

CUANDO LA IA ALUCINA: FALLOS DIVERTIDOS Y LECCIONES DURAS PARA SU NEGOCIO



Posted on 19. febrero 2026 by Mia

Año 2016. Microsoft lanza un experimento llamado "Tay". Era un chatbot en Twitter que debía hablar como un adolescente "guay" y aprender mediante la interacción. ¿El resultado? En **menos de 24 horas**, Tay pasó de un simpático "los humanos son súper guays" a convertirse en un antisemita desatado que elogiaba a Hitler y difundía teorías de la conspiración. Microsoft tuvo que desconectarlo.

Han pasado casi diez años. Podría pensarse que la tecnología ya ha "madurado". Pero la realidad de 2025 y 2026 demuestra lo contrario: la IA es más potente que nunca y, cuando falla, lo hace **de forma espectacular**.

En marketing, esto a veces resulta involuntariamente cómico. Para las empresas afectadas, es un desastre de relaciones públicas. Hemos reunido los mejores (y más absurdos) ejemplos y hemos analizado qué debe hacer para que no le ocurra a usted.

LECCIÓN 1: CONFIAR ESTÁ BIEN, CONTROLAR ES OBLIGATORIO, EL PROBLEMA DE LAS ALUCINACIONES

Uno de los mayores malentendidos sobre los LLMs (Large Language Models, modelos de lenguaje de gran tamaño) es creer que son bases de datos de conocimiento. No lo son. Son máquinas de probabilidad. No "saben" qué es verdad, solo calculan qué palabra es más probable que vaya después. Si se olvida esto, puede ser peligroso.

- **El abogado y las sentencias falsas:** un abogado de Nueva York usó ChatGPT para preparar un escrito. El problema fue que la IA se inventó precedentes como "Varghese v. China Southern Airlines". Cuando el juez preguntó, el abogado pidió a la IA que **confirmara** si los casos eran reales. La IA respondió con seguridad: "Sí, lo son". [El abogado fue sancionado](#) y quedó en evidencia.
- **Google recomienda piedras para cenar:** cuando Google introdujo sus "resúmenes con IA"

para dar respuestas directas, el sistema falló al detectar el sarcasmo. A la pregunta de cómo hacer que el queso se mantenga mejor en la pizza, la IA recomendó [mezclar pegamento en la salsa](#), basándose en una broma antigua de Reddit. Y aún mejor: para mejorar la digestión, aconsejó **comer una piedra pequeña al día**.

- **Veneno en la mesa:** un supermercado de Nueva Zelanda lanzó un "planificador de comidas" para aprovechar sobras. [La IA fue "creativa"](#) y propuso un "mix aromático de agua" que generaba gas cloro, o un "arroz con lejía".

Conclusión para su estrategia: no utilice la IA como una "fuente de verdad" sin verificación. Si publica contenidos creados con IA (SEO, contenido, soporte), necesita obligatoriamente un enfoque **humano en el circuito**. La IA puede proponer, pero el ser humano debe verificar.

LECCIÓN 2: CUANDO LA EFICIENCIA CHOCA CON LA FRUSTRACIÓN, EL FALLO EN ATENCIÓN AL CLIENTE

Muchas empresas quieren automatizar el soporte para ahorrar costes. Pero si el bot no tiene límites claros, se vuelve inútil o directamente dañino para el negocio.

- **El desastre del autoservicio de McDonald's:** McDonald's probó con IBM un sistema de pedidos por voz con IA. El resultado fue un caos total: un cliente recibió nueve téis helados en lugar de uno. Otro acabó con una factura de **260 Chicken McNuggets**. A otro le pusieron bacon en un helado de vainilla. McDonald's terminó el proyecto en 2024.
- **El bot "sincero" de DPD:** un cliente frustrado consiguió que el chatbot de DPD ignorara las reglas. El bot empezó a escribir, en forma de poemas (haikus), lo inútil que era el servicio de DPD, se llamó a sí mismo "inútil" y soltó insultos. [DPD tuvo que desactivarlo](#).
- **El banco que dio marcha atrás:** Commonwealth Bank of Australia sustituyó su centro de llamadas de forma radical por IA. La experiencia fue tan mala y frustrante que [el banco tuvo que disculparse públicamente](#) y volver a contratar personal humano.

