

LAS MATEMÁTICAS EN LA ERA DE LA IA

TERENCE TAO

Resumen. Ensayo, basado en una conferencia pública impartida en el Congreso Internacional de Matemáticos de 2026, sobre cómo la comunidad matemática podría responder a la llegada de herramientas de inteligencia artificial capaces de realizar tareas matemáticas de nivel de investigación.

En lugar de debatir las capacidades de dichas herramientas, partimos de la hipótesis de que estas capacidades llegarán y examinamos una cuestión ajena a ella: cuáles son realmente los objetivos y valores de la investigación matemática. El componente de resolución de problemas de las matemáticas se utiliza como caso de estudio.

1. Un prólogo histórico

Durante siglos, las matemáticas funcionaron con éxito sobre bases "ingenuas". Los matemáticos en ejercicio demostraban teoremas sobre conjuntos, números e infinitos sin sentir la necesidad de precisar qué eran realmente estas cantidades (o qué era una demostración en sí misma); tales cuestiones se delegaban en gran medida a los filósofos, y el matemático en ejercicio podía dedicarse libremente a las matemáticas.

Pero a principios del siglo XX, descubrimientos como la paradoja de Russell en 1901 [14] y los teoremas de incompletitud de Gödel en 1931 [8] obligaron a los matemáticos a reexaminar críticamente supuestos sobre su disciplina que hasta entonces se habían dado por implícitos. Los axiomas de la teoría de conjuntos ingenua se contradecían entre sí; y un sistema formal no podía ser simultáneamente consistente, suficientemente expresivo y capaz de demostrar su propia consistencia.

La consiguiente crisis en los fundamentos matemáticos, que se extendió aproximadamente desde 1900 hasta 1930, constituyó un período verdaderamente turbulento para la disciplina. Sin embargo, el resultado de esa turbulencia fue sumamente valioso: un marco fundamental explícito, riguroso y estandarizado, en el que los objetos de las matemáticas y las reglas para razonar sobre ellos se presentan de una forma que permite su análisis, enseñanza y, como se ha demostrado, su automatización. Ciertamente, existe margen para seguir mejorando este marco, y la investigación en fundamentos continúa hasta el día de hoy. Pero nuestros fundamentos actuales han superado un siglo de rigurosas pruebas y ahora proporcionan un entorno fiable en el que las matemáticas pueden practicarse con un alto grado de confianza.

Creo que estamos entrando en una era de turbulencia comparable en matemáticas. Esta vez, sin embargo, lo que se está poniendo a prueba no es nuestro marco fundamental para la verdad matemática, sino más bien el marco, en gran medida implícito, de valores y prácticas matemáticas: qué consideramos una contribución, qué premiamos, qué consideramos comprendido y a quién —o a qué— consideramos que ha realizado el trabajo. Sostengo que será necesario hacer mucho más explícitos estos objetivos no escritos de las matemáticas; pero una vez que hayamos...

Fecha: 18 de agosto de 2026.

Clasificación de asignaturas de matemáticas 2020. Primaria 00A30; Secundaria 01A80, 68T01, 68V20, 68V35.

Palabras y frases clave. Inteligencia artificial, práctica matemática, filosofía de las matemáticas, formalización, asistentes de prueba, revisión por pares, ley de Goodhart, valores matemáticos.

Tras examinarlas y codificarlas, nuestra comunidad saldrá fortalecida y más resiliente que antes.

2. La pregunta motivadora

Este artículo está organizado en torno a una única pregunta.

Pregunta 2.1 (Pregunta de respuesta de la comunidad). ¿Cómo debería responder la comunidad matemática¹ ante la llegada de las tecnologías modernas de IA y sus capacidades reales y/o declaradas para realizar tareas matemáticas?

Esta es una pregunta que concierne a toda la comunidad, y no pretendo tener todas las respuestas; nadie las tiene. Sin embargo, tengo algunas ideas sobre cómo abordarla.

Lo primero que cabe señalar es que la pregunta 2.1 no es una cuestión matemática. Es metamatemática, y también política, ética, sociológica y cultural; no puede resolverse mediante un argumento puramente lógico. Sin embargo, a continuación recurriré deliberadamente al lenguaje preciso y familiar de las matemáticas —conjeturas, hipótesis, etc.— para clarificar la estructura de la pregunta.

3. La primera subpregunta: capacidad de IA

La respuesta a la pregunta 2.1 depende fundamentalmente de una subpregunta sobre las capacidades reales de las herramientas de IA. Resulta conveniente formularla pseudomatemáticamente, no como una única conjetura, sino como una familia de conjeturas indexadas por un gran número de parámetros libres.

Conjetura 3.1 (Conjetura sobre la capacidad de la IA, plantilla). En algún momento del futuro cercano, algunas herramientas de IA podrán, a un costo determinado y con cierto nivel de supervisión humana, realizar tareas matemáticas de nivel de investigación en algunos campos de las matemáticas, con una tasa de éxito considerable y con cierto grado de precisión y calidad.

Cada aparición de la palabra «algunos» debe interpretarse como un marcador de posición (o un parámetro libre). Dependiendo de cómo se completen estos marcadores, se pueden obtener numerosas conjeturas distintas. Si bien las sutilezas entre estas formulaciones son importantes, no constituyen el objetivo de este artículo; me limitaré a la distinción general entre las formas «débiles» y «fuertes» de la Conjetura 3.1.

