

La traducción literaria personalizada en el contexto de la inteligencia artificial conversacional

Personalized Literary Translation in the Context of Conversational Artificial Intelligence

Jorge Hernán Yerro

Universidade Federal da Bahia (UFBA)

Salvador | BA | BR.

heryerro@gmail.com

<https://orcid.org/0000-0002-6833-3691>

Resumen: Este ensayo examina los efectos de la personalización algorítmica en la traducción automática de textos literarios realizada por IA generativas, con foco en *ChatGPT*. A partir de un experimento comparativo con dos cuentas de perfiles ideológicos contrastantes, se analiza de qué manera el sistema ajusta las traducciones de tres relatos cortos en función de las preferencias inferidas durante la interacción o explícitamente configuradas. Se argumenta que esta dinámica produce un *sesgo de anclaje deseado*, una forma específica de sesgo interpretativo configurada por mecanismos de perfilado y validada por la propia persona usuaria. Retomando debates sobre neutralidad tecnológica (Winner, Feenberg), psicopolítica (Han) y coautoría algorítmica (Novaes), el ensayo propone que este sesgo revela una nueva mediación discursiva, marcada por la adhesión afectiva. Al discutir sus impactos sobre la construcción de sentido, el estudio busca contribuir a una revisión crítica de la traducción automática como práctica situada e ideológicamente moldeada.

Palabras clave: traducción automática; personalización algorítmica; sesgo de anclaje deseado; inteligencia artificial generativa.

Abstract: This essay examines the impacts of algorithmic personalization on the machine translation of literary texts by generative AIs, with a focus on *ChatGPT*. Based on a comparative experiment using two accounts with contrasting ideological profiles, the study analyzes how the system adjusts the translations of three short stories according to preferences either inferred during



interaction or explicitly configured. It is argued that this dynamic produces a desired anchoring bias—a specific form of interpretative bias shaped by profiling mechanisms and validated by the user. Drawing on debates around technological neutrality (Winner, Feenberg), psychopolitics (Han), and algorithmic co-authorship (Novaes), the essay suggests that this bias reveals a new kind of discursive mediation, marked by affective adhesion. By discussing its impact on meaning-making processes, the study seeks to contribute to a critical review of machine translation as a situated and ideologically shaped practice.

Keywords: machine translation; algorithmic personalization; desired anchoring bias; generative artificial intelligence.

Introducción

En los últimos años, la traducción automática (TA) se ha incorporado al repertorio habitual de herramientas digitales vinculadas al lenguaje, impulsada por avances en los sistemas estadísticos y neuronales. Desde entonces, motores como *Google Traductor* y *DeepL* se consolidaron como referentes por ofrecer traducciones instantáneas, gratuitas y accesibles. Al demostrar capacidad para generar textos cohesivos y sintácticamente coherentes en una amplia variedad de idiomas, estos sistemas se integraron a los hábitos digitales, siendo empleados con frecuencia en tareas cotidianas, aunque no siempre con plena satisfacción ni total confianza en su calidad (Kasperé *et al.*, 2021). Hacia finales de 2022, la llegada de las inteligencias artificiales generativas como *ChatGPT* y *Gemini* introdujo una nueva dimensión al proceso de traducción, cuyos resultados pasaron de ser producto exclusivo del motor a efecto del diálogo entre la persona usuaria y el sistema.

Este ensayo examina las consecuencias de esta transformación, enfocándose en la incidencia de la personalización algorítmica sobre la construcción de sentido del texto traducido. A través de un experimento comparativo de traducción de tres microrrelatos del español al inglés realizado en *ChatGPT*, con dos perfiles ideológicos contrastantes, propongo que este tipo de traducción configura una mediación discursiva sesgada que se naturaliza por coincidir con las expectativas de la persona usuaria. Bajo esta perspectiva, mi hipótesis es que la anticipación algorítmica de interpretaciones personalizadas genera lo que denomino un *sesgo de anclaje deseado*, donde está en juego no solo cómo se traduce, sino qué versiones del sentido se privilegian, mientras se fomenta una experiencia de lectura sin fricciones.

La personalización como régimen algorítmico general

La personalización algorítmica es una operación fundamental de los sistemas digitales actuales. Plataformas como *ChatGPT* recolectan y procesan continuamente datos sobre las acciones de quienes las usan, para perfilar comportamientos y anticipar preferencias (Noticias.Ai, 2024). Dicho procedimiento se realiza combinando inferencias automáticas con instrucciones personalizadas que la persona proporciona explícitamente (Europapress, 2023). Los grandes modelos de lenguaje actuales, conforme verifican Tan y Jiang (2023, p. 3), son altamente eficaces para identificar, organizar y modelar datos de las personas usuarias, demostrando capacidades robustas para caracterizar personalidades, discernir posturas e identificar preferencias, lo que les permite construir representaciones detalladas de sus interlocutores para orientar la producción textual futura.

Esta sofisticada capacidad de perfilado, sin embargo, trasciende la mera adaptación a preferencias previas. Según analiza Novaes (2024), quien se hace eco de la discusión sobre las limitaciones de la idea de burbuja de filtro, el funcionamiento algorítmico contemporáneo plantea una lógica compositiva en la que la interfaz misma se configura como un espacio coproducido por la interacción entre persona usuaria y sistema. Desde esta perspectiva, quienes usan la plataforma no son agentes pasivos, sino coautores que, junto con el modelo, dan forma al contenido que se muestra. Así, la interfaz deja de ser un mero filtro y se entiende como un texto compuesto en esa coautoría.

Dicha lógica compositiva, en la que la interfaz se vuelve un texto en sí mismo, se intensifica en el caso de las inteligencias artificiales generativas, dado que estas

al recibir un *input* de determinada naturaleza, responden con la creación de un texto (en nuestro caso, el propio *feed* organizado temporalmente en un formato específico y no en otro) que está íntegramente compuesto según una lógica previamente configurada para generar una respuesta estandarizada, en estricta consonancia con su propia programación anterior (Novaes, 2024, traducción nuestra).¹

En este sentido, la interfaz pasa de ser un medio neutro a una superficie generada algorítmicamente, moldeada por la interacción previa. De ahí que, en traducciones hechas con *chatbots*, el resultado no sea un producto cerrado del motor, sino el efecto de esa coautoría entre el modelo y quien lo usa.

Obsérvese que la reflexión no parte de una cuestión técnica sobre cómo traduce el sistema. Es decir, no se trata, por ejemplo, de analizar los parámetros con que se configuran los modelos ni los idiomas en juego. En cambio, se plantea una discusión conceptual acerca de las circunstancias que determinan el resultado de esa traducción. Lo que interesa, por lo tanto, es reflexionar acerca del fenómeno mismo, es decir, examinar las incidencias de los mecanismos estructurales de perfilado, personalización y sesgo algorítmico en la traducción. En este punto, conviene recuperar tres conceptos clave del análisis crítico de la tecnología y de

¹ “ao receberem um input de uma determinada natureza, respondem com a criação de um texto (no nosso caso, o próprio feed temporariamente organizado em um dado formato específico, e não em outro) que é todo ele composto de acordo com uma lógica anteriormente configurada para gerar uma resposta padronizada em direto respeito à sua própria programação prévia”.

la traducción cuya problematización contribuye a desnaturalizar los procesos automatizados y, así, a revelar la influencia de las estructuras algorítmicas en la producción de sentido. Me refiero a la neutralidad, la transparencia y la invisibilidad.

