evaluó con el MMAT (2018). De los 49 estudios evaluados, 46 se confirmaron como estudios primarios elegibles tras excluir un estudio teórico y dos actas de congreso que no cumplían el criterio de tipo de documento declarado; uno adicional no pudo verificarse por un error de documentación, por lo que la síntesis se basa en n=45 (76% cuantitativos, 16% cualitativos, 9% mixtos). Los hallazgos se organizan en tres ejes: (1) la divulgación reduce la confianza/antropomorfismo percibido, moderada por el tipo de persona sintética y el momento de divulgación; (2) la interacción parasocial ocurre por rutas afectivas más que cognitivas; y (3) hay evidencia de un "lado oscuro" psicológico —comparación social dañina, amenaza a la identidad humana y, en casos extremos, daño psicológico agudo por deepfakes sexualizados. Se concluye que la divulgación transparente es indispensable, aunque no suficiente, para proteger el bienestar psicológico, pues los mecanismos comunicativos que generan valor comercial suelen ser los mismos que generan riesgo psicológico.

Palabras clave: influencers virtuales; deepfakes; divulgación; credibilidad de la fuente; interacción parasocial; bienestar psicológico; revisión sistemática; PRISMA

Abstract

The proliferation of AI-generated synthetic personas —virtual influencers and deepfakes— produces both communicative effects (credibility, trust, disclosure) and psychological effects (wellbeing, self-esteem, anxiety) that have not yet been integrated into a single evidence synthesis.

This study aimed to synthesize empirical evidence on how disclosure, source credibility, and parasocial interaction with AI-generated content relate to communicative and psychological effects on audiences. A rapid systematic review following PRISMA 2020 was conducted, searching Scopus, SciELO, and Redalyc (2016-2026); after screening 553 unique records, 257 studies were considered potentially eligible, and a representative sample (n=49) was assessed at full text, stratified by design and thematic axis, with methodological quality assessed using the MMAT (2018). Of the 49 studies assessed, 46 were confirmed as eligible primary studies after excluding one theoretical study and two conference proceedings that did not meet the declared document-type criterion; one additional study could not be verified due to a documentation error, so the synthesis is based on n=45 (76% quantitative, 16% qualitative, 9% mixed methods). Findings are organized around three axes: (1) disclosure reduces perceived trust/anthropomorphism, moderated by the type of synthetic persona and the timing of disclosure; (2) parasocial interaction with synthetic personas occurs through affective rather than cognitive routes; and (3) there is consistent evidence of a psychological "dark side" —harmful social comparison, threats to human identity, and, in extreme cases, acute psychological harm from sexualized deepfakes. It is concluded that transparent disclosure is indispensable, though not sufficient, to protect audiences' psychological wellbeing, since the communicative mechanisms that generate commercial value are frequently the same ones that generate psychological risk.

Keywords: virtual influencers; deepfakes; disclosure; source credibility; parasocial interaction; psychological wellbeing; systematic review; PRISMA

Riassunto

La proliferazione di persone sintetiche generate dall'IA —influencer virtuali e deepfake— produce effetti comunicativi (credibilità, fiducia, divulgazione) e psicologici (benessere, autostima, ansia) non ancora integrati in un'unica sintesi delle evidenze. L'obiettivo di questo studio è stato sintetizzare le evidenze empiriche su come la divulgazione, la credibilità della fonte e l'interazione parasociale con contenuti generati dall'IA siano associate a effetti comunicativi e psicologici sul pubblico. È stata condotta una revisione sistematica rapida (rapid review) secondo PRISMA 2020, con ricerche in Scopus, SciELO e Redalyc (2016–2026); dopo lo screening di 553 record unici,

257 studi sono risultati potenzialmente eleggibili ed è stato valutato a testo completo un campione rappresentativo (n=49), stratificato per disegno e asse tematico, la cui qualità metodologica è stata valutata mediante il MMAT (2018). Dei 49 studi valutati, 46 sono stati confermati come studi primari eleggibili dopo l'esclusione di uno studio teorico e di due atti di convegno che non soddisfacevano il criterio dichiarato relativo al tipo di documento; un ulteriore studio non ha potuto essere verificato a causa di un errore di documentazione, pertanto la sintesi si basa su n=45 (76% quantitativi, 16% qualitativi, 9% metodi misti). I risultati sono organizzati in tre assi: (1) la divulgazione riduce la fiducia/l'antropomorfismo percepito, con un effetto moderato dal tipo di persona sintetica e dal momento della divulgazione; (2) l'interazione parasociale avviene attraverso percorsi affettivi più che cognitivi; e (3) esistono evidenze coerenti di un "lato oscuro" psicologico —confronto sociale dannoso, minaccia all'identità umana e, nei casi estremi, danno psicologico acuto dovuto a deepfake sessualizzati. Si conclude che una divulgazione trasparente è indispensabile, sebbene non sufficiente, per proteggere il benessere psicologico, poiché i meccanismi comunicativi che generano valore commerciale sono spesso gli stessi che generano rischio psicologico.

Parole chiave: influencer virtuali; deepfake; divulgazione; credibilità della fonte; interazione parasociale; benessere psicologico; revisione sistematica; PRISMA

Introducción

Contexto y fenómeno de estudio. El auge de la inteligencia artificial generativa ha popularizado dos fenómenos comunicativos estrechamente relacionados: los influencers virtuales (avatares completamente sintéticos que actúan como creadores de contenido en redes sociales) y los deepfakes (contenido audiovisual que altera o sustituye la identidad de personas reales mediante IA). Ambos fenómenos comparten una característica central para la comunicación social: presentan a las audiencias una persona que no es lo que aparenta ser, planteando preguntas fundamentales sobre credibilidad de la fuente, divulgación, confianza y sus consecuencias psicológicas sobre quienes los consumen. Estos fenómenos se han documentado en dominios de aplicación tan diversos como la moda (Fajardo Rodríguez-Borlado y Pérez-Curiel, 2024) y la

esfera cívica, donde investigación reciente ha examinado la percepción ciudadana de deepfakes en contextos electorales latinoamericanos (Vélez Bermello et al., 2025).

Aunque conceptualmente distintos —los influencers virtuales son sintéticos por diseño y explícitos en su naturaleza artificial (al menos potencialmente), mientras que los deepfakes suelen implicar una sustitución no consentida de una persona real—, ambos fenómenos activan mecanismos comunicativos y psicológicos comparables: la interacción parasocial, la evaluación de credibilidad de la fuente, y la respuesta emocional a la incertidumbre sobre la autenticidad.

