

Perplexityの新機能「Model Council」の詳細レポート

はじめに

米Perplexity社は2026年2月、新機能「**Model Council**」（モデル評議会）を発表しました^①。Model Councilは**複数の大規模言語モデル（LLM）の知見を合議**することで、より正確で信頼性の高い回答を生成する試みです^①。従来、ユーザーは質問ごとに「どのAIモデルが最適か」を選ぶ必要がありましたが、この機能により「**最強AIたちが裏で会議をしてから回答してくれる**」ような新しい検索体験が実現すると期待されています^②。以下では、Model Councilの技術的仕組みや対応モデル、その効果と活用例、公式発表内容、外部からの評価、将来の展望について詳しく解説します。

1. 技術的な仕組み – 複数モデルの合議による回答生成

Model Councilの根幹は、**複数のLLMが一つの質問に対して同時に回答し、最終的にそれらを統合するプロセス**にあります。ユーザーがModel Councilモードを選択して質問を投げかけると、**3つの異なるAIモデル**（例えば Claude Opus 4.6、GPT-5.2、Gemini 3.0 など）に**並行して同じ質問が送信**されます^③。各モデルは**独立して回答を生成**し、その結果がPerplexity独自の「**シンセサイザー（統合）モデル**」に集められます^④。

Where Models Agree				
Finding	🌀	A	G	Evidence
Bench depth/scoring is a critical weakness	✓	✓	✓	Knicks bench ranked 27th-28th in PPG (~30.7), limiting rotation options knicksfactor +2
Defense has been problematic, especially post-NBA Cup	✓	✓	✓	Ranked 29th in defensive rating (119.6) since December 16 sl +1
Jose Alvarado trade was the right move	✓	✓	✓	Addresses backup PG need, saves money, fills McBride injury void postingandtoasting +2
Second apron cap constraints severely limit trade flexibility	✓	✓	✓	Only \$148K under hard cap, no trade exceptions, must salary-match tightly dailyknicks +1 YouTube
Mitchell Robinson health/availability remains a concern	✓	✓	✓	On load management plan, can't play back-to-backs, expiring contract sny +1
Washington's protected pick will likely NOT convey as a 1st rounder	✓	✓	✓	Top-8 protected; Wizards tanking and resting Trae Young to ensure retention fanspo +1
Draft should prioritize frontcourt depth/rim protection	✓	✓	✓	Robinson's expiring deal and injury history necessitate long-term center solution fanspo +1

合議の結果は上図のように**表形式で提示**されます^④。シンセサイザーモデルが**各モデルの回答内容を精査・比較し、共通点**（全モデルが一致した点）や**相違点**（モデルごとに意見が分かれた点）、**各モデル特有の洞察**を整理して統合します^④。表には質問に対する主要なポイントが箇条書きされ、それぞれのポイントについて**どのモデルが言及・賛同したかがチェックマークで示されています**。また各ポイントには**裏付けとなる情報源（出典）**も併記され、ユーザーは必要に応じてエビデンスを確認できます^⑤。さらに、ユーザーは**各モデルの元の回答全文**も個別に参照可能で、統合結果と併せて詳細を検証することができます^⑤。

重要なのは、単純に多数決で回答を決めているわけではない点です。シンセサイザーモデルは**3つの回答間の矛盾を検出・解決**しながら統合しています⁶。すなわち「2対1で多数決だからこの答え」と機械的に決めるのではなく、**各モデルの主張のどこが一致していて、どこが食い違うのかを分析し、矛盾を調停した上で最適な結論をまとめる仕組み**です⁶。Perplexity社はこのアプローチを「**協調的アンサンブル**」と呼んでおり、いわばAI同士が相互に批評し合うことでより良い解を見出す試みと言えます⁶。実際、MITとユニバーシティ・カレッジ・ロンドンの研究では、**単一のAIエージェントで約70%に留まっていた算術タスクの精度が、3つのエージェントに2ラウンド協調させると約95%に向上したことが報告されています**⁷。また別の研究では、**マルチモデルアプローチによってAIの誤情報（ハルシネーション）を用途に応じ18%から最大90%削減できた**とのデータもあり⁸、独立したモデル同士が同時に全く同じ誤答をする確率は極めて低いため、**合議制は統計的に理にかなったエラー低減策**と位置づけられます⁸。

シンセサイザーモデルによる最終回答は、**モデル間で意見が一致した部分は安心材料として示され、意見が割れた部分は「深掘りすべきポイント」としてハイライト**されます⁹。ユーザーは統合結果から「どの情報は複数モデルがお墨付きを与えたか」「どの点についてモデル間で見解が食い違っているか」を一目で把握でき、信頼性と多角度性を両立した回答を得ることができます¹⁰。

