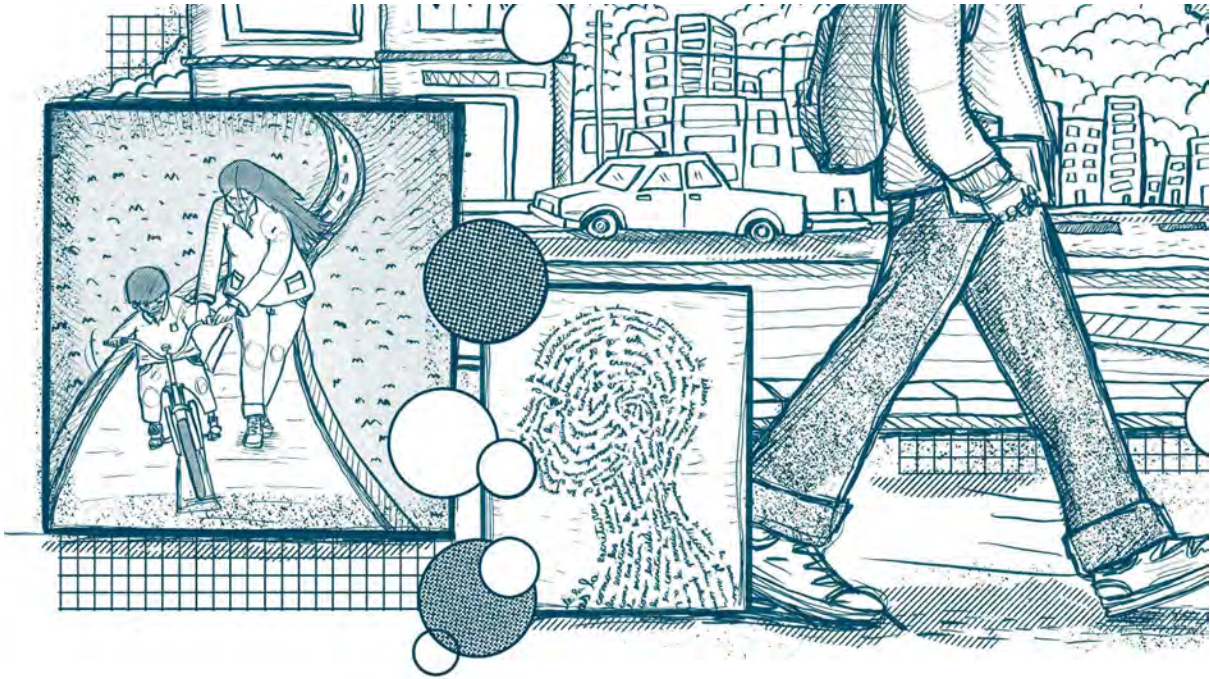




Espacio conceptual

Saber sin estudiar: ¿cómo nos diferenciamos de la inteligencia artificial cuando aprendemos lenguaje?

*Admiróse un portugués
de ver que en su tierna infancia
todos los niños en Francia
supiesen hablar francés.*

Nicolás Fernández de Moratín

En los últimos años, millones de personas han comenzado a interactuar de manera regular con *chatbots* de inteligencia artificial como ChatGPT (lanzado a finales de 2022 y con ya más de cuatro versiones). La característica principal de estos sistemas es que, de manera automática y autónoma, producen texto fluido (e incluso voz) respondiendo a preguntas del usuario de una forma que nos resulta coherente y con sentido.

Sistemas como ChatGPT son ejemplos de lo que se denomina Grandes Modelos de Lenguaje, que abreviamos como LLMs por sus siglas en inglés (*Large Language Models*). Pero ¿qué son exactamente? Los LLMs son modelos entrenados para aprender el lenguaje humano a través de la exposición a volúmenes masivos de datos (libros, artículos, páginas web), y su funcionamiento se basa exclusivamente en la predictibilidad. Su tarea principal es procesar una secuencia de palabras y calcular matemáticamente qué palabra tiene la mayor probabilidad de aparecer a continuación.

