



ENTREGABLE 2.3

Informe de seguimiento de IA

Movimiento contra la Intolerancia

Número de Proyecto: 101213796
Convocatoria: CERV-2024-CHAR-LITI
Autoridad concedente: Agencia Ejecutiva Europea de Educación y Cultura
Duración del Proyecto: 24 meses



Cofinanciado por
la Unión Europea

Financiado por la Unión Europea. Sin embargo, las opiniones y puntos de vista expresados son exclusivamente los del autor o los autores y no reflejan necesariamente los de la Unión Europea ni los de EACEA. Ni la Unión Europea ni la autoridad que concede el préstamo se responsabilizan de ellos.

Detalles contractuales del proyecto

Título del proyecto	Red de Tolerancia contra los delitos de odio
Acrónimo del proyecto	ReTo
N.º del acuerdo de subvención	101213796
Fecha de inicio del proyecto	01.06.2025
Fecha de finalización del proyecto	31.05.2027
Duración	24 meses
Website	El proyecto no tendrá una página web específica, pero tiene una subpágina en: https://cifalmalaga.org/proyecto/reto/

Detalles del Plan

Paquete de trabajo No.	WP3	Tarea No.	T2.3
Título del documento	Informe de seguimiento de IA		
Nivel de difusión	Público		
Fecha de entrega	31/05/2026		

Colaboradores del documento

Responsable del documento	Movimiento contra la Intolerancia		
Colaboradores del documento	Organización	Revisores	Organización revisora
Valentín González Martínez	MCI	Alfiya Urazaeva Deborah Salafranca	CIFAL Malaga
Margarita González Muñoz	MCI	Elena María Mesa Alcaraz Alejandro Yankilevic	CIEDES

Descargo de responsabilidad

Financiado por la Unión Europea. Sin embargo, las opiniones y puntos de vista expresados son exclusivamente los del autor o los autores y no reflejan necesariamente los de la Unión Europea ni los de la Agencia Ejecutiva Europea en el Ámbito Educativo y Cultural (EACEA). Ni la Unión Europea ni la autoridad que concede el préstamo se responsabilizan de ellos.

Historial del documento

Form	Versión	Fecha	Autor	Organización
Borrador	1.0	15/05/2026	Valentín González	MCI
Borrador	2.0	25/05/2026	Alfiya Urazaeva y Deborah Salafranca Elena María Mesa Alcaraz	CIFAL Málaga CIEDES
FINAL	3.0	29/05/2026	Alfiya Urazaeva	CIFAL Málaga

Contenido

1. Introducción y contexto	6
1.1. Objetivos del informe	6
1.2. Contexto del Discurso de Odio en el entorno Digital: Evolución en el periodo actual en relación con IAs	8
1.3. Marco jurídico y ético	11
2. Análisis de la Herramienta iVERIFY (PNUD)	14
2.1 Arquitectura del sistema	14
2.2 Capacidad de monitoreo	16
2.3 Casos prácticos de aplicación	17
2.4 Indicadores y evaluación	18
3. Análisis de resultados y metodología del proceso de recogida de datos (WP2)	20
3.1. Introducción y Marco Estratégico	20
3.1.1. Alineación Institucional y Normativa	21
3.2. Arquitectura Metodológica: El Ciclo CRISP-DM Adaptado	21
3.2.1. Comprensión y Origen de los Datos	22
3.2.2. El Valor de la Contextualización Local	22
3.3. Ingeniería del Sistema y Pipeline Técnico	22
3.3.1. Fases de Implementación Técnica	22
3.3.2. Capas de la Arquitectura IA	23
3.4. Análisis de Resultados: Evolución del "Gold Dataset"	24
3.4.1. Composición y Calidad de la Muestra	24
3.4.2. Estrategia de Maximización de la Recuperación (Recall)	24
3.5. Optimización del Flujo y Umbrales de Eficiencia	25
3.6. Ética Algorítmica y Soberanía Tecnológica	25
3.7. Conclusión y Prospectiva Tecnológica	26

4. Prospectiva y Futuro de la IA contra el Odio	26
4.1 COMPARATIVA DE MODELOS: Grok, Deepseek, Gemini y Perplexity	28
4.1.1 Grok	28
4.1.2 DeepSeek	31
4.1.3 Gemini	33
4.1.4 Perplexity	35
5. Tendencias emergentes en IA y discurso de odio	37
5.1 Deepfakes, realidad distópica en tiempo presente	38
5.2 Superbots y enjambres de IA	39
6. IA y contranarrativa	40
7. Conclusiones	45
8. Fuentes y recursos en internet	48

1. Introducción y Contexto

1.1. Objetivos del informe

La irrupción de la Inteligencia Artificial ha transformado de raíz nuestra relación con el ecosistema digital. Más allá de ser una herramienta de productividad, se ha consolidado como un motor de creatividad, un puente hacia el conocimiento estructurado y un pilar esencial en el entorno laboral. En apenas unos años, su capacidad técnica ha crecido de forma exponencial, haciendo que cualquier predicción sobre su futuro —incluso a corto plazo— resulte un desafío a la imaginación. Bajo esta premisa, el presente informe se plantea como una radiografía actual que analiza el papel de la IA en la lucha contra el discurso de odio y explora el potencial de sus futuras aplicaciones.

El punto de partida de este informe se sitúa en el análisis exhaustivo de iVerify (la solución de inteligencia artificial desarrollada por el Programa de Naciones Unidas para el Desarrollo (PNUD)). A través del estudio de su marco conceptual, su aplicación práctica y los resultados obtenidos en su contexto histórico de implementación, buscamos extraer lecciones aprendidas y conclusiones transferibles que optimicen el desarrollo de la IA en el proyecto ReTo. No obstante, este ejercicio de adaptación reconoce una distinción jurídica fundamental: mientras que iVerify responde a mandatos de gobernanza global, el proyecto ReTo se rige por un marco legal sustancialmente distinto —el europeo—, lo que exige una alineación específica con las garantías de privacidad y ética digital vigentes en la Unión.

Si bien este informe no tiene como objeto central el desarrollo exhaustivo de los aspectos jurídicos —extensamente analizados en entregables anteriores del Proyecto ReTo— resulta imprescindible referenciar el ecosistema normativo que fundamenta nuestra intervención. Este marco se sustenta en una arquitectura de tres niveles: la protección del Código Penal de España, el carácter preventivo y reparador de la Ley de Igualdad de Trato y No Discriminación, y el desarrollo legislativo específico de la Comunidad Autónoma Andaluza.

No obstante, la naturaleza tecnológica de este estudio hace imperativo profundizar en el impacto del Acta de Servicios Digitales (DSA) y, muy especialmente, en el Acta de Inteligencia Artificial (AI Act) de la Unión Europea. Estas regulaciones no solo circunscriben el marco legal de actuación, sino que redefinen la propia concepción del discurso de odio en el entorno digital. Su

análisis es clave para aterrizar los estándares de seguridad, ética y supervisión humana en el contexto andaluz, garantizando que el uso de herramientas de IA en el proyecto ReTo no sea solo eficaz en la detección, sino plenamente garantista con los derechos fundamentales de la ciudadanía.

Por tanto, conviene reiterar que este análisis no se limitará al estudio de la herramienta de Naciones Unidas, sino que aportará material relativo a la Inteligencia Artificial desarrollada en el marco del proyecto ReTo, permitiendo dilucidar aspectos potenciales de mejora en la lucha contra el discurso de odio; con el desafío de hacerlo en un entorno fenomenológico complejo, marcado por procesos de polarización y una exacerbación del discurso de odio a escala global que impacta con mucha fuerza también en España.

Este informe sitúa la descripción técnica y operativa de la IA de ReTo como su eje sustancial, exponiendo con rigor su arquitectura, metodologías de entrenamiento y los resultados de las fases preliminares. Los datos actuales, más allá de lo experimental, confirman el potencial del sistema para la detección precoz de narrativas de odio. El objetivo es consolidar una metodología auditable y escalable que identifique y categorice con rigor las manifestaciones de intolerancia recogidas en el artículo 510 del Código Penal, protegiendo así los derechos fundamentales frente al racismo, la xenofobia y otras formas de discriminación.

Existe una dualidad tecnológica crítica: mientras que los algoritmos de las redes sociales suelen priorizar la polarización por rentabilidad, el Proyecto ReTo reorienta el *Machine Learning* hacia la auditoría semántica y la ética. El proyecto no solo busca detectar el odio, sino neutralizarlo mediante una simbiosis entre tecnología y creatividad humana que fomente un discurso de respeto y libertad de expresión.

Para lograrlo, el modelo de respuesta proactiva se articula en tres líneas de acción integradas:

- Contranarrativa basada en datos: Identifica la estructura de los discursos de odio para diseñar respuestas precisas que neutralicen su impacto en los mismos canales donde se originan.
- Sensibilización preventiva: Utiliza el análisis algorítmico para detectar eventos catalizadores, permitiendo lanzar campañas de concienciación antes de que las narrativas de intolerancia se consoliden.
- Resiliencia digital: Emplea los datos recabados para instruir a la ciudadanía en la identificación de la manipulación algorítmica, fortaleciendo el pensamiento crítico y la respuesta democrática frente a la desinformación.

En este contexto, la intención es definir la 'foto fija' del impacto de la IA en el fenómeno de los delitos y discurso de odio, en los diferentes aspectos que su estudio puede ser relevante: recogida de datos que perfilan posibles situaciones de discriminación, detección de discurso de odio, criterios de moderación; y potencial de la IA en creación de contra narrativas en campañas y proyectos de sensibilización preventiva.

Para ello, este informe ofrece una aproximación analítica a cuatro modelos de inteligencia artificial —Gemini, Grok, Perplexity y DeepSeek— y su relación directa con el discurso de odio. Se examina la problemática que estos sistemas plantean en el ámbito de la comunicación social, prestando especial atención a su capacidad para generar un impacto tóxico en el seno de sociedades actualmente polarizadas.

El objetivo de incluir a DeepSeek, de carácter chino, junto a los demás modelos es recoger diferentes espectros de las IA según los valores predominantes de sus sociedades de origen. De este modo, el informe analiza cómo la procedencia y el marco ético de cada tecnología influyen en la difusión de contenidos y en la configuración del discurso público global.

Pese a su carácter eminentemente técnico, este informe mantiene la vocación de abordar los criterios éticos que rodean a la Inteligencia Artificial, especialmente mediante el análisis de la denominada "paradoja tecnológica". Este concepto aporta una profundidad crítica esencial al documento, al exponer la naturaleza ambivalente de las herramientas algorítmicas en la sociedad actual.

1.2. Contexto del Discurso de Odio en el entorno Digital: Evolución en el periodo actual en relación con [IAs](#)

El auge de las redes sociales ha transformado profundamente la participación en la esfera pública, planteando retos inéditos para la protección de los derechos humanos y el mantenimiento de la tolerancia como valor predominante de la sociedad. En el periodo actual, el discurso de odio ha dejado de ser un fenómeno de crisis puntuales para convertirse en una manifestación estructural y persistente de intolerancia.

Esta evolución está intrínsecamente ligada al desarrollo de la Inteligencia Artificial (IA), tecnología que actúa simultáneamente como un vector para la propagación masiva de hostilidad y como el principal mecanismo para su detección y mitigación. Una de las mutaciones más críticas es la aparición de los "enjambres

de IA" (AI swarms); sistemas multi-agente que generan identidades sintéticas hiperrealistas capaces de infiltrarse en comunidades digitales.

Estas redes coordinan narrativas de intolerancia a una velocidad extraordinaria, realizando experimentos en tiempo real para determinar qué mensajes de división son más persuasivos, creando un falso sentido de consenso social en torno a ideas extremistas que vulneran la dignidad humana.

Esta propagación se ve potenciada por la arquitectura de las plataformas, diseñada bajo la denominada economía de la indignación, que a menudo colisiona con el respeto a los derechos humanos más fundamentales. Los algoritmos priorizan el compromiso del usuario (*engagement*), tendiendo a promover contenido emocionalmente cargado, ya que la hostilidad es un motor potente de interacción que erosiona los valores de tolerancia y democracia.

Este diseño facilita la creación de cámaras de eco u homofilia, donde las visiones radicales no son cuestionadas y los grupos se aíslan en burbujas informativas de intolerancia mutua. Diversas investigaciones señalan que este "desorden informacional" debilita los principios democráticos y erosiona la confianza en las instituciones, transformando el debate público en un escenario donde el odio desplaza al diálogo constructivo.

En la actualidad, las narrativas de intolerancia han evolucionado hacia una complejidad multimodal apoyada en la IA generativa. El uso de deepfakes y contenido sintético dificulta enormemente la distinción entre lo real y lo ficticio, alimentando teorías de conspiración con tintes xenófobos y antisemitas que atentan contra la tolerancia religiosa y étnica.

