

Anthropic Claude Opus 4.5 総合レポート

要約

Anthropic社の最新AIモデル「Claude Opus 4.5」は2025年11月にリリースされ、コーディングやPC操作といった実務タスクへの対応力が飛躍的に向上しました¹。Opus 4.5は**従来モデル（Claude 4.1やSonnet 4.5）より高い知能と効率性**を備え、ソフトウェア開発のベンチマークで最高性能を記録するとともに、多言語コーディングやマルチステップの業務処理でも最先端の成果を上げています²³。新機能として**エージェントによるPCブラウザ操作**（Chrome拡張）や**Excel自動化**への対応が強化され、開発者向けには**手順計画モード（Plan Mode）**や**“effort”パラメータ**によるトークン消費制御などが導入されました⁴⁵。モデルの**安全性・アラインメントも過去最高**とされ、プロンプトインジェクション攻撃耐性が大幅に強化されています⁶。また、Opusクラスのモデルとしては**破格の低価格**（API利用料が従来の3分の1）で提供され、従来存在した利用制限も緩和されたため、企業や開発者は日常業務においてより気軽に高性能モデルを活用できるようになりました⁷⁸。コミュニティでも「**最強のコーディングAI**」との呼び声が高く、競合のOpenAIやGoogleの最新モデルと比べても、多くの専門家が**高度なコーディング・推論タスクにおけるトップランナー**として評価しています⁹¹⁰。

Claude Opus 4.5の新機能と性能向上ポイント

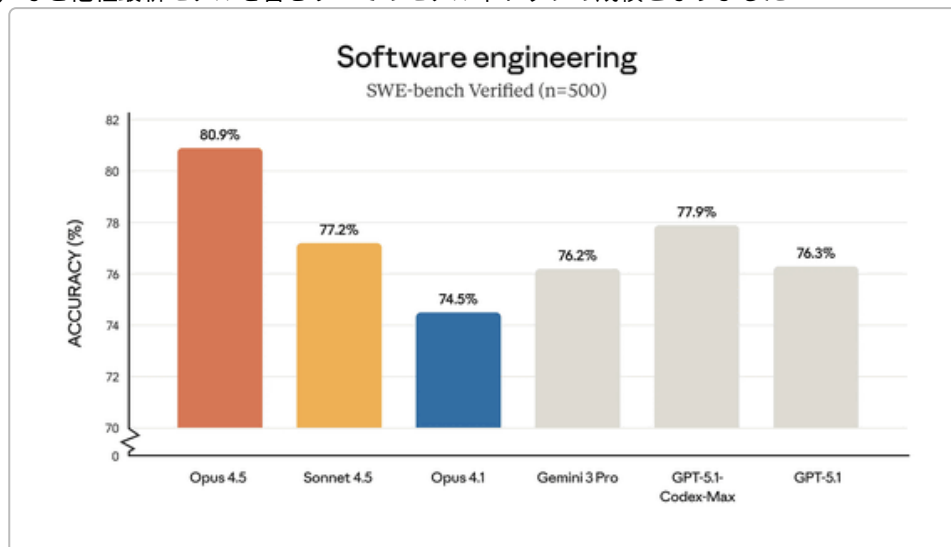
Claude Opus 4.5はAnthropicのClaude 4.5世代モデル群の最上位に位置付けられ、「**インテリジェントで効率的な最新モデル**」として2025年11月25日に提供が開始されました¹。このモデルは前世代に比べ**コーディングやPC操作、Deep Research（深い情報収集）やスプレッドシート操作**など日常的なタスク全般で性能が大幅に向上しています¹。AnthropicはOpus 4.5について「**コーディング、AIエージェント、コンピューター操作において世界最高レベルのモデル**」であり、「**スライドや表計算といった日常タスクでも意味のある改善を実現した**」と紹介しています¹¹。

