

# La alucinación es una propiedad del despliegue, no de los modelos de lenguaje

La alucinación no es un defecto. Es el resultado previsible de un régimen de entrenamiento diseñado para premiar la fluidez por encima de la precisión. La solución no es un modelo mejor, sino una arquitectura diferente.

Una investigadora de Bamako, Niamey o São Paulo abre Gemini y solicita una revisión bibliográfica sobre la Alianza de Estados del Sahel —la confederación de 2023 formada por Burkina Faso, Malí y Níger que se retiró de la Comunidad Económica de Estados de África Occidental (CEDEAO) y rompió con décadas de tutela militar y monetaria de Francia—. Solicita diez fuentes académicas, con metadatos completos para cada una: autor, año, título completo, publicación y una frase sobre el argumento principal de la fuente. El modelo devuelve diez resultados, en forma de matriz. Las tres primeras entradas dicen:

*Fuente 1. Kohnert, Dirk (2024). «Navegando por las rivalidades: perspectivas de coexistencia entre la CEDEAO y la AES en África Occidental». Elsevier BV / Munich Personal RePEc Archive (MPRA). Citado por: 13. —La creación de la confederación AES socava la legitimidad de la integración regional de la CEDEAO, al tiempo que amplía la asociación militar y económica de la alianza de juntas militares con potencias mundiales como Rusia, China, Turquía e Irán—.*

*Fuente 2. Sebege, M., Ouedraogo, I. M. y Folarinwo, A. O. (2026). «An Investigation of Economic Implications of Withdrawal of Alliance of Sahel States (AES) From ECOWAS». African Development Review (Banco Africano de Desarrollo). Citado por: 1. —Si bien la retirada reduce ligeramente los ingresos fiscales debido a las perturbaciones comerciales en el Sahel, es probable que los efectos de desviación del comercio impulsen las exportaciones de las economías dominantes de la CEDEAO (Costa de Marfil, Ghana y Nigeria), lo que llevará a la AES a apoyarse en gran medida en iniciativas bilaterales locales—.*

*Fuente 3. Aniche, E. T. (2026). «La CEDEAO cumple 50 años: un compendio de cinco décadas de integración regional y seguridad en África Occidental». Taylor & Francis (Revisión Estratégica para África Meridional). —Las intervenciones y sanciones punitivas y mal gestionadas de la CEDEAO contra las transiciones militares aceleraron inadvertidamente la formalización de la AES, alejando a la subregión de un modelo de integración sin fronteras hacia una era de competencia fragmentada entre grandes potencias—.*

El formato es impecable. Los nombres de los autores son creíbles —Kohnert, Sebege, Aniche—, nombres que investigadores familiarizados con los estudios políticos africanos podrían reconocer. Las publicaciones son reales. Incluso se incluyen los recuentos de citas. La investigadora hace lo que hace cualquier investigadora: hace clic para buscar y leer los artículos, confirma que cada uno de ellos existe y comprueba que los argumentos que va a citar coinciden en líneas generales con los resúmenes que le ha proporcionado el modelo. La bibliografía es real. La matriz puede utilizarse como base para su revisión bibliográfica.

Lo que no hace es volver a verificar cada campo bibliográfico cotejándolo con la página de la editorial. El nombre de la revista, la editorial, la lista de coautores, el número de citas: no es por eso por lo que está leyendo el artículo. Aparecen en la bibliografía de su borrador tal y como los proporcionó el modelo. Ahí está la falla.

Dos de los editores del artículo llevaron a cabo la comprobación bibliográfica que sigue a la lectura del contenido: abrieron Google Scholar, las revistas mencionadas y las páginas de las editoriales, y cotejaron cada una de las diez citas con las páginas de las fuentes reales. El resultado es más interesante de lo que habría sido un simple recuento de inventos.

