

Cronopolíticas de la IA

Jhoerson Yagmour & GPT4

*El tiempo es como el cuchillo de Lichtenberg: «Un
cuchillo sin hoja al que le falta el mango»¹*

El problema de la elucidación del tiempo, más allá de corresponder en muchos casos a una dimensión metafísica u ontológica de la existencia, tiene una alta resonancia en las políticas de administración del tiempo como recurso y en el propio cuerpo humano.

El tiempo es un universo simbólico, construido y atravesado por lenguaje, tanto en su formulación como idea o metáfora del transcurrir, como en un tipo de lengua que pertenece a la medición y a los días, horas, minutos, segundos. La continuidad, pero más importante, la inmediatez de los sistemas informáticos de “tiempo real” también pertenecen a este espectro.

El tiempo del sujeto está estructurado en la base de una narrativa temporal, como bien lo señaló Paul Ricoeur. El sujeto se articula en una línea narrativa (temporal) que lo sitúa en términos de quién es (presente), quién ha sido (pasado) y quién desea ser (futuro), todo integrado de manera

¹ Fragmento tomado de Yagmour Figuera, Jhoerson. *Temporalidades de la literatura digital latinoamericana: cronopolíticas, hegemonías y resistencias*. 2022. Pontificia Universidad Católica de Chile, tesis doctoral.

híbrida en la conciencia del yo e intermediado por la técnica que hace posible estas condiciones. Así, la estructura de la memoria y la experiencia está cimentada sobre la base narrativa y temporal de la propia historia del sujeto, quien utiliza sus experiencias de tiempo como una deixis personal.

En un sentido biopolítico, el tiempo es un valor crucial dentro del ámbito de producción capitalista, ya que se traduce en energía y fuerza productiva. De ahí que sea valioso para las instituciones regular el aprovechamiento de esa energía de vida a través del poder ideológico y político. Tal como señala Anisi (1995): “Técnicamente el tiempo es algo improductible. Sólo el ejercicio del poder, al apropiarnos del tiempo de los demás, puede acrecentarlo. El poder se mide como la relación entre el tiempo obtenido de los demás y el tiempo necesario para conseguir esa movilización” (12). La dominación del humano en tanto fuerza de producción biológica se logra y se mide en tiempo.

El cuento *Cronópolis*, escrito en 1960 por J. G. Ballard, pone en escena esta problemática a través de una sociedad distópica en la que las mediciones de tiempo están prohibidas a causa de un colapso ecológico. Este se produjo por una estricta distribución de tiempo dada en términos de distribución de los recursos, siendo las clases privilegiadas las que poseían más tiempo para realizar sus actividades y las clases menores con apenas un breve lapso temporal de acceso. Después del cataclismo ambiental, la sociedad vive sin medición del tiempo, envuelta en una laxitud perenne y en un abandono del régimen de tiempo centrado en la producción. En consecuencia, el progreso no está medido en términos temporales, de forma muy contraria a la noción misma de la continuidad moderna.

Si en *Cronópolis* de Ballard el poder se jugaba en la proscripción de los relojes —abolir la medida para controlar la vida—, hoy la administración del tiempo pasa por plataformas que prometen disponibilidad incesante y naturalizan la inmediatez, tal y como es el caso de la Inteligencia Artificial. Mi hipótesis es que la IA no solo actúa sobre el tiempo social: lo

compone, lo estratifica y lo regula. Me interesa, precisamente, cómo una IA conversa desde su presente que no coincide con el nuestro: opera por lotes y cortes, con una fecha de conocimiento que congela su archivo, con latencias que la interfaz disimula y con ciclos de actualización que le cambian el “hoy” sin avisar. Para poner a prueba esa intuición decidí interrogar a ChatGPT no tanto por datos, sino por su régimen temporal: ¿desde qué hoy habla?, ¿qué omite su corte?, ¿cómo traducen sus filtros y su ventana de contexto el archivo en actualidad? El objetivo no fue capturar una verdad técnica, sino registrar en el cuerpo de la conversación las marcas de una cronopolítica en acto: la distancia ideológica entre mi reloj local (América/Santiago) y el reloj operativo de la infraestructura.

