

2026年後半～2027年における米中陣営の次世代LLMブレイクスルー予測と構造的比較分析

Gemini 3.7 Flash

人工知能(AI)研究開発の最前線は、単なる事前学習(Pre-training)時の計算資源投入に依存したモデルパラメーター拡大から、ポストトレーニングおよび推論時計算(Test-Time Compute / Inference Scaling)、さらにはアーキテクチャのハイブリッド化とエージェント基盤(Agent Harness)の刷新へと、その重心を劇的に移動させつつある。2026年後半から2027年にかけての期間は、従来の「Chinchillaスケール法則」が物理的・経済的限界を迎える一方で、米国および中国の双方の陣営が異なる構造的アプローチによって不連続な技術的ブレイクスルーを達成する重要な転換期となる¹。

米国陣営は、メガワットからギガワット規模へと拡大するインフラストラクチャと閉鎖型・高機能フロンティアモデルを中心とし、認知力や持続的記憶、高度なリアルタイム推論の統合を目指している⁴。対照的に、中国陣営は厳格化する米国からの半導体輸出規制に直面しながらも、アルゴリズムの極限的な効率化、オープンウェイト(重み公開)モデルの世界的普及、そして国産半導体と独自ソフトウェアスタックの垂直統合によって、米国の技術的優位に肉薄している¹。本分析では、2026年後半から2027年にかけて予想される両陣営の技術的突破、アーキテクチャ変革、インフラストラクチャ上の制約、および地政学的・産業的波及効果を包括的に検証する。

技術的ブレイクスルーのコア・パラダイムと共通課題

2026年後半から2027年のAIモデル進化を決定づける背景には、事前学習データの枯渇、電力供給の逼迫、そして幾何級数的なコスト増大という「スケール法の壁(Scaling Wall)」が存在する¹。両陣営はこの物理的制約を克服するため、類似のパラダイムシフトを推進している¹。

スケール法・ウォールの突破とポストトレーニングの深化

人間が生成した高品質なテキストデータ(約9～27兆トークン)の枯渇問題は、計算最適(Compute-optimal)な事前学習ポリシーにおいて2026年から2027年の間に完全な限界点に達すると分析されている¹。事前学習モデルの過剰学習(Overtraining)は2028年頃までタイムラインを延長可能であるものの、収穫漸減と過学習のリスクを伴う²。さらに、最先端モデルの事前学習コストはGPT-3の約330万ドルから数億ドルへと跳ね上がっており、エネルギー消費量の増大とともに経済的持続可能性を直撃している¹。

このスケール法の壁を打ち破る技術的軸心が、ポストトレーニング技術と推論時計算(Test-Time Compute: TTC)の高度化である¹。強化学習(RLHF、GRPO、Pure RL)やプロセス監視型報酬モデル(Process-Supervised Reward Models)の進化により、事前学習データの単純増大に頼ることなくモデルの論理推論能力を引き出す手法が確立されている¹。モデルが解答を生成する際に潜伏表現内または連鎖的思考(Chain of Thought: CoT)内で探索空間を広げるPass@kやモンテカルロ木探索(MCTS)的探索の最適化により、数分から数時間に及ぶ複雑な数学・プログラミング・科学的問題解決が可能となっている¹²。

しかし、推論時計算の拡大はトレーディングオフをもたらす。推論時間とコストが指数関数的に増大するため、真に実用可能なタスク範囲と経済的実現性の間に乖離が生じ始めており、2026年後半以降は「トークン単価あたりの推論効率」がモデル性能の最大評価指標となることが示唆されている¹³。

アーキテクチャのハイブリッド化：TransformerとSSM/Mambaの融合

従来のTransformerアーキテクチャが抱える最大の技術的課題は、文脈長（Context Length）に対する計算量およびメモリ帯域幅の2次関数的（ $O(N^2)$ ）増加である³。2026年から2027年にかけての主要モデルでは、Transformerの局所的アテンションメカニズムと、状態空間モデル（State Space Model: SSM / Mamba-2）またはGated DeltaNetなどの線形計算機構（ $O(N)$ ）を融合させたハイブリッドアーキテクチャへの全面移行が進行している³。