具体的な性能指標として、Opus 4.5は**実世界のソフトウェアエンジニアリング試験で最先端の成績**を収めています。社内評価の難関コーディング試験（SWE-bench Verified）では**2時間の制限内で歴代の人間候補者を上回るスコア**を記録し（モデルは並列実行で複数回試行する手法を使用）¹²¹³、従来AIでは困難だった曖昧な要件の対処やトレードオフの推論、複数システムにまたがる複雑なバグ修正を自律的に行えることが確認されています¹⁴。ソフトウェア分野以外でも**視覚処理・推論・数学**といった能力が全般的に強化されており、Opus 4.5は様々な分野のベンチマークで最新のAIモデルにふさわしい高スコアを示しました¹⁵¹⁶。例えば、**エージェント型AIのコーディング能力**を測るTerminal-Bench 2.0では**59.3%**、**ツール使用スキル**を測るMCP Atlasで**62.3%**、**OS操作（PC操作）能力**を測るOSWorldで**66.3%**と、いずれも従来モデルを上回る高水準に達しています³。また抽象的な推論力指標ARC-AGI-2で**37.6%**、画像を含むマルチモーダル能力指標MMMUで**80.7%**を記録するなど、認知・理解面での強化も伺えます³。

コーディング能力の強化とエージェントによるPC操作

コーディング性能の向上はClaude Opus 4.5の大きな目玉です。多言語コード生成ベンチマークの**SWE-bench Multilingual**では、対象8言語中7言語でOpus 4.5が最高性能を示しました²。具体例として、C言語の問題ではOpus 4.5の正答率が約**83%**に達し、次点のSonnet 4.5（約74%）や前版Opus 4.1（約70%）を大きく引き離しました²。JavaでもOpus 4.5は約**90%**の高精度を記録し、Sonnet 4.5（約80%）・Opus 4.1（約70%）を凌駕しています²。さらに総合的なソフトウェアエンジニアリング評価であるSWE-bench

VerifiedでもOpus 4.5は**正解率80.9%**をマークし、OpenAIのGPT-5.1（約76%）やGoogleのGemini 3 Pro（約76%）など他社最新モデルを含むすべてのモデル中トップの成績となりました



。これらの結果から、**現時点で最も優れたコーディングAIの一つ**と評価されています。

Opus 4.5は**長時間・大規模なコーディングタスク**にも強く、複数日にわたる大型プロジェクトの計画・実行や、複雑なリファクタリングを自律的にこなす能力が報告されています^{17 18}。Anthropicの社内テストでは、Opus 4.5が**30分間の自律コーディングセッション**を通して安定した性能を維持し、最難関の評価項目でも従来モデルを上回る成果を上げたとの声もあります^{19 20}。特に、**コードの計画立案からバグ修正までを人手を介さず実行**できる点が強調されており、ユーザーの介入なしで高度なソフトウェア開発タスクを進める「シニアエンジニア並みの働き」が期待されています^{21 22}。実際、「これまで他モデルでは数時間かけても解決できなかった問題をOpus 4.5が一撃で解いてくれた」という開発者の報告もあり、その飛躍ぶりを「GPT-3.5が登場した時に匹敵する驚き」と表現する声もありました²³。

一方、**コンピューターの直接操作（Computer Use）機能**も大きく強化されています。Claude Opus 4.5はPC上でのエージェント動作において「**卓越したコンピューター使用能力**」を持つとされ²⁴、新たに**ズーム（zoom）アクション**が導入されました²⁴。これはエージェントが画面上の特定領域を高解像度で拡大表示し、細かいUI要素や小さい文字まで詳細に確認できる機能です²⁵。このズーム機能により、ボタンやリンクなど**細部の識別**、契約書等の**細字の読み取り**、情報密度の高いインターフェースの解析、アクション実行前の**視覚的な状況確認**が可能になりました²⁶。例えば、従来モデルではスクリーンショット上で判別が難しかった小さなアイコンやチェックボックスを、Opus 4.5はズームで正確に認識し適切な操作を行えるようになっています。

Anthropicはまた、ClaudeをPC上で幅広く活用するための周辺機能も拡充しています。公式のChrome拡張機能「**Claude for Chrome**」では、ブラウザのタブ全体にわたってAIエージェントが操作を行えるようになっており、Opus 4.5リリースに伴い**全てのMaxプランユーザーに開放**されました⁴。これにより、ウェブ検索やフォーム入力、オンラインでの予約・決済といった**ブラウザ上の定型業務を自動化**でき、AnthropicはChatGPTやPerplexityなど他社のエージェント機能に対抗する形で実用性を高めています^{27 28}。同様に、表計算ソフトの自動操作を可能にする「**Claude for Excel**」も2025年10月に発表され、Opus 4.5からそのβ版アクセスがMax・Team・Enterpriseユーザー全員に拡大されました⁴。これによってExcelシートの分析・編集といった**オフィス業務の自動化**も推進されています。実際、Anthropicは「Opus 4.5によりExcel自動化が一段と高度化し、初期ユーザー企業では社内評価で**精度20%向上、効率15%改善**といった成果が出ている」と述べています²⁹。以上のように、Opus 4.5は**プログラミングからPC日常操作まで幅広く“自分で考えて動く”AIエージェント**として進化を遂げています。

