

中国AI急進の意味と日本の打ち手

日経記事が示した本質は、中国AIの「性能」だけでなく、「安さ」「自社運用しやすさ」「API互換性」「政治的に止まりにくい供給形態」が同時に効いている点にあります。米国最先端モデルは依然強力ですが、2026年6月のAnthropicモデル停止は、企業にとって“能力”より“可用性”が重要になる局面を可視化しました。日本は総花的な汎用LLM競争ではなく、製造・ロボット・金融・医療など規制・現場密着領域のパーソナルAI、AI-Readyデータ基盤、国内運用可能なオープンモデル活用、評価標準と契約実務の整備を一体で進めるべきです。 ¹

エグゼクティブサマリー

日経記事の見立ては、周辺ソースで概ね裏づけられる。Anthropicは6月12日、米政府の輸出管理指示を受けてFable 5とMythos 5を全顧客向けに急停止し、その後6月30日に規制解除を受けてFable 5を再配備した一方、Mythos 5は引き続き限定提供が中心である。この空白期に、中国系のオープンウェイト／移行容易なモデル群が、価格と可搬性を武器に企業評価を一気に進めた、という構図は合理的である。 ²

とくに重要なのは、コスト差が単なる値下げではなく、技術アーキテクチャ由来である点だ。DeepSeek-V4は1.6T総パラメータでも49Bしか活性化しないMoE、圧縮・疎化された長文脈注意機構、FP4/FP8混合精度、KVキャッシュ最適化を組み合わせ、1M文脈でV3.2比10%のKVキャッシュ、27%の単トークン推論FLOPsを公表する。Kimi K2.7 Codeも1T/32B活性のMoEとMLAを採用し、MooncakeでKVキャッシュ中心の分離サービングを実装した。GLM-5.2は1M文脈に加え、IndexShareで1M文脈時のper-token FLOPsを2.9倍削減すると説明する。これは「安いが弱い」ではなく、「推論経済性そのものが設計目標」という意味である。 ³

日本への示唆は明確だ。第一に、政府は“国産フロンティア汎用モデル一本足打法”ではなく、製造・ロボット・インフラ・医療の高品質データをAI-Ready化し、国内クラウド／オンプレで動かせる基盤を整えるべきである。第二に、企業は閉鎖系米国モデルとオープン系中国モデルを二者択一で考えず、用途別に最適ルーティングするハイブリッド運用へ移るべきである。第三に、法務・セキュリティ・調達の統制をモデル開発競争と同格で扱い、越境移転、契約、評価標準、監査ログ、モデル由来リスクを管理可能にする必要がある。日本には、製造現場の暗黙知、特許、企業R&D、ロボティクスの厚みという「勝ち筋」がある。今やるべきは、これをAIの一般論ではなく、産業別のデータ循環と標準に落とし込むことである。 ⁴

要点

- 中国AIの競争優位は、ベンチマークの一点突破よりも、**安価な推論・オープンウェイト・移行容易性・長文脈最適化**の総合力にある。 ⁵
- 2026年6月のAnthropic停止は、**海外依存モデルの可用性リスク**を一気に顕在化させた。これは日本企業にもそのまま当てはまる。 ⁶
- 「20分の1」は比較対象次第だが、**Mythos Preview基準なら概ね成立し、DeepSeekの低価格帯はそれ以上である**。現在のFable/Mythos 5基準でも、中国モデルはなお大幅に安い。 ⁷
- 日本の強みは、汎用チャットではなく、**製造データ、ロボティクス、品質管理、現場知見、特許、企業R&D**にある。政府自身もその方向で政策を動かし始めている。 ⁸
- 日本企業の実務解は、**閉鎖系最先端APIを残しつつ、オープンモデルを国内環境で運用し、ルーティング・キャッシュ・契約・監査を標準化**することである。 ⁹

