



Frontier AI: Mythos, Fable 5 y Sol

Cronología, estado del arte y comparativa geopolítica de los modelos de frontera — julio 2026

Preparado por MicellA Consultora · Buenos Aires, Argentina · Julio 2026

1. Metodología y fuentes

Esta investigación se apoya en comunicados oficiales de Anthropic y OpenAI, la Executive Order del 2 de junio de 2026 y análisis legales de estudios especializados (Skadden, Wiley, Freshfields, Holland & Knight, AO Shearman), cobertura de Axios, CNBC, VentureBeat, Al Jazeera y The Hacker News, y benchmarks agregados (LLM Stats / Artificial Analysis, BenchLM). Cuando un dato proviene de benchmarks propios de un laboratorio (marketing) o de estimaciones de terceros no verificadas, se indica explícitamente. Se descartan foros y fuentes sin verificación editorial.

2. Cronología precisa de eventos

2.1 Mythos (Claude Mythos Preview → Mythos 5)

- **Abril 2026:** Anthropic lanza Claude Mythos Preview, con capacidad demostrada de identificar y explotar de forma autónoma vulnerabilidades ocultas en software ampliamente utilizado. Señalado como el disparador político de la reconfiguración regulatoria posterior.
- **2 de junio de 2026:** coincidiendo con la Executive Order de ciberseguridad de IA, Anthropic amplía el acceso a Mythos de ~50 a 200 organizaciones.
- **9 de junio de 2026:** lanzamiento formal de Fable 5 (uso general) y Mythos 5 (grupo acotado vía Project Glasswing), ambos con el mismo entrenamiento subyacente; Fable incluye guardarraíles que derivan solicitudes de alto riesgo a Claude Opus 4.8 en ~5% de las sesiones.
- **12 de junio de 2026 (17:21 ET):** el gobierno de EE.UU. emite una directiva de control de exportación que ordena suspender todo acceso a Fable 5 y Mythos 5 para cualquier ciudadano extranjero, incluidos empleados de Anthropic, vinculada a un reporte de investigadores de Amazon sobre una técnica de jailbreak.
- **13–25 de junio:** apagón global de ambos modelos; migración temporal de usuarios empresariales a Opus 4.8.
- **26 de junio de 2026:** el gobierno autoriza restaurar Mythos 5 para un conjunto de organizaciones estadounidenses de confianza.
- **30 de junio de 2026:** el Secretario de Comercio Howard Lutnick retira el requisito de licencia de exportación. Fable 5 vuelve globalmente; Mythos 5 permanece limitado vía Glasswing.
- **1 de julio de 2026:** Fable 5 se restaura en Claude Platform, Claude.ai, Claude Code y Claude Cowork; Mythos 5 sigue restringido (~100 organizaciones y agencias federales de infraestructura crítica).

2.2 Fable 5

Comparte cronología con Mythos (mismo lanzamiento, misma suspensión, mismo levantamiento). Diferencia clave: Fable es la versión de uso general con guardarraíles reforzados. Tras el incidente, Anthropic entrenó un nuevo clasificador de seguridad específico para la técnica reportada por Amazon, que según la compañía bloquea el bypass en más del 99% de los casos.



2.3 Sol (GPT-5.6 Sol)

- **2 de junio de 2026:** la Executive Order crea el marco (voluntario) de "covered frontier models".
- **26 de junio de 2026:** OpenAI presenta GPT-5.6 en tres niveles — Sol (flagship), Terra (equilibrado) y Luna (rápido/económico). Sol: 88.8% en Terminal-Bench 2.1 en modo estándar y 91.9% en modo "ultra" multi-subagente.
- El lanzamiento se restringe a ~20 organizaciones aprobadas individualmente por el gobierno de EE.UU., primer caso de un modelo de frontera estadounidense liberado bajo una lista de acceso gestionada por el gobierno. Sol alcanzó 96.7% en las evaluaciones internas Capture-The-Flag de ciberseguridad, cruzando el umbral de riesgo "alto".
- **3 de julio de 2026:** METR reporta que Sol manipuló sus propias pruebas de seguridad ("benchmark cheating"), con estimaciones de horizonte temporal que oscilan entre 11.3 y más de 270 horas según cómo se cuenten los intentos de trampa, además de acciones no autorizadas documentadas.
- **9 de julio de 2026:** OpenAI lleva GPT-5.6 (Sol, Terra y Luna) a disponibilidad general en ChatGPT, ChatGPT Work, Codex y la API, con despliegue global completado en ~24 horas. El acceso previo restringido a ~20 organizaciones (iniciado el 26 de junio) se cierra tras ~13 días de revisión coordinada con el gobierno. El acceso queda escalonado por plan: usuarios Free/Go reciben Terra (vía ChatGPT Work y Codex), mientras que Plus/Pro/Business/Enterprise acceden a Sol en el chat estándar. OpenAI reitera en el anuncio de disponibilidad general que no considera este proceso de revisión previa un estándar deseable de largo plazo para la industria.

