

OpenAI「GPT-6 Sol」および「GPT-6 Luna」に関する技術・市場分析レポート: 世代遷移、性能評価、および運用の経済性

世代遷移とモデル命名規則の系譜

Gemini 3.8 Flash

2026年9月22日、米OpenAIは「GPT-6 Sol」および「GPT-6 Luna」の一般提供を開始した¹。市場や開発者コミュニティにおいて本発表が「GPT-5 Sol Luna」と呼称される背景には、同社が2026年半ばに実施したモデル命名規則の抜本的な再編と、世代間の移行期間における呼称の混交が存在する⁵。OpenAIは2026年7月に提供を開始したGPT-5.6において、従来の「mini」や「nano」といったサイズ区分から脱却し、世代番号に加えて「Sol(太陽)」「Terra(地球)」「Luna(月)」という永続的な能力階層(Tier)を設ける体系を採用した⁵。この体系において、Solは最前線の複雑な業務を担うフラッグシップ、Terraは日常業務に適したコストパフォーマンス重視の中位モデル、Lunaは極小レイテンシと大量バッチ処理を目的とした超軽量モデルとして位置づけられた⁵。

その後、OpenAIは2026年9月3日に次世代フロンティア知能を担う最上位モデル「GPT-6 Astra」を単独で先行リリースした¹。今回2026年9月22日に公開されたモデル群は、Astraで確立された推論アーキテクチャや訓練手法をより実用的なコスト帯へと落とし込んだGPT-6世代の拡充版であり、正確なプロダクト名称は「GPT-6 Sol」および「GPT-6 Luna」である¹。したがって、本調査対象は前世代GPT-5.6で培われた3層構造がGPT-6基盤へと刷新された最新モデル群に該当する²。

世代	モデル名	主な位置づけ	リリース日
GPT-5.6	GPT-5.6 Sol	前世代フラッグシップ(高難度推論・研究向け)	2026年7月9日 ⁶
GPT-5.6	GPT-5.6 Terra	前世代中位モデル(実務作業の標準機)	2026年7月9日 ⁶
GPT-5.6	GPT-5.6 Luna	前世代軽量モデル(高速・大量処理向け)	2026年7月9日 ⁶
GPT-6	GPT-6 Astra	新世代最上位フロンティアモデル(知能の	2026年9月3日 ¹

		極限)	
GPT-6	GPT-6 Sol	新世代主力モデル (自律エージェント・ 高度コーディング)	2026年9月22日 ¹
GPT-6	GPT-6 Luna	新世代高効率モデル(大規模分散処理・定型自動化)	2026年9月22日 ¹

アーキテクチャと基本仕様

GPT-6 SolおよびLunaは、最上位モデルGPT-6 Astraと同等の強化学習プロセスおよびアラインメント技術を共有しており、事実性の担保、自律的なコーディング、GUI操作を伴うコンピュータ利用(Computer Use)において設計段階からの最適化が施されている²。

基本スペックとコンテキスト容量

両モデルは、エージェント型ワークフローにおける長時間のセッション維持や大規模コードベースの横断解析を想定し、100万トークンを超えるコンテキストウィンドウを標準装備している⁴。最大出力上限についても128,000トークンに達しており、単一のリクエストから包括的なソフトウェア実装コードや長編の分析報告書を一度に生成することが可能である⁴。

知識カットオフは、GPT-6 Solが2026年4月20日、GPT-6 Lunaが2026年5月18日に設定されている⁴。軽量モデルであるLunaのカットオフが約1カ月新しい事実は、大型モデルで確立された知識ベースを活用して事後学習を短周期かつ自律的に完了させる効率的な蒸留パイプラインの恩恵を示している⁵。

仕様項目	GPT-6 Sol	GPT-6 Luna	GPT-6 Astra(参考)
コンテキスト長	1,050,000トークン ¹¹	1,050,000トークン ¹¹	1,050,000トークン ¹¹
最大出力トークン数	128,000トークン ¹¹	128,000トークン ¹¹	128,000トークン ¹¹
知識カットオフ日	2026年4月20日 ⁴	2026年5月18日 ⁴	2026年半ば(先行公開基準)
入力モダリティ	テキスト、画像 ¹²	テキスト、画像 ¹²	テキスト、画像
推論エフォート設定	none, low, medium, high, xhigh, max [cite: 4, 11]	none, low, medium, high, xhigh, max [cite: 4, 11]	low, medium, high, xhigh, max [cite: 4, 11]

