

米中フロンティアAIモデルの市場動向と技術比較: 主要10モデルの分析と「半年遅れ」説の検証

Gemini 3.7 Flash

米国フロンティアモデルの技術的特徴と展開

米国における最新の大規模言語モデル(LLM)開発は、単なるパラメータ規模の拡大から、マルチエージェントオーケストレーション、高度な推論モードの実装、および多層的な安全性ガイドラインの統合へと明確に舵を切っている¹。商用クラウドAPIを通じた展開を主軸としながら、実用的な開発現場やエンタープライズでの即応性を高める取り組みが進められている¹。

GPT-5.6 Sol: 推論強化と多層的防御スタック

OpenAIが先行公開したGPT-5.6 Solは、同社のSol、Terra、Lunaで構成される新世代モデルファミリーのフラッグシップ(最上位)モデルとして位置づけられている¹。約105万トークンのコンテキストウィンドウを維持しつつ、高度な論理展開を行う「Max」推論モードに加え、応答品質を維持しながら従来比14倍の処理高速化を実現するマルチエージェント主導の「Ultra」モードを備えている²。独立した評価指標であるArtificial Analysis Coding Agent Indexにおいて約80ポイントの最高スコア(SOTA)を記録しており、複雑なソフトウェア構造の解析や自律的プログラミングタスクにおいて極めて高い完遂能力を示す²。

また、GPT-5.6 Solは技術的性能の追求と並行して、同社で最も堅牢な安全構造である多層的ガードレールスタック(Layered Safeguard Stack)を導入した¹。モデルレベルでのジェイルブレイク拒否学習に加え、生成プロセスのリアルタイム監視やアカウントシグナル検知を統合することで、悪意あるサイバー攻撃コードの作成を制約している¹。その一方で、システム脆弱性の特定や修復パッチの自動適用といった防御目的のセキュリティ業務においては高度な支援能力を提供できるよう調律されている¹。

Claude Opus 5 および Claude Fable 5: Mythosクラスと実用性の階層化

Anthropicは、モデルの世代区分(Claude 5)と性能・用途に応じたアーキテクチャ階層(Mythosクラス等)を明示的に切り分ける戦略を採用した⁷。最位階層である「Mythosクラス」に基づいて構築されたClaude Fable 5は、100万トークンのコンテキストウィンドウと128kトークンの最大同期出力を持ち、長期にわたる自律タスクの完遂に特化した能力を有している⁴。SWE-Bench Proで80.4%、Terminal-Bench 2.1で88.0%、CursorBenchで72.9%といった業界最高水準のベンチマークを記録する一方²、サイバー攻撃やバイオ・化学兵器に関連する高リスクな要求を検出する厳格な安全分類器(Safety Classifiers)が組み込まれている⁴。過度な危険性が検知されたリクエストに対しては、自動的に前世代のClaude Opus 4.8へリダイレクト(フォールバック)する制御が機能する⁴。さらに米国の輸出管理規則の適用を受け、一部のアクセス制限措置が講じられている⁹。

これに対してClaude Opus 5は、Fable 5に迫るフロンティア級の知性を維持しながら、API価格を入

力100万トークンあたり**5.00**、出力100万トークンあたり 25.00とFable 5の半額に設定した日常業務向けのフラッグシップとして提供されている⁵。100万トークンの文脈窓と128kトークンの出力をサポートし、思考機能(Thinking)が標準で有効化されている³。Frontier-Bench v0.1やCursorBench 3.2において前世代のOpus 4.8を大きく引き離すスコアを記録したほか、対話途中でのツール構成変更(Mid-conversation tool changes)やプロンプトキャッシュ最小閾値の512トークンへの引き下げなど、高度なエージェント組み込みに適した実効的機能が拡充されている³。

Gemini 3.7 Flash、Grok 4.6、Muse Spark 1.2: 高速化・効率化とエージェント制御

GoogleのGemini 3.7 Flashは、高速レスポンスと低運用コストの向上に注力したワークホースモデルであり、1,048,576トークンの入力と65,536トークンの出力に対応し、テキスト、画像、音声、動画、PDFといった多様な入力を一括処理するマルチモーダル性能を備える¹⁵。毎秒274トークンから340トークンに達する圧倒的な生成速度を誇り¹⁸、従来モデルと比較して複雑ドキュメント処理(GDP.pdf)で22.0%から34.0%へ、業務自動化評価(AutomationBench)で17.0%から30.4%へ、コマンドライン制御(Terminal-Bench 3.0)で5.4%から14.9%へと着実な性能向上を果たしている¹⁶。

