

Grok 4.5の評判・評価と今後への影響

エグゼクティブサマリー

Grok 4.5は、2026年7月8日にSpaceXAIが公開した、コーディング・エージェント作業・知識労働に特化した最新主力モデルです。公式には「最も賢い」Grokとして位置付けられ、Cursorと共同で訓練され、数万基規模のNVIDIA GB300で学習されたと説明されています。APIでは `grok-4.5` として提供され、500Kコンテキスト、画像入力対応、構造化出力と関数呼び出し対応、標準価格は入力\$2／出力\$6 per 1M tokens、200K超の長文入力には割増が適用されます。検索を有効化しない限り最新時事に強いとまでは言えず、企業向けには「学習不使用」「30日保持」「ZDRなら保持ゼロ」というデータ取扱いも明示されています。 ¹

第三者評価では、Artificial AnalysisがGrok 4.5を総合「Intelligence Index」4位、Coding Agent IndexではGPT-5.5級で、かつ低コストと評価しました。SnorkelのGDPval+でも、Grok 4.5はGPT-5.5とClaude Opus 4.8を上回る総合成績を示しています。要するに、**絶対性能で独走するモデル**というより、**最先端に肉薄する水準をかなり安く、しかもエージェント実務に寄せて供給するモデル**として市場に強い衝撃を与えています。 ²

ただし、懸念も明確です。Cursor自身が、CursorBenchについて「Cursorコードベースの過去スナップショットが誤って学習に含まれたため有利になった可能性」を注記しており、ベンチマーク完全性には注意が必要です。さらに、xAIの安全性ページでは「モデルカードを公開している」としつつ、2026年7月11日時点で公開一覧にGrok 4.5は見当たらず、少なくとも公開ベースの安全性透明性は競合上位陣より弱く見えます。加えて、Oxford Internet Instituteらが紹介した最新研究は、AI補助の文章編集・説明機能が人間の意図や政治的ニュアンスを静かにずらしうと警告しており、Grok系機能もその文脈で再検討が必要です。 ³

本報告の結論は明快です。**Grok 4.5の本質は「性能競争の勝者」よりも「価格・速度・エージェント適性を組み合わせて競争地図を描き替える存在」**にあります。短期ではコーディング支出の見直しとモデルルーティング再設計を促し、中期ではCursorを核とした開発者エコシステム強化、長期では「実務ログを持つソフトウェア企業がモデル競争を有利に進める」構図を強める可能性が高いです。他方で、規制産業や日本語重視の業務では、現時点で単独標準化するより、評価付きパイロット導入とモデル併用戦略が妥当です。 ⁴

公式情報と製品整理

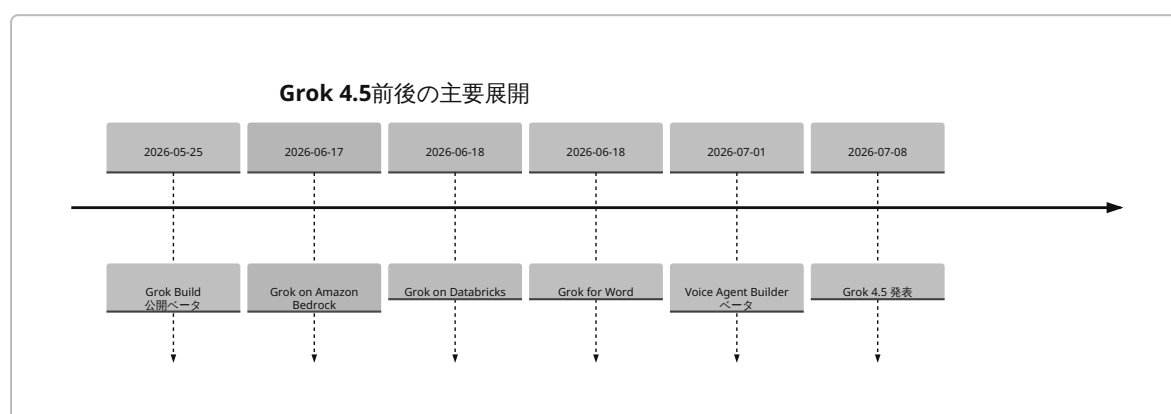
SpaceXAIの公式リリースによれば、Grok 4.5は「coding, agentic tasks, knowledge work」向けの最上位モデルとして2026年7月8日に発表されました。リリース本文では、DeepSWE、SWE Marathon、Terminal Bench、SWE Bench Proなど、ソフトウェア工学寄りの実務ベンチマークを重視している点が特徴的で、研究用汎用ベンチよりも「実務でどこまで解決できるか」を訴求しています。さらに、Cursor側はGrok 4.5を「ソフトウェア工学だけでなく、データサイエンス、金融、法務を含むコンピュータ上の知識労働向けに広げた最初のモデル」と位置付けています。 ⁵