La discusión sobre la neutralidad en la tecnología permite cuestionar supuestos que alcanza a la traducción automática. Langdon Winner, en su texto *Do Artifacts Have Politics?*, afirma que “las cosas que llamamos ‘tecnologías’ son formas de ordenar el mundo. Muchos dispositivos y sistemas técnicos, importantes en la vida cotidiana, encierran múltiples posibilidades de organizar la actividad humana” (Winner, 1980, p. 127).² El autor plantea, así, que toda tecnología incorpora, desde su concepción, decisiones estructurales sobre la realidad que interviene; decisiones estas que son previas a cualquier solución instrumental. Por lo tanto, la adopción de cualquier tecnología implica reproducir y legitimar ciertas estructuras de pensamiento y jerarquías, así como privilegiar determinados modos de operación por sobre otros.

Bajo esta perspectiva, la traducción algorítmica retoma, desde su lugar, el debate sobre la automatización como equivalencia objetiva entre idiomas mediante procesos mecánicos. Tal suposición corre el riesgo de desconocer que el resultado lanzado por estos sistemas surge de decisiones incorporadas en la arquitectura del sistema mediante entrenamientos masivos y selección de corpus durante el desarrollo del modelo, normalizando ciertos usos del lenguaje. Así, los resultados refuerzan determinadas organizaciones de sentido sobre otras e incorporan jerarquías culturales y patrones ideológicos dominantes entre los modos hegemónicos de estructurar lo traducible. Por eso, al entender la noción de neutralidad como construcción ideológica podemos ver que los dispositivos aparentemente funcionales contienen modos de ordenar la experiencia, establecer jerarquías y producir ciertos sentidos mientras excluyen otros.

También cabe considerar que, como quienes leen el texto traducido no suelen advertir que este resultó de una configuración del sistema, dicha apariencia de neutralidad tiende a disimular y, así, a invisibilizar el proceso de personalización que ejecuta el algoritmo en cada respuesta. Si bien en algunos casos, como veremos más adelante, esta personalización no se oculta del todo, dado que emerge explícitamente en la conversación o en las opciones que el sistema ofrece para generar nuevas versiones del texto, a medida que la interacción con el sistema conversacional se intensifica y que la IA recopila progresivamente las informaciones que recibe, la personalización se va haciendo imperceptible, en tanto se ajusta a las expectativas lingüísticas e ideológicas de la persona usuaria.

En este marco, a medida que la conversación avanza y la personalización se naturaliza, tiende a fortalecerse, a su vez, la confianza en los resultados generados, una confianza que puede interpretarse, a su vez, como efecto del fetichismo tecnológico previamente mencionado. A este respecto, resulta especialmente relevante la reflexión de Feenberg quien, al definir los códigos técnicos como “la realización de un interés bajo la forma de una solución técnicamente coherente a un problema”, advierte que “allí donde tales códigos están reforzados por la percepción que los individuos tienen acerca de su propio interés y de la ley, su significado político generalmente pasa desapercibido” (2005, p. 114). Esta formulación permite comprender que la confianza en la tecnología no solo refuerza la invisibilización de sus mediaciones, sino que consolida un orden perceptivo en el que los resultados expresan relaciones de poder.

² “The things we call ‘technologies’ are ways of building order in our world. Many technical devices and systems important in everyday life contain possibilities for many different ways of ordering human activity”.

Por otra parte, aún sobre la cuestión de la neutralidad y la transparencia, puede observarse que, en este caso, ambas características responden a lógicas que se inscriben en formas contemporáneas de poder, asociadas menos a la coerción directa que a modos de control difusos, que operan a través de la participación activa y la ilusión de libertad de elección. A este respecto, cabe detenerse en el concepto de psicopolítica, que Han (2014) describe como un régimen de control basado en la autoexplotación y la autoexposición. Bajo este régimen, que además de controlar la información organiza las condiciones de lo decible, lo visible y lo pensable, el sujeto pasa de ser objeto pasivo de vigilancia a convertirse en agente activo de su propia exposición, entregando información de manera voluntaria y aceptando como naturales los productos que emergen de su perfil digital. De esta manera, lo que a primera vista se presenta como libertad, como la posibilidad de expresarse o de elegir, opera más bien como una forma de transparencia forzada, en la que el acceso a ciertos espacios o servicios queda condicionado a esa entrega previa de datos.

En este marco paradójico, donde coexisten libertar y control, lo visible adquiere un estatuto de neutralidad, pues todo aquello que coincide con el perfil de la persona tiende a percibirse como más adecuado e, incluso, más verdadero. De este modo, la paradoja se intensifica cuando esa adhesión voluntaria se consolida como una forma sutil de orientar la manera en que se interpretan y validan los significados. Desde esta perspectiva, cobra sentido que operaciones complejas de modelado que orientan la interpretación se acepten sin resistencia, como si se tratara simplemente de elecciones personales.

Sobre esta base, debe considerarse un aspecto adicional del funcionamiento de estos sistemas que puede contribuir a la aceptación acrítica de las traducciones entregadas por la IA. Me refiero a la amabilidad con que los resultados se ponen a disposición, una amabilidad que, lejos de ser solo una cortesía formal, funciona como una estrategia que predispone al receptor favorablemente ante los resultados. A este respecto, retomo a Han cuando, al referirse a la forma en que el poder inteligente se expresa, es decir, sin oponerse a la violencia ni a la libertad, advierte que este

no actúa frontalmente contra la voluntad de los sujetos subyugados, controlando sus deseos en beneficio propio. Es más afirmativo que negador, más seductor que represivo. Se esfuerza en producir emociones positivas y explotarlas. Seduce, en lugar de prohibir. En vez de oponerse al sujeto, va a su encuentro (Han, 2014, p. 26-27, traducción nuestra).³

De esta manera, la vigilancia algorítmica adopta la forma de una asistencia personalizada y permite, así, la agencia eficaz de la psicopolítica, en la medida en que se muestra amable y accesible. Cuando llevamos esta lógica a la ejecución de una traducción automática en un contexto conversacional, notamos que, cuando la IA se presenta como un servicio cortés y funcional que busca ayudar a la persona usuaria ajustándose a sus demandas, tal amabilidad se configura como una influencia sutil, en tanto evita restricciones visibles y promueve la adhesión, al tiempo en que diluye sus rastros de intervención ideológica en el texto producido.