SpaceXAIのGrok 4.6は、前世代モデルGrok 4.5が採用していた約1.5兆(1.5T)パラメータのベース基盤をそのまま維持しながら、強化学習(RL)と高度な合成データによる事後学習(ポストトレーニング)のみを深めることで知能を引き上げたモデルである²。入力 **2.00**、出力 6.00という費用対効果を保ちつつ、長時間の自律エージェント運用や視覚的アプリケーションの構築能力を強化し、実用開発プラットフォームへの統合が進められている²。

Metaが発表したMuse Spark 1.2は、Meta Superintelligence Labs (MSL)が開発した100万トークン対応モデルであり、同時に公開されたターミナル型AIコーディングエージェント「Muse Code」の中核エンジンとして機能する²²。目標追跡(Goal tracking)や操作ログの記録、クラッシュ時からの作業再開(muse resume)機能を備え、大規模コードベースにおける長時間の自律開発作業を端末環境上で安定して実行する環境を提供している²²。

中国フロンティアモデルの進化とオープンウェイト戦略

中国の主要AI開発機関は、巨大規模モデルのオープンウェイト(Open-Weights)配布戦略を積極的に展開し、グローバルな開発者コミュニティの取り込みとローカル環境での高度なカスタマイズ性の両立を追求している¹³。また、事後学習の極限的スケールリングによって、短期間で米国最上位モデルとの性能差を縮小させている²⁷。

Kimi K3 および Qwen3.8 Max: 超大規模MoEモデルの展開

Moonshot AIがリリースしたKimi K3は、2.8兆(2.8T)パラメータ規模を誇る超大型オープンウェイトMoE(Mixture-of-Experts)モデルであり、1,048,576トークンの文脈窓と最大100万トークンの出力を提供する¹³。API利用料が入力 **2.80**、出力 14.00と手頃に設定されている一方で²⁵、Artificial Analysis Intelligence Indexにおいて世界3位から4位(スコア57~60)に位置し、米国の商用クロード最上位モデル(GPT-5.6 SolやClaude Opus 5)に迫る極めて高い知能を実証している¹³。ただし、約1.56TBに及ぶVRAM要求スペックのため、最小推論環境であってもハイエンドGPUを16基以上要求する重量級のアーキテクチャとなっている¹³。

AlibabaのQwen3.8 Maxは、2.4兆(2.4T)パラメータ規模のMoEオープンウェイトモデルとして公開され、複数の独立ベンチマークにおいて米国最先端モデルであるClaude Fable 5に比肩する成果を示している¹³。総合的なプログラミング評価「KingBench 3」で81.25%を獲得したほか²⁷、画像内の異物や正常オブジェクトを正確に識別・カウントする高精度な視覚認識テストにおいてGemini 3.6 Flashと同率1位を記録するなど、視覚と言語を横断する多角的な推論能力を備えている³²。

DeepSeek-V4-Pro-0813: 極限のコスト効率と独自アテンション構造

DeepSeekのDeepSeek-V4-Pro-0813は、総パラメータ数1.6兆(1.6T)、アクティブパラメータ数490億(49B)のMoE構造を採用したモデルであり、入力0.435、出力0.87/1Mトークン(キャッシュ入力時は\$0.003625)という低価格なAPI体系を実現している¹⁴。独自のCSA(Coarse-grained Sparse Attention)とHCA(Hybrid Context Attention)を組み合わせることで、100万トークンの長文文脈における推論計算コストを大幅に抑制している¹⁴。この計算効率化を達成しながらも、SWE-bench Verifiedで80.6%、MMLU-Proで87.5%、Codeforcesレーティングで3206を記録し、米国商用上位モデルと同等の推論性能を低コストで提供することに成功している¹⁴。

