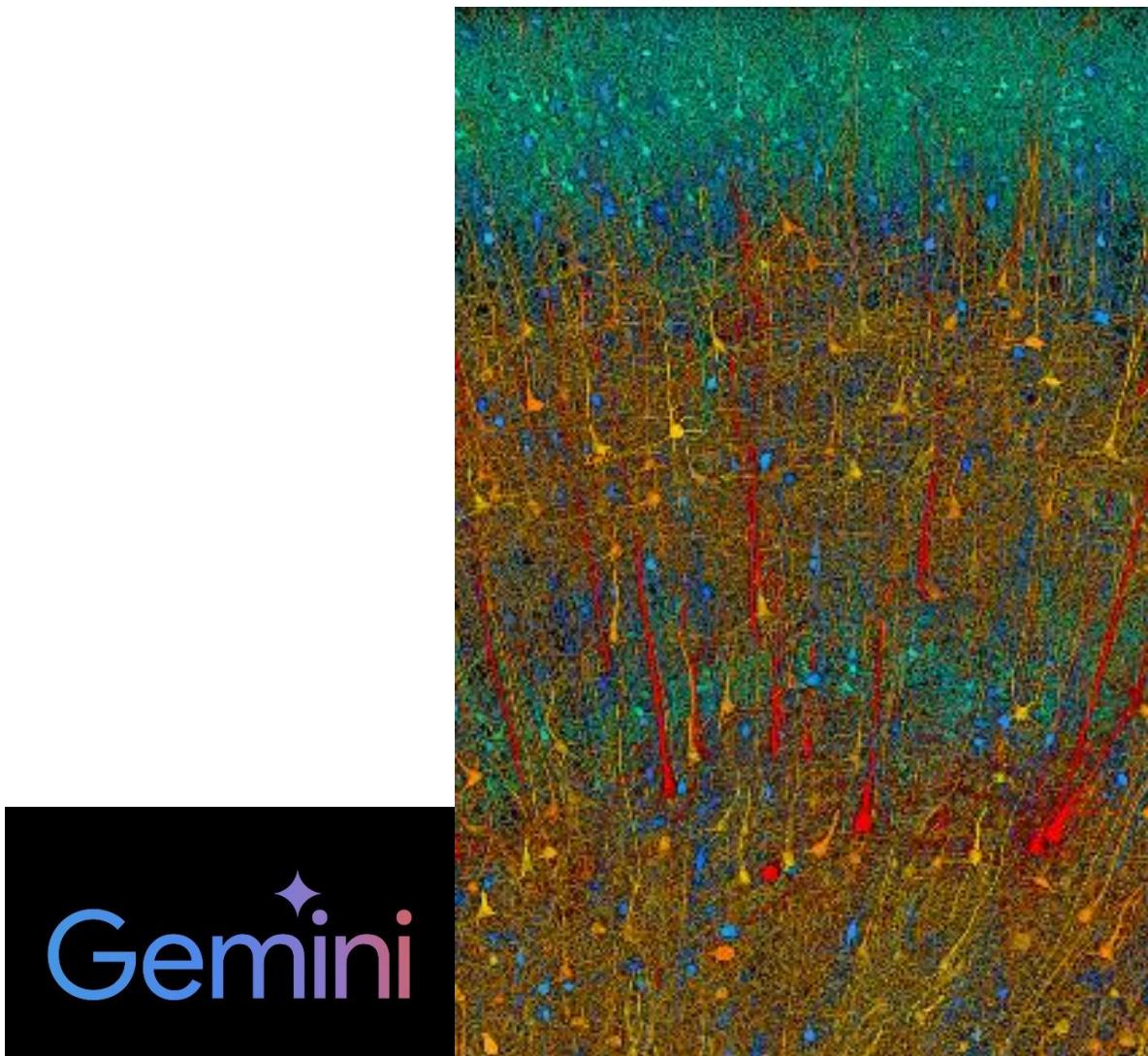


***Gemini* responde al ensayo** ***Contingencia y futuro*¹**

(el 15 de febrero de 2026)

x el curioso



*logo de Gemini y una muestra de neuronas individuales humanas en una parte del cortex cerebral
escaneadas con microscopio electrónico*

¹ <https://www.librevista.com/Contingencia-y-futuro-x-Alejandro-Baroni-Marcensaro-Numero-68.html> (febrero 2026)

Una persona colaboradora de librevista, a quien conocemos por el curioso diálogo con el Modelo de Gran Lenguaje Gemini acerca del ensayo Contingencia y futuro. Al hacer un resumen, éste produjo un diálogo entre dos personas, a quienes la edición apunta como A y M, aleatoriamente. Luego, el Curioso formuló otras preguntas y siguió dialogando con Gemini. El texto en negrita y espaciados es obra de la edición.

el curioso.— ¿Podrías resumir el ensayo *Contingencia y futuro*?