2. 対応しているモデルの種類と役割

Model Councilで利用できるモデルは、現時点で業界をリードする**最先端の大規模言語モデル**が中心です³。公式に明示された例としては、Anthropic社の「**Claude（クロード）**」、OpenAI社の「**GPTシリーズ**」（GPT-4や社内テスト中のGPT-5.2など）、Google DeepMindの新モデル「**Gemini（ジェミニ）**」、そしてPerplexity社内の高機能モデル（コードネーム: Opus 4.6 など）などが挙げられています³¹¹。さらに報道や有識者の言及からは、オープンソースの有力モデルである「**Mistral（ミストラル）**」や、他社の最新モデル（例：Elon Musk氏のXAIが開発した「Grok」やMeta社の「Llama系モデル」）も念頭に置かれていることが伺えます¹²。ユーザーはModel Council実行時に**使用する3モデルの組み合わせを自分で選択**できるようになっており、UI上で候補モデルのオン/オフ切替が可能です¹³（Maxプランでは標準で先述のトップモデル群が利用可能とみられます）。

では具体的に、これら各モデルにはどのような特徴・強みがあり、合議においてどのような**役割**を果たすと考えられているのでしょうか。以下に主要なモデルを挙げ、その特色をまとめます。

- **GPT-4（OpenAI）** – OpenAIが提供するGPTシリーズのモデルで、極めて**汎用性が高い高性能モデル**です。幅広い知識と高度な推論能力を持ち、プログラミング支援や文章生成、要約など多用途に活用されています¹⁴。開発者向けの拡張性（プラグインやAPI、コード生成能力）にも優れており、「AIの何でも屋」的存在と言えます¹⁴。合議体においては、その**総合的な知識と安定した回答品質**で他モデルの基盤を支え、**複雑な内容を平易に説明**する役割で貢献するとされています¹²。実際、ある分析ではGPT-4（ChatGPT）が他モデルの出力を**分かりやすく噛み砕いて説明**する「解説役」を担う傾向があると指摘されています¹²。
- **Claude 2（Anthropic）** – Anthropic社が開発するモデルで、**安全性と長文処理能力**に定評があります。人権や倫理に配慮したチューニングが施されており、**有害な出力を抑制**するガードレールの堅牢さが特徴です¹⁵。また大容量のコンテキスト（長い会話履歴や文書）を扱う能力にも優れ、複雑な推論タスクにも強みを発揮します¹⁵。Model CouncilではClaudeは**慎重で信頼性の高い視点**を提供し、議論をまとめる「**結論づけ役**」になることが多いとも言われます¹²。実際ある実験では、Claudeは提示された情報から**最終判断を下す役**を自然と引き受ける傾向が見られました¹²。例えばアーキテクチャ設計の議論では、ClaudeがROI（費用対効果）などを踏まえて「～すべきである」と**意思決定を主導**する回答を示したケースが報告されています¹⁶¹²。
- **Mistral（ミストラル）** – 近年登場したオープンソースの大規模言語モデルです。比較的**軽量で動作が高速**なモデルでありながら、高い性能を発揮するよう最適化されています¹⁷。開発者が自前でカス

タマイズしたり追加訓練したりできる**柔軟性**も特徴です¹⁷。Mistralは公開モデルゆえ内蔵の制約が少なく、技術的・専門的な詳細に踏み込んだ回答や、データのありのままを出力する回答を得やすい利点があります。合議に加えるモデルとしては、**軽量ゆえの高速応答で下調べを素早く提示**したり、**他の大規模モデルが抽象化・要約してしまった技術的ディテールを補完**する役割が期待できます

¹²。実際に複数モデルでの実験では、Mistralは**技術的な精密さを提供するポジション**を担い、ChatGPTやClaudeが見落としした専門的視点を示す場面があったと報告されています¹²。

- ・**Gemini（ジェミニ、Google DeepMind）** – Googleが開発した次世代の大規模モデルで、特に**マルチモーダル（画像・音声なども扱う）能力**とGoogleの膨大な知識基盤への統合で注目されています

¹⁸。Google Workspaceなど業務ツールとのシームレスな連携も視野に設計されており、**ビジネス分野への適用**が期待されます¹⁹。Geminiは高度なパターン認識能力を活かし、合議体では**全体像の俯瞰とパターンの抽出**を得意とするでしょう²⁰。例えば長期的な戦略立案に関する問いでは、Geminiが膨大な情報から**傾向を総合してパターンを見出し**、他モデルが見落としがちな大局的視点を提供する役を果たしたとされています²⁰。要するに、Geminiは「**パターンの統合役**」として、多角的な情報をまとめ上げる力を持っています²⁰。