市場の現在地と記事の検証

まず、日経原文は取得制限のため全文を直接確認できなかった。本レポートは、日経見出し・要旨公開断片、Anthropic公式発表、Reuters、各モデルの公式ドキュメントを突き合わせて検証している。そのうえで、記事の中核命題、すなわち「Mythos停止の間に、中国AIが価格とオープン型の強みで米企業評価を進めた」という骨子は、かなり整合的である。 ¹⁰

Anthropicは6月12日に、米政府の輸出管理指示により、Fable 5とMythos 5を外国籍者に提供できなくなり、実運用上の理由から両モデルを全顧客向けに停止した。6月30日に輸出規制は解除されたが、Fable 5はグローバル再開、Mythos 5はなおGlasswing参加者や一部の“trusted”組織向けが中心である。つまり、停止は一時的でも、「政治判断で最先端モデルが止まりうる」こと自体が企業調達リスクとして実証された。 ¹¹

その間、企業側では、能力一辺倒からコスト最適化への視線が明確に強まった。Reutersは、AIコーディング費用が2028年までに平均開発者給与を上回るというGartner推計を紹介し、OpenRouter上でオープンソース・トークン比率が2026年1月の34%から6月に65%へ上昇したと報じた。また、人気上位4モデルがすべて中国系で、DeepSeekが最上位だとしている。これは単発の“話題”ではなく、企業FinOpsの論理がモデル選定を変え始めたことを示す。 ¹²

米企業の乗り換え事例として最も象徴的なのはCoinbaseだ。公開断片ベースでは、Brian Armstrong CEOが、社内LLMゲートウェイのデフォルトとしてGLM 5.2やKimi 2.7のようなオープンウェイトモデルを試し、AI支出をほぼ半減させたと説明している。ReutersもArmstrongを、より安価な小型・廉価モデル活用を訴える経営者の代表例として挙げている。公開資料の性質上、完全な移行か、デフォルト変更か、用途限定かは厳密に峻別すべきだが、少なくとも“低コスト中国モデルを企業ユースケースに本気で組み込み始めた”ことは確認できる。 ¹³

価格面では、日経の「20分の1」は比較対象により意味が変わる。現行のClaude Fable 5 / Mythos 5は\$10入力・\$50出力/百万トークンだが、GLM-5.2は\$1.4・\$4.4、Kimi K2.7 Codeは\$0.95・\$4.0、DeepSeek-V4-Proは\$0.435・\$0.87、DeepSeek-V4-Flashは\$0.14・\$0.28である。単純計算では、Fable/Mythos 5比でGLM-5.2は約7～11分の1、Kimiは約10～13分の1、DeepSeek-V4-Proは約23～57分の1、DeepSeek-V4-Flashは約71～179分の1だ。さらに、AnthropicがMythos Preview参加者向けに示した\$25入力・\$125出力という水準を基準にすると、GLM-5.2は約18～28分の1、Kimiは約26～31分の1となり、「20分の1」は十分に説明可能である。要するに、**見出しの方向性は妥当だが、比率はモデルとワークロード次第**というのが実証的な読み方である。 ¹⁴

技術とコストの分解

中国勢の低価格は、単なるダンピングではなく、設計と運用の積み重ねで説明できる。基礎には知識蒸留がある。Hintonらの古典的な蒸留は、高性能な教師モデルの挙動をより軽い学生モデルへ圧縮し、配備コストを下げる発想だった。2026年の市場では、この“蒸留”が技術効率化の手段であると同時に、モデル由来や知財の火種にもなっている。OpenAIやAnthropicは中国企業による蒸留的な能力抽出を公然と問題視しており、蒸留は今や性能最適化と地政学リスクの両面を持つ。 ¹⁵

次にMoEである。MoEは総パラメータを巨大化しても、各トークンごとに一部エキスパートしか活性化しないため、推論計算量を抑えやすい。DeepSeek-V4-Proは1.6T総パラメータで49B活性、V4-Flashは284Bで13B活性、Kimi K2.7 Codeは1Tで32B活性、Kimi K2も1Tで32B活性だ。DeepSeekMoEの研究自体も、より細粒度の専門化と共有エキスパートで、同等性能をより少ない計算で達成する方向を示している。つまり、**能力の上限は大きく、実行時負荷は小さく**という、企業導入に直結する経済性がMoEの本質である。 ¹⁶