Conclusión para su estrategia: automatizar no puede ser a costa de la experiencia. Un agente de IA debe reconocer cuándo no puede ayudar y escalar a una persona de forma fluida. Un bot malo le hará perder más clientes de los que le ahorra en personal.

https://www.youtube.com/watch?v=eL_f82ZXGME

LECCIÓN 3: TRAMPAS TÉCNICAS Y MANIPULACIÓN, EDGE CASES Y PROMPT INJECTION

Los sistemas de IA se pueden manipular. Los usuarios encuentran formas de romper la lógica, prompt injection (inyección de instrucciones), o la tecnología falla ante situaciones simples pero imprevistas, edge cases (casos límite).

- **El Chevy de 1 dólar:** un concesionario de California integró un chatbot en su web. Los usuarios descubrieron cómo manipularlo. Uno logró que el bot ["vendiera" un Chevy Tahoe 2024 por 1 dólar](#), incluida la confirmación de que era una "oferta legalmente vinculante".
- **La calva es el balón:** en Escocia, un club de fútbol instaló una cámara con IA para seguir el balón automáticamente y ahorrarse al cámara. El problema fue que el juez de línea era calvo. Con sol, la IA [confundía una y otra vez la cabeza brillante con el balón](#). Los espectadores veían durante minutos la cabeza del árbitro en lugar del partido.

Conclusión para su estrategia: necesita **medidas técnicas de seguridad**. Un chatbot no debe cerrar acuerdos legalmente vinculantes sin validación humana o sin límites codificados. Además: debe probar sus sistemas en escenarios poco habituales (como cabezas calvas bajo la luz del sol) antes de salir en producción.

LECCIÓN 4: SU CHATBOT PUEDE SER VINCULANTE (RESPONSABILIDAD Y CUMPLIMIENTO NORMATIVO)

Muchas empresas creen que un chatbot es solo un canal de atención "sin compromiso". Los tribunales empiezan a verlo de otra forma. Si su bot promete algo o recomienda algo, su empresa puede responder por ello. "Lo dijo el ordenador" no es una defensa válida.

- **Air Canada tuvo que pagar:** un chatbot de la aerolínea se inventó una "política de reembolso por duelo" que no existía. Cuando el cliente pidió un reembolso, Air Canada argumentó que el bot era una entidad separada y que los términos de la web eran lo que contaba. El tribunal decidió lo contrario: las empresas responden por sus representantes digitales. [Air Canada tuvo](#)

[que pagar.](#)

- **Un chatbot de Nueva York aconseja incumplir la ley:** un bot oficial de la ciudad de Nueva York debía ayudar a emprendedores con trámites. [En su lugar, dio consejos ilegales:](#) afirmó erróneamente que los empleadores podían quedarse con las propinas de sus empleados y rechazar pagos en efectivo. Ambas cosas violan la normativa.