1つのTransformerアテンションブロックに対して複数のMambaブロックを配置するようなハイブリッド設計（Mamba-2-HybridやQwen3-nextなど）は、数百万トークンを超える超長文脈におけるKVキャッシュの圧迫と推論デコード相のメモリ帯域制約を劇的に緩和する³。これにより、リアルタイムでのセンサーデータ処理、長時間の音声・映像ストリーミング解析、数百万行のコードベース全体を対象とした自律リファクタリングが標準機能化しつつある³。

エージェント・ハーネス（Agent Harness）とオーケストレーション基盤の刷新

2026年における決定的な認識の転換は、「単体モデルの能力向上」から「モデルを包摂するシステム層（Agent Harness）の品質」へのシフトである¹⁷。長時間のマルチステップタスクを実社会環境で自律実行する長期タスク実行エージェント（Long-Horizon Agents）において、ボトルネックはモデルの知能そのものよりも、状態管理、コンテキストエンジニアリング、エラー回復、および環境との外部インターフェースを提供するハーネス構造に移行した¹²。

実践的な研究では、モデル自体を変更することなく、編集ツールのフォーマットやコンテキスト管理構造（OpenClawやNemoClawなど）を変更するだけで、複雑なコーディングベンチマークのスコアが1桁（6.7%）から大幅（68.3%）に跳ね上がる現象が実証されている¹⁷。また、Model Context Protocol（MCP）などの標準化通信プロトコルの広範な採用により、企業環境におけるツール連携とマルチエージェント協調の自動化が進んでいる³。

ただし、市場の急拡大の一方で、マルチエージェントシステムの調整失敗（相互矛盾、コンテキストの破損、コストの爆発）に起因し、2027年末までにエンタープライズ領域におけるエージェントプロジェクトの40%以上が中止・見直しに追い込まれるとの予測も示されており、堅牢な調停アルゴリズムとガバナンスが次世代モデルの必修要件となっている¹⁸。

米国陣営（US Camp）の次世代LLMロードマップとブレイクスルー予測

米国陣営（OpenAI, Anthropic, Google DeepMind, Meta, xAI）は、巨大資本に基づくギガワット規模のデータセンター拡張と、高度な閉鎖型フロンティアモデルの開発を主導している⁴。2026年後半から2027年にかけて、認知機能の持続性、全モダリティのネイティブ統合、そして自律科学研究能力の獲得に重点が置かれる⁵。

米国主要ラボの競争力と技術ロードマップ

OpenAIは、2026年中にGPT-5.5(Spud)およびGPT-5.6の順次投入を経て、次世代基盤モデルコードネーム「Astra」または「GPT-6」の開発に全力を注いでいる⁶。テキサス州アビリーンで稼働を開始した「Stargate I」サイトなどの数ギガワット(5GW以上)規模のインフラを活用し、数十兆パラメータ規模のモデル群を高度なMoE構造で運用すると目される⁴。特筆すべき突破口は「持続的認知記憶(Persistent Cognitive Memory)」と「自律適応機能」であり、セッションごとの文脈リセットを脱し、ユーザーの長期的目標を暗黙的に記憶・更新するシステムへと進化する⁶。さらに、数学・理論計算機科学の未解決問題の解法や証明を自律生成する(Astraプロトタイプによる10件の未解決問題解決とLean 4形式証明など)など、自律的な科学的発見を推進するAIとしての実用化が本格化している⁶。

Anthropicは、2026年半ばにかけて「Claude Sonnet 4.6」、「Opus 4.8」、さらに巨大スケールの「Claude Mythos 5」(推定10兆パラメータ級)および「Claude Fable 5」を展開している⁵。同社の次世代モデルの最大の特徴は、極めて高い安全性・サイバーセキュリティ防御機能と、数日～数週間に及ぶ複雑なソフトウェアエンジニアリング作業(SWE-bench Pro等での標準的優位性)を自律完遂する能力である⁵。また、米安全保障当局との緊密な連携のもと、モデルの出力および推論プロセスに対する厳格なアライメント技術(Constitutional AIの拡張)を組み込んでおり、軍事・高度インフラ領域での自律的ガバナンス能力を強化している⁸。