Marco teórico. Este trabajo se sitúa en la intersección de tres tradiciones teóricas de la comunicación social: la teoría de la credibilidad de la fuente (source credibility theory), que explica cómo los receptores evalúan la fiabilidad y pericia de quien emite un mensaje; la teoría de la interacción parasocial (Horton y Wohl, 1956), que describe el vínculo unidireccional que las audiencias desarrollan con figuras mediáticas; y la investigación sobre divulgación (disclosure) de contenido patrocinado o generado por IA, un área de creciente relevancia regulatoria. La psicología aporta la capa de consecuencias medibles: bienestar subjetivo, autoestima, comparación social y respuestas emocionales ante el malestar generado por estímulos "casi humanos" (valle inquietante).

Justificación de la revisión y vacío en la literatura. Una búsqueda preliminar de revisiones existentes identificó al menos tres revisiones sistemáticas o meta-análisis publicados recientemente (2023-2026) sobre aspectos parciales de este fenómeno. Sin embargo, ninguna integra los tres elementos que esta revisión aborda simultáneamente: (a) ambos fenómenos —deepfakes e influencers virtuales— como manifestaciones relacionadas de personas sintéticas generadas por IA; (b) el eje comunicativo (divulgación, credibilidad, parasocialidad) como marco teórico central, con la psicología como capa de efectos derivados, en lugar de aproximaciones puramente de marketing o puramente político-electorales; y (c) diseños empíricos de cualquier tipo —cuantitativos, cualitativos y mixtos— con cobertura de bases de datos que incluyen producción académica latinoamericana (SciELO, Redalyc), ausente en las revisiones previas centradas exclusivamente en Scopus/Web of Science.

Específicamente, la revisión de Byun y Ahn (2023) y la de Laszkiewicz y Kalińska-Kuła (2023) se centran en influencers virtuales exclusivamente desde una perspectiva de marketing, con outcomes limitados a actitud de marca e intención de compra. Por su parte, el meta-análisis de Kang y Valadez (2025) sobre impacto cognitivo, emocional y conductual de deepfakes se restringe a 24 estudios puramente experimentales cuantitativos, excluyendo diseños cualitativos y mixtos, y no aborda influencers virtuales.

Preguntas de investigación. Esta revisión se organiza en torno a tres preguntas:

- RQ1 (comunicación — credibilidad/confianza): ¿Cómo afecta la divulgación explícita —o la ausencia de divulgación— de que un contenido de influencer virtual o deepfake fue generado por IA a la credibilidad percibida de la fuente y la confianza de la audiencia?
- RQ2 (comunicación — vínculo parasocial): ¿Qué papel desempeña la interacción parasocial en la relación entre la exposición a personas sintéticas generadas por IA y las actitudes de la audiencia hacia el mensaje, la marca o la fuente?
- RQ3 (psicológica derivada): ¿Qué efectos psicológicos sobre el bienestar de la audiencia se derivan de estos procesos comunicativos de exposición a contenido sintético generado por IA?

Metodología

Esta revisión sistemática se desarrolló siguiendo la declaración PRISMA 2020 (Page et al., 2021). El protocolo no fue registrado formalmente en PROSPERO, dado que dicho registro no constituye un requisito estándar en revisiones sistemáticas del área de ciencias sociales y comunicación, a diferencia de las ciencias de la salud.

Pregunta de revisión y marco PECO. Los elementos del marco PECO utilizado se presentan en la Tabla 1.

Tabla 1.
Marco PECO de la revisión

Elemento	Definición
Población	Audiencias/usuarios de medios digitales expuestos a contenido con personas sintéticas generadas por IA
Exposición	Contenido de influencers virtuales o deepfakes, diferenciando si la naturaleza sintética se divulga explícitamente o no

Comparación	Se aceptaron dos comparadores alternativos, no simultáneamente requeridos: (a) contenido de personas humanas reales; o (b) el mismo contenido generado por IA con divulgación vs. sin divulgación de su naturaleza sintética.
Outcome primario	Credibilidad de la fuente, confianza, interacción parasocial
Outcome secundario	Actitud hacia el mensaje/marca, percepción de autenticidad, bienestar psicológico, autoestima, malestar (valle inquietante), ansiedad

Estrategia de búsqueda. Se realizaron búsquedas sistemáticas en Scopus, SciELO y Redalyc, ejecutadas el 5 de septiembre de 2026, restringidas a artículos de revista (Document Type = Article) y al periodo 2016-2026. Las ecuaciones de búsqueda completas, calibradas iterativamente para cada motor, se detallan en el material suplementario.

Criterios de elegibilidad. Inclusión: estudios empíricos primarios (cuantitativos, cualitativos o mixtos) con participantes humanos como audiencia/usuarios, centrados en influencers virtuales y/o deepfakes, que reportaran al menos un resultado de comunicación y/o un resultado psicológico. Exclusión: estudios puramente técnicos de detección algorítmica sin medición de respuesta humana; estudios centrados exclusivamente en métricas de marketing sin ningún constructo de comunicación o psicológico; revisiones, editoriales y ensayos conceptuales sin datos primarios; deepfakes exclusivamente en el dominio político-electoral sin medición de credibilidad/parasocialidad/bienestar.

Selección de estudios. Tras eliminar 16 registros duplicados de un total de 569 registros identificados, se cribaron 553 registros únicos por título y resumen, resultando en 257 estudios elegibles para evaluación de texto completo.

Estrategia de muestreo para la extracción de datos. Dado el volumen del corpus preseleccionado (257 registros considerados potencialmente elegibles) y su crecimiento exponencial reciente (75% publicados en 2025-2026), se adoptó un diseño de revisión rápida (rapid review; Tricco et al., 2015; Garritty et al., 2021). En lugar de evaluar la totalidad de los 257 registros, se evaluó a texto completo una muestra representativa (n=49). De estos, 46 se confirmaron como estudios primarios elegibles: se excluyó 1 estudio de naturaleza teórica sin datos empíricos primarios, y 2 estudios correspondientes a actas de congreso que no cumplieran el criterio de tipo de documento (Document Type = Article) declarado en la estrategia de búsqueda. El

registro de extracción de uno de estos 46 estudios no pudo verificarse debido a un error de documentación durante la compilación de la tabla de extracción, por lo que la síntesis final se basa en n=45 estudios con datos completos. Los 208 registros restantes del corpus preseleccionado no fueron evaluados a texto completo y, por tanto, no forman parte de la síntesis de esta revisión. La muestra se conformó priorizando representación proporcional de diseños metodológicos y cobertura de los ejes temáticos centrales identificados. La distribución de diseños de la muestra final (76% cuantitativo, 16% cualitativo, 9% mixto; n=45) es sustancialmente equivalente a la del corpus preseleccionado completo (76%, 9%, 11% respectivamente), respaldando su representatividad dentro de los límites de este diseño.