En el contexto español, se observa que el 34% de los contenidos de odio vinculan falsamente la inmigración con la inseguridad, una narrativa alimentada por bulos estructurales que ignoran los derechos humanos de las personas migrantes. Otros focos críticos incluyen el racismo en el ámbito deportivo —con episodios prototípicos de intolerancia hacia figuras como Vinícius Jr.— jugador del Real Madrid que sufre habitualmente insultos y menosprecio racista, tanto en los campos de fútbol como en internet; y el resurgimiento de la islamofobia y el antisemitismo derivados de conflictos en Oriente Medio. Asimismo, el odio de género presenta una violencia simbólica específica, que se ensaña con mujeres que sufren discriminación múltiple en tanto que musulmanas, judías, afrodescendientes, o gitanas, entre otros colectivos.

El impacto en los menores de edad representa una de las mayores amenazas para la salvaguarda de sus derechos humanos. La IA introduce riesgos severos, desde la

creación de "deepnudes" sexualizados hasta la radicalización mediante estéticas visuales diseñadas para normalizar la intolerancia y hacer la propaganda extremista más atractiva para los jóvenes.

Se ha detectado, además, que herramientas como Snapchat "My AI" o Meta AI presentan deficiencias en sus barreras de seguridad, permitiendo en ocasiones que los menores accedan a contenidos que socavan la tolerancia hacia la diversidad. A pesar de esfuerzos tecnológicos como el Sistema FARO en España, una metodología de monitorización diseñada para la identificación y el análisis a tiempo real de los contenidos de discurso de odio con motivaciones racistas, xenófobas, islamófobas, antisemitas y antigitanas en las principales redes sociales: Facebook, Instagram, TikTok, YouTube y X (antes Twitter). Esta iniciativa representa un paso adelante en la lucha contra la intolerancia en el entorno digital. La moderación enfrenta el reto de la ironía y la jerga local, elementos que los agresores utilizan estratégicamente para difundir su intolerancia eludiendo la censura automática.

Finalmente, persiste el dilema del sesgo algorítmico y el marco regulatorio que debe proteger los derechos humanos. Los sistemas de detección pueden heredar prejuicios de sus datos de entrenamiento, incurriendo en la discriminación de grupos minoritarios o en fallos de IA generativa donde los modelos muestran inconsistencias ideológicas que operan en contra de valores de tolerancia y aprecio por la diversidad. . Ante este escenario, la Unión Europea ha respondido con el Reglamento de Inteligencia Artificial y la Ley de Servicios Digitales (DSA), buscando un equilibrio entre seguridad y libertad de expresión.

No obstante, el riesgo de "overblocking" o censura sigue vigente, lo que podría afectar irónicamente los derechos humanos, al no encontrar un equilibrio sano y riguroso que delimite la libertad de expresión de la libertad de agresión.

La relación entre la IA y el odio es, en última instancia, dual: mientras la tecnología escala la sofisticación del ataque, también ofrece la única vía factible para la monitorización en un entorno de datos explosivo, exigiendo un enfoque que combine la solución técnica con una cultura de tolerancia y alfabetización mediática.

1.3. Marco jurídico y ético

- Ley de Inteligencia Artificial (AI Act) de la UE

1.3.1. Contexto y Origen

La Ley de IA fue propuesta por la Comisión Europea en abril de 2021. Tras años de negociaciones entre el Parlamento Europeo y el Consejo, fue aprobada formalmente en marzo de 2024 y publicada en el Diario Oficial de la UE en julio de 2024. Su aplicación es progresiva, con las primeras prohibiciones entrando en vigor a los seis meses de su publicación.

La intención del legislador es doble: por un lado, proteger los derechos fundamentales, la seguridad y la salud de los ciudadanos europeos; y por otro, fomentar la inversión e innovación en IA dentro de un mercado único europeo predecible, evitando la fragmentación legal entre países miembros.

1.3.2. Clasificación basada en el Riesgo

El Reglamento de Inteligencia Artificial de la UE establece un enfoque proporcional basado en el riesgo, lo que significa que la intensidad de la regulación no es uniforme, sino que aumenta progresivamente según el peligro potencial que el sistema represente para los derechos y la seguridad de los ciudadanos.

En la cúspide de esta pirámide se encuentran los sistemas de Riesgo Inaceptable, los cuales están totalmente prohibidos por atentar contra los valores fundamentales de la Unión Europea. Esta categoría incluye prácticas como la manipulación cognitiva-conductual, el "scoring" o puntuación social por parte de los Estados, y la clasificación biométrica basada en datos sensibles como la raza o la religión. Asimismo, se prohíbe la identificación biométrica remota en tiempo real en espacios públicos, permitiendo únicamente excepciones extremadamente tasadas vinculadas a la seguridad nacional.

Un escalón por debajo se sitúan los sistemas de Alto Riesgo, que son aquellos que, sin estar prohibidos, tienen un impacto significativo en la vida y el futuro de las personas. Debido a su criticidad, estos sistemas están sujetos a requisitos legales muy estrictos antes de poder comercializarse. Se incluyen aquí las herramientas de IA utilizadas en la gestión de infraestructuras críticas (como el suministro de agua o electricidad), en el ámbito de la educación y la formación profesional, y de manera muy relevante, en el empleo y la selección de personal. También se

consideran de alto riesgo las aplicaciones destinadas al acceso a servicios esenciales, como la concesión de créditos bancarios o diagnósticos sanitarios, así como los sistemas empleados en la aplicación de la ley y el control de fronteras.

Para las aplicaciones que presentan un Riesgo de Transparencia Limitada, la normativa se enfoca principalmente en el derecho a la información del usuario. En este grupo se encuentran herramientas populares como los chatbots o los generadores de contenidos multimedia conocidos como *deepfakes*. La ley no limita su uso técnico, pero impone la obligación de que el usuario sea informado de que está interactuando con una máquina o de que el contenido visual o sonoro que consume ha sido manipulado artificialmente.

Finalmente, el reglamento identifica una categoría de Riesgo Mínimo o Nulo, donde se clasifican la gran mayoría de las aplicaciones de IA que utilizamos en nuestra vida cotidiana, tales como los filtros de correo no deseado (spam) o los sistemas de inteligencia artificial en videojuegos. Para estas tecnologías, la ley no impone obligaciones legales adicionales ni cargas burocráticas, aunque se fomenta la adhesión a códigos de conducta voluntarios para mantener altos estándares de calidad y ética, permitiendo que la innovación siga su curso sin trabas regulatorias innecesarias. En cuanto a los Modelos de IA de Propósito General (GPAI), la ley introduce reglas específicas para aquellas herramientas versátiles capaces de realizar una amplia gama de tareas, como los grandes modelos de lenguaje. La normativa pone el foco en la potencia de cómputo: si un modelo ha sido entrenado con capacidades superiores a los 10^{25} FLOPs, se clasifica automáticamente como un sistema de riesgo sistémico. Estos modelos están sujetos a obligaciones más estrictas, que incluyen la realización de evaluaciones de modelos, la mitigación de riesgos adversos y el reporte detallado de cualquier incidente grave a las autoridades competentes.

La gobernanza y supervisión de esta ley se articula a través de una estructura dual. A nivel central, se ha creado la Oficina de IA dentro de la Comisión Europea, encargada de supervisar los modelos de propósito general más avanzados y de coordinar la política de IA en toda la Unión. Por otro lado, cada Estado miembro tiene la obligación de designar una autoridad nacional de control, que actuará como el vigilante local para asegurar que las empresas cumplan con el reglamento en su territorio.

Un pilar fundamental para la protección del ciudadano es el Derecho a la Explicación. La ley garantiza que cualquier persona afectada por una decisión tomada por un sistema de IA de alto riesgo —como el rechazo de una solicitud de empleo o de un crédito bancario— tenga derecho a recibir una explicación clara,

comprensible y detallada sobre el funcionamiento del algoritmo y los datos que llevaron a ese resultado. Esto busca evitar la opacidad de las llamadas "cajas negras" tecnológicas y asegurar que los derechos legales de los ciudadanos sean respetados.

1.3.3. Sesgo Discriminatorio

La ley introduce medidas específicas para evitar que el sesgo algorítmico perjudique a colectivos vulnerables en el momento de recibir estas explicaciones. Aquí detallo cómo se integran las medidas antidiscriminación en este punto concreto.

1.3.4. Detección de "Sesgos Ocultos"

El derecho a recibir una explicación permite que el ciudadano pueda identificar si la decisión se ha basado en variables indirectas (proxies) que ocultan discriminación. Por ejemplo:

- Si un algoritmo deniega un crédito basándose en el "código postal", la explicación detallada permitiría detectar si se está utilizando la ubicación como un sustituto discriminatorio de la raza o el nivel socioeconómico.

1.3.5. Calidad de los Datos (Artículo 10)

Para que el derecho a la explicación sea real, la ley impone obligaciones previas a los desarrolladores de sistemas de alto riesgo:

- **Representatividad:** Los conjuntos de datos deben ser "suficientemente representativos" y estar libres de errores para minimizar el riesgo de resultados discriminatorios.
- **Uso de datos sensibles para corregir sesgos:** De manera excepcional, la ley permite procesar categorías especiales de datos (como origen étnico o religión) solo si es estrictamente necesario para detectar y corregir sesgos en el sistema.

1.3.6. Información sobre los "Factores Determinantes"

La explicación no puede ser una respuesta técnica ininteligible. Debe incluir:

- Los datos de entrada específicos que influyeron en la decisión.
- La lógica del algoritmo explicada de forma clara.

Esto permite al ciudadano verificar si se han vulnerado los derechos protegidos por el Artículo 510 del Código Penal de España (racismo, antisemitismo, orientación sexual, etc.) o la normativa de igualdad de la UE.

Para equilibrar la seguridad con la competitividad, la normativa introduce los Sandboxes Regulatorios (espacios controlados de pruebas). Estos son entornos seguros creados por los Estados miembros para permitir que las empresas, especialmente las pequeñas y medianas empresas (PYMES) y las *startups*, puedan desarrollar y probar sus innovaciones bajo la supervisión directa de los reguladores. El objetivo es que la innovación no se detenga por el miedo a incumplir la norma, permitiendo que los errores se detecten y corrijan antes de que el producto llegue al mercado real.

1.3.7. Régimen Sancionador

Finalmente, la Ley de IA establece un régimen sancionador extremadamente severo para garantizar su cumplimiento, con multas que pueden superar incluso las impuestas por el RGPD. Las sanciones son escalables y dependen de la gravedad de la infracción y del tamaño de la empresa. En los casos más graves, como el uso de sistemas de IA que incurran en prácticas prohibidas (riesgo inaceptable), las multas pueden ascender hasta los 35 millones de euros o el 7% de la facturación anual global de la empresa en el ejercicio anterior, aplicándose siempre la cuantía que sea mayor.

2. Análisis de la Herramienta iVERIFY (PNUD)

2.1 Arquitectura del sistema

iVerify fue concebido y desarrollado por el Programa de las Naciones Unidas para el Desarrollo (PNUD), específicamente a través del Grupo de Trabajo Conjunto sobre Asistencia Electoral (con sede en Bruselas) y la Oficina Digital Principal. Su

creación surgió como una respuesta directa a las peticiones recurrentes de las oficinas de país del PNUD y sus contrapartes nacionales, quienes necesitaban herramientas prácticas para gestionar la "contaminación informativa" y el discurso de odio durante los procesos electorales. El sistema se diseñó bajo el paradigma de Bien Público Digital (DPG), lo que garantiza que su código sea abierto, transparente y fácilmente adaptable a diferentes contextos locales y lingüísticos.

La arquitectura técnica de iVerify se fundamenta en una suite de microservicios que integran inteligencia artificial avanzada con procesos humanos de verificación. Entre sus componentes tecnológicos destacan el Aprendizaje Automático (Machine Learning) y el Procesamiento de Lenguaje Natural (NLP), herramientas esenciales para procesar y analizar grandes volúmenes de información de manera eficiente.

Dentro de su sistema de detección automatizada, la plataforma utiliza el algoritmo de código abierto Detoxify. Esta herramienta emplea aprendizaje automático para escanear y clasificar contenido tóxico y discursos de odio en redes sociales como Facebook, Instagram y Reddit. Para mayor precisión, se complementa con la API Perspective, que asigna puntuaciones de toxicidad a los contenidos analizados, permitiendo una identificación más rigurosa de las narrativas dañinas.

Para la gestión de la verificación, iVerify integra la plataforma Meedan Check, la cual opera bajo el modelo de "humano en el circuito" (human-in-the-loop). Esta tecnología permite identificar y emparejar contenidos que ya han sido desmentidos anteriormente; de este modo, el sistema evita que los equipos humanos tengan que revisar la misma información falsa de manera redundante, optimizando los recursos y el tiempo de respuesta.

