

サム・アルトマン伝：人工知能革命の開拓者

目次、プロローグ、7部、20章



キム・ギョンジン、キム・ギョンラン

Japanese PDF edition

サム・アルトマン伝：人工知能革命の開拓者

キム・ギョンジン、キム・ギョンラン

サム・アルトマンの少年時代、起業、Y Combinator、OpenAI、ChatGPT、解任と復帰、AI時代の責任をたどるオンライン伝記です。

目次

プロローグ

1.1 幼少期と学生時代

1.2 19歳の青年が300億ウォンを集める

1.3 Y Combinatorのリーダーシップ、スタートアップ界のキングメーカー

2.1 AI革命を予感し、OpenAIを創設する

2.2 モデル開発から製品化までの歩み

2.3 イーロン・マスクとの対立

3.1 CEO解任、5日間の悪夢

3.2 社員たちの劇的な忠誠

3.3 危機を越えて残った教訓

4.1 組織と開発の哲学

4.2 OpenAIの未来のビジネスモデル

4.3 生産性と仕事の管理

5.1 モデルはどこへ進化するのか

5.2 エージェントとコーディングの役割

5.3 人間とAIの相互作用の未来

6.1 規制と国家安全保障競争

6.2 インフラとエネルギーの課題

6.3 AIと人間社会の共存

7.1 個人的な刺激とビジョン

7.2 最後の助言と展望

プロローグ

次の章 →

2022年11月30日、ChatGPTが公開されました。30分ほど会話を交わしたあと、背筋が凍りつきました。根本的に異なる時代へと踏み込んだような感覚でした。5日で100万人、2か月で1億人を超えました。多くの変化を見てきましたが、これほど速い変化は初めてでした。

冗談のようにこんなことを言います。「BCはBefore Christではなく、Before ChatGPTだ」と。笑い飛ばすこともできますが、まったく外れた話でもありません。ChatGPT以前と以後は、まるで別の世界です。

AI政策を研究するなかで、多くの人々と会いました。アメリカは「イノベーションを妨げるな」と言い、ヨーロッパは「安全が先だ」と言い、中国は「我々が主導する」と言いました。

しかし、たった一つだけ、皆が同じことを言っていました。「ChatGPT以前には戻れない」と。

サム・アルトマンは「知能の時代はエネルギーの時代だ」と言います。最初は意味が分かりませんでした。AI政策を研究するうちに気づきました。AIを動かすには電力が必要で、エネルギーが未来の限界を決めるのだということ。

生きてきた中で、才能のある人をたくさん見てきました。しかし、機会を得られなかった人をもっと多く見てきました。ソウルに住んでいないから、良い学校に行けないから、親の経済力のせいで、塾の費用がないから。胸が痛みました。

サム・アルトマンが語る「人類の潜在力を最大限に引き出す」とは、こうした問題を解決することです。AIがすべての人に最高水準の知識と教育を届けられるとしたら？地方に住む子どもが、ソウル・江南の子どもと同じ教育を受けられるとしたら？あなたの8歳の子に世界中の知識を教えてくれる家庭教師がいるとしたら？がんの診断を10倍速くできるとしたら？

これが本当の革命です。

この本をまず子どもたちに見せました。小学生は「私もサム・アルトマンみたいになれる？」と聞きました。中学生は「AIって怖いけど、不思議でもある」と言いました。

どちらも正解です。

私たちの子どもたちが育つ世界は、今とはまるで違うものになるでしょう。怖いこともあるかもしれない。不思議なこともあるかもしれない。面白いこともあるかもしれない。

サム・アルトマンのように変化をつくる人になることもできるし、その変化を賢く使いこなす人になることもできます。どちらでも、それでいいのです。

このプロローグを起点に、私たちは一人の人物の人生を通じて新しい時代を理解する旅を始めます。

キム・ギョンジン、キム・ギョンラン 2025年

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

1.1 幼少期と学生時代

第1部 屋根裏のコンピューター少年からシリコンバレーのキングメーカーへ

← 前の章

次の章 →

屋根裏のコンピュータ少年、世界とつながる。

1985年4月22日、シカゴのごく普通の家庭に始まります。母のコニー・ギブスタインは皮膚科医で、父のジェリー・アルトマンは不動産業者でした。4人きょうだいの長男として生まれたサムは、幼い頃に家族とともにミズーリ州セントルイス郊外のクレイトンへ引っ越しました。1989年、4歳の時のことです。