- ・**Perplexity独自モデル（Opusなど）** – Perplexity社内で運用するモデルも存在します。社内名称で**Opus 4.5/4.6**や**Sonnet**といったモデルが言及されており²¹²²、おそらく他社モデル（GPT-4等）をベースに追加チューニングや独自工夫を凝らした**検索特化・推論強化型のモデル**と推測されます。例えばPerplexityには**検索結果から証拠を提示する機能**がありますが、これはOpusやPerplexityモデルが**外部知識ソースを引用付きで統合**する役割を果たしているためです²³。合議においては、自社モデルが**Web検索結果や引用情報を検証・補強する「エビデンス検証役」**として組み込まれており、他の汎用LLMが生成した回答に対して**事実関係の裏付け**を与えるポジションを担っています²⁰。

これらのモデルはそれぞれ**得意分野や応答傾向が異なる**ため、Model Councilではその多様性を活かして相互補完しています。実際の運用例として、ある専門家は「**ChatGPTは説明役、Claudeは決断役、Mistralは技術的な正確さを補い、Geminiはパターンを統合する**」といった形で、モデルごとに人間のチームにおける役割分担のような振る舞いが見られると述べています¹²。例えば「**アーキテクチャ設計の検討**」という問いでは、ChatGPT（GPT-4）が説明を担当し、Mistralが技術面の精密な指摘を行い、Claudeが意思決定（結論提示）を行い、Grok（XAIのモデル）がストレステスト（反論提示）をする、といった具合です¹²。また「**事業戦略の検討**」では、Geminiがパターンを総括し、Copilot（コード生成特化モデル）が実装上の現実性を指摘し、Perplexityモデルが証拠を提示して検証し、最後にClaudeが全体を要約して決定案を示すという流れが観察されたといいます²⁰。このように、Model Councilでは各モデルの強みを役割として引き出すことで、**単一モデルでは得られない深みとバランスを持った回答**を生成しているのです。

3. 代表的な活用例と効果 – 検索精度の向上・利便性・回答の一貫性と多様性

Model Councilの導入によって期待される利点は多岐にわたります。ここでは公式発表で示された代表的なユースケースとともに、その効果（検索精度や利便性の向上、回答の一貫性・多様性への影響）を解説します。

- ・**検索精度の向上（正確性・信頼性）**：複数モデルの相互検証により、**事実誤認や見落としのリスクを低減**できます。単一モデルではどうしても訓練データ由来の偏りや知識の盲点があり、「自信満々に誤った情報（ハルシネーション）を出力する」問題が付きまといま。しかしModel Councilなら、他のモデルが誤りに気づいて修正したり、異論を提示することで**誤答を潰していく効果**が期待できます⁸。実例として、先述のように複数モデルによる合議は算術問題の正答率を70%台から95%近くまで引き上げ⁷、ハルシネーション由来のエラーを大幅に削減したという報告もあります⁸。Perplexity社自身、「複数の視点からクロスバリデーションすることで、**研究や意思決定で重要な質**

問に対してより信頼性の高い回答が得られる」と述べています²⁴。ユーザーから見れば、**3つのモデル全員が一致している回答**であればかなり安心感が増し、逆に食い違っている場合は「鵜呑みにせず検証が必要」という注意喚起が得られるため、**結果的に検索・分析の精度と信頼性が向上**します。

- ・**ユーザー利便性の向上**: Model Councilはユーザーの手間を大幅に軽減します。従来、ある質問に対してChatGPTやClaudeなど複数のAIチャットボットを行き来し回答を見比べるのはユーザーの負担でした⁵。特に答えがモデルによって異なる場合、その違いを自分で分析して結論を出すのは時間がかかります。Model Councilではこうした作業を裏側で自動化し、数秒で要点を比較整理した結果を提示します⁵。ユーザーは短時間で全体像を把握し、回答の妥当性を検証しやすくなります⁴。例えば、「この法律問題についてGPT-4とClaudeで答えが異なるが、どちらを信用すべきか？」と悩む代わりに、Model Councilなら一つの統合回答の中でその差異が明示されるため、追加の判断材料が得られます。加えて各ポイントには引用元が付くため、ユーザー自身が深掘りして裏付けを確認することも容易です⁵。このように**答え合わせと比較検討がシームレス**になることで、ユーザーはAI回答を以前より**安心して参考**にできるようになります。