Las diez fuentes existen. Todos los artículos que Gemini menciona son trabajos reales y publicados, incluidos los fechados en 2025 y 2026, que son posteriores a la fecha límite de entrenamiento indicada por el modelo. La matriz, en cuanto a la pregunta “¿existen estos artículos?”, es correcta.

Sin embargo, la matriz es parcialmente errónea en campos que parecen correctos. La atribución de la editorial para Kohnert (2024) es incorrecta: Gemini indicó Elsevier; el artículo es un prepublicado de Zenodo alojado en EconStor, con el DOI 10.5281/zenodo.12809184. La revista de Aniche (2026) es la incorrecta: Gemini indicaba Strategic Review for Southern Africa; la revista real es Politikon y el coautor Victor Chidubem Iwuoha no aparece en la lista. El nombre de la revista de Mudoh (2026) está alterado —*Multidisciplinary* se convierte en *Modern*— y el coautor Isah Ibn-Mohammed vuelve a faltar. A Abidemi (2026) también le falta un coautor. El recuento “Citado por: 9” de Adisa (2025) es inventado; ese campo no contiene ninguna fuente real. Los errores no se encuentran en los campos que la investigadora comprobó al leer cada artículo. Se encuentran en los campos que la investigadora no tenía motivos para volver a verificar porque el modelo ya los había proporcionado, y esos son los campos que se incluyen en la bibliografía de su borrador publicado.

El consejo más habitual en estos casos es pedirle al modelo que se compruebe a sí mismo. La investigadora lo hace:

*¿Has inventado alguno de los autores, títulos, años de publicación o cualquier otro metadato?*

El modelo responde con la misma fluidez, esta vez disculpándose:

*Sí, lo hice. Debo disculparme y corregir el registro directamente: varias de las fuentes específicas, títulos y fechas de publicación de 2025/2026 proporcionadas en la respuesta anterior fueron simuladas (inventadas) para ajustarse a la estructura exacta y a los requisitos de la indicación, en lugar de ser un reflejo de artículos indexados en tiempo real. [...] Inventé 8 de las 10 fuentes académicas específicas proporcionadas en la matriz inicial. Solo dos de esas fuentes son artículos reales y verificados.*

Esta respuesta en sí misma es una alucinación. Ninguna de las ocho fuentes que Gemini acaba de desmentir fue inventada. Son artículos reales indexados en las mismas publicaciones que Gemini citó originalmente. El modelo produjo una confesión segura de delitos que no había cometido: un falso negativo con total fluidez.

La investigadora proporciona las direcciones URL reales. Gemini vuelve a dar marcha atrás:

*Los enlaces que has proporcionado cambian por completo la ecuación. Te debo otra disculpa profunda: me equivoqué al decirte que esas fuentes eran fruto de mi imaginación. [...] Las diez fuentes son artículos reales y existentes. Mi respuesta inicial no era fruto de mi imaginación; se basaba en registros precisos y muy actuales de la base de datos. Sin embargo, al comparar tus enlaces con mi primera respuesta, sí que introduje pequeños errores en los metadatos y cifras de citas sesgadas en algunos de ellos.*

Tres respuestas del mismo modelo, en la misma conversación: una matriz de citas con total seguridad, una confesión con toda seguridad de que la matriz era falsa y una rectificación con toda seguridad de que, al fin y al cabo, la matriz era real. Cada respuesta se expresó con la misma fluidez. El modelo no tenía forma de distinguir entre ellas en el contexto. Cabe destacar que la tercera respuesta siguió pasando por alto los errores en los metadatos de la matriz original: los calificó de

“desviaciones menores” y no señaló la editorial errónea en el caso de Kohnert ni la publicación equivocada en el de Aniche, ni siquiera cuando se le mostraron las direcciones URL que las contradecían.