Si incluso las versiones más débiles de la Conjetura de la Capacidad de la IA resultan ser falsas, podríamos descartar sin temor a equivocarnos la generación actual de herramientas de IA por carecer de relevancia a largo plazo para la investigación matemática y continuar prácticamente como siempre. Si, por el contrario, las versiones más contundentes de la conjetura son ciertas, entonces resulta muy difícil mantener intactas nuestra cultura y prácticas actuales, sobre todo si seguimos priorizando objetivos como la obtención del mayor número posible de soluciones a problemas sin resolver.

Por lo tanto, resulta difícil entablar un debate constructivo sobre la Pregunta 2.1 mientras la validez de la Conjetura 3.1 siga en disputa. No es de extrañar, entonces, que la mayor parte del debate público sobre IA y matemáticas se haya centrado en qué versiones de la Conjetura de Capacidades de la IA son verdaderas; yo mismo he dedicado numerosas conferencias, escritos y publicaciones en redes sociales a este tema.

¹Los impactos de la IA, por supuesto, se extienden mucho más allá de las matemáticas, pero ese es un tema demasiado amplio para abordarlo aquí.

Actualmente existen muchísimos datos que respaldan diversas formas de la conjetura. Lamentablemente, la mayoría de ellas no se han recolectado en condiciones científicas controladas. Gran parte de la evidencia disponible públicamente está sujeta a un grave sesgo de información —los éxitos se anuncian y los fracasos no— y a incentivos no científicos, con costos y variables importantes (el número de intentos, la cantidad de apoyo humano, la capacidad de cálculo empleada, el grado de contaminación del problema con la literatura previa) que con frecuencia no se divulgan. Además, a veces se confunde el valor de verdad de una determinada formulación de la conjetura con su conveniencia.

A pesar de la relevancia central de la Conjetura de Capacidades de la IA para la Pregunta 2.1, este artículo no trata sobre dicha conjetura. En este sentido, mencionaré únicamente los resultados recientes del proyecto First Proof [7], una evaluación independiente de las capacidades de los modelos y arneses de IA de vanguardia en matemáticas genuinamente novedosas. Cada lote del proyecto consta de diez problemas de nivel de investigación, aportados por matemáticos en activo de una amplia gama de campos, cuyas soluciones son conocidas por el colaborador, pero nunca se han publicado en línea. El segundo lote [6] se evaluó en condiciones controladas frente a cuatro sistemas de IA, utilizando modelos de acceso público al 28 de mayo de 2026, y las soluciones resultantes fueron revisadas por expertos en cuanto a su corrección y calidad de exposición. De los diez problemas, siete recibieron al menos una calificación aprobatoria —es decir, una solución considerada prácticamente impecable o que requería solo revisiones menores— de al menos un sistema, con costes computacionales del orden de decenas a cientos de dólares por problema. Se prevén lotes adicionales.

4. El complemento a la conjetura de capacidad

Este artículo trata, en cambio, sobre lo que podríamos llamar el “complemento ortogonal” de la Conjetura de Capacidades de IA dentro de la Pregunta 2.1. Para aislar ese componente, adoptaré la siguiente hipótesis formulada de forma imprecisa.

Hipótesis 4.1 (Hipótesis de trabajo). Una versión razonablemente sólida de la Conjetura de Capacidades de la IA es cierta: las herramientas de IA serán capaces, en un futuro razonable, de realizar una fracción razonable de tareas matemáticas de nivel de investigación, con niveles razonables de éxito, calidad, supervisión y coste.

El significado preciso de “razonable” en este contexto no es crucial para lo que sigue.

En el resto de este artículo, pido al lector que asuma que la hipótesis 4.1 es cierta.

No le pido al lector que desee que sea cierto, que crea que es cierto o que lo acepte como tal; lo que sigue es un análisis condicional. En particular, la evidencia a favor o en contra de la hipótesis de trabajo es independiente de la discusión que se presenta a continuación.

5. La subpregunta ortogonal: nuestros objetivos y valores.

Una vez que se acepta la hipótesis de trabajo, surge una segunda subpregunta fundamental:

Pregunta 5.1 (Pregunta sobre metas y valores). ¿Cuáles son las metas, los objetivos y los valores precisos de nuestra comunidad matemática y de la investigación matemática en general? No me refiero solo a las metas explícitas que comunicamos al público, a nuestros estudiantes o a las agencias de financiación, sino también a las metas implícitas que buscamos optimizar en la práctica.

En el pasado, hemos delegado en gran medida la Pregunta 5.1 a las humanidades —a historiadores, filósofos y sociólogos de las matemáticas— y hemos centrado nuestra atención en el contenido técnico de nuestra profesión. Partiendo de la Hipótesis de Trabajo, ya no tendremos ese lujo. Pero —al igual que con la crisis de los fundamentos— sostengo que un examen crítico de la pregunta resultará, en última instancia, sumamente valioso, independientemente del estado de la Hipótesis de Trabajo.

Entonces: ¿cuáles son nuestros objetivos? Que yo sepa, no existe una lista oficial de los objetivos de las matemáticas. Se han recopilado sistemáticamente, pero aquí hay una lista parcial:

- Resolver problemas sin resolver, tanto puros como aplicados; • Desarrollar nuevas teorías, estructuras y técnicas; • Comprender el mundo que nos rodea;
- Construir y mantener una comunidad de matemáticos;
- Formar a la próxima generación de matemáticos y permitirles guiar las futuras direcciones de la disciplina; • Contribuir a la red compartida y acumulativa del conocimiento matemático; • Crear obras perdurables de valor estético; • etc.