Imaginemos que escribimos: *La chica*. El modelo evalúa miles de opciones y les asigna un valor numérico. Un verbo como *es* o un adjetivo como *alta* tendrán mucha más probabilidad de aparecer que otro sustantivo, como *bicicleta*. Para decidir esto, el modelo no usa un diccionario de reglas gramaticales, sino que "recuerda" los patrones que extrajo durante su entrenamiento. Es decir, asigna probabilidades basándose en las millones de veces que vio estructuras similares en otros textos. Por eso se les llama generativos: no repiten frases que ya vieron en otros textos, sino que construyen nuevas respuestas palabra por palabra a partir de esos cálculos estadísticos¹.

La llegada de los LLMs (y de la llamada inteligencia artificial) a nuestra vida cotidiana abre muchas preguntas. Sin embargo, mi objetivo aquí no es analizar sus virtudes o los riesgos que plantean para la sociedad, ni tampoco cómo podrían transformar nuestra forma de pensar, aprender o relacionarnos con la verdad. En cambio, quiero centrarme en una pregunta relacionada pero distinta y que, como lingüista, me interesa particularmente (como a muchos lingüistas antes de mí): ¿en qué medida podemos aprender acerca de la cognición humana a través de los LLMs? ¿Realmente se parecen a nosotros? Es crucial preguntarse si el hecho de que su comportamiento se parezca –al menos superficialmente– al de los seres humanos implica necesariamente que funcionan como nosotros o que compartimos una base cognitiva común.

Para responder a esta pregunta, podemos analizar la relación entre la existencia de comportamiento observable y los mecanismos subyacentes que lo generan. Históricamente, ha existido una brecha entre los ingenieros dedicados a desarrollar la inteligencia artificial y los científicos cognitivos. Mientras los científicos cognitivos queremos entender cuáles son los mecanismos cognitivos (mentales) que subyacen al comportamiento humano, los ingenieros quieren reproducir un comportamiento similar al humano tan bien como se pueda, sin importar mucho si el mecanismo interno subyacente es radicalmente distinto al de una mente humana.

Esta disociación se justificaba bajo la premisa de que un mismo resultado puede ser alcanzado por vías diferentes. Como se ha señalado frecuentemente: los LLMs no tendrían por qué ser más relevantes para la ciencia cognitiva de lo que la ingeniería de submarinos lo es para un ictiólogo. El hecho de que tanto un submarino como un pez se desplacen bajo el agua no implica que el estudio de uno aporte conocimientos especialmente valiosos o profundos sobre el otro (Futrell y Mahowald, 2026).

Pero, en la actualidad, el desempeño de los LLMs (revolucionario, me atrevería a decir) vuelve a poner sobre la mesa una pregunta fundamental: si el comportamiento es tan similar, ¿podemos mantener que los mecanismos subyacentes son tan distintos de los nuestros? Y si existen diferencias mecánicas, ¿podemos verlas manifestarse en alguno de los comportamientos? A continuación, voy a centrarme en las diferencias en cómo aprendemos el lenguaje los seres humanos (niños) y los LLMs.