- ・**回答の一貫性・多様性への影響**: Model Councilは回答内容の一貫性（Consistency）と多様性（Diversity）のバランスにもユニークな効果をもたらします。一見すると「複数モデルの合意に基づく回答」は妥協点のようで、多様な意見が均されてしまう懸念もあります。しかし実際には、Model Councilの回答には**一致点と相違点の両方が明示されるため、共通する事実はより確からしいものとして強調**されつつ、**意見の分かれた部分は多面的な視座として残**されます⁹。すなわち、一貫して**正確な部分とモデル間で異なる見解の両方**をユーザーが識別できるのです。これは回答の質に二重のメリットを与えます。一つは、各モデルが個別に答えたとき生じがちな些末な不一致（言い回しの違いや細部のズレ）が統合によって解消され、**コアとなる結論部分の一貫性が高まる**ことです。もう一つは、単一モデルでは出てこなかった**創造的なアイデアや異論**が統合回答内で紹介され、**回答に深みと多様性**が加わることです。Perplexity社も「複数のAIの観点を同時に反映することで、広い視野に基づいた回答が得られる」と強調しています²⁵。例えば投資助言の質問で、あるモデルは強気シナリオを、別のモデルは弱気シナリオを提示した場合、統合回答では両シナリオの根拠を並記しつつ「ここが見解の分かれ目です」と指摘する形になるでしょう。これはユーザーにとって**一方的な答えより有益な判断材料**となります。総じて、Model Councilは**正確性（一貫性）の向上と視点の多様化**を両立し、回答の質的向上に寄与しています。

- ・**典型的なユースケース**: Perplexity公式ブログや発表では、Model Councilの効果が特に発揮されるユースケースとして以下が挙げられています⁴²⁶：

- ・**投資リサーチ** – 株式や市場に関する分析で、複数モデルの**バイアスを相殺し合ったバランスの良い見解**を得る²⁷。金融分野では一つの判断ミスが大きな損失につながりかねないため、**複数AIのセンサスで高い確信度**をもって判断材料を提供できる点が評価されています²⁷。
- ・**複雑な意思決定** – 高額な買い物やキャリアの選択、ビジネス戦略など**複合的要因を考慮する意思決定**で、異なるモデルの推論を付き合わせて検討できる²⁸。様々な角度から利点・欠点が提示されるため、見落としを減らし「後悔のない選択」に資するとされています²⁸。
- ・**創造的ブレインストーミング** – プレゼントや旅行プラン、コンテンツアイデアの発想など**クリエイティブな問い**で、各モデルの得意分野から多彩なアイデアを引き出す²⁹。例えばあるモデルはユニークな発想を、別のモデルは定番だが堅実な案を提示するかもしれず、**多様なアイデアの中からベストを選ぶ**のに役立ちます²⁹。
- ・**情報の検証・クロスチェック** – 医療・法律・学術など**正確さが極めて重要な質問**で、**複数モデルに事実確認させ相互に裏付けを取る**³⁰。一つのモデルでは見逃したエビデンスも、別のモデルが拾ってきて補完することで、回答の信頼性を高めます³⁰。間違いやデマ情報のリスクを抑え、より安全なアウトプットを得られるメリットがあります。

以上のように、Model Councilは高精度な検索・Q&Aが求められる場面で威力を発揮します。実際、あるメディアはこの機能について「ユーザーが複数モデルの回答を**手作業で読み比べなくても**、AI側で**自動的に食い**

違いを洗い出して集約してくれるので、重要な情報の見落としが減り、より良質な意思決定がスピーディーに可能になる」と評価しています³¹³²。特に調査や分析業務では、AIの一貫性のなさに悩まされるケースが多かったため、この問題に正面から取り組むModel Councilのアプローチは新たな標準となりうる可能性を秘めています。

4. 公式発表・開発者ブログでの言及

Perplexity社はModel Councilの公開にあたり、公式ブログ記事（Introducing Model Council）やヘルプセンターの記事でその概要を説明しています³³³⁴。発表日は2026年2月5日で、同社の最高経営責任者（CEO）アラヴィンド・スリニバス氏のSNS投稿などでも告知が行われました³³³⁵。以下に公式発表で言及されたポイントをまとめます。