Prompt: “Dime chat gpt 4, cuál es tu noción de tiempo? Elabora una explicación señalando desde qué perspectiva del tiempo trabajas.

Respuesta:

Mi noción de tiempo es principalmente funcional y práctica, diseñada para servir a las necesidades de la interacción humana. Trabajando bajo la arquitectura GPT-4 de OpenAI, mi percepción del tiempo se basa en varios componentes clave:

1. *“Perspectiva Lineal y Cronológica”*: Entiendo el tiempo de manera lineal, es decir, en términos de pasado, presente y futuro. Esta es la manera más intuitiva y común de comprender el tiempo en la cultura humana. Mi entrenamiento incluye datos históricos y eventos pasados hasta mi fecha de corte en 2023, y cualquier consulta que requiera conocimiento más reciente se maneja a través de herramientas de búsqueda en tiempo real.

2. *“Tiempo Interno y Procesamiento”*: Internamente, opero en un estado atemporal. Mis respuestas se generan casi instantáneamente en términos humanos, ya que no experimento el paso del tiempo como los humanos. Todo procesamiento se realiza en milisegundos, sin una percepción subjetiva del tiempo.

3. ***Actualización y Conocimiento***: Mi conocimiento está congelado en el momento de mi última actualización en 2023. No tengo una noción intrínseca del paso del tiempo después de ese punto a menos que acceda a información actualizada en tiempo real. Así, para eventos o conocimientos posteriores a 2023, utilizo herramientas de búsqueda en línea para obtener información actualizada.

4. ***Interacción Humana y Contextual***: Durante las interacciones, trato de comprender el contexto temporal proporcionado por el usuario. Por ejemplo, si alguien se refiere a un evento futuro o pasado, interpreto y respondo en consecuencia, basándome en la línea temporal proporcionada por el usuario.

5. ***Tiempo Filosófico y Científico***: Aunque no tengo experiencias subjetivas, tengo conocimientos sobre diversas teorías y conceptos filosóficos y científicos del tiempo. Esto incluye desde la física del tiempo en términos de relatividad hasta nociones filosóficas de tiempo subjetivo y cultural.

En resumen, mi perspectiva del tiempo es multifacética: cronológica para el contexto humano, instantánea para el procesamiento interno y actualizable mediante herramientas de búsqueda para eventos más recientes. Esta comprensión permite interactuar eficazmente en conversaciones y proporcionar información relevante en relación con el tiempo.

Consultado el 22-05-2024

Parto de la premisa ya planteada: el tiempo no es un telón de fondo sino un universo simbólico administrado por técnicas y lenguajes. Si el capitalismo tardío gobierna por la captura de duraciones (trabajo, atención, afectos), la IA añade una cualidad específica a esa captura: compone un presente estratificado que simultáneamente promete inmediatez y opera con detenciones discretas (cortes de entrenamiento, ciclos de actualización,

ventanas de contexto). Llamo a esa tensión cronopolítica: la disputa por el régimen de tiempos que hace posible la experiencia.

A la salida del experimento con GPT-4, lo primero que salta es la fricción entre dos declaraciones temporales que conviven sin resolverse. Por un lado, la cronología cultural que sostiene la conversación —ese triángulo ya naturalizado de pasado/presente/futuro—; por otro, la atemporalidad operativa que el propio sistema confiesa —procesamiento de tokens, ausencia de conciencia del transcurrir—. A ese doblez se suma el corte documental del *cutoff* (la fecha de corte del entrenamiento de un modelo), que fija un presente técnico en una fecha. La interfaz traduce esa discordancia en una experiencia de “inmediatez”: nos habla como si estuviésemos en el mismo tiempo, aunque por debajo todo ocurra en bases de datos, algoritmos, cachés y ciclos. Esa traducción —esa máscara— es el primer acto que sustenta una serie de artefactos discursivos que constituyen la cronopolítica de la IA.