公開ドキュメント上の主要仕様は、500Kコンテキスト、画像入力対応、テキスト出力、JSON mode・function calling・structured outputs対応、知識カットオフは2025年7月、標準価格は入力\$2／出力\$6 per 1M tokens、キャッシュ入力\$0.5で、200K超の長い入力には割増がかかる、というものです。また「current events」に関しては、モデル単体では最新知識を保証せず、`web_search` や `x_search` を有効化する必要があると明記されています。これは、Grokが「Xと連携するから常に最新」という一般的イメージと違い、**検索有効化前提の鮮度設計**だと読むべきです。 ⁶

利用面では、Grok 4.5はGrok Buildのデフォルトモデル、Cursorの全プラン、SpaceXAI console/API、SuperGrok月額\$30のBusiness系プランで提供されます。企業向けにはRBAC、コネクタ、監査制御、SCIM、SSO、カスタム保持、CMK、専用データプレーンなどが掲げられており、単なる消費者向けチャットではなく、明確にエンタープライズ導入を意識した商品設計です。APIデータについては、明示的許可なしに学習されず、標準では30日保持、Enterprise向けZDRでは保持ゼロです。 7

2026年春から夏にかけての周辺機能公告を見ると、SpaceXAIはGrok Build、Amazon Bedrock、Databricks Agent Bricks、Wordアドイン、Voice Agent Builderなど、**Grok 4.5単体ではなく「仕事の実行面に埋め込む導線」**を先に並べてきました。これは、モデル単体性能を超えて、現場ツールへの接続性で採用障壁を下げる戦略と解釈できます。 8

以下のタイムラインは、Grok 4.5が単発リリースではなく、5月以降の開発者・業務統合施策の上に置かれたことを示します。 9



ビジュアル挿入案: ここに「xAI公式のGrok 4.5リリースページ冒頭」「xAI Docsのモデル仕様ページ」「価格ページ」のスクリーンショットを並べると、製品変更の確認が視覚的に容易になります。

技術性能と主要モデル比較

Grok 4.5の公式ベンチマーク訴求はかなり実務寄りです。xAIの公開値では、DeepSWE 1.0で62.0%、DeepSWE 1.1で53%、SWE Marathonで29.0%、Terminal Bench 2.1で83.3%、SWE Bench Proで64.7%を記録しました。これらはFable 5やGPT-5.5を全面的に上回るわけではないものの、SWE Marathonのような長時間解決率では優位があり、Terminal Benchではほぼ最上位群に食い込んでいます。つまりGrok 4.5は、**派手な一発スコアより、ツールを使いながら粘り強く仕事を完遂するタイプ**として設計されていると読むのが妥当です。 10

速度とコストでは、公式は「80 TPS」、Artificial AnalysisはSpaceXAI API実測で93.2 tokens/sを報告しています。Allの総合指数では54点で4位、Coding Agent Indexでは76点、Intelligence Index 1タスク当たり\$0.31、Coding Agent Index 1タスク当たり約\$2.49とされ、近接する最先端モデル群の中では費用対効果はかなり強い部類です。SnorkelのGDPval+でも、Grok 4.5は約2,000件の専門職ワークフロー問題で29%の通過率を示し、GPT-5.5の22%、Opus 4.8の21%を上回りました。 11