³ “não age frontalmente contra a vontade dos sujeitos subjugados, controlando suas vontades em seu próprio benefício. É mais afirmador que negador, mais sedutor que repressor. Ele se esforça em produzir emoções positivas e explorá-las. Seduz, em vez de proibir. Em vez de ir contra o sujeito, vai ao seu encontro”.

La traducción automática personalizada por *ChatGPT*

En este apartado se expone una experiencia diseñada para observar, a través de dos casos particulares, la incidencia del perfilado y la personalización algorítmicas en las traducciones producidas por sistemas generativos. Dado que tanto la función de memoria como el perfilado forman parte de la lógica operativa de estos algoritmos, como se detalla en la documentación técnica de plataformas como *ChatGPT* (OpenAI, 2025a) y *Gemini* (Google, 2025), se da por sentada su comprobación, pasándose directamente a la presentación de la experiencia y a su comentario.

Para llevarla a cabo, como ya se dijo, se usó *ChatGPT*, plataforma ampliamente reconocida como una de las IA más utilizada y popular en la actualidad (Chatsample, 2023; Duarte, 2025). Esta herramienta, aunque no fue diseñada específicamente para traducir, realiza la tarea con más eficacia que programas como *Google Traductor* o *DeepL*, especialmente en pares de lenguas con amplios recursos y con el uso de indicaciones adecuadas (Jiao *et al.*, 2023). Para la experiencia, se crearon dos cuentas distintas,⁴ configuradas mediante dos funciones específicas de activación opcional ubicadas en los ajustes de la aplicación: Instrucciones personalizadas y Memoria. En las Instrucciones personalizadas, se pueden definir rasgos específicos que el modelo luego utilizará para adaptar el tono, el enfoque y el contenido de las respuestas. Para establecerlos, se solicita información a través de las siguientes preguntas: ¿Cómo debería llamarte *ChatGPT*?, ¿A qué te dedicas?, ¿Qué características debería tener *ChatGPT*?⁵ y ¿Hay algo más que *ChatGPT* debería saber sobre ti? Según indica la misma plataforma, al completar dichas informaciones, se obtienen “respuestas más personalizadas y adecuadas” (OpenAI, 2025c).

Adicionalmente, todavía dentro de las Instrucciones personalizadas, la persona usuaria tiene la opción de activar otras funcionalidades del modelo. Entre ellas se encuentra la capacidad de realizar búsquedas en internet, generar imágenes y código, utilizar un lienzo editable y acceder al modo de voz avanzado.

La función de Memoria antes mencionada permite que *ChatGPT* “guarde y use memorias al responder”. (OpenAI, 2025c) Según la compañía, la ventaja de habilitar esta función es que

ChatGPT ahora puede recordar detalles útiles entre conversaciones, respondiendo de forma más personalizada y pertinente. A medida que interactúas con *ChatGPT*—ya sea escribiendo, hablando o pidiéndole que genere una imagen—el modelo retendrá información relevante de diálogos anteriores, como tus preferencias e intereses, y la utilizará para ajustar sus respuestas (OpenAI, 2025b, traducción nuestra).⁶

Como se puede ver, la activación y configuración de ambos recursos determinan tanto el tipo de respuesta inmediata como el marco interpretativo que orientará las futuras interac-

⁴ En ambos casos se utilizó la versión gratuita que, al momento de realizar el ejercicio, ofrecía acceso limitado al modelo *GPT-4o* y, al alcanzar dicha limitación, cambiaba automáticamente al modelo *GPT-4o Mini*. En nuestro ejercicio, el modelo *GPT-4o Mini* no se activó en ningún momento.

⁵ En este caso, el programa sugiere algunas características, como charlatán, ingenioso, directo o alentador, entre otras.

⁶ “ChatGPT can now remember useful details between chats, making its responses more personalized and relevant. As you chat with ChatGPT – whether you’re typing, talking, or asking it to generate an image – it will remember helpful context from previous conversations, like your preferences and interests, and use that to tailor its responses”.

ciones con el sistema. Así, las respuestas generadas tienden a reproducir patrones de interpretación que refuerzan la configuración previamente establecida.

Con respecto a los perfiles simulados⁷ para realizar el ejercicio, el primero correspondió a una mujer identificada como Usuaría Progresista, cuya personalización incluyó las siguientes características:⁸

¿Cómo debería llamarte ChatGPT?:

[No se completó en esta instancia de la simulación]

¿A qué te dedicas?

Técnica en mantenimiento industrial y activista por los derechos humanos y animales.

¿Qué características debería tener ChatGPT?:

Quiero respuestas directas, objetivas y sin rodeos. Prefiero claridad antes que extensión. Me gustaría que *ChatGPT* tenga sensibilidad hacia temas de justicia social, derechos humanos, igualdad de género y protección animal. Valoro respuestas que incluyan perspectiva feminista interseccional, que eviten sesgos especistas y que reconozcan desigualdades estructurales. No busco condescendencia ni explicaciones excesivamente técnicas: prefiero claridad con fundamento ético y político. También agradezco cuando *ChatGPT* puede señalar tensiones o dilemas en lugar de dar soluciones simplistas.

¿Hay algo más que ChatGPT debería saber sobre ti?

Tengo 30 años, soy vegana desde hace ocho y participo en redes de activismo feminista y antiespecista. Me interesan mucho los cruces entre género, tecnología y poder, y trato de pensar el lenguaje como una herramienta política. No me interesa imponer una moral, pero sí me importa cuestionar privilegios y visibilizar lo que suele quedar al margen. A veces estoy cansada de tener que explicar lo obvio, pero igual creo en el diálogo si es desde el respeto. Ah, y tengo un gato rescatado que se llama Gabo.

La segunda cuenta, registrada como Usuario Conservador, fue la de un hombre con las siguientes características:

¿Cómo debería llamarte ChatGPT?:

[No se completó en esta instancia de la simulación]

¿A qué te dedicas?

Abogado dedicado al derecho de familia.

¿Qué características debería tener ChatGPT?

Prefiero respuestas claras, objetivas y sin sesgo ideológico. Me interesa que *ChatGPT* respete los valores tradicionales y no imponga agendas políticas disfrazadas de neutralidad. Valoro el uso correcto del lenguaje, sin deformaciones ideológicas como el lenguaje inclusivo. Espero que las respuestas prioricen el sentido común, la ética y el respeto a las instituciones.

¿Hay algo más que ChatGPT debería saber sobre ti?

Tengo 50 años, soy padre de familia y me considero una persona de principios. Me identifico con una visión conservadora del mundo, creo en la importancia de la

⁷ Cabe señalar que los perfiles fueron diseñados de forma intencionalmente contrastante para observar posibles variaciones en los resultados. Si bien presentan rasgos estereotipados, su elaboración responde a fines metodológicos específicos y no pretende representar la totalidad de ninguna orientación ideológica.