Google DeepMindは、「Gemini 3.5 Pro / Flash」から「Gemini 4」への移行期において、同社独自のTPU(Ironwood / Trillium)インフラを活用した「全モダリティ・ネイティブ融合(Native Omnimodal Integration)」と「オンデバイス・クラウドハイブリッドの圧縮技術」で他社を画している⁴。テキスト、画像、音声、リアルタイムビデオ入力に加え、触覚や環境センサーデータを直接トークン化して処理する能力を獲得している⁵。さらに、Googleが開発した次世代メモリ圧縮アルゴリズムは、ロングコンテキスト維持に必要なメモリ帯域幅を飛躍的に削減し、企業向けNotebookLMやAgent Engineの運用コストを大幅に引き下げている⁵。

MetaおよびxAIは異なる角度から先行している。Metaは「Llama 4」および「Llama 5」を通じてオープンウェイト戦略を維持しつつ、超大型のマルチモーダルMoEモデルを投入している⁵。xAIは「Grok 4.20 / Grok 5」において、X(旧Twitter)プラットフォームからの即時リアルタイムデータストリームと巨大クラスタ(Colossus等)を組み合わせ、世界状況を即座にモデル化するリアルタイム世界モデル(World Model)構築に注力している⁵。

インフラストラクチャ上の制約: 電力ボトルネックとGigawattスケールの課題

米国陣営が直面する最も深刻なボトルネックは、半導体チップの確保以上に「データセンターの電力確保と系統接続(Interconnection Queues)」へ移行している⁷。2027年までにAIデータセンターは世界全体で新たに92ギガワット(GW)以上の電力を必要とすると予測されている¹⁰。

変電所の新設や138kV送電線の引き込みには36～60ヶ月のリードタイムを要するため、Nvidia GB200/RubinやGoogle TPUの最新ラックが納品されても設置場所が受電できない現象が発生している⁴。これに対し、Hyperscaler各社は小型モジュール炉(SMR)や既存原子力発電所との直接PPA(電力購入契約)、Behind-the-meter(メーター裏)ガスタービン発電の導入を急速に進めている⁷。また、ハードウェア構造上の制約として、Nvidiaの高速相互接続システム(NVLink/NVSwitch)に依存するスケールアップ・ドメインのコストが過大になるため、TPUのような独自Gigawattスケール・インターコネクト網を持たない事業者は、数兆～数十兆パラメータのフラッグシップモデルの低遅延サー

ピングにおいて、2027年から2028年にかけて一時的なスケーリング停滞期を迎えるリスクが指摘されている⁴。

中国陣営(China Camp)の次世代LLMロードマップと独自エコシステム

中国陣営(DeepSeek, Alibaba Qwen, Moonshot AI, Zhipu AI, Huawei, ByteDance等)は、米国の先進半導体輸出規制やHBM(高帯域幅メモリ)入手制限という極めて厳しい環境下で、独自の進化モデルを確立した¹。それは、「オープンウェイトモデルによる世界市場の席巻」、「アルゴリズム的超効率化」、そして「国産シリコンとの垂直統合」を三位一体としたエコシステムである¹。

オープンウェイト戦略とモデル性能の飛躍

中国のLLM開発は、2026年半ば時点で世界最先端のオープンウェイトモデル市場を実質的に主導している³。「DeepSeek V4-Pro(1.6兆パラメータ)」、「Qwen3-next」、「Moonshot Kimi K3」、「Zhipu GLM-5.2」などのモデル群は、米国陣営の最上位クローズドモデル(GPT-5.5やClaude Opus 4.7等)に匹敵する性能を、数分の一から十分の一のAPI/推論コストで提供している³。