- ・**信頼性向上へのソリューション**: Perplexityは「AIの回答を本当に信じてよいのか」というユーザーの不安に応える形でModel Councilを位置付けています¹。公式ブログでは、複数の最先端モデルから**“AIコンセンサス”**を得ることで新しい信頼性の基準を提示すると述べられました¹。「3つのモデルを同時実行し、回答の一致点と相違点を明示することで、高い確信度を持てる回答を提供する」というのが公式の謳い文句です³⁶。特に、重要な質問では単一モデルの即答よりも、合議による**慎重で検証済みの回答**が望ましいことを強調しています。
- ・**技術的仕組みの説明**: 公式発表ではModel Councilの動作を次のように説明しています。「ユーザーの質問を3つの選択されたAIモデルに**並列実行**し、それらの回答を**比較**します。次に別のモデル（**シンセサイザー**）が全回答を**精査・統合**し、モデル間で**どこが一致しどこが異なるか**をハイライトした一つの回答を生成します」³⁶³⁷。またUI上で回答を並べて比較閲覧できること、**各モデルが最終回答にどう寄与したかが分かる**ことにも触れられています³⁸。これらは前述した仕組みそのものですが、公式として「単なるマルチモデルのルーティングではなく、まるで**少人数の専門家パネルが議論して答えをまとめるようなアプローチ**」であると説明されています³⁹⁴⁰。Perplexityはこの機能を**Structured deliberation（構造化された討議）**と表現し、単一モデルでは得られない洞察を引き出せる点を強調しています⁴¹。
- ・**想定ユースケースの提示**: 公式ブログやヘルプでは、Model Councilが有用な具体例として先述の**投資調査、複雑な意思決定、創造的発想、情報検証**の4つが列挙されました⁴²⁶。例えば「株式の分析でバイアスの異なる複数モデルからバランスの取れた意見を得る」「家や車の購入判断で異なる推論アプローチを照らし合わせる」「新製品のアイデア出して複数のモデルから提案を募る」「医療情報をクロスチェックする」といった具体例が示されています²⁷⁴²。これらは前述した内容と重なりますが、公式に「Model Councilは**正確性や多角的視点が求められる場面に適しています**」と説明されていることがわかります⁴³。
- ・**提供プランと対象ユーザー**: Model Councilはまず**Perplexity Max**（同社の最上位有料プラン、月額約200ドル）加入者向けにウェブ版で提供開始されました⁴⁴⁴⁵。また企業向けの**Enterprise Max**にも含まれています³⁴⁴⁶。**無料ユーザーやProプランには未開放**であり、リリース当初は限定的な提供でした⁴⁶。CEOのスリニバス氏はX（旧Twitter）への投稿で「将来的にProユーザーにも提供する方法を模索している」と述べつつ、**計算コストの高さゆえに利用制限（回数制限など）が必要になる可能性**にも言及しています³⁵。その場合は「頻繁に使うユーザーはMaxプランへのアップグレードが合理的だろう」と説明しており、現状では事実上**トッププラン向けのプレミアム機能**として位置づけられているようです³⁵。なお、モバイルアプリでのサポートも近日中に開始予定とされています⁴⁴⁴⁷。

以上が公式発表や開発者ブログで明らかにされたポイントです。総じてPerplexity社は、Model Councilを**AI活用の新たなパラダイム**と位置づけ、「これまで『どのモデルが一番優れているか』を問うていたが、今後は『**複数モデルの合意をどう活用するか**』が重要になる」と示唆しています⁴⁸⁴⁹。

5. 外部メディアの評価・専門家の意見・ユーザーの反応

Model Councilに対しては、AI業界の専門家や一般ユーザーから様々な声が寄せられています。総合すると、多くはこの試みを高く評価しつつ、一部に課題や要望も見られます。

肯定的な評価: - AI専門メディアや有識者の多くは、「Model CouncilはジェネレーティブAIの信頼性問題に対する有力な解決策だ」と歓迎しています。例えば、とあるデータサイエンティストは「開発者たちは以前から複数のLLMに同じ質問を投げて合意を見ることで信頼性を統計的に高めてきた。だがそれは裏方の技術であって消費者が直接目にはなかつた。Perplexity Model Councilはそれを初めて消費者向けツールに実装したものだ」と指摘しています⁵⁰。彼女は「複数モデルの出力を比較させるこの手法のおかげで、ユーザーがより批判的に考えることを後押しできるのが素晴らしい」と評価し、モデルの透明性が人間側の能動的な推論を促す点を称賛しています⁵¹。実際、3つの異なる視点が示されればユーザーも自ずと「なぜ違うのか？」と考えるため、AI任せにせず主体的に情報を評価する習慣につながるという指摘です⁵¹。