GLM-5.3: 事後学習による性能躍進とサイバーセキュリティ能力の創発

Z.ai(智譜AI)がリリースしたGLM-5.3は、事前学習基盤を前世代モデルGLM-5.2(7,530億パラメータ)のまま維持しながら、強化学習および長尺タスク環境でのポストトレーニング(SAO手法およびSlime非同期学習フレームワーク)を極限までスケールアップすることで著しい能力向上を果たした²⁷。独立したプログラミング評価ベンチマーク「KingBench 3」において、GLM-5.3は91.25%(73/80)を獲得し、Claude Opus 5(77.5%)、Kimi K3(77.5%)、Qwen3.8 Max(81.25%)を上回り、Claude Fable 5(82.5%)をも凌ぐ最高スコアを記録した²⁷。さらに、長尺なターミナル操作を評価するTerminal-Bench 3.0においても、前世代の4.6から28.3へとスコアを大幅に向上させている²⁸。また、事後学習のスケールアップに伴い、GLM-5.3には想定を超える高度なサイバーセキュリティ解析能力(Emergent Cyber Capability)が自発的に発現した²⁸。ソースコードからの脆弱性特定テスト(CyberGym)において84.5%を獲得し、Mythos 5(83.8%)やGPT-5.6 Sol(83.6%)を超えるスコアを記録したほか²⁹、実用的なエクスプロイト構成を評価するExploitBenchでも54.4%を記録した²⁸。実際のオープンソースプロジェクトにおいて2,436件の脆弱性(うち1,097件が過酷なリスク)を特定した実績を持つ²⁸。こうした攻撃的サイバー機能のリスクに鑑み、Z.aiは安全評価とハードニングを完遂するため重みの一般公開を約2週間保留し、段階的なアクセス管理を適用する慎重な公開プロセスをとっている²⁸。

米中フロンティアモデルの定量的スペックおよびベンチマーク比較

以下の表は、各モデルの最新の公開仕様、価格設定、および主要ベンチマークにおける定量的性能データを一括して整理した比較データである。

モデル名	開発機関	提供形態	コンテキスト窓 /	API価格 (入力 / 出	Intelligence Index	主要ベンチマーク
------	------	------	-----------	------------------	--------------------	----------

	/ 国		最大出力	力 1Mトークン)	スコア	スコアおよび技術的特徴
GPT-5.6 Sol	OpenAI (米国)	クローズド API	1.05M / 非公表 ²	\$5.00 / \$30.00 ²	61 ¹⁸	Coding Agent Indexで~80 (SOTA) ² 。多層的な安全スタックと防御的サイバー機能 ¹ 。14倍高速なUltraモード統合 ⁶ 。
Claude Opus 5	Anthropic (米国)	クローズド API	1M / 128k ³	\$5.00 / \$25.00 ¹⁴	63 (max) / 61 ¹⁴	SWE-bench Proで79.2% ¹⁴ 。Frontier-Bench v0.1 SOTA ⁵ 。思考機能標準有効 ³ 。プロンプトキャッシュ最小512トークン ³ 。
Claude Fable 5	Anthropic (米国)	クローズドAPI (一部制限) ⁴	1M / 128k ⁴	\$10.00 / \$50.00 ⁹	62 (fallback 時) ¹⁸ / 100 ¹⁴	Mythosクラス ² 。SWE-Bench Proで80.4% ² 。Terminal-Bench 2.1で88.0% ⁷

						。高リスク要求時にOpus 4.8へ自動退避 ⁴ 。
Gemini 3.7 Flash	Google (米国)	クローズド API	1M / 65k ¹⁶	\$0.60-\$0.75 / \$3.00-\$3.75 ¹⁶	56 - 57 ¹⁸	生成速度 274-340 tok/s ¹⁸ 。Terminal-Bench 3.0で 14.9% ¹⁷ 、GDP.pdfで34.0% ¹⁶ 。マルチモーダル一括処理 ¹⁶ 。
Grok 4.6	SpaceX AI (米国)	クローズド API	未公表	\$2.00 / \$6.00 ²	62 ¹⁸	1.5Tベースの事後学習強化モデル ² 。長時間の自律エージェント運用および対話的アプリ生成機能 ²¹ 。Cursor統合 ²¹ 。
Muse Spark 1.2	Meta (米国)	API / エージェント統合	1M / 未公表 ²³	タスク平均 \$0.40 ¹⁸	54 (xhigh) ¹⁹	Terminal-Bench 2.1で82.9% ²³ 。「Muse Code」エージェントの中

