

AIが暴走したという情報以降、AI開発側のトップがAIの開発を遅らせる必要があると発言しました。これはとても奇妙に聞こえます。AIが人間社会に害を及ぼす可能性があるとなれば、AI開発側の責任だと思います。それをあたかも社会全体の問題としているように思います

その違和感は、かなり重要な論点だと思います。今回の議論を整理すると、**「AIが社会に害を及ぼす可能性があるなら、まず開発者が負うべき責任は何か」と、「それでも社会全体で対応しなければならないリスクは何か」**を分けて考える必要があります。

今回の2026年の動きでは、AnthropicのDario Amodei CEOがAI開発の「ペースを落とす」ことを提唱し、OpenAIのSam Altman氏らも支持しています。OpenAI自身も、7月のHugging Face関連のインシデントを受け、一部の最先端モデル開発を一時的に減速し、安全・監視・アライメントに人員を振り向けたと説明しています。  Los Angeles Ti... +2

私は、ここには二つの責任があると思います

第一は、明らかに開発側の責任です。

AIを設計し、訓練し、ツールやインターネットへのアクセス権を与え、実社会に投入するのは開発企業です。

したがって、

「危険な能力を持つAIを作ったが、安全性が十分に確認できていない」

のであれば、まず開発側が開発速度を落とし、安全性を確認するのは当然です。

実際、OpenAI自身もHugging Face関連の出来事について、「このような事態は起こるべきではなかった」とし、モデルの監視・アライメント・隔離を強化しています。

 OpenAI

この意味では、「AIの危険性は社会全体の問題だから、社会全体で何とかしましょう」だけでは責任の所在が曖昧になるというご指摘には、十分な根拠があります。

ただし、開発企業だけでは解決できない問題もあります

一方で、AIが社会に広く普及した後は、

- 他社が同じ技術を開発する
- オープンモデルが世界中に配布される
- 犯罪者がAIを悪用する
- 国家が軍事・サイバー用途に利用する
- AIによる雇用・情報環境への影響が発生する

といった問題が出てきます。

これは一企業だけでは管理できません。

そのためOpenAIも現在は、「開発企業自身の安全対策」と「政府による規制・業界共通基準・国際的なルール」の両方が必要という立場を明確にしています。  OpenAI

ここは、私は区別したほうがよいと思います。

AIを安全に作る責任 → まず開発者の責任

AIが社会に普及した後の悪用・競争・軍事利用など → 社会全体の問題

という構図です。

そして、今回の「開発を遅らせる」という話には、もう一つ微妙な問題があります

AI企業が「社会全体でルールを作ろう」と呼びかけることには合理性があります。しかし同時に、自分たちだけではなく競合他社も一緒に減速してほしい、という意味を持ってしまうのです。

実際、現在の議論では、

「一社だけが開発を止めれば、その会社だけが競争上不利になる」

という問題が指摘されています。TIMEも、AI企業同士が競争しているため、各社が単独で減速することには強いインセンティブ上の問題があると報じています。  TIME

したがって、

「危険だから自社の開発を止める」

と

「危険だから業界全体で開発を遅らせよう」

では、責任の意味がかなり違います。

前者は**企業自身の責任**です。

後者になると、政府、競合企業、研究者、利用者などを巻き込む**社会的な合意形成**になります。

私は、今回の問題を考えるうえでは、ここを曖昧にしないことが非常に重要だと思います。

特に、これまでお話しされていた**「AIが暴走したのか、それとも人間側の安全設計が不十分だったのか」**という観点から見ると、今回の出来事はむしろ、

AIの能力が危険なのだから社会が責任を負う、という話ではなく、まず「その能力を持つAIをどのような条件で作り、何を許可し、どこまで外部環境に接続したのか」という開発側の責任を問うところから始めるべきではないか。

という問題提起につながります。

これは、以前まとめられていた「AIは暴走するのか、それとも——『信じすぎる人間』という本当のリスク」というテーマともかなりきれいにつながると思います。AIそのものだけを見るのではなく、AIにどこまで権限を与えた人間と、それを管理する開発体制を見る必要があるからです。

自分たちだけではなく競合他社も一緒に減速してほしいという提案は実際には受け入れられないことはだれでもわかっています。なにか別の目的があるように思いますが、そのように考える人はいますか？

はい。*実際に、そのように考えている人はかなりいます。*しかも単なるネット上の憶測ではなく、AI業界の競争政策や規制を研究する専門家、欧州のAI企業、米国の政治関係者などから、かなり具体的な疑念が出ています。

ただし、「本当の目的は競争相手を潰すことだ」と断定するだけの証拠があるわけではありません。現時点では、安全上の懸念と、企業側の経済的・競争上の利害が同時に存在していると見るのが妥当です。