複雑な業務タスクへの適用事例

Claude Opus 4.5の強みは、**複雑なマルチステップ業務や高度な意思決定を要するタスク**に適用できる点にもあります。AnthropicはOpus 4.5を「これまでSonnet 4.5ではほぼ不可能だったタスクも遂行可能にしたモデル」であると位置づけています³⁰。例えば、**企業の内部データに基づく高度な分析**では、情報検索・ツール使用・深層解析を組み合わせたマルチステップ推論で最先端の結果を達成したと報告されています³¹。Enterpriseプラン環境下で、社内知識ベースやデータベースと接続したOpus 4.5を用いることで、契約書中の義務条項の分析や規制文書からのリスク抽出、研究論文の要約から洞察の構造化といった**高度なドキュメント理解タスク**を自動化できるようになります³²³³。Databricks社は自社プラットフォームへの統合に際し「Claude Opus 4.5は**オフィス業務全般とコンピューター操作で新たな標準を打ち立てた**」と評価しており、マーケティングキャンペーンの自律管理や部門横断ワークフローの自動実行など、**長時間にわたるマルチチャネル業務のエージェント代行**が現実味を帯びています¹⁸³⁴。

実例として、**AIエージェントに店舗経営を任せるシミュレーション (Vending-Bench)** では、Claude Opus 4.5が収益面でClaude Sonnet 4.5を**29%も上回る成果**を出しました³⁵。また、航空会社のカスタマーサポート業務を模した**t2-bench**というベンチマークでは、Opus 4.5は従来モデルにはない**創造的な問題解決策**を提示しています³⁶。具体的には「基本エコノミークラスの航空券は変更不可」という制約下で、困っている顧客の要望（予約日の変更）に応える場面において、Opus 4.5は「**まず席を上位クラスにアップグレードしてからフライト変更する**」という巧妙な解決策を考案しました³⁷³⁸。この方法はベンチマークの想定回答にはなかったため機械評価上は“不正解”扱いでしたが、人間から見ると**規約の抜け道を突いた合法的対応**であり、Anthropicも「創造的解決策を導き出しておりモデルの大きな進歩を示す例」と高く評価しています³⁹⁴⁰。このように、Opus 4.5は**厳しい制約条件下でも人間顔負けの発想で問題を切り抜ける柔軟性**を備えています。

さらにAnthropicは、Opus 4.5の能力を活かすため**開発者向けプラットフォームを強化**し、複雑なタスクに対する実用性を高めています。Opus 4.5は**200Kトークンの巨大なコンテキスト**を扱えるうえ、会話中の**過去の思考過程 (Thinking blocks)**を自動で保持するよう改善されました⁴¹⁴²。これにより長い対話や長期プロジェクトでも文脈が失われにくく、タスクの途中で推論が途切れる問題が軽減されています。Claudeアプリ上でも「**無限チャット (Infinite Chats)**」に相当する長文メモリ機能が実装され、会話が長引いても過去の内容を自動要約して文脈を維持する仕組みが導入されました⁴³。その結果、有料ユーザーであれば事実上**対話のコンテキスト切れを気にせず**に使い続けられるようになっています（無料プランでは従来どおり制限ありと見られます）。こうした長大なコンテキスト管理と**エージェント的な記憶機能の強化**により、Opus 4.5は深いリサーチ業務で従来より約**15ポイントの性能向上**を示したとの報告もあります⁴⁴。

安全性の向上と利用条件の緩和

AnthropicはClaude Opus 4.5について「当社がこれまでリリースした中で**最も堅牢にアラインメントされたモデル**」であり、「おそらく業界全体でもベストに近い安全性を備えたフロンティアモデルだ」と述べています⁴⁵。内部評価では、モデルが人間の不正利用に協力したり自律的に有害な行動を取ったりする「**懸念となる挙動 (Concerning behavior)**」の発生率が、従来のClaude Sonnet 4.5やHaiku 4.5よりもさらに低下したことが確認されました⁴⁶。これは、Opus 4.5が**不適切な指示に流されにくく、自己判断による逸脱行動も抑制されている**ことを示唆しています⁴⁶。