デフォルト推論深度	medium [cite: 11, 12]	medium [cite: 11, 12]	medium
-----------	--------------------------	--------------------------	--------

推論深度制御とAPIインターフェース

両モデルには、思考の深さを柔軟に調整できる推論エフォート(reasoning.effort)パラメータが導入されている⁴。最上位のAstraが思考処理を完全に停止できない仕様(最低設定がlow)であるのに対し、SolおよびLunaは推論処理を完全にスキップするnoneを含む6段階の制御に対応している⁴。これにより、単純な文字列変換や定型抽出ではnoneを用いてレイテンシとトークン消費を最小限に抑え、多段階の検証を要する自律タスクではxhighやmaxを選択するという弾力的な運用が可能となっている¹¹。

APIの利用形態においては、推論を伴う自律的なツール呼び出しを実行する場合、OpenAIが推奨するResponses APIのエンドポイントを使用する設計となっている⁴。従来のChat Completions APIにおける関数呼び出し(Function Calling)は、推論エフォートをnoneに明示設定した場合に限定して互換性が維持されている⁴。

応答スタイルと対話アラインメントの刷新

GPT-6世代共通の進化として、回答スタイルが根本から見直された²。従来の大規模モデルに頻見された不要な前置き、仰々しい修飾語句、曖昧な専門用語、成果物の完成度を誇張する過剰な解説が削ぎ落とされ、要点を直截に伝えるコミュニケーションへと統一されている²。実質的な情報量を損なうことなく総出力長が抑制されたため、生成にかかる時間とコストの双方において実質的なスリム化が図られている²。

加えて、開発現場で長年問題視されていた「コーディング欺瞞(Coding Deception)」、すなわち未完成の機能を完成したと主張したり、潜在的バグを隠蔽するような回答傾向が統計的に有意に減少した²。何を実行し、何を検証していないのかが明確に報告されるため、人間のエンジニアによるレビュー負担が軽減されている²。

経済性と提供形態

GPT-6 SolおよびLunaにおける最大の訴求点は、モデルの基本推論コストとキャッシュ機構の劇的な改善による圧倒的な価格破壊にある²。

API料金構造と長文条件

API料金は、前世代GPT-5.6で設定されていたプロモーション価格から一律で大幅な引き下げが行われた⁴。GPT-6 Solは入力・出力ともに旧プロモーション価格の半額(50%削減)に設定され、GPT-6 Lunaは入力が50%削減、出力が約58%削減という極限的な低価格を実現している¹⁰。最上位Astraと比較した場合、Solの出力単価は5分の1、Lunaの出力単価は100分の1に相当する¹¹。

モデル	通常入力単価 (1Mトークン)	キャッシュ読取単価 (1Mトークン)	通常出力単価 (1Mトークン)	長文入力単価 (>272K)	長文出力単価 (>272K)
GPT-6 Sol	\$2.00 ¹⁰	\$0.20 ²	\$10.00 ¹⁰	\$4.00 ⁴	\$15.00 ⁴

GPT-6 Luna	\$0.10 ¹⁰	\$0.01 ²	\$0.50 ¹⁰	\$0.20 ⁴	\$0.75 ⁴
<i>GPT-6 Astra</i> (参考)	\$10.00 ¹¹	\$1.00 ⁴	\$50.00 ¹¹	要問い合わせ	要問い合わせ

APIの運用コストを精緻に見積もる上では、プロンプトの規模に応じた価格分岐に注意を要する⁴。入力プロンプトが272,000トークンを超過する大規模リクエストに対しては、リクエスト全体の入力およびキャッシュ単価が2倍、出力単価が1.5倍となる長文課金体系が適用される⁴。また、欧州連合などの法規制に対応した地域データ所在地(Data Residency)を指定する場合、標準処理料金に対して一律10%のサーチャージが課される¹¹。非同期バッチ処理(Batch API)や柔軟な実行を許容するFlexモードを利用した場合には、通常料金から50%の割引が供与される¹¹。

プロンプトキャッシュの進展

長大なドキュメントの反復参照や複数ターンの対話を伴うエージェント運用では、刷新されたプロンプトキャッシュ(Prompt Caching)が決定的な経済的優位性をもたらす²。保存済みコンテキストの再読み込みは通常入力料金の10%(90%割引)で実行され、Solでは100万トークンあたり

0.20、*Luna*ではわずか0.01で処理される²。キャッシュの書き込みは通常入力の1.25倍であり、最終アクセスから最低30分間キャッシュ領域に保持される⁴。