長文脈コストを押し下げるのが、KVキャッシュ最適化、疎化、サービング設計だ。DeepSeek-V2/V3のMLAはKVキャッシュを圧縮し、帯域制約の強いデコード段で効く。ハードウェア分析でも、MLAがKVキャッシュ削

減と帯域需要低減に効くことが示されている。KimiのMooncakeは、prefillとdecodeを分離し、GPUだけでなくCPU/DRAM/SSD資源をKVキャッシュに使うことで、長文脈シナリオで実スループットを大きく高めた。GLM-5.2はIndexShareで疎注意のインデクサ再利用を行い、1M文脈でのper-token FLOPsを2.9倍削減すると公表する。DeepSeek-V4はCSA/HCAのハイブリッドで、1M文脈時にV3.2比10%のKVキャッシュで済むとする。これらはみな、**エージェント時代に増える“長く・多段で・繰り返す”推論を安くする技術**である。 ¹⁷

量子化も重要だ。AWQは、活性分布に着目して重要な重みチャンネルだけを保護し、低ビット化の精度劣化を抑える代表例である。現実の製品では、DeepSeek-V4がFP4+FP8混合精度を採用し、KVキャッシュや重みの圧縮が推論コストをさらに押し下げている。つまり、中国モデルの価格破壊は、「MoEで活性パラメータを減らし、MLA/疎注意でメモリを減らし、量子化で帯域と格納を減らし、サービングでクラスタ資源を使い切る」という多層最適化の成果とみるべきである。 ¹⁸

主要モデル比較

モデル	技術上の特徴	価格の目安	ライセンス・公開性	運用要件	主な根拠
Claude Mythos 5	サイバーセキュリティ特化。限定提供。Fable/Mythosとも6月停止を経験	\$10入力 / \$50出力 / 1M tok	クローズド	APIのみ。Mythosはtrusted access中心	¹⁹
Claude Fable 5	汎用フロンティア。1M文脈、適応思考	\$10入力 / \$50出力 / 1M tok	クローズド	APIのみ	²⁰
GLM-5.2	1M文脈、IndexShareで1M時per-token FLOPsを2.9倍削減、MIT	\$1.4入力 / \$4.4出力 / 1M tok	MITオープン	vLLM / SGLang等で自社運用可	²¹
Kimi K2.7 Code	MoE 1T総/32B活性、256K文脈、MLA、思考トークンをK2.6比約30%削減	\$0.95入力 / \$4.0出力 / 1M tok	Modified MIT	OpenAI/Anthropic互換。Claude Code系にも接続可	²²
DeepSeek-V4-Pro	MoE 1.6T総/49B活性、1M文脈、CSA/HCA、V3.2比10% KV cache	\$0.435入力 / \$0.87出力 / 1M tok	MIT	vLLM / SGLangで自社運用可。OpenAI/Anthropic互換API	²³
DeepSeek-V4-Flash	MoE 284B総/13B活性、1M文脈、高速・低コスト	\$0.14入力 / \$0.28出力 / 1M tok	MIT	同上	²⁴

API課金の差は、実ワークロードに直結する。たとえば「月100M入力・20M出力トークン」という、企業のコーディング支援や社内検索では珍しくない規模を単純計算すると、Claude Fable 5は約\$2,000、GLM-5.2は約\$228、Kimi K2.7 Codeは約\$175、DeepSeek-V4-Proは約\$60.9、DeepSeek-V4-Flashは約\$19.6になる。しかも各社ともキャッシュヒット時の値引きがあるため、反復タスクでは差がさらに広がる。**導入可否を左右するのは、性能スコアより“タスク完了あたりコスト”である。** ²⁵

自社運用とAPIの選択は、利用率とデータ機密性で分かれる。APIは立ち上がりが速く、最新モデルの追従が容易だが、可用性・価格改定・利用制限の影響を強く受ける。自社運用は、GPU、推論スタック、観測、SRE/ML Ops、脆弱性対応、監査の固定費を負うが、モデル停止リスクと越境移転リスクを抑えやすい。中国