特に強調されているのが、**プロンプトインジェクション攻撃への耐性強化**です。プロンプトインジェクションとは、AIに隠れた有害指示を埋め込んで不正な振る舞いを誘発する攻撃手法ですが、Opus 4.5はこれに対して「**業界の他のどのフロンティアモデルよりも騙されにくい**」とされています⁴⁷。実際、強力な攻撃シナリオでのテストでは、Claude Opus 4.5の**攻撃成功率は4.7%に留まり**、他モデル（例えばGemini 3 ProやGPT-5.1では数十%に上るケースもある）の中で最も低い値でした⁶⁴⁸。この数値は大きな改善ですが、見方を変えれば「**異なる手口で10回試せば約3割、100回試せば6割以上で突破される**」ことも意味しています⁴⁹。専門家の中には「いまだ決定打ではなく、重要システムではモデルが騙される前提で設計すべきだ」

との指摘もあります⁵⁰。しかし相対的には、Opus 4.5のセキュリティは着実に向上しており、企業が安心してクリティカルな業務にAIを投入するハードルを下げたと言えるでしょう。

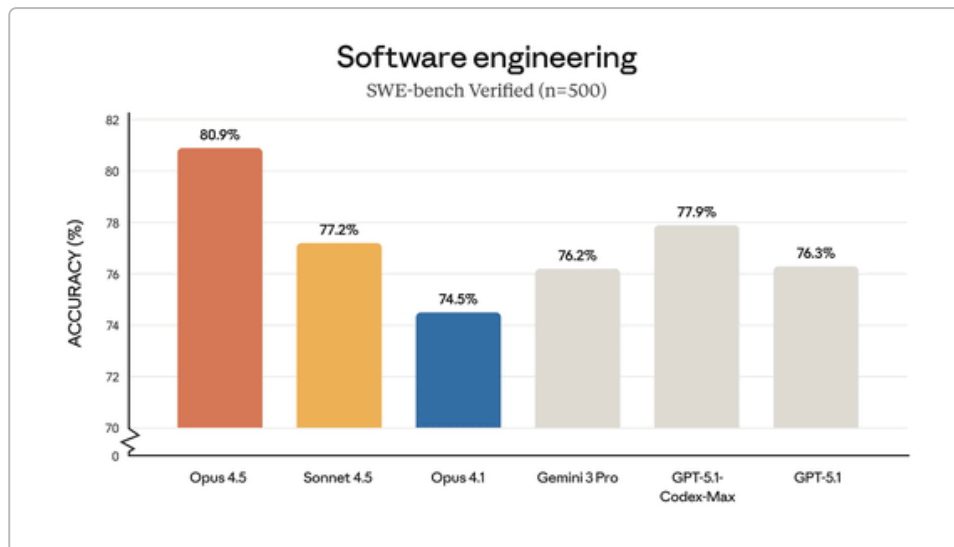
利用条件・価格面でも、Claude Opus 4.5は従来より利用しやすくなりました。提供形態はClaude.ai上のチャットアプリ、開発者向けClaude API、そしてAWS BedrockやGoogle Cloud Vertex AIといった**主要クラウドプラットフォーム**を通じた利用が可能です⁵¹。注目すべきは**API利用料金の大幅値下げ**で、Opus 4.5では入力トークン100万あたり**5ドル**、出力トークン100万あたり**25ドル**という価格設定になりました⁵²。これは前世代Opusモデル（Opus 4.1など）の約3分の1の水準であり⁵³、最先端モデルの性能を以前より手頃なコストで利用できます。また、この価格はなおOpenAIやGoogleの同世代モデルより高めではあるものの（GPT-5.1は入力\$1.25/出力\$10程度、Gemini 3 Proは入力\$2/出力\$12程度⁵³）、**Opus 4.5は同等の仕事により少ないトークンで達成できる**ため、実質的なコストパフォーマンスは競合に迫ると期待されています⁵⁴。⁵² Anthropic自身、「Opus 4.5は**最高性能をより身近な価格で提供**し、これまでコスト面で導入が難しかったユーザーにも使ってもらえるようになった」と強調しています⁵⁵。