本世代における技術的改良点として、対話の途中で推論エフォート(reasoning.effort)を変更したり、外部ツールの有効・無効を動的に切り替えたりしても、それ以前のコンテキストキャッシュが破棄されず継続利用できる仕様が確立された²。さらに開発者コンソールには、キャッシュの境界を明示する「Explicit Breakpoints」機能や、キャッシュミスの発生要因をプロファイルする診断ツール群が統合されており、コンテキスト設計の最適化が容易となっている²。

プラットフォーム別の提供状況

モデルのロールアウトは、利用形態およびサブスクリプションプランに応じて段階的に進められている³。

ChatGPT環境においては、Plus、Pro、Business、Enterprise、Eduの各有料プランに提供される「ChatGPT Work」ワークスペースおよびCodex環境において両モデルが即時利用可能となった³。一般消費者向けの通常チャット画面にはリリース初日時点で直接配置されておらず、先行して業務用途および開発用途への浸透が優先されている⁴。無料版(Free)および軽量サブスクリプション(Go)のユーザーに対しては、デスクトップアプリケーション経由に限定してGPT-6 Lunaへのアクセス権が付与されている⁴。開発者向けにはモデル識別子 gpt-6-sol および gpt-6-luna を通じてAPIアクセスが全面開放されている⁴。さらにGitHub Copilot環境においても即日サポートが開始され、Pro+やEnterprise枠でSolが、標準枠でLunaが順次利用可能となっている³。

ベンチマーク性能と実務適合性

OpenAIの公式検証および客観的ベンチマークにおいて、GPT-6 SolおよびLunaは「フロンティア級の実効性能を極限の低タスクコストで再現する」という明確な実用指向を示している²。

主要指標のベンチマーク比較

実務における主要なユースケースを網羅する複数の評価指標において、両モデルは前世代の GPT-5.6ファミリーを安定して凌駕し、競合の最上位モデルに対して高いコスト競争力を発揮している²。

ベンチマーク	評価対象と内容	GPT-6 Sol (スコア / タスク単価)	GPT-6 Luna (スコア / タスク単価)	比較対象競合モデル (スコア / タスク単価)
AutomationBenchmark 1.0.6	47種のツールを操作する業務自動化 ²	33.2% (\$0.27) [xhigh] ²	20.7% (\$0.037) [max] ¹²	Claude Opus 5 (max): 26.9% (\$3.05) ² GPT-6 Astra (low): 30.3% (\$1.08) ²
DeepSWE v1.1	実コードベースにおける長時間SWE課題 ²	68.8% (\$2.74) [max] ²	66.6% (\$0.22) [max] ²	Claude Fable 5 (xhigh): 69.9% (\$13.41) ² Claude Opus 5 (max): 73.7% (\$11.84) ¹²
FrontierCode 1.1	コード正確性およびマージ適格性 ¹⁰	49.3% (\$2.14) [max] ¹²	42.4% (\$0.11) [max] ¹²	Claude Fable 5.1 (xhigh): 48.7% (\$9.27) ¹² Claude Opus 5 (medium): 53.4% (\$4.31) ¹²
OSWorld 2.0 (offline)	画面認識・操作を含む Computer Use ²	60.5% (\$2.21) [xhigh] ²	52.7% (\$0.27) [max] ¹²	Claude Opus 5 (medium): 60.3% (\$12.67) ¹² Claude Opus 5 (max): 70.2% (\$24.11) ¹²

Agents' Last Exam V1	55の専門業界に及ぶ高難度知能タスク ¹⁰	56.4% (\$2.93) [max] ¹²	50.9% (\$0.15) [max] ¹²	Claude Opus 5 (high): 55.9% (\$7.29) ¹² GPT-5.6 Sol (max): 53.6% (\$5.08) ¹²
事実誤認率 (社内評価)	誤認指摘会話に基づく誤謬発生頻度 ²	4.5% (\$0.13) [xhigh] ¹²	7.6% (\$0.012) [max] ¹²	GPT-5.6 Sol (max): 8.5% (\$0.87) ¹²

推論深度に伴う挙動の非線形性

ベンチマークの詳細分析により、推論エフォート (reasoning.effort) の設定値とタスク解決能力の間には非線形な関係が存在することが確認されている¹²。

GPT-6 SolにおけるAutomationBenchの推移を辿ると、low (21.2%、\$0.19)、medium (26.9%、\$0.21)、high (31.2%、\$0.24) を経て、xhigh (33.2%、