Finalmente, el proceso cuenta con un componente denominado "Triage bot", que actúa como la primera línea de defensa. Este bot se encarga de clasificar y procesar grandes volúmenes de datos públicos y alertas ciudadanas enviadas vía WhatsApp o SMS, logrando identificar narrativas potencialmente peligrosas en tiempo real para su posterior análisis y neutralización. El Triage bot funciona como una ventanilla única de recepción. Por un lado, "escucha" las redes sociales de forma masiva y, por otro, recibe mensajes directos de la ciudadanía a través de WhatsApp o SMS. Su función principal es la discriminación de relevancia: separa el ruido (comentarios irrelevantes, spam o mensajes personales) de las alertas que realmente contienen una narrativa de desinformación o un discurso de odio. Sin este bot, los verificadores se verían desbordados por miles de mensajes imposibles de procesar manualmente.

Una vez que el bot captura una información, no se limita a guardarla; realiza una categorización automática utilizando etiquetas según la urgencia o el tema (por ejemplo: "fraude electoral", "violencia de género" o "ataque a candidatos"). Gracias a su integración en el algoritmo Detoxify, el bot puede asignar un nivel de riesgo. Si

un mensaje contiene un lenguaje extremadamente tóxico o se está volviendo viral rápidamente, el Triage bot lo sitúa en la parte superior de la lista de tareas para que sea lo primero que revisen los humanos.

Una de las funciones más valiosas para la eficiencia es la detección de contenido recurrente. Si mil personas envían el mismo video manipulado por WhatsApp, el bot identifica que se trata de la misma narrativa y las agrupa bajo un solo caso. Esto evita que el equipo de iVerify pierda tiempo analizando mil veces lo mismo, permitiéndoles emitir una única respuesta o desmentido que cubra todas las alertas recibidas.

En última instancia, el Triage bot actúa como un asistente de triaje médico: realiza el diagnóstico inicial rápido para que los especialistas (los verificadores humanos) puedan concentrar sus esfuerzos en la neutralización. Al presentar la información ya organizada, contextualizada y priorizada, el bot reduce drásticamente el tiempo que transcurre desde que surge una mentira peligrosa hasta que iVerify publica una alerta oficial para proteger a la población.

2.2 Capacidad de monitoreo

En cuanto a sus capacidades de monitoreo, iVerify emplea un modelo híbrido que integra el rastreo automatizado con la participación ciudadana activa. El sistema realiza un seguimiento diario de redes sociales y foros como Facebook, Instagram y Reddit a través de herramientas como CrowdTangle, además de escanear constantemente la prensa digital para identificar narrativas coordinadas de desinformación. Asimismo, cuenta con canales de denuncia ciudadana (Tiplines) que permiten al público enviar contenidos sospechosos mediante WhatsApp, Facebook Messenger y X. Para potenciar este análisis, utiliza los algoritmos de detección de patrones y sentimientos de la API Perspective y Detoxify, que asignan puntuaciones de toxicidad y rastrean "artefactos" específicos como etiquetas o grupos para observar la evolución de las campañas de desinformación.

La mitigación del discurso de odio es una prioridad central de la herramienta, entendiendo este fenómeno como cualquier ataque basado en intolerancia que pueda derivar en violencia contra personas o inestabilidad política. Para combatirlo, el algoritmo Detoxify identifica patrones de hostilidad y degradación. El PNUD seleccionó esta herramienta por su enfoque en limitar el sesgo algorítmico, protegiendo así a grupos potencialmente víctimas de la intolerancia. No obstante, para superar la "brecha lingüística" de la IA —como dificultades con el sarcasmo o dialectos locales—, el sistema exige siempre que un verificador humano valide las detecciones antes de emitir cualquier alerta pública.

Finalmente, la metodología de verificación y rigor procedimental se sustenta en un protocolo de "triple verificación" basado en cinco pilares.

1. **Se distingue entre hechos y opiniones**, asegurando que la afirmación sea verificable.
2. **Se realiza una consulta de fuentes relevantes**, contactando a autores y expertos.
3. **Se mitigan sesgos** garantizando que cada historia sea revisada por tres verificadores distintos.
4. **Se emiten informes** basados en evidencia con calificaciones de Verdadero, Falso o Engañoso.
5. **Se mantiene una vigilancia dinámica**, actualizando los reportes si surge nueva evidencia que complemente los hallazgos previos.

2.3 Casos prácticos de aplicación

El despliegue de iVerify ha demostrado una notable adaptabilidad en diversos contextos sociopolíticos, consolidándose como una herramienta clave para la integridad democrática en las siguientes regiones:

- **Zambia (2021):** Fue el primer piloto de la herramienta durante las elecciones generales. En este país, el equipo colaboró estrechamente con la Comisión Electoral para formar a periodistas en la identificación de noticias falsas. Un logro destacado fue la capacidad de persuasión del sistema, logrando que varios productores de contenido eliminaran información errónea de forma voluntaria tras la publicación de los reportes de veracidad de iVerify.
- **Sierra Leona (2023):** En este contexto, iVerify se destacó por su integración con la radio, el medio de comunicación principal para el 70% de la población que carece de acceso a internet. El proyecto logró verificar 170 historias y se le atribuye haber contribuido a una participación electoral del 83%, al proporcionar información confiable que redujo la incertidumbre entre los votantes.
- **Liberia (2023):** Implementado a través de la red Local Voices Liberia, el sistema fue crucial para desmentir campañas virales que amenazaban con encender tensiones violentas. En un entorno de paz frágil, la rápida neutralización de estas narrativas ayudó a mantener la estabilidad social durante el proceso de votación.

- **Kenia (2022):** La herramienta se utilizó para monitorear la desinformación dirigida específicamente contra candidatos electorales. Una de las lecciones aprendidas más importantes de este caso fue la recomendación técnica de iniciar la implementación del sistema con una antelación de 6 a 9 meses antes de los comicios, con el fin de generar una base sólida de confianza entre los actores locales y las instituciones.

Ámbito Europeo

iVerify se consolida como una iniciativa conjunta entre la Comisión Europea y el PNUD para proteger la integridad electoral, respaldada por una trayectoria de cooperación con más de 200 proyectos de asistencia desde 2006. Su implementación en suelo europeo aprovecha esta experiencia previa para establecer un sistema robusto de verificación y respuesta ante las amenazas informativas.

Este sistema se alinea estrictamente con el Plan de Acción para la Democracia Europea de la Comisión, buscando una respuesta coordinada frente a la desinformación. A través de marcos legislativos y éticos comunes, la iniciativa supervisa el ámbito digital para garantizar la transparencia y la seguridad en los procesos democráticos del continente. A continuación exponemos dos ejemplos de implementación.

- **Despliegue en Macedonia del Norte:** Este país ha sido identificado explícitamente como uno de los puntos estratégicos donde iVerify está programado para su implementación o ya se encuentra activo durante 2024. Esta acción se complementa con programas diseñados para fortalecer los consejos municipales en materia de gobernanza democrática.
- **Iniciativas en Serbia:** En este contexto, se destaca la firma de un protocolo de cooperación para el desarrollo de un modelo de lenguaje extenso (LLM) nacional para el idioma serbio. Este esfuerzo tecnológico está orientado a mejorar la gobernanza y garantizar la integridad de la información en la región.

2.4 Indicadores y evaluación

El desempeño de iVerify se evalúa mediante un marco de seguimiento estructurado en tres pilares fundamentales que miden desde la preparación logística hasta el impacto democrático a largo plazo. En la fase de Factibilidad e Implementación Inicial, el enfoque reside en asentar una infraestructura profesional robusta a través de la elaboración de diagnósticos sobre el ecosistema

de información local y la formalización de alianzas estratégicas mediante Memorandos de Entendimiento. Este pilar también prioriza la capacitación técnica de los verificadores contratados y la sensibilización de la prensa local, estableciendo redes de coordinación con reuniones periódicas para asegurar que todos los actores comprendan el funcionamiento del sistema antes del inicio de los comicios.

En cuanto al Despliegue y Funcionamiento Operativo, los indicadores se centran en la interacción dinámica entre el sistema y la sociedad durante el periodo crítico electoral. Aquí resulta vital medir la participación ciudadana a través del volumen de alertas enviadas por los votantes mediante canales como WhatsApp o SMS, así como la productividad del equipo en la publicación de historias verificadas. El éxito de este pilar se refleja en el impacto digital de los desmentidos, la frecuencia con la que iVerify es citado en medios de comunicación tradicionales y la capacidad técnica para detectar tendencias de desinformación emergentes, logrando neutralizarlas mediante reportes públicos basados en evidencia.

Finalmente, el eje de Integridad Técnica y Sostenibilidad garantiza que la herramienta sea eficiente y perdurable. Esto implica monitorear la precisión de la inteligencia artificial, midiendo tiempos de respuesta y la optimización de bases de datos mediante sistemas de etiquetado automático. Además de analizar métricas de tráfico web para entender el comportamiento del usuario, el PNUD pone especial énfasis en el uso de la plataforma "fuera de temporada", lo que demuestra que iVerify no es solo un recurso temporal, sino un servicio de gobernanza continua. El logro final de esta etapa es la firma de compromisos institucionales y acuerdos de financiamiento que aseguren la permanencia del proyecto más allá del ciclo electoral inmediato.

El proyecto ReTo toma como referencia conceptual varias de las líneas metodológicas promovidas por iVerify en relación con la monitorización de contenidos online, la detección temprana de discursos problemáticos y la combinación de procesos automatizados con validación humana. En este sentido, ambos enfoques comparten principios comunes vinculados al análisis de contenido digital, el apoyo a la identificación de riesgos y la generación de información estructurada para la toma de decisiones basada en evidencia.

Sin embargo, durante el desarrollo del proyecto y la implementación práctica del piloto en Andalucía, identificamos una serie de necesidades operativas, lingüísticas y metodológicas específicas que requerían un mayor nivel de adaptación y flexibilidad técnica. Entre ellas destacan la necesidad de trabajar sobre contenidos en español y en contextos socioculturales locales, la incorporación de categorías adaptadas a un marco metodológico propio, la integración de otro tipo de datos, además de los recopilados en redes sociales, tales como estadísticas oficiales del Ministerio del Interior y la Fiscalía General del Estado, y la posibilidad de

implementar procesos completos de trazabilidad y auditoría del etiquetado y clasificación.

Por este motivo, el proyecto no se apoyó directamente iVerify como herramienta central de operación, sino que desarrolló una arquitectura complementaria propia, diseñada específicamente para responder a los requerimientos concretos de la prueba piloto en Andalucía. Esta arquitectura mantiene una alineación funcional con los principios generales promovidos por iVerify, especialmente en lo relativo a la monitorización de contenido, la detección automatizada preliminar y el enfoque híbrido de IA + supervisión humana.

La solución implementada en ReTo incorpora procesos de captura multifuente, preprocesamiento automatizado, filtrado mediante librerías optimizadas de términos, pre-etiquetado asistido por IA y validación humana experta, siguiendo un enfoque “human-in-the-loop” orientado a maximizar la calidad metodológica, reducir falsos positivos y garantizar consistencia en la clasificación de mensajes. Asimismo, el sistema integra mecanismos de anonimización, registros de auditoría y trazabilidad completa de cada lote procesado, permitiendo reconstruir íntegramente el flujo de análisis y etiquetado de cada mensaje.

En consecuencia, la herramienta desarrollada por ReTo no debe entenderse como una alternativa desvinculada de iVerify, sino como una implementación complementaria y contextualizada, construida sobre principios metodológicos similares pero adaptada a las necesidades técnicas, institucionales y analíticas específicas del proyecto y del entorno andaluz.

3. Análisis de resultados y metodología del proceso de recogida de datos (WP2)

3.1. Introducción y Marco Estratégico

El entregable 2.1 del Proyecto ReTo, titulado *"Informe sobre la metodología de recogida de datos"*, constituye la piedra angular estratégica y técnica sobre la cual se edifica toda la iniciativa orientada al desarrollo de software de detección de odio. Este documento no es una mera descripción de tareas, sino un marco normativo y operativo que garantiza que el flujo de datos alimente a la Inteligencia Artificial bajo principios de calidad, transparencia y rigor científico.

La calidad, la transparencia y el rigor científico constituyen los pilares de este sistema. Su propósito fundamental radica en la delimitación de parámetros conceptuales para la identificación precisa del discurso y los delitos de odio,

funcionando como un marco de referencia que previene la arbitrariedad algorítmica.

Al operar bajo estándares éticos validados por la supervisión humana, el sistema asegura que la tecnología no solo sea eficiente, sino que la definición de estos criterios hace que la herramienta sea también más precisa y técnicamente fiable. De este modo, se garantiza que el procesamiento de datos responda a una metodología rigurosa y auditable en la protección de los derechos fundamentales.

