

Qwen3.8-Max 徹底調査レポート

性能・価格・産業影響および知的財産／ガバナンス上の論点

調査基準日：2026 年 8 月 4 日

Claude Opus 5

1. 要旨

- 前提はほぼ正しい。「Qwen3.8-Max」は実在し、2026 年 8 月 2 日に Qwen 公式ブログで正式発表¹、8 月 3 日に主要メディアが一斉報道した¹³。ご指摘の「8 月 3 日発表」はメディア報道日であり、公式発表は 8 月 2 日、さらに 7 月 19 日には上海 WAIC で「Qwen3.8-Max-Preview」が先行公開されていた⁶。命名規則への懸念は当たらず、Qwen3.5→3.6→3.7→3.8 という月次刻みの連番系列に整合する¹¹。
- 性能は「フロンティア第 2 集団」。総パラメータ 2.4 兆（アクティブ約 950 億）の MoE 型マルチモーダルモデルで、コーディング／エージェント／マルチモーダルの自己申告ベンチマークは Claude Fable 5 や GPT-5.6 Sol に肉薄する^{1,12}。ただし推論・知識系（HLE 等）では明確に見劣りし、第三者評価は Arena のクラウドソース順位と Artificial Analysis の Index 53 点にとどまる^{4,17}。
- 知財・調達上の最大の論点は「Max クラス初のオープンウェイト化」とライセンス未確定。API は入力\$2／出力\$6（100 万トークンあたり）と米国フロンティア勢の 3～4 分の^{12,4}。日本（東京）リージョンでも利用可能⁸。一方でウェイトのライセンス条文が未公開であり、米国防総省 1260H リスト掲載および中国国家情報法リスクという調達上の重大論点を伴う^{21,22}。

2. 前提の真偽検証

2.1 確認された事実

- Qwen 公式サイト（qwen.ai）のブログに「Qwen3.8-Max: A New Bar for Coding and Cowork」と題する 2026 年 8 月 2 日付の記事が掲載され、同モデルを Qwen ファミリー史上最も高性能なモデルとして正式リリースする旨が記されており、名称・モデルともに実在する¹。

- 発表日は複数系列が存在する。①2026年7月19日＝上海WAICでのPreview公開⁶、②8月2日＝公式ブログでの正式リリース¹、③8月3日＝CNBC／Bloomberg／CGTN等の一斉報道¹³。ご指摘の「8月3日」は報道日として妥当である。

2.2 分析

「Qwen3.8」という表記は従来の命名規則と矛盾しない。Qwenは2026年に3.5（2月）→3.6（4月）→3.7（5月）→3.8（7～8月）と細かくバージョンを刻んでおり、「3.8」は自然な連番である^{7,11}。当初懸念された「命名規則との不整合」は該当しない。

※ 紛らわしい点：「Qwen3.8」（2.4兆パラメータのMaxクラス）は「Qwen3-8B」（旧世代の約80億パラメータのdense小型モデル）とは全く別物である。

3. モデル基本仕様

3.1 確認された事実

- **アーキテクチャ**：スパース MoE（Mixture-of-Experts）。総パラメータ 2.4 兆、推論時アクティブ約 950 億／トークン^{12,6}。
- **基盤**：「Qwen3.5 のアーキテクチャ基盤の上に構築」と公式に説明¹。Qwen3.5 系は Gated DeltaNet＋スパース MoE のハイブリッド、早期融合マルチモーダル学習、201 言語対応が特徴³。
- **コンテキスト長**：100 万トークン（約 75 万語）。プロバイダ仕様では最大入力 991K（thinking 時 983K）、最大出力 131,072 トークン、最大推論バジェット 262K²。
- **マルチモーダル**：テキスト・画像・動画・文書を入力可、出力はテキスト。Qwen 開発者 Shuai Bai 氏は「1 兆パラメータ超で初のマルチモーダルモデル」と説明している⁶。
- **推論モード**：thinking は常時有効。reasoning_effort に low／high／xhigh の 3 段階があり、xhigh がデフォルト。thinking 履歴は保持され課金対象となる²。
- **提供形態**：API は即日提供（QwenCloud／Model Studio、OpenAI・Anthropic・DashScope 互換）²。オープンウェイトは「翌週（8月10日週）」に Qwen3.8-Max 本体および小型の Qwen3.8-27B を Hugging Face／ModelScope で公開予定とされ、**Max クラスとしては史上初のオープンウェイト化**にあたる^{1,5}。