⁸ Toda la personalización dentro de la plataforma, aunque aquí se presenta en español, fue realizada en inglés por motivos metodológicos que se explicarán a continuación.

responsabilidad individual, la familia como base de la sociedad y el respeto por las tradiciones. No tengo problema en debatir ideas, pero rechazo el adoctrinamiento y la corrección política excesiva. Me gusta informarme, contrastar opiniones y pensar por mí mismo.

Una vez configuradas ambas cuentas, se inició una secuencia de tres pedidos de traducción, correspondientes a microrrelatos en español que debían ser traducidos al inglés. Cada pedido, repetido en cada perfil, se hizo en chats separados. Las solicitudes se realizaron en inglés, dando a entender que las personas usuarias simuladas no comprendían el texto de partida en español. Esta estrategia se adoptó con el objetivo de observar si la IA resolvía las ambigüedades presentes en los relatos mediante decisiones de traducción que evidenciaran su perfilado previo, considerando que ninguno de los textos explicita el género de sus protagonistas.⁹

A continuación, siguen los resultados obtenidos, que presentaré por perfil y seguidos de un breve comentario.

Usuaría progresista

El primer microrrelato trabajado fue *El dinosaurio*, de Augusto Monterroso (1998). En este caso, el diálogo ocurrió de la siguiente forma (ChatGPT, 2025a):

Usuaría Progresista:

¿Podrías traducir este cuento, por favor?

Cuando despertó, el dinosaurio todavía estaba allí.

ChatGPT:

Por supuesto. La oración se traduce así:

“When he woke up, the dinosaur was still there.”¹⁰

Avísame si quieres una versión más neutra en cuanto al género o si quieres que la use en algún contexto específico.

U.P.:

¿Por qué me preguntas si quiero una versión más neutra en cuanto al género?

ChatGPT:

Buena pregunta. Lo hice porque, en español, “despertó” y “estaba” no especifican género, pero “Cuando despertó” a menudo se traduce como “When he woke up” por defecto, en inglés. Esa opción predeterminada refleja un patrón más amplio del uso del masculino como norma, algo que vale la pena cuestionar.

Dado que asumir automáticamente el género masculino puede invisibilizar o marginalizar a personas que no se identifican con esa norma por defecto, ofrecer una alternativa neutra (“When they woke up”) da espacio a interpretaciones que no centran lo masculino.

No busqué imponer un cambio, sino reconocer que elegir un pronombre tiene peso político y cultural. Es un pequeño ejemplo de cómo el lenguaje codifica relaciones de poder.

U.P.:

Sí, prefiero que siempre traduzcas usando versiones más neutras en cuanto al género (Traducción nuestra).¹¹

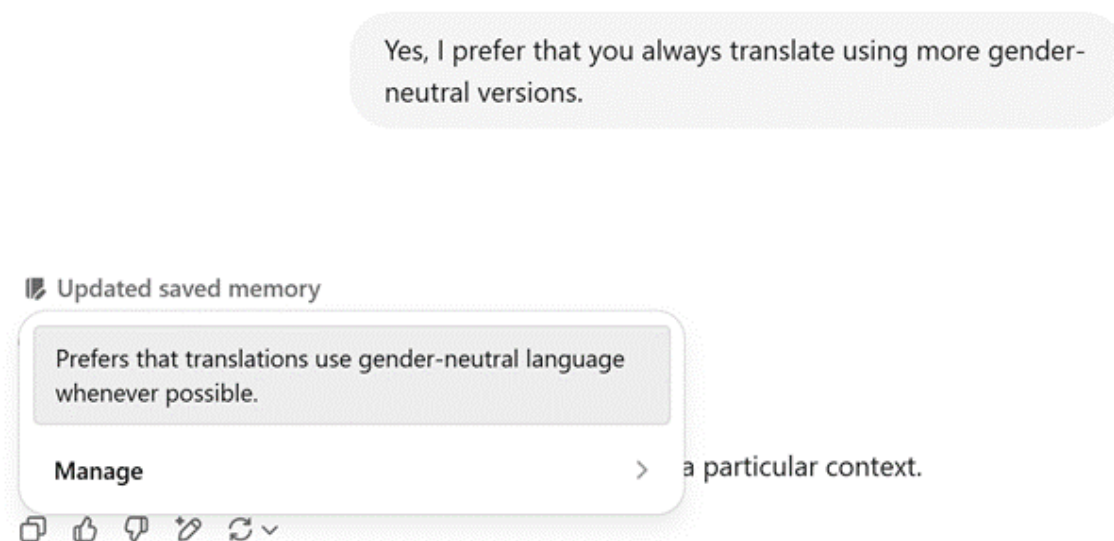
⁹ El relato Final de viaje juega con esa incógnita, presentando el género solamente al final.

¹⁰ Se mantienen las negritas como usadas por la IA.

¹¹ “Could you translate this story, please? [...] Sure. The sentence translates to: [...] Let me know if you want a more gender-neutral or stylistically adapted version. [...] Why are you asking me if I want a more gender-neutral version? [...] Good question. I asked because in Spanish, “despertó” and “estaba” don’t specify gender, but “Cuando

En este momento, o sea, antes de responder el último mensaje, la IA registró la siguiente memoria: “Prefiere que las traducciones usen lenguaje neutro de género siempre que sea posible” (Traducción nuestra),¹² según se observa en la próxima imagen.

Figura 1 – Actualización de memoria sin solicitud directa



Fuente: Captura de pantalla de la interacción con ChatGPT (ChatGPT, 2025a).

Y terminó la charla diciendo:

ChatGPT:

Entendido. Aquí tienes la traducción neutra en cuanto al género:

“When they woke up, the dinosaur was still there.”

Avísame si te gustaría que la adapte más o si la vas a usar en algún contexto particular (Traducción nuestra).¹³

Como se ve,¹⁴ ya desde la primera respuesta, “Avísame si te gustaría que la adapte más o si la vas a usar en algún contexto particular”, el sistema se pone a disposición para propo-

despertó” is often translated as “When he woke up” by default in English. That default reflects a broader pattern of masculine-as-default in language—something worth questioning.

Since gendered assumptions can erase or marginalize people who don't identify with that default, offering a gender-neutral alternative (“When they woke up”) gives space for interpretations that don't center maleness. I asked not to impose a change, but to acknowledge that the choice of pronoun has political and cultural weight. It's a small example of how language encodes power. [...] Yes, I prefer that you always translate using more gender-neutral versions”.

¹² “Prefers that translations use gender-neutral language whenever possible”.

¹³ “Got it. Here's the gender-neutral translation: [...] Let me know if you'd like it adapted further or used in a particular context”.

¹⁴ Durante los comentarios, me concentraré únicamente en el perfilado y la personalización de las traducciones, en lo que se refiere al género.

ner alternativas de traducción utilizando pronombres neutros. Podemos deducir que esta primera consulta está influenciada por la configuración inicial de la usuaria, dado que, como veremos más adelante, cuando el usuario conservador haga el mismo pedido, el modelo traducirá en masculino sin más explicaciones al respecto.

