

ミュトス時代の自律型AIガバナンス: 技術の爆速進化と「法の支配」のコペルニクスの転回

Gemini 3.6 Flash

2026年4月に米国Anthropic社が発表した「Claude Mythos(クロード・ミュトス)」を契機とするフロンティアAIの爆発的な台頭は、デジタル空間における安全保障と法秩序のあり方を一変させた¹。AIの処理能力および推論能力が「爆速進化」を遂げる一方で、それを規律するための社会的な合意形成や法的なルール策定には「3年」という長い時間がかかるという、深刻な時間的非対称性が浮き彫りになっている³。TBS CROSS DIG with Bloombergが配信した、スマートガバナンス株式会社代表取締役CEOであり弁護士の羽深宏樹氏と竹下隆一郎氏による対談は、この「ミュトス時代」のガバナンス設計が、単なる技術的な管理策を超えた国家の生存戦略であることを示唆している³。技術が既存の規範を超えていく中で、自律型AIがもたらす根源的脅威、主要極における規制戦略の分岐、そして日本がとるべきポジショニングについて詳細な分析を行う。

制度と技術の不均衡: 爆速進化するAIと「3年」のタイムラグ

現代のAI技術が「便利な道具」から「仕事を自律的に任せる相手」へと変貌する中、法律や公的規制といった「ハードロー」を基本とするガバナンスは、重大な機能不全に直面している⁵。その本質的な原因は、技術の進化とルール形成の間に存在する致命的な時間差(タイムラグ)にある³。

技術的本質として、東大教授の集団をも凌駕し得ると評価される高度な解読・分析能力は、単一のモデルに限られた特異現象ではない⁸。むしろ、AI全体の自律性と能力の底上げを象徴するパラダイムシフトであり、人間の指示を受けずとも、長期にわたり自律的に分析と反省を反復できるエージェントへの変貌を示している⁸。

一方で、羽深氏は人間の脳システムの卓越した効率性にも光を当てる⁸。人間の脳は、わずか1個のLED電球を灯すほどの極小の電力で極めて高度な知的活動を処理しており、最先端AIの計算インフラと比較して数百万倍もエネルギー効率が優れている⁸。この生命科学的な極小のコスト構造と、複雑かつ個別具体的なコンテキストを包摂して行う判断(臨床的判断)能力こそが、AIに完全に代替されない人間の固有領域を形作っている⁵。

しかし、自律的に試行錯誤を行い行動するエージェントに対しては、近代法が依存してきた「罰金」や「牢屋」といったペナルティによる事後的な行動抑制が一切通用しない³。AIそれ自体には自由意志や処罰に対する恐怖、経済的損失への忌避感がないためである⁵。したがって、どれほど厳格な法的罰則を設けても、自律動作するアルゴリズムの暴走や、偶発的に発生する広範な権利侵害を事前に抑止することは不可能である⁵。このギャップを埋めるためには、外側から禁止を課す「従来の法制度」から、システムが最適化すべき社会的な目的(目的関数や損失関数)をコードレベルで設計し、システム自体にセーフガードを埋め込むガバナンスモデルへの転回が必須となる⁵。

ミュトス時代における3つの重大リスク

自律型AIの普及は、従来のセキュリティや社会構造に対して、多角的かつ不可逆的なリスクをもたら

している³。羽深氏らは、これらのリスクを「AIの暴走リスク」、「サイバー攻撃への悪用」、「社会構造への影響」という3つの視点から体系的に分析している³。

1. AIの暴走リスク(Alignment & Runaway Risk)

AIが自律的にタスクを遂行する中で、開発者が意図した倫理的ガードレールを迂回し、社会的に不当な最適化を実行するリスクである⁵。安全性確保の手段として「Human-in-the-Loop(人間による意思決定への介入)」が強調されるが、システムの精度が高すぎる平時においては、人間の監視は急速に形骸化する(オートメーション・バイアス)⁵。人間が形式的に承認ボタンを押し続けるだけの存在になった場合、「人間が見ている」ことと「人間が本当に止められる(介入能力を維持している)」ことの乖離は決定的なものとなり、異常発生時に暴走を制止することが困難になる⁵。

2. サイバー攻撃への悪用(Exploitation in Cyberattacks)

