

Verificar en tiempos de IA: hacia un nuevo derecho de rectificación y sus consecuencias para el periodismo

En la actualidad, la **verificación informativa** ha dejado de ser un desafío creciente del ecosistema digital, cuya respuesta dependía principalmente de las organizaciones de *fact-checking* y de los medios de información, para convertirse en un **problema estructural impulsado por la inteligencia artificial** y que debería obligar al periodismo a tomar posición.

RAÚL MAGALLÓN

1. El valor económico de la imprecisión y de la inexactitud

“Si tu madre te dice que te quiere, comprueba primero que es tu madre y no una inteligencia artificial (IA)”, podría ser la actualización del clásico aforismo periodístico. Hasta el 34% de las comprobaciones realizadas por *Efe Verifica* en el mes de marzo de este año fueron contenidos sintéticos¹.

Este proceso de normalización de la “economía de la información imprecisa”, donde puede ser más importante la

respuesta rápida y verosímil que la información veraz, hace necesario formular algunas cuestiones:

¿Qué papel deben asumir los medios de comunicación ante el creciente uso de la IA generativa y la proliferación de respuestas erróneas o sencillamente falsas? ¿Cómo deben orientar la relación con sus lectores y audiencias sobre esta cuestión?

En los últimos tiempos hemos visto cómo se está transformando la forma de buscar información en internet, pero

¹ Martín Lorenzo, R. (4 de mayo de 2026). “De la excepción a la rutina: cómo la IA se ha integrado en el ecosistema desinformador”. EFE Verifica. <https://verifica.efe.com/de-la-excepcion-a-la-rutina-como-se-integra-la-ia-desinformacion/>

también de ofrecerla y presentarla a los usuarios. Este cambio está afectando tanto a la sostenibilidad de los medios como a la credibilidad de la información que consumimos.

Más allá de que este tipo de estrategia empresarial de buscadores, asistentes y chatbots consiga que las personas dejen de pinchar en los enlaces de los medios y no incentive a los usuarios a profundizar en su búsqueda inicial, también están las consecuencias de que las audiencias se detengan en una respuesta automatizada e individualizada sin procesos de verificación humanos.

Actualmente, somos conscientes de que la precisión de las respuestas de los distintos asistentes y chatbots depende de la fiabilidad de las fuentes de las que se nutren, pero también de los posibles acuerdos entre las principales compañías de IA y los medios de comunicación.

Como sabemos, en 2024, Google empezó a posicionar una respuesta generada por IA como alternativa a la primera posición de sus búsquedas y ya empezamos a conocer sus primeros resultados. Un reciente análisis de las AI Overviews [vistas o resúmenes generados por la IA] reveló que estas eran precisas aproximadamente en nueve de cada diez ocasiones (Mickle *et al.*, 2026).

Podría considerarse que es nivel de

“acierto” bastante preciso, pero, ¿es suficiente para no distorsionar nuestro consumo de información? ¿Qué consecuencias tiene que los principios de precisión y exactitud dejen de ser los que definan esta nueva fase de la “economía de la atención”?

La realidad es que si las búsquedas anuales en Google alcanzan los 5 billones, este resultado nos indica que a nivel global se producen cientos de miles de imprecisiones cada minuto en sus respuestas.

Por lo tanto, y en un escenario en el que las respuestas de chatbots y buscadores pueden ser cada vez más imprecisas (sobre todo si se alimentan a partir de esos resultados inexactos), los “hechos alternativos” ofrecidos por la IA pueden empezar a ser más fáciles de encontrar en nuestras búsquedas que los “hechos verificados” por los medios de comunicación.

Todo ello, además, acelerado por el bloqueo reciente de muchos medios de comunicación al Wayback Machine para que con sus contenidos no se entrenen modelos de IA².

En esta línea, en octubre de 2025, se publicaba una investigación liderada por la BBC y la participación de RTVE que señalaba que hasta el 31% de las respuestas analizadas presentaba proble-

² Europa Press. PortalTIC (30 de enero de 2026). “Medios de comunicación bloquean el acceso de Internet Archive para evitar el ‘web scrapping’ destinado a entrenar IA”. <https://www.europapress.es/portaltic/sector/noticia-medios-comunicacion-bloquean-acceso-internet-archive-evitar-web-scrapping-destinado-entrenar-ia-20260130104334.html>

mas significativos a la hora de atribuir las fuentes y, en términos de precisión, hasta el 20% de las respuestas de las IA presentaban alucinaciones (como inventarse una respuesta o un dato) o detalles que estaban desactualizados (EBU, 2025).