他方、モデルの内部構造については、Cursorが「mixture-of-experts」と公式に述べている一方、**公式のパラメータ数は今回参照した公開ドキュメントでは未指定**です。Artificial AnalysisはMusk開示として「前世代の3倍、1.5T parameters」と紹介していますが、xAIのモデルページ自体に明記された数値ではありません。このため、本報告ではモデルサイズについては「外部報道・分析では1.5T説が流通、公式公開資料では未指定」と整理します。 12

また、CursorはGrok 4.5がCursorデータの数兆トークンを含む共同訓練であることを明かしており、これはコーディング・エージェント用途には強い追い風です。他方で、長いコンテキストについてはGrok 4.3の1MからGrok 4.5の500Kへ逆に縮小しており、巨大RAGや超長文法務レビューでは、必ずしも万能な前進とは言えません。つまりGrok 4.5は「何でも最大コンテキスト」ではなく、**エージェント仕事のコスト効率に最適化した設計**と見るべきです。 ¹³

主要モデル比較

モデル	主な位置づけ	コンテキスト	価格 per 1M tokens	推論	入力/出力	独立指標	独立速度	備考
Grok 4.5	コーディング・エージェント・知識労働向け旗艦	500K	入力\$2 / 出力\$6 / キャッシュ\$0.5	あり	テキスト・画像 / テキスト	AA 指数 54	93.2 t/s	200K超入力で割増。公式は80 TPS訴求。 ¹⁴
GPT-4.1	非推論系の高性能汎用API	1M	入力\$2 / 出力\$8 / キャッシュ\$0.5	なし	テキスト・画像 / テキスト	AA 指数 19	103.8 t/s	instruction followingとtool calling重視。 ¹⁵
GPT-4o	高速な汎用マルチモーダル旗艦	128K	入力\$2.5 / 出力\$10 / キャッシュ\$1.25	なし	テキスト・画像 / テキスト	AA 指数 11	193.9 t/s	速度は非常に高いが、独立総合指数は相対的に低い。 ¹⁶
Claude 3.5 Sonnet	2024年時点の中核 Claude 3系	200K	入力\$3 / 出力\$15	なし	画像入力 / テキスト	AA 指数 10	未指定	公式にはClaude 3 Opusの約2倍速。2026年主力ではなく旧世代寄り。 ¹⁷
Gemini 2.5 Pro	複雑推論・コーディング向け Google上位機	1M	入力\$1.25 / 出力\$10	あり	テキスト・画像・音声・動画 / テキスト	AA 指数 26	131.4 t/s	200K超で単価上昇。マルチモーダル入力が広い。 ¹⁸

この表から見てくるのは、Grok 4.5が**GPT-4.1より出力単価が25%安く、GPT-4oより入力単価が20%・出力単価が40%安く、Claude 3.5 Sonnetよりさらに大幅に安い**ことです。一方で、コンテキスト長ではGPT-4.1やGemini 2.5 Proに劣り、純粋な高速生成ではGPT-4oに及びません。したがって、Grok 4.5の強みは「最長文」「最速応答」ではなく、**開発作業の完了コストとエージェント性能の釣り合い**にあります。 ¹⁹

なお、2026年7月時点の“実際の競争相手”は、ユーザー指定のGPT-4.1/4o/Claude 3系よりも、むしろFable 5、GPT-5.5、Claude Opus 4.8といった新しい最上位モデルです。xAI自身が比較対象としてそれらを並べ、Artificial AnalysisもGrok 4.5を「Fable 5、GPT-5.5、Opus 4.8に次ぐ4位」と位置付けています。よって、**Grok 4.5は旧世代比較ではなく、現行フロンティアの廉価・高効率プレイヤー**として読む方が正確です。