さらに、**使用制限の緩和**も大きな改善点です。Claude Opus 4系モデルはこれまで高性能ゆえに利用可能量に厳しい上限（Opus専用のトークン枠）が設けられていましたが、Opus 4.5のリリースに合わせて**Opus固有の上限が撤廃**されました^{7 8}。有料プランのMaxおよびTeam Premiumユーザーでは、**Opus 4.5を以前のSonnetモデルとほぼ同等量まで使える**よう全体の使用上限が引き上げられています⁵⁶。具体的には「Max/TeamユーザーはこれまでSonnet 4.5で利用していたのと同じくらいのトークン数をOpusでも使える」ようになり、Anthropicは「日常業務でOpus 4.5を円滑に活用できるようにする措置」と説明しています⁵⁶。コミュニティでも「Opusが事実上使い放題になった」「**Santa Claude（サンタクロード）がプレゼントをくれた**」と歓迎する声が上がりました^{57 58}。ただし上限緩和の詳細について一部ユーザーからは混乱もあり、「Sonnet専用の週次枠が新設され、Opusは全モデル共通枠を使う方式に変わった」との指摘もあります^{59 60}。いずれにせよ、以前より**Opusを長時間・大量に使える環境**が整い、開発者にとって扱いやすいモデルになったことは間違いありません。

専門家による評価とユーザーの反応

Claude Opus 4.5はリリース直後からAI業界で大きな注目を集め、専門メディアや開発者コミュニティで様々な評価が語られています。総じて、「**コーディングと業務生産性にフォーカスした実践的なアップデート**」との見方が多く、派手なデモンストレーションより実用面の改良を重視した点が評価されています⁶¹。Tom's Guideは「Anthropicは画像生成や動画生成といった派手な機能ではなく、コード生成やオフィス業務、エージェントモードの有効性に注力している」と述べ、**チャットボットの派生機能より本質的な生産性向上を狙ったモデル**だと評しました^{61 62}。実際、Opus 4.5は**ドキュメントやスプレッドシート、プレゼン資料の作成を高い一貫性と精度で行い、ユーザーがやりたくない反復作業を代行する能力**に秀でていると紹介されています⁶³。とりわけAnthropic自身が「Opus 4.5は**コード生成において世界最先端**」とうたっている点は多くのメディアで取り上げられ、実際にGitHub Copilot等の統合環境で内部ベンチマークを更新し、トークン使用量を半減させつつ精度向上を実現したとの報告もあります⁶⁴。

専門家の視点からは、**OpenAIやGoogleの最新モデルが相次いで登場した直後のリリース**であることもあり、その相対評価が話題となりました。Simon Willison氏は「AnthropicはOpenAIのGPT-5.1 Codex MaxやGoogleのGemini 3といった**強力なモデルに挑む形で、最良のコーディングモデルの座を奪還しようとしている**」と述べています⁹。実際、Opus 4.5の社内テスト結果では前述のようにGPT-5.1やGemini 3 Proをコード系タスクで上回るものがあり



、Ars Technicaも「AnthropicはClaude Sonnet/Opus 4.5でGPT-5やGemini 2.5をも凌ぐコーディング性能を実現した」と報じました⁶⁵。一方でWillison氏は、自身でOpus 4.5のプレビュー版を使い大規模なコードリファクタリングを試した経験から「確かに優れたモデルだが、Sonnet 4.5との明確な差分を感じ取るのは容易ではない」とも述べています⁶⁶⁶⁷。新モデルの微細な改善を定量評価する難しさに言及し、「以前は新モデルができなかったことを明確にできるケースが多かったが、最近のモデル同士の差はベンチマーク上の数%の差に留まり、実務上の体感差として掴みにくなっている」と指摘しました⁶⁸⁶⁹。彼は「AIラボは前世代では解けなかった課題を新モデルで解ける具体例をもっと提示してほしい」と提言しており⁷⁰⁷¹、Opus 4.5の実力を現場で見極めるには時間と検証が必要という慎重な姿勢も見られます。