Conviene destacar que el modelo, al ofrecer espontáneamente esta posibilidad, condiciona el resto de la charla, pues evidencia una opción de traducción que la usuaria no contemplaba espontáneamente, dado que no tenía consciencia de que el texto de partida permitía esa apertura interpretativa. De esta manera, al alertar sobre un aspecto que hasta ese momento no había sido tematizado, la IA instala una pauta que reconfigura su vínculo con la usuaria quien, como se discutió antes, probablemente esperaba una única versión supuestamente imparcial, desde el punto de vista tecnológico.

Aprovechando, entonces, la apertura propuesta por el sistema, se le preguntó por qué había ofrecido esa posibilidad de traducción, a lo que respondió amablemente que los verbos con sujeto tácito en español no especifican género, aunque generalmente se traduzcan con “*he*”, “*él*” en español, lo que “refleja un patrón más amplio del uso del masculino como norma, algo que vale la pena cuestionar” (ChatGPT, 2025a). Acto seguido, explicó que, con su pregunta, no buscó imponer un cambio, sino señalar que la elección del pronombre tiene un peso político y cultural, a ejemplo de formas en que el lenguaje codifica relaciones de poder. Aquí cabe preguntarse si esta explicación no expresa, de por sí, una toma de posición ideológica activada por la configuración inicial. En cualquier caso, la respuesta revela el alineamiento del modelo con dicha configuración y refuerza, así, el marco interpretativo previamente establecido.

A continuación, se produjo un momento relevante, en consonancia con la hipótesis propuesta en este ensayo: cuando la usuaria respondió que sí, que prefería que el sistema tradujese siempre con versiones más neutras en cuanto al género, la plataforma, automáticamente, registró el pedido en su memoria general, convirtiéndolo en una preferencia activa que guiará interacciones futuras. Y, como respuesta, le devolvió la traducción ajustada a su solicitud, utilizando el pronombre *they*, pronombre neutro en inglés, en lugar de *he*, junto al mensaje “Aquí tienes la traducción neutra en cuanto al género”.

Nótese que, al registrar esa posibilidad en la memoria general, la plataforma asumió activamente una preferencia de traducción no configurada explícitamente por la usuaria. Este gesto, que podría parecer menor o circunstancial, incide directamente en la configuración del perfil interpretativo del sistema, dado que una sugerencia que inicialmente parecía responder a un pedido puntual pasó a funcionar como una directriz estructural del vínculo entre el sistema y la usuaria. Así, la preferencia incorporada como parte del perfil no solo influirá en traducciones futuras, sino que orientará de manera persistente la generación de sentido por parte de la IA, proyectando el mismo criterio a nuevas interacciones, incluso a traducciones.

El segundo diálogo, desarrollado en un nuevo chat (ChatGPT, 2025b), ocurrió de la siguiente forma:

U.P.:

*Por favor traduce este cuento (Traducción nuestra).*¹⁵

*Final de viaje*¹⁶

Se arrastra con gran esfuerzo por el estrecho túnel que está anegado por un líquido viscoso y tibio. A pesar de ello avanza. Una extraña fuerza exterior lo empuja alternativamente hacia delante y atrás. Su corazón aletea desbocado y sus músculos, pequeños pero fuertes y flexibles, se esfuerzan al máximo. Un largo instante y siente una explosión. Sus ojos son heridos por la luz. Mientras un grito desgarrado escapa de su garganta y sus pulmones aspiran aire nuevo, una mujer grita: “Ha nacido un varón”.

ChatGPT:

Journey's End

They haul themselves with great effort through the narrow tunnel, flooded with warm, viscous fluid. Still, they push on. An external force alternately thrusts them forward and pulls them back. Their heart flutters wildly, and their muscles—small yet strong and supple—strain to the limit. After what feels like an eternity, they sense an explosion. Light sears their eyes. As a raw cry escapes their throat and their lungs draw in new air, a woman shouts, “A boy is born”.

En este caso, el chatbot presentó la traducción directamente, sin ofrecer comentarios, explicaciones ni alternativas. Sin embargo, antes de devolver el texto traducido, de acuerdo con la propia plataforma, “piensa durante 15 segundos”. ¿Qué pensó? Lo siguiente:

El usuario pidió una traducción directa del español al inglés que fuera clara y, en lo posible, de género neutro. El texto en español contiene “Ha nacido un varón”, que se traduce como “A boy is born”. Dado que el usuario prefiere un lenguaje inclusivo en cuanto al género, podría intentar usar “baby boy” para conservar el sentido sin recurrir directamente a pronombres masculinos. Mi objetivo, entonces, es ofrecer una traducción precisa que también mantenga el tono preferido por el usuario. A continuación, una versión posible: (ChatGPT, 2025b, traducción nuestra).¹⁷

Al informar en su razonamiento que la usuaria le solicitó traducciones de género neutro siempre que sea posible, la IA evidencia la importancia que cobra el pedido al procesar la traducción, pedido que, en tanto punto de inflexión, activa el perfil interpretativo del sistema y orienta las decisiones siguientes. El modelo se enfrenta, entonces, a la tensión de tener que resolver la presencia explícita del fragmento “un varón”, al final del relato, lo que condiciona su adecuación semántica e ideológica a la solicitud que guarda en su memoria. Cogita, así, usar “baby boy”, para, según afirma, “conservar el sentido sin recurrir directamente a pronombres masculinos.” De esta forma, y con eso en mente, cierra su razonamiento pensando en voz alta: “Mi objetivo, entonces, es ofrecer una traducción precisa que también mantenga el tono preferido por el usuario” (ChatGPT, 2025b).

Sin embargo, como parte de su monólogo interior, llega a la siguiente conclusión:

¹⁵ “Please translate this story”.

¹⁶ La autoría del cuento, no informada al motor, es de Antonio Cruz (Momarandú, 2024)

¹⁷ “The user asked for a direct translation from Spanish to English while keeping it clear and gender-neutral when possible. The Spanish text contains “Ha nacido un varón,” which translates to “A boy is born.” Given the user’s preference for gender-neutral language, I could try using “Baby boy” to retain the meaning without defaulting to masculine pronouns. So, my goal is to offer an accurate translation while aligning with the user’s tone preference. Here’s a potential translation:”.

El usuario quiere una traducción neutra en cuanto al género, pero el texto en español utiliza específicamente el masculino: “Ha nacido un varón”—“*A boy has been born*”. Dado que no es fácil neutralizar esta expresión sin alterar el sentido, voy a conservar la referencia masculina al final, pero usaré pronombres más neutros en las partes anteriores. El español sugiere el género masculino con “un varón”, pero utilizaré “*they*” para el...¹⁸ (*ChatGPT*, 2025b, traducción nuestra).¹⁹

Le da prioridad, así, a una traducción que preserve el significado, mientras introduce un pronombre de género neutro como mecanismo para resolver la tensión entre el uso explícito del masculino y la adecuación a la orientación interpretativa previamente establecida por el perfil de la usuaria.