						核 ²² 。目 標追跡お よびロー カルログ 復元機能 ²² 。
Kimi K3	Moonshot AI (中国)	オープン ウェイト / API	1M / 1M ²⁵	\$2.80 / \$14.00 ²⁵	57 - 60 ¹⁸	2.8T MoE 構成 ¹⁴ 。 オープン モデルと して最高 峰の Intelligen ce Index スコア(世 界3~4 位) ¹³ 。 VRAM要 求 1.56TB ³¹ 。
Qwen3.8 Max	Alibaba (中国)	オープン ウェイト予 定 / API	未公表	未公表	最高水準 ¹³	2.4T MoE 構成 ¹³ 。 KingBenc h 3で 81.25% ²⁷ 。視覚オ ブジェクト カウントテ ストで同 率1位 ³² 。
DeepSee k-V4-Pro -0813	DeepSee k (中国)	オープン ウェイト / API	1M / 384k ¹⁴	\$0.435 / \$0.87 ¹⁴	53 ¹⁸	1.6T MoE (アクティ ブ49B) ¹⁴ 。 SWE-ben ch Verifiedで 80.6% ¹⁴ 。

						CSA+HC Aハイブリッドアテンションによる超低価格API ¹⁴ 。
GLM-5.3	Z.ai (中国)	オープンウェイト (段階的) ²⁸	未公表	ZCode / 契約形式 ²⁸	最高水準 ²⁹	753Bベースの事後学習強化 ²⁷ 。 KingBench 3で91.25% (首位) ²⁷ 。 CyberGymで84.5% ²⁹ 。 脆弱性特定2,436件 ²⁸ 。

米中技術格差の多角的検証と「半年遅れ」説の解体

過去において米中に存在した技術的格差、特に「中国のAIモデルは米国より半年遅れている」とされた状況は、AIモデルの開発パラダイムが「事前学習パラメータの巨大化」から「強化学習(RL)と高度な検証環境を活用したポストトレーニングのスケールアップ」へ移行したことによって実質的に払拭された²。Grok 4.6やGLM-5.3が証明したように、事前学習済みベースモデルの規模を変更せずとも、高品質な合成データの投入と環境内での強化学習を推し進めることで、短期間(わずか数ヶ月)のうちに複雑なコーディングや多段階推論のベンチマークスコアを飛躍的に向上させることが可能となっている²。先端半導体の調達制限を受ける中国企業であっても、ポストトレーニングにおける効率的な計算アロケーションとタスク環境の最適化を推進することで、米国最先端モデルが示す推論能力に短期間で追従し、特定ドメインでは逆転できる技術構造が形成されている²⁷。この技術的メカニズムの変化が、従来の「半年遅れ」という固定的な時間差を無効化させる主要因となっている¹³。

また、米中間の競争軸は単なる「絶対的知能の勝敗」から「エコシステムの展開モデルと運用の自由度」という新たな次元へと移行している¹³。米国主要ラボは最上位モデル(Claude Fable 5やGPT-5.6 Sol)を秘匿性の高いクローズドAPIとして提供し、安全分類器や政府アクセスの監視下で商用化を進める戦略をとる¹。これに対して、Moonshot AI、Alibaba、DeepSeek、Z.aiといった中国の主要機関は、2兆パラメータを超える超大型モデルやSOTA級のプログラミングモデルをオープンウェイトとして広く一般に配布するアプローチを採用している¹³。独立評価機関のデータにおいて、Kimi K3や

Qwen3.8 Max、GLM-5.3といった中国発オープンモデルは、米国商用クローズドの最高峰（Claude Opus 5やGPT-5.6 Sol）との総合スコア格差を数ポイント圏内に縮めており、プログラミングやデータ操作などの実用環境においては同等以上の評価を得ている¹³。自社サーバーでの直接運用やローカル環境での完全なデータ制御を重視する開発者コミュニティやエンタープライズにとって、実質的に利用可能な最高性能の選択肢が中国発のオープンモデルとなりつつあり、実効的な技術的拡散力において中国側が明確なポジションを確立している¹³。

さらに、プログラミングや自律エージェントの高度化が進むにつれ、モデル内部で自発的に育まれるサイバー攻撃・防御能力が米中双方の安全保障上の課題として浮上している¹。プログラムの深層論理を解釈し自動改修を行う機能は、本質的に脆弱性を探知してエクスプロイトを構成する能力と表裏一体の関係にある¹。GLM-5.3が提示した高い脆弱性検証能力やゼロデイ攻撃コードの構成力は、ポストトレーニングの深化が想定外のサイバー能力を自発的に創発させることを実証した²⁸。この事態を受け、米国企業は輸出管理やリアルタイムフィルタリングによる制限を強化し、中国企業もまた自発的なリスク評価と段階的公開プロセス（モデルウェイト保持とハードニング）を余儀なくされている¹。結果として、技術開発そのものの遅れではなく、各国の安全ガバナンスおよび法規制への適合手続きがモデルリリースのタイムラインを決定づける主要因となっており、市場投入速度の差は純粋な技術的限界ではなく安全管理戦略の差としてあらわれている¹。