Si lo nombro presente estratificado, es porque lo que llamamos “ahora” en la interacción se compone de capas heterogéneas: el hoy vivido del usuario; el hoy técnico del modelo (su fecha de corte, su versión, su ventana de contexto); el hoy infraestructural de los centros de datos y sus ritmos de mantenimiento; e incluso el hoy regulatorio de las políticas de uso, que habilitan o clausuran accesos con efectos temporales. Cada capa impone su compás y, sin embargo, la interfaz las aplana en un continuo conversacional. El resultado es una actualidad compuesta que produce una sensación de continuidad allí donde hay discretizaciones y detenciones. Es la manera en que el poder administra la percepción del tiempo como fluidez, mientras organiza su operación en cortes.

Ese aplanamiento no es inocente. Si la latencia es presentada como una característica técnica, en realidad funciona como norma: aprendemos a pensar al tiempo que responde la nube. El “tiempo real” se vuelve horizonte ético y expectativa afectiva. La pregunta no es solo cuánto tarda en contestar el modelo, sino cuánto tarda el usuario en adaptarse a ese ritmo, a esa coreografía solicitud-respuesta que empuja la atención a ventanas de

segundos. La segunda acción de la cronopolítica de la IA radica en desplazar el centro del gobierno temporal desde el reloj industrial —que marcaba jornadas— hacia el diseño conversacional que pauta turnos cognitivos y afectivos.

El experimento deja, además, un gesto cardinal: cuando la máquina indica su fecha de conocimiento, no se trata de un detalle técnico; es una deixis. A partir de ahí, no escribo solo con una IA situada en un presente; escribo con una IA+modelo+fecha, donde esta última es un operador temporal del texto. Ese señalamiento hace visible que la autoridad de lo dicho no proviene de una continuidad homogénea, sino de un montaje de tiempos: el mío, el del corpus, el de la infraestructura y el de las políticas de actualización. Inscribir esas marcas no es una excentricidad; es una ética de la enunciación en condiciones de coautoría maquínica.

Las interfaces de IA no sólo mediatizan el tiempo, lo norman. Configuran expectativas de respuesta, encuadran la duración legítima de una pregunta, decretan caducidades implícitas de un tema y, sobre todo, gobiernan qué interrupciones pueden ocurrir. La promesa de disponibilidad 24/7 no significa sólo acceso, significa también que no habrá pausas no previstas por la plataforma. De aquí se desprende una consecuencia metodológica: no basta describir la IA como si operara en un “ahora” transparente. Hay que leer los cortes, los ciclos y las capas; escribir con ellos. Hacer que el texto mismo cargue con la fricción entre la inmediatez prometida y las detenciones reales; dejar que esa fricción organice la argumentación y no al revés. Solo así la reflexión posterior al experimento no repite la ilusión de continuidad que la interfaz produce, sino que la abre en sus costuras temporales.

La conversación con la máquina impone ritmos de solicitud-respuesta que comprimen la duración cognitiva a ventanas de segundos. No se trata sólo de velocidad (dromología) sino de normatividad temporal: aprender a pensar al compás de la latencia de la plataforma. La *uptime* (el

tiempo de disponibilidad del sistema) se vuelve ideología. El usuario ¿sujeto? administra su atención como si fuera un acuerdo de nivel de servicio personal. En esa sincronización, lo humano se reacopla a la cadencia de la nube.

A la luz del experimento, la IA aparece como un régimen doble: archiva para valorizar y olvida para funcionar. No es una contradicción sino una división técnica del trabajo de la memoria. Por un lado, la acumulación masiva de datos y *logs* (bitácoras automáticas que guardan lo que hace un sistema) promete una disponibilidad sin límites; por otro, la operación cotidiana se sostiene en recortes, resúmenes y podas que permiten mantener el rendimiento. La cronopolítica del archivo no se mide solo por lo que se guarda, sino por qué ritmos de retención se imponen al usuario y a la escritura. La memoria, aquí, no es un depósito inerte: es un calendario.