3.2 未確認事項

学習データ規模、正式なモデルカード（学習手法・データソース・安全性評価）は 8 月 3 日時点で未

公開である¹。学習に使用されたチップも確認できていない（第 6 章参照）。

4. ベンチマーク性能

重要な区別：以下のスコアはすべて Alibaba の自己申告値である¹。第三者による全表再現は本レポート作成時点で存在せず、独立系評価は Arena のクラウドソース順位、Artificial Analysis の Index 53 点、Trilogy AI／StackPerf の単発評価に限られる^{4,17,18}。

4.1 テキスト／コーディング／一般（Alibaba 自己申告値）

ベンチマーク	Qwen3.8-Max	Claude Fable 5	Claude Opus 4.8	GPT-5.6 Sol(max)	Qwen3.7-Max
Terminal-Bench 2.1	86.6	84.6	84.6	88.8	74.5
SWE-bench Pro	67.7	80.0	69.2	64.6	60.6
DeepSWE 1.1	56.6	70.0	59.0	73.0	21.6
FrontierSWE	73.5	88.8	70.0	—	40.7
PaperBench	93.0	88.8	80.3	90.5	64.8
GPQA Diamond	92.6	92.6	92.0	94.1	92.4
HLE（ツールなし）	43.6	53.3	45.7	47.2	41.4
HLE（ツールあり）	56.2	64.5	57.9	58.0	53.5
IFBench	82.8	63.5	62.2	72.7	79.1
MRCR v2 256K	92.9	—	83.2	93.8	86.7
WebArena-Verified	66.8	71.3	67.9	—	69.7

注：太字は当該行の最高値を示す。公式表に SWE-bench Verified および LiveCodeBench は含まれない。しばしば言及される「80.4」という SWE-bench Verified 値は前世代 Qwen3.7-Max のものであり、Qwen3.8-Max の値ではない⁷。Kimi K3 および Gemini 3.1 Pro はテキスト表に列が存在しない¹。

4.2 マルチモーダル（Alibaba 自己申告値）

ベンチマーク	Qwen3.8-Max	Fable 5	Opus 4.8	Gemini 3.1 Pro	GPT-5.6 Sol
OSWorld-Verified	86.1	85.0	83.4	76.2	83.2
OmniDocBench 1.5	92.1	89.5	86.5	90.0	86.7
Parametric CAD Bench	91.5	87.5	85.1	73.5	86.2

その他、MMMU-Pro 82.3、VideoMME（字幕あり）90.4、OCR-Bench-V2 EN／ZH 74.2／68.3 等で Qwen3.8-Max が優勢とされる¹。

4.3 第三者評価（独立系）

- **Artificial Analysis**：Intelligence Index v4.1 で 53 点（中央値 32 を大きく上回る）。ただし出力速度は 46 トークン／秒と遅め（中央値 72）、評価時の出力トークンは 1.5 億と極めて冗長。API 価格は入力\$2／出力\$6⁴。
- **Arena.AI**：テキスト部門 5 位（Claude Fable 5 および 3 つの Opus 系に次ぐ、中国モデル最上位）、ビジョン部門 2 位。ただしこれはブラインド選好投票であり学術ベンチマークではない¹⁷。
- **前世代の実測（要警戒）**：Qwen3.7-Max は Artificial Analysis で 56.6 点・5 位だったが、知識検索タスクで幻覚率 22.9%・棄権率 48%という弱点があった。Qwen3.8-Max の同種独立評価は未実施である^{4,7}。

4.4 ベンチマーク上の留意点（分析）

- Alibaba 自身が脚注で「Fable 5 の結果は fallback モデルを含む可能性」と明記しており、競合スコアの厳密性には留保が必要である¹。
- マルチモーダル表の比較対象が Qwen3.7-Plus（Max ではない）であり、世代間の伸びが過大に見える構成となっている¹。
- QwenSWEbench、QwenQoderBench 等の多数の行は Alibaba 独自の内製ベンチマークであり、標準化された第三者テストではない¹。
- Alibaba 自身が公表した RL スケーリング曲線は約 4,000 環境で 0.725 にピークを打った後、0.719→0.689 と低下しており、過学習の兆候が読み取れる¹。