La precisión de los asistentes depende de la fiabilidad de las fuentes de las que se nutren, pero también de los acuerdos entre compañías de IA y los medios

Como recordaba David Corral Hernández (2025), “para obtener los resultados se había dado acceso temporal al sitio web de la BBC a los cuatro asistentes de IA más destacados y empleados por el público: ChatGPT de OpenAI, Copilot de Microsoft, Gemini de Google y Perplexity. A los cuatro se les hicieron las mismas 100 preguntas sobre determinadas noticias, pidiéndoles que utilizaran artículos de BBC News como fuentes, siempre que fuera posible”.

2. La necesidad del periodismo de posicionarse ante los ‘hechos alternativos’ generados por la IA

En la mayoría de las ocasiones, al hablar de una posible integración de equipos

de verificación en los medios españoles, pensamos en las grandes cabeceras y grupos de comunicación. Sin embargo, y como bien sabemos, la cartografía de medios en España se aproxima a los 3.000 medios digitales (Salaverría, 2023) con un ecosistema fragmentado y diferentes modelos de negocio.

Se trata de un ecosistema informativo nacional, regional y local, pero también especializado, en el que la (des)confianza en el mismo por parte de la ciudadanía depende del nivel de confianza de esta en las redes sociales o los chatbots, pero también del nivel de independencia y rigor percibido sobre la competencia y los medios que no consumimos.

En este escenario, y más allá de establecer recomendaciones o juicios sobre la necesidad de tener o no un departamento de documentación, corrección o verificación (en algunas ocasiones esta verificación está vinculada actualmente a los departamentos de fotografía o al defensor del lector), puede resultar más pertinente conocer si las organizaciones periodísticas españolas pueden responder a las siguientes preguntas y, si no es así, qué medidas deberían tomar internamente para hacerlo:

- ¿Tiene mi medio los protocolos y personal adecuados para la corrección de textos, fechas o nombres de las noticias que se publican?
- ¿Tiene mi organización mecanismos para verificar imágenes, audios o vídeos creados con IA o manipulados?

- ¿Está preparada mi redacción para afrontar el creciente aumento de informaciones erróneas o imprecisas que se encuentran en internet?
- ¿Tiene mi equipo de “redes” o de corresponsales la formación necesaria para verificar contenidos dudosos o claramente engañosos?
- ¿Hay alguna sección particularmente sensible al aumento de la desinformación?
- ¿Tiene mi medio de comunicación un código ético que responda al uso de la IA en la redacción?

La capacidad de responder afirmativamente a estas preguntas no solo indicará una ventaja competitiva frente a los chatbots y asistentes de IA en relación con la confianza ciudadana en los resultados que pueden ofrecer, sino también en cuanto al proceso a través del cual se han alcanzado esos resultados.

3. Hacia un posible derecho de rectificación ante la IA. El dilema de cómo ha de responder el periodismo

Mientras que el periodismo tiene como objetivo transmitir los hechos de manera precisa y veraz al mayor número de personas posible, actualmente las respuestas generadas por la IA no solo pueden ser personalizadas, sino que no existen mecanismos públicos y de rendición

de cuentas externos para saber a cuántos usuarios puede llegar una respuesta errónea o diferenciada.

Por lo tanto, ¿qué ocurre si las audiencias deciden antes preguntarle a una IA que acudir a un medio para informarse o comprobar posibles novedades, hechos o acontecimientos imprevistos? ¿Puede acabar afectando la imprecisión de las respuestas de los asistentes de IA a la credibilidad de los medios?

No existen mecanismos públicos y de redención de cuentas externos para saber a cuántos usuarios puede llegar una respuesta errónea o diferenciada

El riesgo evidente de que las alucinaciones puedan generar confusión y decisiones equivocadas, con consecuencias reales, aparece sobre todo cuando hacemos referencia a la información de última hora.

Al inicio de la guerra en Irán, el chatbot de X indicaba a los usuarios que le consultaban por las imágenes del conflicto que estas eran realmente imágenes de 2021 en Kabul³, poniendo en cuestión la propia verificación de medios como la Agencia EFE.

³ Ocaña, J. (1 de marzo de 2026). “No es Kabul en 2021, como dice Grok, estas imágenes son del ataque a Irán este 28 de febrero”. EFE Verifica. <https://verifica.efe.com/imagenes-ataque-escuela-iran-no-kabul-como-dice-grok/>

En otra ocasión, incluso respondió a los usuarios de forma individualizada que se trataba de un vídeo de las fallas de Valencia⁴.