クレイトンは、アメリカ中西部の典型的な住宅街でした。静かな通り、手入れの行き届いた芝生、古いレンガ造りの家々が並ぶ、そういった街並みです。表向きはごく普通のアメリカ中産階級の家庭の姿でした。夕食時には家族が集まり、卓球やビリヤードをしながら時間を過ごしていました。

サムの人生に何か特別なことが起きたのは、彼が8歳のときでした。

両親がプレゼントしてくれたApple Macintosh。幼いサムにとって、それはまったく新しい世界への扉でした。ほかの子どもたちが外で野球をしたり自転車に乗ったりしているとき、サムは屋根裏部屋に腰を下ろし、コンピュータの前で時間を過ごしていました。夜遅くまで屋根裏の窓に明かりが灯っている日が多くありました。

サムはコンピュータでゲームをしていたわけではありませんでした。分解したのです。ハードウェアがどう動くのか知りたかったのです。ネジを一本一本外し、部品を眺め、また組み立てる作業を繰り返しました。8歳の子どもが、です。

プログラミングを学び始めました。

コードを入力すると画面に何かが現れるのが不思議でした。命令を下すとコンピュータがそれに従うのは、まるで魔法のようでした。サムは自分を"コンピュータ・ナード"と呼んでいました。友人たちも彼をそう呼んでいました。口数は少ないけれど、一度何かにはまるとレーザーのように集中する子でした。

幼稚園に通っていた頃、市外局番の体系を理解していたという逸話があります。ふつうの子どもなら数を数えるのも難しい年齢なのに、サムはアメリカ全土の市外局番がどんな規則で割り当てられているかを把握していたのです。数字やパターン、システムを読み取る能力が際立っていました。

サムにとってコンピュータは単なる機械ではありませんでした。世界とつながる窓でした。AOLのチャットルームで、サムは自分と似た人々と出会いました。匿名で会話し、悩みを打ち明け、慰めを受けました。Macintoshは彼にとって、世界へと続く命綱だったのです。

"AOLのチャットルームを見つけたことが、私の人生を変えました。"

屋根裏部屋で過ごした無数の夜。コンピュータの画面の前でコードを打ち込み、インターネットを探索し、新しい世界を発見し続けた時間。その経験がサムを作りました。好奇心旺盛で、粘り強く、世界を変えられると信じる人間として。

サムは後にこう語っています。"私は屋根裏部屋に座って夜遅くまでプログラミングを学んでいた子どもでした。あの頃から今のOpenAIまで、一本の直線でつながっていると思っています。"

セントルイスのあの屋根裏部屋がなければ、ChatGPTも生まれなかったかもしれません。幼いサムがコンピュータの前で過ごした無数の時間が、人類の歴史を変えるAI革命の出発点だったのです。

高校時代：水球、水泳、フェンシング、そしてポーカー

ジョン・パロウズ・スクールという私立学校に入学しました。セントルイスのレイド地区に位置するその学校は、アメリカ屈指の名門私立高校でした。年間学費が3万ドルを超える、裕福な家庭の子女が通う場所です。進歩的で自由な教育理念で知られていました。

サムは学校でも相変わらず"コンピュータおたく"でした。学校のウェブサイトを自分で作り、授業の時間割を管理するソフトウェアまで開発しました。15歳、16歳の子どもがです。友人たちは彼を"勉強虫"と呼んでいましたが、サムは勉強だけしている子どもではありませんでした。

スポーツが好きでした。

なかでも水球に熱心に取り組みました。水球は水中で行う激しいスポーツです。体力も必要で、チームワークも重要です。サムは水球チームでプレーし、兄弟たちとともに練習しました。放課後はフォレスト・パークへ行って自転車に乗ったり、水泳をしたりしました。

フェンシングもしていました。フェンシングは瞬時の判断力が問われるスポーツです。相手の動きを読み、素早く反応し、戦略を立てなければなりません。サムが後にビジネスで見せた素早い意思決定能力は、もしかするとフェンシングで培ったものかもしれません。