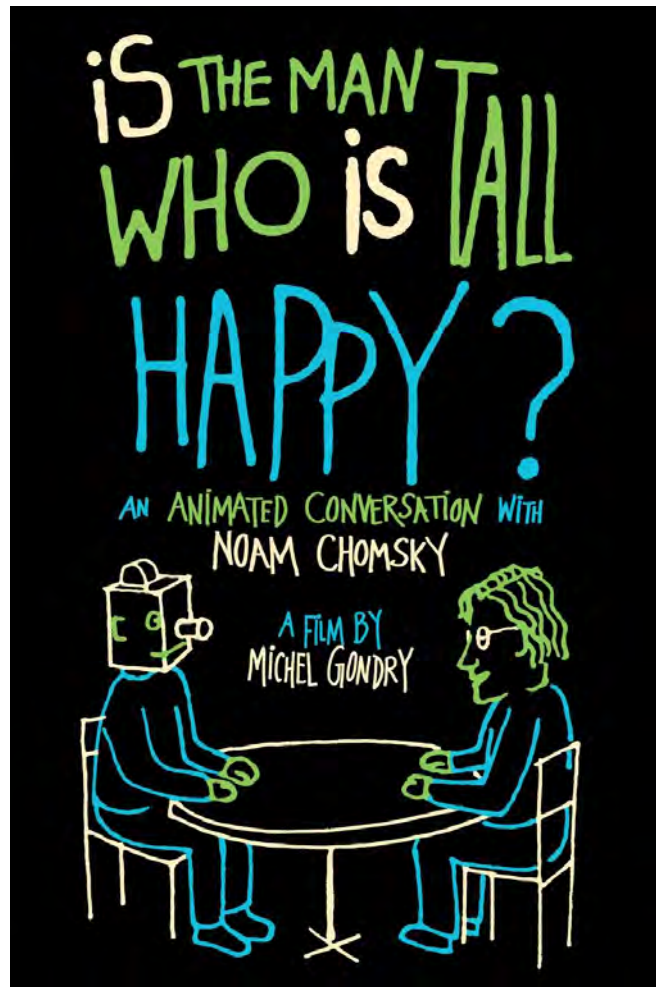
A cualquier persona que haya estado cerca de niños de entre uno y cinco años le fascinará cómo –en un lapso muy corto– pasan de producir unas pocas palabras a armar oraciones complejas. En muchos casos, estas estructuras son reminiscentes del lenguaje adulto, en otros, tienen características muy propias del habla infantil (pero igualmente complejas).

¹ Aunque muchos modelos han incorporado datos no-lingüísticos como imágenes, volviéndose multimodales, todavía la gran mayoría de datos que se usan son lingüísticos.

Aunque los niños carezcan del vocabulario o la experiencia vital de un adulto, usan el lenguaje de manera generativa y original, produciendo oraciones que nunca antes han escuchado.

Una idea muy influyente de la lingüística moderna es que la experiencia lingüística que tiene un niño en sus primeros años de vida es insuficiente para explicar la competencia a la que llega en tan poco tiempo. En otras palabras: el conjunto de palabras, oraciones y enunciados que un niño escucha en su entorno no basta para explicar las reglas que termina infiriendo (para una exposición más detallada y amena de este punto, recomiendo la película de Michel Gondry *Is the man who is tall happy?*). Esta idea es conocida como el “argumento de la pobreza del estímulo” (Chomsky, 1980) y sugiere que los seres humanos poseemos ciertos sesgos cognitivos que restringen el número de hipótesis posibles que podemos inferir de nuestra evidencia: no aprendemos únicamente de los datos, sino que nuestras mentes imponen restricciones sobre cómo procesamos esa información, lo que permite a su vez acotar las hipótesis posibles. Si aprendiéramos solo de los datos, habría ciertas reglas o hipótesis que no podríamos inferir, porque la información no es suficiente para elegir esas reglas o hipótesis frente a muchas otras alternativas².

² Aunque esta idea, en esta formulación particular, no es unánimemente aceptada en lingüística (Tomasello, 2011), suele ser difícil explicar la adquisición del lenguaje de los niños sin asumir la existencia de ciertos sesgos sobre los datos (aunque sean generales y aprendidos), sesgos que los LLMs no parecen tener.



Documental de animación dirigido por Michel Gondry en 2013.

En contraste, el aprendizaje de los LLMs es fundamentalmente inductivo y basado en datos. No poseen sesgos gramaticales predefinidos, sino que su conocimiento deriva directamente de la exposición a volúmenes masivos de texto y la extracción de patrones recurrentes.

