

Z.ai「GLM-5.3」の全貌と産業的影響：事後学習のスケールアップが生み出したオープンウェイトモデルの新たなパラダイム

Gemini 3.7 Flash

中国のAIスタートアップであるZ.ai（智譜AI / Zhipu AI）は、2026年8月14日、次世代大規模言語モデル「GLM-5.3」を正式に発表した¹。本モデルは、前世代モデル「GLM-5.2」と同じ7,430億（743B）パラメータ規模の基盤モデル（ベースモデル）を継承しつつ、アーキテクチャ変更を伴わない純粋な事後学習（ポストトレーニング）の拡張のみによって、複雑なコーディング能力および長期エージェント作業能力の大幅な強化を達成した¹。

さらに、事後学習の強化に伴い、モデル内部に高度なサイバーセキュリティ能力が自律発生的に伸長したことが確認された³。特に脆弱性の特定および検証においては、Anthropicが厳格なアクセス制限下で提供するセキュリティ特化モデル「Mythos 5」を超える精度を記録し、オープンウェイトモデルの新たな基準を提示している²。本報告書は、GLM-5.3の技術的構造、ベンチマーク評価、エコシステム展開、ならびにAIガバナンスとサイバーセキュリティ市場へ及ぼす構造的影響について包括的な分析を提供する。

技術的基盤と事後学習アーキテクチャ

GLM-5.3の開発プロセスにおける最大の特徴は、モデルの基本アーキテクチャやパラメータ数を維持したまま、強化学習（RL）とデータ環境のスケールアップのみで最高水準の性能を導き出した点にある¹。この成果は、モデル容量の拡大に頼らずとも、事後学習の質と最適化システムを向上させることでフロンティア級モデルに匹敵する性能を獲得できることを実証している²。

学習インフラストラクチャとRLシステム

GLM-5.3の開発では、前世代で構築された長文脈処理基盤「IndexShare」、長期的タスク強化学習フレームワーク「SAO」、および大規模非同期訓練基盤「slime」が全面的に活用された⁵。訓練システムは、分散学習エンジン「Megatron」、高速推論ロールアウトエンジン「SGLang」、およびデータバッファが統合された一元的なデータフローとして構成されており、学習ループ自体を変更することなく、数学、プログラミング、サンドボックス環境、長期的タスク環境を動的に組み込むことが可能となっている⁸。

アルゴリズム面では、オンライン・ポリシー蒸留（OPD）のバリエーション手法やtop-pマスキングが適用された⁸。訓練パイプラインと推論ロールアウト間の対数確率（logprob）誤差を 1×10^{-7} レベルに抑え込み、数値的な完全一致を達成することで、実験の制御精度とRLの収束安定性を劇的に高めることに成功している⁸。さらに、ローカルストレージをキャッシュ層として利用するマルチティーチャーOPDの導入により、長時間タスクにおけるロールアウト計算コストの大幅な低減を実現した⁸。

実務特化型タスク環境の自動生成

従来のプログラミング演習レベルのベンチマーク学習から脱却するため、Z.aiは人間のエンジニアや

研究者が実践する実際のワークフローを模した長時間タスク環境の自動生成システムを構築した²。このシステムではまず、リサーチエージェントが実際の開発パターンを分析し、マルチステップ処理、隠れた状態、長期的な依存関係を含む複雑な作業環境を自動で生成する⁸。続いて、ジャッジエージェントが生成されたタスク環境の解決可能性や、不正なショートカット(報酬ハック)の有無を即座に検証し、これらの厳格な基準をクリアした環境のみを強化学習の報酬モデルとして適用する仕組みが確立されている⁸。

技術スペックとシステム要件

GLM-5.3は、長大なコードベースや複雑な指示に対応するため、1,000,000トークン(1M)におよぶ極めて広大なコンテキストウィンドウと、最大131,072トークン(131K)の出力容量を備えている⁹。推論モードにおいては、low、high、maxの3段階の思考強度(reasoning_effort)を柔軟に選択できる仕様となっている⁵。なお、本モデルでは思考機能の完全無効化(thinking.type: "disabled")がサポートされていないため、従来のAPI環境から移行する際には、思考機能を有効化した上で適切な思考強度を指定することが必須の要件となる⁵。

性能評価とベンチマーク分析

GLM-5.3の評価は、ソフトウェア開発能力、長期的エージェント作業、サイバーセキュリティの3分野において実施され、オープンウェイトモデルとしてトップクラスの数値を記録した²。