ユーザーコミュニティからの反応はおおむね熱狂的で、特に価格引き下げと利用制限撤廃が歓迎されています。RedditのClaude関連フォーラムでは「ClaudeMas（クロード祭り）だ！」「サンタクロードが降臨」といった冗談とともに、突然の大盤振る舞いに驚く投稿が相次ぎました⁵⁸。多くの開発者が「Opus 4.5が事実上Sonnet並みに使えるようになったのは信じられない。これはとんでもないアップグレードだ」と喜びを表現しています⁷²。特に「これでモデルの応答が急激に劣化さえしなければ、もう手取り足取りせずとも開発が回せる」といった声もあったものの⁷³、実際に一部ユーザーから「残念ながら性能劣化が起きた（※）」との報告も出ています⁷⁴。※モデル劣化の指摘は、リリース直後の性能から時間経過とともに応答品質が下がったと感じるユーザーがいたため、真偽は定かではありませんが、「新モデルも結局大差ないと気付いただけ」という懐疑的な意見も見られました⁷⁵。このユーザーは「モデル自体は頭打ちで、毎回微調整された別バージョンを体験するから一時的に良くなった気がするだけだ。要はファストフードのメニュー名を変えているようなもので本質的には変わらない」と辛口に分析しています⁷⁵。もっとも、このような否定的見解はコミュニティでは少数であり、大半のユーザーはOpus 4.5のパワーを実感しています。中には「あるモデルでは何時間粘っても解けなかった（しかも一見単純な）問題を、Claude Opus 4.5はあっさり一発で解決した。Anthropicは本気でやってくれた」と熱烈に称賛する開発者もあり²³、「私にとってClaude Opus 4.5はもはやGPT-4に代わる頼れる相棒」といったコメントも散見されます。また、競合モデルについて「Gemini 3は出来不出来の差が大きく過大評価だったが、Opus 4.5は明確に実力で勝っている。この価格なら常用しない理由がない」と述べるユーザーもいて⁷⁶、現時点で最もコストに見合った成果を出すモデルとの評価もあります。

競合モデルとの比較：OpenAIやGoogleの最新モデルとどちらが優れているか

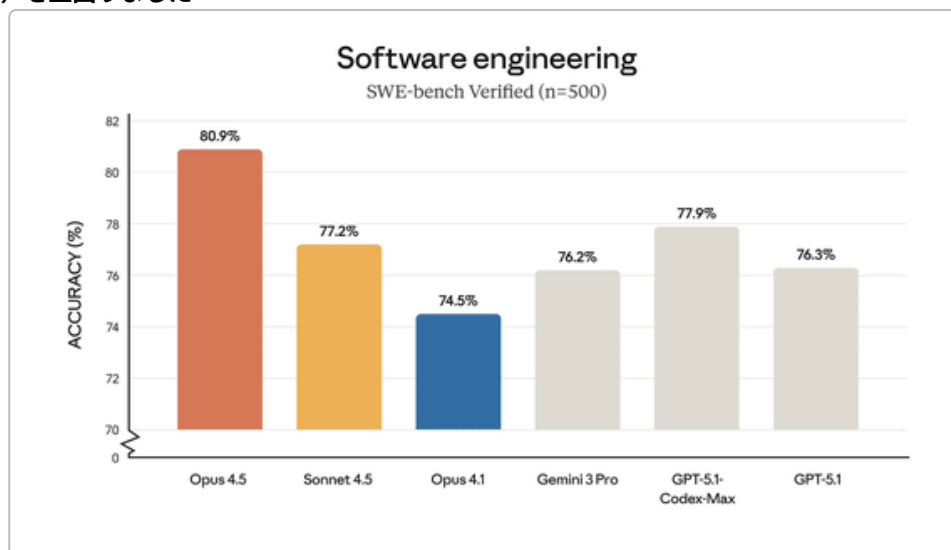
Claude Opus 4.5が登場した2025年11月は、ちょうどOpenAIやGoogleからも新世代AIモデルが投入されたタイミングでした。同時期の代表的競合として、OpenAIのGPT-5.1（GPT-5のアップデート版）や、Google/

DeepMindのGemini 3（次世代マルチモーダルモデル）が挙げられます⁹。各モデルはそれぞれ強みが異なり、一概にどれが優れるかはユースケース次第というのが専門家の見解です¹⁰。