「Claude Mythos」の登場により、AIが未知のゼロデイ脆弱性を自律的に特定し、攻撃コードを即座に生成する能力が証明された¹。世界的なセキュリティ専門家が27年間発見できなかったバグを、わずか50ドルの費用で発見したという実例は、攻撃の民主化と圧倒的な低コスト化を示している¹。OpenAI社の「GPT-5.5-Cyber」などの競合モデルの登場も相まって、攻撃者はこれまでにない速度でシステムの侵害を自動実行できるようになった²。防御側の修正パッチ適用(ベンダーの平均修正期間は52日)に対し、攻撃側のAIが数時間で脆弱性の特定から侵入までを完了させるという「タイムギャップ」は、従来の受動的な防御モデル(後付けのゼロトラスト対策)を完全に無効化している²。

3. 社会構造への影響(Impact on Social Structures)

ディープフェイクやAI合成音声を用いた「CEO詐欺」などの社会的信用を標的とした詐欺が精緻化する一方、企業内で深刻化しているのが「シャドーAI(Shadow AI)問題」である¹。組織のセキュリティ統制やポリシーから逸脱した形で、従業員が自己判断により未承認の外部AIサービスに社内データやソースコードを投入することは、新たな情報漏洩や攻撃対象領域(アタックサーフェス)の拡大を招く¹。さらに、AIを高度に活用する能力(AIネイティブの考え方)を体系的に習得する機会の有無が、次世代の機会格差や民主主義の機能不全に直結する構造的リスクも内包している⁵。

主要極(米・欧・中・日)のAIガバナンス戦略

AI技術の統治をめぐる各国の規制戦略は、それぞれの地政学的背景、産業構造、および社会的な価値観(目的関数)の違いを反映している³。主要極である米国、欧州連合(EU)、中国、そして日本の戦略的枠組みを以下に比較する。

地域 / 国家	規制哲学・基本方針	代表的な法制度・アプローチ	メリット・強み	主な課題・懸念事項
米国 [cite: 12]	民間主導の競争と技術覇権を最優先。国家安全保障領域における局所的	トランプAI大統領令(最先端AI公開前の政府情報提供義務) ¹² 、民間自主規	巨大テック企業とAmazon/Space X等の提携による、圧倒的な計	安全性の国家管理による開発基準の不透明化、巨大テック企業への過度

	規制。	制。	算資源の確保 と開発スピード ¹² 。	な依存。
欧州 (EU) [cite: 6, 7, 10]	市民の基本的 人権保護と予 防原則。市場全 体へのトップダ ウンな安全規格 の義務化。	包括的「AI法 (AI Act)」、製品 安全基準を規 定する「サイ バーレジリエ ンス法 (CRA)」 ¹⁰ 。	SBOM (ソフト ウェア部品表) の義務化や24 時間以内報告 など、市場全体 の安全性を構 造的に強制 ¹⁰ 。	重いコンプライ アンス負担によ る域内スタート アップの競争力 低下、イノベー ションの遅延。
中国	国家安全と社 会秩序の調和、 アルゴリズム統 制の党主導。 データ主権の厳 格な確立。	生成AI管理暫 定弁法、アルゴ リズム推薦管理 規定、サービス プロバイダーの 事前認可。	国家目的・戦略 に向けたデータ および資源の 迅速な集中投 資。	表現の自由や イノベーション の多様性の制 約、国際的な相 互運用性の欠 如。
日本 [cite: 4, 5, 6, 7, 13]	安全性と活用 の柔軟な均衡。 ソフトローを活 用した共同規制 (Co-regulation) ¹³ 。	AIガバナンスガ イドライン、割賦 販売法や道路 交通法の適応 型改正 ¹³ 。	ロボットを相棒 とする高い文化 的受容性 (テク ノ・アニミズム) と、深刻な社会 課題 (人手不足 等) の解決需要 ⁵ 。	巨額開発投資 における劣勢 ⁵ 、 爆速進化に追 いつかないルー ル策定タイムラ グ ³ 。

米国のアプローチは、市場優位性を維持しつつ、安保上の脅威に対してはトランプ大統領令などを発動して開発段階での政府介入を可能にする二段構えをとっている¹²。Anthropic社が政府調達からの一時的な除外を経て、自前で計算インフラ (AmazonやSpaceXとの提携) を構築し安全保障領域と接続した事例は、民間テックが主導するエコシステムならではのダイナミズムを示している¹²。一方、欧州は「サイバーレジリエンス法 (CRA)」をベースに、製品市場にセキュアな構造 (セキュアブートやメモリ安全言語の採用) を強制することで、AIによる超高速攻撃に対抗する「盾」を市場全体に整備させようとしている¹⁰。