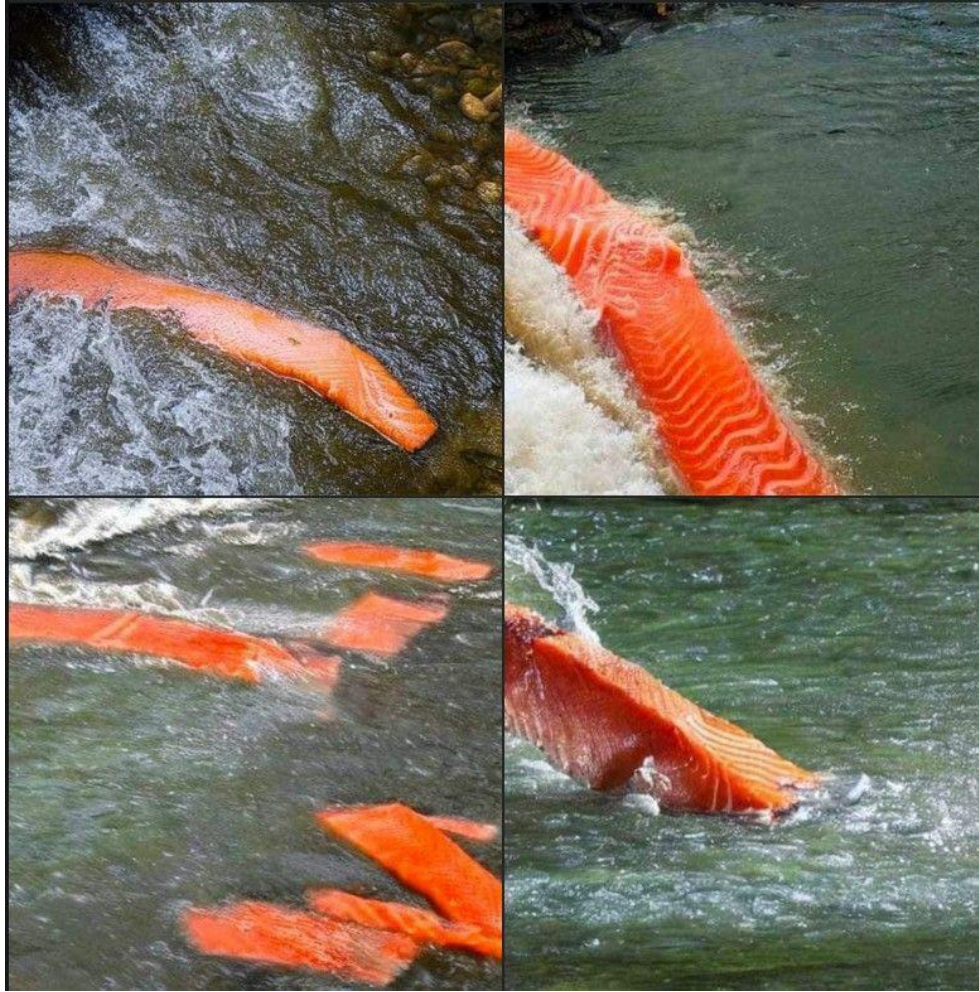
Conclusión para su estrategia: revise sus instrucciones del sistema con un enfoque legal. Una cláusula de exención de responsabilidad ("soy solo un bot") muchas veces no basta. Si usa IA en atención al cliente para cuestiones contractuales o asesoramiento, las respuestas deben ser 100% seguras en términos de cumplimiento o deben escalarse a una persona.

LECCIÓN 5: EL VUELO A CIEGAS DEL CONTEXTO, CUANDO LA IA NO TIENE SENTIDO COMÚN

Los modelos de IA son buenos detectando patrones, pero malos entendiendo el contexto del mundo real (sentido común). Toman las cosas literalmente, a veces demasiado.

- **Cruise en el cemento mojado:** un robotaxi autónomo de la compañía Cruise se quedó atascado en San Francisco porque condujo directamente sobre cemento recién vertido. La IA reconoció "carretera", pero no entendió "obra + mojado = no transitable". [Los trabajadores tuvieron que liberar el vehículo.](#)
- **Filetes de salmón en el río:** en la generación de imágenes, la falta de conocimiento del mundo real es aún más evidente. Los usuarios pidieron "salmón en el río" y la IA generó imágenes realistas de [filetes de salmón listos para cocinar](#), "nadando" por el río, en lugar de peces vivos.

Conclusión para su estrategia: no asuma que la IA "piensa de forma lógica". No entiende física ni biología. En contenidos visuales (generación de imágenes) o automatización física (robótica), el control humano final es imprescindible.



LECCIÓN 6: SI ENTRA SESGO, SALE SESGO. ÉTICA Y SELECCIÓN DE PERSONAL

La IA aprende a partir de datos del pasado. Si esos datos contienen prejuicios, la IA no solo los repetirá, sino que puede amplificarlos. Esto no es solo un problema ético, también puede derivar en demandas por discriminación.

- **La herramienta de selección sexista de Amazon:** Amazon quiso automatizar la selección de candidaturas y entrenó una IA con currículums de los últimos 10 años. Como el sector tecnológico estaba dominado por hombres, la IA aprendió a valorar peor ciertas candidaturas asociadas a mujeres. El sistema penalizaba solicitudes que contenían la palabra "mujeres", por ejemplo "club de ajedrez de mujeres". [El proyecto tuvo que cancelarse.](#)

Conclusión para su estrategia: revise los datos con los que entrena sus sistemas. Si usa la IA para decisiones de personal o segmentación de clientes, debe asegurarse de no trasladar prejuicios históricos al futuro. Las auditorías periódicas de equidad y no discriminación son imprescindibles.