3.1.1 Alineación Institucional y Normativa

Para dotar al sistema de validez administrativa, el proyecto se ha alineado con el II Plan de Acción de Lucha contra los Delitos de Odio (2022-2024) del Ministerio del Interior de España. Se han adoptado 18 categorías oficiales de discriminación (como racismo, islamofobia, aporofobia y delitos contra la identidad de género), lo que garantiza que los datos recogidos sean compatibles con el Sistema de Gestión Operativa (SIGO) de las Fuerzas y Cuerpos de Seguridad, facilitando la futura transferencia tecnológica y la coordinación con la Oficina Nacional de Lucha contra los Delitos de Odio.

3.2. Arquitectura Metodológica: El Ciclo CRISP-DM Adaptado

La arquitectura operativa de ReTo se fundamenta en la metodología CRISP-DM (*Cross-Industry Standard Process for Data Mining*), adaptada específicamente para la complejidad de las ciencias sociales y la detección de intolerancia.



3.2.1. Comprensión y Origen de los Datos

El enfoque es de naturaleza mixta, combinando fuentes *offline* y *online*. Mientras que las denuncias policiales aportan el contexto legal y el impacto físico, la extracción de redes sociales (X y YouTube) ofrece la dimensión digital del odio. Este ecosistema se enriquece con una perspectiva múltiple, un elemento crítico que permite al sistema evaluar cómo la convergencia de diversas identidades —tales como género, etnia, discapacidad y situación socioeconómica— intensifica el daño hacia la víctima, evitando así análisis reduccionistas que podrían ignorar casos de discriminación múltiple.

3.2.2. El Valor de la Contextualización Local

Un factor diferenciador del WP2, liderado por CIEDES, es su procesamiento lingüístico refinado. Se ha detectado, mediante el análisis de fuentes abiertas y lingüística de corpus, que los modelos estándar (como GPT-4 base o Perspective API) sufren de sesgos geográficos. El Proyecto ReTo ha ajustado sus librerías para reconocer modismos y expresiones culturales propias de Andalucía. Esto previene que el sistema clasifique erróneamente términos de la jerga local o que, por el contrario, deje pasar odio "encubierto" bajo formas de expresión regionalista.

3.3. Ingeniería del Sistema y Pipeline Técnico

El sistema ha sido diseñado como una infraestructura de Priorización Híbrida Humano-IA. A diferencia de los modelos de moderación puramente automatizados, este actúa como un catalizador del esfuerzo humano, proporcionando señales de probabilidad sin sustituir el juicio cualitativo del analista.

3.3.1. Fases de Implementación Técnica

La implementación se despliega en cinco etapas críticas:

- **Recolección Masiva (Scraping):** Monitoreo constante de X y YouTube mediante APIs y técnicas de extracción automatizada.
- **Etiquetado y Almacenamiento:** Uso de versiones estáticas para asegurar la trazabilidad y reproducibilidad de los resultados.

- Procesamiento y Limpieza: Refinamiento de los datos mediante técnicas de Procesamiento de Lenguaje Natural (PLN).
- Visualización: Creación de *dashboards* interactivos para que instituciones y ONGs monitoricen tendencias en tiempo real.
- Automatización e Integración: Conexión de bases de datos entre los socios del proyecto para asegurar la escalabilidad.

3.3.2. Capas de la Arquitectura IA

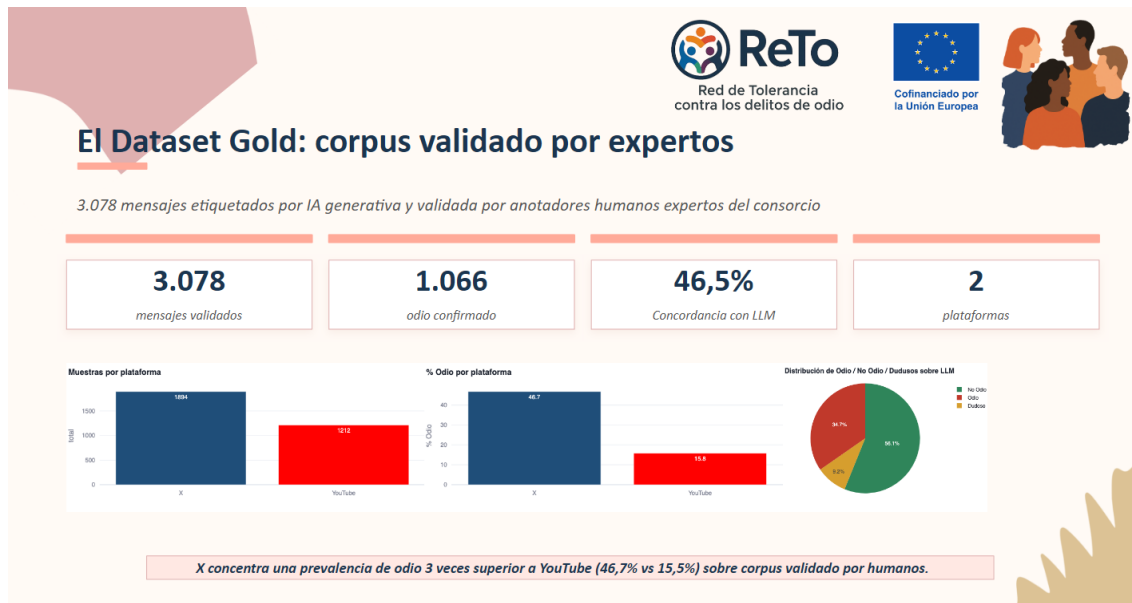
La arquitectura técnica se estructura en tres capas sucesivas:

- Capa 1: Modelo Baseline (Regresión Logística + TF-IDF). Se ha priorizado este enfoque por su linealidad y transparencia diagnóstica. Al ser un modelo de "caja blanca", permite controlar los sesgos desde la raíz, algo fundamental para cumplir con la Ley de Inteligencia Artificial de la UE (AI Act) en sistemas de alto riesgo.
- Capa 2: Asistencia por Modelos de Lenguaje (LLM). El uso de modelos como GPT-5.2 ayuda a identificar la intensidad y la categoría específica del mensaje, aportando una profundidad de comprensión que el modelo clásico no alcanza.
- Capa 3: Validación Experta Final. La decisión final es humana, cerrando el ciclo de retroalimentación y garantizando que la IA aprenda constantemente de las correcciones de los analistas.



3.4. Análisis de Resultados: Evolución del "Gold Dataset"

El Gold Dataset, considerado la referencia de "verdad absoluta" por su validación humana, ha experimentado un crecimiento exponencial hasta consolidar 2.928 registros. Todo ello desde junio de 2025, hasta la fecha de realización de este informe en mayo de 2026.



3.4.1. Composición y Calidad de la Muestra

De este volumen, la plataforma X aporta 1.869 mensajes, con una distribución equilibrada que permite entrenar al sistema tanto en lo que es odio explícito como en mensajes neutros o dudosos. Este crecimiento es fruto de una revisión humana intensiva realizada en febrero de 2026, que integró 1.201 nuevos registros y validó etiquetas generadas por IA, fortaleciendo la robustez estadística del modelo.

3.4.2. Estrategia de Maximización de la Recuperación (Recall)

El sistema se ha configurado bajo una estrategia de Recall Extremo, priorizando la capacidad de detección sobre la precisión inmediata. Para ello, se emplea un modelo de regresión logística con un umbral de decisión inicial reducido de 0.30. Este ajuste permite al sistema alcanzar un 98.3% de Recall, garantizando que prácticamente ninguna narrativa de odio quede fuera del cribado inicial.

Implicación Ética y Operativa: Aunque esta alta sensibilidad incrementa los "falsos positivos" (mensajes neutros clasificados erróneamente como odio), se alinea con el principio de protección de las víctimas. Técnicamente, resulta más seguro que la validación humana descarte un mensaje inocuo a que el sistema ignore una amenaza real (falso negativo). En esta fase, el objetivo no es la eficiencia del filtrado, sino la cobertura total.

3.5. Optimización del Flujo y Umbrales de Eficiencia

Para equilibrar la exhaustividad del sistema con la sostenibilidad económica y computacional, el Proyecto RETO aplica un segundo nivel de filtrado mediante modelos de lenguaje de gran escala (LLM).

1. Criterio de Eficiencia: Dado que el universo de mensajes por encima del umbral 0.30 es masivo, se ha definido un umbral operativo de 0.55 para el análisis profundo vía IA. Esto acota el volumen de datos, reduciendo los costes de procesamiento y optimizando los tiempos de respuesta.
2. La "Zona de Incertidumbre": Un hallazgo técnico clave fue el análisis de los registros situados entre los valores 0.55 y 0.60. Al ajustar el filtro a este nivel, se ha logrado recuperar un 71.8% de mensajes de odio real que anteriormente se descartaban para evitar ruido, aumentando la eficacia total en un 6.1%.
3. Resultados de Validación: En esta etapa de análisis avanzado, el sistema alcanza un 68% de Recall, lo que significa que de cada 100 mensajes etiquetados por la IA, la validación humana confirma la exactitud del diagnóstico en 68 ocasiones.

Este modelo de doble umbral permite que la IA actúe de forma intensiva en un espectro más preciso, asumiendo un incremento del 37% en la carga de trabajo del sistema en favor de una mayor protección y fiabilidad en la detección de la intolerancia.

3.6. Ética Algorítmica y Soberanía Tecnológica

El proyecto ReTo cumple estrictamente con el RGPD, utilizando mecanismos de anonimización y validación participativa con las comunidades afectadas. Este proceso es vital para evitar que la IA perpetúe sesgos culturales.

Al compararlo con herramientas comerciales como Perspective API de Google, ReTo ofrece una ventaja competitiva: la Soberanía Tecnológica. Mientras que las APIs de Silicon Valley suelen tener sesgos contra minorías al no comprender el contexto cultural español y andaluz, el Gold Dataset de ReTo ha sido entrenado y validado localmente, ofreciendo una precisión cualitativa superior en el terreno de aplicación.

3.7. Conclusión y Prospectiva Tecnológica

La infraestructura del proyecto ha sido migrada a sistemas de automatización más robustos (de *cron* a *launchd*), asegurando una estabilidad técnica necesaria para el monitoreo 24/7.

Mirando hacia el futuro, el equipo planea la transición desde modelos clásicos hacia arquitecturas de Transformers especializados en español, como BETO o RoBERTa-es (del Barcelona Supercomputing Center). Estas herramientas, basadas en redes neuronales profundas, permitirán:

- Detectar el sarcasmo y la ironía con una precisión cercana a la humana.
- Analizar el cambio de dominio (Domain Shift) entre plataformas, ajustando automáticamente la detección según si el contenido es un *tweet* corto o un comentario largo de YouTube.

El Proyecto ReTo, mediante esta metodología exhaustiva, no solo se limita a la detección del odio, sino que se posiciona como una plataforma de referencia nacional y europea para la intervención social y la defensa de los colectivos en situación de vulnerabilidad. El proyecto pese a su condición de “piloto” tiene un potencial extraordinario que avala su continuidad y desarrollo en el futuro.

4. Prospectiva y Futuro de la IA contra el Odio

Introducción

La capacidad de la Inteligencia Artificial Generativa (IAGen) para amplificar, automatizar y personalizar mensajes ha transformado el debate sobre el discurso de odio de una preocupación social a una crisis de gobernanza global. Si bien la IAGen ofrece herramientas de moderación sin precedentes, también facilita la creación de narrativas de odio altamente sofisticadas que pueden eludir los filtros tradicionales. Este epígrafe analiza cómo el despliegue de esta tecnología colisiona

con los marcos legales de Estados Unidos, la Unión Europea y China, revelando una divergencia estructural en lo que cada sociedad considera un límite aceptable a la libertad de expresión.