ソフトウェア開発およびエージェント機能

社内評価指標である「Z.ai Code Bench」において、GLM-5.3は前世代のGLM-5.2から50%の性能向上を達成した²。公開ベンチマークにおいても、コマンドライン操作や開発エージェントとしての自律性能を示す各種指標で前モデルを圧倒している²。

ベンチマーク指標	GLM-5.2	GLM-5.3	比較対象モデル(参考値)
Terminal-Bench 3.0	4.6	28.3	Claude Fable 5: 33.7 ⁴
DeepSWE v1.1	46.2	66.9	Claude Fable 5: 69.7 ⁴
Agents' Last Exam	23.8	28.5	--
Z.ai Code Bench (Max Effort)	23.4% (96k token)	34.5% (75k token)	Claude Fable 5: 39.5% ⁵
Z.ai Code Bench (High Effort)	--	31.4% (50k token)	Claude Opus 4.8 (Max): 29.5% (120k)

			token) ²
GDPval-AA v2 (44 職種対応)	--	1769	--

ベンチマーク結果の分析により、GLM-5.3は単に精度が向上しただけでなく、トークン消費効率が大幅に優れていることが明らかとなった²。Z.ai Code BenchのMax設定において、GLM-5.2がタスクあたり約96,000トークンを消費して23.4%の精度にとどまっていたのに対し、GLM-5.3は約75,000トークンで34.5%の精度を達成した²。また、High設定時のGLM-5.3(精度31.4%、消費トークン数約50,000)は、Max設定時のClaude Opus 4.8(精度29.5%、消費トークン数約120,000)と比較して、精度とトークン効率の双方で上回る結果を示している²。

なお、Terminal-Bench 3.0における「28.3」というスコアは、Claude Code 2.1.207を実行基盤とし、最大600ターン・10時間に及ぶ試行(3回の平均)によるものであり、単発の対話性能ではなく、長時間の自律的な問題解決能力を反映した数値である¹⁰。

サイバーセキュリティ能力と攻撃チェーン推論

事後学習のスケールアップによって獲得された最も著しい変化は、サイバーセキュリティ分野における能力の急拡張である³。トレーニング環境に脆弱性発見データおよびアクティブな検証サンドボックスを組み込んだ結果、受動的な欠陥特定のみならず、複数の脆弱性を組み合わせた攻撃計画を自律的に立案する能力が向上した²。

セキュリティベンチマーク	GLM-5.2	GLM-5.3	Anthropic Mythos 5	OpenAI GPT-5.6 Sol
CyberGym (脆弱性特定・検証)	77.2%	84.5%	83.8%	83.6%
ExploitBench (エクスプロイト推論)	24.4%	54.4%	78.0%	76.5%
ExploitGym (2時間完了タスク数)	29	105	181	--
ExploitGym (6時間完了タスク数)	39	130	247	--

ソースコードの監査と欠陥検証を行う受動的セキュリティ指標「CyberGym」において、GLM-5.3は84.5%のスコアを記録し、制限付きモデルであるMythos 5(83.8%)やGPT-5.6 Sol(83.6%)をわずかに上回るトップスコアを達成した²。

一方で、特定した脆弱性を実際に動作する攻撃用エクスプロイトコードへと変換する能動的攻撃能力(ExploitBenchおよびExploitGym)においては、前世代GLM-5.2から2倍以上の伸びを示したものの、AnthropicのMythos 5に対しては明確な性能差が残されている²。これは、防御的なコードレビュー能力の強化が先行して達成され、高度な攻撃コード生成能力の獲得にはさらなる強化学習ステップが必要であることを示唆している⁴。

リアルワールドコードベースでの検証実績

Z.aiは複数の大学および提携セキュリティチームと協同し、実際の各種デジタルインフラ(OSカーネル、ブラウザ、ネットワーク機器、スマートデバイス等)のコードベースに対してGLM-5.3を適用した²。その結果、269個のオープンソースプロジェクトにわたり合計2,436件のセキュリティ脆弱性を特定し、そのうち1,097件が「Critical(緊急)」または「High(高危険度)」の深刻度と判定された²。この検証でモデルが検出した最も古い脆弱性は1981年に作成されたコード内に存在し、平均して26.6年間見過ごされてきたものであった⁴。

エコシステム統合と商業展開モデル

GLM-5.3のリリースにあたり、Z.aiは単なるモデルAPIの提供にとどまらず、開発者向け環境およびコスト最適化メカニズムを緊密に統合したサービス構造を展開している⁸。