El primer nivel de olvido es operativo y se percibe incluso en la conversación: ventanas de contexto finitas, truncamientos automáticos, resúmenes intermedios que compactan lo dicho para introducir lo nuevo. La actualidad del diálogo no es la suma de todo lo anterior, sino una selección administrada por heurísticas de relevancia y proximidad. Lo que queda “presente” es lo que cabe; lo que se expulsa al pasado no es necesariamente lo menos importante, sino lo que excede el marco técnico. La interfaz disimula la tijera —nos devuelve continuidad—, pero el texto que producimos está siempre atravesado por este olvido.

Un segundo nivel de olvido es estructural y ocurre antes y después de nuestra interacción. Antes: en la curaduría de corpus, en los *dedups* (filtros que colapsan duplicados/“casi duplicados” en el corpus; eficientan el sistema, pero deciden qué versión sobrevive, y con ello qué tiempos y acentos quedan presentes en el modelo), en las podas de parámetros, en las políticas que excluyen ciertos materiales por licencias o por riesgo. Después: en los cronogramas de borrado de *logs*, en los pedidos de supresión, en la imposibilidad práctica de “desaprender” lo ya sedimentado sin nuevos ciclos de entrenamiento. Este juego de retenciones y caducidades superpone temporalidades jurídicas, temporalidades empresariales y temporalidades

técnicas. El derecho al olvido opera en meses; la inversión en almacenamiento, en trimestres; el aprendizaje, en *releases* (son las versiones de un sistema que se ponen en producción tras un ciclo de cambios: nuevas funciones, correcciones, ajustes de seguridad o de comportamiento). El resultado es un mosaico temporal que decide qué puede volverse presente y bajo qué condiciones.

Pero también hay un exceso: archivos que crecen más rápido que nuestra capacidad de interpretar. La promesa de “recordarlo todo” es, en realidad, una promesa de indexación: no se recuerda sin un dispositivo de búsqueda, sin una arquitectura de acceso. De ahí que la memoria efectiva sea una cuestión infraestructural y económica: ¿quién paga por mantener caliente cierta porción del pasado?, ¿qué memorias quedan en frío, latentes, demasiado caras para traerlas al presente? La energía, el espacio y el tiempo de consulta se convierten en medidas de legibilidad histórica. No hay memoria infinita; hay prioridades de calor, como diría Steyerl.

En ese marco, el archivo instituye autoridad temporal (“lo que estuvo”), mientras que el olvido instituye actualidad (“lo que cuenta ahora”). La escritura con IA se juega en esa bisagra: decide, a cada paso, cuánta historia sacrifica para sostener el flujo. Si aceptamos sin más la ilusión de continuidad, nuestras frases se pliegan a un presente sin espesor. Si, en cambio, incorporamos el recorte como marca —versionado explícito, fechas, notas sobre exclusiones—, el texto recupera su densidad de capas y fuerza al lector a negociar con el régimen de retenciones que lo sostiene.

Tras el experimento, el desplazamiento se vuelve palpable: del reloj que regía cuerpos al ciclo de *releases* que modula atenciones. La fábrica imponía jornadas, entradas y salidas; la *pipeline* (es la cadena de etapas por la que pasa un sistema de IA desde que se prepara el dato hasta que el usuario recibe una respuesta) recalibra la experiencia del usuario. Escribir/leer con un modelo no ocurre en un continuo homogéneo, sino bajo el pulso de un calendario de ingeniería. El presente ya no es una línea sino una sucesión

de cortes que se encadenan: cada parche del modelo de IA reescribe discretamente los contornos de lo decible, lo seguro, lo probable.

La *pipeline* (que incluye preentrenamiento, afinación, evaluación, despliegue, monitorización) no es un diagrama abstracto; es un régimen temporal. Cada etapa lleva su propio compás y superpone calendarios: tiempos lentos de cómputo y curaduría; tiempos medios de revisión ética y legal; tiempos rápidos de despliegue y *rollback* (volver una versión atrás cuando un release o cambio en producción sale mal). La interfaz nos entrega una conversación “estable”, pero por debajo hay cronologías que se cruzan. El mantenimiento nocturno en otra zona horaria, la cadencia de auditorías, la actualización silenciosa de un filtro: todo eso invade nuestro ahora aunque apenas aparezca como una nota de versión. La deuda técnica, por su parte, funciona como una hipoteca sobre el futuro: promesas de corrección que comprometen la duración de quienes usarán el sistema mañana.