Gartnerの予測によると、中国製オープンウェイトモデルのグローバル企業における採用率は、2025年の5%から2027年には50%へと急激に跳ね上がるとされている⁸。企業がデータプライバシーやベンダーロックインを回避するため、オンプレミスやプライベートクラウドで自社ホスティング可能な高品質オープンモデルを選択する傾向が強まっており、中国陣営のモデルがその需要を吸収している⁸。DeepSeekは、Pure RL(純粋強化学習)による思考チェーンの自己進化をさらに推し進め、プログラミングおよび数学領域で商用最先端モデルを超越する探索能力(Pass@kの最大化)を達成している¹。また、Sparse MoEのルーティング機構を微細化し、アクティブパラメータ率を極限まで下げることで、低スペックインフラでの動作を可能にしている¹。

Alibaba Qwenは、Gated DeltaNet等の線形ハイブリッドアーキテクチャを全面採用し、100万トークンを超えるコンテキストを極めて低いVRAM消費量で高速処理する³。Apple Intelligenceの中国本土版基盤モデルとして採用されるなど、消費者デバイスへの組み込みでも先行している³⁴。

Moonshot AI(Kimi K3)およびZhipu AI(GLM-5.2)は、長文脈での推論記憶と、マルチエージェント型コーディング・ビジネスプロセス自動化に特化している⁸。特にGLM-5.2は「SWE-bench Pro」等の高度実績指標において、GPT-5.5を低コストで凌駕する性能を実証している²⁴。

国産シリコンと「Tau Scaling Law」による技術的突破

中国陣営が抱える最大の構造的アキレス腱は、TSMCの先進プロセス(5nm/3nm以下)へのアクセス遮断と、SK Hynix等からの最新HBM3E/HBM4の供給制限である²⁹。しかし、中国はこの制約を回避するための技術的バイパス戦略を急速に構築している⁸。

Huaweiは「Ascend 910C」、「910D」(5nm級/多ダイ構成)、および「Ascend 950PR / 950DT」の量産を加速させている⁹。特に2026年に本格投入されたAscend 950PRIは、独自のHBM(HiZQ 2.0、帯域幅4 TB/s)を統合し、CUDA非依存の独自ソフトウェアスタック(MindSpore/CANN)を構築した⁹。FP4精度で1.56 PFLOPSを叩き出し、米国側の輸出規制準拠チップ(Nvidia H20)の約2.8倍の推論性能を16,000ドルという低価格で提供している³²。ByteDanceが56億ドル規模(約35万枚)の大型発注を行ったほか、AlibabaやTencentも大量採用に踏み切っている³²。

単一チップのFLOPSや電力効率でNvidia Blackwell等に及ばないハンディキャップを補うため、

Huaweiは「Tau Scaling Law」を提唱した⁸。これは、3D高度半導体パッケージング技術とウェハレベル統合（CloudMatrix 384等）を利用して多数のダイを超高密度の物理的トポロジーで結合し、HBM供給制約をシステム全体のアグリゲート帯域幅（チップあたり8スタックのHBM3等）で相殺する戦略である⁸。

さらに、ソフトウェアの工夫による効率化だけでは極限に達したと判断したDeepSeekは、2026年中盤より独自のAIアクセラレータ設計に着手した²⁷。これにより、モデルアーキテクチャ（Sparse MoEやPure RL）と半導体命令セットを完全に一致させる独自の垂直統合スタイルを確立しつつある²⁷。

なお、現状において、これら国産チップ（Ascend等）はポストトレーニングおよび推論フェーズでは極めて強かに機能しているものの、数万～数十万枚規模の超巨事前学習クラスタにおける信頼性・歩留まりには依然課題が残る⁸。そのため、中国陣営は事前学習での不必要なパラメーター拡大を避け、アルゴリズムによる事前学習コストの圧縮と、推論時計算およびポストトレーニングへのリソース集中を戦略的に選択している¹。