- 一部のメディアは、Model CouncilがAIの「幻覚問題（ハルシネーション）」への対抗策として有効だと報じています。Livemintのテックニュース記事では「この新機能は回答の正確さと深度を高め、ユーザーが複数モデルの答えを一度にクロスチェックできるようにすることで、矛盾を減らしより良い意思決定を素早く行えるよう支援する」と紹介されています⁵²⁵³。また別の専門家は「複数モデルが同じ誤情報を同時に生成する確率は極めて低いため、合議は理にかなっている。選挙のように虚偽情報対策が重要な局面で、こうした機能がもっと普及してほしい」と期待を述べています⁵⁴。要するに、「AI版ピアレビュー」とも呼べるこのアプローチが、AIの信頼性向上に寄与すると肯定的に受け止められています。
- 初期ユーザーからも「Google BardやChatGPT単独より信頼できる回答が得られた」「自分で複数AI試す手間が省けて便利」といった好意的な反応がSNS上で見られました。特に技術系コミュニティでは「GPT・Claude・Geminiが一堂に会して答えてくれるなんて画期的だ」と注目を集め、YouTube上でも「NEW Perplexity Model Council is INSANE!」といったレビュー動画が登場しています（※Insane=「すごい」の俗語）。こうした反応から、AIに詳しいユーザー層ほど本機能の価値をいち早く見出している様子がうかがえます。

指摘・懸念されている点: - 一方で、Model Councilには課題や限界も指摘されています。まず大きいのはコストと速度の問題です⁵⁵。3つの高性能モデルを同時に動かすため、計算資源の消費は単一モデルの少なくとも3倍になり、応答に時間がかかる場合もあります⁵⁵。事実、一般ユーザー向けにはMaxプラン（高額）のみに提供という形になっており、多くの人にとっては「高嶺の花」状態です⁵⁶。LinkedIn上でも「素晴らしい機能だが最上位ティア限定で残念。多くの人には届かないのは惜しい」という声が上がっています⁵⁶。Perplexity社も企業努力でProプラン以下への展開を検討中ですが、無料で気軽に使える機能とするには技術的コストのハードルが高い現状があります³⁵。

- モデル間のバイアスの問題も懸念されています。複数モデルとはいえ、現在の主流モデルは共通して巨大量学習データ（主にインターネット由来）に依存しているため、例えば3モデルすべてが類似の偏見を内包していた場合、合議によってそのバイアスが却って強化されてしまうリスクがあります⁵⁷。つまり、多数決で同じ誤った先入観が「正解」とされてしまう恐れです。合議制だからといって必ずしも中立・公正になるわけではなく、むしろ評価プロセスの透明性確保がより重要になるだろうと指摘されています⁵⁷。この点、Model Councilはユーザーに食い違いポイントを見せてはくれますが、「なぜその結論になったか」という各モデル間の討論プロセス自体はブラックボックスです¹⁶。ある専門家は「現状のModel Councilは、モデル同士の議論そのものは内部で省略してしまっており、結論の収束過程は人間には見えない。いわば認識論的利便性（epistemic convenience）を提供しているに過ぎず、真の意味での『AIガバナンス（AI同士が議論しベストを選ぶ統治）』には達していない」と手厳しく評価しています¹⁶。彼はむしろ、モデル同士が何度か反論し合い最終回答を競わせる「オープンな討議型」の仕組みこそが望ましいと提案しています¹⁶⁵⁸。

- ・**ユーザーの初期反応**としては、「確かに慎重な答えになるが、その分回答が冗長に感じる」という声も一部にあります。3モデル分のエッセンスが詰まるため、回答が長めになりがちで「簡潔な答えだけ欲しい時には不向き」という意見です。ただしこの点については、シンセサイザーが重複を排し簡潔にまとめる工夫がなされており、時間とともに洗練されていく可能性があります。またRedditのPerplexityコミュニティでは「面白いデモだが、本当に生産性が上がるかは質問内容による」との議論もありました⁵⁹。「普段は単一モデルで十分な質問も多いし、合議が役立つケースは限定的では？」という慎重な意見です⁶⁰。これに対し他のユーザーから「**高リスクな場面ではお金を払ってでも使う価値がある**。単純な質問なら1モデルでいいが、例えば研究論文のレビューや重要な契約書のチェックではCouncilが威力を発揮する」といった反論も出ています。総じて現時点では、「ポテンシャルは非常に高いが、使い所を見極めたい」「まだ高額なので自分には手が出ないが技術的には面白い」という反応が多い印象です。