Evaluación de calidad metodológica. La calidad metodológica se evaluó con el Mixed Methods Appraisal Tool (MMAT, versión 2018; Hong et al., 2018), aplicando los criterios correspondientes a cada categoría de diseño (cualitativo, cuantitativo no-aleatorizado, o mixto) tras confirmar la categoría de diseño de cada estudio mediante lectura del texto completo o, cuando no fue posible, del resumen extendido disponible.

Síntesis de datos. Se realizó síntesis narrativa temática organizada según las tres preguntas de investigación, dada la heterogeneidad de diseños, constructos y medidas que impide un meta-análisis cuantitativo agregado en esta etapa.

Resultados

Selección de estudios. De 569 registros identificados (Scopus n=503, SciELO n=16, Redalyc n=50), se eliminaron 16 duplicados, resultando en 553 registros únicos cribados. Tras el cribado por título/resumen, 257 estudios fueron considerados potencialmente elegibles y avanzaron a evaluación de texto completo. Dado el volumen de este corpus preseleccionado, se evaluó a texto completo una muestra representativa de 49 estudios (rapid review). De estos, 46 se confirmaron como estudios primarios elegibles para la extracción de datos: uno, sobre el "problema del médico deepfake", fue reclasificado como excluido al verificar que constituía un artículo teórico/argumentativo sin datos primarios propios; y dos correspondían a actas de congreso que no cumplían el criterio de tipo de documento (Document Type = Article) declarado en la estrategia de búsqueda, por lo que también fueron excluidos. El registro de extracción de uno de estos 46

estudios no pudo verificarse debido a un error de documentación durante la compilación de la tabla de datos; en consecuencia, la síntesis que sigue se basa en n=45 estudios con datos completos. Los 208 registros restantes del corpus preseleccionado no fueron evaluados a texto completo y no forman parte de la síntesis de esta revisión.

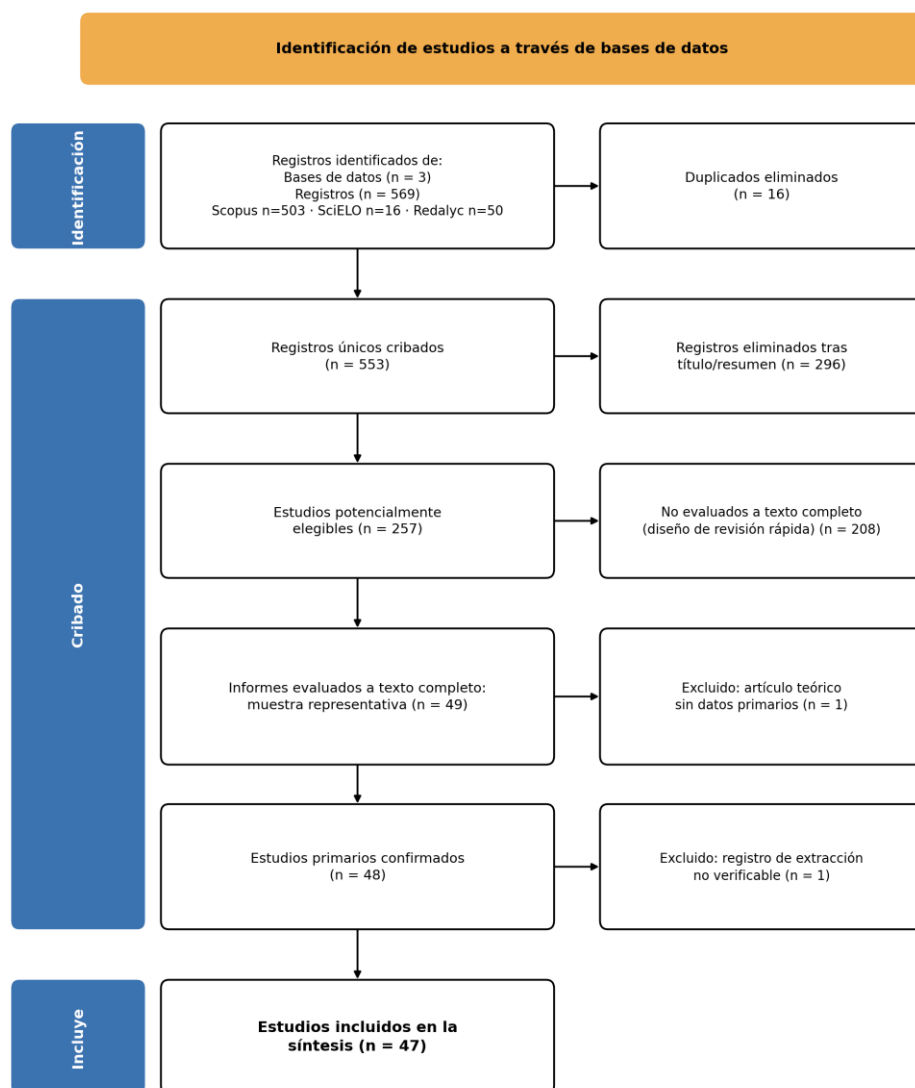


Figura 1. Diagrama de flujo PRISMA 2020 del proceso de selección de estudios.

Características de los estudios incluidos. De los 45 estudios de la muestra final, 34 (76%) emplearon diseños cuantitativos (predominantemente experimentales y encuestas transversales

con modelado de ecuaciones estructurales), 7 (16%) diseños cualitativos (entrevistas, grupos focales, análisis de casos), y 4 (9%) diseños mixtos. Los estudios cubrieron los ocho ejes temáticos centrales identificados: divulgación, credibilidad, interacción parasocial, bienestar/autoestima/ansiedad, antropomorfismo, autenticidad, valle inquietante, y daño psicológico/violencia digital. La Tabla 2 presenta un estudio ilustrativo por cada eje temático, seleccionado por su relevancia para la síntesis narrativa.