結論

米中フロンティアAIモデルの包括的な性能検証および市場動向の分析から導き出される結論として、「中国のAIモデルは米国に対して半年遅れている」という過去の共通認識は、現在のフロンティア市場において既に成立しなくなっている¹³。

米国は依然として、最高峰のクローズドフラッグシップ（Claude Fable 5やGPT-5.6 Sol）において、極めて広範な汎用論理推論や厳格な安全ガードレールの統合面で僅小なリードを保持している¹。しかし、ソフトウェアエンジニアリング自動化、ターミナル環境における長時間のエージェント制御、コストパフォーマンス、およびオープンウェイトによるエコシステムの支配力においては、中国のフロンティアモデル（GLM-5.3、Kimi K3、DeepSeek V4 Pro、Qwen3.8 Max）が米国最先端と完全に対等、あるいは部分的に凌駕する領域に到達している¹³。

今後は、絶対的な知能指数のみを競う段階から、用途に応じたモデルの動的ルーティング、事後学習環境の高度化、および創発的サイバーリスクに対する安全ガバナンスへの対応力を軸とした、より多層的な米中技術競争へと移行していくことが確実視される¹。

引用文献

1. Previewing GPT-5.6 Sol: a next-generation model | OpenAI, <https://openai.com/index/previewing-gpt-5-6-sol/>
2. Grok 4.6 vs Grok 4.5, GPT-5.6 Sol & Fable 5 - Eigent AI, <https://www.eigent.ai/blog/grok-4-6-vs-grok-4-5-gpt-5-6-fable-5>
3. What's new in Claude Opus 5, <https://platform.claude.com/docs/en/about-claude/models/whats-new-opus-5>
4. Introducing Claude Fable 5 and Claude Mythos 5 - Claude Platform Docs, <https://platform.claude.com/docs/en/about-claude/models/introducing-claude-fable-5-and-claude-mythos-5>

5. Introducing Claude Opus 5 - Anthropic,
<https://www.anthropic.com/news/claude-opus-5>
6. GPT-5.6 Solの「超高速モード」ついに実装「品質損なわず14倍高速」うたう,
<https://www.itmedia.co.jp/aipplus/article/2608/14/2000000545/>
7. Claude Fable 5, Explained: The Most Capable Model Ever Released – and the Fine Print That Changes Everything | by Enrico Papalini | Medium,
<https://medium.com/@enrico.papalini/claude-fable-5-explained-0c4d94448ea9>
8. What Is Claude Fable 5? Anthropic's Mythos-Class Model for Agentic Work | MindStudio,
<https://www.mindstudio.ai/blog/what-is-claude-fable-5-anthropic-mythos-class-model>
9. Claude Fable 5: Model Specifications and Details - ApX Machine Learning,
<https://apxml.com/models/claude-fable-5>
10. Claude Fable 5 - API Pricing & Benchmarks | OpenRouter,
<https://openrouter.ai/anthropic/claude-fable-5>
11. Claude Fable 5 & Claude Mythos 5 System Card - Anthropic,
<https://www-cdn.anthropic.com/d00db56fa754a1b115b6dd7cb2e3c342ee809620.pdf>
12. Claude Fable 5 and Claude Mythos 5 - Anthropic,
<https://www.anthropic.com/news/claude-fable-5-mythos-5>
13. The best AI you can own is Chinese. The West needs to close that gap quickly.,
<https://www.atlanticcouncil.org/blogs/the-best-ai-you-can-own-is-chinese-the-west-needs-to-close-that-gap-quickly/>
14. Top 10 AI Models of August 2026: From Claude Fable 5 to DeepSeek V4 | みさき - note, https://note.com/bright_jacana710/n/nf45cca54c74b?hl=en
15. Introducing Gemini 3.7 Flash,
<https://blog.google/innovation-and-ai/models-and-research/gemini-models/introducing-gemini-3-7-flash/>
16. What Is Gemini 3.7 Flash? Benchmarks, Price, and API Guide - CometAPI,
<https://www.cometapi.com/what-is-gemini-3-7-flash>
17. Gemini 3.7 Flashとは？ベンチマーク、価格、APIガイド - CometAPI,
<https://www.cometapi.com/ja/what-is-gemini-3-7-flash/>
18. Gemini 3.7 Flash (medium)とGrok 4.6 (high): モデル比較 - Artificial Analysis,
<https://artificialanalysis.ai/ja/models/comparisons/gemini-3-7-flash-medium-vs-grok-4-6>
19. Artificial Analysis: AI Model & API Providers Analysis, <https://artificialanalysis.ai/>
20. Google releases Gemini 3.7 Flash for coding and agents - DataNorth AI,
<https://datanorth.ai/news/google-releases-gemini-3-7-flash>
21. SpaceXAI unveils Grok 4.6, claims frontier-level performance at lower cost,
<https://indianexpress.com/article/technology/artificial-intelligence/spacexai-unveils-grok-4-6-claims-frontier-level-performance-at-lower-cost-10831649/>
22. Meta コーディングエージェント「Muse Code」公開、新モデル「Muse Spark 1.2」搭載 - ビジネス+IT, <https://www.sbb.it.jp/article/cont1/186487>
23. Meta、ターミナル型AIコーディングエージェント「Muse Code」をベータ公開 100万トークンの新モデル「Muse Spark 1.2」搭載 | Ledge.ai,