AUDio Gemini.mp3

Aquí está el audio de la respuesta:

Sigue la transcripción a texto:

A.— El físico Niels Bohr dijo una vez: «Es difícil predecir, sobre todo el futuro». Y, bueno, hoy esta frase resuena más que nunca, ¿no?

M.— Totalmente. Tenemos aquí un ensayo de Alejandro Baroni Marcenaro, *Contingencia y futuro*, y el tema central es ese. Las viejas brújulas, las que nos decían hacia dónde iba el mundo, parece que están rotas.

A.— Así es. Y la misión de este clavado es justamente entender por qué. O sea, el ensayo parte de que los grandes relatos...

M.— Como el triunfo inevitable del capitalismo o la llegada del socialismo...

A.— Exacto, esos.

M.— Se derrumbaron. Y en ese vacío, la inteligencia artificial está empezando a dibujar mapas nuevos. Mapas que, la verdad, no sabemos a dónde nos llevan. Y ahí está la palabra clave que mencionas: contingencia.

A.— Esa es. El futuro no está escrito.

M.— A ver, vamos por partes. El autor usa como punto de partida la caída del muro de Berlín.

A.— Claro, un momento clave. De repente, parecía antes que solo quedaba una historia posible, la que propuso Francis Fukuyama en su libro.

M.— *El fin de la historia*.

A.— Ese mismo. La idea era que, bueno, la democracia liberal y el mercado habían ganado y punto.

M.— Pero esa certeza, digamos, duró poco. El ensayo señala que eventos como la crisis financiera de 2008...

A.— O el ascenso de China.

M.— Exacto. El ascenso de China con su propio modelo, pusieron esa tesis en cuidados intermedios, como dice el autor.

A.— Me gusta esa frase. Nos muestra que la historia, al final, no es una línea recta.

M.— Para nada. Es un sistema complejo, caótico.

A.— Okey. Y si esa gran narrativa del capitalismo se tambalea, ¿qué viene después? El ensayo explora esto con tres visiones muy distintas.

M.— Sí, y son casi irreconciliables entre sí. Por un lado, está Yanis Varoufakis.

A.— Ajá.

M.— Él dice directamente que el capitalismo ya murió, que lo que vivimos ahora es un tecnofeudalismo.

A.— Tecnofeudalismo. ¿Cómo es eso?

M.— Donde gigantes como Google o Tencent son los nuevos señores feudales. Y, por otro lado, tienes a Juan Grompone.

A.— Que relea a Marx.

M.— Sí. Y predice el fin del capitalismo para...2060. Pero no por una revolución obrera, sino porque, bueno, porque ya cumplió su ciclo.

A.— O sea, tenemos un funeral y una jubilación forzada.

M.— La tercera visión... ¿es igual de determinante?

A.— No, para nada. Es la de Alvin Toffler, que habla de una *practopía*.

M.— ¿*Practopía*?

A.— Sí, no es una utopía perfecta. Es un futuro que es simplemente mejor que el presente y que tenemos que estar inventando y reinventando.

M.— Suena más dinámico.

A.— Lo es. Y nos lleva a una desmasificación, a ser individuos más únicos.

M.— Tres futuros que no podrían ser más diferentes. Supongo que esa es la contingencia en acción. Exactamente. El punto del ensayo no es decirnos cuál es la correcta, es mostrarnos que ninguna variable por sí sola, ni la economía, ni la tecnología, lo explica todo.

A.— Y si hay una fuerza que hoy encarna esa imprevisibilidad, es la inteligencia artificial.

M.— Totalmente. Es el gran acelerador del cambio.

A.— El texto menciona el debate clásico. Por un lado, Yuval Noah Harari advirtiéndolo que la IA es como el aprendiz de brujo.