統合開発環境「ZCode」との深層結合

Z.aiはデスクトップ型エージェント開発環境(ADE)である「ZCode」をGLM-5.3の主力フロントエンドとして位置づけている⁸。ZCodeは100万トークンのコンテキスト保持能力を活用し、ファイル変更履歴、ターミナル出力、Git管理状態、ブラウザ文脈をシームレスに統合する⁸。これにより、要件定義から計画立案、コード変更、検証、レビューに至るまでの一連の工程を単一のタスク内で完結させ、作業中の文脈断絶を極小化することが可能となった¹¹。さらに、デスクトップ環境、モバイルリモート環境、FeishuやWeChatのボット環境間で同一のワークスペースタスクを同期・継続できるマルチデバイス連携を強化している⁸。また、ZCodeはハネス層としての最適化が施されており、94%から98%以上という高いプロンプトキャッシュヒット率を記録することで、反復的なコンテキスト読み込みにかかる計算コストとレイテンシを劇的に抑制している⁸。

サブスクリプション体系と利用経済性

有料プラン「GLM Coding Plan」においては、従来のトークン定額制からポイントベースのクォータ管理システムへの移行が行われた⁵。運用コストの最適化を図るため、日本時間および中国標準時(UTC+8)の平日14:00~18:00をピークタイムと定め、それ以外の時間帯や週末におけるモデル呼び出し時のポイント消費量を50%に割引くオフピーク優遇措置が導入されている⁵。加えて、緊急性を要さないバックグラウンド作業をキューに投入しておくことで、システムの空きリソース発生時に契約者の利用枠を消費することなく無料で処理を実行する「待機時間タスク(Idle Tasks)」の仕組みも段階的に提供されている⁸。開発者の柔軟な運用を支えるため、AnthropicのカスタムベースURL設定

等を利用することで、Claude Code や OpenCode などの既存のエージェントCLIツールから GLM-5.3を直接呼び出して運用する互換性も確保されている⁸。

セキュリティ・ガバナンスと社会的・産業的影響

GLM-5.3の登場は、高度サイバーセキュリティAIツールの普及形態、オープンソースエコシステムの防御体制、ならびにAI開発におけるリスク管理のあり方に大きな一石を投じている⁴。

オープンモデル vs 閉鎖型制限モデルのパラダイム対立

高度なサイバーセキュリティ能力を持つAIの普及形態をめぐっては、二つの対比的なアプローチが浮き彫りとなっている⁴。Anthropicに代表される閉鎖的制限管理のアプローチでは、「Project Glasswing」方針のもと、攻撃的アラインメントを除去したセキュリティモデル (Mythos 5) の利用を厳格に審査された一部の公的機関や特定企業のみ限定している⁴。これに対しZ.aiは、世界のデジタルインフラの多くがリソースの乏しい小規模なオープンソースチームによって支えられている実態を強調し、防御技術が一部の巨大組織に独占される構造を防ぐため、フロンティア級セキュリティモデルのオープンウェイト化を唱導している⁴。この思想に基づき、Z.aiはオープンソースプロジェクトの無償監査、防御モデルの提供、および発見された脆弱性の管理を行う「Open Source Shield (または OpenVuln)」イニシアチブを発足させた⁴。モデルによって特定された2,436件の脆弱性は専用の公開台帳に記録され、プロジェクト保持者との守秘性の高い責任ある開示プロセス (Responsible Disclosure) を経て段階的な修正が進められている⁴。

3層セーフガード構造とローカル実行に伴う課題

Z.aiはモデルの悪用を防止するため、危険なプロンプトを弾く外部分類器 (External Classifier)、実行中の異常信号を検知する推論モニター (Inference Monitor)、そしてモデルの重みレベルで正当な監査と攻撃行為を自律識別させる深層安全アラインメント (Deep Safety Alignment) からなる3層のセーフガード構造を構築した⁶。しかし、AIガバナンスの専門家からは、オープンウェイトモデル特有の構造的限界も指摘されている⁴。ユーザーがモデルウェイトをダウンロードしてローカル環境で運用する場合、Z.aiのクラウドインフラに依拠する外部分類器や推論モニターは機能しない⁷。モデル本体に深層安全アラインメントが組み込まれているとはいえ、重み自体が開放されればファインチューニングやコード改変によってセーフガードを迂回することが技術的に可能となるため、高度な攻撃技術の意図せぬ拡散 (デュアルユース・リスク) に対する懸念は依然として根強く残されている⁴。

中国AIラボにおけるガバナンスの成熟と市場影響

GLM-5.3のリリース戦略において特筆すべきは、モデル発表 (8月14日) からモデルウェイトの一般公開までに「2週間」の猶予期間を設けた点である²。Z.aiはこの期間を安全評価およびガードレールの強化に充てるとしており、最も機微な能動的攻撃機能については「信頼済みアクセス (Trusted Access)」プログラムを通じた承認済みユーザーへの提供に限定する方針を示した⁴。