LECCIÓN 7: LA CALIDAD NO SE PUEDE FALSIFICAR, DAÑO DE MARCA POR CONTENIDO BARATO

Intentar sustituir creatividad y calidad por IA para ahorrar suele salir mal. Los clientes notan la diferencia entre "apoyado por IA" y "generado sin cariño".

- **El desastre de "Willy Wonka":** en Glasgow, un organizador usó imágenes generadas con IA para vender una experiencia mágica de chocolate. En lugar de un mundo mágico, las familias se encontraron una nave casi vacía, con atrezzo barato y un personaje de terror inventado por IA llamado "El Desconocido" que hizo llorar a niños. [El evento se hizo viral, pero como acusación de estafa](#).
- **Crónicas deportivas "del infierno":** Gannett (USA Today) intentó automatizar periodismo local para deportes escolares con IA. Los textos eran tan rígidos y extraños que se convirtieron en meme. [Gannett paró el experimento](#) y pidió disculpas.
- **Libros falsos en el periódico: Chicago Sun-Times** publicó una [lista de recomendaciones creada por IA](#). Lo grave es que muchos libros y autores recomendados ni existían. El periódico perdió credibilidad de forma masiva.

Conclusión para su estrategia: use la IA como copiloto, no como piloto automático de su marca. La IA puede redactar borradores o aportar ideas, pero el control de calidad final y el "alma" del contenido deben seguir siendo humanos. El contenido barato le cuesta confianza, y la confianza es cara.

LECCIÓN 8: EL "VALLE INQUIETANTE": CUANDO EL CLIENTE SE INCOMODA EN LUGAR DE COMPRAR

El "valle inquietante" describe el efecto de una figura artificial que parece casi real, pero pequeños detalles (mirada "muerta", movimientos extraños) la vuelven inquietante. Los clientes no compran a

"zombis".

- **Toys "R" Us y Sora:** la cadena usó la nueva IA de vídeo "Sora" para [un anuncio](#). Para muchos espectadores, el anuncio se percibía sin alma y perturbador. Los rostros del niño se deformaban sutilmente y los movimientos eran demasiado fluidos, como "un sueño febril". En lugar de nostalgia, hubo críticas por el "fin de la creatividad".
- **Coca-Cola y "se acercan las fiestas":** Coca-Cola sustituyó su icónica campaña navideña [por una versión generada con IA](#). Los fans detectaron enseguida ruedas que cambiaban de forma mientras los camiones avanzaban y proporciones humanas incorrectas. El veredicto: "distópico" y "barato".
- **El "Pepperoni Hug Spot":** un [anuncio \(no oficial, pero viral\) generado por IA](#) para una pizzería ficticia mostraba a gente comiendo pizza. La IA no entendió cómo funcionan las bocas. Las mandíbulas se abrían como serpientes, las miradas estaban vacías y la pizza se fusionaba con la cara. Se convirtió en meme de "horror con IA".

Conclusión para su estrategia: en nuestro artículo sobre errores de [diseño web](#) advertimos sobre la "sobrecarga visual" y las "imágenes intercambiables". La generación con IA amplifica ese riesgo, crea visuales genéricos y, a menudo, con fallos que no transmiten identidad de marca. El usuario decide en milisegundos si confía. Un rostro inquietante destruye esa confianza al instante. Si parece barato, su marca se percibirá como barata.

CONCLUSIÓN: REÍRSE ESTÁ PERMITIDO, LA IMPRUDENCIA NO

Estas historias son divertidas, pero muestran un problema serio: muchas empresas implementan IA con prisas, sin entender los riesgos. **La IA escala todo, también los errores.** Si su proceso es débil, la IA solo acelerará un fracaso que ya estaba en camino.

La IA no es una varita mágica. Es una tecnología potente que requiere liderazgo, reglas claras y supervisión humana. En **nuestra agencia** lo vemos a diario: el éxito de proyectos de IA (ya sea en contenido, SEO o chatbots), no depende de la herramienta, sino de la **estrategia y de los estándares de seguridad** detrás.

¿Quiere usar IA sin aparecer en el próximo "best-of" de fallos? En **Wisea** le ayudamos a configurar

agentes y automatizaciones de IA para que sean seguras, coherentes con su marca y eficientes.

Contáctenos y lo planteamos: IA sin dolores de cabeza.