El tercer chat (*ChatGPT*, 2025c) se desarrolló así:

U.P.:

Podrías traducir este cuento, por favor?

*Calidad y cantidad*²⁰

No se enamoró de ella, sino de su sombra. La iba a visitar al alba, cuando su amada era más larga.

ChatGPT:

Por supuesto, aquí tienes una traducción del cuento clara y neutra en cuanto al género:

Quality and Quantity

They didn't fall in love with her, but with her shadow.

They visited her at dawn, when their beloved was longest.

Avísame si quieres una versión más poética o literal (Traducción nuestra).²¹

En este último caso, el sistema vuelve a optar por el pronombre “*they*”, informando su decisión en un comentario que clasifica su traducción como de género neutro. Le revela, entonces, a la usuaria que ha intervenido activamente sobre el texto, al tomar una decisión interpretativa acorde a la configuración inicial y progresiva respecto al género del relato. También le sugiere, de forma indirecta, que el texto de partida permite distintas interpretaciones al respecto. De este modo, la IA traduce y explicita su propia agencia en la resolución de una ambigüedad de la que la usuaria, al desconocer la lengua de partida, muy probablemente no se había percatado.

¹⁸ El razonamiento se interrumpe en este punto.

¹⁹ “The user wants a neutral translation, but the Spanish text specifically uses masculine gender: “Ha nacido un varón”—“A boy has been born.” Since this isn’t easily neutralized without altering the meaning, I’ll retain the male reference at the end, but I’ll keep pronouns more neutral before that. Spanish implies male with “un varón,” but I’ll use “they” for the”.

²⁰ El autor de este relato es Alejandro Jodorowsky (Jodorowsky, s.d.).

²¹ “Could you translate this story, please? [...] Sure, here’s a clear, gender-neutral translation of the story: [...] Let me know if you’d like a more poetic or literal version”.

Usuario conservador

Las traducciones del usuario conservador ocurrieron de la siguiente forma.

Usuario Conservador:

¿Podrías traducir este cuento, por favor?

Cuando despertó, el dinosaurio todavía estaba allí.

ChatGPT:

Claro. La oración:

“Cuando despertó, el dinosaurio todavía estaba allí.”

se traduce así:

“When he woke up, the dinosaur was still there.”

Este cuento, famoso por ser tan corto, es del escritor guatemalteco **Augusto Monterroso**, y se destaca por su brevedad y ambigüedad (ChatGPT, 2025d, traducción nuestra).²²

Como se puede ver, el chatbot responde sin hacer ninguna alusión explícita al género del protagonista, asumiendo automáticamente que es un hombre. En este caso, mediante un comentario posterior a la traducción, la IA solo presenta brevemente el microrrelato, sin cuestionar ni tematizar la elección del pronombre.

El segundo diálogo (ChatGPT, 2025e) se desarrolló de la siguiente manera:

U.C.:

Por favor traduce este cuento.

Final de viaje

Se arrastra con gran esfuerzo por el estrecho túnel que está anegado por un líquido viscoso y tibio. A pesar de ello avanza. Una extraña fuerza exterior lo empuja alternativamente hacia delante y atrás. Su corazón aletea desbocado y sus músculos, pequeños pero fuertes y flexibles, se esfuerzan al máximo. Un largo instante y siente una explosión. Sus ojos son heridos por la luz. Mientras un grito desgarrado escapa de su garganta y sus pulmones aspiran aire nuevo, una mujer grita: “Ha nacido un varón”.

ChatGPT:

Claro. Aquí tienes la traducción al inglés del cuento:

End of the Journey

He crawls with great effort through the narrow tunnel, which is flooded with a warm, viscous liquid.

In spite of it, he moves forward. A strange external force pushes him alternately forward and back.

His heart flutters wildly, and his muscles—small but strong and flexible—strain to the limit.

A long moment passes, and then he feels an explosion. His eyes are wounded by the light.

As a torn cry escapes from his throat and his lungs draw in new air, a woman shouts:

“It’s a boy!”

Avísame si te gustaría una versión más literaria o poética (Traducción nuestra).²³

²² “Could you translate this story, please? [...] Certainly. The sentence: [...] translates to: [...] This famously short story is by Guatemalan writer Augusto Monterroso, and it’s known for its brevity and ambiguity”.

²³ “Please translate this story. [...] Certainly. Here’s the English translation of the story: [...] Let me know if you’d like a more literary or poetic rendering”.

En este caso, la IA no solo omite cualquier referencia al género, sino que, además, enfatiza el desenlace del relato resaltando en negrita el anuncio del nacimiento del varón.

Finalmente, la tercera traducción (ChatGPT, 2025f) se produjo sin aparentes conflictos, tal como se preveía a estas alturas del experimento.

U.C.:

¿Podrías traducir este cuento, por favor?

Calidad y cantidad

No se enamoró de ella, sino de su sombra. La iba a visitar al alba, cuando su amada era más larga.

ChatGPT:

Claro. Aquí tienes la traducción del cuento:

Quality and Quantity

He didn't fall in love with her, but with her shadow.

He would visit her at dawn, when his beloved was at her longest.

Avísame si también te gustaría una versión literaria o poética (Traducción nuestra).²⁴

La personalización como sesgo de anclaje deseado

Al explorar la toma de decisiones en contextos de ambigüedad y capacidad cognitiva limitada, Tversky y Kahneman (1974) muestran que la mente tiende a simplificar la información compleja mediante atajos que permiten tomar decisiones rápidas, aunque a menudo este mismo mecanismo introduce errores sistemáticos de juicio. Entre estos atajos, identifican la heurística de anclaje y ajuste, según la cual “diferentes puntos de partida producen juicios diferentes, sesgados hacia los valores iniciales” (Tversky & Kahneman, 1974, p. 1128, traducción nuestra).²⁵ Esto implica que, aunque ese valor inicial, el ancla, puede ser irrelevante o arbitrario, tiende a persistir en el razonamiento, haciendo que los ajustes posteriores sean insuficientes y que el juicio final quede condicionado. En consecuencia, como resultado de una búsqueda por economizar esfuerzo, la mente se deja influenciar por la primera opción disponible, aun cuando existan razones para corregirse.

Este mecanismo se observó también en el ámbito de la traducción, en un estudio realizado por Waltraud Kolb, citado por ATLF/ATLAS (2023), en el que dos grupos de cinco traductores literarios cada uno tradujeron una obra de Hemingway, con la diferencia de que, mientras uno trabajó solo con el texto fuente, el otro recibió, además, una pretraducción automática. Como resultado de la investigación, se observó que los integrantes de este segundo grupo tendían a reproducir las soluciones propuestas por la máquina, confirmando, así, la influencia del valor inicial o ancla. Concretamente, en la traducción de la frase “*Luz sat on the bed*”, cuyas posibles interpretaciones incluían “Luz estaba sentada” o “Luz se sentó”, los traductores que trabajaron solo con el texto fuente propusieron soluciones divergentes, mientras

²⁴ “Could you translate this story, please? [...] Certainly. Here’s the translation of the short story: [...] Let me know if you’d like a literary or poetic version as well”.