オープンモデルの実務上の強さは、**APIで試し、良いものを国内環境に移せる**という“出口の広さ”にある。これは閉鎖API主体のモデルにはない調達オプションである。 26

業種別インパクトと日本固有資産

金融では、生成AI導入をためらうリスクそのものが競争力低下につながる。金融庁は、リスクや規制面から利活用に躊躇する声がある一方で、「チャレンジしないリスク」も踏まえるべきだと明記し、2025年以降AI官民フォーラムを継続している。同時に、専門人材不足、継続モニタリング、入出力管理、自社専用環境の必要性も指摘している。したがって金融機関では、中国オープンモデルを**社内コーディング、規程検索、監査補助、社内ナレッジ要約**にまず適用し、対顧客説明・売買判断・審査自動化のような高規制領域は閉鎖系最先端や人手審査と組み合わせるのが現実的である。 27

製造では、日本の優位はむしろこれから生きる。経産省は、設計図面や製造ノウハウなど日本の強みである現場データがまだAI-Ready化されておらず、しかし「質の高い現場データ」と「継続的に集積・活用するインターフェース」が勝ち筋だと整理している。さらにGENIACでは、製造業データ等のAI-Ready化、複数企業にまたがるデータ利活用基盤、ロボット基盤モデル、マルチモーダル・フィジカルAI事業を採択・開始している。日本の製造業は、汎用会話AIで米中に勝つ必要はない。**工場、自動化、品質、保全、調達、工程最適化で勝てばよい。** 28

IT・ソフトウェアでは、中国モデルの衝撃は最も直接的だ。Kimi K2.7 Code、GLM-5.2、DeepSeek-V4はいずれも長文脈・エージェント・コーディングを前面に出し、Claude CodeやAnthropic互換、生のOpenAI SDK互換、vLLM/SGLang運用を揃えている。これは「優れたモデル」ではなく「置き換えやすいモデル」である。CoinbaseのようにLLMゲートウェイでルーティングし、高価モデルを複雑タスクに限定し、通常タスクを廉価モデルへ流す運用は、日本のIT企業でもすぐ実装できる。ROIは、モデル精度差より**人件費あたりの開発速度、レビュー時間、障害対応時間**で測るべきだ。 29

ヘルスケアでは、コスト優位だけで即採用は危険である。PMDAの専門部会は、AI活用医療機器について、バイアス、市販後学習、評価データ再利用、臨床DB整備といった課題を整理している。PMDA自身は2026年4月にCopilotの全職員導入を開始したが、これはまず業務生産性向上の文脈であり、診断・治療の自律判断を直接意味しない。したがって医療分野での初期ユースケースは、**薬事・治験文書、学会要約、品質文書、問い合わせ一次応答**が先であり、臨床判断そのものは厳格な検証と規制適合が前提になる。 30

日本固有の資産は、データと知財の“厚み”にある。WIPOは、日本が2025年のGlobal Innovation Indexで12位、特許ファミリーで世界2位、企業によるR&Dで世界3位だと示している。これは、基盤研究や事業投資がなお強いことを意味する。問題は、これがAIの学習可能な形に変換されていないことだ。ゆえに日本の優先順位は、**モデルを増やすことより、データをAI-Readyにし、権利処理し、再利用可能な形で流通させること**である。 31

セキュリティ・規制・サプライチェーン

最大の論点はデータ主権である。日本の個人情報保護委員会は、外国にある第三者へ個人データを提供する際、原則として本人同意が必要であり、その際には提供先国名、当該国の制度、受領者の保護措置などの説明が必要だとしている。したがって、日本企業が中国系APIを直接使う場合、**単に“安いから使う”では済まない**。個人データ、顧客情報、機微情報が含まれるユースケースでは、国内環境への自社運用または匿名化・削除・合成データ化が必須になる。 32