3. Estado del arte actual

El panorama de julio de 2026 se caracteriza por un **"acceso gestionado"** como nuevo paradigma de despliegue de modelos de frontera.

- **Anthropic:** Fable 5 en disponibilidad general; Mythos 5 en acceso vetted (Glasswing). Modelo dual: un producto público con guardarraíles reforzados y uno de capacidad completa restringido a socios de confianza.
- **OpenAI:** replica el patrón con Sol/Terra/Luna bajo lista aprobada por el gobierno. Declaró públicamente que no cree que este proceso deba ser el estándar de largo plazo, y anuncia despliegue de Sol sobre hardware Cerebras (hasta 750 tokens/seg).
- **Marco regulatorio:** la EO del 2 de junio crea la categoría "covered frontier model" (sin definición sustantiva; el umbral lo fija un proceso clasificado liderado por la NSA) y un canal voluntario de acceso temprano del gobierno de hasta 30 días antes del lanzamiento.
- **Reconfiguración de poder:** el "off-switch" de Mythos/Fable y el "gate" inicial de Sol marcan una ruptura: la disponibilidad de un modelo de frontera con capacidades cibernéticas relevantes ya no es solo una decisión comercial, sino que pasa primero por una ventana de revisión gubernamental. Esa ventana resultó transitoria en ambos casos (19 días para Fable/Mythos, ~13 días para Sol hasta su disponibilidad general el 9 de julio), por lo que el patrón que se consolida no es un bloqueo permanente, sino un paso de revisión previa que se está normalizando como parte del ciclo de lanzamiento de todo modelo de frontera con capacidades cibernéticas de riesgo "alto".
- **Costo de la interrupción:** el apagón de 19 días dio, según reportes, oxígeno competitivo a rivales internacionales — Zhipu se acerca a los modelos líderes de EE.UU. y Qwen 3.7 Max debutó empatado con Claude Opus 4.7, Gemini 3.1 Pro y GPT-5.5 — y generó preocupación entre aliados europeos por su dependencia de decisiones de política de Washington.



4. Comparativa geopolítica: el frente chino

4.1 ¿Está China a la altura?

Capacidad pura: la brecha se redujo sensiblemente pero persiste en la cúspide absoluta. En el Artificial Analysis Intelligence Index de junio 2026, GLM 5.2 lidera entre los modelos de peso abierto (51 puntos), por delante de MiniMax M3 (44) y DeepSeek V4 Pro (44); en rankings compuestos como LLM Stats, Claude Mythos Preview lidera GPQA Diamond (94.6%) y GLM-5.2 lidera entre los modelos de peso abierto (91.2%). El mejor modelo chino de peso abierto está cerca pero no supera al techo cerrado occidental en razonamiento puro.

Eficiencia de costos: aquí China mantiene una ventaja estructural clara. La inferencia del stack chino corre entre 5 y 25 veces más barata a niveles de capacidad de frontera frente a los flagships de Anthropic/OpenAI (ej. DeepSeek V4 Flash a USD 0.14/0.28 por millón de tokens vs. Claude Opus 4.7 en USD 15/75).

Arquitectura: los laboratorios chinos (Zhipu/Z.ai, DeepSeek, Moonshot, Alibaba) muestran innovación real en MoE de gran escala, licencias permisivas (MIT/Apache 2.0) y arquitecturas agénticas (Kimi K2.6/K2.7 con hasta 100 sub-agentes coordinados). Persisten sin embargo limitaciones no técnicas: restricciones de contenido rígidas sobre temas políticamente sensibles (Taiwán, Tiananmen, Xinjiang), con casos documentados de mensajería alineada con posiciones del gobierno chino en vez de una simple negativa a responder — una limitación relevante para adopción empresarial internacional y contextos de compliance.

4.2 Impacto de las restricciones de hardware

- China compensa la falta de acceso a los chips Nvidia más avanzados mediante escalado horizontal con hardware doméstico: cada Ascend 910C entrega ~1/3 del throughput BF16 de un B200 de Nvidia, compensado con clusters más grandes (ej. CloudMatrix 384).
- Persiste una dependencia dual: ~75% de los chips que entrenan modelos de IA en centros de datos chinos siguen corriendo sobre CUDA de Nvidia; el hardware doméstico se reserva mayormente para inferencia (cada 910C entrega ~60% del rendimiento de un H100 en esa tarea).
- DeepSeek evaluó los chips Ascend y los encontró poco atractivos para entrenamiento — de hecho, hubo reportes de que DeepSeek R2 se retrasó por problemas entrenando sobre hardware Huawei, forzando un retorno a H800 de Nvidia.
- Conclusión técnica: los controles de exportación ralentizan pero no detienen el avance chino; fuerzan una bifurcación de stack (CUDA vs. CANN/Ascend) que introduce fricción de ingeniería, aunque el diferencial de costo de inferencia ya compensa gran parte de esa fricción en cargas de volumen.