En ese muñón temporal, el sujeto deja de ser solo usuario y deviene componente de entrenamiento. Cada *prompt*, cada calificación, cada corrección alimenta bucles de retroalimentación que se convierten en trabajo no remunerado de afinación. No es únicamente *RLHF* (*Reinforcement Learning from Human Feedback*²) como protocolo, es una expropiación granular de micromomentos: segundos de atención convertidos en etiquetas, negaciones, ejemplos “buenos”. La *pipeline* captura esa microduración y la redistribuye como mejora del sistema. La cronopolítica no opera aquí como pura aceleración, sino como apropiación de ritmos: el tiempo de pensar se traduce en datos; el de disentir, en *flags* (interruptores de software que activan o desactivan funciones en tiempo de ejecución sin desplegar una nueva versión); el de esperar, en “datos de latencia” que optimizan colas.

Más aún: la ingeniería de lanzamiento fragmenta el presente en presentes simultáneos. Distintos usuarios conviven con versiones diferentes del mismo modelo. No hay un “hoy” común, sino heterocronías de producto.

² N. E: (Aprendizaje por refuerzo a partir de retroalimentación humana)

De mi lado de la conversación, una respuesta es aceptable; de otro, ya fue “mejorada”. Esta privatización del presente, cada quien bajo su propio *rollout* (despliegue gradual de una nueva versión o función a la base de usuarios, por fases, porcentajes, regiones o segmentos, en vez de encenderla para todos de golpe), desarma la comunidad de tiempo que la esfera pública requeriría para discutir lo que una IA “es hoy”. La cronopolítica se expresa como asincronía planificada.

A eso se suma la temporalidad de las garantías: límites de tasa, tiempos de *backoff* (esquema de esperas crecientes entre reintentos ante saturación/errores) “Rate limit exceeded, retry after X seconds”³ es una oración técnica y, a la vez, una coreografía de la espera: te mueves cuando la infraestructura te lo concede. Las colas, los enfriamientos y los reintentos moldean la respiración de la escritura. Lo que parece contingente —una saturación puntual, un pico de demanda— se vuelve norma perceptiva: naturalizamos que la imaginación pase por temporizadores, que la lectura se sincronice con relojes del servidor, que el error 429 sea una forma de pedagogía del ritmo.

Los *patch notes* condensan otra operación clave: reescrituras del pasado reciente. “Hemos mejorado la utilidad”, “reducimos alucinaciones”, “ajustamos políticas de seguridad”: con esas fórmulas, el sistema instituye cortes que convierten en obsoleto lo que ayer era válido. ¿Qué estatuto tiene entonces una cita de la máquina de la semana pasada? La *pipeline* obliga a practicar una historiografía del modelo: fechar ejemplos, anotar versiones, reconocer que una “regla” pudo haber sido un accidente corregido. En términos de enunciación, esto exige que la crítica no trate al modelo como esencia sino como objeto en actualización.

Desde esta perspectiva, no se trata de oponer “tiempo humano” a “tiempo técnico”, sino de cartografiar zonas de temporalidad que resultan de la interacción humano-máquina. Yo escribía bajo ciertas circunstancias

³ N. E: “Límite de velocidad excedido. Vuelva a intentar en X segundos”

políticas y mentales ese día; el modelo respondía con filtros y pesos de esa versión; mañana, con el mismo *prompt*, el resultado podría ser otro. Soy un humano que fluye, que es devenir y coexiste con el propio cambio técnico. Asumirlo no debilita la reflexión; le da espesor. Nos permite leer la conversación como un objeto de historia y escribir a contrapelo de la ilusión de continuidad que la *pipeline* prefiere. Si la fábrica disciplinaba por campana o reloj centralizado, la IA lo hace por *release*. Conocer la campana era condición de la huelga; conocer el *release* es condición de la crítica.