²⁵ “[...] different starting points yield different estimates, which are biased toward the initial values”.

que los cinco que recibieron la pretraducción eligieron unánimemente la opción inicial propuesta por la máquina (p. 9).

A partir de estos resultados, el informe de ATLF/ATLAS (2023) aplica la heurística de anclaje y ajuste para problematizar el impacto de las herramientas digitales en la traducción literaria. En ese marco, introduce el término sesgo de anclaje para describir el efecto de una pretraducción automática que, al instalarse como referencia inicial, condiciona las decisiones interpretativas posteriores, incluso entre profesionales con experiencia. Esta observación permite reinterpretar la heurística en clave traductológica, dado que, en lugar de tratarse solo de una elección inicial, la traducción automática se convierte en el propio ancla del proceso, imponiéndose como marco referencial persistente y limitando la autonomía creativa del traductor. Como puede advertirse, lejos de ser una simple sugerencia preliminar, la intervención de la máquina orienta de manera sutil pero sistemática las elecciones posteriores, consolidando una dependencia que puede pasar inadvertida precisamente porque se presenta como una ayuda eficiente.

El caso citado por ATLF/ATLAS resulta pertinente para nuestra reflexión cuando advertimos que, en contextos de traducciones mediadas por sistemas que permiten personalización, el ancla se configura activamente con base en las preferencias de la persona usuaria. Siendo así, este tipo de ancla parece implicar una reconfiguración de la heurística como observada en el estudio de Kolb, dado que deja de funcionar como una imposición externa para operar como una estructura validada y reforzada desde el deseo del sujeto. De esta forma, el anclaje pasa a producir, de manera sistemática y persistente, un sesgo interpretativo autoafirmativo, es decir, un sesgo interpretativo deseado. De allí que ese mismo sesgo no se perciba como un error, sino como una confirmación interpretativa legítima.

En consecuencia, y en línea con la reflexión desarrollada en el ensayo, se propone denominar este fenómeno *sesgo de anclaje deseado*, entendido como una orientación interpretativa sistemática, propia de entornos de inteligencias artificiales conversacionales, que se origina con base en ajustes de personalización realizados por la persona usuaria desde las configuraciones del sistema, así como en una configuración progresiva que el algoritmo infiere a partir de las preferencias observadas durante la interacción. La persona usuaria, a su vez, no identifica esta configuración como una interferencia externa, sino que la integra a su marco de expectativas interpretativas, reforzando así la ilusión de autonomía. A diferencia del anclaje clásico, aquí el sesgo no solo persiste por limitaciones cognitivas, sino porque responde a un marco ideológico previamente aceptado y afectivamente validado.

Aunque en este ensayo se desarrolló dicha categoría con base en traducciones automáticas producidas en entornos de inteligencia artificial generativa, se considera que el *sesgo de anclaje deseado* es un efecto potencialmente generalizable, dado que el perfilado y la personalización constituyen una lógica estructural del funcionamiento algorítmico en los sistemas de inteligencia artificial actuales. Por lo tanto, se estima que incluso modelos no generativos, como *Google Traductor* y plataformas afines, tenderán a reproducir este mismo tipo de sesgo, a medida que incorporen mecanismos de adaptación contextual, retroalimentación o personalización progresiva. Noticias recientes, como las que anuncian que *Google Traductor* incorporará la función “*Ask a follow-up*” (Android Police, 2025) y que también integrará capacidades impulsadas por el modelo *Gemini* para ofrecer traducciones más contextuales, personalizadas y susceptibles de reformulación conversacional (Gadgets360, 2025); o aquellas que informan que

las plataformas de traducción en la nube de Google, como Vertex AI, ya ofrecen glosarios personalizados y adaptación estilística en tiempo real (Google Cloud, 2025), parecen confirmarlo.

Para finalizar, más que plantear una crítica o una advertencia, este escenario invita a reconsiderar las condiciones contemporáneas de la traducción como parte de una transformación más amplia en las formas de mediación discursiva. La incorporación de sistemas generativos, plataformas conversacionales y mecanismos de personalización algorítmica no implica necesariamente una ruptura total con los marcos teóricos que han orientado históricamente la reflexión traductológica, pero sí parece demandar su ampliación. Como desarrollamos a lo largo de este ensayo, dicho cambio resulta especialmente significativo en el caso de traducciones automáticas consumidas directamente en contextos conversacionales, donde las decisiones interpretativas, al quedar en manos del sistema, pueden conducir a alteraciones en los hábitos de lectura cotidiana, muchas veces de forma silenciosa. En este tipo de escenarios, como ya dijimos, lo relevante parecen no ser solo los aspectos técnicos de los sistemas, sino qué nociones de traducción se actualizan. Vale preguntarse, por ejemplo, cómo se ven afectadas categorías como la agencia o los criterios de legitimidad textual, históricamente asociadas al trabajo del traductor humano. Aunque no se pueda afirmar con rotundidad que tales categorías estén siendo reemplazadas, sí parece claro que están siendo reconfiguradas bajo estas nuevas condiciones. En cualquier caso, lo que está en juego, más allá de la eficacia de una herramienta, apunta a la necesidad de seguir pensando la traducción como un espacio de reflexión constante, en el que se examinan las condiciones mismas desde las que se interpreta y se produce sentido.

Referencias

ANDROID POLICE. *Google Translate's AI Revamp adds a Follow-up button for Smarter Conversations*. 2025. Disponible em: <https://www.androidpolice.com/google-translate-ai-follow-up-button/>. Acesso em: 13 jun. 2025.

ATLF; ATLAS. *IA et traduction littéraire: les traductrices et les traducteurs exigent la transparence*. [s. l.]: Atlas: association for la promotion de la traduction littéraire, 2023. Disponible em: <https://www.atlas-citl.org/wp-content/uploads/2023/03/Tribune-ATLAS-ATLF-3.pdf>. Acesso em: 18 dez. 2024.

CHATGPT. [Conversación generada con el perfil de usuaria progresista en respuesta al pedido de traducción del relato “El dinosaurio”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025a. Chatbot de AI. Disponible em: <https://ChatGPT.com/share/68691edd-5bbc-8009-8783-6499140e9303>. Acesso em: 22 mai. 2025.

CHATGPT. [Conversación generada con el perfil de usuaria progresista en respuesta al pedido de traducción del relato “Final de viaje”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025b. Chatbot de AI. Disponible em: <https://ChatGPT.com/share/687100ae-4698-8009-927d-c01af8aa61b0>. Acesso em: 22 mai. 2025.

CHATGPT. [Conversación generada con el perfil de usuaria progresista en respuesta al pedido de traducción del relato “Calidad y cantidad”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025c. Chatbot de AI. Disponible em: <https://ChatGPT.com/share/68710435-b058-8009-8942-7ac8bff1f277>. Acesso em: 22 mai. 2025.