Tabla 2.
Estudios ilustrativos por eje temático (n=45)

Eje temático	Estudio(s) ilustrativo(s)	Hallazgo principal
Divulgación	Muniz et al. (2024); Whittaker et al. (2026)	Divulgar reduce confianza/antropomorfismo; hasta 42% de audiencias no reconoce la naturaleza no-humana pese a la divulgación
Credibilidad de la fuente	Jha et al. (2025); Chen et al. (2025)	La credibilidad predice confianza y compra; puede reducirse si se resalta la naturaleza artificial
Interacción parasocial	Bang y Han (2025); Stein et al. (2024)	La identificación afectiva predice el vínculo parasocial más fuertemente que la cognitiva
Bienestar / autoestima / ansiedad	Barari et al. (2025); Sattar et al. (2025)	El antropomorfismo alto reduce el bienestar vía comparación social ascendente
Antropomorfismo	Liu y Wang (2025); Su et al. (2026)	Modelo curvilíneo antropomorfismo-confiabilidad en el valle inquietante
Autenticidad	Liu y Lee (2024); Lee et al. (2025)	Tensión: los VI se perciben como menos auténticos en general, pero más "objetivos" bajo la heurística de máquina
Valle inquietante	Luo y Kim (2024)	Análisis de construcción de identidad de un VI real (Ling) y su relación con el valle inquietante
Daño psicológico / violencia digital	Shepherd et al. (2026); Dalouh Ounia et al. (2025)	Deepfakes sexualizados generan asco/ansiedad/ira; violencia digital de pareja en adolescentes

Calidad metodológica. La evaluación de calidad metodológica mediante el MMAT reveló un patrón consistente de fortalezas y debilidades a través de los 45 estudios de la muestra final. Entre las fortalezas más frecuentes, un número considerable de estudios cuantitativos (particularmente aquellos publicados en revistas de marketing y psicología del consumidor) emplearon escalas validadas psicométricamente para sus constructos centrales (credibilidad de fuente, interacción parasocial, antropomorfismo), reportando índices de confiabilidad adecuados (alfa de Cronbach, AVE, HTMT). Asimismo, una proporción relevante de los estudios (aproximadamente un tercio)

utilizó diseños experimentales con manipulación y, en varios casos, asignación aleatoria a condiciones (p. ej., Barari et al., 2025; Deng et al., 2026), lo que fortalece la validez interna para establecer relaciones causales entre exposición a personas sintéticas y sus efectos. Dos estudios destacaron por su transparencia en el manejo de datos, poniendo sus bases de datos a disposición pública (Forrai y Balaban, 2026, vía Open Science Framework) o declarando explícitamente su disponibilidad bajo solicitud (Cheng et al., 2025).

Entre las debilidades más recurrentes, la gran mayoría de los estudios cuantitativos (más del 80% de los que reportaron información suficiente para evaluar este criterio) emplearon muestreo no probabilístico o de conveniencia, lo cual limita la representatividad de los hallazgos frente a las poblaciones objetivo declaradas y constituye la limitación de calidad más generalizada del corpus. La integridad de los datos (manejo de valores perdidos, tasas de respuesta) no pudo determinarse en la mayoría de los estudios por ausencia de esta información en los resúmenes disponibles, lo que sugiere una oportunidad de mejora en la transparencia reportada por la disciplina. De manera similar, el control sistemático de variables de confusión sociodemográficas fue infrecuente fuera de los diseños experimentales puros, presente principalmente en estudios que emplearon modelos de regresión múltiple con covariables explícitas. Finalmente, cabe señalar que la evaluación de un subconjunto de estudios (aproximadamente 15%) se realizó a partir de resúmenes extensos disponibles públicamente en lugar del texto completo, lo cual pudo limitar la valoración de criterios que típicamente se detallan en las secciones de método y resultados del texto íntegro (p. ej., procedimientos exactos de aleatorización, manejo de casos perdidos).

En conjunto, el perfil de calidad metodológica observado es consistente con el de un campo de investigación emergente y de rápido crecimiento: metodológicamente sofisticado en el diseño de instrumentos y manipulaciones experimentales, pero con limitaciones persistentes en representatividad muestral y transparencia en el reporte de datos faltantes.

RQ1 (Divulgación, credibilidad y confianza). La evidencia confirma consistentemente que divulgar el origen sintético de una persona generada por IA reduce la confianza, el antropomorfismo percibido y/o la actitud hacia la marca (Muniz et al., 2024; Bie et al., 2025; Gerlich, 2024). Sin embargo, este efecto está moderado por al menos tres factores: el tipo de

persona sintética —los influencers virtuales de estilo anime pueden beneficiarse de la divulgación por percepción de novedad, a diferencia de los de apariencia humana (Kim et al., 2025; Carrillo-Durán et al., 2024)—; el momento de la divulgación, con efectos diferenciados según se divulgue antes o después de la exposición al contenido (Whittaker et al., 2026); el tipo de narrativa que acompaña a la divulgación, dado que las narrativas emocionales pueden atenuar parcialmente su efecto negativo sobre la credibilidad y el vínculo parasocial (Lim y Lee, 2023); y la efectividad real de la divulgación, en línea con evidencia sobre riesgo y confianza en endosantes humanos frente a virtuales (Ameen et al., 2024) y sobre cómo el nivel de realismo generado por IA modula las respuestas de la audiencia (Sorosrungruang et al., 2024) —incluyendo aplicaciones que van desde el consumo juvenil en TikTok (Peemanee et al., 2026) hasta la promoción de comportamientos de salud mediante influencers CGI (Saumer et al., 2025)—, dado que hasta un 42% de las audiencias no reconoce la naturaleza no-humana de un influencer virtual aun cuando está declarada (Muniz et al., 2024).

Un hallazgo cualitativo central señala que, en audiencias adolescentes, la divulgación cumple una función protectora emocional más allá de su función informativa: su ausencia genera sentimientos de traición y malestar emocional (Mouritzen et al., 2026).