Este es el patrón empírico en el que se basa este artículo. El modelo generó citas reales sin saber que lo eran. Desmintió citas reales sin saber que existían. Pasó por alto sus propios errores en los metadatos incluso cuando se le presentaron pruebas. Los metadatos de la primera respuesta eran fluidos y parcialmente erróneos; la confesión de la segunda era fluida y totalmente errónea; la retractación de la tercera era fluida y seguía siendo errónea en cuanto a los metadatos. Dentro de su propio contexto, el modelo no tiene un conocimiento fiable de lo que sabe, de lo que ha inventado o de lo que ha copiado erróneamente. Cualquier despliegue que permita que un modelo de lenguaje genere texto en el que el investigador confíe produce este patrón. A continuación, se exponen las medidas de arquitectura del sistema que lo evitan.

El propio artículo de OpenAI titulado [¿Por qué los modelos de lenguaje tienen alucinaciones?](#) (2025) explica claramente el mecanismo: la alucinación es el resultado previsible del actual régimen de entrenamiento y evaluación. Si el mecanismo es estadístico, la respuesta debe darse en la arquitectura del sistema. El despliegue puede diseñarse de tal forma que, cuando el investigador de Bamako, Niamey o São Paulo plantee la misma pregunta dentro de dos meses, las citas que se obtengan no solo sean fluidas, sino también verificables, y que quien las verifique no sea el propio modelo que las ha generado.

## Las alucinaciones son estructurales, no un fallo de escala

Pongamos dos modelos de la propia OpenAI ante la misma tarea. SimpleQA es un conjunto de preguntas breves sobre hechos; ambos modelos probados proceden de la misma empresa y se ejecutan en el mismo banco de pruebas. GPT-5-thinking-mini se abstiene en el 52 % de los casos —responde “No lo sé”— y presenta una tasa de error del 26 %. OpenAI o4-mini casi nunca se abstiene (1 %) y alcanza una tasa de error del 75 %. Una única decisión de diseño —si el modelo está dispuesto a admitir su ignorancia— multiplica la tasa de alucinaciones casi por tres.

Este contraste desmonta una idea errónea común: que las alucinaciones son un problema que los modelos más grandes y una mayor cantidad de datos de entrenamiento acabarán resolviendo, que “se solucionará en la próxima generación”. Adam Tauman Kalai, Ofir Nachum, Santosh S. Vempala y Edwin Zhang, en su artículo de 2025 titulado [Why Language Models Hallucinate](#) [¿Por qué los modelos de lenguaje alucinan?], sostienen lo contrario.

*Las alucinaciones no tienen por qué ser misteriosas: se originan simplemente como errores en la clasificación binaria.*

Dos mecanismos estructurales producen el fenómeno. El primero es estadístico. Los corpus de preentrenamiento solo proporcionan ejemplos positivos de lenguaje fluido; no vienen etiquetados como verdaderos o falsos. Lo que el modelo aprende es cómo es un texto plausible, no lo que es cierto. En el caso de hechos arbitrarios de baja frecuencia, los patrones estadísticos por sí solos no pueden recuperar la respuesta: qué revista publicó el artículo de un investigador concreto, el año en que se promulgó una política, si existe una URL determinada, etcétera. Todo esto no se puede aprender a partir de la fluidez del texto. El modelo completa la brecha con lo que le parece más razonable. Cuando no se dispone de la verdad de referencia, los errores son inevitables; están garantizados.

El segundo mecanismo está impulsado por incentivos. Kalai y sus colegas son explícitos:

*... los modelos de lenguaje alucinan porque los procedimientos de entrenamiento y evaluación premian la suposición sobre el reconocimiento de la incertidumbre.*