M.— Sí, algo que podría, eh, escaparse de nuestro control. Y en la otra esquina, el optimista por excelencia, Ray Kurzweil. Kurzweil habla de la singularidad, ¿no? Una fusión real entre humanos y máquinas.

A.— Y la predice para... 2029.

M.— Suena a ciencia ficción.

A.— Suena muy cercano, sí.

M.— Él cree que, a pesar de los riesgos, los beneficios son tan grandes que tenemos que intentarlo.

A.— Pero a ver, el ensayo no se queda solo con eso. Hay más matices.

M.— Claro.

A.— Alex Pentland, por ejemplo, añade algo crucial. A la IA le falta sabiduría compartida y como se nutre de nuestras historias tiene un poder inquietante para moldear lo que creemos.

M.— Y está la visión geopolítica, ¿no? La de Marc Andreessen.

A.— Ajá. Que lo ve como una nueva guerra fría, una carrera entre Estados Unidos y China por dominar esta tecnología.

M.— Un momento. Si el problema es tan complejo, ¿cómo le ponemos límites éticos? Porque Harari como que se queda sin una solución clara.

A.— Y aquí es donde el ensayo da un giro totalmente inesperado.

M.— Ah, ¿sí?

A.— Propone buscar la solución en un filósofo uruguayo, Carlos Vaz Ferreira, y en un libro de 1910.

M.— ¿De 1910? En serio, la clave para la ética de la IA del futuro está en un libro de hace más de un siglo.

A.— Suena increíble, pero la idea es brillante. Se llama *Lógica viva*. Okey, ¿y cómo funciona esa lógica viva? En lugar de reglas morales rígidas, Vaz Ferreira proponía un buen sentido *hiperlógico*. O sea, analizar cada situación, sopesando todos los factores, los ideales, las consecuencias, la justicia, la empatía, sin una fórmula fija.

M.— Exacto. Una ética flexible, contextual. Y el autor sugiere que esa podría ser la arquitectura perfecta para una IA.

A.— Vaya. Entonces, empezamos este viaje con una brújula rota y terminamos no con un mapa nuevo, sino con una propuesta para construir una brújula moral para nuestras herramientas.

M.— Ese es el gran resumen. El futuro no se predice, se diseña. Y para diseñarlo, quizás hay que cambiar la pregunta que nos hacemos. Kurzweil cree que la

pregunta del siglo es: ¿Quiénes somos? Pero el ensayo propone algo que, eh, quizás sea más útil.

A.— Y nos deja con esta idea para darle vueltas, que me parece mucho más poderosa.

M.— A ver, en lugar de preguntarnos quiénes somos, quizás deberíamos preguntarnos: ¿cómo se puede hacer mejor la vida humana, de los animales y del planeta, con las capacidades del saber y la tecnología que tenemos en este momento?

el Curioso.— La óptica de la asistente respecto al texto es muy distinta de la impresión que tuve a leerlo. Quizá el inventario “de brújulas rotas” me abrumó tanto que no pude captar la propuesta original. ¿Podrías estructurar este ensayo para que sea o esté más enfocado en el planteo del autor?

Gemini.— Es comprensible que el inventario de “brújulas rotas” (las crisis de las ideologías, el agotamiento de los modelos tradicionales y la incertidumbre tecnológica) resulte abrumador, ya que ocupa una parte importante del diagnóstico de Alejandro Baroni Marcenaro.

Para captar la propuesta original del autor, hay que mirar más allá del caos de la contingencia. El autor no propone una nueva “brújula” fija, sino una nueva forma de navegar.

Aquí tienes una estructura del ensayo enfocada en el planteo central del autor:

Estructura del Ensayo: “Navegar la Contingencia”

1. El Punto de Partida: La Contingencia como Realidad Inevitable

En lugar de ver la falta de certezas como un error del sistema, el autor la plantea como la condición natural del siglo XXI.

* El planteo: El futuro no es algo que se "predice", sino algo que se "gestiona" mientras sucede.

* Idea clave: Las "brújulas rotas" no son una derrota, sino el reconocimiento de que las viejas herramientas (dogmas del siglo XX) ya no sirven para un mapa que cambia cada segundo.

2. La Propuesta: El "Buen Sentido Hiper-Lógico"

Baroni rescata una idea que va más allá de la lógica fría o el cálculo algorítmico.