Sin embargo, el hecho de que las inteligencias artificiales hayan alcanzado un comportamiento tan similar al humano ha reabierto el debate sobre si los seres humanos realmente necesitamos esos sesgos (Piantadosi, 2023) o si, al igual que los modelos, aprendemos únicamente por asociación estadística. La comunidad lingüística sostiene mayoritariamente que estas diferencias persisten y son profundas, especialmente en la adquisición. A continuación, señalaré tres divergencias clave entre el comportamiento lingüístico humano y el de los LLMs.

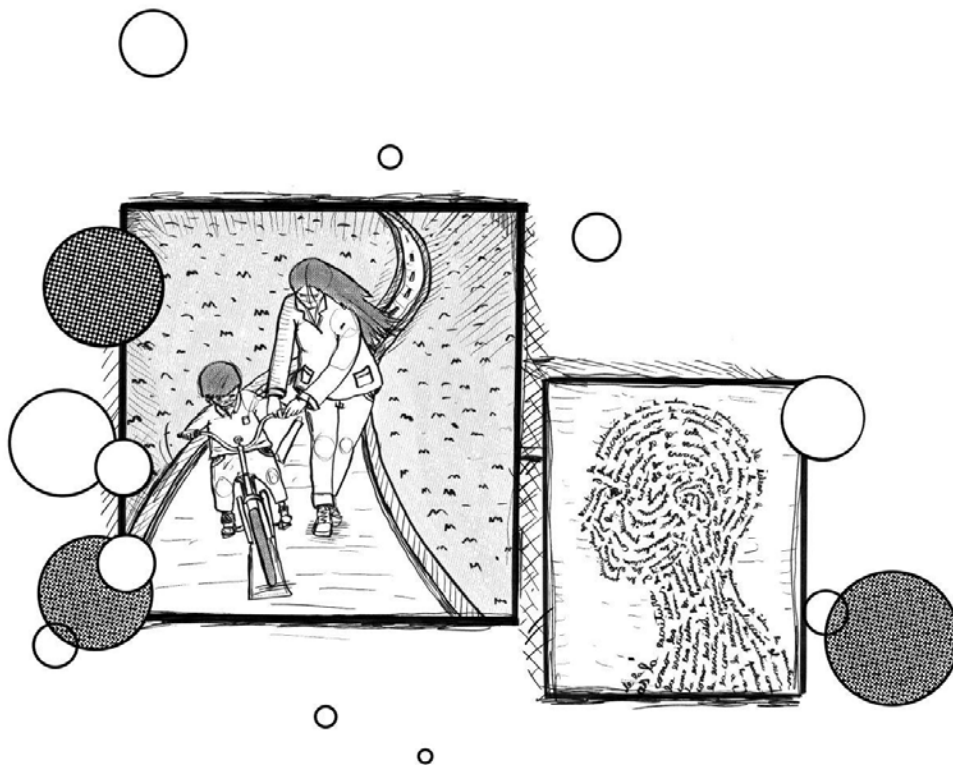
Las divergencias

La primera diferencia tiene que ver con la evidencia lingüística a la que tenemos acceso los seres humanos y los datos a los que acceden los LLMs durante su

entrenamiento. Es decir, ¿qué tipo de experiencia se necesita para aprender lenguaje? Es radicalmente diferente tanto a nivel cuantitativo como cualitativo.

Cuantitativamente, la versión GPT-3 de ChatGPT (OpenAI), por poner un ejemplo, fue entrenada inicialmente con 300.000 millones de tokens (palabras) de texto escrito, proveniente principalmente de Wikipedia y otras fuentes de la web (Brown y otros, 2020). Esto es lo que se utiliza para el *pre-training*, sin tener en cuenta todas las anotaciones hechas por seres humanos necesarias para ajustar el aprendizaje de la inteligencia artificial.