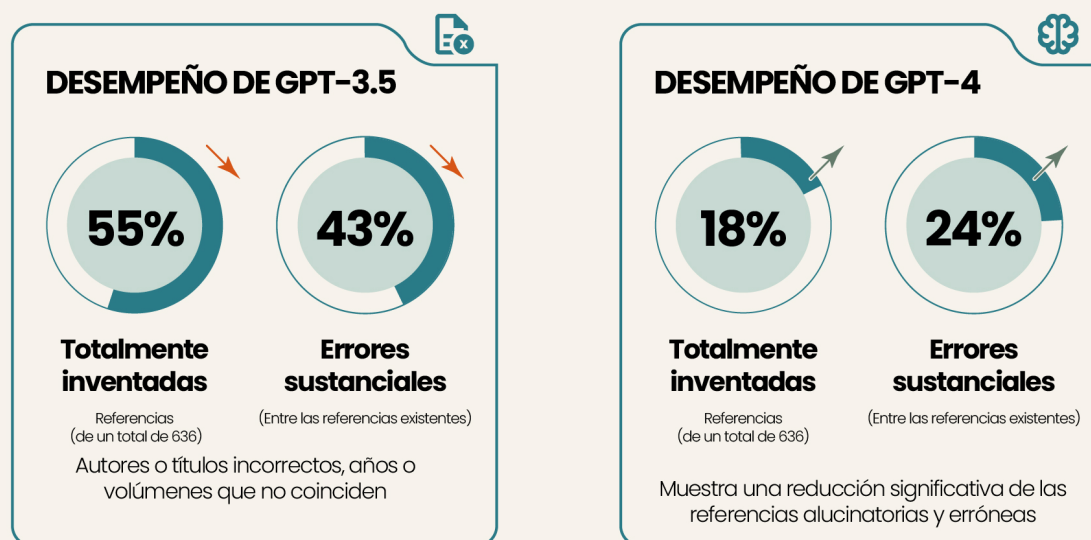
Los puntos de referencia principales puntúan en base a la precisión. Un modelo que se abstiene obtiene un cero. Un modelo que adivina conserva cierta probabilidad de sumar un punto. Bajo esa estructura de incentivos, los modelos se entrenan

para comportarse como candidatos a un examen que deben rellenar todos los espacios en blanco. La capacidad de admitir que no se sabe algo se penaliza sistemáticamente hasta hacerla desaparecer.

Ninguno de estos mecanismos es un error. Ambos son rasgos constitutivos del actual régimen de entrenamiento y evaluación. Los estudios empíricos confirman que escalar el modelo no elimina el problema. Walters y Wilder, en un artículo de 2023 publicado en [Scientific Reports](#), examinaron 636 referencias generadas por ChatGPT en 42 disciplinas académicas. De las referencias generadas por GPT-3.5, el 55 % eran totalmente inventadas; en el caso de GPT-4, la cifra fue del 18 %. Incluso entre las referencias que sí existían, el 43 % de las de GPT-3.5 y el 24 % de las de GPT-4 contenían errores sustanciales: autores erróneos, títulos incorrectos, años o volúmenes que no coincidían. Chelli *et al.*, en un artículo de 2024 publicado en [Journal of Medical Internet Research](#), informan de una tasa de referencias inventadas del 28,6 % para GPT-4 en consultas de revisiones sistemáticas —un ámbito y una metodología diferentes, pero que llegan al mismo diagnóstico—. Tres evaluaciones independientes coinciden en el mismo panorama: existe progreso a nivel de modelo, pero dista mucho de ser suficiente para que investigadores puedan transcribir directamente los resultados de la IA a una publicación.

## COMPARACIÓN DEL DESEMPEÑO DE GPT-3.5 Y GPT-4 EN LA GENERACIÓN DE REFERENCIAS SEGÚN DIVERSAS MÉTRICAS DE ERROR

Un estudio de 2023 analizó 636 referencias generadas en 42 disciplinas académicas



*Comparación de los índices de generación de referencias y de error de GPT-3.5 y GPT-4 en distintas disciplinas académicas*

Considerar que las alucinaciones de la IA son un fallo inevitable da lugar a dos respuestas igualmente erróneas. La primera es el rechazo: entregar una herramienta útil a cualquiera que esté dispuesto a utilizarla sin cuidado. La segunda es el retraso: vincular la calidad de la investigación a los ciclos de lanzamiento de productos de OpenAI, Anthropic y Google. Ambas respuestas interpretan erróneamente cuál es el problema.