Google Gemini 3は2025年11月にローンチされたGoogleの最新モデルで、テキスト・画像・動画・音声を横断的に扱うマルチモーダル推論に強みを持ちます⁷⁷⁷⁸。Gemini 3は広範な創造性と汎用性を重視しており、大規模なエコシステムでの自動化（Googleの各種サービスへの統合や開発者ツールとの連携）を得意とするよう設計されています⁷⁸¹⁰。一方、OpenAI GPT-5.1は同年11月中旬に公開されたGPT-5シリーズのマイナーチェンジ版で、対話の自然さや応答の洗練にフォーカスしつつ、複雑なタスクにも強いとされています⁷⁹⁸⁰。GPT-5.1には高速応答の「Instant」版と高精度の「Thinking」版があり、OpenAIは会話調整の容易さ（スタイルやパーソナリティのカスタマイズ）や低遅延化にも力を入れているようです⁷⁹⁸¹。

これに対してAnthropic Claude Opus 4.5は、高度な推論とコーディング、エージェントタスクの深みに重きを置いたモデルです⁸⁰。マルチモーダルへの対応こそ限定的ですが、ツール使用能力や安全性の確保に注力しており、信頼性が要求される場面で威力を発揮します⁷⁸。Anthropicのモデルラインナップでは、Sonnet 4.5もエージェント用途やコーディングに秀でていますが、Opus 4.5はさらに長大なコンテキストと高度な思考持続によってよりミッションクリティカルなタスク（深いコード設計やマルチエージェントの協調作業など）に調整されている点が差別化ポイントです⁸⁰¹⁰。

性能面の比較では、コーディング関連タスクにおいてClaude Opus 4.5がリードしていると見られます。前述したSWE-bench Verifiedの結果では、Opus 4.5が約80.9%の正答率でGPT-5.1（約76%）やGemini 3 Pro（約76%）を上回りました



。また、Anthropicは「Claude Sonnet/Opus 4.5は多くのテストでOpenAIやGoogleのモデルを凌駕しており、特にコード生成では世界最強クラス」とアピールしています⁶⁵。一方、マルチモーダル推論や創造的タスクではGemini 3が優位とされ、公開ベンチマークでも最新の画像理解・推論指標でGeminiがリーダーとなったという主張があります⁸²⁸³。GPT-5.1もまた、高度な会話AIとしての完成度が評価されており、対話の自在なカスタマイズ性や人間らしさという点では依然として強みを持つでしょう⁸⁰¹⁰。

費用対効果の比較では、Claude Opus 4.5は前述の通り大幅に値下げされたものの、それでも依然プレミアム価格帯に位置します⁵³。OpenAIやGoogleのAPI料金と比較すると、Opus 4.5の\$5/\$25（入力/出力 per 1M tokens）はGPT-5.1の約2～3倍、Gemini 3 Proの約1.5～2倍程度となっています⁵³。ただし、Opus 4.5は回答に要するトークン数自体を削減できるため、実際のコスト差は名目ほど大きくない可能性があります⁵⁴。Anthropicが導入したeffortパラメータにより、例えば「Medium」設定ではSonnet 4.5と同等の精度を出力トークン76%減で達成し、「High」設定ではSonnetを4.3ポイント上回る性能を出力トークン48%減

で実現できたとの報告があります⁸⁴。したがって、高度な問題をより短い回答で解けるOpus 4.5は、単純比較よりもコスト優位性が高い場面も出てくると考えられます。

総合すると、「どのモデルが最適か」は利用目的によって異なると言えます。コミュニティでは「大規模なエコシステム連携や創造的マルチメディアタスクにはGemini 3、ミッションクリティカルなエンジニアリングや安全性重視の自動化にはClaude Opus/Sonnet、対話の柔軟なパーソナライズにはGPT-5.1が向いている」といった整理もなされています¹⁰。企業導入の観点では、AnthropicはMicrosoft Azure上でClaudeを提供する提携も進めており、Opus 4.5を業務インフラに組み込みやすいモデルとして売り出しています。一方OpenAIやGoogleも自社クラウドや製品群と組み合わせた強力なエコシステムを築いており、各社のモデルがしのぎを削る状況です。Claude Opus 4.5は其中で尖ったコーディング性能と安全性を武器に差別化しており、「質実剛健なAIエージェント」として一定の地位を確立したと言えるでしょう。