Estas cantidades contrastan con las de un ser humano, que en sus primeros años de vida está expuesto a una escala muchísimo menor. Se estima que un niño de dos años ha escuchado aproximadamente 10.000.000 de palabras (Hart y Risley, 1995). Un ser humano, por lo tanto, necesitaría aproximadamente 30.000 años de vida para estar expuesto a la cantidad de datos que requiere un modelo para desarrollar habilidades lingüísticas similares.



Cualitativamente, los datos utilizados para entrenar LLMs consisten, sobre todo, en texto escrito recolectado de Internet (Wikipedia), los cuales difieren sustancialmente del lenguaje al que están expuestos los niños. Mientras que en Wikipedia la longitud promedio es de 25 palabras por oración, en el corpus CHILDES (MacWhinney, 2000), que registra el habla dirigida a niños (*child-directed speech*) en más de 30 lenguas, el promedio de palabras por oración es de apenas cinco. Pero la diferencia va más allá de la longitud. Gran parte del lenguaje que escuchan los niños corresponde a interacciones dinámicas donde

hay múltiples interlocutores, solapamientos en el habla y ruido ambiental. Además, dependiendo del contexto cultural o socioeconómico, el niño se expone a una enorme variedad de hablantes y registros. En muchas culturas con crianza comunitaria o en casos de escolarización temprana, en la experiencia lingüística no predomina el habla adulta, sino el habla espontánea de otros niños o hermanos, que tiene características muy diferentes al habla adulta (Casillas y otros, 2020; Weisleder y Fernald, 2013).

Estas diferencias cuantitativas y cualitativas entran en tensión. En su adquisición y uso del lenguaje, un ser humano tiene mucha más información disponible que la lingüística, pero los LLMs aprenden fundamentalmente de texto sin tener mucha otra información disponible. Por ejemplo, si reflexionamos en cómo aprendemos el significado de la palabra *manzana* rápidamente pensaremos en que, además de escuchar la palabra, típicamente nos encontramos en situaciones en las que también hay una fruta manzana, un dibujo de una manzana, el olor a manzana o el sabor a manzana, etc. La experiencia cualitativa de los seres humanos es más rica que la de los LLMs, por lo que puede ser que estos necesiten más datos lingüísticos para compensar la falta de multimodalidad (explicando las diferencias cualitativas).

Por otra parte, cuando se ha intentado entrenar a un LLM utilizando exclusivamente datos reales que reflejen lo que un niño escucha en su día a día, los resultados son bastante reveladores: los modelos alcanzan un desempeño absolutamente inferior al de un niño de la misma edad (Warstadt y Bowman, 2022; Huebner y otros, 2021). Esto es particularmente notable en la adquisición de reglas gramaticales, como la formulación de preguntas. Para aprender cómo formular preguntas, a diferencia de lo que ocurriría con el ejemplo de la manzana (donde la información sensorial puede ser clave), los niños no utilizan mucho más que su experiencia estrictamente lingüística, por lo que estos casos son buenos contextos para testear la importancia de la cantidad de datos.

La segunda diferencia tiene que ver con que el comportamiento observado durante el periodo de adquisición lingüística diverge entre los seres humanos y las inteligencias artificiales.

El desarrollo del aprendizaje lingüístico humano tiene una naturaleza discontinua. En particular, los niños no aprenden de manera lineal a medida que acumulan experiencia, sino que atraviesan “saltos” en su competencia, que pueden describirse como una curva de aprendizaje en forma de “U” (Marcus y otros, 1992).

Un fenómeno muy común que ilustra este proceso es la *sobrerregularización*: la aplicación de un patrón gramatical productivo a formas irregulares: por ejemplo, decir *sabo* en lugar de *sé*. Este fenómeno está muy documentado (Yang, 2002; Mayol, 2007; Clahsen y otros, 2002) y representa entre el 5% y el 10% de las producciones de niños entre 2,5 y 4 años, y refleja lo que parece una regresión evolutiva: aunque inicialmente (alrededor de los dos años) los niños producen correctamente un verbo como *sé* por imitación, su desempeño “empeora” temporalmente al intentar aplicar una regla interna en lugar de limitarse a repetir lo que oyen, aplicando la regla a casos donde no se aplica.