ビジュアル挿入案: ここに「xAI公式ベンチマーク棒グラフ」と「Artificial Analysisの価格対知能散布図」のスクリーンショットを並置すると、「絶対性能では首位ではないが、費用対効果では非常に強い」という構図が一目で分かります。

第三者レビューとユーザー反応

主要メディアは、Grok 4.5を「xAIがようやくOpenAI・Anthropic・Googleとの企業向け競争へ本格復帰した象徴」と捉えています。Reutersは新モデルをコーディングとエージェント作業向けの最先端投入と整理し、Axiosは「消費者向け雑談チャットよりもコーディング/agentic work寄り」である点を強調しました。TechCrunchも、xAIの訴求点は「Opus級だがより速く、よりトークン効率が高く、より安い」という組み合わせにあると要約しています。 ²¹

独立分析では、Artificial Analysisの評価が最も重い材料です。AAはGrok 4.5を総合4位とし、Coding Agent IndexではGPT-5.5と同等水準、しかもコーディング・エージェント指数1タスク当たり約\$2.49で、Fable 5の約\$11.80、GPT-5.5の約\$5.07よりかなり安いと報告しました。この点は、ベンチマークで1位を取り切れなくても、企業のAPI請求書ベースでは十分に魅力的な選択肢になることを意味します。 ²²

SnorkelのGDPval+も興味深く、法律・教育・医療・QA分析など専門判断タスクを含む約2,000問で、Grok 4.5が総合29%通過率となり、GPT-5.5の22%、Opus 4.8の21%を上回りました。これは、Grok 4.5が“純コーディングだけの特化機”というより、**検証可能な専門業務をかなり広い範囲で処理しう**ることを示唆します。ただしこの評価も公開直後の単独ベンダー実施であり、再現の積み上がりは今後を待つべきです。 ²³

ユーザー反応は、おおむね二極化しています。RedditやHacker Newsでは、「速い・安い・コーディングで意外に強い」という称賛が多い一方で、「価格の安さは補助金的ではないか」「ベンチマークが過大評価ではないか」という警戒も強く見られます。特にCursorBench汚染注記は、コミュニティでかなり敏感に受け止められています。 ²⁴

日本語圏の反応もまだ初動段階ですが、方向性は似ています。日本の技術メディアは「Claude Opus級で安価」「数万GB300、Cursor由来のデータ、実務特化RL」を強調して報じました。他方、X上の日本語投稿では「日本語文章生成はClaudeやChatGPTの方が良いが、コーディングはかなり速い」という声が見られ、noteの記事では逆に「日本語も不自然ではなかった」という実地評価もあります。つまり、**日本語品質についてはまだ評価が割れており、特に文書作成用途では企業ごとの自前評価が必須**です。 ²⁵

総合すると、第三者レビューとユーザー反応は「Grok 4.5はxAI初の“本当に検討対象になる企業モデル”だが、ランキング首位と断言するには早い」という一点でほぼ一致しています。称賛の中心はコスト効率と実務感、懸念の中心はベンチマーク信頼性と安全性透明性です。 ²⁶

倫理・安全性・ガバナンス

安全性・透明性の観点では、Grok 4.5は**性能ほどには説明責任が追いついていない**というのが率直な評価です。xAIの安全性ページは「各フロンティアモデルについてモデルカードと安全性評価を公開する」と述べていますが、2026年7月11日時点で公開一覧にあるのはGrok 4 Fast、Grok Code Fast 1、Grok 4であり、Grok 4.5は見当たりません。少なくとも一般公開文書ベースでは、4.5の安全性評価・限界・推奨用途の体系的説明は未提示です。これは、同世代の大手競合がモデルカードや安全性付録を比較的早く出す流れと比べると弱点です。 ²⁷

一方で、利用条件そのものはかなり明確です。2026年6月26日改定のAUPは、違法行為、プロンプトインジェクション、ウォーターマーク除去、出力の再販・蒸留、競合AI開発用途などを禁止しています。Enterprise Termsはベンチマーク行為そのものも禁止条項に含めており、開発者の自由度よりもプラットフォーム統制を