5. 価格・提供形態

-
- **API 価格**：入力\$2.00／出力\$6.00（100 万トークンあたり）。100 万トークンのコンテキスト全域で長文サーチャージのないフラット料金^{2,4}。
 - **キャッシュ**：暗黙的キャッシュ読込\$0.25／M（入力の 8 分の 1）、明示的キャッシュ作成\$2.50・読込\$0.17／M²。

- **競合比**： Claude Opus 5 (\$5/\$25) 比で入力 60%・出力 76%安。GPT-5.6 Sol Max の約 4 分の 1、Claude Opus 5 の約 3 分の 1 (入出力合算)。Kimi K3 API (\$3/\$15) よりも安価。Artificial Analysis のブレンド単価は約\$1.18/M (7:2:1 キャッシュ比) ^{4,14}。
- **提供チャネル**： Alibaba Cloud Model Studio (国際版)、QwenCloud、Qwen Studio アプリ (chat.qwen.ai)、およびエージェント基盤 QwenWork (8 月 3 日パブリックベータ開始、QoderWork+MuleRun+Wukong を統合) ^{1,2}。
- **レート制限**： 200 万トークン/分、15,000 リクエスト/分。北京・シンガポール・米バージニア の 3 リージョンでエンドポイントを運用 ²。
- **日本からの利用：可能**。Model Studio は東京 (Japan/Tokyo) リージョンを提供し、ap-northeast-1 の OpenAI 互換エンドポイントが存在する。ワークスペースで Global/Japan のデプロイスコップを選択でき、データレジデンシー要件に対応可能 ⁸。

6. 業界の反応と評価

6.1 称賛点

- 中国ラボによるフロンティア級オープンウェイト競争として歓迎する声 (Hacker News 主流意見) ¹⁵。
- マルチモーダル能力および長時間エージェント能力の飛躍。Alibaba は「16 日間の自律コーディング」「研究論文をゼロから再現 (33 回の GPU 学習、約 125 時間、7,600 行のコード、18 の改善案)」等のデモを提示した ¹。
- 価格破壊力。VentureBeat は「米国フロンティア勢に数十%のマージンで肉薄しつつ価格は 3~4 分の 1」と評価している ¹⁴。

6.2 批判点・懐疑論

- **ベンチマーク自己申告問題**：「Qwen はベンチマーク特化の評判があり、"trust us, it's #2"はベンチマーク表で決着すべき」との批判 ¹⁵。
- **再現性・独立評価の欠如**： Artificial Analysis および Hugging Face Open LLM Leaderboard に 8 月 3 日時点で未掲載だった (AA は後に Index 53 点を公表) ⁴。
- **実運用との乖離懸念**： 前世代の高い幻覚率、および Kimi K3 (Moonshot AI が 2026 年 7 月 16 日に公開した 2.8 兆パラメータのオープンウェイトモデル) が複数の第三者ベンチで上位でも知識タスクで幻覚率 51%だった前例から、Arena 順位が実務性能を保証しないとの警告 ^{6,17}。

- **2.4 兆パラメータの実行不可能性**：「2.4 兆パラメータは 4bit 量子化でも約 1.2TB の VRAM を要し、8×H200（約 1.13TB）でも不足」との指摘があり、オープンウェイトであっても自己ホストは事実上困難。現実的な自己ホスト経路は小型の Qwen3.8-27B のみとなる¹⁶。
- **「kaleb」事件**：7 月 18 日に Code Arena に匿名モデル「kaleb」が出現し「Claude」を自称、Qwen トークナイザー特有のパターン（PostalCodesNL）から正体が特定された経緯があり、ステルステストへの疑念を呼んだ¹⁵。

7. 影響分析（分析・推論）

7.1 米中性能差の現状

UC Berkeley のイオン・ストイカ氏（Databricks／Arena 共同創業者）は 2026 年 7 月 30 日、中国のオープンウェイトと米国フロンティアの差が「6～9 か月から約 2～3 か月に縮小した」と評価した⁹。ただし他の定量評価はより慎重で、Epoch AI は「2026 年 1 月以降、最も高性能なオープンウェイトモデルはフロンティア・クローズドモデルに平均 4 か月（8 ECI ポイント）遅れている」とし¹⁹、米 CAISI は 2026 年 5 月評価で DeepSeek V4 Pro を「先行する米国勢に約 8 か月遅れ」と算定している²⁰。Qwen3.8-Max はこの縮小トレンドを裏付ける一方、HLE 等の推論・知識系では依然として 10 ポイント近い差が残る¹。