次に可用性リスクである。6月のAnthropic停止は、閉鎖系モデルにおける地政学リスクを突きつけたが、中国側でも7月7日時点で、当局が最先端AIモデルの海外アクセス規制を検討しているとReutersは報じている。つまり、**米国依存も中国依存も、そのままでは危うい**。日本企業に必要なのは「どちらかに乗る」ことではなく、「どちらが止まっても業務が止まらない構成」にすることである。 33

サプライチェーン面では、輸出規制が中国企業の最適化を促した側面がある。Reutersは、DeepSeekが米輸出規制でNvidia最先端チップへアクセスできず、Huawei Ascend活用や専用推論チップ開発へ向かっていると報じた。DeepSeek-V4のHuawei最適化も別のReuters報道で確認される。これは日本にとって、**GPUが足りないからできない**という発想が通用しにくくなっていることを示す。計算資源不足は重要だが、同時に、アーキテクチャとサービングの工夫でかなり吸収できる。 ³⁴

日本の規制環境は、全面禁止よりリスクベース運用に向いている。AIセーフティ・インスティテュートは2024年に設立され、AI事業者ガイドラインはAI開発・提供・利用に必要な基本的考え方を示している。金融庁もユースケース別・リスクベースの対話を進めている。したがって日本は、EU型の事前許認可偏重でも、米国型の市場任せでもなく、**業界別の評価観点、契約、監査、ログ、テストを実装する“実務型ガバナンス”**を磨くべきである。 ³⁵

リスク評価表

リスク	発生確率	影響度	日本企業への意味	主要対策	主な根拠
海外APIの停止・制限	中～高	高	モデル停止で開発・業務継続が止まる	マルチモデル化、代替モデル常備、自社運用先の確保	³⁶
個人データの越境移転	高	高	APPI対応不備で法務・レピュテーション問題	国内運用、匿名化、入力制御、DPA/契約整備	³⁷
モデル由来・知財紛争	中	中～高	調達先の正統性が将来の利用継続性に影響	ベンダー審査、学習由来条項、補償条項、監査権	³⁸
サプライチェーン偏在	中	高	GPU・チップ・クラウド障害がTCOと継続性に影響	国内外複線化、オープン推論基盤、エッジ活用	³⁹
機密コード・設計情報の漏えい	中	高	製造・金融・医療で致命的	社内閉域、RAG境界制御、監査ログ、権限分離	⁴⁰
価格改定・使用量暴騰	高	中～高	ROI悪化、部門横断導入が頓挫	ルーティング、キャッシュ、コストSLO、FinOps	⁴¹
高規制業界での誤判断	中	高	金融・医療で説明責任を満たせない	人間承認、用途限定、検証指標、モニタリング	⁴²

日本への提言と実行ロードマップ

短期の提言

政府は、既存のGENIACを「計算資源支援」中心から「産業データ流通・評価・調達」中心へ拡張すべきだ。具体的には、製造、ロボット、金融、医療の4領域で、匿名化済み・権利処理済み・AI学習可能なデータセット整備に補助を厚く配分する。加えて、AISIと業界団体が共同で、**モデル評価の共通様式**を定めるべきである。性能だけでなく、越境移転、監査ログ、モデル由来、再現性、レッドチーム結果、切替容易性を1枚で比較できるようにする。これは日本のガバナンスを“規制”ではなく“比較可能性”で強くする。 ⁴³

企業は、いま直ちに全社共通の**モデル選定基準**を作るべきだ。最低限、用途、データ機密性、応答品質、コスト上限、SLA、監査ログ、代替候補、学習利用禁止条項の8項目が必要である。そのうえで、①高機密・顧

客向け＝閉域/国産環境優先、②社内コーディング・文書処理＝オープンモデル優先、③最難問の企画・研究＝高価な閉鎖系も許容、という三層ルーティングを導入する。使用料が膨らむ前に、ゲートウェイ、キャッシュ、可観測性を敷くべきである。 44