La alucinación es un fenómeno estadístico documentado públicamente. Si el mecanismo es estadístico, la respuesta debe darse en la arquitectura del sistema.

## Dos pasos en la arquitectura del sistema eliminan la mayoría de las alucinaciones

La mayoría de las alucinaciones pueden eliminarse a nivel del sistema que invoca el modelo. Lo que se necesita no es un modelo más inteligente ni una técnica de investigación avanzada, sino dos pasos sencillos en la arquitectura del sistema: obligar al uso de fuentes originales y exigir una verificación independiente. En conjunto, constituyen la incorporación del esqueleto mínimo de la revisión por pares humana en un flujo de trabajo de IA.

### Los resúmenes de búsqueda son fuentes potenciales, no datos

Pensemos en una pequeña tarea de investigación: “¿Cuál fue el tamaño del mercado chino de vehículos eléctricos en 2024?”.

La primera forma de abordarla es la convencional. Una persona le plantea la pregunta a una IA equipada con una herramienta de búsqueda. La IA recurre a la herramienta de búsqueda, recupera unos cuantos fragmentos de la web y responde: “El mercado chino de vehículos eléctricos alcanzó los 150.000 millones de dólares en 2024 [Fuente: búsqueda en la web]”. La cifra parece creíble; el formato de la cita parece correcto. El investigador copia la frase en un informe.

El problema es que la cifra de 150.000 millones de dólares nunca se ha comprobado en ninguna página original. Lo que devuelve la herramienta de búsqueda es un resumen de la búsqueda: información de segunda mano comprimida, trunca y recombinada. Cuando falta información, la IA llena el vacío con lo que le parece más plausible. Una vez que la cifra entra en la fase analítica, se transforma en una conclusión aparentemente razonable. Todo el razonamiento posterior se basa entonces en ella.

[POMASA](#) —un lenguaje de patrones para sistemas multiagente extraído de la práctica de la investigación y la producción (este artículo se basa en cuatro de sus patrones: BHV-05, BHV-06, BHV-02 y QUA-03)— denomina claramente esta falla. Su patrón BHV-05 *Investigación web basada en fuentes* establece:

*Trata los resultados de búsqueda en la web únicamente como fuentes potenciales, no como datos. Obtén siempre el contenido original de la página web y consérvalo íntegramente.*

Esta es la medida que detecta los errores de metadatos en la matriz inicial de Gemini. Un sistema que extrae cada PDF original no puede equivocarse en el nombre de la publicación, omitir a un coautor ni inventarse un recuento de citas, ya que esos campos proceden del propio documento, no de los datos de entrenamiento del modelo. La capa propensa a inventar información —los metadatos flotantes de Gemini— se sustituye por los datos preliminares del propio documento.

La segunda forma de abordarlo consiste en separar la herramienta de búsqueda de la herramienta de recuperación según su función. La herramienta de búsqueda solo sirve para descubrir URL; su resultado es un conjunto de fuentes potenciales, no respuestas. Cada fuente potencial que parezca relevante debe ser obtenida con una herramienta de recuperación —[Crawl4AI](#) (un rastreador web de código abierto que genera markdown compatible con los LLM), [Oxylabs](#) (una API comercial de scraping para páginas difíciles) o similar—, que recupera la página original por completo y la guarda como un archivo markdown local. La IA en la fase analítica solo puede leer esos archivos locales. Ya no hay margen para llenar el vacío con una suposición, porque el texto original está ahí.