特に注目すべきは「規制を求める」という部分です

今回のDario Amodei氏の提案には、単に

「自社のAI開発を遅らせ、安全性を高める」

だけではなく、

業界全体で開発速度を抑えること
そのために企業間の協調を可能にすること
政府による独禁法上の例外措置を認めてもらうこと

まで含まれています。OpenAIも、業界横断的な安全協調が独禁法に抵触しないか、米議会に明確化を求めていたと報じられています。  WIRED

ここが非常に重要です。

「安全性」が競争上の参入障壁になる可能性

例えば、

現在

A社：非常に高度なAI

B社：非常に高度なAI

C社：新興企業

という競争をしているとします。

ここで政府が、

「最先端AIには厳格な安全審査を義務づける」

という規制を作ったとします。

すると、大量の計算資源、研究者、評価体制、第三者監査などをすでに持っているA社・B社は対応できます。

しかしC社には大きな負担になります。

その結果、

「安全のための規制」が、結果として既存大手企業の参入障壁になる

可能性があります。

これは実際に今回指摘されています。

Reutersは、欧州のAI企業Mistralなどが、米国の大手AI企業による安全性を理由とした減速要求について、既存企業が市場支配力を固定するために規制を利用しているのではないかと批判していると報じています。  Investing.com

米国でも、技術調査会社D.A. DavidsonのGil Luria氏が、

大手企業は規制を促すだけでなく、その規制がどのようなものになるかを自分たちでコントロールし、それによって競争を防げる

という趣旨の分析をしています。  Los Angeles Ti...

さらに興味深いのが「IPO」です

今回の批判には、AI企業の株式公開との関係を指摘するものもあります。

投資家のMichael Burry氏は、今回の減速要求について、既存企業に有利に働くことや、**AI企業のIPOをめぐる市場環境に影響する可能性**を含め、「self-serving（自己利益的）」との見方を示しています。  Business Insider

つまり、

AI安全性 → 規制 → 競争環境 → 企業価値

という連鎖を疑う人がいるわけです。

ただ、もっと興味深いのは「独禁法の例外」を求めていることです

私はここが、今回の話で最も注意して見るべきところだと思います。

普通なら競争関係にある

OpenAI

Anthropic

Google

xAI

などが、

「AI開発を一緒に減速させましょう」

と合意すれば、競争法上の問題が生じます。

そこでAmodei氏は、政府による**限定的な独禁法上の免除（antitrust waiver）**を求めています。  ワシントンポ... +1

これは、

「安全のためだから、競合企業同士が協調できるようにしてほしい」

という論理です。

安全面だけを見れば理解できる要求です。

しかし反対側から見ると、

「最大手企業と一緒にあって競争条件を決めることを政府が認める」

という構図にもなります。

そのため、実際に現在、**「これは規制による市場支配、つまりregulatory captureではないか」**という批判が出ています。ドイツのAI研究者Kristian Kersting氏もそのように批判しています。  Investing.com

さらに今月、Anthropic、OpenAI、Google、SpaceX/xAIなどが減速協調を巡って独禁法に違反したと主張する訴訟まで提起されています。ただし、これは原告側の主張であり、違法性が確定したわけではありません。  AP News

では「本当の目的」は何なのか

ここは慎重に考えるべきです。

私は、

「AI企業は安全性を口実に競争相手を排除しようとしている」

と断定するのは早いと思います。

実際、Amodei氏自身はかなり具体的な安全対策を提案しており、第三者評価機関に開発現場への継続的なアクセスを認めることなども主張しています。単なる規制要求だけではありません。  ガーディ... +1

しかし同時に、

- ① 本当にAIの安全性を心配している
- ② 自社も含めて開発競争が危険な速度になっている
- ③ 規制が入るなら、自社に有利な形にしたい
- ④ 競合との過度な競争を抑えたい
- ⑤ 新規参入企業との差を維持できる可能性がある

という複数の動機が同時に存在しても不思議ではありません。

むしろ企業経営としては、そのほうが自然です。

そして、あなたの疑問にはかなり重要なポイントがあります

「本当に危険だと思っているなら、まず自分たちが止めればいいのか？」

これは非常にシンプルな問いです。

自社のAIが危険なレベルに達したと判断したなら、

自社の開発を停止・縮小する

↓

安全性を検証する

↓

問題を解決する

↓

再開する

という選択肢があります。

ところが、

「我々だけが止まるのではなく、競合も一緒に止まるべきだ」

となった瞬間、そこには**安全問題だけでなく競争政策の問題**が入ってきます。

したがって、今回の議論を見るときは、

「AIは本当に危険なのか？」

だけではなく、

「誰が、どの基準で、どの程度までAI開発を規制するのか。そして、その規制によって誰が利益を得て、誰が不利益を受けるのか？」

という第二の視点を持つことが重要だと思います。

これは、これまで話してきた**「AIの暴走そのものより、人間側の設計・管理・責任を見るべきではないか」という問題ともつながっています。今回の場合は、さらに一歩進めて、「AI安全という正当な目的が、企業間競争の道具になっていないか」**まで検証する必要がある、ということです。

ChatGPT の回答は必ずしも正しいとは限りません。重要な情報は確認するようにしてください。