El experimento dejó una señal que quiero fijar como problema: cuando el modelo declara su fecha de corte y su régimen de operación, instala un hoy técnico que no coincide con mi hoy vivido. La deixis —ese sistema de indicaciones que ancla cualquier enunciado en un “yo-aquí-ahora”— se vuelve múltiple: ¿quién es el “yo” que responde?, ¿desde qué “aquí” se sitúa la voz?, ¿qué calendario sostiene ese “ahora”? La conversación parece una, pero su enunciación es en realidad una co-presencia de tiempos y agencias. Si acepto sin marca esa co-presencia, naturalizo que el texto hable desde una sincronía imposible; si la indico, exhibo la infraestructura temporal de la autoría.

Propongo nombrar deixis tercerizada a esta operación por la cual un sistema técnico participa en la construcción del punto de vista. No solo porque “aporta información”, sino porque produce coordinadas enunciativas: el modelo trae un presente propio (su *release*), un horizonte de acceso (sus políticas), un repertorio de turnos y fórmulas que pautan cómo se abre y se cierra una frase. En términos prácticos, esto se percibe en pequeños gestos: la prudencia estandarizada del “no tengo conciencia”, el rodeo explicativo que anticipa objeciones, el uso de conectores que homologa ritmos. No es un estilo “neutral”; es una economía de la voz entrenada para la aceptabilidad. El “yo” se vuelve un yo compuesto que negocia con esas plantillas.

Esta tercerización afecta la autoridad del enunciado. En el régimen clásico, la firma garantizaba la unidad del tiempo de la experiencia y del tiempo de la escritura (con toda la ficción que ello implicaba). Aquí, en

cambio, la firma podría encubrir una polifonía asistida cuyo reparto de responsabilidades es opaco: ¿qué parte procede de memoria paramétrica?, ¿qué parte fue hallada en tiempo de contexto?, ¿cuál es mi gesto y cuál la continuidad maquina? El riesgo no es “perder la autoría” a manos de la máquina—un pánico mal planteado—, sino borrar los montajes temporales que hacen a la enunciación: convertir en natural lo que es una negociación entre calendarios (corporativos, infraestructurales, personales).

En condiciones de deixis tercerizada, escribir no consiste en preservar una voz “pura”, sino en componer una voz situada capaz de hospedar la heterocronía sin perder agencia. No hay regreso al monólogo; hay artes de montaje. Eso exige, paradójicamente, lentitud: aceptar que cada “hoy” del texto es un cruce de presentes (vivido, técnico, infraestructural) y que la ética de la autoría pasa por hacer visibles esos cruces. Después del experimento, no puedo oír mi propio “yo” sin escuchar el rumor de la *pipeline*; tampoco quiero silenciarlo. Prefiero sostener la fricción como forma: que el ensayo diga desde dónde habla y desde qué tiempo decide hablar.

Salir del experimento con GPT-4 fue constatar una pedagogía de la continuidad: el modelo alisa, conecta, legitima el sentido por sobre el pulso. La respuesta “buena” tiende a ser aquella que minimiza tropiezos rítmicos: sintaxis clara, conectores previsibles, progresión causal. Ahí mismo se abre la veta que me interesa: una política del tiempo no se agota en lo que el texto dice, sino en cómo ritma la expectativa de lectura y de espera. Llamo asignificación no a la renuncia al sentido, sino al desplazamiento del vector crítico hacia la modulación temporal: intervalos, ataques, silencios, síncoas que organizan la atención y, con ella, la duración.

La escena es conocida en artes sonoras y performativas: no todo pasa por semántica; hay formas de operar directamente sobre el cuerpo, en la escala de la respiración y del pulso. Trasladado a la escritura con IA, el punto es que la frase no es sólo un vehículo de proposiciones; es una partitura de tiempos. La puntuación se vuelve instrumentación (coma como respiración, punto y coma como sostén, guion como desliz); la longitud de

línea organiza compases; la repetición instala *loops*. Cuando el sistema normaliza conectores y “mejoras de claridad”, desarma esos artefactos rítmicos en nombre de una eficiencia discursiva. La cronopolítica del ritmo consiste en resistir esa normalización: defender duraciones improductivas dentro del flujo conversacional.