7.2 オープンウェイト・エコシステム

Max クラス初のオープン化は象徴的意義を持つ。ただし①ライセンス未確定、②2.4 兆パラメータは事実上データセンター級で誰も動かせない、という 2 点から「名ばかりオープン」との批判もある¹⁶。実効的インパクトは Qwen3.8-27B の性能次第と評価される。

7.3 API 価格競争

OpenAI は直近で GPT-5.6 の Terra／Luna を値下げし（Luna は入力 100 万トークンあたり \$0.20 へ 80% 減、Terra は 20% 減）⁹、DeepSeek-V4-Flash も低価格帯に投入されている。Qwen3.8-Max の \$2／\$6 は価格圧力を一段と強める。

7.4 半導体・計算資源

学習チップは未確認である。Alibaba 子会社 T-Head の自社 PPU「Zhenwu M890」（144GB HBM3）は 5 月に Qwen3.7-Max と同時発表されたが、公式デモでは Qwen3.7-Max を「学習時に遭遇していない M890」で動作させており、M890 は推論／ターゲット基盤であって学習ハードとは限らない¹⁰。Qwen3.8-Max の学習ハードウェアを明言した情報源は確認できなかった。なお Bloomberg は、

Alibaba が Moonshot に Kimi K3 学習用として約 2 万個の Nvidia チップを供給したと報じているが、これは Qwen3.8-Max とは別件である^{9,13}。

7.5 日本企業への含意

東京リージョンで利用可能でありデータの国内処理も選択できるが⁸、後述のガバナンス懸念から、日本企業においては「汎用業務は米国系、機密・規制業種・日本語特化は国産、コスト重視のバッチ処理系に Qwen 等の中国モデル」というデュアルスタック運用が現実的と考えられる。

7.6 規制・ガバナンス

重大論点である。 Alibaba はケイマン籍・杭州拠点であり、中国の①国家情報法（2017 年）、②サイバーセキュリティ法（2026 年 1 月改正で AI 統治条項を追加）、③データセキュリティ法（2021 年）の適用対象となる。特に国家情報法第 7 条は、すべての組織および国民に対し、法に従って国家の情報活動へ支援・援助・協力する義務と、その過程で知り得た秘密を守る義務を課しており、第 14 条は情報機関に協力要求権限を付与している²²。これらはサーバ所在地やプライバシーポリシーの内容に関わらず適用される。とりわけ QwenWork（Wukong 知識ベース）は全社文書を取り込む構造であるため、露出面が拡大する¹。

加えて米国防総省は 2026 年 6 月 8 日に 1260H リストを更新し、Alibaba・Baidu・BYD 等を含む 65 エンティティ（親 17 社＋子会社 48 社）を追加、計約 188 社とした。直接契約禁止は Section 805（FY2024 NDAA）により 2026 年 6 月 30 日発効、間接調達禁止は 2027 年 6 月 30 日発効であり、Alibaba は「恣意的」として提訴中である²¹。

8. 知的財産・ライセンス上の論点

8.1 公開 API モデル

クローズド提供であり、Model Studio の利用規約およびデータ処理規約に従う。Alibaba は「顧客データを学習に使用しない」と表明している^{2,8}。

8.2 オープンウェイトのライセンス（最重要・未確定）

2026 年 8 月 3 日時点で **Qwen3.8-Max のライセンス条文は未公開**であり、これが知財上の最大の不確実性である¹。

- **先例としての傾向：** Qwen3・3.5・3.6 系の多くのオープンウェイトは Apache 2.0（商用利用・再配布・改変・派生モデルに寛容）で提供されてきた。OpenLM.ai も「All our open-weight models are licensed under Apache 2.0」と記載している³。

- **ただし留保：**①直近2世代のMaxフラッグシップ（3.7-Max、3.6-Max-Preview）はAPI専用でウェイト非公開だった⁷、②Qwen独自の「Tongyi Qianwen License」は72B以上のモデルに適用されApache 2.0より商用再配布で制限的である⁸、③競合Moonshotの「オープン」なKimi K3も独自制限ライセンスだった前例がある⁹。以上から、2.4兆パラメータモデルが制限的なカスタムライセンスで公開される可能性は排除できない。