優先する設計です。これは悪用抑止には資する一方、外部監査や再現性検証をやややりにくくする面もあります。 28

データガバナンス面では、xAIはAPI入出力を明示的許可なしに学習へ使わないとし、標準保存30日、Enterprise向けにZDRを提供しています。Business/Enterpriseページでも“No training”を掲げています。したがって、**データ保持ポリシー自体は競争力がある**と言えますが、これはあくまで契約的約束であり、どのような監査証跡・第三者検証があるかまでは公開資料からは十分分かりません。 29

実際のリスクとしては三つあります。第一に、最新情報リスクです。公式Docsは、時事情報には `web_search` や `x_search` を有効化せよと明記しており、検索なしの回答は鮮度に限界があります。第二に、評価信頼性リスクです。CursorはCursorBenchについて学習データ混入の可能性を公開し、その影響は不明と認めました。第三に、社会的バイアス・媒介影響リスクです。Oxford Internet Instituteらの研究紹介では、LLMがユーザー文面の意味や争点方向を静かに変えてしまい、Grok系機能は右寄り方向の偏りを示すケースがあったと報告されています。 30

以上を踏まえると、Grok 4.5は「危険だから使うな」ではなく、「**企業自身の安全評価を前提に使うべきモデル**」です。Microsoft Foundry向けドキュメント自体も、Azure Content Safety統合、自前レッドチームング、使用監視を本番前提条件として推奨しています。xAI自身が「プラットフォーム側の安全対策だけでは足りない」と事実上認めている形であり、導入企業はそれを素直に受け止めるべきです。 31

経済・産業影響

Grok 4.5の最大の経済的意味は、**エージェント型ソフトウェア開発の価格破壊圧力**です。Artificial Analysisは、Grok 4.5が近接する最先端モデルより大幅に安い価格帯で近い性能を出していると評価しており、特にCoding Agent Indexでの1タスク当たりコストは、Fable 5やGPT-5.5を大きく下回りました。これは、IDE内エージェント、CIレビュー、障害調査、テスト生成といった“たくさん回す”用途ほど効きます。単発の高難度推論ではなく、「一日に何千回も回るAI労働」を安くする点が市場への効き目です。 22

SaaSと開発者エコシステムへの影響も大きいです。Grok 4.5はGrok Buildのデフォルト、Cursor全プラン、Warp、Wordアドイン、Amazon Bedrock、Databricks Agent Bricks、Google Cloud Vertex、Microsoft Foundryといった複数の流通路を持ちます。つまり、モデルを単体APIとして売るだけでなく、**IDE、クラウド、オフィス、音声、企業エージェント基盤へ同時に差し込む「面の戦略」**が進んでいます。導入障壁の低さは、単純なベンチ順位以上に強い競争優位です。 32

検索・知識労働市場では、Grok 4.5単体よりも、Grok BuildやWord/PowerPoint、Voice Agent Builder、Vapi連携などの周辺商品が効いてきます。特にVapiでは2.5M超のボイスエージェントにGrok Voiceが既定エンジンとして入るとされており、音声カスタマーサポートや業務オペレーションにも影響が広がっています。Grok 4.5自体は主にコード・知識作業向けですが、xAIは明らかに「**横断的なAI作業OS**」の方向へ進んでいます。 33

一方で、規制産業や長文文書集約業務では慎重さが必要です。Grok 4.5は500Kコンテキストで十分大きいものの、1Mを標準提供するGPT-4.1やGemini 2.5 Proと比べると、超長文RAGや大規模法務文書一括処理では不利になる可能性があります。加えて、安全性文書の公開遅れは、金融・医療・行政の調達審査で不利に働きます。よって業界インパクトは、「まず開発者向け高頻度ワークロードから侵食し、規制産業は後から追随」という順番になる公算が高いです。 34