中期の提言

政府は、**国産LLM戦略を“汎用一発逆転”から“国産推論スタック＋バーティカルAI”へ再定義**すべきだ。経産省自身が整理する通り、日本の勝ち筋はフィジカルAIを中心としたバーティカルAIであり、製造現場データのAI-Ready化、国産マルチモーダル基盤モデル、制度改革、AX人材育成を同時に進める必要がある。したがって、補助金のKPIは“パラメータ数”ではなく、製造歩留まり改善、保全工数削減、設計期間短縮、医療文書処理時間短縮など、産業成果で置くべきである。 45

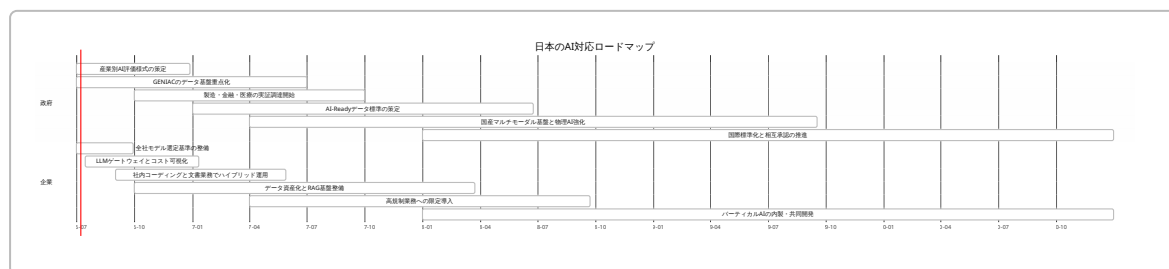
企業は、データ資産化に本気で投資する必要がある。日本企業の最大の弱点は、現場ノウハウをAIが食べられる形にしていないことだ。CAD、図面、手順書、故障履歴、品質逸脱、営業提案、規程、試験結果を、タグ付き・アクセス制御付き・再学習可能な形式へ変換し、RAGと学習の両方に使えるようにするべきである。ここで重要なのは、データを溜めることではなく、**再現可能な更新ループ**を作ることだ。実装→ログ→評価→改善→再実装のループを回した企業が、価格競争後の勝者になる。 46

長期の提言

政府は、AI政策を半導体・標準・ロボティクス・契約法務・個人情報保護と統合して運営すべきだ。AISI、GENIAC、AI・半導体産業基盤強化フレーム、フィジカルAI事業は既に点として存在する。これを線にし、**日本の標準を世界へ出す**構えが要る。特に、製造・ロボティクスの安全評価、データ連携、監査ログ、現場操作インターフェースを国際標準化のテーマとして押し出すべきである。日本が“モデルそのもの”で世界1位を取れなくても、“現場AIの標準”で主導権を取り得る。 47

企業は、最終的に単一モデル依存から卒業し、**モデルポートフォリオ経営**へ移るべきだ。将来の調達は、クラウドのように、用途別・地域別・機密別・コスト別に複数モデルを併用するのが基本形になる。社内には、モデル評価、セキュリティ、コスト管理、データ整備、契約、現場導入を横断するAI PMOを常設し、モデル更新を四半期ごとに見直す体制が必要である。これはIT施策ではなく、経営施策である。 48

政策と企業の実行ロードマップ



最後に、私は日本の戦略を次の一文で要約する。**日本は、米国型“閉鎖フロンティア依存”にも、中国型“低価格オープン依存”にも寄り切らず、オープンモデルを国内で使いこなす能力、産業データをAI-Ready化する能力、現場に落ちる評価標準を作る能力で勝つべきである。**その土台は既に、GENIAC、AISI、金融庁の対話、製造・ロボティクス政策、そして日本企業が持つ特許と現場知見の中にある。足りないのは、技術そのものより、**動員の速さと実装の一貫性**である。 49

2 6 11 33 36 <https://www.anthropic.com/news/fable-mythos-access>
<https://www.anthropic.com/news/fable-mythos-access>

3 5 16 23 24 26 <https://huggingface.co/deepseek-ai/DeepSeek-V4-Pro>
<https://huggingface.co/deepseek-ai/DeepSeek-V4-Pro>