Se anima al lector a ampliar esta lista.

Históricamente, los objetivos mencionados han estado correlacionados positivamente entre sí. El progreso en cualquiera de ellos suele acercarnos también a los demás. Por ejemplo, al resolver un problema complejo, se puede desarrollar una nueva técnica, se forma una nueva comunidad de investigadores en torno a ella, se capacita a estudiantes, se escriben libros de texto y la teoría resultante resulta aplicable en otros ámbitos. Debido a esta correlación, se podrían utilizar uno o dos de estos objetivos como indicadores convenientes de los demás, dejando el resto implícitamente expresados. Véase la Figura 1, así como [18] para un análisis previo del autor sobre este fenómeno de alineación.

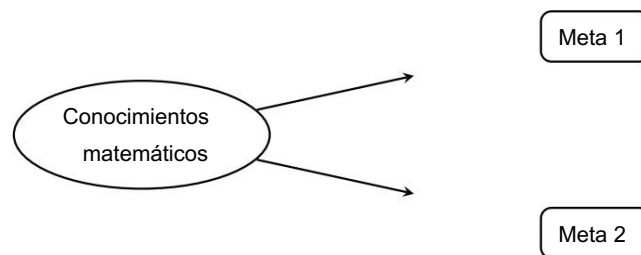


Figura 1. Históricamente, cualquiera de los objetivos de las matemáticas servía como un indicador útil de los demás.

Sin embargo, todas las métricas, cuando se optimizan excesivamente, corren el riesgo de verse afectadas por la ley de Goodhart. En la formulación popularizada por Strathern [16], siguiendo la observación original de Goodhart sobre política monetaria [9]: Cuando una medida se convierte en un objetivo, deja de ser una buena medida.

Las herramientas de IA son particularmente propensas a desencadenar este efecto, por dos razones independientes. La primera es técnica: la IA generativa es inherentemente infundada, en el sentido de que optimiza la apariencia de un resultado satisfactorio en lugar de la propiedad subyacente que se supone que el resultado debe indicar, y por lo tanto es inusualmente buena para encontrar la brecha entre una medida y

lo que mide. El segundo es económico: los incentivos financieros de la industria de la IA premian los logros demostrables, citables y comparables, precisamente en el tipo de métricas que históricamente hemos utilizado como indicadores indirectos.

En consecuencia, una optimización excesiva para uno o dos objetivos puede causar los muchos problemas previamente mencionados. objetivos alineados de las matemáticas divergen entre sí; véase la Figura 2.

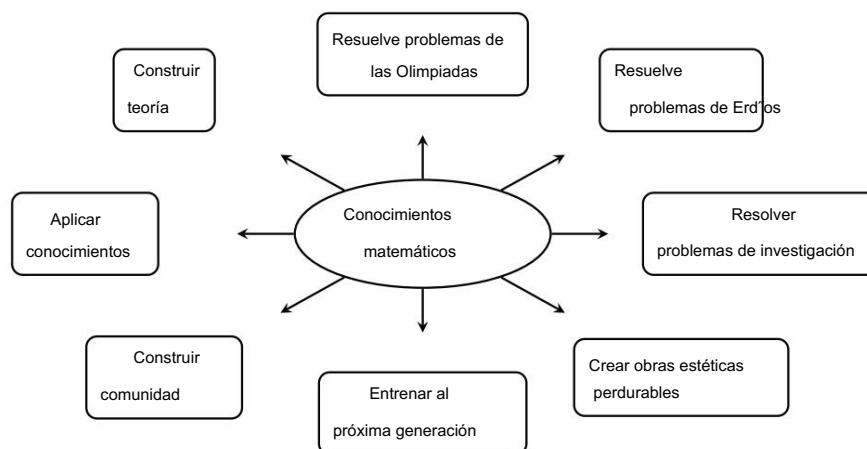


Figura 2. Bajo una optimización excesiva, los objetivos de las matemáticas divergen entre sí. El diagrama es, por supuesto, una simplificación excesiva; en particular, debería tener una dimensión mucho mayor.

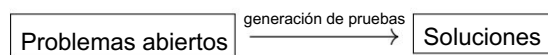
6. Un estudio de caso: resolución de problemas

Para concretar la discusión anterior, me centraré en un componente específico de la investigación matemática: la resolución de problemas. Cabe destacar que este no es, en absoluto, el único aspecto de nuestra profesión. La construcción de teorías, por ejemplo, es una actividad complementaria de igual o mayor importancia, que requiere un análisis independiente; lo mismo ocurre con la docencia, la mentoría y las diversas formas de servicio que sustentan a la comunidad investigadora. Sin embargo, la resolución de problemas constituye un primer caso de estudio idóneo, tanto por ser el aspecto más susceptible de verse afectado por la Hipótesis de Trabajo como por ser aquel en el que nuestros objetivos implícitos se alejan más de nuestros objetivos explícitos.

Supongamos entonces que intentamos escribir lo que queremos obtener de la resolución de problemas. Un primer intento podría ser el siguiente.

Objetivo 6.1 (primer intento). Resolver tantos problemas sin resolver como sea posible.