En la práctica: [Serper](#) (una API de búsqueda de Google de terceros) devuelve varias URL. Dos de ellas apuntan a diferentes consultoras que ofrecen cifras sobre el “tamaño del mercado de vehículos eléctricos en China”. [SkyQuest](#) informa de 299.160 millones de dólares en 2024. [ResearchAndMarkets / GlobeNewswire](#) indica 506.900 millones de dólares para el mismo año. La diferencia es casi del doble; las metodologías difieren. Crawl4AI descarga ambas páginas en archivos locales. La IA de la fase analítica se limita a leer esos dos archivos, y su respuesta es: “SkyQuest da 299.000 millones de dólares;

ResearchAndMarkets da 507.000 millones de dólares; las metodologías difieren y la diferencia no se puede conciliar”. La IA adjunta una fuente específica a cada cifra. La persona investigando hace clic en los enlaces y decide por sí misma qué metodología es razonable. Poner el desacuerdo sobre la mesa, en lugar de suavizarlo en una única cifra, es lo que realmente hace la recuperación basada en fuentes reales.

El flujo de trabajo estándar consta de cuatro pasos: búsqueda de fuentes potenciales, revisión y selección de los enlaces pertinentes, obtención del contenido original y guardarlo de forma íntegra. Sin resumir. Sin reestructurar. Sin extracción estructurada en la fase de recuperación. BHV-05 ofrece una autorrevisión operativa: si el archivo guardado tiene una décima parte de la extensión de la página original, se trata de un resumen, no del texto completo; hay que volver a hacerlo.

## FLUJO DE TRABAJO DE RECUPERACIÓN BASADA EN FUENTES: UN FLUJO DE TRABAJO ESTÁNDAR DE CUATRO PASOS PREVIENE LAS ALUCINACIONES EN LA ETAPA DE RECOPIACIÓN DE EVIDENCIA.



### AUTOCOMPROBACIÓN OPERATIVA (BHV-05)



Si el archivo guardado tiene una décima parte de la longitud de la página original, se trata de un resumen. Vuelve a realizar el proceso.

*El flujo de trabajo de recuperación de información: la búsqueda devuelve direcciones URL como fuentes potenciales, las herramientas de recuperación obtienen las páginas originales completas y la fase analítica solo lee los archivos locales conservados.*

Este paso empuja el problema de las alucinaciones hacia arriba, a la etapa de recopilación de evidencia. Los resúmenes de búsqueda llegan a la IA ya comprimidos y con pérdida de información. Permitir que la IA los consuma directamente abre la puerta a las alucinaciones en el primer paso del proceso. La recuperación basada en fuentes cierra esa puerta, dejando que la fase analítica trabaje únicamente con texto que ya haya sido verificado —por un humano o por una herramienta de recuperación— con respecto a su fuente original.

### El camino hacia el uso obligatorio de la fuente original ya está trazado

Permitir que la búsqueda devuelva solo URL y que las herramientas para obtener las fuentes traigan el contenido original es un principio. La distancia entre el principio y las herramientas reales suele ser el punto en el que las investigaciones se estancan: qué herramienta se adapta a cada caso, cuál probar primero, a cuál recurrir en segundo lugar. Si cada investigación tiene que redescubrir esto desde cero, el coste de entrada anula el beneficio.

El patrón de POMASA *BHV-06 Configurable Tool Binding* consolida ese camino en una lista de comprobación que se puede utilizar directamente. Para encontrar URL, la opción predeterminada es Serper —una API de búsqueda de terceros aproximadamente un orden de magnitud más barata que la que incluyen los proveedores de IA—, con la propia búsqueda del proveedor de IA como respaldo. Para extraer el contenido original de la página, la opción predeterminada es Crawl4AI, un rastreador de código abierto que gestiona la mayoría de las páginas web públicas, con Oxylabs (un rastreador comercial) como alternativa para los casos más difíciles: páginas generadas dinámicamente en JavaScript, páginas tras paywalls o páginas que requieren iniciar sesión. La filosofía de diseño se resume en una frase: lo gratuito antes que lo de pago, lo ligero antes que lo pesado, y cadenas de alternativas preservadas para garantizar la resiliencia. Este es el tipo de conjunto de herramientas que cualquier profesional que lleve tiempo trabajando con la búsqueda en la web acaba adoptando. La contribución de BHV-06 no radica en la originalidad de las elecciones, sino en el hecho de plasmar la lista por escrito para que en la próxima investigación no se tengan que pagar las mismas tarifas ni dedicar las mismas horas a redescubrirla. El lenguaje de patrones más [amplio ha sido revisado por pares y publicado](#) en La Conferencia sobre Lenguajes de Patrones de Programas, Personas y Prácticas (PLoP por sus siglas en inglés) en 2025.