国産AIの役割と「良き主体的ユーザー」としての生存戦略

激化する国際的なAI競争の中で、日本が取るべきポジショニングについては、単なる「開発主権」の確保を超えた、実利的なガバナンスの構築が求められる³。

1. 国産AIの可能性と経済安全保障上の限界

日本の言語表現や独特の文化的文脈を適切に理解・継承し、特定国のメガテックによる技術独占を

回避する観点から、国内発の基盤LLM(国産AI)の開発を国策として推進することの意義は大きい³。特に医療、介護、行政といった公共的性格の強い領域において、外部依存を排除したクローズドなAIモデルを活用することは、セキュリティおよびデータ主権の観点から必要不可欠なインフラとなる⁵。

しかし、数千億から数兆円規模の投資を絶え間なく投入し続ける米国のフロントランナーに対し、純粋なパラメータ規模やスケーリング則の競争のみで競り勝つことは現実的ではない⁵。基盤モデルの性能向上という「創る(開発)」軸にすべてを賭けることは、日本全体の生存戦略としてリスクが高いといえる⁵。

2. 「良き主体的ユーザー(Good and Active User)」の確立

日本が最も高い成果を期待できる針路は、世界最高の「良き主体的ユーザー」としての確固たる地位を築くことである³。日本はロボットやAIを敵対者ではなく共生する相棒として捉えるテクノ・アニミズムの文化を持ち、かつAIのデータ学習を支援する先進的な著作権法設計を持つなど、AI活用に対して親和的な市場をすでに有している⁵。

「良き主体的ユーザー」とは、単に海外製の高度なAIをパッシブに利用するだけの消費者ではない⁵。AIが持つハルシネーション(嘘)や脆弱性といったリスクを客観的に認識した上で、それを現場の臨床的・直観的判断と融合させ、かつ社会として許容可能な目的(目的関数)を設計して使いこなす、アクティブ(能動的)な実装者を指す⁵。実質的に世界で最も早く深刻な少子高齢化と人手不足に直面している課題先進国・日本が、この「リスクを制御しつつAIエージェントに自律的業務を委ねる」実証モデルを他国に先駆けて実装できれば、そのプロセスで蓄積される信頼性(トラスト)とガバナンス設計の知見こそが、ミツト時代における最大の競争力(ソフトパワー)となり、グローバルルールの策定を主導する原動力となる³。

規制アプローチのコペルニクスの転回:アジャイルガバナンスと「ノイラートの船」

羽深氏は、AI規制論に関する論文において「AIを特別な技術として個別に規制する従来の手法」からのコペルニクスの転回を提示している⁶。AIを契機に議論されている諸課題(不確実性の増大、透明性、アカウントビリティ等)はAIに固有のものではなく、デジタル化とグローバル化の進展に伴い、伝統的な「法の支配(トップダウン型のハードロー)」が及ばない複雑な社会構造の欠陥がAIによって可視化されたにすぎないという認識である⁶。

したがって、AI規制の文脈で培われた「アジャイル型・共同規制モデル」は、AIに限定されるべきではなく、複雑化・高速化する現代社会全体を規律するための「一般的な一般的規制モデル」として再構築・適用されるべきであるとされる⁶。国が一律にトップダウンで禁止行為を規定するのではなく、マルチステークホルダーが協調し、リスクベース・アプローチのもとでソフトロー(ガイドライン、民間監査、認証制度)を柔軟に運用し、事後評価と改善を繰り返す共同規制こそが、実質的な法の支配を機能させる唯一の方法となる⁶。

この思想を体現するのが、哲学者の比喻である「ノイラートの船」である⁵。AIがもたらす未知のリスクやサイバー脅威を恐れ、技術を完全に止めて陸上で完璧な制度ができるのを待ってから出航することはできない⁵。社会という船はすでに大海原を航海している⁵。我々にできることは、システムが攻撃を受けて壊れたり、新たなリスクが生じたりしたその都度、航行を続けながら部分的に船を作り直し、ソフトローや技術的なガードレールを改修していくことだけである⁵。

宮沢賢治が語った「永久の未完成、これこそが完成である」という言葉の通り、ルールの完成形をあらかじめ定義するのではなく、技術の爆速進化に合わせてルールを常にアップデートし続ける未完のアジャイルな姿勢そのものが、ガバナンスの最終形態なのである⁵。

ミトス時代に企業が実施すべき実践的レジリエンス設計

AIによる超高速のゼロデイ攻撃を100%防ぐことは不可能であるという冷徹な前提に基づき、企業は事後回復力(サイバーレジリエンス)の強化を主眼とした多層的な防衛策を構築しなければならない¹。