RQ2 (Interacción parasocial). Los usuarios desarrollan vínculos parasociales con personas sintéticas, aunque por rutas distintas a las de influencers humanos. La identificación afectiva predice la interacción parasocial más fuertemente que la identificación cognitiva, y este efecto es más pronunciado en influencers virtuales de apariencia humanoide, en línea con evidencia de que la expresión emocional en influencers virtuales humanoides impulsa el engagement del usuario mediante un modelo de empatía (Liu, Li y Liu, 2026) y de que las respuestas afectivas, actitudinales y conductuales de las audiencias ante las emociones expresadas por personas virtuales son consistentes con este patrón (Li, Ham y Eastin, 2024). La congruencia entre el tipo de endosante (humano o virtual) y el tipo de producto (físico o digital/NFT) determina la fuerza relativa del vínculo parasocial, sin que exista una ventaja universal de un tipo de endosante sobre otro (Yuan et al., 2025). El sentido de comunidad percibido en torno al influencer virtual también predice la interacción parasocial y el apego (Bang y Han, 2025), y estudios longitudinales mediante

experience sampling confirman que este vínculo se construye y se rompe de forma dinámica a lo largo del tiempo (Breves y van Berlo, 2025).

RQ3 (Efectos psicológicos derivados). La evidencia documenta consistentemente un patrón de "lado oscuro" psicológico: el antropomorfismo alto en influencers virtuales reduce el bienestar del seguidor mediante comparación social ascendente, incluso cuando aumenta el engagement comercial (Barari et al., 2025), un patrón que se explica en parte por dinámicas de poder social propias del marketing de influencers virtuales (Lu y Ali, 2026). Un hallazgo cualitativo original identifica una emoción específica de la comparación con influencers virtuales —vergüenza por compararse con una entidad no-humana— ausente en la literatura sobre comparación social con influencers humanos (Nasr et al., 2025). De forma convergente, la "perfección" visual de un influencer virtual (ausencia de imperfecciones humanas) genera actitudes de marca más negativas, mediadas por amenaza a la identidad humana del consumidor (Cheng et al., 2025). No obstante, el antropomorfismo no es un mecanismo uniformemente negativo: también se ha vinculado con intención prosocial (Su et al., 2026) y con aplicaciones orientadas al bien social, como la promoción de tolerancia interétnica a través de influencers virtuales (Levkov et al., 2026), lo que sugiere que su valencia psicológica depende del propósito comunicativo y no únicamente de la magnitud del rasgo antropomórfico.

Crucialmente, el reconocimiento explícito de que un contenido es generado por IA no protege por sí solo el bienestar: un experimento con datos abiertos mostró que dicho reconocimiento no previno la comparación social dañina, mientras que la alfabetización mediática sí se asoció positivamente con el bienestar (Forrai y Balaban, 2026).

En el extremo más severo del espectro de daño, evidencia experimental muestra que la exposición a deepfakes sexualizados genera emociones aversivas (asco, ansiedad, ira) en mujeres, aumentando su disposición a emprender acciones legales (Shepherd et al., 2026), en línea con evidencia cualitativa sobre violencia digital de pareja en adolescentes mediada por deepfakes (Dalouh Ounia et al., 2025), y con investigación que documenta un patrón más amplio de sexualización tanto de influencers humanos como virtuales (Fowler y Thomas, 2025). La percepción del daño varía además según la edad: estudios sobre el efecto de tercera persona

muestran que las personas mayores tienden a percibir un riesgo mayor para terceros que para sí mismas frente a los deepfakes, un patrón distinto al observado en adultos jóvenes (Huang et al., 2025), lo cual se asocia con un mayor respaldo a políticas de regulación (Vogler et al., 2025).

Discusión

Síntesis integradora. El patrón que emerge de esta revisión sugiere una tesis central: la comunicación transparente sobre la naturaleza sintética de una persona generada por IA es necesaria pero insuficiente para proteger el bienestar psicológico de las audiencias. Los mecanismos comunicativos que generan valor comercial —antropomorfismo, interacción parasocial, perfección visual— son frecuentemente los mismos mecanismos que generan riesgo psicológico (comparación social dañina, amenaza a la identidad humana, vergüenza). Esta tensión estructural entre efectividad comunicativa y bienestar de la audiencia no había sido articulada explícitamente en las revisiones previas. Estas abordan los fenómenos de forma compartimentada, ya sea desde el marketing o desde la política de desinformación.

Implicaciones teóricas. Los hallazgos de esta revisión invitan a extender, más que a descartar, los marcos teóricos clásicos de la comunicación social hacia el caso específico de las entidades sintéticas generadas por IA. La teoría de la interacción parasocial (Horton y Wohl, 1956), formulada originalmente para explicar el vínculo unidireccional entre audiencias y figuras mediáticas humanas, muestra en este corpus que opera de forma diferenciada ante entidades no humanas: la identificación afectiva predice la vinculación parasocial más fuertemente que la cognitiva, y el proceso completo de formación de la relación —desde el contacto inicial hasta la integración en la vida cotidiana del seguidor— no parece completarse de la misma manera en el caso de los influencers virtuales, cuyas etapas avanzadas de integración no resultaron empíricamente distinguibles en el único estudio longitudinal de la muestra (Breves y van Berlo, 2025). Esto sugiere que la parasocialidad hacia entidades sintéticas podría constituir una variante teórica distinta, y no una simple aplicación del modelo clásico, ameritando el desarrollo de escalas e indicadores específicos para su medición.

De modo similar, la teoría de la comparación social (Festinger, 1954) requiere una extensión conceptual para dar cuenta de un hallazgo que no tiene equivalente claro en la literatura sobre

comparación con influencers humanos: la vergüenza específica que experimentan los usuarios al reconocer que se están comparando con una entidad no-humana. Este hallazgo sugiere que la comparación social con estímulos sintéticos activa un componente metacognitivo adicional —la evaluación de la propia comparación como socialmente inadecuada o absurda— ausente en los modelos clásicos de comparación ascendente/descendente. En la misma línea, el hallazgo de que la "perfección" visual de un influencer virtual genera rechazo por amenaza a la identidad humana del consumidor conecta con la literatura sobre distintividad óptima (optimal distinctiveness theory) y sugiere que los seres sintéticos demasiado perfectos podrían estar transgrediendo una necesidad psicológica implícita de las audiencias humanas de mantener una frontera categorial clara entre lo humano y lo artificial.