6. 今後の展望 – 標準化の可能性、他サービスとの比較、エージェント的進化、商用利用

Model Councilの登場は、今後のAI検索・AIアシスタントの潮流にどのような影響を与えるでしょうか。いくつかの観点から展望を述べます。

- ・**AI検索・回答の新標準となる可能性**: 複数AIの合議によるコンセンサス取得は、長期的にはAI業界におけるデファクトスタンダードになる可能性があります⁶¹。2025年時点で企業の78%がAIを業務利用している一方、77%がAIの出力する誤情報に不安を感じていたという調査があります⁶¹。実際、AIプロジェクトの70～85%は期待された成果を上げられていないとも報告されており⁶¹、**信頼性のギャップ**が大きな課題でした。Model Councilのような技術はこのギャップを埋める鍵となるかもしれません⁶¹。今後、他社のAI検索サービス（例えばMicrosoft BingのAIチャットやGoogle Bardなど）も類似のマルチモデル評価機能を導入する可能性があります。現時点では主要サービスはいずれも**単一の大規模モデル**（BingはGPT-4系、BardはPaLM系など）で応答していますが、Perplexityが示した成果次第では「**複数AIで相互検証してから回答する**」という流れが広がるでしょう。あるメディアは「検索や分析の軸足が『どのAIを選ぶか』から『複数AIの合意をどう活用するか』へ移りつつある」と評しており⁴⁹、この統合アプローチが**意思決定支援の新しい選択肢として定着していく可能性を示唆**しています⁴⁹。
- ・**他の検索サービスとの比較優位**: 現状、Perplexity Model Councilは競合サービスにはない独自機能であり、**AI検索プラットフォームにおける差別化要因**となっています。ChatGPTやBing、Bardなど大手は一モデル回答+ウェブ検索補助に留まるため、**複数モデルの同時活用による高精度回答**という点でPerplexityは先行しています。実際、Perplexityは以前から**検索エンジン機能とAI回答の融合**（ウェブ検索結果を引用する回答）で独自路線を歩んできました²³。Model Councilの追加により、より一層「**正確さ重視の調査向けAI**」としての地位を強めた形です。これは専門職ユーザー（リサーチャーやアナリスト等）にアピールする強みであり、短文の雑談や創作より**信頼性が重視されるユースケース**でPerplexityが選ばれる理由になるでしょう。もちろん、他社もいずれ追随する可能性があります。現時点では**Perplexityが合議AIのパイオニア**と言えます。
- ・**エージェント的機能の進化**: Model Councilは将来的に、より高度なAIエージェントの協調へと発展する可能性があります。現在は3モデルが独立に1ステップ回答して統合していますが、これを複数ラウンドに拡張し、モデル同士が**対話形式で議論・反論し合う**ようにすれば、まさに**AI同士の会議を人間がモデレート**するような形になります¹⁶。専門家の中には、複数LLMを組み合わせることで認知的な役割分担を与え、一種の「AIチーム」を構成することを提唱する声もあります⁵⁸⁶²。例えば「**リサーチ担当AI**」「**批評担当AI**」「**決定案をまとめるAI**」のように役目を明確化し、段階的な討議プロセスをデザインするアプローチです⁶²。実際、Petros Bountis氏（エンジニアリング部長）は自身の実験から、8種類のモデルにそれぞれ**リーダー・説明役・チェレンジャー・技術担当・実装担当・検証担当・要約担当**などの役割を与えることで、人間の会議さながらの意思決定プロセスを再現できたと

報告しています⁶³⁶²。彼は「重要なのはモデルの良し悪しではなく、適切な認知的役割を割り当てることだ」と述べており⁶⁴、この方向性はまさにAIエージェントのチーム化です。将来、PerplexityがModel Councilを拡張して複数モデルが対話しながら問題解決に当たるエージェントを実現すれば、単なるQ&Aに留まらない高度な意思決定支援AIが登場するかもしれません。例えば、Perplexityがまず必要情報を検索で集め（調査役）、Councilがその情報について議論し合い（批評役）、最後に別のモデルが結論を文書化する（執筆役）——といったワークフロー全体をAIが担うビジョンも考えられます⁶²。そのようなエージェント的進化は、AIがより包括的に人間の知的作業を支援・拡張する未来につながるでしょう。

・**商用利用の可能性**: Model Councilの商業的インパクトも見逃せません。まず、Perplexity自身はこの機能をエンタープライズ向けにも提供開始しており³⁴、今後企業内ナレッジ検索や専門領域での利用が見込まれます。特に金融・法律・医療のようなミスが許されない分野では、多少コストを払ってでも信頼度を上げたいニーズがあります⁶⁵。実際、Perplexityも投資リサーチを主要ユースケースに挙げているのはこのためです⁶⁶。企業はModel Councilを使うことで、AIから得た情報に対する説明責任（なぜそう答えたか）を取りやすくなります。複数モデルの合意という形で説明されれば、上層部や顧客にも説得力を持って提示できるでしょう。さらに、将来的には他社のソリューションへの組み込みも考えられます。例えば調査会社のリサーチツールや、専門家向けの情報サービスにこの技術をライセンス提供するなど、プラットフォーム化の展望もあります。現状ではPerplexity独自の強みとして囲い込んでいますが、いずれAPI経由で他サービスからModel Council機能を呼び出せるようになる可能性もゼロではありません（※実際、PerplexityはAPI提供も行っています）。とはいえ、当面はPerplexityプラットフォームへの集客戦略としてこの機能を活用するでしょう。競合が追随してくれば、精度や速度、対応モデル数などでさらなる改良競争が起き、市場全体の水準が引き上がると考えられます。