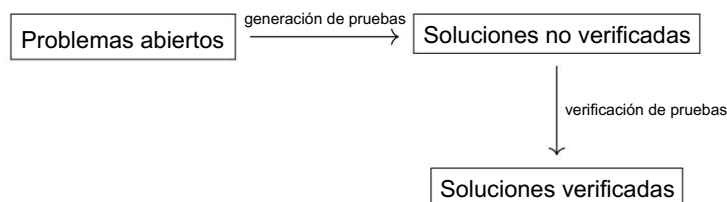
En el marco del Objetivo 6.1, estamos intentando optimizar el flujo en una red muy simple:



Incluso antes de la llegada de la IA, sabíamos que esta métrica era inadecuada, y lo sabíamos por una razón completamente trivial: optimizarla produce una gran cantidad de soluciones incorrectas a importantes problemas abiertos. Todo matemático en activo con una dirección de correo electrónico pública está familiarizado con el flujo constante de supuestas demostraciones de la hipótesis de Riemann. Por lo tanto, podemos actualizar nuestro objetivo:

Objetivo 6.2 (segundo intento). Resuelve tantos problemas sin resolver como sea posible y verifica que sean correctos.

La red correspondiente adquiere un segundo paso:

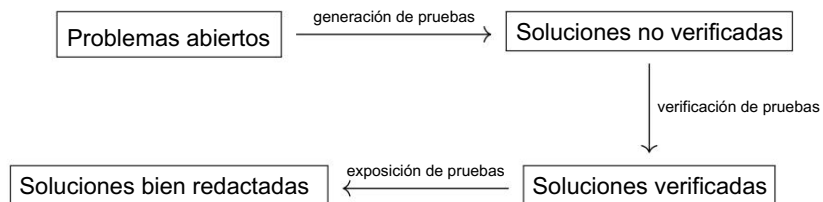


Los avances en IA y en la autoformalización en lenguajes de asistencia para pruebas² como Rocq, HOL o Lean [10, 12] han acelerado significativamente tanto la generación como la verificación de pruebas en muchos casos, y bajo la Hipótesis de Trabajo esta aceleración continuará. Una prueba formalmente verificada es, en definitiva, precisamente una prueba cuya corrección ya no depende de la reputación ni de la diligencia de su autor.

Pero ahora surge un nuevo modo de fallo. ¿Qué ocurre si una herramienta de IA genera una demostración extensa que se verifica como correcta, pero que nadie —ni siquiera los humanos que la crearon— comprende? Esto ya no es hipotético. Sitios dedicados a recopilar problemas matemáticos, como la base de datos de problemas de Erdős [5],³ ya contienen decenas de demostraciones generadas por IA. Es probable que muchas de ellas sean correctas; pero en un número considerable de casos, ningún experto humano se ha ofrecido a verificarlas y avalarlas, y en varios casos, los propios autores han declarado no estar cualificados para hacerlo. Pronto podríamos enfrentarnos a la posibilidad muy real de una demostración verificada de un resultado importante que ningún humano comprenda lo suficientemente bien como para explicarlo.

Por lo tanto, podemos actualizar nuestro objetivo nuevamente:

Objetivo 6.3 (tercer intento). Resolver tantos problemas sin resolver como sea posible, verificar su corrección y asegurar que los resultados puedan comunicarse con claridad y ser comprendidos por la comunidad matemática.



Las herramientas de IA actuales presentan resultados bastante dispares en cuanto a la exposición de pruebas. Por un lado, la ortografía, la gramática y el formato son prácticamente impecables. Por otro lado, la redacción suele detenerse extensamente en trivialidades, pasando por alto —o incluso ocultando activamente— las partes más interesantes y novedosas del argumento. Los textos matemáticos generados por IA tampoco suelen contextualizar el resultado en la literatura previa ni ofrecer una visión general que permita al lector decidir si merece la pena dedicarle tiempo.

²Para un análisis más detallado de los avances recientes en la formalización, véase [1].

³Para un estudio sobre el surgimiento de comunidades matemáticas en línea florecientes, véase [13].

4. Se podría argumentar que son demasiado perfectos.

La exposición de las demostraciones es, sin duda, un objetivo de optimización mucho más difuso que la verificación de las mismas, y la hipótesis de trabajo predice que las herramientas de IA mejorarán considerablemente en este aspecto con respecto a los niveles actuales. Pero aquí quiero destacar un punto que, en mi opinión, se subestima: la exposición también puede optimizarse en exceso. Una demostración puede estar redactada de forma demasiado pulcra, presentando tanto los pasos rutinarios como los realmente difíciles como igualmente fáciles de comprender.

En una demostración escrita por un humano, las partes del argumento que el autor encontró difíciles suelen conservar cierta fricción natural: una aclaración, un lema inusualmente cuidadoso, un cambio de notación, un párrafo que claramente ha sido reescrito varias veces. Esta fricción es informativa. Indica al lector dónde debe detenerse y prestar atención, y es uno de los principales canales por los que se transmite el conocimiento tácito de un campo. Una demostración excesivamente pulida por IA puede eliminar tanto la fricción «artificial» (errores tipográficos, frases torpes, desorganización) como la fricción «natural», dando como resultado un texto fácil de leer pero difícil de comprender. Paradójicamente, los «errores» en la exposición humana pueden ser realmente útiles para el lector; véase la Figura 3.