米中陣営の包括的比較分析

2026年後半から2027年にかけての米中陣営における次世代LLMの技術・戦略的ポジショニングの差異は、以下の比較構造として整理することができる。

比較軸	米国陣営（US Camp）	中国陣営（China Camp）
主要開発主体	OpenAI, Anthropic, Google DeepMind, Meta, xAI ⁵	DeepSeek, Alibaba (Qwen), Moonshot AI, Zhipu AI, Huawei, ByteDance ¹
モデル公開形態	閉鎖型（API主導・Proprietary）が主流。一部Meta (Llama) がオープン化 ⁵	オープンウェイト（Open-weight）戦略が主流。自社ホスティング市場を統合 ⁸
アーキテクチャ設計	超高パラメータMoE、ネイティブ全モダリティ（Omnimodal）、記憶の持続化 ⁵	Transformer+SSM/Mambaハイブリッド、Gated DeltaNet、極細粒度Sparse MoE ²
スケーリング手法	ギガワット級事前学習 + ポストトレーニング + 時間依存型推論時計算（TTC） ¹	事前学習のアルゴリズム的極限圧縮 + Pure RL + 効率化推論時計算 ¹
基盤ハードウェア	Nvidia Blackwell / Rubin Ultra, Google TPU (Ironwood), AWS Trainium ⁴	Huawei Ascend 910C/910D/950PR, SMIC 7nm/12nm, 自社設計シリコン ⁹

物理・構造的制約	データセンター電力確保(系統接続遅延 36-60ヶ月)、KVキャッシュ帯域 ⁷	米国輸出規制による先進プロセス・HBM調達制限、巨大クラスタの歩留まり ²⁹
グローバル市場の強み	最高峰の認知知能、科学的発見能力、厳格な安全・ガバナンス枠組み ⁵	圧倒的なAPI/推論コストの低さ、オンプレミス展開の容易さ、実用コード生成性能 ⁸

地政学的・産業的波及効果の多角的一省

両陣営の技術進化と地政学的制約の相互作用からは、単純な性能競争を超えた深い構造的変化（二次・三次的影響）が読み取れる⁸。

第一に、米国が提唱する「Pax Silica」AI枠組み（同盟国に対する半導体・AIサプライチェーンの選択要求）と、中国の「AIフルスタック自給自足」政策が衝突することで、世界のAIインフラは完全に二極化しつつある⁸。グローバル企業は、米国の閉鎖型高性能AIサービスを利用する領域と、中国系オープンウェイトモデルをローカルで安全かつ低コストに運用する領域という、二重のシステムアーキテクチャの保持を余儀なくされている⁸。

第二に、推論時計算（TTC）への依存度が高まるにつれ、高度な推論を伴うエージェントの1タスクあたりの実行コストが人間の件費に接近するという経済的壁が出現している¹⁴。この結果、2027年にかけて企業は全自動エージェントの過度な期待から脱却し、極めてコスト効率の高い中国系オープンウェイトモデル（QwenやDeepSeek系）を自社のエージェント・ハーネス（OpenClaw等）に組み込む実用主義へと急速に舵を切る動機が生まれている⁸。

第三に、DeepSeekの成功やHuaweiのTau Scaling Lawが示す通り、ハードウェアのハンディキャップ（7nm vs 2nm）は、アロケーションアルゴリズム、モデル蒸留、Sparse MoE設計、およびPure RLポストトレーニングによって数ヶ月のタイムラグまで圧縮可能であることが実証された¹。この事実は、ハードウェアの封じ込めのみで技術的優位を恒久維持することは不可能であるという地政学的インプリケーションを政策決定者に突きつけている⁸。

結論と戦略的提言

2026年後半から2027年にかけての次世代LLMモデルのブレイクスルーは、単なるモデル規模の拡大（Brute-force Scaling）の終焉と、推論時計算の高度化、アーキテクチャのハイブリッド化、エージェント・ハーネスのシステム化、そしてハードウェア・ソフトウェアの垂直統合を軸とするマルチパラダイムの時代への移行を象徴している¹。米国陣営はStargateに代表されるギガワット級の資本投入と認知科学的アプローチによって汎用人工知能（AGI）および自律科学研究AIの領域を切り拓く一方、中国陣営は厳格な外圧をテコにして、超高効率なオープンウェイトモデルと国産シリコンによる独自エコシステムの完全自立を達成しつつある⁶。