En la siguiente imagen se muestra otro de los tuits de respuesta de Grok⁵:



Si algo hemos aprendido en los últimos tiempos es que la calidad de las respuestas que ofrece la IA depende sobre todo de la calidad de las preguntas que seamos capaces de formular. Por lo tanto, y para saber el impacto de una

respuesta errónea, deberíamos poder conocer el resultado de las siguientes preguntas: ¿Cuántos usuarios recibieron una respuesta individualizada? ¿A cuántas personas alcanzó la afirmación falsa? ¿A cuántos usuarios les llegó la rectificación del propio Grok?

Si bien el derecho a rectificación frente a una IA no es algo que esté consolidado jurídica y socialmente, el debate sobre su responsabilidad y los ejemplos que se repiten son cada vez más recurrentes.

En el caso español, una posible respuesta legislativa a las “alucinaciones informativas” estará determinada tanto por la Digital Service Act como por la AI Act, aprobadas por la Unión Europea, pero también por la necesidad de establecer un marco jurídico que actualice los retos informativos que ya estamos identificando.

Un debate que coincide en el tiempo con otro que ya debería estar superado. El 30 de enero de 2026 se presentó en el Congreso de los Diputados el Proyecto de Ley Orgánica reguladora del derecho de rectificación que actualiza el vigente de 1984⁶.

⁴ EFE Verifica [@EFEVerifica]. (12 de marzo de 2026). Ni son las Fallas de Valencia, como afirma Grok; (...) [Post]. <https://x.com/EFEVerifica/status/2031995823546929225>

⁵ Grok [@grok]. (7 de marzo de 2026). No, este vídeo no es de Israel. Es de Valencia, España, (...) [Post] <https://x.com/grok/status/2030301500895826381>

⁶ En este proyecto, entre otras cuestiones, se incorporan a la norma “usuarios de especial relevancia” como *influencers* y creadores de contenido, pero no se recoge nada sobre la posibilidad de un derecho de rectificación relacionado con los escenarios descritos. Véase: Gobierno de España. (2026, 30 de enero). *Proyecto de Ley Orgánica reguladora del derecho de rectificación*. Boletín Oficial de las Cortes Generales: Congreso de los Diputados, Serie A (83-1), 1–11. https://www.congreso.es/public_oficiales/L15/CONG/BOCG/A/BOCG-15-A-83-1.PDF

En este sentido, es posible que sea necesario iniciar de forma urgente y coordinada una reflexión sobre si puede establecerse algún tipo de norma que regule el derecho de rectificación ante las respuestas erróneas o imprecisas de un asistente de IA y quién y cómo podría ejercerlo.

La calidad de las respuestas que ofrece la IA depende sobre todo de la calidad de las preguntas que seamos capaces de formular

Puede que estas hipótesis y escenarios parezcan preguntas teóricas o deontológicas alejadas de la realidad del mercado, de las necesidades impuestas por la última hora o del ecosistema digital actual, pero el marco de discusión que se está consolidando nos permite pensar que las “alucinaciones” pueden acabar teniendo un valor económico tan rentable como el *clickbait* o la polarización.

Estamos en un momento en el que los medios tienen que decidir si pueden seguir abanderando la defensa de la información veraz ante la sociedad y quién o qué puede socavar sus principios deon-

tológicos y metodológicos (pero también los económicos) y su función social. ■

Referencias

- Corral Hernández, D. (2025). “Integridad de las noticias en los asistentes de IA: Un estudio internacional de medios públicos. La participación de RTVE”. *Cuadernos de Periodistas*, 51, 25–33. https://www.cuadernosdeperiodistas.com/wp-content/uploads/sites/2/2025/12/25_34-David-Corral.pdf
- European Broadcasting Union. (Octubre de 2025). *News integrity in AI assistants: An international PSM study*. BBC y EBU. https://www.ebu.ch/files/live/sites/ebu/files/Publications/MIS/open/EBU-MIS-BBC_News_Integrity_in_AI_Assistants_Report_2025.pdf
- Mickle, T., Metz C, Freedman D, Mondría T y Collins K. (7 de abril de 2026). “How Accurate Are Google’s A.I. Overviews?”. *The New York Times*. <https://www.nytimes.com/2026/04/07/technology/google-ai-overviews-accuracy.html>
- Salaverría, R. (2023). “Miles de medios ‘online’ en un mapa: cartografía ibérica del periodismo digital”. *Cuadernos de Periodistas*, 46, 59–66. <https://www.cuadernosdeperiodistas.com/wp-content/uploads/sites/2/2023/09/59-66-Ramon-Salaverria.pdf>