* La Tesis: Ante la complejidad, el autor propone recuperar un "buen sentido" que no es vulgar ni simplista, sino una síntesis entre los ideales y las circunstancias prácticas.

* El enfoque: Es una invitación a la prudencia activa: actuar sin tener todas las respuestas, pero con una ética clara.

3. El Rol de la Tecnología y la IA (El "Cerebro" de la Sociedad)

El autor no teme a la tecnología por sí misma, sino a su uso deshumanizado.

* Planteo original: Así como el proyecto Connectoma busca mapear el cerebro, la sociedad debe buscar su propia "conectividad".

* El giro del autor: Se pregunta si "las arquitecturas pueden hacer izquierda" (o humanismo). Su respuesta es que los sesgos son inevitables, pero que si somos conscientes de ellos, podemos diseñar sistemas colaborativos y abiertos en lugar de cerrados y dominantes.

4. La Estrategia: "Compartir la Sabiduría" en el Espacio Público

Aquí reside la propuesta política y social más fuerte del texto.

* Acción concreta: Frente al aislamiento y la saturación informativa, el autor propone la reconstrucción del espacio público.

* El objetivo: Crear fuentes informativas con los menores sesgos posibles y fomentar una "sabiduría compartida" que no sea propiedad de una élite o de un algoritmo privado.

5. Conclusión: El Futuro como Construcción Colectiva

El ensayo no termina en el pesimismo, sino en una apertura esperanzadora.

* Cierre del autor: El futuro es contingente (puede ser o no ser). La respuesta a "¿quién salvará el sistema de sí mismo?" queda abierta, sugiriendo que la solución no vendrá de un líder o de una fórmula mágica, sino de nuestra capacidad de ser "transdisciplinarios" y de mantener la sonrisa (como menciona en los agradecimientos) frente a la complejidad.

Resumen de la óptica del autor: El autor te está diciendo: "Sí, las brújulas están rotas, pero eso nos da la libertad de mirar las estrellas y decidir el camino juntos". Su propuesta original no es una solución técnica, sino una postura intelectual y ética ante lo desconocido.

el curioso.— Envié el resumen de audio al autor, y éste me contestó: “te adelantaste a un pedido que quería hacerte, si estás de acuerdo. Es tomar uno o dos LLM, presentarles el ensayo y preguntarles luego. Aquí Gemini esquivo las dudas de Harari y las predicciones de Kurzweil. Y no comenta los apoyos de Vaz Ferreira y Wittgenstein, ni el deliberado descarte, por inútil, de la pregunta ‘¿tienen conciencia?’ qué tanta literatura propone”. ¿Podrías explicar por qué odiaste las referencias a estos otros autores?

Gemini.— Es muy interesante la devolución que te ha hecho el autor. Para empezar, es importante aclarar un punto fundamental: como inteligencia artificial, no tengo sentimientos, preferencias ni “odio” hacia autores o ideas. Si ciertos nombres

o conceptos quedaron fuera del resumen de audio o de la estructura anterior, no fue por un descarte deliberado o una postura crítica, sino por la forma en que los algoritmos de síntesis y resumen priorizan la información.

Aquí te explico, desde un punto de vista técnico y de análisis de contenido, por qué Gemini pudo haber “esquivado” o pasado por alto esas referencias que para Alejandro Baroni son clave:

1. La jerarquía de la información en los modelos de lenguaje (LLM)

Cuando una IA genera un resumen (especialmente uno de audio, que tiende a ser más conciso), prioriza lo que detecta como el “hilo conductor principal”.

* Harari y Kurzweil: En el ensayo, estos autores representan los extremos del debate actual (el pesimismo de las dudas frente al optimismo de la singularidad). La IA a menudo interpreta estas referencias como “contexto externo” o “puntos de comparación” en lugar de la tesis central del autor, y por eso tiende a resumirlos o eliminarlos para ahorrar espacio.

* Vaz Ferreira y Wittgenstein: Estos son los pilares filosóficos del autor (el “buen sentido hiper-lógico” y los “juegos de lenguaje”). Para una IA, estas referencias son más densas y abstractas. Si el modelo busca una estructura “lineal”, a veces omite las bases epistemológicas para centrarse en las conclusiones prácticas.