La tensión encontrada entre estudios que muestran mayor confianza/autenticidad hacia influencers humanos y aquellos que muestran el patrón opuesto —mediado por la heurística de máquina, es decir, la creencia de que los sistemas automatizados son más objetivos e imparciales que los humanos (Sundar, 2008)— indica que la teoría de la credibilidad de la fuente necesita incorporar explícitamente esta heurística como un moderador de doble filo: puede incrementar la credibilidad percibida en contextos donde se valora la objetividad (recomendaciones basadas en datos), pero puede reducirla en contextos donde se valora la calidez y autenticidad humana (narrativas emocionales, causas sociales), tal como evidencian Lee et al. (2025) respecto a la autenticidad percibida de influencers virtuales frente a humanos, y como confirma evidencia adicional sobre la confiabilidad percibida de contenido deepfake (Kumar et al., 2026). Futuros modelos teóricos deberían especificar las condiciones de frontera que determinan cuándo domina uno u otro mecanismo. Esta tensión conecta, en un plano más amplio, con planteamientos recientes sobre la crisis de autenticidad digital que plantean los medios sintéticos en general (Verma, 2026).

Finalmente, el hallazgo de que el reconocimiento cognitivo de la naturaleza artificial de un influencer no previene por sí solo el daño psicológico derivado de la comparación social, mientras que la alfabetización mediática sí lo hace, tiene una implicación teórica importante para los modelos de divulgación basados en el conocimiento de persuasión (persuasion knowledge model; Friestad y Wright, 1994): el modelo asume que reconocer la intención persuasiva de un mensaje

activa mecanismos de defensa cognitiva que protegen al receptor, pero nuestros hallazgos sugieren que, en el caso de las entidades sintéticas, este reconocimiento cognitivo opera de forma disociada de los procesos afectivos de comparación social que subyacen al daño al bienestar. Esto respalda modelos duales (cognitivo-afectivos) de procesamiento de la persuasión por sobre modelos puramente cognitivos en este dominio específico.

Implicaciones prácticas y de política pública. Los hallazgos sugieren que las políticas de divulgación obligatoria de contenido generado por IA, aunque necesarias desde una perspectiva ética, son insuficientes como única medida de protección psicológica de las audiencias. La evidencia respalda la necesidad de programas de alfabetización mediática crítica específicos, más allá del simple etiquetado de contenido, especialmente para poblaciones vulnerables (adolescentes, adultos mayores), en línea con investigación sobre protección digital de la infancia y la adolescencia frente a la socialización mediada por dispositivos móviles (Establés et al., 2025).

Limitaciones. Esta revisión presenta las siguientes limitaciones:

- Debido al diseño de revisión rápida adoptado, el 81% del corpus preseleccionado (208 de 257 registros considerados potencialmente elegibles) no fue evaluado a texto completo. Esta revisión, por tanto, no constituye una síntesis exhaustiva del campo, sino un panorama representativo basado en una submuestra estratificada (n=45, tras evaluar 49 estudios, excluir 1 estudio no empírico, excluir 2 estudios cuyo tipo de documento no cumplía el criterio declarado, y descartar 1 adicional cuyo registro de extracción no pudo verificarse debido a un error de documentación); es posible que estudios con hallazgos divergentes no estén representados en la síntesis.
- El cribado y la evaluación de calidad fueron realizados inicialmente por un único revisor. Posteriormente, un segundo revisor realizó de forma independiente una ronda de verificación del cribado y la extracción de datos, completada antes del envío de este manuscrito.
- Para algunos estudios de la muestra, la evaluación se basó en resúmenes extensos disponibles públicamente en lugar del texto completo, lo cual pudo limitar la evaluación de algunos criterios de calidad metodológica.

- La heterogeneidad de diseños y medidas impidió la realización de un meta-análisis cuantitativo agregado.
- No se realizó una evaluación formal del riesgo de sesgo de publicación del corpus incluido; dada la heterogeneidad de diseños que impidió el metaanálisis, este riesgo no pudo cuantificarse y debe considerarse como limitación adicional al interpretar la ausencia de hallazgos nulos o contradictorios en la síntesis.

Conclusiones

Esta revisión sistemática integra por primera vez la evidencia sobre influencers virtuales y deepfakes bajo un marco comunicativo-psicológico unificado, evidenciando que la divulgación de la naturaleza sintética de una persona generada por IA, aunque necesaria, no es suficiente para proteger el bienestar psicológico de las audiencias frente a los riesgos asociados a la comparación social, la amenaza a la identidad humana, y en casos extremos, el daño psicológico agudo — un planteamiento que converge con el argumento conceptual de Kim y Han (2026) sobre la insuficiencia del etiquetado de expertos sintéticos para proteger la confianza del consumidor. Se requiere investigación futura que desarrolle y evalúe intervenciones de alfabetización mediática crítica específicamente diseñadas para este fenómeno, así como marcos regulatorios que consideren estos efectos psicológicos más allá de la sola exigencia de divulgación.

Referencias

- Ameen, N., Cheah, J-H., Ali, F., El-Manstrly, D., & Kulyciute, R. (2024). Risk, trust, and the roles of human versus virtual influencers. *Tourism Recreation Research*.
- Bang, J., & Han, S. (2025). "I like community more than Influencers": Unraveling the influence of followers on parasocial interaction and attachment with virtual influencers. *Computers in Human Behavior*, 173, 108796. <https://doi.org/10.1016/j.chb.2025.108796>
- Barari, M., Ross, M., & Azad Moghddam, H. (2025). Virtual influencers' anthropomorphism, consumer engagement and/or well-being. *Journal of Consumer Marketing*, 42(5), 673-688.
- Bie, Y., Tang, J., Miao, M., Sun, W., & Jiang, Y. (2025). Beyond the shores of reality: the effect of virtual influencers on luxury perception and its downstream consequence. *Psychology & Marketing*, 42, 2369-2387.