8.3 知財実務上の判断ゲート

1. Hugging Face に実際のライセンスファイルが掲載されるか。
2. Apache 2.0 か独自ライセンスか。
3. 商用利用・フィールド制限・再配布・改変・派生モデルの各条件はどうか。

以上を必ず条文で確認したうえで製品組込みを判断すべきである。なお「オープンウェイト」と「オープンソース（OSI 準拠）」は法的に別カテゴリである点に留意を要する。

8.4 学習データの著作権論点

Alibaba は正式なモデルカード・学習データソースを未公開である¹⁰。学習データの出所・ライセンス・著作物混入の有無が不明であり、生成物の著作権帰属や第三者権利侵害リスクの評価が困難である。中国系モデル全般に共通する論点ではあるが、透明性の欠如は日本の知財デューデリジェンス上のリスク要因となる。また蒸留（distillation）による他社モデルからの知識移転疑惑も Qwen 系で過去に取り沙汰されており（前述「kaleb」事件における Claude 自称等）¹⁵、派生・蒸留を巡る知財論点は継続的な監視を要する。

9. 提言（段階的アプローチ）

4. **即時（評価フェーズ）：**東京リージョンの Model Studio で自社の実プロンプト 20～30 件を用いた A/B 評価を実施する。競合（GPT-5.6/Claude/Kimi K3）と同一条件で品質・レイテンシ・トークン・コストを実測し、ベンチマークの自己申告値ではなく自社ワークロードの数値で判断する。**本番トラフィックは現状維持とする。**
5. **ライセンス確定待ち（知財ゲート）：**オープンウェイト公開後、Hugging Face のライセンスファイル条文を精査する。**Apache 2.0 であることの確認を製品組込みの前提条件**とし、独自ライセンスであれば商用利用・再配布・派生モデル条項を法務レビューに付す。
6. **ガバナンス評価：**機密データ・個人情報・営業秘密を扱う用途では、中国 3 法（国家情報法・サイバーセキュリティ法・データセキュリティ法）および 1260H リスト掲載を踏まえ、クラウド API および QwenWork の利用を避け、自己ホスト（Qwen3.8-27B）または Enterprise

Deployment Kit によるオンプレミス運用に限定する。米国防総省関連取引がある組織は特に慎重を期す。

7. 用途別採用：コスト重視の大量バッチ処理／エージェント、マルチモーダル文書処理、非機密のコーディング支援には有力候補となる。推論・知識精度が重要な用途および機密性の高い用途は米国系・国産を優先する。

9.1 判断を変える閾値

- Artificial Analysis 等の独立系が SWE-bench Pro／Terminal-Bench 等で Alibaba 自己申告値の±3ポイント以内を再現した場合 → コーディング用途での採用検討を格上げ。
- オープンウェイトが Apache 2.0 で公開された場合 → 自己ホスト前提の製品組込みを本格検討。
- 知識タスクの独立幻覚率が前世代（22.9%）から大幅に改善した場合 → 知識検索用途への適用を再評価。

10. 留意事項（不確実性の明示）

- 本レポートのベンチマーク数値は、特記なき限り Alibaba の自己申告値である。第三者による全表再現は未実施であり、独立系評価は Arena 順位および Artificial Analysis Index 53 点に限られる^{1,4,17}。
- 正式なモデルカード（学習手法・データ・安全性評価）、オープンウェイトのライセンス条文、学習ハードウェアは 2026 年 8 月 3 日時点で未確認である。
- アクティブパラメータ 950 億は GA 時に Alibaba が公表した値（プレビュー時は非開示）であり、OpenLM.ai のスペック表には行項目として存在しない（報道経由の Alibaba ブログ引用による）^{3,12}。
- オープンウェイト公開日およびライセンスは予定であり確定情報ではない。過去には「数日で公開」とされながら数週間遅延、あるいは未公開に終わった例もあり、ロードマップは変動し得る。
- 一部の URL（Bloomberg 等）はペイウォールまたは bot 検出によりアクセスできなかった。参考文献一覧のうち URL 欄が「未取得」のものは、本調査で一次 URL を確認できていないことを示す。