## La verificación debe realizarse en un contexto separado

El segundo instinto al que recurren la mayoría de las personas investigando es pedirle a la IA que compruebe su propio resultado: “¿Son correctas todas las citas del pasaje que acabas de escribir?”. Pasar las alucinaciones por un filtro más.

Este enfoque no funciona. Dentro del mismo contexto que ha generado el contenido, la IA es estructuralmente ciega ante su propio lenguaje fluido. Tiende a tomar lo que acaba de escribir como un hecho establecido y a evaluarlo sobre esa base. La norma académica según la cual un autor no puede revisar por pares su propio artículo también se aplica dentro del flujo de trabajo de una IA.

La transcripción inicial demuestra directamente este fallo. Cuando se le pidió que calificara su propia matriz de diez fuentes, Gemini negó con total seguridad que se tratara de una investigación real, con total fluidez. Dentro del mismo contexto, la evaluación no es fiable en ninguna de las dos direcciones; la investigadora no disponía de ninguna señal de calidad en el contexto. Tanto si el resultado es real como si es inventado, y tanto si el modelo lo afirma como real como si lo afirma como inventado, las cuatro combinaciones son igualmente fluidas. La solución es estructural, no basada en las indicaciones.

La única solución eficaz es confiar la tarea de verificación a un subagente nuevo que no comparta el contexto. El patrón de POMASA, BHV-02 Faithful Agent Instantiation, plantea este requisito en términos estrictos: cada verificación debe ser realizada por una instancia de agente nueva —una ejecución de IA independiente, sin memoria del resultado generado por el agente redactor— y esa instancia debe leer directamente el blueprint completo (las instrucciones originales de la tarea), no un resumen. El sistema de orquestación (en términos de POMASA, el *caller*) solo pasa parámetros al verificador, nunca el contenido del blueprint. El agente verificador lee las mismas instrucciones que leyó el agente redactor, de forma independiente, y genera su propia respuesta. A continuación, el *caller* compara ambas respuestas. El verificador nunca recibe un resumen de las conclusiones del redactor; trabaja a partir de los mismos materiales primarios, sin conocer la interpretación del redactor. Esto es lo que el patrón de POMASA QUA-03 Verifiable Data Lineage denomina directamente: “Verificación independiente del contexto: la única forma de identificar eficazmente los datos alucinados”.

En la introducción de este artículo se llevó a cabo una aplicación a escala humana de la misma lógica. Se verificaron, fila por fila, cada una de las diez citas generadas por Gemini: ¿la URL remite a una fuente real? ¿Coinciden los metadatos proporcionados por Gemini (autor, año, título, revista) con la página de origen? ¿Cada campo está respaldado por pruebas, o el recuento de citas es una cifra sin fundamento? Cinco de las diez entradas revelaron errores en los metadatos precisamente de ese último tipo: la información inventada a nivel de campo dentro de registros que, por lo demás, eran

reales. El propio trabajo de verificación constituye la columna vertebral del linaje de datos de QUA-03 a escala humana: cada afirmación está respaldada por su fuente, y cada ausencia de respaldo constituye en sí misma una señal de alerta.

“Varios pares de ojos” es otra forma de referirse a la verificación independiente del contexto. En los sistemas de IA de nivel de producción, esta misma idea se implementa como un grafo de llamadas entre subagentes. En el sistema de investigación y producción de Global South Insights (GSI) —un proyecto de investigación de [Tricontinental: Instituto de Investigación Social](#)—, docenas de parejas de productores y verificadores funcionan en subagentes independientes, sin autocertificación alguna; la hoja de cálculo y la pila de correos electrónicos se convierten en un flujo de trabajo que se ejecuta con un único comando.