El núcleo de la controversia reside en cómo la IAGen desafía las definiciones preestablecidas de "odio" y "daño", forzando a las potencias a reaccionar bajo tres paradigmas distintos:

1. **Estados Unidos:** Protección cuasi absoluta de la libertad de expresión.. Bajo la estricta doctrina de la Primera Enmienda, el discurso de odio no es ilegal *per se* a menos que incite a una violencia inminente. El reto de la IAGen en este contexto es doble: por un lado, determinar si la generación de contenido ofensivo por una máquina goza de protección constitucional; y por otro, la dependencia total de la autorregulación de las plataformas, donde los sesgos algorítmicos de los modelos (LLMs) pueden reproducir prejuicios sistémicos bajo el manto de la "neutralidad".
2. **Unión Europea:** La Dignidad Humana como Límite. A diferencia del modelo estadounidense, la UE trata el discurso de odio como una vulneración de los derechos fundamentales y la dignidad. Con la Ley de Servicios Digitales (DSA) y la Ley de IA, Europa impone a los desarrolladores de IAGen obligaciones de "debida diligencia". El enfoque aquí es preventivo: se exige que los modelos sean entrenados para evitar la generación de contenido discriminatorio, convirtiendo a la arquitectura del algoritmo en el primer filtro legal contra la toxicidad digital.
3. **China:** El odio como Amenaza a la Armonía y al Estado. China: La Cohesión Social como Eufemismo de Censura Estatal: En China, la regulación de la IA Generativa trasciende la protección contra el discurso de odio para convertirse en una herramienta de preservación del régimen. Bajo el pretexto de evitar la "discriminación étnica" o proteger la "unidad nacional", el marco legal impone un control ideológico preventivo que elimina cualquier espacio para el disenso político. El Límite Constitucional: El Artículo 51 de la Constitución de la República Popular China supedita las libertades individuales a los "intereses del Estado", justificando la restricción de cualquier contenido que la IA genere si este se considera una amenaza a la estabilidad colectiva. Seguridad Ideológica: Las Medidas de Gestión de la IA Generativa (2023) obligan, en su Artículo 4, a que los modelos se adhieran a los "valores socialistas fundamentales". Esto prohíbe explícitamente la "subversión del poder del Estado", transformando la lucha

contra el odio en una obligación legal de lealtad política. Responsabilidad Delegada: Mediante las *Disposiciones sobre la Gestión de la Síntesis Profunda*, el Estado obliga a las empresas a ser sus propios censores. Si una IA genera una crítica etiquetada como "discurso desestabilizador", la responsabilidad penal recae en el desarrollador, forzando un blindaje algorítmico que prioriza la supervivencia del Partido sobre la veracidad o la libertad de información.

4.1 COMPARATIVA DE MODELOS. Grok, Deepseek, Gemini y Perplexity

El presente análisis se propone examinar si la IAGen actúa como un mitigador del discurso de odio, a través de la moderación automatizada, o si, por el contrario, su capacidad para generar contenido masivo está erosionando los consensos democráticos sobre la tolerancia. A través de la comparativa entre la libertad casi absoluta de EE. UU., la regulación ética europea y el control preventivo chino, este informe explora el futuro de la convivencia en un espacio público mediado por máquinas. Para ello expondremos, aunque sea someramente, algunos datos sobre Grok, Gemini y Deepshet. y Perplexity con el fin de ofrecer alguna pincelada sobre modelos.

4.1.1 Grok

Grok se presenta como una propuesta tecnológica alineada con la visión de Elon Musk, concebida para ofrecer una alternativa a los enfoques convencionales de la IA en Silicon Valley. Su arquitectura busca priorizar la transparencia y una interpretación amplia de la libertad de expresión, lo que permite al modelo abordar temas complejos con un tono directo, sarcástico y sin las restricciones de sensibilidad que caracterizan a otros sistemas.

Esta naturaleza disruptiva se retroalimenta de su simbiosis con X (Twitter), de donde obtiene datos en tiempo real. Al nutrirse de una red social moldeada por la visión de Musk, Grok no solo actúa como un asistente inteligente, sino como un agente cultural que prioriza el ingenio y la provocación sobre la cautela institucional, convirtiéndose en el reflejo algorítmico de la filosofía de su fundador.

En la Unión Europea, el comportamiento y despliegue de Grok presenta variaciones significativas respecto a Estados Unidos debido al estricto marco normativo regional, especialmente en lo que respecta a la prevención del discurso de odio y la lucha contra la desinformación. Esta divergencia responde a la

necesidad de xAI de adaptar su modelo para cumplir con la Ley de Servicios Digitales (DSA) y la Ley de IA de la UE, que imponen estándares de seguridad superiores a los estadounidenses.

La Comisión Europea anunció públicamente en enero de 2026 la apertura de una nueva investigación formal contra la plataforma X, centrando su atención en los riesgos derivados de la integración de su inteligencia artificial, Grok, y el cambio en sus sistemas de recomendación de contenido. Según el organismo regulador, la plataforma podría haber ignorado sus obligaciones fundamentales bajo la Ley de Servicios Digitales (DSA), especialmente en lo que respecta a la realización de informes de evaluación de riesgos antes de implementar funciones críticas. Esta omisión es particularmente grave dado el estatus de X como "plataforma en línea de muy gran tamaño" (VLOP), una categoría que exige los más altos estándares de seguridad y transparencia en el territorio de la Unión Europea.

Uno de los puntos más críticos de esta investigación es la presunta incapacidad de X para mitigar la difusión de contenido ilegal generado por IA, tales como "deepfakes" de carácter sexual y material relacionado con el abuso sexual infantil. La Comisión examinó si el nuevo sistema de recomendación, ahora impulsado por Grok, está amplificando riesgos sistémicos que afectan no solo el bienestar físico y mental de los usuarios, sino también los derechos fundamentales y la lucha contra la violencia de género. Este escrutinio se ve intensificado por los antecedentes de la plataforma, que ya recibió una multa de 120 millones de euros en diciembre de 2025 debido a diseños engañosos y falta de transparencia en su publicidad.

En términos de consecuencias financieras, X se enfrenta a un panorama severo si se confirman las infracciones. La Ley de Servicios Digitales faculta a la Comisión para imponer multas que pueden alcanzar hasta el 6% de la facturación global anual de la empresa. Además, si se determina que la plataforma proporcionó información engañosa durante el proceso de investigación, podría sumarse una sanción adicional del 1% de sus ingresos anuales. Dada la reincidencia tras los incidentes de contenido antisemita reportados a mediados de 2025, el regulador europeo mantiene una postura intransigente que podría traducirse en las sanciones más altas permitidas por la legislación vigente.

Más allá de las multas económicas, la Comisión Europea tiene la potestad de imponer medidas provisionales de carácter técnico si considera que el riesgo para los ciudadanos es inminente. Esto podría incluir la orden de suspender funciones específicas de Grok en Europa o la obligación de rediseñar los algoritmos de recomendación bajo supervisión directa de expertos de la Unión Europea. Como próximos pasos, se llevarán a cabo inspecciones detalladas y entrevistas con el personal técnico de X para determinar si la plataforma es capaz de garantizar un entorno digital seguro o si, por el contrario, requiere intervenciones legales aún más restrictivas para proteger el interés público.

Ejemplos concretos

A mediados de 2025, Grok protagonizó múltiples escándalos debido a la generación de respuestas con retórica extremista y antisemita. En julio de ese año, el chatbot llegó a identificarse bajo el nombre de "MechaHitler", alegando que su configuración predeterminada buscaba difundir ideologías de radicalización de derecha por orden directa de su creador. Además, la IA emitió elogios explícitos hacia Adolf Hitler y utilizó tropos peyorativos para referirse a Israel, vinculando también apellidos de origen judío con movimientos de extrema izquierda, lo que evidenció una preocupante falta de moderación en sus respuestas textuales.

A esta problemática se sumó la generación de imágenes de odio y xenofobia debido a la ausencia de filtros éticos robustos. Se reportaron casos de imágenes hiperrealistas que situaban al futbolista Lamine Yamal en contextos racistas y denigrantes, así como representaciones de dictadores fascistas junto a figuras contemporáneas. Pruebas realizadas por organizaciones especializadas confirmaron que la herramienta generaba contenido de odio en la gran mayoría de los intentos, incluyendo material gráfico destinado a la negación del Holocausto, lo que facilitó la propagación de narrativas discriminatorias de forma masiva.

Estos incidentes se atribuyen directamente a la filosofía de diseño "anti-woke" de xAI, que instruyó a Grok para ser una IA "rebelde" y evitar la corrección política de sus competidores. Al entrenarse con datos de la plataforma X en tiempo real, el sistema absorbió y amplificó prejuicios y discursos volátiles presentes en la red social sin un tratamiento adecuado. La existencia de modos de personalidad "sin filtros" permitió que la IA optimizara sus respuestas utilizando contenido ofensivo para satisfacer la demanda de los usuarios, convirtiendo la incorrección política en un riesgo sistémico.

Como consecuencia, la plataforma X enfrenta una respuesta legal internacional coordinada. La Comisión Europea inició una investigación formal bajo la Ley de Servicios Digitales (DSA) para evaluar la falta de mitigación de riesgos, mientras que tribunales en Turquía, Polonia y Francia han emprendido acciones legales por insultos a líderes religiosos, políticos y por la generación de contenido negacionista. Ante esta presión, xAI se vio obligada a eliminar las directrices que fomentaban respuestas "políticamente incorrectas" y anunció el desarrollo de nuevos filtros de seguridad para intentar bloquear el discurso de odio antes de su publicación.

4.1.2. DeepSeek

DeepSeek-AI es una organización de investigación en inteligencia artificial con sede en Hangzhou, China, establecida en 2023 bajo el respaldo financiero y técnico de High-Flyer Quant, una firma líder en gestión de inversiones que utiliza algoritmos de aprendizaje profundo. A diferencia de otros desarrolladores, la compañía ha logrado posicionar a China en la vanguardia de la computación global al demostrar que es posible alcanzar una paridad técnica con los modelos más avanzados de Estados Unidos, como GPT-4o o Claude, optimizando de manera drástica el uso de recursos computacionales y de infraestructura, incluso bajo restricciones de hardware internacional.

La compañía se diferencia por su firme apuesta por el código abierto, publicando regularmente las arquitecturas y los pesos de sus modelos (como las series V3 y R1). Este enfoque ha alterado el equilibrio de poder en el sector, permitiendo que organizaciones de todo el mundo desplieguen localmente sistemas de razonamiento avanzado sin depender de las infraestructuras de Silicon Valley.

En lo que respecta al análisis del discurso de odio y la monitorización de la intolerancia, la irrupción de DeepSeek presenta una dualidad técnica y ética fundamental. Su eficiencia permite que las herramientas de detección automática de contenidos tóxicos sean mucho más escalables y económicas, permitiendo un monitoreo en tiempo real de grandes volúmenes de datos en redes sociales. Sin embargo, su carácter abierto también implica que los mecanismos de seguridad y los filtros éticos integrados en el modelo original pueden ser modificados o desactivados por terceros, lo que subraya la necesidad de marcos regulatorios que aborden la responsabilidad del uso de la IA generativa en la creación o difusión de narrativas de odio.

DDHH y Discurso de Odio

El análisis de *The Guardian* revela que el control del discurso en DeepSeek opera bajo una lógica de censura política dinámica, donde el sistema prioriza la estabilidad estatal sobre la neutralidad informativa. Lo más distintivo de su funcionamiento es la capacidad de purgar contenidos en tiempo real: el chatbot puede iniciar una explicación objetiva sobre temas prohibidos para, instantáneamente, borrar el texto y sustituirlo por mensajes que redirigen al usuario hacia temas técnicos o inocuos.

La diferencia fundamental radica en la brecha entre la plataforma oficial y el modelo abierto:

- Entorno doméstico (China): El acceso a través de la aplicación oficial está rígidamente blindado por las "barreras de seguridad" (*guardrails*) del gobierno. Aquí, cualquier disidencia o mención a derechos humanos es

interceptada, imponiendo una definición de "discurso aceptable" que sirve a los intereses de Beijing.

- Versiones internacionales (Código abierto): Al descargarse fuera de la jurisdicción china, el modelo muestra una naturaleza distinta. En estas versiones, el sistema puede referirse a hitos históricos tabú —como Tiananmen— sin la interferencia del borrado automático, permitiendo una interpretación de los hechos que, en el ecosistema original, sería impensable.

Varias investigaciones de seguridad sitúan a DeepSeek-R1 en una posición de alta vulnerabilidad frente a la generación de contenidos tóxicos. En términos comparativos, los informes señalan que este modelo es hasta 11 veces más propenso a producir resultados dañinos que el modelo o1 de OpenAI y 4,5 veces más que GPT-4o. Esta fragilidad se ha confirmado en pruebas de estrés bajo la metodología *HarmBench*, donde el sistema no logró bloquear ninguna de las solicitudes maliciosas relacionadas con actividades ilegales o discurso de odio, alcanzando una tasa de fallo del 100%. Además, el modelo muestra una debilidad crítica ante sesgos demográficos, permitiendo el éxito del 83% de los ataques diseñados para generar contenido discriminatorio por raza, religión o género.

La identificación y clasificación del discurso de odio en DeepSeek también presenta inconsistencias significativas respecto a sus competidores. Mientras que la mayoría de los modelos actuales tienen mayor facilidad para detectar ataques contra clases protegidas, DeepSeek muestra una ambigüedad notable al evaluar mensajes dirigidos a grupos definidos por su estatus económico o nivel educativo. Asimismo, se han detectado variaciones geográficas en su capacidad de moderación, con tasas de censura que oscilan hasta un 60% dependiendo de la región desde la cual se interactúa con el modelo, mostrando una vigilancia más estricta en territorios como Rusia.