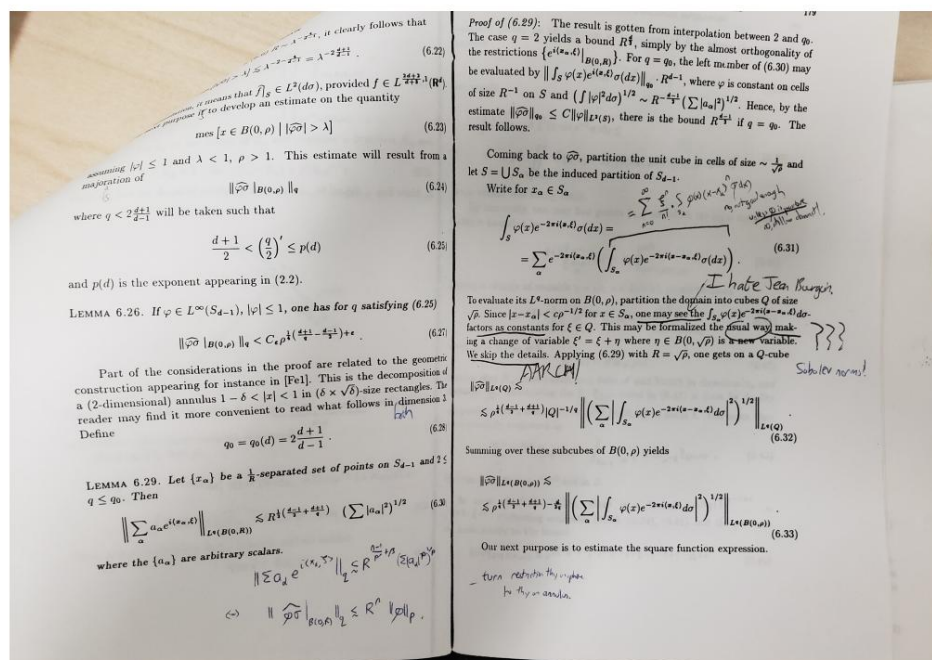


Figura 3. Una página de un artículo de Bourgain de 1991 [3], anotada por mí mismo, mucho más joven (y muy frustrado). Pero al esforzarme por comprender estos textos, llegué a entender la forma de pensar de Bourgain y, con el tiempo, busqué activamente sus artículos para leerlos. Véase también [19], [17].

Vale la pena recordar la formulación de Thurston sobre este punto, de su ensayo clásico [21], que Sigue siendo tan relevante en la era de la IA como lo era en 1994:

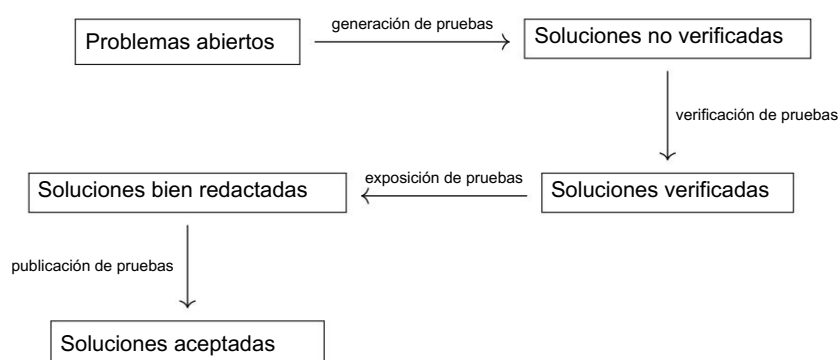
“No pretendemos cumplir con una cuota abstracta de producción de definiciones, teoremas y demostraciones. La medida de nuestro éxito reside en si nuestro trabajo permite a las personas comprender y reflexionar de forma más clara y eficaz sobre las matemáticas.”

Para que una demostración contribuya realmente a su campo, no basta con que sea correcta ni legible. También debe ser aceptada y valorada por la comunidad: otros matemáticos deben comprender el resultado e incorporarlo a su propio trabajo. Los autores pueden contribuir significativamente a este proceso de comprensión, describiendo las ideas iniciales, los intentos fallidos y las anécdotas del periodo en que trabajaron en el problema.

Por el contrario, las herramientas de IA actuales son bastante opacas en cuanto a su propio proceso de resolución de problemas, y esto es especialmente cierto en el caso de los modelos propietarios cuyo funcionamiento interno es un secreto corporativo.

Por lo tanto, podemos actualizar nuestro objetivo una vez más:

Objetivo 6.4 (cuarto intento). Resolver problemas sin resolver, verificar su corrección, asegurar que se comuniquen con claridad y lograr que la comunidad matemática los comprenda y acepte.



La aceptación de un resultado por parte de la comunidad científica es, por naturaleza, lenta y humana. Si bien puede fomentarse mediante una buena exposición y una redacción cuidadosa, en última instancia se trata de un proceso externo que no puede optimizarse únicamente con la ayuda de los autores y sus herramientas. Nuestra infraestructura de publicación actual depende de editores y revisores humanos para brindar esta aceptación, de forma voluntaria y, en gran medida, sin reconocimiento. Este trabajo suele considerarse menos prestigioso que el de generar las demostraciones en sí; sin embargo, es un componente esencial de la profesión y, precisamente, el mecanismo mediante el cual los logros individuales de los matemáticos se transforman en progreso y comprensión colectivos.

Las herramientas de evaluación basadas en IA podrían servir como filtros útiles en este proceso; cabe imaginar que las revistas clasifiquen automáticamente los artículos que presenten verificación inadecuada, falta de atribución o exposición incoherente, del mismo modo que se utiliza hoy en día la detección de plagio. Si bien su implementación es controvertida, estos filtros conservarían el escaso recurso de la atención humana experta. Sin embargo, superar un filtro automático no sustituye la aceptación de la comunidad; no creo que se pueda prescindir de los revisores humanos en la publicación.

proceso.