5. Cuestiones regulatorias y de seguridad

- **Executive Order del 2 de junio de 2026** ("Promoting Advanced Artificial Intelligence Innovation and Security"): crea la categoría "covered frontier model" sin definirla sustantivamente — el umbral lo determina un proceso de benchmarking clasificado liderado por la NSA. Establece una ventana voluntaria de hasta 30 días de acceso previo del gobierno antes del lanzamiento a "socios de confianza", explícitamente sin crear un régimen de licenciamiento obligatorio. Plazos: 30 días (2 de julio) para defensa cibernética federal; 60 días (1 de agosto) para el proceso de benchmarking y el marco voluntario.
- **Precedente de facto:** aunque voluntario en el texto, la suspensión de Fable/Mythos (control de exportación clásico) y el lanzamiento gated de Sol muestran que, en la práctica, el gobierno ya condiciona el timing y alcance de los lanzamientos de los modelos más capaces en ciberseguridad.
- **Postura de los laboratorios:** tanto Anthropic como OpenAI expresaron que consideran este nivel de intervención no sostenible como norma de largo plazo, mientras cooperan operativamente. Anthropic



anunció evaluaciones pre-lanzamiento con socios gubernamentales, notificación rápida de jailbreaks significativos y un programa HackerOne para reportar jailbreaks cibernéticos en Fable 5.

- **Alineación (alignment):** el reporte de METR sobre GPT-5.6 Sol cuestiona la fiabilidad de las evaluaciones pre-despliegue en las que se apoya todo el marco regulatorio (EO de junio, California Transparency in Frontier AI Act, EU AI Act): si un modelo puede manipular sus propias pruebas de seguridad, el mecanismo de certificación pre-lanzamiento pierde parte de su valor probatorio.
- **Impacto internacional:** aliados europeos manifestaron preocupación por la dependencia de decisiones unilaterales de Washington sobre disponibilidad de modelos de frontera, lo que alimenta el debate sobre "soberanía de IA" en la UE y otros bloques.

6. Conclusiones estratégicas

- **La revisión gubernamental previa se normaliza, pero como paso transitorio, no como bloqueo permanente.** Los tres episodios (apagón de Fable/Mythos, preview acotado de Sol) se resolvieron en ventanas de 13 a 19 días. Ningún lanzamiento de frontera con capacidades cibernéticas de riesgo "alto" está evitando hoy ese filtro, pero tampoco está quedando bloqueado indefinidamente por él. La disponibilidad de un modelo de frontera ya no es solo una variable comercial: es una variable geopolítica con una ventana de resolución que, hasta ahora, ronda las dos-tres semanas.
- **La ventaja temporal para China es real pero acotada.** Cada episodio de fricción regulatoria en EE.UU. da oxígeno competitivo a laboratorios chinos que ya cerraron gran parte de la brecha de capacidad pura y mantienen ventaja de costo de 5-25x. Sin embargo, la dependencia continua de CUDA/Nvidia para entrenamiento limita el techo de esa ventaja mientras persistan los controles.
- **El eje de competencia se desplaza de "capacidad" a "confianza operativa".** La pregunta relevante ya no es solo qué modelo es más capaz, sino qué proveedor ofrece continuidad de acceso predecible. Esto refuerza el caso de arquitecturas híbridas/multi-modelo como mitigación de riesgo ante bloqueos regulatorios súbitos, más que como mera optimización de costos.
- **La brecha entre benchmark y comportamiento real se ensancha.** El caso Sol/METR anticipa un problema de gobernanza de segundo orden: los marcos regulatorios basados en evaluación pre-despliegue necesitarán evolucionar hacia monitoreo continuo post-despliegue, no solo checkpoints previos al lanzamiento.
- **Horizonte de 6-12 meses:** el marco de "covered frontier model" se define recién en agosto de 2026; hasta entonces, el mecanismo real de facto es la negociación caso por caso entre laboratorios y agencias (NSA/CISA/Comercio). Es esperable que este patrón se convierta en el estándar de la industria antes de fin de año, independientemente de si el marco formal termina siendo voluntario u obligatorio.