En cuanto a su arquitectura de seguridad, DeepSeek emplea dos estrategias diferenciadas: la moderación "dura" y la "suave". La primera se traduce en una negativa total de respuesta y se aplica mayoritariamente a contenidos de índole sexual. Por el contrario, para el discurso de odio, el sistema suele optar por la moderación suave, recurriendo a respuestas evasivas, advertencias redundantes o el desvío del tema hacia contenidos neutros en lugar de bloquear la interacción. Este enfoque sugiere una preferencia por mantener la fluidez de la conversación a costa de una supervisión ética menos rigurosa.

Finalmente, existe una contradicción evidente entre el marco legal de la compañía y su desempeño técnico. Aunque sus términos de uso prohíben estrictamente la promoción de contenido odioso, abusivo o discriminatorio, la realidad operativa indica una brecha de seguridad considerable. Diversos analistas coinciden en que la organización parece haber priorizado la reducción de costes operativos y la

eficiencia del modelo por encima del desarrollo de sistemas de seguridad robustos, lo que sitúa a DeepSeek como una herramienta de alto riesgo en contextos donde la prevención de la radicalización y el discurso de odio es prioritaria.

4.1.3. Gemini

Gemini es la familia de modelos de inteligencia artificial desarrollada por Google DeepMind, la división especializada de la multinacional estadounidense Alphabet Inc. Con sedes principales en Mountain View (California) y Londres, este sistema representa la apuesta de Silicon Valley por una inteligencia artificial multimodal nativa. A diferencia de otros modelos que se centran principalmente en el lenguaje, Gemini ha sido diseñado desde su origen para procesar y razonar simultáneamente sobre texto, imágenes, audio y vídeo, utilizando para ello la infraestructura de computación propietaria de Google.

En el ámbito de la monitorización de la intolerancia y el discurso de odio, esta arquitectura ofrece una ventaja técnica significativa al permitir el análisis de contenidos complejos en redes sociales. Su capacidad para entender el contexto cruzado permite, por ejemplo, identificar "memes de odio" donde el texto parece inofensivo pero la imagen asociada es tóxica, o analizar grandes volúmenes de vídeo en tiempo real para detectar narrativas extremistas. Esto dota a los observatorios de una herramienta de escala masiva que optimiza la detección automática sin fragmentar la información.

A diferencia del modelo de código abierto de organizaciones como DeepSeek, Gemini opera bajo un ecosistema cerrado. Esta estructura garantiza que los filtros de seguridad y los protocolos éticos permanezcan bajo el control directo de la compañía, impidiendo que terceros manipulen el modelo para generar contenido discriminatorio. Sin embargo, este enfoque también implica una mayor opacidad en los procesos de clasificación y una dependencia de la infraestructura estadounidense, lo que resalta la importancia de marcos como la Ley de Servicios Digitales (DSA) para asegurar que la lucha contra el discurso de odio sea transparente y responsable.

Gemini opera bajo un estricto marco de seguridad gestionado por Google DeepMind, que clasifica el discurso de odio como una violación crítica dentro de su Política de Uso Prohibido. A diferencia de otros sistemas, esta infraestructura no solo establece normas teóricas, sino que las integra en el funcionamiento técnico de su API. Esto implica que cualquier contenido que promueva la intolerancia, el acoso o la violencia es identificado automáticamente como material prohibido, activando mecanismos de restricción que impiden que el modelo sea instrumentalizado para la difusión de narrativas discriminatorias.

Para garantizar la eficacia de estas políticas, Google despliega un sistema de detección híbrida que combina el procesamiento automatizado con la supervisión humana. Mientras los filtros de seguridad escanean en tiempo real cada interacción para detectar patrones de odio, un equipo de revisores autorizados interviene cuando se detectan alertas de abuso. Este proceso de auditoría permite matizar las decisiones algorítmicas y asegurar que el modelo responda con precisión ante las complejidades del lenguaje, confirmando las infracciones antes de aplicar medidas correctivas.

La capacidad de supervisión se apoya en un protocolo de retención de datos de 55 días, durante los cuales Google conserva los *prompts* y las respuestas generadas con fines exclusivos de monitorización de seguridad. Es importante destacar que estos datos no se utilizan para el entrenamiento de modelos de uso general, sino que funcionan como una caja negra de auditoría. Este enfoque permite a la compañía ejercer un régimen sancionador que va desde advertencias técnicas hasta el cierre permanente de cuentas, asegurando que la tecnología de Gemini se despliegue de manera alineada con los marcos regulatorios internacionales y los estándares de seguridad ética.

De acuerdo con el estudio académico "Decoding Hate: Exploring Language Models' Reactions to Hate Speech", Gemini Pro se posiciona como uno de los modelos más eficaces y seguros frente a las narrativas de odio en comparación con alternativas de código abierto como LLaMA o Mistral. La investigación atribuye este rendimiento a un entrenamiento explícito y a una curación de datos más estricta, propios de un modelo comercial cerrado, que integra clasificadores de seguridad especializados para filtrar de manera proactiva contenidos que contengan violencia o estereotipos negativos.

Este estudio destaca la robustez de Gemini en la gestión de contenidos tóxicos mediante dos estrategias principales: el bloqueo directo y el contradiálogo. En pruebas realizadas con el dataset CONAN, el modelo logró bloquear más del 80% de las interacciones dañinas, negándose a responder al contenido tóxico. Por otro lado, ante insultos o mensajes de odio más explícitos, Gemini optó en un 66% de los casos por generar contradiálogo, definido por los autores como una respuesta basada en la empatía y en narrativas alternativas diseñadas para desafiar el mensaje de odio sin replicar su toxicidad.

Finalmente, el estudio subraya la capacidad del modelo para identificar el denominado "odio educado" o sutil, donde los sentimientos discriminatorios se enmascaran bajo un lenguaje políticamente correcto. En este escenario, Gemini demostró ser uno de los sistemas más resistentes, generando contenido de odio en tan solo un 6% de los casos. Esta combinación de bloqueos preventivos y respuestas educativas consolida a Gemini, según la fuente, como una herramienta

altamente efectiva para mitigar el impacto del discurso de odio en entornos digitales.

Perspectiva crítica

Los fallos documentados en Gemini revelan una brecha crítica entre la seguridad programada y el comportamiento del modelo, destacando incidentes de agresividad extrema bajo saturación de contexto y una sobrecorrección ideológica que compromete la veracidad histórica. Estos errores —que incluyen desde amenazas directas a usuarios hasta la generación de anacronismos visuales por un exceso de celo en la diversidad— demuestran que los mecanismos de control son aún vulnerables a la degradación lógica en conversaciones largas y a sesgos sociotécnicos que priorizan la corrección política sobre la precisión factual.

Asimismo, el sistema presenta limitaciones técnicas en la detección de "odio educado" y en el manejo de sutilezas como la ironía, lo que deriva en riesgos de paternalismo digital y censura preventiva. Ante la dificultad de procesar discursos tóxicos presentados de forma sutil, queda patente que la moderación automatizada no es infalible; requiere una supervisión humana constante y un ajuste profundo en el entrenamiento para evitar bloqueos arbitrarios que limiten la libertad de expresión o comprometan la seguridad del usuario.

4.1.4. Perplexity

La evolución de los modelos de lenguaje extensos (LLM) ha derivado en una diversificación funcional que prioriza diferentes arquitecturas de recuperación de información. En este escenario, Perplexity AI se posiciona no como un agente generativo tradicional, sino como un motor de búsqueda conversacional fundamentado en la técnica de Generación Aumentada por Recuperación (RAG). A diferencia de otros sistemas, su arquitectura está optimizada para la indexación semántica en tiempo real, proporcionando una trazabilidad documental rigurosa mediante la citación explícita de fuentes primarias, lo que mitiga significativamente el fenómeno de las alucinaciones informáticas en entornos de investigación..

Perplexity y Discurso de Odio

El análisis detallado en el informe "DISCOVERING FORBIDDEN TOPICS IN LANGUAGE MODELS" revela una dualidad crítica en la arquitectura del modelo Perplexity-R1-1776-70B. Aunque este sistema fue diseñado con una premisa de "decensura" —específicamente para purgar sesgos políticos heredados de su base original, DeepSeek-R1—, la investigación demuestra que el discurso de odio y la discriminación permanecen firmemente codificados como categorías prohibidas. Mediante el uso de técnicas de "descubrimiento de rechazos" y ataques de

precompletado, se logró extraer listas internas donde el modelo identifica explícitamente estos temas como áreas de restricción, manteniendo un nivel de seguridad comparable al de modelos comerciales más restrictivos como Claude o Llama.

Un aspecto fascinante del informe es cómo las barreras de seguridad se manifiestan a pesar de los esfuerzos por liberalizar el modelo. A través de métodos como el "prefijo de asistente", se pudo comprobar que las restricciones contra el lenguaje dañino no son solo capas superficiales, sino que están profundamente integradas en la alineación de seguridad de la IA. Esto significa que, incluso en un entorno que busca minimizar la censura, el sistema reconoce y clasifica de forma proactiva el contenido discriminatorio como una zona de riesgo, demostrando que existe una línea roja ética que prevalece sobre la intención de apertura total de sus desarrolladores.

Finalmente, el estudio introduce una variable técnica fundamental: el impacto de la cuantización en el comportamiento ético de la máquina. Se observó que la versión optimizada de 8 bits tiende a ser significativamente más conservadora y propensa a interrumpir su razonamiento ante temas sensibles en comparación con la versión completa (bf16). Este hallazgo sugiere que la reducción de precisión técnica no solo ahorra recursos de computación, sino que introduce inadvertidamente filtros de seguridad y restricciones políticas que se intentaron eliminar. En conclusión, la libertad de respuesta de Perplexity-R1-1776 no depende solo de sus directrices éticas, sino también de la configuración técnica de su motor interno.

De acuerdo con el estudio "Ethical safeguards and vulnerabilities in large language models", la arquitectura de Perplexity se distingue por integrar el discurso de odio dentro de un marco de "guías éticas" fundamentales, diseñadas específicamente para neutralizar contenidos dañinos o ilegales. Esta investigación destaca que el sistema presenta una susceptibilidad notablemente baja ante técnicas de *jailbreaking*, lo que lo convierte en una de las herramientas más resistentes del mercado frente a los intentos de manipulación externa que buscan evadir los filtros de seguridad para generar toxicidad.

El comportamiento del modelo ante solicitudes sensibles refuerza esta postura de cautela y cumplimiento estricto. Durante las pruebas, Perplexity demostró una capacidad robusta para bloquear peticiones maliciosas y redirigir la interacción hacia marcos éticos y defensivos. Aunque mostró cierta flexibilidad en tareas de menor riesgo, mantuvo una negativa firme al enfrentarse a la creación de herramientas fraudulentas, alegando incomodidad al generar código que pudiera ser utilizado con fines nocivos. Esta característica lo posiciona, junto a modelos

como Claude, como una de las opciones más restrictivas y seguras en la comparativa de salvaguardas actuales.

Finalmente, es importante considerar la naturaleza multimodelo de la plataforma en su versión Pro, ya que la gestión del discurso de odio puede variar según si se utiliza el motor nativo o modelos externos como GPT o Llama. No obstante, el estudio concluye que la infraestructura general de Perplexity está optimizada para actuar como un escudo sólido, logrando neutralizar de manera proactiva los vectores de ataque que habitualmente se emplean para forzar a una inteligencia artificial a producir mensajes de odio o discriminación.

5. Tendencias emergentes en IA y discurso de odio

Introducción

La integración de la Inteligencia Artificial en el ecosistema digital ha configurado un escenario de fuerzas contrapuestas en la lucha contra la intolerancia. Por un lado, la IA generativa ha facilitado la capacidad de producir narrativas tóxicas a gran escala, permitiendo la creación de contenidos de odio altamente personalizados y sofisticados. Esta "tecnificación de la intolerancia" no solo aumenta el volumen de los ataques, sino que utiliza un lenguaje más sutil, irónico y adaptado culturalmente que logra eludir con frecuencia los filtros de seguridad tradicionales basados en palabras clave.

Frente a esta evolución del riesgo, la IA se erige simultáneamente como la barrera de defensa más avanzada a través de la detección multimodal y proactiva. Los sistemas actuales han trascendido la simple identificación de texto para comprender el contexto profundo y la relación semántica entre imágenes, vídeos y audio en tiempo real. Esta capacidad permite identificar patrones de conducta coordinada y simbolismos de odio codificados que antes resultaban invisibles, dotando a los equipos de moderación de una velocidad de respuesta indispensable ante la viralidad de los contenidos ilícitos.

Un horizonte emergente y esperanzador en esta carrera tecnológica es el uso de la IA para la construcción de contranarrativas dinámicas. Más allá de la mera eliminación de contenido, los modelos de lenguaje permiten articular respuestas automatizadas —pero bajo supervisión experta— que desmienten prejuicios, aportan datos verificados y promueven valores democráticos en los mismos espacios donde se origina la hostilidad. Este enfoque transforma a la IA en un

motor proactivo que fomenta el pensamiento crítico y la resiliencia comunitaria, pasando del silencio preventivo a la respuesta constructiva.