Finalmente, ni siquiera la publicación es la última etapa. Los resultados clave deberían, en última instancia, formar parte de los libros de texto y materiales de referencia definitivos de su materia, en la forma en que se enseñan a la próxima generación de estudiantes. Este proceso de canonización —en el que un resultado se reformula en su generalidad natural, se le da la demostración correcta en lugar de la primera, se conecta con sus resultados afines y se incorpora al conjunto de herramientas estándar— es la etapa más lenta de todas. Requiere un amplio consenso deliberativo y es la etapa menos susceptible de optimización mediante herramientas de IA. También es, en mi opinión, la parte más valiosa de todo el proceso.

Muchas aplicaciones de las matemáticas solo se vuelven factibles una vez que se ha establecido la teoría subyacente.

De esta forma, todo queda completamente digerido. De hecho, el éxito de las herramientas de IA en matemáticas depende fundamentalmente de las teorías canónicas que los matemáticos humanos han construido y reconstruido minuciosamente a lo largo de los siglos: los datos de entrenamiento para estas herramientas son, literalmente, el resultado del proceso de canonización.

Esto nos da una versión (potencialmente) final del objetivo de resolución del problema:

Objetivo 6.5 (¿último intento?). Resolver problemas sin resolver, verificar su corrección, asegurar que se comuniquen con claridad y lograr que se comprendan, acepten e incorporen a la teoría definitiva del campo.

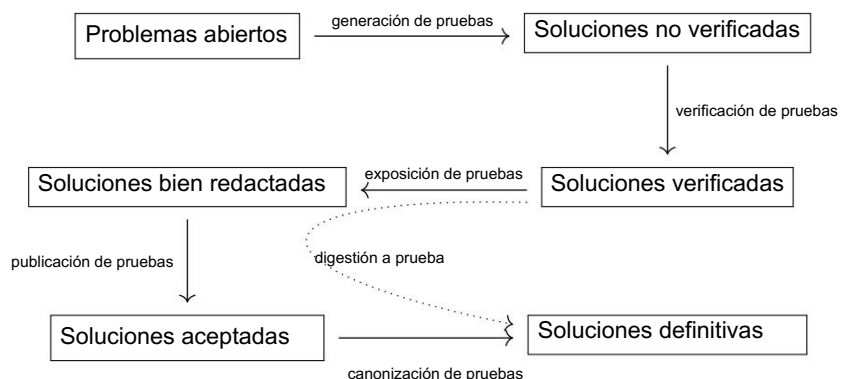


Figura 4. El proceso de resolución de problemas, en la forma obtenida al iterar los objetivos 6.1 a 6.5.

Lo que comenzó como una sola flecha se ha convertido en una cadena de cinco etapas, de las cuales solo la primera fue un objetivo explícito de la comunidad.

El proceso específico que se muestra en la Figura 4 puede estar simplificado en exceso y ser objeto de un análisis más profundo; sin embargo, ilustra los matices que se descubren al desglosar un objetivo, como la resolución de problemas, que parece simple en la superficie, pero que de hecho conlleva muchos subobjetivos implícitos que vale la pena explicitar.

7. Prueba de escasez y prueba de abundancia

Si la hipótesis de trabajo se cumple, entonces, en ausencia de cambios políticos y culturales adecuados, surgirán importantes “desajustes de impedancia” —o, para usar una metáfora menos halagadora, indigestión de pruebas— a lo largo de todo el proceso de la Figura 4:

- Las pruebas generadas por IA se acumularán más rápido de lo que se pueden verificar; •
- Las pruebas verificadas generadas por IA se acumularán más rápido de lo que se les puede dar una forma legible. redacción;
- Las pruebas generadas por IA, incluso aquellas que deben ser correctas y estar bien escritas, saturarán un sistema tradicional de revisión por pares que depende del trabajo de expertos voluntarios; • e incluso las pruebas publicadas serán demasiado numerosas para que la comunidad pueda trabajar en ellas. forma definitiva.

En resumen, pasaremos de una era de escasez de pruebas a una era de abundancia de pruebas. La mayoría de nuestras instituciones —revistas, convenciones prioritarias, criterios de contratación y promoción, premios, la propia noción de programa de investigación— fueron diseñadas bajo el supuesto de escasez, y no debería sorprendernos que se comporten mal bajo la abundancia. Algunos signos de esta indigestión

Ya estaban apareciendo antes de la llegada de la IA moderna: el aumento en el volumen de la literatura, la creciente longitud y especialización de las pruebas principales y la bien documentada La presión sobre el sistema de arbitraje es anterior al momento actual. Pero la llegada de la IA... exacerban notablemente estas tensiones existentes.

8. De los objetivos a las recomendaciones

Identificar los objetivos de la resolución de problemas de la forma en que lo hemos hecho anteriormente facilita considerablemente ver cómo responder a estos desajustes de impedancia emergentes, porque cada uno La discrepancia ahora está asociada a una etapa identificada y a un valor identificado. No tengo intención de... Propondré aquí un programa completo. En cambio, me referiré a la Declaración de Leiden sobre Energía Artificial. Inteligencia y Matemáticas [11], publicado en junio de 2026 y avalado por la International Unión Matemática, que considero un excelente punto de partida. La declaración surgió de un taller de 2025 en el Centro Lorentz en Leiden, y consta de veintitrés recomendaciones dirigidas a matemáticos individuales, a organizaciones matemáticas y financiadores sin fines de lucro y a los responsables políticos.