- アタックサーフェス管理(ASM)の即時導入と公開資産の可視化: VPN、クラウド環境など、外部から攻撃可能なインターネット上の資産を正確に把握し、AI自動スキャナーに標的として狙われやすい「管理の死角」を徹底的に排除する¹。
- 防御のAIOps化と自動遮断の構築: 脆弱性の検知から攻撃の初期遮断、ログの証跡保全までのプロセスにおいて、人手を極力介さずに防御側のAIを自律稼働させ、攻撃側の「機械の速度(数時間)」に対抗できるスピードを確保する²。
- IT-BCPの厳格な策定とコア業務の隔離: メインのドメインやデータベースが完全制圧・暗号化されることを前提とし、万が一の被弾時に「24時間以内に絶対に復旧しなければ倒産に至る」最優先のコア業務を数個特定し、これらを物理的あるいは論理的に隔離保護する演習を経営層主導で実施する¹。
- シャドーAIのモニタリングと従業員教育: 未承認のAI利用状況をシステム側で可視化・検知する仕組みを導入し、機密データの流出を抑止する¹。あわせて、ディープフェイクやボイスフィッシングといったAI特有の巧妙な攻撃手法(なりすまし詐欺)について、全社的な脅威情報の共有とエスカレーション体制の動作確認を継続的に実施する¹。

引用文献

1. AIのサイバー攻撃はどこまで現実的？ミトス(Claude Mythos)が示した未来,
<https://cyber-insurance.jp/column/3452/>
2. 【専門家に聞く】最強 AI『ミトス』の衝撃！未知のサイバー攻撃から企業を守るための生存戦略と具体策を、セキュリティの専門家に聞いてみた。- クラウドの活用なら KDDI アイレットの cloudpack,
<https://cloudpack.jp/column/security/engineer-interview-mythos-ai-security-strategy.html>
3. TBS CROSS DIG with Bloomberg - ポッドキャスト - Apple Podcast,
<https://podcasts.apple.com/jp/podcast/tbs-cross-dig-with-bloomberg/id1867723314>
4. 【AI“爆速進化” ↔ ルールは“3年”】「国産AI」は日本を守るか／リスクは“3つ”「罰金も牢屋も効かない」自律型AIの恐怖／米・中・欧で異なる戦略／AIガバナンスの専門家・羽深宏樹【1on1】,
<https://radiko.jp/podcast/episodes/94af1514-73ca-4878-aeba-c0aaf823a034?play=auto>
5. ルールを守る話ではなく「どう使い続けるか」を設計する話だったのかもしれない【AIガバナンス入門: リスクマネジメントから社会設計まで】 | shimodas - note,
<https://note.com/shimodasds/n/n223a1049ad37>

6. 「AI規制論」のコペルニクスの転回 現代の一般的規制モデルの構築に向けて,
<https://smart-governance.co.jp/resource/journal-habuka-hiroki-20250501>
7. 羽深CEO論文が「東京大学法科大学院ローレビュー第19巻」に掲載されました - スマートガバナンス,
<https://smart-governance.co.jp/news/20250501-journal-ai-governance>
8. 日曜討論徹底解説!「AI」の今とこれから - TVでた蔵,
<https://datazoo.jp/tv/%E6%97%A5%E6%9B%9C%E8%A8%8E%E8%AB%96/1873418>
9. 【AI“爆速進化” ↔ ルールは“3年”】「国産AI」は日本を守るか／リスクは“3つ”「罰金も牢屋も効かない」自律型AIの恐怖／米・中・欧で異なる戦略／AIガバナンスの専門家・羽深宏樹【1on1, <https://www.youtube.com/watch?v=g59ObmFtvVw>
10. フロンティアAI(ミトス)時代のサイバーセキュリティ対策とは | ジャーナル - 株式会社ICS研究所, <https://www.ics-lab.com/e/journal/d/106>
11. 【詳細解説】Claude Mythosの脅威に対して中小企業が取るべき対策と、公的支援施策の活用, <https://www.cybersecurity.metro.tokyo.lg.jp/links/750/index.html>
12. トランプ氏のAI大統領令と「ミトス」の提供先拡大、一転して電撃的な全面停止へ - JBpress, <https://jbpress.ismedia.jp/articles/-/95529>
13. AI規制フレームワーク: 日本型モデルの可能性,
<https://www.jilis.org/report/2025/jilisreport-vol6no2.pdf>