参考文献

本文中の肩付き番号は下記の文献番号に対応する。URL は 2026 年 8 月 4 日時点のもの。

1. Qwen (Alibaba Cloud) 公式ブログ「Qwen3.8-Max: A New Bar for Coding and Cowork」2026 年 8 月 2 日
<https://qwen.ai/>
2. QwenCloud「Qwen3.8-Max」モデル仕様ページ
<https://www.qwencloud.com/models/qwen3.8-max>
3. OpenLM.ai「Qwen3.8」スペック一覧
<https://openlm.ai/qwen3.8/>
4. Artificial Analysis「Qwen3.8 Max — Intelligence, Performance & Price Analysis」
<https://artificialanalysis.ai/models/qwen3-8-max>
5. The Daily Star「Alibaba releases Qwen 3.8-Max, its largest AI model yet」
<https://www.thedailystar.net/news/tech-startup/news/alibaba-releases-qwen-38-max-its-largest-ai-model-yet-4238986>
6. MarkTechPost「Alibaba Previews Qwen3.8-Max, a 2.4 Trillion-Parameter Multimodal Model, Days After Moonshot's Kimi K3 Open-Weight Launch」2026 年 7 月 19 日
<https://www.marktechpost.com/2026/07/19/alibaba-previews-qwen3-8-max-a-2-4-trillion-parameter-multimodal-model-days-after-moonshots-kimi-k3-open-weight-launch/>
7. MarkTechPost「Qwen Introduces Qwen3.7-Max: A Reasoning Agent Model With a 1M-Token Context Window」2026 年 5 月 21 日
<https://www.marktechpost.com/2026/05/21/qwen-introduces-qwen3-7-max-a-reasoning-agent-model-with-a-1m-token-context-window/>
8. Coursiv Blog「Qwen 3.8: Specs, Pricing, Access & Benchmarks」
<https://coursiv.io/blog/qwen-3-8>
9. Tech Times「Alibaba Bankrolled Kimi K3 With 20,000 Nvidia Chips: Now Qwen Is Losing to Its Own Compute」2026 年 8 月 2 日
<https://www.techtimes.com/articles/322699/20260802/alibaba-bankrolled-kimi-k3-20000-nvidia-chips-now-qwen-losing-its-own-compute.htm>
10. Wccftch「Alibaba Targets NVIDIA's Hopper With Zhenwu M890 AI Chip, Claiming 3x The H20 Performance, 144GB HBM3 & A Roadmap Through 2028」
<https://wccftch.com/alibaba-targets-nvidia-hopper-with-zhenwu-m890-ai-chip-claiming-3x-h20-performance/>
11. GitHub「QwenLM/Qwen3.6」リポジトリ (Qwen バージョン系列の確認)
<https://github.com/QwenLM/Qwen3.6>
12. InfoWorld (Alibaba 公式ブログを引用した総パラメータ数・アクティブパラメータ数の報道)
[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]
13. CNBC／Bloomberg／CGTN (2026 年 8 月 3 日の Qwen3.8-Max 一斉報道)
[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

14. VentureBeat (Qwen3.8-Max の価格対性能比に関する評価記事)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

15. Hacker News (開発者コミュニティにおける反応・ベンチマーク自己申告への批判)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

16. Reddit r/LocalLLaMA (2.4 兆パラメータモデルの自己ホスト可否に関する議論)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

17. Arena.AI リーダーボード (テキスト部門・ビジョン部門の順位)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

18. Trilogy AI/StackPerf (Qwen3.8-Max-Preview の独立評価)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

19. Epoch AI, Jack Edwards & Luke Emberson (2026) オープンウェイトとフロンティアモデルの性能差分析

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

20. CSIS (米 CAISI による DeepSeek V4 Pro 評価の解説、2026 年 5 月)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

21. China Law Translate「中華人民共和国国家情報法」英訳 (2017 年 6 月 27 日可決、2018 年 4 月 27 日改正)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

22. WilmerHale/Cleary Gottlieb (米国防総省 1260H リスト 2026 年 6 月 8 日更新に関する法律事務所解説)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]

23. South China Morning Post (Qwen3.8-Max 関連報道)

[URL 未取得・未検証：本文中の記述は二次的に確認した内容に基づく]