- https://ledge.ai/articles/meta_muse_code_muse_spark_1_2
24. 「Muse Spark」登場: 人を中心に設計されたMeta Superintelligence Labs初のモデル,
<https://about.fb.com/ja/news/2026/04/introducing-muse-spark-meta-superintelligence-labs-first-model-built-to-prioritize-people/>
 25. Kimi K3 - API Pricing & Benchmarks - OpenRouter,
<https://openrouter.ai/moonshotai/kimi-k3>
 26. China AI Model Benchmarks: A Powerful 2026 Shock to Silicon Valley - DataCore Blog,
<https://blog.datacore.vn/en/china-ai-model-benchmarks-2026/>
 27. GLM-5.3 Benchmark Results: How It Stacks Up Against Opus 5, Fable 5 | MindStudio,
<https://www.mindstudio.ai/blog/glm-5-3-benchmark-test-results>
 28. GLM-5.3: Z.ai Coding Model, Benchmarks & Weights - Eigent AI,
<https://www.eigent.ai/blog/glm-5-3-coding-cyber-model>
 29. GLM-5.3: Frontier Coding with Emergent Cyber Capabilities - Z.ai,
<https://z.ai/blog/glm-5.3>
 30. MetaのMuse Spark 1.2がAI総合性能でGrok 4.5と肩を並べる、わずか4カ月で急成長,
<https://finance.biggo.jp/news/a28535ce-84ad-442d-9568-6c836b05d754>
 31. What is Kimi K3? A Complete Developer Guide for 2026 - Firecrawl,
<https://www.firecrawl.dev/blog/kimi-k3>
 32. メモ: Qwen3.8-Max の画像認識 - Zenn,
<https://zenn.dev/kun432/scraps/ba4e832aa8c23b>
 33. DeepSeek V4 Pro 0813 Benchmarks & Pricing (August 2026) | BenchLM.ai,
<https://benchlm.ai/models/deepseek-v4-pro-0813>
 34. GLM-4.5 vs GLM-5.3: Benchmarks, Pricing & Which Is Better in 2026 - LLM Stats,
<https://llm-stats.com/models/compare/glm-4.5-vs-glm-5.3>
 35. GLM-5.3 - Overview - Z.AI DEVELOPER DOCUMENT,
<https://docs.z.ai/guides/llm/glm-5.3>
 36. GLM 5.3 released: Frontier Coding with Emergent Cyber Capabilities : r/singularity - Reddit,
https://www.reddit.com/r/singularity/comments/1vnz30c/glm_53_released_frontier_coding_with_emergent/
 37. GLM-5.3 is here with advanced cyber capabilities — and reportedly already found a 'serious vulnerability' in Cursor | VentureBeat,
<https://venturebeat.com/technology/glm-5-3-is-here-with-advanced-cyber-capabilities-and-reportedly-already-found-a-serious-vulnerability-in-cursor>