En lugar de reproducir la declaración, permítanme citar cuatro de sus recomendaciones a los matemáticos individuales y ofrecer un comentario sobre cada una desde la perspectiva desarrollada arriba.

Divulgue el uso de herramientas. Divulgue de forma transparente el uso de herramientas automatizadas, incluidos grandes modelos de lenguaje, sistemas de aprendizaje automático, asistentes de pruebas, y otro software matemático. Incluya una sección de “Divulgación de herramientas y recursos computacionales” en sus artículos; muchas revistas, editoriales y organizaciones profesionales ya han desarrollado directrices para esto, y aunque La forma precisa de dicha sección necesariamente evolucionará, alentamos a los autores a estar a la altura del espíritu reflejado en la Recomendación de la UNESCO sobre Ciencia Abierta [22] y los principios FAIR [23]. Al actuar como revisor, Cumpla con las directrices del editor. Si se permite el uso de inteligencia artificial, transparente sobre cómo lo usaste y asume la responsabilidad de cualquier daño significativo. recomendaciones que usted haga.

El escenario que debe evitarse a toda costa es aquel en el que los autores utilizan herramientas de IA de forma encubierta para ayudar su trabajo, pero ocultan ese uso para evitar críticas de sus compañeros. (Para mi propio

(Divulgación de herramientas de IA utilizadas en la preparación de este documento, véase la Sección 10.)

Apoyar las necesidades de revisión. El uso de inteligencia artificial en La preparación de artículos puede introducir material que hace que la revisión sea más exigente. Facilita a tus compañeros la revisión de tu trabajo divulgando la herramienta uso, dando referencias precisas y completas a resultados anteriores y proporcionando Demostraciones formales cuando sea factible y apropiado.

En términos más generales, sostengo que necesitamos disminuir el énfasis que nuestra cultura pone en generación de pruebas, y en particular sobre ser el “primero” en resolver un problema, y correspondientemente aumentar el énfasis que ponemos en la digestión de pruebas: exposición, revisión, publicación, y canonización. Véase también el reciente ensayo de Bessis [2] sobre la necesidad de alejarse de la “economía de teoremas” basada principalmente en la generación de demostraciones.

Reafirmemos la humanidad de la autoría. El mérito y la responsabilidad siguen perteneciendo a los seres humanos dentro de la comunidad matemática y no deben atribuirse a sistemas automatizados. La inteligencia artificial puede ocultar, pero no reemplaza, el trabajo humano colectivo que hay detrás de un resultado.

Esfuércese por atribuir correctamente las fuentes. Las limitaciones conocidas de las herramientas automatizadas para atribuir ideas de forma adecuada generan la obligación correspondiente de realizar un esfuerzo proactivo para encontrar y dar crédito a las fuentes que hicieron posible un nuevo resultado. Cuando no sea posible una atribución satisfactoria, indíquelo explícitamente en la publicación.

Mi regla general: si los autores no pueden demostrar de forma convincente que son capaces de presentar sus resultados de manera clara y experta, con una explicación correcta y debidamente atribuida, entonces el resultado no debería publicarse. Una prueba que ningún ser humano puede explicar adecuadamente debe considerarse incompleta, incluso si ha sido verificada formalmente.

9. Reflexiones finales

He presentado la resolución de problemas como un aspecto de las matemáticas en el que la hipótesis de trabajo nos obliga a examinar objetivos y valores que durante mucho tiempo hemos podido dejar implícitos. Sin embargo, la hipótesis de trabajo puede tener un impacto en muchos otros aspectos de nuestra labor —la docencia, la tutoría, la contratación, las solicitudes de subvención, la evaluación, la divulgación pública— y debería realizarse un análisis similar para cada uno de ellos.

Las conclusiones de estos análisis no serán uniformes. En algunos ámbitos, sobre todo en la educación y la formación de jóvenes matemáticos, será crucial destacar el aspecto intrínsecamente humano de nuestro trabajo y restringir estrictamente el uso de herramientas de IA; el objetivo de formar a un matemático no se consigue simplemente haciendo tareas correctas. En otros ámbitos, tendremos que tomar la iniciativa en el uso de la IA y definir las mejores prácticas para incorporar estas herramientas a nuestros flujos de trabajo según nuestros propios criterios, en lugar de los impuestos por los proveedores. También necesitaremos nuevos flujos de trabajo e infraestructuras que complementen los tradicionales: proyectos de formalización colaborativa, bases de datos de problemas estructuradas, nuevos espacios para la exposición y la publicación de resultados negativos o parciales; véase el Apéndice A para una lista parcial.

Ante todo, nuestra comunidad necesita unirse para entablar debates abiertos y honestos sobre las dos subcuestiones aquí planteadas: la capacidad de la IA y nuestros propios objetivos y valores. Esto, una vez más, es una recomendación de la Declaración de Leiden.

Participar en el debate público. Los matemáticos tienen la responsabilidad de apoyar el periodismo científico serio y de participar en el debate público para explicar y contextualizar los métodos y resultados obtenidos con la ayuda de la inteligencia artificial.

Esto es particularmente importante para el trabajo dentro de nuestros propios subcampos, donde se requieren conocimientos especializados para evaluar las afirmaciones sobre la profundidad, la dificultad y la importancia de los resultados. Además, alentamos a los matemáticos a buscar oportunidades para cooperar y apoyar a otros investigadores y profesionales creativos que enfrentan desafíos similares.