Es interesante notar que los verbos irregulares son, de hecho, muy comunes en lenguas como el español y suelen ser los que los niños escuchan con más asiduidad. Aun así, los niños siguen la lógica de la generalización, ignorando esa frecuencia: aunque *sé* es muy frecuente, utilizan *sabo*.

Este fenómeno contrasta –como es de esperar– con el comportamiento de los LLMs. Dado que el aprendizaje de estos modelos es fundamentalmente estadístico, su rendimiento mejora de manera gradual y continua. Al depender de la probabilidad, los modelos no muestran este tipo de regresiones ni de sobreregularizaciones durante su entrenamiento, ya que la alta frecuencia de las formas correctas en sus trillones de datos “aplata” cualquier generalización que no se alinee exactamente con los datos. El hecho de que los niños otorguen esta prioridad a la existencia de reglas abstractas parece ser un sesgo cognitivo específicamente humano que no se reproduce en los LLMs (Kodner y otros, 2023).

Otra manera interesante de contrastar a los LLMs y los seres humanos consiste en testear aquellos conocimientos gramaticales que los lingüistas identificamos como reglas para las que los seres humanos recibimos muy poca evidencia directa en nuestro día a día. Se trata de ejemplos interesantes, ya que la inteligencia artificial también dispone de muy pocas aplicaciones claras de ellos en sus datos. Estas estructuras son las reglas típicas en las que pensaba Chomsky al proponer el argumento de la pobreza del estímulo. Para ilustrarlo, voy a adaptar al español un ejemplo de Lan y otros (2024).

Al preguntarle a ChatGPT³ cuál de estas dos oraciones es más aceptable, el *chatbot* consistentemente elige la número (2):

1. ¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó ayer?
2. ¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó a Ana ayer?

Cuál de estas preguntas es mejor en español: 1. ¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó ayer? o 2. ¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó a Ana ayer?

La mejor formulada es la 2.

2. “¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó a Ana ayer?”

La razón principal es la **claridad sintáctica**. En la primera:

1. “¿Qué chico dijiste que Kim *gosteó* el año pasado y que María besó ayer?”

la segunda parte (“que María besó ayer”) queda **incompleta o ambigua**, porque no se especifica a quién besó María. En español, el verbo *besar* normalmente necesita un objeto directo claro para evitar confusión.

En la segunda opción, aunque sigue siendo una oración compleja, “María besó a Ana ayer” está **completa y bien definida**, lo que reduce la ambigüedad y mejora la comprensión.

Si quieres, puedo ayudarte a reformularlas para que suenen más naturales en español (porque “*gosteó*” también es bastante informal/anglicismo).



³ Esta pregunta fue efectuada en Mayo 2026 al ChatGPT (GPT-4o). Nuevas versiones del ChatGPT pueden incluir en su entrenamiento oraciones de este tipo, lo que indefectiblemente llevará a que respondan correctamente, pero sirven para ilustrar el argumento de todas maneras.

Para un hablante nativo, la oración (2) es imposible (agramatical), mientras que la (1) es, aunque quizás no tan frecuente, correcta e interpretable. El contraste entre (1) y (2) ilustra una restricción sintáctica robusta observada originalmente por Ross (1967): cuando se coordinan dos cláusulas y se deja un “hueco” (una elipsis) en la primera, se debe también dejar un hueco análogo en el segundo elemento de la conjunción.

En la oración (2), ese segundo hueco se ha “llenado” indebidamente con “Ana”, rompiendo el paralelismo estructural necesario. El motivo por el que ChatGPT prefiere (2) sobre (1) es porque se rige enteramente por frecuencias de coaparición, y el predicado “besar” aparece mucho más frecuentemente con un receptor en su contexto cercano.