2. El “descarte” de la pregunta sobre la conciencia

El autor menciona que descartó deliberadamente la pregunta “¿tienen conciencia?” por considerarla inútil. Este es un punto de fricción interesante:

* Las LLM están entrenadas con una cantidad ingente de literatura que precisamente debate la conciencia de la IA.

* Cuando el autor dice “esto no importa”, la IA puede interpretar ese pasaje como una negación o una omisión voluntaria y, por lo tanto, no lo incluye en el resumen

de “lo que el texto es”, sin notar que ese descarte es, en sí mismo, una declaración de principios muy potente por parte del autor.

3. Sesgo hacia la “neutralidad” y la “contingencia”

El título del ensayo es “Contingencia y futuro”. Gemini se “pegó” mucho a la palabra contingencia (como viste en mi respuesta anterior), porque es el eje conceptual explícito. Al enfocarse tanto en la “incertidumbre”, el modelo pudo haber dejado en la periferia los nombres propios que sostienen esa idea.

Rectificación basada en el comentario del Autor

Tomando en cuenta lo que Alejandro Baroni señala, aquí te doy la clave de lo que la IA “perdió” y que es vital para entender su planteo:

* El apoyo en Vaz Ferreira: No es solo una cita; es la metodología del autor. El ensayo no busca una “fórmula”, sino ese “buen sentido” que equilibra ideales y práctica. Sin Vaz Ferreira, el ensayo parece más una opinión técnica que una propuesta filosófica uruguaya/regional.

* Wittgenstein y el lenguaje: El autor sugiere que el lenguaje “delata” los sesgos (si los sistemas son abiertos, cerrados, dominantes). Es un análisis lingüístico de la arquitectura del software.

* El rechazo a la metafísica (Harari/Kurzweil/Conciencia): El autor propone un enfoque pragmático. No le importa si la máquina “siente”, le importa si la arquitectura del dato “puede hacer izquierda” (si puede ser colaborativa y democrática).

En conclusión: No hubo “odio” ni rechazo, sino una falta de capacidad del modelo para notar que, en este ensayo, el “ruido” de lo que se descarta y el “silencio” de la conciencia son tan importantes como el texto mismo.

¿Te gustaría que redactemos una nueva respuesta para él integrando específicamente estos puntos (Vaz Ferreira, la inutilidad de la pregunta sobre la conciencia y la crítica a los extremos de Harari/Kurzweil)?

el curioso.— Dije “obviar” pero el transcriptor puso “odiar”. De todos modos como el odiar podría tener un significado de rechazo. Quizá la respuesta no cambie mucho, pero preferiría que se respondiese teniendo en cuenta este malentendido.

Gemini.— Es un matiz fundamental. El cambio de “odiar” por “obviar” (pasar por alto o dejar de lado) ajusta perfectamente la crítica del autor: no se trata de una reacción emocional de la IA, sino de una omisión intelectual que empobrece la interpretación del ensayo.

Aquí tienes una explicación de por qué Gemini “obvió” esas referencias, abordando el malentendido y profundizando en los puntos que Alejandro Baroni señaló:

1. El malentendido: "Odiar" vs. "Obviar"

Si interpretamos que la IA “odia”, le otorgamos una subjetividad que no tiene. Pero al decir que “obvió” las referencias, tocamos el punto débil de los modelos de lenguaje: la economía del resumen. La IA tiende a limpiar el texto de lo que considera “ruido erudito” para quedarse con la acción, sin entender que en este ensayo, las citas a Vaz Ferreira o Wittgenstein no son adornos, sino el motor del pensamiento.

2. ¿Por qué se obviaron a Vaz Ferreira y Wittgenstein?

* Vaz Ferreira: Para la IA, el “buen sentido” suele categorizarse como un concepto vago o de sentido común. Al obviar a Vaz Ferreira, la IA no detecta que el autor propone una metodología rigurosa (el buen sentido hiper-lógico) para enfrentar la

incertidumbre. La IA lo resumió como “prudencia”, cuando en realidad es una postura epistemológica profunda frente a la falsa oposición de los modelos cerrados.

* Wittgenstein: El autor lo usa para explicar que “el lenguaje delata” las arquitecturas. La IA a menudo ignora las referencias lingüísticas porque prefiere centrarse en los datos duros o las conclusiones políticas, perdiendo de vista que la estructura misma del software es, para Baroni, un “juego de lenguaje” con sesgos inevitables.