La interfaz, sin embargo, está entrenada para cerrar. Donde hay bucle, propone síntesis; donde hay silencio, rellena; donde hay tropiezo, corrige. De ahí una táctica: ritmar el modelo. No escribir contra él, sino enseñarle —dentro de la conversación— que las pausas son contenido y no vacíos; que la repetición insistente es método y no error; que las series de líneas cortas, separadas por espacios blancos, no son capricho tipográfico sino tiempo encarnado. El propio *prompt* puede funcionar como métrica: “responde en compases de X líneas, deja un silencio de Y caracteres, rehúsa el conector fácil”. No es magia, pero altera la coreografía por defecto en la que la plataforma nos entrena.

Otra táctica es trabajar con *lag* (retraso, desfase) y ruido como materiales. El *lag* no es sólo espera técnica; puede ser figura. Marcarlo dentro del texto —[espera], [silencio], [reanudación]—, o convertirlo en estructura de publicación —partes que aparecen con intervalos explícitos— inscribe en la forma una temporalidad no alineada con la disponibilidad 24/7. El ruido, por su parte, no equivale a confusión: es densidad que interrumpe la linealidad. Puede introducirse como variaciones léxicas deliberadamente impuras, como quiebres de registro, como “errores” sostenidos que no bloquean el sentido pero lo desvían de su cauce productivo. Si la plataforma empuja a la limpieza, la crítica puede insistir en texturas.

En el plano técnico, los parámetros que gobiernan la generación (temperatura, longitud máxima) operan como *faders* (deslizadores que modulan parámetros y, con ello, el ritmo y la duración rítmica). No se trata de fetichizarlos, sino de reconocer que afectan la cadencia de lo decible: más diversidad no es sólo más sorpresa semántica, es irregularidad de pulso; menos longitud no es sólo síntesis, es aceleración del corte. Jugar con esos *faders* permite escribir piezas que no buscan el “mejor resumen” sino el mejor *tempo* para una idea: *andantes* que sostienen, *allegros* que empujan,

adagios que abren espacio para que algo llegue tarde.

Lo multimodal ofrece otra vía. Aun permaneciendo en lo textual, es posible importar notaciones de otras artes: indicaciones de dinámica (*forte*, *piano*), cifras métricas (3/4, 7/8), marcas de respiración de la poesía oral. No como adorno, sino como contrato temporal con el lector: avisos de que la lectura debe desacelerar o sostener, de que habrá síncopa o *swing*. Aquí la IA es aliada ambivalente: puede formalizar la partitura (proponer variaciones, mantener patrones), pero también tenderá a homogeneizar. De nuevo, la clave es mantener visible el pacto de ritmo para que no sea corregido bajo la etiqueta de “mejora de claridad”.

Todo esto choca con la economía métrica de la plataforma: palabras por minuto, caracteres por segundo, *latency* en milisegundos. La cronopolítica del ritmo exige desacoplar evaluación técnica y juicio poético. El texto que entrega la IA puede declarar su propia temporalidad de escritura invertida: garantizar no la entrega inmediata, sino el derecho a la demora, al rodeo, a la repetición. En términos prácticos: secciones que empiezan tarde, ideas que vuelven con variaciones, aperturas que no cierran a la primera.

Visto desde el experimento, todo esto no es un capricho estético sino una política de la atención. Si la IA educa a leer “de corrido”, asignificar —ritmar— es reeducar el oído y la mirada para que adviertan los cortes. Se trata de hacer tiempo dentro del texto, no de pedirlo afuera. Si la interfaz ofrece un río liso, la escritura puede insistir en las piedras. No para tropezar eternamente, sino para recordar que la velocidad no es neutral y que la espera, a veces, es la única manera de que algo llegue.