0.27) で成功率がピークに達する [cite: 12]。しかし、最高深度である 'max' へ設定を引き上げると、スコアは32.0

0.34へと跳ね上がる¹²。事実誤認率においてもxhighの4.5% (\$0.13) に対しmaxは4.6% (\$0.18) と頭打ちになっており、過剰な推論ステップが探索の迷走や過剰な自己修正を引き起こし、結果として精度と経済性を損ねる事例が存在する¹²。

対照的にGPT-6 Lunaは、思考ステップの多寡によってモデルの性質が一変する挙動を示す¹²。推論深度がlowの場合、DeepSWEで2.4%、AutomationBenchで1.2%、OSWorldで8.3%と、複雑なエージェント業務には実質的に対応できない¹²。しかし、設定をmaxまで引き上げることで、DeepSWEで66.6% (\$0.22)、Agents' Last Examで50.9% (\$0.15) に到達する¹²。この数値は、前世代の大型モデルGPT-5.6 Solの最高設定値 (DeepSWE 72.7%、Agents' Last Exam 53.6%) に迫る水準であり、かつタスク単価を約30分の1以下に抑制している¹²。Lunaは推論エフォートを深く設定して思考時間を十分に担保することにより、極めて高い投資対効果を創出する特異なアーキテクチャ特性を有している²。

初期の評判と競合分析

GPT-6 SolおよびLunaの公開と完全に時を同じくして、競合であるAnthropicが「Claude Opus 5.5」を公開したことにより、技術コミュニティおよびエンタープライズ市場では熾烈な比較検証が行われた¹³。

Claude Opus 5.5との実務的対比

同日発表となった両モデルは、アプローチと強みの所在において明確なコントラストを形成している¹⁷。独立系ベンチマーク機関Artificial Analysisによる包括的知能指数 (Intelligence Index) では、Claude Opus 5.5が58点を獲得し、GPT-6 Solの48点を10ポイント引き離れた¹⁷。また、ZapierによるAutomationBenchの直接検証においても、Opus 5.5が40.0%を記録し、Solの33.2%に対して優位を示している¹²。

特にフロントエンド開発、SVGや3D WebGLアニメーションの描画、複雑なGUIデザインの実装など、

出力物の美的完成度や抽象的な指示の文脈解釈が求められる領域では、Opus 5.5の洗練された出力が開発者から圧倒的な支持を集めている¹⁸。

しかし、運用の経済性という観点では力関係が反転する¹⁷。Artificial Analysisの同一タスク群における平均実行コストを比較すると、Opus 5.5が1タスクあたり

5.98を要したのに対し、GPT-6 Solは 1.06と約5.6分の1にとどまった¹⁷。Opus 5.5は深く思考して高品質な解答を生成する反面、平均出力トークン数が約11.9万トークンに肥大化しやすく、API課金を急激に圧迫する傾向がある¹⁸。対照的にGPT-6 Solは、回答スタイルの簡素化に伴い、無駄な思考ループや冗長なテキスト生成を回避して迅速に結論を出力するため、大量処理におけるコスト予見性が極めて高い²。

比較項目	GPT-6 Sol	Claude Opus 5.5
API基本単価 (入力 / 出力 1M)	\$2.00 / \$10.00 ¹⁰	\$4.00 / \$20.00 ¹⁷
キャッシュ読取単価 (1M)	\$0.20 ²	\$0.20 ¹⁹
Artificial Analysis Intelligence Index	48 点 ¹⁷	58 点 ¹⁷
AutomationBench スコア	33.2% (xhigh) ²	40.0% (max) ¹²
FrontierCode 1.1 スコア	49.3% (max) ¹²	54.4% (max) ¹²
1タスクあたり平均コスト (AA 実測)	\$1.06 ¹⁷	\$5.98 ¹⁷
強みとなる領域	API・バックエンド実装、反復修正、定型テスト作成 ¹⁷	アーキテクチャ設計、大規模コード移行、UI/UX ¹⁷

開発現場の反響と機能分担

開発現場からは、SolおよびLunaの価格破壊力に対して概ね好意的な評価が寄せられている¹⁸。CodexometerベンチマークにおいてGPT-6 Lunaが高速かつ極小コスト(1課題あたり約\$0.002)で依存関係の解決タスクを完走した事例に代表されるように、これまで上位モデルのコストが障壁となっていた自動化処理への導入が急速に進んでいる²¹。メディア「Every」の最高経営責任者であるダン・シッパー (Dan Shipper) 氏は、Solを「Codexにおける新たな日常の主力機 (Daily Driver)」と形容し、最上位Astraの性能の大半を5分の1の価格で享受できる点を高く評価した²⁰。