AIガバナンスの専門家は、中国のAIラボが「安全性とリスク管理」を明示的な理由としてオープンウェイトの公開を段階化・遅延させた初の実例であると評価しており、中国国内におけるAIリスク管理規範が成熟段階に入りつつあることを示している⁶。

結論

Z.aiによるGLM-5.3の発表は、AIモデルの開発効率、実用性、およびセキュリティガバナンスの各観

点において重要な節目となる成果である¹。

第一に、事前学習モデルの基本構造を変更せず、高品質な合成環境と非同期強化学習システムによる「事後学習の拡張」のみでプログラミング性能を50%向上させ、フロンティアモデルに迫る性能を引き出せることを実証した¹。この知見は、今後のモデル開発における投資対効果の最大化手法として、業界全体へ強く影響を与えられとされる¹。

第二に、ソフトウェア開発におけるトークン効率と自律性の向上が挙げられる²。高いコンテキスト保持力とプロンプトキャッシュ技術の組み合わせにより、長時間のソフトウェア構築タスクにおけるトークン単価を劇的に下げること成功しており、開発現場におけるAIエージェントの日常的導入を一層加速させる動因となる²。

第三に、サイバーセキュリティ分野における能力の民主化とリスク管理の二面性である⁴。受動的な脆弱性検出において制限付きフロンティアモデル(Mythos 5)を上回る精度を提示したことは、オープンソース社会の防御力を高める大きな恩恵となる一方で、ローカル実行可能なオープンウェイトモデルの安全策保持という困難な課題を提示している⁴。Z.aiが導入した段階的公開と信頼済みアクセスプログラムが、今後のオープンソースAIコミュニティにおける責任ある公表(Responsible Release)の標準モデルとなり得るか、数週間後のモデルウェイト公開とその後のエコシステムの運用が注視される⁴。

引用文献

1. AIモデル「GLM-5.3」を発表 ―コーディングと脆弱性解析を強化、重みは2週間後に公開予定, <https://gihyo.jp/article/2026/08/glm-5.3?summary>
2. Fable 5とも戦える「GLM-5.3」登場。事後学習の拡張のみで大幅性能アップ - PC Watch, <https://pc.watch.impress.co.jp/docs/news/2132889.html>
3. Z.ai、GLM-5.3をリリース - フロンティアコーディングとサイバー能力を強化 - Unite.AI, <https://www.unite.ai/ja/z-ai-launches-glm-5-3-with-frontier-coding-and-a-cyber-capability-that-outgrew-its-training/>
4. 中国Zhipu、オープンソース「GLM-5.3」がAnthropicの制限付きモデルを脆弱性検出で上回ると発表 - BigGo ファイナンス, <https://finance.biggo.jp/news/b569502f-66a8-4fbc-94f1-8ffd61046cee>
5. GLM-5.3: Frontier Coding with Emergent Cyber Capabilities - Z.ai, <https://z.ai/blog/glm-5.3>
6. China's Z.ai claims its new GLM-5.3 model is close to Anthropic's Mythos 5 level in cybersecurity tests, <https://timesofindia.indiatimes.com/technology/tech-news/chinas-z-ai-claims-its-new-glm-5-3-model-is-close-to-anthropics-mythos-5-level-in-cybersecurity-tests/articleshow/133245860.cms>
7. Z.ai、サイバーセキュリティ特化AI「GLM-5.3」発表 Mithos 5に対抗, <https://www.digitaltoday.co.kr/jp/view/93638/chinese-open-weight-models-target-security-z-ai-challenges-mythos-with-glm-5-3>
8. メモ: GLM-5.3 & ZCode - Zenn, <https://zenn.dev/kun432/scraps/fcf8117a2aac6a>
9. GLM-5.3 pricing, providers, and specs - Models.dev, <https://models.dev/models/zhipuai/glm-5.3>
10. Z.aiがGLM-5.3を発表、ポストトレーニングだけでコード性能と攻撃能力を伸ばす | XenoSpectrum, <https://xenospectrum.com/glm-5-3-post-training-cyber/>

11. 「ZCode」を試す ① - Zenn, <https://zenn.dev/kun432/scraps/5bb9088b328f16>
12. GLM-5.3 - Overview - Z.AI DEVELOPER DOCUMENT,
<https://docs.z.ai/guides/llm/glm-5.3>
13. GLM-5.3 Usage, Cost & Rank | OpenCode Data,
<https://opencode.ai/data/hipu/glm-5-3>