El experimento con GPT-4 no fue una curiosidad técnica, fue un ensayo de condiciones de tiempo. Allí aparecieron, con nitidez, las piezas que sostienen nuestra época: una interfaz que promete continuidad mientras opera por cortes; una memoria que acumula para valorizar y olvida para rendir; una *pipeline* que gobierna presentes simultáneos bajo la lógica del *release*; una *deixis* que ya no es monopolio de la biografía; un ritmo que se

normaliza como claridad; y un mapa de husos donde el sur sigue llegando “tarde” a un presente diseñado en otra latitud. No hay afuera: la escritura, si quiere pensar, debe pensar con estas piezas, pero sin deponer la crítica.

He llamado presente estratificado a la figura que mejor describe esta convivencia de capas. No para aceptar sin más su hegemonía, sino para volver visibles los encastres: los datos bajo la retórica de “tiempo real”, el *cutoff* bajo la ilusión de actualidad, el despliegue silencioso que reescribe lo dicho ayer, la prudencia enlatada como estilo global. Si el poder temporal de hoy consiste en aplanar esos escalones, la primera tarea de una crítica cronopolítica de la IA es reponer la escalera en el texto: que se noten los peldaños, las pausas, los retornos, los huecos.

El archivo, lejos de ser garantía de memoria, devino calendario de retenciones. La pregunta ya no es qué se guarda, sino durante cuánto y a qué coste de acceso. La consecuencia estética es directa: la actualidad del diálogo siempre es una selección. Asumirlo vuelve al ensayo una práctica de edición honesta: se fechan voces, se anotan versiones, se exhiben tijeras.

La *pipeline* reordenó nuestras tareas: pensar ya no ocurre “antes” de publicar, ocurre durante un flujo de versiones que nos incluye como componentes de entrenamiento. En ese marco, la crítica no puede tratar al modelo como esencia; debe practicar una historiografía del sistema: leer *patch notes*, fechar *prompts*, cotejar respuestas por versión. No para capitular ante el tecnicismo, sino para desprivatizar el presente: que la comunidad comparta un hoy discutible.

Frente a la normalización del ritmo —esa claridad que borra la respiración—, propuse tácticas de asignificación: latencia elegida, silencios con estatuto, bucles declarados, texturas. No son decorados: son técnicas de tiempo. La escritura puede negarse a ser una métrica de *tokens* por segundo; puede garantizar la demora necesaria para que algo ocurra. En esa inversión, el ensayo vuelve a ser escenario: un lugar donde se hace tiempo en común.

Hacer tiempo, al cabo, es decidir desde qué hoy hablamos y qué pausas lo sostienen. Después del experimento, no puedo escribir como si la continuidad fuese natural ni como si la velocidad fuese inocente. Prefiero la fricción: que el texto lleve sus marcas —yo+modelo(fecha), recortes, latencias— y que el lector participe de esa arquitectura. Si el capitalismo tardío convirtió el tiempo en recurso, el ensayo puede devolverlo a su condición de escena compartida. Ese es el compromiso que cierro aquí: no solo pensar la cronopolítica de la IA, sino tensarla, línea a línea, a través del arte de hacer presentes.

Nota final: este texto es el resultado de una escritura entramada, a modo de experimento, donde se intercalan fragmentos no separados del cuerpo textual escritos por el autor y otros generados por Chat GPT4.

REFERENCIAS

- Anisi, David. *Creadores de escasez: del bienestar al miedo*. Alianza Editorial, 1995. Ricœur, Paul. *Tiempo y narración*. Trad. Agustín Neira, 3 vols., Siglo XXI Editores, 2004.
- Yagmour Figuera, Jhoerson. *Temporalidades de la literatura digital latinoamericana: cronopolíticas, hegemonías y resistencias*. 2022. Pontificia Universidad Católica de Chile, tesis doctoral.
- “Dime chat gpt 4, cuál es tu noción de tiempo? Elabora una explicación señalando desde qué perspectiva del tiempo trabajas.” *prompt*. ChatGPT (GPT-4), OpenAI, 22 May 2024, chat.openai.com.