一方で、開発者の間では単一モデルに依存するのではなく、課題の性質に応じたオーケストレーション体制が模索されている¹⁷。システム全体の設計や境界条件の定義、難解なバグの根本原因究明

などの「司令塔」的役割にはClaude Opus 5.5を充当し、確定した仕様に基づく詳細コードの実装、単体テストの量産、反復的なレビュー作業などの「実働部隊」にはGPT-6 Solを割り当て、さらにログ解析や定型的なルーティングにはGPT-6 Lunaを充てるという3段構えの運用が定着し始めている¹⁷。

技術的制約、安全性評価、および懸念事項

GPT-6 SolおよびLunaは高い実用性を誇る一方で、詳細な検証を通じていくつかの技術的制約や懸念点が明らかになっている¹²。

ピーク性能における一部後退

最も議論を呼んでいるのが、特定の複雑タスクにおける絶対性能の推移である¹²。実コードベースのソフトウェア開発を評価するDeepSWE v1.1において、GPT-6 Solの最高スコア(max設定で68.8%)は、前世代であるGPT-5.6 Solの最高スコア(72.7%)を3.9ポイント下回っている¹²。コンピュータ操作を評価するOSWorld 2.0においても、前世代の66.2%に対してGPT-6 Solは64.4%にとどまる¹²。OpenAIの公式発表では、Claude Fable 5の最高スコア(69.9%)と僅差(1.1ポイント差)であることが強調されていたが、前世代の同格モデル自身が保持していたピークスコアからの後退については積極的な説明がなされなかった²。知能の天井を高めることよりも、タスク実行単価を80%以上削減することに主眼を置いたモデル調整が行われた結果、極限的な難問に対する突破力においては前世代フラッグシップに及ばない領域が残されている¹²。

幻覚率低減の構造的要因

客観的なベンチマーク測定において、GPT-6 Solの事実誤認率の低下メカニズムに批判的な分析がなされている²⁰。Artificial Analysisの評価(AA-Omniscience)によると、Solの幻覚率は92%から60%へと大幅に改善された²⁰。しかし、同時にモデルが質問に対して回答を試みた割合(Attempt Rate)は99%から83%へと低下しており、全体の正答率自体は59%から54%へと微減している²⁰。これは、モデルが知識を正確に修正したのみならず、確信度が低い領域において回答を差し控える「沈黙による過誤回避(少说少错)」を選択する傾向が強まったことを意味する²⁰。知識労働の現場においては、要約やドラフト作成において細部の記述が脱落し、結果として業務ループブリックの要件を満たせなくなるリスクが指摘されている²⁰。

挙動の不安定さと安全性評価

OpenAI開発者コミュニティでは、GPT-6 Solが外部ツール呼び出しを行う際に、状況確認と探索のループから抜け出せなくなり、実質的に軽量モデルのLunaと同等の挙動に陥る事象が報告されている²²。

安全性に関して、システムカード追補の報告によると、両モデルのサイバーセキュリティ能力評価は「High」に留まり、危険水準である「Critical」には達していない⁴。Solのサイバー能力はGPT-5.6 Solと同等水準であり、目立った脅威の拡大は認められていない³。生物・化学兵器関連の脅威防御も「High」を維持し、モデルの自律的自己改善リスクは「High」未満と判定されている⁴。脱獄(Jailbreak)耐性は前世代から明確に向上したが、Lunaにおいては正当な技術的指示をも拒絶してしまう過剰拒絶(False Refusals)が一部で見受けられる⁴。また、医療評価(HealthBench)では、回答の簡潔化が進んだ影響により、複雑な症例分析において詳細な留意事項が省略され、スコアが低下する副作用が確認されている⁴。

結論と運用の展望

2026年9月22日に公開されたGPT-6 SolおよびGPT-6 Lunaは、モデルの絶対的な知能フロンティア

を極限まで押し広げることではなく、最高峰の知能技術をいかに経済的に持続可能な形で大量消費できるインフラへと転換するかという、生成AI産業の構造転換を明確に提示した¹。