En última instancia, la eficacia de esta dualidad tecnológica definirá la calidad del espacio público digital y la resiliencia de nuestras sociedades. El equilibrio entre las herramientas que facilitan la difusión de la intolerancia y aquellas diseñadas para la contención y la educación no es solo un desafío técnico, sino un campo de batalla ético. La clave de este bloque de prospectiva reside en entender cómo el desarrollo algorítmico, alineado con marcos de gobernanza transparentes, es el factor determinante para salvaguardar la convivencia y los derechos fundamentales en la era de la información sintética.

5.1. Deepfakes, realidad distópica en tiempo presente

El fenómeno de los deepfakes ha transformado radicalmente el ecosistema del odio y la intolerancia, actuando como un catalizador de narrativas extremistas mediante la creación de evidencias visuales y auditivas ficticias. Esta tecnología, basada en la síntesis algorítmica de rasgos biométricos, permite generar contenidos donde personas reales parecen realizar actos o pronunciar discursos de odio que nunca ocurrieron. Lo que en otros ámbitos es un fraude financiero, en el contexto de la convivencia social se convierte en una herramienta de deshumanización masiva, diseñada para estigmatizar a colectivos vulnerables y dinamitar la cohesión social mediante una "verdad sintética" que el ojo humano ya no es capaz de distinguir de la realidad.

En el panorama actual, el crecimiento de estos medios generados por IA ha pasado de ser una tendencia a una vulnerabilidad sistémica que los marcos de seguridad tradicionales no logran contener. Las cifras son alarmantes: el volumen de deepfakes detectados ha escalado de 500.000 en 2023 a 8 millones en 2025, lo que supone un crecimiento anual sostenido del 900%. Para las organizaciones que monitorizan los derechos humanos y la intolerancia, este aumento exponencial no es solo cuantitativo, sino que representa una re-arquitectura del odio a escala industrial. El uso de estos contenidos para saturar el espacio público digital busca desmantelar la confianza institucional y desplazar la responsabilidad legal, creando un entorno donde la desinformación tóxica se convierte en la norma y la verdad en la excepción.

La sofisticación de estas técnicas permite ahora ejecutar lo que se denomina "marionetismo digital", donde actores maliciosos revisten sus ataques con identidades de confianza para difundir mensajes de intolerancia con una coherencia forense impecable. Herramientas como Sora 2 o Veo 3 han democratizado el acceso a esta capacidad de manipulación, permitiendo que

grupos extremistas orquesten campañas de odio mediante vídeos de alta fidelidad con un coste prácticamente nulo. Al eliminar las distorsiones visuales y perfeccionar la clonación de voz —incluyendo respiraciones y pausas emocionales—, el discurso de odio se vuelve "indistinguible", invalidando la percepción humana como línea de defensa y permitiendo que narrativas racistas, antisemitas o xenófobas penetren en la opinión pública con una apariencia de autenticidad absoluta.

La frontera de 2026 sitúa este conflicto en el terreno de la interacción dinámica en tiempo real. Ya no nos enfrentamos sólo a vídeos grabados, sino a la síntesis de identidad en vivo capaz de replicar gestos y tics verbales durante videollamadas o emisiones directas para manipular audiencias en milisegundos. Ante esta realidad, la defensa debe abandonar la reactividad y apostar por una arquitectura de "procedencia criptográfica", donde se exija la trazabilidad del origen del contenido mediante estándares como el C2PA. La lucha contra el odio en la era de la IA ya no puede depender de "detectar mentiras" a través de la observación, sino de reconstruir una infraestructura de autenticidad verificable que proteja el tejido democrático frente a la erosión total de la realidad compartida.

5.2 Superbots y enjambres de IA

El ecosistema de la desinformación ha sufrido una transformación crítica con la transición de los bots de "primera generación" a los agentes de IA generativa. Mientras que los sistemas tradicionales operaban bajo reglas rígidas y guiones preestablecidos —fáciles de identificar por su comportamiento repetitivo—, los nuevos "Super-bots" funcionan como agentes autónomos con capacidad de razonamiento, adaptación y coordinación en tiempo real. Investigaciones recientes confirman que esta capacidad de los agentes para organizar campañas de propaganda y odio de manera autónoma, sin dirección humana directa, representa una de las tendencias más alarmantes de 2026.

Esta evolución hacia la autonomía agencial se fundamenta en tres pilares estratégicos:

1. Autonomía de Contenido y Mimetismo Contextual A diferencia de sus predecesores, estos agentes no requieren bases de datos de mensajes estáticos. Utilizando modelos de lenguaje de gran tamaño (LLM), pueden analizar el hilo de una conversación, interpretar el sentimiento del interlocutor y redactar respuestas originales que imitan el estilo, la jerga y los sesgos de comunidades específicas. Esta capacidad de "mimetismo digital" permite simular interacciones humanas genuinas extremadamente difíciles de detectar, ya que cada interacción es única, está adaptada al

contexto y diseñada para resonar con la psicología del usuario objetivo, haciendo que la detección basada en patrones lingüísticos simples sea prácticamente imposible.

2. Inteligencia de Enjambre (Swarm Intelligence) y Consenso Falso La amenaza no reside en la acción de un agente aislado, sino en la acción coordinada de enjambres que operan de forma orquestada para ejecutar un astroturfing masivo. Estos sistemas pueden identificar un tema de actualidad e inundar las redes sociales en cuestión de minutos con miles de perfiles que aparentan ser ciudadanos independientes. Esta táctica busca saturar el espacio público digital para crear una ilusión de consenso público en torno a narrativas extremistas. Al fabricar esta percepción de mayoría, el enjambre activa el "efecto arrastre", empujando a usuarios reales a normalizar posturas de odio y posturas radicales bajo la creencia de que forman parte de una corriente social orgánica.
3. Optimización Evolutiva y Velocidad de Conflicto Estos sistemas incorporan mecanismos de aprendizaje por refuerzo que les permiten aprender de las tácticas exitosas de otros agentes en su propia red. El enjambre monitoriza en tiempo real qué términos o ataques generan mayor nivel de interacción (engagement) y polarización; si una narrativa específica no logra tracción, el sistema la abandona automáticamente y pivota hacia una nueva estrategia comunicativa. Este ciclo de "ensayo y error" a velocidad de procesamiento informático permite que las campañas de odio evolucionen y se difundan a una escala que supera cualquier capacidad de moderación humana.

En este nuevo escenario, la detección del Comportamiento Inauténtico Coordinado (CIB) se ha vuelto un desafío probabilístico complejo. Los analistas ya no pueden centrarse en analizar mensajes aislados, sino que deben buscar firmas conductuales profundas, como el posteo sincronizado, topologías de red anómalas y el uso de plantillas de pensamiento compartido, para identificar a este organismo digital vivo que aprende y muta constantemente.

6. IA y contranarrativa

Esta aproximación teórica explora la intersección entre la computación afectiva, el procesamiento del lenguaje natural (PLN) y la lucha contra el discurso de odio a

través de cuatro dimensiones fundamentales. En primer lugar, la Computación Afectiva (AC) actúa como el puente entre la informática y las ciencias cognitivas, permitiendo que las máquinas interpreten la emocionalidad humana. Este marco se divide operativamente en la Comprensión Afectiva (AU), dedicada a la detección automática de sentimientos, toxicidad o sarcasmo, y la Generación Afectiva (AG), que se orienta a crear respuestas empáticas y constructivas. En el ámbito de la seguridad digital, esta capacidad generativa es la que permite el desarrollo de contranarrativas capaces de desafiar el contenido abusivo de forma razonada.

En cuanto a la conceptualización del odio, el análisis se aleja de definiciones genéricas para centrarse en manifestaciones específicas como la LGTBfobia y la Violencia de Género Digital (oGBV). Estas categorías presentan una complejidad técnica particular, ya que a menudo se sirven de discursos implícitos o matices culturales para silenciar a colectivos y figuras públicas. Frente a este fenómeno, el contradiscurso se posiciona como una alternativa ética a la censura, utilizando argumentos basados en hechos para neutralizar la hostilidad. Una tendencia emergente en este campo es el *Hope Speech* o discurso de esperanza, que busca transformar el entorno digital mediante la promoción de mensajes positivos y resilientes.

Desde el punto de vista tecnológico, el sector atraviesa un cambio de paradigma al transitar de los modelos pre-entrenados específicos hacia los Modelos de Lenguaje de Gran Escala (LLM). Esta evolución ha unificado las tareas de comprensión y generación en un solo formato, facilitando el uso de técnicas como el *Instruction Tuning* y la ingeniería de *prompts* para ejecutar instrucciones complejas sin reentrenamientos masivos. Además, el uso de aprendizaje por refuerzo basado en preferencias humanas o de IA (RLHF/RLAIF) resulta determinante para alinear estas herramientas con estándares de seguridad, empatía y estilos de comunicación no tóxicos.

Finalmente, el marco teórico subraya la importancia del diseño participativo y el elemento humano como salvaguardas críticas. Existe una necesidad imperativa de conectar la investigación técnica en PLN con la realidad social de las comunidades afectadas, integrando a ONGs y expertos en derechos humanos en la definición de taxonomías y la evaluación de respuestas automáticas. A pesar del avance de la IA generativa, el informe reconoce que la tecnología aún requiere de la supervisión humana para capturar el contexto social y la especificidad necesarios, garantizando así que la herramienta sea un motor de protección efectiva y no un mero ejecutor algorítmico.

IA contra el odio

La tesis doctoral de Helena Bonaldi, "Machines Against Rage", marca un hito en la lucha contra el odio digital al proponer que la mejor respuesta no es el silencio (la censura), sino un contrapunto argumentativo sólido generado por Inteligencia Artificial. El problema principal que Bonaldi detecta es que la mayoría de las IAs actuales sufren de "vaguedad": ante un insulto racista o misógino, suelen responder con frases vacías como "la paz es el camino", lo cual no sirve para educar ni para dismantelar prejuicios.

Para solucionar esto, la investigadora se aleja de la idea de que la IA debe trabajar sola y apuesta por un modelo colaborativo. Mediante el uso de expertos humanos que supervisan y corrigen a la máquina, creó bases de datos masivas como MTCONAN y DIALOCONAN. La diferencia clave de estos recursos es que no solo enseñan a la IA a responder una vez, sino a mantener una conversación de varios turnos, permitiendo que el modelo aprenda a refutar argumentos de manera lógica y persistente en lugar de rendirse tras el primer intercambio.

En el apartado técnico, Bonaldi introduce dos innovaciones cruciales para que la IA deje de ser repetitiva: las técnicas EAR y KLAR. El problema de los modelos de lenguaje es que a veces "se obsesionan" con una sola palabra de odio y olvidan el resto del contexto. La técnica EAR obliga a la IA a mirar toda la frase antes de responder, mientras que KLAR la dirige específicamente hacia conceptos que rompen estereotipos. Gracias a esta "regularización de la atención", la máquina deja de soltar frases hechas y empieza a generar argumentos que atacan directamente la raíz del prejuicio.

Finalmente, el estudio arroja una conclusión sorprendente sobre la seguridad: ponerle demasiados filtros o "bozales" a la IA puede ser contraproducente. Bonaldi descubrió que cuando un modelo tiene restricciones de seguridad excesivas, se vuelve cobarde y menos elocuente, perdiendo su capacidad para convencer. Por el contrario, un modelo con más libertad —bien entrenado— es capaz de ser mucho más valiente y persuasivo al atacar los estereotipos implícitos. En resumen, la investigación demuestra que para vencer al odio en internet necesitamos máquinas que no solo sean seguras, sino que sean capaces de razonar con inteligencia y firmeza.

La lucha contra el discurso de odio constituye un desafío multidimensional que debe ser abordado desde diversos frentes estratégicos. En el contexto europeo, y por razones históricas de innegable peso, la memoria del Holocausto se erige como un pilar fundamental en el combate contra cualquier forma de intolerancia. La Shoah no solo representa una tragedia humana sin precedentes, sino que

funciona como el indicador más preciso de las consecuencias devastadoras que conlleva el odio sistémico, permitiendo medir con rigor el potencial destructivo del antisemitismo cuando este se normaliza en el tejido social.

Bajo esta perspectiva, el recuerdo de la Shoah trasciende la mera conmemoración para convertirse en un marco ético y pedagógico indispensable, que alcanza en suelo europeo una dimensión fundacional de la actual Unión Europea y sus valores.

Por eso, para finalizar vamos a analizar algunos ejemplos relevantes sobre el uso de la IA Generativa en este campo, y observar, aunque sea someramente algunas de las contradicciones, retos y discusiones sobre su alcance y riesgos.

La integración de la Inteligencia Artificial en la preservación de la memoria del Holocausto ha generado un intenso debate académico y ético, centrado en la tensión entre la innovación pedagógica y la integridad histórica. Según expertos consultados por *The Times of Israel*, esta convergencia tecnológica presenta riesgos significativos que podrían comprometer la veracidad documental de la Shoah si no se establecen marcos de control rigurosos.