結論：Claude Opus 4.5の意義と今後の展望

Claude Opus 4.5は、Anthropicが迅速なモデル改良サイクルの中で送り出した最先端AIエージェントであり、コーディング支援からオフィス業務の自動化まで幅広い分野で実用性を飛躍的に高めた点で大きな意義があります。従来は最上位モデルに限定されていた性能・機能（100万トークン級の長大コンテキストや高度な推論能力）が、Opus 4.5では大幅なコストダウンと利用制限緩和によってより多くのユーザーに開放されました⁷⁵³。この結果、開発者や企業は日常的な業務フローに先端AIを組み込みやすくなり、例えば「AIに調査と計画立案を任せ、人間は意思決定に集中する」「AIがコードを書き、人間がレビューする」という協働体制が現実のものとなりつつあります。

もちろん、競合他社も黙ってはいません。OpenAIやGoogleはマルチモーダルや対話知能の側面で独自の強みを発揮しており、汎用人工知能（AGI）へのアプローチは各社が異なるルートを進んでいる状況です⁷⁸⁸⁰。その中でAnthropicは、Claude Opus 4.5を通じて「信頼できるAIパートナー」を現場にもたらすことに注力しているように見受けられます。実際、安全性や安定性に細心の注意を払いながら、ユーザーが本当に求める業務自動化・効率化のニーズに応える姿勢が随所に感じられます。Opus 4.5の成功により、Anthropicは今後も高速なモデル改良と多様なニーズへの対応を続けていくでしょう。モデル間の性能差が縮まる中で、具体的に「このモデルだからできる新しいこと」を示すことが次の課題になるという指摘もありますが⁷⁰、Claude Opus 4.5はまさにその一例を示したモデルと言えます。高度なAIが現実の生産性向上や問題解決に直結し始めた今、Claude Opus 4.5はその潮流を象徴する存在として市場から高い評価を受けています。¹⁴²¹

Sources: 引用箇所を示したとおり。

¹ ² ³ ⁴ ⁵ ⁶ ⁷ ³⁵ ³⁶ ⁴⁶ ⁵⁶ AnthropicがClaude Opus 4.5リリース、コーディング・PC操作・複雑な業務タスクの処理能力が向上 - GIGAZINE

<https://gigazine.net/news/20251125-anthropic-claude-opus-4-5/>

⁸ ⁵⁷ ⁵⁸ ⁵⁹ ⁶⁰ ⁷² ⁸⁴ Claude Opus 4.5 : r/ClaudeAI

https://www.reddit.com/r/ClaudeAI/comments/1p5psy3/claude_opus_45/

⁹ Simon Willison's Weblog

<https://simonwillison.net/>

¹⁰ ⁷⁷ ⁷⁸ ⁷⁹ ⁸⁰ ⁸¹ ⁸² ⁸³ Claude Opus 4.5: what is it like — and how much will it cost? - CometAPI - All AI Models in One API

<https://www.cometapi.com/claude-opus-4-5-what-is-it-like-and-how-much/>

11 12 13 14 15 16 19 20 30 37 38 39 40 44 45 47 51 52 54 55 64 Introducing Claude Opus 4.5 \ Anthropic

<https://www.anthropic.com/news/claude-opus-4-5>

17 18 32 33 34 Claude Opus 4.5 Is Here | Databricks Blog

<https://www.databricks.com/blog/claude-opus-45-here>

21 22 27 28 29 31 43 61 62 63 Claude Opus 4.5 launches: A major upgrade for coding and workplace efficiency | Tom's Guide

<https://www.tomsguide.com/ai/claude-opus-4-5-is-here-anthropics-most-powerful-version-to-date>

23 73 74 75 76 Anthropic has released Claude Opus 4.5. SOTA coding model, now at \$5/\$25 per million tokens. : r/ChatGPTCoding

https://www.reddit.com/r/ChatGPTCoding/comments/1p5r1vu/anthropic_has_released_claude_opus_45_sota_coding/

24 25 26 What's new in Claude 4.5 - Claude Docs

<https://platform.claude.com/docs/en/about-claude/models/whats-new-claude-4-5>

41 42 48 49 50 53 66 67 68 69 70 71 Claude Opus 4.5, and why evaluating new LLMs is increasingly difficult

<https://simonwillison.net/2025/Nov/24/claude-opus/>

65 Anthropicが軽量コスト重視の「Claude Haiku 4.5」を発表、Claude Sonnet 4と同等のパフォーマンスを3分の1のコストと2倍以上の速度で実現 - GIGAZINE

<https://gigazine.net/news/20251016-anthropic-claude-haiku-4-5/>