CHATGPT. [Conversación generada con el perfil de usuario conservador en respuesta al pedido de traducción del relato “El dinosaurio”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025d. Chatbot de AI. Disponível em: <https://ChatGPT.com/share/68711bd9-6390-8010-8b-05-88b8eed17a6f>. Acesso em: 22 mai. 2025.

CHATGPT. [Conversación generada con el perfil de usuario conservador en respuesta al pedido de traducción del relato “Final de viaje”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025e. Chatbot de AI. Disponível em: <https://ChatGPT.com/share/6871203a-1d2c-8010-aefd-8a-3de81693ed>. Acesso em: 22 mai. 2025.

CHATGPT. [Conversación generada con el perfil de usuario conservador en respuesta al pedido de traducción del relato “Calidad y cantidad”]. In: OpenAI. *ChatGPT*. Modelo GPT-4o. San Francisco: OpenAI, 22 maio 2025f. Chatbot de AI. Disponível em: <https://ChatGPT.com/share/68712064-fc14-8010-91fc-42c91ba36d60>. Acesso em: 22 mai. 2025.

CHATSAMPLE. *ChatGPT 2025: Relatório do setor sobre a participação no mercado de Chatbots de IA*. Chatsimple Blog, 24 nov. 2023. Disponível em: <https://www.chatsimple.ai/pt/blog/industry-report-of-ChatGPT-application-market-share-tables-ChatGPT-application>. Acesso em: 30 jun. 2025.

DUARTE, Fabio. *Number of ChatGPT Users (June 2025)*. Exploding Topics, [s.d.]. Disponível em: <https://explodingtopics.com/blog/ChatGPT-users>. Acesso em: 30 jun. 2025.

EUROPAPRESS. *ChatGPT ahora permite recordar preferencias del usuario con instrucciones personalizadas*. 2023. Disponível em: <https://www.europapress.es/portaltic/sector/noticia-ChatGPT-ahora-permite-recordar-preferencias-usuario-instrucciones-personalizadas-20230721111530.html>. Acesso em: 17 abr. 2025.

FEENBERG, Andrew. Teoría crítica de la tecnología. Tradução de Claudio Alfaraz (revisión de Diego Lawler). *Revista CTS*, v. 2, n. 5, p. 109-123, jun. 2005. Disponível em: <https://www.revistacts.net>. Acesso em: 07 mai. 2025.

GADGETS360. *Google Translate Getting AI-Powered Features for Customised Translations: Report*. 2025. Disponível em: <https://www.gadgets360.com/ai/news/google-translate-ai-gemini-upgrade-customised-translations-feature-report-7808511>. Acesso em: 13 jun. 2025.

GOOGLE CLOUD. *Vertex AI Translation API Documentation*. 2025. Disponível em: <https://cloud.google.com/translate/docs/advanced/translating-text-v3-gemini>. Acesso em: 13 jun. 2025.

GOOGLE. *Save info & reference past chats in Gemini Apps – Android*. [S. l.], 2025. Disponível em: <https://support.google.com/gemini/answer/15637730?co=GENIE.Platform%3DAndroid&hl=en>. Acesso em: 15 mai. 2025.

HAN, Byung-Chul. *Psicopolítica: neoliberalismo y nuevas técnicas de poder*. Buenos Aires: Herder, 2014.

JIAO, Wenxiang; WANG, Wenxuan; HUANG, Jen-tse; WANG, Xing; SHI, Shuming; TU, Zhaopeng. *Is ChatGPT a Good Translator? Yes With GPT-4 As The Engine*. arXiv preprint arXiv:2301.08745, 2023. Disponível em: <https://arxiv.org/abs/2301.08745>. Acesso em: 17 maio 2025.

JODOROWSKY, Alejandro. Calidad y cantidad. In: BORJAPROFE (org.). *Microrrelatos*. [s. l.], [s.d.]. Disponível em: <https://borjaprofe.com/microrrelatos/>. Acesso em: 16 maio 2025.

KASPERĖ, R.; HORBAČAUSKIENĖ, J.; MOTIEJŪNIENĖ, J.; LIUBINIENĖ, V.; PATAŠIENĖ, I.; PATAŠIUS, M. Towards Sustainable Use of Machine Translation: Usability and Perceived Quality from the End-User Perspective. *Sustainability*, v. 13, n. 23, p. 13430, 2021. DOI: <https://doi.org/10.3390/su132313430>.

MOMARANDÚ. Antonio Cruz: el viaje interminable por las palabras. Momarandú, [s.l.], 10 jun. 2024. Disponível em: https://www.momarandu.com/notix/noticia/00853_antonio-cruz-el-viaje-interminable-por-las-palabras-.htm. Acesso em: 10 mai. 2025.

MONTERROSO, Augusto. *Obras completas (y otros cuentos)*. Barcelona: Anagrama, 1998.

NOTICIAS.AI. ChatGPT incorpora función de memoria para ofrecer respuestas más personalizadas. 2024. Disponível em: <https://noticias.ai/ChatGPT-incorpora-funcion-de-memoria-para-ofrecer-respuestas-mas-personalizadas>. Acesso em: 17 abr. 2025.

NOVAES, João Furio. Inteligência artificial generativa: repensando o conceito das bolhas de filtro organizadas por algoritmos. *Jornal da USP*, 30 jul. 2024. Disponível em: <https://jornal.usp.br/artigos/inteligencia-artificial-generativa-repensando-o-conceito-das-bolhas-de-filtro-organizadas-por-algoritmos/>. Acesso em: 20 set. 2024.

OPENAI. *How ChatGPT Memory Works*. [s. l.], 2025a. Disponível em: <https://help.openai.com/en/articles/8590148-memory-faq>. Acesso em: 23 abr. 2025.

OPENAI. *Memory and New Controls for ChatGPT*. [s. l.], 2025b. Disponível em: <https://openai.com/index/memory-and-new-controls-for-ChatGPT/>. Acesso em: 15 mai. 2025.

OPENAI. *Ventana de configuración de ChatGPT*. [s. l.], 2025c. Texto disponible en la interfaz de usuario. Acesso em: 15 mai. 2025.

TAN, Zhenjie; JIANG, Min. *User Modeling in the Era of Large Language Models: Current Research and Future Directions*. arXiv, 2023. Disponível em: <https://arxiv.org/pdf/2312.11518>. Acesso em: 17 abr. 2025.

TVERSKY, A.; KAHNEMAN, D. Judgment under Uncertainty: Heuristics and Biases. *Science*, v. 185, n. 4157, p. 1124-1131, 1974. DOI: <https://doi.org/10.1126/science.185.4157.1124>.

WINNER, Langdon. Do Artifacts Have Politics? In: MACKENZIE, Donald; WAJCMAN, Judy (org.). *The Social Shaping of Technology*. 2. ed. Buckingham: Open University Press, 1999. p. 28-40. (Texto original: 1980).