・**課題と展望のまとめ**: Model Councilはまだ始まったばかりの機能であり、計算コストや応答遅延、バイアス補正、透明性など課題も抱えます⁵⁵⁵⁷。しかし、高リスク領域ではそのコストを上回る価値があるとの指摘も多く⁶⁵、今後の技術革新次第で広範な普及が期待できます。ゆくゆくは、「AIに質問するときは複数のAIエージェントが協議して答えを出すのが当たり前」という時代が来るかもしれません⁴⁹。AIがより高度化するほど、その推論過程を相互チェックさせる仕組みの重要性も増すでしょう。Perplexity Model Councilは、その先鞭を付けた革新的な一歩として評価されます。今後の発展に注目が集まるとともに、我々人間がAIの合議をどう活用し意思決定に役立てるかという、新たな課題も浮かび上がっています。技術が進歩しても最後に判断を下すのは人間です⁶⁷。Model Councilが提示する「AIコンセンサス」を賢く使いこなし、真に価値ある知見を得られるかどうかは、ユーザー次第とも言えるでしょう。

参考文献・出典: - Perplexity公式ブログ: What We Shipped – Feb 6, 2026 (Model Councilリリース)⁶⁸
⁶⁹ - Perplexityヘルプセンター: What is Model Council? (機能説明と利用方法)¹¹⁷⁰ - Perplexity CEOの発言 (LinkedIn/X経由)³⁵ - 各種ニュースサイト (Innovatopia, Plus Web3) による発表解説⁷¹⁷²
⁷³ - 専門家の分析 (Medium記事 “The Council of AIs” など)¹²¹⁶ - LinkedIn上の有識者コメント (Dr. Amanda J. Williamson氏の投稿)⁷⁴ - Livemint他テックメディアの記事⁵³³²

¹ ³ ⁶ ⁷ ⁸ ⁴⁴ ⁴⁸ ⁶¹ ⁶⁵ ⁶⁶ ⁷¹ Perplexity、Model Council発表—3つのAIモデルを同時実行し回答精度を向上

<https://innovatopia.jp/ai/ai-news/79678/>

² ²⁵ ³³ ⁴⁷ ⁴⁹ ⁵⁵ ⁵⁷ ⁷² ⁷³ 複数AIの合議で意思決定が変わる Perplexity「Model Council」が回答生成の新基準に | Plus Web3 media

https://plus-web3.com/media/latestnews_1000_7662/

4 5 35 43 複数AIの視点を統合する検索AIのPerplexity、回答を比較検証する新機能「Model Council」を発表

<https://www.atpartners.co.jp/news/2026-02-10-perplexity-ai-a-search-ai-that-integrates-multiple-ai-perspectives-announces-a-new-feature-called-model-council-for-comparing-and-verifying-answers>

9 10 11 13 24 26 34 45 46 70 What is Model Council? | Perplexity Help Center

<https://www.perplexity.ai/help-center/en/articles/13641704-what-is-model-council>

12 16 20 58 62 63 64 The Council of AIs | The Socratic Leadership

<https://medium.com/the-socratic-leadership/the-council-of-ais-0740d56240a4>

14 15 17 18 19 23 Claude vs Gemini vs ChatGPT vs Mistral vs Perplexity vs CoPilot: The AI Showdown - DEV Community

<https://dev.to/nermine-slimane/claude-vs-gemini-vs-chatgpt-vs-mistral-vs-perplexity-vs-copilot-the-ai-showdown-50o0>

21 22 27 28 29 30 42 68 69 Perplexity Changelog

<https://www.perplexity.ai/changelog/what-we-shipped---february-6th-2026>

31 32 52 53 Perplexity unveils Model Council to compare answers across AI models: How it works | Mint

<https://www.livemint.com/technology/tech-news/perplexity-unveils-model-council-to-compare-answers-across-ai-models-how-it-works-11770651545326.html>

36 37 38 50 51 54 56 67 74 Engineers Reduce GenAI Hallucinations with Model Consensus | Dr Amanda J. Williamson posted on the topic | LinkedIn

https://www.linkedin.com/posts/amandaondata_ai-genai-machinelearning-activity-7425655002800693249-BJXH

39 40 41 59 60 Perplexity's Model Council : r/perplexity_ai

https://www.reddit.com/r/perplexity_ai/comments/1qxklpf/perplexitys_model_council/