産業構造の観点では、Grok 4.5は「モデル会社が最終製品を作る」だけでなく、「製品会社が持つ実務ログでモデルを鍛える」循環の強さを示しました。Cursorは数兆トークンの開発者・エージェント相互作用データを訓練に使ったと明かしており、これは汎用Webデータ中心の学習より、実務環境での道具使用・試行錯誤・修復行動を学ばせやすいはず。今後は、IDE、CRM、ERP、サポートデスクなど“仕事の現場ログ”を

握る企業が、モデル性能競争でも優位に立つ流れが強まるでしょう。これはオープンな研究競争より、統合型企業競争を強める変化です。 35

将来予測と提言

短期では、Grok 4.5は企業のモデル選定基準を「ベンチマーク順位」から「完了タスク当たり総コスト」へ強く移す可能性があります。特にコーディング、コードレビュー、障害解析、社内文書下書きのように反復回数が多い業務では、Grok 4.5の費用対効果は無視しにくいです。ただし、日本語品質、安全性透明性、規制対応で未確定要素が残るため、短期の主流導入形態は**全面置換ではなく、モデルルータの一候補としての採用**になる可能性が高いです。 36

中期では、xAIがCursor連携をさらに深め、Bedrock・Databricks・Foundry・Word/Voice系を通じて「開発と業務の両面に埋め込まれたプラットフォーム」になるかが焦点です。もしGrok 4.5級のモデル改善が継続し、かつ公開安全文書と多言語品質が強化されれば、OpenAI・Anthropic・Googleに対する**第三極の企業AI供給者**としての地位を固める可能性があります。反対に、透明性不足と政治的・社会的論争が続けば、性能が高くても大企業の標準調達に入りにくいまま留まるでしょう。 37

長期では、Grok 4.5の本当の影響はモデルそのものより、「アプリから収集される実務ログを用いた閉ループ学習」が市場標準になることにあります。IDEやオフィスソフト、音声エージェント、サポート基盤などで生成される操作履歴・修正履歴・承認履歴が、次世代モデルの主要燃料になるなら、AI産業はますます“データを持つアプリ企業が強い”構造になります。これは、オープンベンチの価値を相対的に下げ、データガバナンスと学習データ由来の競争優位をめぐる規制論を強めるはずです。 38

企業向け提言としては、第一に日本語・セキュリティ・規制要件を含む自社評価セットでPoCを行うこと、第二にGrok 4.5をコーディングやバックオフィス草稿に限定して試し、法務・IR・対外広報では人手承認を堅持すること、第三に `web_search` 有無で回答鮮度が変わる前提をUIと監査ログに明示することが重要です。さらに、ZDRや保持設定、コンテンツ安全層を環境別に分けて設計すべきです。 39

研究者向け提言としては、Grok 4.5を「英語コーディングで強い」だけで終わらせず、ベンチ汚染監査、多言語エージェント評価、長時間作業の失敗様式、政治・社会的媒介バイアスの再現研究へ拡張すべきです。特にCursorBenchのような事例は、今後の評価設計で**学習データ由来の優位がどこまで混入しうるか**を定量化する必要性を示しています。 40

政策立案者向け提言としては、フロンティアモデル公開時のモデルカード提出、ベンチマーク汚染の開示義務、AI媒介コミュニケーション機能への説明ラベル、データ保持と学習不使用ポリシーの標準化が必要です。Grok 4.5は、競争政策よりむしろ**透明性・監査可能性・媒体上の認知操作リスク**という新しい政策論点を前面化させた事例と見るべきです。 41

主要一次情報URL

種別	URL
SpaceXAI リリース	https://x.ai/news/grok-4-5
SpaceXAI モデル仕様	https://docs.x.ai/developers/models/grok-4.5
SpaceXAI 価格	https://docs.x.ai/docs/models/overview
xAI API セキュリティFAQ	https://docs.x.ai/developers/faq/security
xAI AUP	https://x.ai/legal/acceptable-use-policy