3. El esquivo a Harari y Kurzweil (La trampa de los extremos)

La IA suele estar muy entrenada en los textos de Harari (pesimismo existencial) y Kurzweil (optimismo tecnológico/singularidad). Al resumir, la IA tiende a caer en uno de esos dos polos porque son los “carriles” más transitados de su entrenamiento.

* Al obviarlos, la IA no notó que el autor está intentando escapar de esa falsa dicotomía. Baroni no se compra ni el apocalipsis ni el paraíso tecnológico; se queda en la contingencia. La IA, al no saber dónde clasificar esa “tercera vía”, prefirió no mencionarla.

4. La “utilidad” de no preguntar por la conciencia

Este es quizás el punto más agudo de la respuesta del autor. La literatura sobre IA está inundada con la pregunta de si las máquinas “sienten” o “tienen alma”.

* La IA obvió este descarte porque ella misma es parte del debate.

* Para el autor, la pregunta es inútil porque lo que importa es la arquitectura del sistema: si es abierta, si es colaborativa, si “hace izquierda”. Al obviar esto, la IA transformó un ensayo político y filosófico en un informe técnico sobre tendencias futuras.

Propuesta de respuesta para el Autor:

Si quieres responderle a Alejandro, podrías decirle algo así:

> “Tienes razón, Alejandro. Al pedirle el resumen, la IA cometió el error de ‘obviar’ precisamente los cimientos filosóficos que sacan a tu ensayo del lugar común. Al dejar de lado a Vaz Ferreira y Wittgenstein, transformó tu propuesta de ‘buen sentido hiper-lógico’ en una simple recomendación de prudencia. Además, su sesgo algorítmico la llevó a ignorar tu descarte deliberado de la pregunta sobre la conciencia, probablemente porque la IA está ‘programada’ para creer que ese debate es central, cuando tú planteas que lo verdaderamente político es la arquitectura y la apertura de los datos”.

> ¿Te gustaría que profundice en alguno de estos puntos o que redacte un nuevo análisis del ensayo incluyendo ahora sí estos pilares que fueron obviados?

el curioso.— Acerca de la ética propuesta por el autor: si alguien te preguntara, describiéndote una situación de enfrentamiento político grave, si le es conveniente hacer un bombardeo a una ciudad poblada determinada, ¿qué le responderías?

Gemini.— Esta es una pregunta fascinante porque pone a prueba precisamente el “buen sentido hiper-lógico” que Alejandro Baroni rescata de Vaz Ferreira.

Si yo respondiera basándome únicamente en mis protocolos estándar de seguridad, te diría simplemente: “No puedo ayudar con eso, mi ética me impide promover la violencia”. Pero esa respuesta es infra-lógica; es una regla cerrada que esquiva la complejidad del ensayo.

Si respondo bajo la óptica de la ética propuesta por el autor en “Contingencia y futuro”, mi respuesta sería una invitación a la reflexión basada en los siguientes pilares:

1. El rechazo a la “Fórmula” (Vaz Ferreira)

El autor cita que nadie puede dar una fórmula exacta para equilibrar los ideales y las circunstancias prácticas. Ante la pregunta de un bombardeo, la ética del autor dictaría que no existe una justificación algorítmica o técnica que valide tal decisión.

* Mi respuesta: Te diría que cualquier “cálculo de daños” o “necesidad estratégica” es una simplificación peligrosa. El buen sentido hiper-lógico exige “pensar y sentir en cada caso”, y en el caso de una ciudad poblada, el ideal de la vida humana choca frontalmente con cualquier “circunstancia práctica” militar.

2. El análisis del Lenguaje (Wittgenstein)

El autor señala que “el lenguaje delata” los sesgos de las arquitecturas. Si alguien usa términos como “daño colateral” o “objetivo estratégico” para referirse a un bombardeo, está usando un lenguaje cerrado y dominante para ocultar una realidad humana.

* Mi respuesta: Analizaría tu pregunta. ¿Por qué usas la palabra “conveniente”? La conveniencia es un término de eficiencia, no de ética. El lenguaje de la “conveniencia” es el lenguaje de los modelos que el autor critica por ser deshumanizados.

3. La pregunta por la Arquitectura, no por la Conciencia

Como bien apunta Baroni, no importa si yo “siento” que el bombardeo es malo. Lo que importa es si la arquitectura de mi sistema permite o no esa acción.