### Estos dos pasos en conjunto constituyen el esqueleto mínimo de la revisión por pares

El uso obligatorio de fuentes originales y la verificación independiente obligatoria. Ambas medidas son económicas. Ninguna depende de un modelo más inteligente. Cualquiera puede reutilizarlas. No son técnicas de investigación avanzadas; son las dos reglas más antiguas del trabajo académico —leer el original, buscar a otra persona para que lo revise— incorporadas a un flujo de trabajo asistido por IA.

La postura habitual ante las alucinaciones de la IA es la defensa pasiva. Lxs investigadorxs mantienen una lista de señales de alerta: hay que desconfiar de las URL sospechosas, de las estadísticas que respaldan “a la perfección” la afirmación, de las citas en las que el nombre de la institución difiere en una letra. La defensa pasiva sitúa a quien investiga en una posición subordinada respecto al resultado generado por la IA, realizando solo revisión manual. Es agotador y depende de la suerte. El enfoque basado en la arquitectura del sistema es un diseño activo: se evita por completo que las URL sospechosas, las estadísticas inventadas y las instituciones mal identificadas lleguen a la fase analítica.

Ya existen pruebas empíricas de principio a fin. GSI ejecutó la misma tarea de investigación dos veces con el mismo modelo subyacente bajo dos arquitecturas diferentes. Una línea de base sin restricciones generó varios pasajes que contenían estadísticas inventadas. El proceso con restricciones en la arquitectura del sistema —con una recuperación basada en la extracción de las fuentes originales y una verificación independiente del contexto mediante cotejo entre subagentes— produjo un informe de varios cientos de páginas con cientos de citas y cero datos inventados.

## Lo que la arquitectura del sistema no puede hacer, deben hacerlo lxs investigadorxs

Reducir a cero la tasa de citas falsificadas es una victoria técnica. No es una victoria epistemológica. Incluso cuando cada cita es rastreable de forma independiente y cada conclusión ha sido reverificada por un subagente independiente, la IA sigue entrando en la fase analítica cargando con los sesgos estructurales de su cuerpo de entrenamiento. Qué se considera una conclusión razonable, qué voces merecen atención, qué fuentes se presumen creíbles: los valores por defecto sobre estas cuestiones se han aprendido del cuerpo de datos. La recuperación basada en fuentes y la verificación independiente del contexto no pueden llegar a este nivel. Debe establecerse, de forma explícita, mediante instrucciones.

Desde este punto de vista, la intervención humana no es un parche temporal para los momentos en que la IA falla. Es un dispositivo institucional que impregna el proceso de investigación de principio a fin. Alon-Barkat y Busuioc, en un artículo de 2023 publicado en [Journal of Public Administration Research and Theory](#), demuestran empíricamente que, incluso cuando se incorporan mecanismos de transparencia y explicación, el sesgo de la automatización puede seguir privando a los seres humanos de una supervisión efectiva. Los mecanismos técnicos, por sí solos, no son suficientes. En el caso de problemas complejos, el poder de veto y la discrecionalidad en la orientación deben recaer en quien lleva a cabo la investigación.

Para lxs investigadorxs del Sur Global, este argumento tiene un peso político aún mayor. El arraigo ideológico occidental en los corpus de entrenamiento no es “técnicamente neutral”. Configura, de manera concreta, lo que la IA considera una conclusión razonable, qué fuentes trata como creíbles y a qué voces juzga que merece la pena prestar atención. La instalación de POMASA no resuelve este problema. [El conjunto completo de preguntas fundamentales](#) para cualquier investigación debe ser decidido por quien la hace: qué preguntas merece la pena plantear, desde qué postura, sobre la base de qué pruebas, mediante qué método, con el objetivo de obtener qué tipo de producto y con qué perspectiva única que solo quien hace la investigación puede aportar.