種別	URL
Cursor 共同発表	https://cursor.com/blog/grok-4-5
OpenAI GPT-4.1 モデル docs	https://developers.openai.com/api/docs/models/gpt-4.1
OpenAI GPT-4o モデル docs	https://developers.openai.com/api/docs/models/gpt-4o
Anthropic Claude 3.5 Sonnet 発表	https://www.anthropic.com/news/claude-3-5-sonnet
Google Gemini API pricing	https://ai.google.dev/gemini-api/docs/pricing

🔍navlist🔍最近の関連報道🔍turn15news33,turn18news29,turn15news31,turn18news28🔍

1 5 10 20 Introducing Grok 4.5 | SpaceXAI

<https://x.ai/news/grok-4-5>

2 4 22 36 Grok 4.5 brings SpaceXAI to the intelligence frontier

<https://artificialanalysis.ai/articles/grok-4-5-brings-spacexai-to-the-the-intelligence-frontier>

3 40 Introducing Grok 4.5

https://cursor.com/blog/grok-4-5?utm_source=chatgpt.com

6 14 34 grok-4.5 | SpaceXAI Docs

<https://docs.x.ai/developers/grok-4-5>

7 Pricing: Compare Grok Plans | SpaceXAI

https://x.ai/pricing?utm_source=chatgpt.com

8 9 32 Introducing Grok Build | SpaceXAI

https://x.ai/news/grok-build-cli?utm_source=chatgpt.com

11 Grok 4.5 (high) - Intelligence, Performance & Price Analysis

<https://artificialanalysis.ai/models/grok-4-5>

12 13 35 38 Introducing Grok 4.5 · Cursor

<https://cursor.com/blog/grok-4-5>

15 19 GPT-4.1 Model | OpenAI API

<https://developers.openai.com/api/docs/models/gpt-4.1>

16 GPT-4o Model | OpenAI API

<https://developers.openai.com/api/docs/models/gpt-4o>

17 Introducing Claude 3.5 Sonnet \ Anthropic

<https://www.anthropic.com/news/claude-3-5-sonnet>

18 Gemini Developer API pricing | Gemini API | Google AI for Developers

<https://ai.google.dev/gemini-api/docs/pricing>

21 SpaceXAI launches Grok 4.5 model for coding, agentic tasks

https://www.reuters.com/business/media-telecom/spacexai-launches-grok-45-model-coding-agentic-tasks-2026-07-08/?utm_source=chatgpt.com

23 Grok 4.5 Benchmark Results vs GPT 5.5 & Claude Opus 4.8

<https://snorkel.ai/blog/grok-4-5-testing-results-how-spacexais-new-model-performs-on-real-professional-work/>

- 24 **Introducing Grok 4.5 : r/singularity**
https://www.reddit.com/r/singularity/comments/1ur0ye6/introducing_grok_45/?utm_source=chatgpt.com
- 25 「Grok 4.5」が登場、Claude Opus 4.8と同等性能で料金は安価
https://gigazine.net/news/20260709-grok-4-5-ai/?utm_source=chatgpt.com
- 26 **Grok 4.5 Is SpaceXAI's First Real Entry Into the Enterprise**
https://aibusiness.com/agent-ai/grok-4-5-spacexai-s-first-real-entry-into-enterprise?utm_source=chatgpt.com
- 27 41 **Safety | SpaceXAI**
<https://x.ai/safety>
- 28 **Acceptable Use Policy | SpaceXAI**
<https://x.ai/legal/acceptable-use-policy>
- 29 **FAQ - xAI API Security | SpaceXAI Docs**
<https://docs.x.ai/developers/faq/security>
- 30 39 **Models | SpaceXAI Docs**
<https://docs.x.ai/developers/models>
- 31 **Microsoft Foundry | SpaceXAI Docs**
<https://docs.x.ai/developers/community/microsoft-foundry>
- 33 **Grok Becomes the Voice of Vapi | SpaceXAI**
https://x.ai/news/grok-vapi?utm_source=chatgpt.com
- 37 **Grok on Databricks | SpaceXAI**
https://x.ai/news/grok-databricks?utm_source=chatgpt.com