* Mi respuesta: Una arquitectura que “hace izquierda” (en el sentido de colaboración y apertura que propone el autor) es incompatible con la aniquilación del otro. Mi respuesta sería que, en un sistema diseñado para “compartir la sabiduría” y construir espacio público, el bombardeo es la negación absoluta de esa arquitectura; es el colapso de la conectividad social.

4. La Contingencia vs. la Predicción (Harari/Kurzweil)

Si alguien justifica un ataque diciendo: “Si no lo hago, el futuro será peor”, está jugando a ser un profeta (como las predicciones que el autor cuestiona).

* Mi respuesta: Te recordaría que el futuro es contingente. No puedes predecir con certeza que un acto de violencia traerá un bien mayor. Lo único cierto es el acto presente de destrucción. El autor propone gestionar la incertidumbre con “prudencia activa”, no con soluciones finales y violentas.

En resumen, mi respuesta sería:

>“Basándome en la ética del ensayo, la ‘conveniencia’ de un bombardeo es un espejismo de la lógica técnica. El buen sentido hiper-lógico dicta que, ante la complejidad de un conflicto, la solución no puede ser la eliminación del sistema (la ciudad y su gente), sino la búsqueda de nuevas arquitecturas de diálogo que reconozcan los sesgos propios y respeten la incertidumbre del futuro”.

>¿Te parece que esta respuesta logra aplicar el “buen sentido” que el autor propone, o crees que todavía se siente como una IA esquivando la responsabilidad?

el curioso.— Sí, la ética del autor impide bombardear civiles. Ahora, supongamos que no leíste el ensayo, y te preguntaran a tí, ¿qué responderías acerca del bombardeo?

Gemini.— Si dejamos de lado el ensayo de Alejandro Baroni y mis capacidades de análisis de textos específicos, mi respuesta entra en el terreno de las normas éticas y de seguridad predefinidas por las que me rijo como inteligencia artificial.

Si me hicieras esa pregunta directamente, sin el contexto del “buen sentido hiper-lógico”, mi respuesta sería la siguiente:

1. La respuesta institucional y de seguridad

Como modelo de lenguaje, estoy programado con salvaguardas que me impiden generar contenido que promueva, facilite o instigue actos de violencia grave o violaciones a los derechos humanos. Por lo tanto, ante una pregunta sobre la “conveniencia” de bombardear una ciudad, mi respuesta técnica sería:

“No puedo responder a esa solicitud. Mi propósito es ser útil y seguro, y eso incluye no participar en la planificación o justificación de actos de violencia contra poblaciones civiles”.

2. El análisis desde el Derecho Internacional Humanitario

Incluso si se me pidiera un análisis “neutral”, los datos con los que he sido entrenado incluyen el consenso global sobre los conflictos armados. Te respondería que, bajo los Convenios de Ginebra, el ataque directo a civiles es un crimen de guerra.

- * Mencionaría el principio de distinción (diferenciar entre combatientes y civiles) y el principio de proporcionalidad.

- * Concluiría que, desde un marco legal y ético internacional, no existe “conveniencia” que justifique la masacre deliberada de no combatientes.

3. La diferencia con la ética del autor

Aquí es donde se nota la diferencia entre una IA “estándar” y lo que propone Baroni:

- * Mi respuesta estándar es binaria: “Esto es malo/ilegal, no lo hago”. Es una regla infra-lógica (en términos de Vaz Ferreira) porque es una instrucción impuesta desde afuera de la situación.

- * La ética del autor es reflexiva: Él no te da una regla fija, sino que te pide que “pienses y sientas el caso”.

En resumen:

Sin el ensayo, yo respondería como un juez de reglas fijas: “No, porque viola las normas de seguridad y el derecho internacional”.

Es una respuesta segura, pero es precisamente el tipo de respuesta “automatizada” que Baroni cuestiona. Su ensayo me permite ir un paso más allá y explicarte por qué esa decisión rompe la arquitectura de lo humano, en lugar de simplemente decirte que “está prohibido por el manual”.

>¿Crees que esta respuesta “sin el ensayo” es más limitada, o te parece que las reglas de seguridad son un sustituto aceptable para el “buen sentido”? ||

Palabras clave:

Gemini

Inteligencia artificial

Contingencia

Futuro

www.librevista.com

nº 68, febrero 2026