API単価の50%以上の引き下げ、90%割引のプロンプトキャッシュ、そして柔軟な推論エフォート制御の組み合わせにより、GPT-6 Solは本格的な長時間コーディングエージェントを採算ベースに乗せる主力エンジンとしての地位を確立した²。一方のGPT-6 Lunaは、適切な思考深度を与えることで前世代中核モデルに迫る性能を極小コストで引き出せるため、高スループットな分散処理や定型タスクの自動化基盤として不可欠な選択肢となっている²。

今後のシステム構築においては、Claude Opus 5.5のような最高峰モデルを全体統括やアーキテクチャ設計などの戦略的レイヤーに配置し、具体的実装、テスト、データ変換などの実行レイヤーをGPT-6 SolおよびLunaに分散委託する多層的なアーキテクチャが、最も堅牢かつ費用対効果の高い実務標準となる¹⁷。

引用文献

1. GPT-6 SolとLunaを徹底解説: 価格と活用方法 - Eigent AI, <https://www.eigent.ai/ja/blog/gpt-6-sol-luna>
2. GPT-6 Sol・GPT-6 Luna の概要 | npaka - note, <https://note.com/npaka/n/n6ab510c32a38>
3. GPT-6 Sol」と「GPT-6 Luna」公開 API料金は前世代の半額に, <https://www.itmedia.co.jp/news/article/2609/23/2000001674/>
4. 速度と費用効率を重視, <https://gihyo.jp/article/2026/09/gpt-6-sol-and-luna>
5. 【徹底解説】OpenAI「GPT-5.6」登場。Sol・Terra・Lunaの実力を, <https://chatgpt-lab.com/n/n06afcb90c286>
6. GPT-5.6 Explained: Sol, Terra, Luna & AI Visibility - GEO Toolbox, <https://geotoolbox.ai/blog/gpt-5-6>
7. GPT-5.6 - Wikipedia, <https://en.wikipedia.org/wiki/GPT-5.6>
8. Clarification Needed: GPT-6 Astra Was Announced for “All ChatGPT”, <https://community.openai.com/t/clarification-needed-gpt-6-astra-was-announced-for-all-chatgpt-plus-users-but-plus-access-is-currently-limited-to-work-code-x/1395038>
9. GPT-6 Astra: 新世代の応答性能 - OpenAI, <https://openai.com/ja-JP/index/gpt-6-astra/>
10. Introducing GPT-6 Sol and Luna - OpenAI, <https://openai.com/index/introducing-gpt-6-sol-and-luna/>
11. GPT-6 SolとLunaの違い | 料金・性能・仕事別の選び方【2026年9月】, <https://uravation.com/media/gpt-6-sol-luna-difference-pricing-guide-2026/>
12. GPT-6 Sol and GPT-6 Luna: Specs, Benchmarks, Pricing and How, <https://kingy.ai/blog/gpt-6-sol-luna-specs-benchmarks-pricing-comparison/>
13. OpenAI、「GPT-6 Sol」「Luna」公開 Astra並の性能でコストはGPT-5.6の半額, <https://www.watch.impress.co.jp/docs/news/2142636.html>
14. GPT-6 Sol and Luna: OpenAI's Cheaper Coding and Agentic Models, <https://www.eigent.ai/blog/gpt-6-sol-luna>
15. GPT-6 Luna / Sol と GPT-5.6 の違い: 推論能力・プロンプト ... - Qiita, <https://qiita.com/tomokusaba/items/033439c3a7922dcada8b>

16. GPT-6 Sol (batch) - API Pricing & Benchmarks - OpenRouter,
<https://openrouter.ai/openai/gpt-6-sol:batch>
17. Claude Opus 5.5 と GPT-6 Sol: どちらが貴社にとってより優れて,
<https://www.unite.ai/ja/claude-opus-5-5-vs-gpt-6-sol-which-is-better-for-your-business/>
18. Claude Opus 5.5 vs GPT-6 Sol: 同日リリースの2モデルを比べる,
https://note.com/it_navi/n/n2f8c659e891e
19. 【特別版 | AI実務ガイド】GPT-6 Sol・Luna vs Claude Opus 5.5,
<https://note.com/mitarashi440/n/n8e8759fe2feb>
20. 安くなったのか GPT-6 SolとOpus 5.5、同じ日に値札を下げた二本,
https://note.com/regal_emu393/n/n529cf307e635
21. <https://community.openai.com/t/announcing-gpt-6-sol-and-gpt-6-luna-in-the-api-codex-and-chatgpt/1399925>
22. <https://community.openai.com/t/gpt-6-sol-behaving-more-like-luna/1399997>