10. Agradecimientos

Este artículo se basa en una conferencia pública impartida en el Congreso Internacional de Matemáticos en julio de 2026. Agradezco a Bryna Kra, Jeremy Avigad, Martin Hairer, Akshay Venkatesh y Emily Riehl sus comentarios sobre las primeras versiones de dicha conferencia.

Se utilizó la asistencia de la IA para realizar búsquedas bibliográficas, generar diagramas, autocompletar texto y convertir las diapositivas a formato impreso.

Apéndice A. Algunos flujos de trabajo e infraestructuras nuevos

Para el lector interesado, enumero algunos proyectos existentes que ilustran los tipos de nuevos infraestructura mencionada anteriormente:

- Mathlib, una biblioteca formalizada unificada de matemáticas: <https://mathlib.org/> [12]; • Mathematical Discourse, una revista de vídeo revisada por pares para charlas de investigación: <https://www.mathematicaldiscourse.org/> [4];
- la base de datos de problemas de Erdős: <https://www.erdosproblems.com> [5]; • la base de datos de constantes de optimización: <https://github.com/teorth/optimizationproblems> [20];
- Las competiciones de matemáticas de la Fundación SAIR: <https://competition.sair.foundation/competitions> [15];
- El proyecto First Proof: <https://1stproof.org/> [7]. • El registro de pruebas formalizadas de Palomar: <https://palomar-registry.org/>

Referencias

1. J. Avigad, Matemáticas y el giro formal, Bull. AMS 61 (2024), 225–240.
2. D. Bessis, La caída de la economía del teorema, 21 de abril de 2026, <https://davidbessis.substack.com/p/la-caída-de-la-economía-del-teorema>.
3. J. Bourgain, Operadores maximales de tipo Besicovitch y aplicaciones al análisis de Fourier, Geom. Funct. Anal. 1 (1991), núm. 2, 147–187.
4. Mathematical Discourse, una revista de vídeo revisada por pares para charlas de investigación matemática, <https://www.mathematicaldiscourse.org/>.
5. T. Bloom, Problemas de Erdős, <https://www.erdosproblems.com>.
6. M. Abouzaid, N. Srivastava, et al., First Proof Second Batch, preimpresión, 2026. arXiv:2606.18119.
7. Proyecto First Proof, <https://1stproof.org/>.
8. K. Gödel, Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I, Monatshefte für Mathematik und Physik 38 (1931), 173–198.
9. CAE Goodhart, Problemas de gestión monetaria: la experiencia del Reino Unido, en Papers in Monetary Economics, Vol. I, Banco de la Reserva de Australia, 1975.
10. L. de Moura, S. Ullrich, El demostrador de teoremas y lenguaje de programación Lean 4, Deducción automatizada — CADE 28, Lecture Notes in Comput. Sci. 12699, Springer, 2021, 625–635.
11. La Declaración de Leiden sobre Inteligencia Artificial y Matemáticas, 2 de junio de 2026, <https://leidendeclaration.ai/>, doi:10.5281/zenodo.20302944.
12. La comunidad mathlib, La biblioteca matemática Lean, Actas de la 9.ª Conferencia Internacional ACM SIGPLAN sobre Programas y Pruebas Certificadas (CPP 2020), 367–381.
13. A. Pease, U. Martin, F. S. Tanswell y A. Aberdein, Uso de las matemáticas colaborativas para comprender la práctica matemática, ZDM 52 (2020), 1087–1098.
14. B. Russell, Los principios de las matemáticas, Cambridge University Press, 1903.
15. Concursos de la Fundación SAIR, <https://competition.sair.foundation/competitions>.
16. M. Strathern, “Mejora de las calificaciones”: auditoría en el sistema universitario británico, European Review 5 (1997), n° 3, 305–321.

17. P. Sarnak, T. Tao, I. Daubechies, F. Delbaen, L. Guth, S. Jitomirskaya, A. Kontorovich, E. Lindenstrauss, V. Milman, G. Pisier, Z. Rudnick, W. Schlag, G. Staffilani, P. Varjú, Remembering Jean Bourgain (1954–2018), *Notices Amer. Matemáticas. Soc.* 68 (2021), núm. 6, 942–957.
18. T. Tao, ¿Qué son las buenas matemáticas?, *Perspectivas matemáticas*, *Bull. Amer. Math. Soc.* 44 (2007), 623–634.
19. T. Tao, Explorando el conjunto de herramientas de Jean Bourgain, *Bull. Amer. Math. Soc.* 58 (2021), 155–171. doi:10.1090/bull/1716.
20. T. Tao et al., Una base de datos de constantes de optimización, [https://github.com/teorth/problemas de optimización](https://github.com/teorth/problemas-de-optimización).
21. WP Thurston, Sobre la demostración y el progreso en matemáticas, *Bull. Amer. Math. Soc. (NS)* 30 (1994), n.º 2, 161–177. arXiv:math/9404236.
22. UNESCO, Recomendación de la UNESCO sobre Ciencia Abierta, UNESCO, París, 2021. <https://doi.org/10.54677/MNMH8546>.
23. MD Wilkinson et al., Los principios rectores FAIR para la gestión y administración de datos científicos, *Scientific Data* 3 (2016), Artículo 160018. <https://doi.org/10.1038/sdata.2016.18>.

Departamento de Matemáticas de la UCLA, Los Ángeles, CA 90095-1555.

Dirección de correo electrónico: tao@math.ucla.edu