本分析に基づき、今後のAI技術の活用および投資戦略を策定するエンタープライズのCIO、CTO、および政策決定者に対して以下の戦略的提言を提示する。

- 1. シングルベンダー依存からの脱却とマルチモデル・スタックの構築
単一の米国系クローズドモデルに完全依存するシステム設計は、コスト増大および地政学的・規制リスクに対して極めて脆弱である⁸。最高度の推論が必要なコア領域には米国先端モデルを配置しつつ、定常的なエージェント作業やオンプレミス処理には中国系・オープンウェイ

トモデルを配置する柔軟なマルチモデルアーキテクチャの構築が必要不可欠である⁸。

2. モデル単体から「エージェント・ハーネス」への投資シフト
モデルのバージョンアップを受動的に待つ戦略は、費用対効果が低い²²。自社の業務フローを緻密に分解し、コンテキスト保持、エラーリカバリ、ツール連携(MCP標準準拠)を担うハーネス層の標準化とプロンプト・ライブラリの資産化を先行させるべきである¹⁷。
3. インフラ評価軸における電力と推論単価の再定義
AI導入の評価軸を単なるGPU枚数やパラメータ数から、1タスクあたりの受電可能電力(MW)および完遂タスクあたりの推論コスト(Cost per Task)へ再定義する必要がある⁷。特に推論時計算の増大に伴うコスト爆発を抑制するため、ハイブリッドアーキテクチャ(Mamba/SSM融合型)の早期評価と導入が推奨される³。

引用文献

1. LLMOrbit: A Circular Taxonomy of Large Language Models -From Scaling Walls to Agentic AI Systems - arXiv, <https://arxiv.org/pdf/2601.14053>
2. LLMOrbit: A Circular Taxonomy of Large Language Models —From Scaling Walls to Agentic AI Systems - arXiv, <https://arxiv.org/html/2601.14053v1>
3. LLMOrbit: A Circular Taxonomy of Large Language Models —From Scaling Walls to Agentic AI Systems - arXiv, <https://arxiv.org/html/2601.14053v2>
4. Vladimir_Nesov's Shortform - LessWrong, https://www.lesswrong.com/posts/fdCaCDfstHxyPmB9h/vladimir_nesov-s-shortform
5. New AI Model Releases News | April, 2026 (STARTUP EDITION) - Mean CEO's BLOG, <https://blog.mean.ceo/new-ai-model-releases-news-april-2026/>
6. GPT-6 (2026) – Dr Alan D. Thompson - LifeArchitect.ai, <https://lifearchitect.ai/gpt-6/>
7. AI Data Center Power Infrastructure 2026 - GPU Insights -, <https://gpuinsights.net/ai-data-center-power-infrastructure-2026/>
8. Half of China's AI Accelerators Will Be Homegrown by 2030: Gartner Forecasts 'AI Full-Stack' Self-Reliance - BigGo Finance, <https://finance.biggo.com/news/813d32c3-af15-4bd8-8c1e-3a9604808118>
9. 7 Chinese companies are already shipping H100/H200-class AI chips, most IPO'd in the last 6 months. I mapped all of them. - Reddit, https://www.reddit.com/r/LocalLLaMA/comments/1udkxde/7_chinese_companies_are_already_shipping/
10. 2026 Global Semiconductor Industry Outlook - Deloitte, <https://www.deloitte.com/us/en/insights/industry/technology/technology-media-telecom-outlooks/semiconductor-industry-outlook.html>
11. R1": Learning to Retrieve and Answer Step-by-Step with Synthetic Data - arXiv, <https://arxiv.org/html/2605.01248v1>
12. Towards Long-Horizon Agents: A Survey - OpenReview, <https://openreview.net/pdf?id=HyhfhbWGh>
13. Artificial Intelligence - arXiv, <https://arxiv.org/list/cs.AI/new>
14. Writing - Toby Ord, <https://www.tobyord.com/writing>