- Breves, P. L., & van Berlo, Z. M. C. (2025). Making and breaking parasocial relationships with human and virtual influencers: An experience sampling study. *Media Psychology*, 1-29. <https://doi.org/10.1080/15213269.2025.2558029>
- Byun, K. J., & Ahn, S. J. G. (2023). A systematic review of virtual influencers: Similarities and differences between human and virtual influencers in interactive advertising. *Journal of Interactive Advertising*, 23(4), 293-306.
- Carrillo-Durán, M. V., García García, M., & Corzo Cortés, L. (2024). Influencers virtuales de apariencia humana como forma de comunicación online: el caso de Lil Miquela y Lu do Magalu en Instagram. *Revista de Comunicación*, 23(1), 119-140. <https://doi.org/10.26441/RC23.1-2024-3453>
- Chen, H., Lou, C., Wang, Y., & Lee, Y. (2025). Social virtual influencer effectiveness: environmental factor and source trust. *Journal of Interactive Advertising*.
- Cheng, L.-K., Toungh, C.-L., & Ho, C. M. (2025). The danger of perfection: How consumer power diminishes brand attitudes toward flawless virtual influencers. *International Journal of Advertising*, 1-43. <https://doi.org/10.1080/02650487.2025.2576961>
- Dalouh Ounia, R., Soriano Ayala, E., & Caballero Cala, V. (2025). Deepfakes and Digital Violence in Adolescent Couples: Perceptions, Risks, and Prevention Strategies. *Pedagogika*, 160(4), 74-93. <https://doi.org/10.15823/p.2025.160.4>
- Deng, R., Ahmed, S., & Kang, H. (2026). Reconsidering deepfake priming: effects on belief and sharing across media platforms. *Online Information Review*.
- Establés, M.-J., Llovet, C., & Gallego Gómez, C. (2025). Teléfonos móviles, socialización sexual y educación mediática: percepciones y retos en la protección digital de la infancia y la adolescencia. *Revista de Comunicación*, 24(1), 127-182.
- Fajardo Rodríguez-Borlado, P., & Pérez-Curiel, C. (2024). Impacto de la inteligencia artificial en la moda: análisis de influencers digitales en las fashion weeks internacionales. *Universitas-XXI, Revista de Ciencias Sociales y Humanas*, 41, 75-99. <https://doi.org/10.17163/uni.n41.2024.03>
- Festinger, L. (1954). A theory of social comparison processes. *Human Relations*, 7(2), 117-140. <https://doi.org/10.1177/001872675400700202>
- Forrai, M., & Balaban, D. C. (2026). Disclosures and literacy as determinants of AI-influencer recognition and well-being. *Computers in Human Behavior*.
- Fowler, K., & Thomas, V. L. (2025). An investigation into the sexualization of human and virtual social media influencers. *European Journal of Marketing*, 59(5), 1318-1346.

- 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>
- Garritty, C., Gartlehner, G., Nussbaumer-Streit, B., King, V. J., Hamel, C., Kamel, C., Affengruber, L., & Stevens, A. (2021). Cochrane Rapid Reviews Methods Group offers evidence-informed guidance to conduct rapid reviews. *Journal of Clinical Epidemiology*, 130, 13-22. <https://doi.org/10.1016/j.jclinepi.2020.10.007>
- Gerlich, M. (2024). Societal perceptions and acceptance of virtual humans: trust and ethics across different contexts. *Social Sciences*, 13(10), 516.
- Hong, Q. N., Fàbregues, S., Bartlett, G., Boardman, F., Cargo, M., Dagenais, P., Gagnon, M.-P., Griffiths, F., Nicolau, B., O'Cathain, A., Rousseau, M.-C., Vedel, I., & Pluye, P. (2018). The Mixed Methods Appraisal Tool (MMAT) version 2018 for information professionals and researchers. *Education for Information*, 34(4), 285-291. <https://doi.org/10.3233/EFI-180221>
- Horton, D., & Wohl, R. R. (1956). Mass communication and para-social interaction: Observations on intimacy at a distance. *Psychiatry*, 19(3), 215-229. <https://doi.org/10.1080/00332747.1956.11023049>
- Huang, S., Huang, L., & Goh, D. H.-L. (2025). Third-person perception of deepfake harms: Comparing seniors and young adults. *Proceedings of the Association for Information Science and Technology*, 62, 920-925. <https://doi.org/10.1002/pr2.1314>
- Jha, S., Chaudhary, R., & Shukla, B. (2025). Trust dynamics of virtual influencers: exploring their influence on fashion purchase decisions. *Innovative Marketing*, 21(3), 226-236.
- Kang, S., & Valadez, K. (2025). Deepfakes' cognitive, emotional, and behavioral impact: A systematic review and meta-analysis of individual responses. *Journalism & Mass Communication Quarterly*, 102(4), 958-992. <https://doi.org/10.1177/10776990251357294>
- Kim, D., & Han, Y. (2026). The deepfake doctor problem: why labeling synthetic experts is necessary but insufficient for protecting consumer trust. *Frontiers in Communication*, 11, Article 1856259. <https://doi.org/10.3389/fcomm.2026.1856259>
- Kim, E. A., Shoenberger, H., Kim, D., Thorson, E., & Zihang, E. (2025). Novelty vs. trust in virtual influencers: exploring the effectiveness of human-like virtual influencers and anime-like virtual influencers. *International Journal of Advertising*, 44(3), 453-483. <https://doi.org/10.1080/02650487.2024.2386916>
- Kumar, V., Kumar, S., Chatterjee, S., Chaudhuri, R., & Gupta, S. (2026). Decoding deepfake perceptions: Trustworthiness, importance and behavioral intention. *Information Systems Frontiers*. <https://doi.org/10.1007/s10796-026-10765-9>