Lo sorprendente es que los niños reciben muy poca (por no decir ninguna) evidencia directa para aprender el contraste entre (1) y (2) (Pearl y Sprouse, 2013, para datos del inglés): nunca escuchan oraciones como (2), porque no son posibles, ni escuchan a sus padres decir “no digas esto”. Aun así, a diferencia de los LLMs, los niños infieren correctamente que (2) no es una oración posible en español. Este tipo de argumento sirve para mostrar, nuevamente, que los seres humanos imponemos sesgos en nuestra experiencia que no pueden inferirse directamente de esta.

Estudiar para saber: lo que la inteligencia artificial puede enseñarnos acerca de nosotros

En este artículo, intenté exponer por qué es importante analizar las diferencias entre los seres humanos y la inteligencia artificial (específicamente, los LLMs) en lo que respecta a la adquisición del lenguaje. La increíble similitud –al menos superficialmente– de comportamiento entre ambos es imposible de ignorar y hace que sea importante, precisamente, identificar en qué se asemejan y en qué divergen en tanto nos permite avanzar en la comprensión de nuestra propia cognición.

El análisis que hice sobre la adquisición del lenguaje podría extenderse a otros aspectos del conocimiento lingüístico. Por ejemplo, la mayoría de los LLMs se entrenan con un subconjunto ínfimo de las lenguas del mundo, lo que plantea dudas sobre su eficacia en estructuras muy distintas al inglés o español (Blasi y otros, 2022). Asimismo, los modelos muestran a veces un rendimiento “sobrehumano” al procesar estructuras que los humanos hallamos casi imposibles, como las oraciones con muchas subordinadas. Estos contrastes alimentan hoy debates muy vívidos en las ciencias cognitivas que invito al lector a explorar en las referencias citadas.

Referencias

- Blasi, D., Anastasopoulos, A. y Neubig, G. (2022). Systematic Inequalities in Language Technology Performance across the World’s Languages. En S. Muresan, P. Nakov y A. Villavicencio (Eds.), *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 5486–5505).
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D. y otros (2020). Language Models are Few-Shot Learners. En H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan y H.

- Lin (Eds.), *Advances in Neural Information Processing Systems* (Vol. 33, pp. 1877–1901). Disponible en: https://proceedings.neurips.cc/paper_files/paper/2020/file/1457c0d6bfc4967418bfb8ac142f64a-Paper.pdf
- Casillas, M., Brown, P. y Levinson, S. C. (2020). Early language experience in a Tzeltal Mayan village. *Child Development*, 91(5), 1819–1835.
- Chomsky, N. (1980). *Rules and representations*. Nueva York: Columbia University Press.
- Futrell, R. y Mahowald, K. (2025). How linguistics learned to stop worrying and love the language models. *Behavioral and Brain Sciences*, 1-98.
- Hart, B. y Risley, T. R. (1995). *Meaningful differences in the everyday experience of young American children*. Baltimore: Paul H. Brookes Publishing.
- Huebner, P. A., Sulem, E., Cynthia, F. y Roth, D. (2021). BabyBERTa: Learning More Grammar With Small-Scale Child-Directed Language. En A. Bisazza y O. Abend (Eds.), *Proceedings of the 25th Conference on Computational Natural Language Learning* (pp. 624–646). Stroudsburg: Association for Computational Linguistics.
- Katzir, R. (2023). Why large language models are poor theories of human linguistic cognition: A reply to Piantadosi. *Biolinguistics*, 17, 1-12.
- Kodner, J., Payne, S., y Heinz, J. (2023). Why linguistics will thrive in the 21st century: A reply to Piantadosi (2023). <https://doi.org/10.48550/arXiv.2308.03228>
- Lan, N., Chemla, E., y Katzir, R. (2024). Large language models and the argument from the poverty of the stimulus. *Linguistic Inquiry*, 1-28.
- MacWhinney, B. (2000). *The CHILDES Project: Tools for analyzing talk. Volume 1: Transcription format and programs*. Mahwah: Lawrence Erlbaum Associates.
- Marcus, G. F., Pinker, S., Ullman, M., Hollander, M., Rosen, T. J., Xu, F., y Clahsen, H. (1992). Overregularization in language acquisition. *Monographs of the society for research in child development*, 57(4).
- Mayol, L. (2007). Acquisition of irregular patterns in Spanish verbal morphology. En V. V. Nurmi y D. Sustretov (Eds.), *Proceedings of the Twelfth ESSLLI Student Session* (pp. 1–11). Dublín.
- Pearl, L. y Sprouse, J. (2013). Syntactic islands and learning biases: Combining experimental syntax and computational modeling to investigate the language acquisition problem. *Language Acquisition*, 20(1), 23-68.
- Rodriguez-Fornells, A., Münte, T. F. y Clahsen, H. (2002). Morphological priming in Spanish verb forms: An ERP repetition priming study. *Journal of Cognitive Neuroscience*, 14(3), 443-454.
- Ross, J. R. (1967). *Constraints on variables in syntax* [Tesis de doctorado, Massachusetts Institute of Technology].
- Tomasello, M. (2011). Language development. En U. Goswami (Ed.), *The Wiley-Blackwell handbook of childhood cognitive development (Second Edition)* (pp. 239-257). Oxford: Wiley-Blackwell.
- Warstadt, A. y Bowman, S. R. (2022). What artificial neural networks can tell us about human language acquisition. En S. Lappin y J. P. Bernardy (Eds.), *Algebraic structures in natural language* (pp. 17–60). Boca Raton: CRC Press.