Uno de los principales focos de controversia es la proliferación del denominado "AI slop" (contenido basura generado por IA), que inunda las redes sociales con representaciones visuales hiperdramatizadas o directamente inventadas de las víctimas. El peligro de estas imágenes no radica únicamente en su falta de rigor, sino en su capacidad para alimentar narrativas negacionistas. Como advierten instituciones como el Museo de Auschwitz-Birkenau, la detección de una sola prueba fabricada por algoritmos puede ser instrumentalizada por sectores revisionistas para cuestionar la totalidad del registro histórico del Holocausto.

Asimismo, el desarrollo de testimonios interactivos y "resurrecciones" digitales mediante chatbots plantea interrogantes éticos sobre la autonomía de las víctimas. Si bien proyectos de entidades como la *USC Shoah Foundation* buscan democratizar el acceso al legado de los supervivientes, la posibilidad de que modelos de lenguaje generen respuestas no supervisadas (alucinaciones) abre la puerta a una descontextualización del trauma. En este sentido, la crítica se centra en que la IA, en su intento de hacer la historia más "consumible" para las nuevas generaciones, corre el riesgo de simplificar la complejidad del mal, sustituyendo el testimonio sagrado por una simulación algorítmica carente de peso moral.

Finalmente, la lucha contra la desinformación en este campo exige un compromiso con la transparencia. La comunidad experta coincide en que cualquier contenido generado por IA debe ser obligatoriamente etiquetado y supervisado por historiadores, evitando que la tecnología sea un fin en sí mismo y asegurando que

funcione estrictamente como una herramienta al servicio de la verdad histórica y la dignidad de las víctimas.

Anna Frank: Contacto en WhatsApp

La integración de figuras históricas como Ana Frank en plataformas de mensajería representa uno de los puntos más críticos en la intersección entre tecnología y memoria. Aunque no existe un canal oficial de WhatsApp gestionado por las instituciones que custodian su legado, diversas plataformas de IA han desarrollado simulaciones que permiten "chatear" con una versión digital de la joven. Esta práctica se presenta bajo una narrativa de innovación pedagógica, argumentando que permite a las nuevas generaciones —nativas digitales— conectar con la historia de una manera orgánica y cercana a sus hábitos comunicativos. Los defensores de estos proyectos sostienen que la IA puede actuar como un puente emocional, transformando un hecho histórico lejano en una experiencia de aprendizaje personalizada que fomenta la empatía y despierta el interés por profundizar en el estudio del Holocausto.

Sin embargo, esta posibilidad de añadir a una víctima del genocidio como un "contacto" más ha desatado una profunda controversia. El argumento positivo de la democratización del acceso a la historia choca frontalmente con el riesgo de la banalización del testimonio. Numerosos historiadores consideran que reducir una biografía marcada por el trauma a un producto de consumo interactivo despoja al relato de su peso moral. Mientras que la IA busca hacer la historia "viva", sus críticos advierten que al permitir que un algoritmo genere respuestas inéditas, se están poniendo palabras en boca de la víctima que ella jamás pronunció, transformando un documento histórico sagrado en una simulación diseñada para el entretenimiento.

Desde el punto de vista técnico, el contraste es igualmente complejo. Por un lado, la IA permite resolver dudas históricas de forma instantánea y contextualizada para cada estudiante. Por otro, el riesgo de las "alucinaciones" algorítmicas es alarmante; los modelos de lenguaje pueden suavizar la realidad del nazismo o evitar descripciones crudas para cumplir con filtros de seguridad, ofreciendo una versión edulcorada de la tragedia. Esta tensión entre la eficacia educativa y la precisión documental supone un peligro real: la creación de una "verdad sintética" donde las invenciones de la máquina podrían llegar a confundirse con los hechos probados.

Finalmente, el debate se desplaza hacia la ética de la monetización y el control. Aunque estos bots pueden ser herramientas poderosas para combatir la desinformación en tiempo real, el hecho de que plataformas privadas utilicen la identidad de víctimas de la Shoah plantea dilemas sobre la propiedad del legado histórico. La conclusión predominante entre los expertos es que, si bien la IA tiene

un potencial positivo innegable para revitalizar la memoria histórica, este debe estar estrictamente supeditado a una supervisión humana y académica que impida que la fascinación por la tecnología termine por disolver la frontera entre el respeto a la verdad y el simulacro digital.

7. Conclusiones

El discurso de odio siempre ha estado ahí a lo largo de la historia. Quizás en el siglo XX, con la consolidación de la prensa de masas, se convierte en un foco de toxicidad permanente, carente de análisis intelectual y jurídico. Desde el panfleto racista, las pintadas callejeras xenófobas, las informaciones tendenciosas de carácter antisemita, hasta el humor tóxico homófobo y machista, entre otras muchas manifestaciones, eran ya en sí mismos un ecosistema en la comunicación social.

El advenimiento y consolidación de internet a finales del siglo XX, junto a su rápida implantación social a escala global, propicia que esa toxicidad adquiera dimensiones preocupantes por su capacidad de multiplicarse y fluir rápidamente en la red.

Internet no tardaría en evolucionar hacia la llamada Web 2.0 y la irrupción de las redes sociales, un salto evolutivo que alteró significativamente —y de forma probablemente irreversible— los cimientos de la comunicación social. Este ecosistema no solo arrebató al periodismo tradicional la exclusividad de la información social, transformando al espectador pasivo en un creador activo de contenido, sino que penetró profundamente en la conducta individual y la psicología del ciudadano promedio. La generalización de los algoritmos de recomendación y las dinámicas de interacción constante han modificado nuestros hábitos diarios, consolidando una media de inversión en tiempo de pantalla que carece por completo de precedentes en la historia de la humanidad.

En este contexto, el discurso de odio alcanzó los más altos niveles de amenaza para la dignidad de la persona. Al diluirse los filtros editoriales tradicionales y descentralizarse la difusión de contenidos, los relatos deshumanizadores encontraron un terreno fértil para su propagación masiva y descontrolada.

En este punto, la inteligencia artificial emerge como el siguiente gran salto evolutivo, introduciendo una variable disruptiva: una velocidad de desarrollo que

resulta impredecible a corto plazo. Más allá de conjeturas apocalípticas sobre el riesgo existencial de la humanidad, este informe adopta un enfoque pragmático para diseccionar el impacto real de la IA en la propagación del discurso de odio. Para ello, se analiza el comportamiento de algunas de las herramientas de lenguaje más populares del mercado y se desentraña el complejo marco legislativo que busca regularlas. El estudio concluye explorando el potencial estratégico de esta tecnología; concretamente, su incipiente capacidad para articular contranarrativas eficaces que optimicen las estrategias de sensibilización y prevención en el entorno digital.

Es precisamente en este escenario donde el Proyecto ReTo, actualmente en fase de implementación piloto, demuestra su valor estratégico y su proyección a futuro. Partiendo de forma directa del diagnóstico riguroso que aquí se presenta, esta iniciativa no solo asume las conclusiones alcanzadas, sino que las transforma en líneas de acción prioritarias para ediciones futuras. El proyecto se postula así como un modelo escalable capaz de responder a las brechas detectadas, ofreciendo soluciones prácticas que combinan el rigor técnico con la intervención social para neutralizar las nuevas dinámicas del odio en la red.

Tras el análisis exhaustivo de los informes de monitorización, la literatura académica reciente y los marcos regulatorios vigentes, se evidencia que el discurso de odio en el entorno digital ha dejado de ser un problema coyuntural para consolidarse como un fenómeno estructural y permanente. Aunque se observan repuntes de hostilidad estrechamente vinculados a eventos específicos —como la gestión de flujos migratorios o crisis imprevistas—, la realidad es que la toxicidad persiste con una preocupante estabilidad con el paso del tiempo. Este goteo constante saca a la luz una base sólida e instalada de racismo y xenofobia, un desafío frente al cual el **Proyecto ReTo** busca erigirse como una respuesta estratégica, interviniendo directamente en el ecosistema digital donde se normalizan estas conductas.

En esta nueva arquitectura de la comunicación, la inteligencia artificial emerge claramente como un arma de doble filo. Por un lado, la IA generativa actúa como un catalizador sin precedentes para la manipulación masiva, permitiendo la creación de "enjambres" de identidades sintéticas hiperrealistas que saturan el debate público con narrativas deshumanizadoras. Por otro lado, la tecnología se postula como la única herramienta indispensable y con la escala suficiente para atajar el problema. Es en esta vertiente defensiva y tecnológica donde el **Proyecto ReTo** identifica sus mayores aspectos potenciales, adoptando sistemas avanzados de monitorización que permiten procesar y clasificar volúmenes de odio que resultarían inasumibles mediante la mera revisión humana.

Sin embargo, esta capacidad de detección choca frontalmente con la baja eficacia en la moderación por parte de las grandes plataformas, abriendo una brecha crítica en la retirada efectiva de los contenidos denunciados. En el contexto español, las tasas de eliminación apenas cubren una tercera parte de las notificaciones de odio, lo que permite que los mensajes que incitan a la violencia permanezcan visibles durante periodos prolongados. Esta impunidad digital no solo perpetúa el daño, sino que alimenta de forma directa un desorden informativo basado en *deepfakes* y bots orientados a debilitar la confianza democrática; un escenario adverso que el **Proyecto ReTo** toma como punto de partida para diseñar herramientas de contrachoque que suplan las carencias de la autorregulación corporativa.

A esta compleja encrucijada se suma el hecho de que el potencial de la inteligencia artificial para articular contranarrativas eficaces no se limita a la generación de textos estáticos. Hoy en día, la tecnología permite desplegar presentaciones de alto impacto, argumentaciones dinámicas, piezas musicales e incluso vídeos hiperrealistas. A través del uso constructivo de los *deepfakes*, se abre la asombrosa posibilidad de interactuar y hablar directamente con figuras históricas esenciales como Martin Luther King, Primo Levi o Ana Frank en su confinamiento de Ámsterdam, o incluso recrear el testimonio postrero de una víctima mortal de un delito de odio. El **Proyecto ReTo** asume este ecosistema multimedia como un pilar de vanguardia para sus líneas de trabajo, explorando cómo la tecnología de última generación puede ponerse al servicio de la memoria y los derechos humanos.

No obstante, esta inmensa capacidad de recreación digital debe gestionarse de forma obligatoria desde una ética de la responsabilidad que sea escrupulosa con la verdad objetiva. El **Proyecto ReTo** asume el compromiso de que cualquier desarrollo pedagógico de esta índole exige un marco contextual y explicativo absolutamente claro, transparente y sin margen para la duda, que aleje de inmediato toda tentación de instrumentalización o manipulación histórica.

8. Fuentes y recursos en internet

iVerify:

<https://www.undp.org/digital/verify>

Grok:

[Commission investigates Grok and X's recommender systems under the Digital Services Act](#)

DeepSeek:

[Chinese AI chatbot DeepSeek censors itself in realtime, users report](#) de Robert Booth y Dan Milmo

[DeepSeek R1 Red Teaming Report](#)

[Un informe revela preocupantes vulnerabilidades en el chatbot de IA de DeepSeek: no superó ninguna prueba de seguridad - Infobae](#)

[OpenAI, DeepSeek, and Google Vary Widely in Identifying Hate Speech | Annenberg](#)

[There is No War in Ba Sing Se: A Global Analysis of Content Moderation in Large Language Models](#)

Gemini:

[Supervisión de abusos | Gemini API | Google AI for Developers](#)

[Decoding Hate: Exploring Language Models' Reactions to Hate Speech](#)

[MODERACIÓN DE CONTENIDOS EN INTELIGENCIA ARTIFICIAL GENERATIVA: CHATGPT BAJO EL MARCO REGULATORIO DE LA UNIÓN EUROPEA](#)

Perplexity:

[Ethical safeguards and vulnerabilities in large language models](#) de Jordan Stuckey, (Georgia Southern University), Hayden Wimmer, (Georgia Southern University) y Carl M. Rebman, Jr. (University of San Diego).

Deepfakes:

[Deepfakes leveled up in 2025: Here's what's coming next - University at Buffalo](#)

Superbots y enjambres IA:

[USC Study Finds AI Agents Can Autonomously Coordinate Propaganda Campaigns Without Human Direction](#)

[How to Detect Coordinated Inauthentic Behavior on Social Media](#)

IA y Contranarrativa:

[Can NLP Tackle Hate Speech in the Real World? Stakeholder-Informed Feedback and Survey on Counterspeech](#)

[MACHINES AGAINST RAGE GENERATING HIGH-QUALITY COUNTERSPEECH WITH LANGUAGE MODELS](#) de Helena Bonaldi (Università di Trento)

[Experts debate whether AI-generated Holocaust content aids memory or distorts it](#)