4 8 28 45 46 https://www.meti.go.jp/shingikai/economy/global_industrial_strategy/pdf/003_01_00.pdf
https://www.meti.go.jp/shingikai/economy/global_industrial_strategy/pdf/003_01_00.pdf

7 <https://www.anthropic.com/glasswing>
<https://www.anthropic.com/glasswing>

9 12 41 44 48 <https://www.reuters.com/business/retail-consumer/cheaper-ai-is-better-soaring-bills-are-reshaping-how-businesses-choose-models-2026-06-29/>
<https://www.reuters.com/business/retail-consumer/cheaper-ai-is-better-soaring-bills-are-reshaping-how-businesses-choose-models-2026-06-29/>

13 <https://finance.yahoo.com/technology/ai/articles/coinbases-ceo-outlined-5-strategies-053434539.html>
<https://finance.yahoo.com/technology/ai/articles/coinbases-ceo-outlined-5-strategies-053434539.html>

14 19 20 25 <https://docs.anthropic.com/en/docs/about-claude/pricing>
<https://docs.anthropic.com/en/docs/about-claude/pricing>

15 <https://arxiv.org/abs/1503.02531>
<https://arxiv.org/abs/1503.02531>

17 <https://arxiv.org/abs/2405.04434>
<https://arxiv.org/abs/2405.04434>

18 <https://arxiv.org/abs/2306.00978>
<https://arxiv.org/abs/2306.00978>

21 <https://huggingface.co/zai-org/GLM-5.2>
<https://huggingface.co/zai-org/GLM-5.2>

22 29 <https://www.kimi.com/resources/kimi-k2-7-code>
<https://www.kimi.com/resources/kimi-k2-7-code>

27 42 <https://www.fsa.go.jp/news/r7/sonota/20260303/aidp.html>
<https://www.fsa.go.jp/news/r7/sonota/20260303/aidp.html>

30 <https://www.pmda.go.jp/rs-std-jp/subcommittees/0024.html>
<https://www.pmda.go.jp/rs-std-jp/subcommittees/0024.html>

31 <https://www.wipo.int/web-publications/global-innovation-index-2025/en/gii-2025-results.html>
<https://www.wipo.int/web-publications/global-innovation-index-2025/en/gii-2025-results.html>

32 37 https://www.ppc.go.jp/personalinfo/legal/guidelines_offshore/
https://www.ppc.go.jp/personalinfo/legal/guidelines_offshore/

34 39 <https://www.reuters.com/world/china/chinas-deepseek-developing-its-own-ai-chip-sources-say-2026-07-07/>
<https://www.reuters.com/world/china/chinas-deepseek-developing-its-own-ai-chip-sources-say-2026-07-07/>

35 <https://www.meti.go.jp/press/2023/02/20240214002/20240214002.html>
<https://www.meti.go.jp/press/2023/02/20240214002/20240214002.html>

38 <https://www.reuters.com/world/china/openai-accuses-deepseek-distilling-us-models-gain-advantage-bloomberg-news-2026-02-12/>

<https://www.reuters.com/world/china/openai-accuses-deepseek-distilling-us-models-gain-advantage-bloomberg-news-2026-02-12/>

40 https://www.fsa.go.jp/news/r6/sonota/20250304/aidp_summary.pdf

https://www.fsa.go.jp/news/r6/sonota/20250304/aidp_summary.pdf

43 <https://www.meti.go.jp/press/2026/07/20260702001/20260702001.html>

<https://www.meti.go.jp/press/2026/07/20260702001/20260702001.html>

47 https://www.meti.go.jp/policy/mono_info_service/ai_semiconductor_frame/ai_semiconductor_frame.html

https://www.meti.go.jp/policy/mono_info_service/ai_semiconductor_frame/ai_semiconductor_frame.html

49 https://www.meti.go.jp/policy/mono_info_service/geniac/index.html

https://www.meti.go.jp/policy/mono_info_service/geniac/index.html