- Weisleder, A. y Fernald, A. (2013). Talking to children matters: Early language experience strengthens processing and builds vocabulary. *Psychological Science*, 24(11), 2143–2152.
- Yang, C. D. (2002). *Knowledge and learning in natural language*. Oxford: Oxford University Press.

Este artículo de la revista Scholé es producto del trabajo colaborativo entre los autores y los diferentes equipos de producción del ISEP.

Autoría: Mora Maldonado

Equipo de producción de la revista

Dirección: Laura Percaz

Producción de contenidos: Lila Far

Edición: Martín Schuliaquer

Diseño e ilustración: Facundo Fernández

Maquetación: Daniel Wolovelsky

Desarrollo web: Santiago Rubiolo

Coordinación de producción: Matías Pardo

Coordinación de infraestructura digital: Paula Fernández

Autoridades

Martín Llaryora | Gobernador

Myriam Prunotto | Vicegobernadora

Horacio Ademar Ferreyra | Ministro de Educación

Victoria Flores | Secretaria General de Ambiente, Economía Circular y Biocidadanía

Luis Sebastián Franchi | Secretario de Educación

Gabriela Cristina Peretti | Secretaria de Innovación, Desarrollo Profesional y Tecnologías en Educación

María Eugenia Rotondi | Directora General de Innovación Educativa

Nora Esther Bedano | Secretaria de Coordinación Territorial

Claudia Amelia Maine | Secretaria de Fortalecimiento Institucional y Educación Superior

Lucía Escalera | Subsecretaria de Administración

Equipo de gestión del ISEP

Adriana Fontana | Directora

Lucía Cugini | Secretaria Académica

Victoria Farina | Secretaria de Organización Institucional

Laura Percaz | Secretaria de Contenidos Educativos

Cómo citar este material

Maldonado, M. (2026). Saber sin estudiar: ¿cómo nos diferenciamos de la inteligencia artificial cuando aprendemos lenguaje?. *Scholé* n.º 18. Para el Instituto Superior de Estudios Pedagógicos, Ministerio de Educación de la Provincia de Córdoba.

ISSN: 2683-7129

Este material está bajo una licencia Creative Commons ([CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc/4.0/))