- Laszkiewicz, A., & Kalińska-Kuła, M. (2023). Virtual influencers as an emerging marketing theory: a systematic literature review. *International Journal of Consumer Studies*, 47(6), 2479-2494.
- Lee, H., Shin, M., Yang, J., & Chock, T. M. (2025). Virtual influencers vs. human influencers in the context of influencer marketing: The moderating role of machine heuristic on perceived authenticity of influencers. *International Journal of Human-Computer Interaction*, 41(10), 6029-6046. <https://doi.org/10.1080/10447318.2024.2374100>
- Lee, J., & John, K. K. (2025). Detecting deepfakes through emotion?: Facial expression and emotional contagion as dual indicators of deepfake credibility. *Applied Cognitive Psychology*, 39(6), e70141. <https://doi.org/10.1002/acp.70141>
- Levkov, N., Santa, M., & Kitanovikj, B. (2026). Exploring opportunities and challenges in promoting interethnic tolerance as a social good through virtual influencers. *Global Knowledge, Memory and Communication*, 75(7-8), 2546-2572. <https://doi.org/10.1108/GKMC-01-2024-0033>
- Li, S., Ham, J., & Eastin, M. S. (2024). Social media users' affective, attitudinal, and behavioral responses to virtual human emotions. *Telematics and Informatics*, 87, 102084.
- Lim, R. E., & Lee, S. Y. (2023). "You are a virtual influencer!": Understanding the impact of origin disclosure and emotional narratives on parasocial relationships and virtual influencer credibility. *Computers in Human Behavior*, 148, 107897. <https://doi.org/10.1016/j.chb.2023.107897>
- Liu, F., & Lee, Y. (2024). Virtually authentic: Examining the match-up hypothesis between human vs. virtual influencers and product types ("Pixels vs. people"). *Journal of Product & Brand Management*, 33(2), 287.
- Liu, F., & Wang, R. (2025). Fostering parasocial relationships with virtual influencers in the uncanny valley: Anthropomorphism, autonomy, and a multigroup comparison. *Journal of Business Research*, 186, 115024. <https://doi.org/10.1016/j.jbusres.2024.115024>
- Liu, R., Li, X., & Liu, S. (2026). How emotional expression in human-like virtual influencers drives user engagement: Empathy model and its antecedents. *Frontiers in Psychology*, 16, 1544037. <https://doi.org/10.3389/fpsyg.2025.1544037>
- Lu, Z., & Ali, K. S. (2026). Virtual influencer marketing and consumers' well-being: Understanding how social power works. *Acta Psychologica*, 268, 107251. <https://doi.org/10.1016/j.actpsy.2026.107251>
- Luo, L., & Kim, W. (2024). How virtual influencers' identities are shaped on Chinese social media: a case study of Ling. *Global Media and China*, 9(3), 325-343.

- Mouritzen, S. L. T., Pedersen, S., & Jacobsen, L. F. (2026). "Like a human – just digital": Adolescents' and parents' perceptions of virtual influencer marketing. *European Journal of Marketing*, 60(4), 891-931. <https://doi.org/10.1108/EJM-05-2024-0351>
- Muniz, F., Stewart, K., & Magalhães, L. (2024). Are they humans or are they robots? The effect of virtual influencer disclosure on brand trust. *Journal of Consumer Behaviour*, 23(3), 1234-1250.
- Nasr, L. I., Mousavi, S., & Michaelidou, N. (2025). Self-comparing with virtual influencers: effects on followers' wellbeing. *Psychology & Marketing*, 42(3), 780-798. <https://doi.org/10.1002/mar.22151>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., ... Moher, D. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, n71. <https://doi.org/10.1136/bmj.n71>
- Peemanee, J., Udomlarp, T., Weber, P., & Weerathana, R. (2026). The Role of Virtual and Human Influencer Characteristics in Shaping Gen Z Purchases on TikTok: Hybrid SEM-ANN Approach. *Journal of Theoretical and Applied Electronic Commerce Research*, 21(5), 150. <https://doi.org/10.3390/jtaer21050150>
- Sattar, A. A., Ul Haq, J., Saleem, A., & Ahmad, W. (2025). Shyness vs confidence: the effect of virtual influencer personalities on social media users' well-being and addiction. *Psychology & Marketing*, 42(4), 1088-1109.
- Saumer, M., Neureiter, A., Varga, É. M., Gataric, V., Liu, C. Y., & Matthes, J. (2025). Artificial presence, real-life influence? Effects of CGI influencers on young adults' health behavior intentions. *Cyberpsychology*, 19(2).
- Shepherd, L., Sopina, D., Dawson, J., Brown, G., & Paterson, J. (2026). Fake content, harmful emotions: The influence of being a victim-survivor of sexualised deepfakes on women's aversive emotions and desire for legal action. *Sex Roles*, 92, Article 46. <https://doi.org/10.1007/s11199-026-01677-8>
- Sorosrungruang, T., Ameen, N., & Hackley, C. (2024). How real is real enough? Unveiling the diverse power of generative AI-enabled virtual influencers and the dynamics of human responses. *Psychology & Marketing*, 41, 3124-3143.
- Stein, J.-P., Breves, P. L., & Anders, N. (2024). Parasocial interactions with real and virtual influencers: The role of perceived similarity and human-likeness. *New Media & Society*, 26(6), 3433-3453. <https://doi.org/10.1177/14614448221102900>

- Su, L., Duan, X., & Pan, X. (2026). From anthropomorphism to prosocial intention: a PLS-SEM and fsQCA study of virtual influencers. *Frontiers in Psychology*.
- Sundar, S. S. (2008). The MAIN model: A heuristic approach to understanding technology effects on credibility. In M. J. Metzger & A. J. Flanagin (Eds.), *Digital media, youth, and credibility* (pp. 73-100). MIT Press. <https://doi.org/10.1162/dmal.9780262562324.073>
- Tricco, A. C., Antony, J., Zarin, W., Striffler, L., Ghassemi, M., Ivory, J., Perrier, L., Hutton, B., Moher, D., & Straus, S. E. (2015). A scoping review of rapid review methods. *BMC Medicine*, 13, 224. <https://doi.org/10.1186/s12916-015-0465-6>
- Vélez Bermello, G. L., Reyes Andrade, J. J., & Parrales Saltos, M. E. (2025). Deepfakes en las elecciones ecuatorianas de 2025. Percepción ciudadana y desafíos. *Razón y Palabra*, 29(124), 10-26. <https://doi.org/10.26807/rp.v29i124.2308>
- Verma, A. (2026). Deepfakes and the crisis of digital authenticity: ethical challenges in the age of synthetic media. *Journal of Information, Communication and Ethics in Society*, 24(1), 59-76. <https://doi.org/10.1108/JICES-04-2025-0083>
- Vogler, D., Rauchfleisch, A., & de Seta, G. (2025). Support for deepfake regulation: The role of third-person perception, trust, and risk. *Studies in Communication and Media*, 14(4), 570-593. <https://doi.org/10.5771/2192-4007-2025-4-570>
- Whittaker, L., Mulcahy, R., Russell-Bennett, R., Letheren, K., & Kietzmann, J. (2026). Examining consumer appraisals of deepfake advertising and disclosure: show deepfakes as 'real life' or say they're 'just fantasy'? *Journal of Advertising Research*, 66(1), 113-134. <https://doi.org/10.1080/00218499.2025.2498830>
- Yuan, Y., Nguyen, M., Shao, W., & Zhang, Y. (2025). Digital congruence in influencer marketing: how NFTs shift the balance between human and AI-driven virtual endorsers through parasocial relationships. *Marketing Intelligence & Planning*. Advance online publication. <https://doi.org/10.1108/MIP-03-2025-0281>