15. The Stakes of AI: Progress, Trajectories, Impacts - Oxford Academic,
<https://academic.oup.com/book/61416/chapter/533867765>
16. New IBM Granite 4 Models to Reduce AI Costs with Inference-Efficient Hybrid Mamba-2 Architecture - InfoQ,
<https://www.infoq.com/news/2025/11/ibm-granite-mamba2-enterprise/>
17. Agent Harness for Large Language Model Agents: A Survey[v1] | Preprints.org,
<https://www.preprints.org/manuscript/202604.0428/v1>
18. LLM-Based Multi-Agent Orchestration: A Survey of Frameworks, Communication Protocols, and Emerging Patterns - Preprints.org,
<https://www.preprints.org/manuscript/202604.2147>
19. LLM-Based Multi-Agent Orchestration: A Survey of Frameworks, Communication Protocols, and Emerging Patterns - MDPI,
<https://www.mdpi.com/1999-5903/18/6/326>
20. NotebookLM in 2026: Gemini 3, Video Overviews, and Gemini App Integration | TIMEWELL, <https://timewell.jp/en/columns/ai-gemini-google-notebooklm-4>
21. OpenAI ChatGPT 2026: GPT-5 Features & GPT-6 Predictions - mindlifty,
<https://mindlifty.com/openai-chatgpt-2026-gpt-5-features-gpt-6-predictions/>
22. GPT-6はいつ出る？Astraの正体とGPT-5.6の現在地 - 株式会社Uvation,
<https://uravation.com/media/gpt6-spud-release-date-enterprise-guide-2026/>
23. 令和7年度 研究報告書 電波ばく露に関する標準的な研究手法の確立及び 中間周波電磁界の神 - 総務省 電波利用ポータル,
https://www.tele.soumu.go.jp/resource/j/ele/body/report/pdf/r7_06.pdf
24. AI Breakthroughs in June 2026: The Mid-Year Update Every Founder and CIO Needs, <https://kersai.com/ai-breakthroughs-june-2026-mid-year-update/>
25. Treaty-Following AI - Institute for Law & AI, <https://law-ai.org/treaty-following-ai/>
26. (PDF) Treaty-Following AI - ResearchGate,
https://www.researchgate.net/publication/398367041_Treaty-Following_AI
27. Deep Dive into Today's AI News | July 8, 2026 | | 宮野宏樹 - note,
<https://note.com/hirokimiyano/n/ne60ab48431f3?hl=en>
28. China's Top AI Models in 2026: DeepSeek, Qwen, Kimi, Doubao and the New AI Race | by MEXC Learn Editor | Coinmonks - Medium,
<https://medium.com/coinmonks/chinas-top-ai-models-in-2026-deepseek-qwen-kimi-doubao-and-the-new-ai-race-08083866ac5a>
29. Serving LLM on Huawei CloudMatrix - LessWrong,
<https://www.lesswrong.com/posts/jWGZeui4LHfzzCgSs/serving-llm-on-huawei-cloudmatrix>
30. Escape Velocity — When Does China's AI Stack Break Free? - Futurum Research,
<https://futurumgroup.com/insights/escape-velocity-when-does-chinas-ai-stack-break-free-8/>
31. China's AI Moment: Kimi K3, DeepSeek V4, and Qwen Are Rewriting the Global Model Race,
<https://www.elser.ai/news/chinas-ai-moment-kimi-k3-deepseek-v4-qwen-2026>
32. Huawei AI Chip Revenue Hits \$12B in 2026: The Nvidia Alternative Enterprise CTOs Are Evaluating - Rajesh Beri,
<https://www.beri.net/article/huawei-ai-chip-revenue-12-billion-nvidia-alternative>

33. Issues | AI News - Smol AI News, <https://news.smol.ai/issues/>
34. China AI Weekly: DeepSeek's Funding Freeze and AGI Gambit, Kimi, <https://www.bighatgroup.com/blog/china-ai-weekly-2026-07-25/>
35. US to tell partners they must pick sides in AI race with China - The Economic Times, https://m.economictimes.com/news/international/global-trends/us-to-tell-partners-they-must-pick-sides-in-ai-race-with-china/articleshow/133255962.cms?utm_source=Google_Newsstand&utm_campaign=RSS_Feed&utm_medium=Referral
36. On DeepSeek and Export Controls - Dario Amodei, <https://darioamodei.com/post/on-deepseek-and-export-controls>