友人たちはサムを"親切だが、ものすごく粘り強い人物"として記憶していました。目標を定めたら最後まで追い求める性格でした。コンピュータのプロジェクトであれ、スポーツであれ、何でも同じでした。

高校を卒業する頃には、サムはもう普通の10代ではありませんでした。コンピュータプログラミングの腕前は大学生レベルで、オンラインビジネスの経験もあり、社会的な問題への自分なりの意見も持っていました。スポーツで鍛えた体力と精神力もありました。

サムには粘り強さがありました。何かを始めたら最後までやり抜く性格。障害があれば30通りの別の方法で突破する根気。この特性は、後にLooptを創業し、Y Combinatorを率い、OpenAIを立ち上げる際にもそのまま現れました。

ジョン・パロウズ・スクール時代のサム・アルトマン。

屋根裏のコンピュータ少年は、いよいよ外の世界へ踏み出す準備を整えていました。

スタンフォード進学、そして運命の選択

2003年、19歳のサム・アルトマンがスタンフォード大学に入学しました。

スタンフォード。シリコンバレーの中心に位置する、世界屈指の名門大学です。Google、Yahoo、Hewlett-Packard、Instagram.....数多くのテクノロジー企業がスタンフォードの寮や教室から生まれました。サムにとってスタンフォードは"当然の選択"でした。

専攻はコンピュータサイエンスでした。当然の選択です。8歳からプログラミングをしてきた子どもでしたから。スタンフォードのCSプログラムは世界最高水準でした。優秀な教授陣、最新の研究施設、そして何より"一緒に学びたいと思える人たち"が集まっている場所でした。

サムはスタンフォードでコーディングだけを学んだわけではありませんでした。

"あらゆる科学の授業"を受けました。物理学、生物学、化学.....ライティングの授業も受けました。人文学の授業も。友人たちがなぜコンピュータ専攻の学生がそんな授業を取るのかと尋ねると、サムはこう答えました。

"後から振り返れば、こういった授業のほうがCSの授業よりずっと役に立つはずです。"

彼の言葉は正しかったのです。OpenAIを率いるサムは、単なる技術者ではありませんでした。哲学者でもあり、経済学者でもあり、政策立案者でもありました。AIが人類に与える影響を理解するには、技術だけでは足りなかったのです。

サムはこう語っています。"大学で最も重要だったのは、あらゆる面白いアイデアを追い求める本当に賢い人たちのそばにいられたことです。"

物理学専攻の友人たちが最も興味深かったといいます。宇宙の本質について議論し、複雑な数学の問題を夜通し解き、世界を違う目で見える方法を教えてくれました。

サムには、別の「授業」がありました。

ポーカーです。

毎晩、友人たちとポーカーに興じました。サムは後にこう打ち明けています。

「ポーカーで大学の生活費を稼ぎました。講義よりポーカーから多くを学びました。」

何を学んだのでしょうか。

パターン認識。相手の行動を観察し、心理を読み、次の一手を予測する能力。不確実性の中で意思決定する方法。リスクを計算し、時には大きな賭けに出る勇気。感情をコントロールし、失敗しても立ち上がる回復力。

サムはこう語っています。「ビジネスについて多くを教えてくれたものが二つあります。ポーカーとエンジェル投資です。どちらも強くお勧めします。」

2018年のあるポッドキャストで、彼はこう付け加えました。「ポーカーは誰にでも向いているわけではありませんが、世界やビジネス、心理学、リスクについて学ぶ良い方法です。」

スタンフォードの2年生になった年、サムは重大な決断を下しました。

学校を辞めることにしたのです。

友人たちは驚きました。スタンフォードを中退するというのは、世界最高峰の大学を学位なしで去るというのは、サムは確信に満ちていました。

「私は2年生で、（今すぐ）行きます。」

Y CombinatorのポールGrahamが投資をためらったとき、サムが送ったメッセージです。サムは自身初のスタートアップ、Looptを創業することを心に決め、学校はただの障壁でしかありませんでした。

ポール・グレアムのパートナー、ジェシカ・リビングストンは、サムと初めて会ったときをこう振り返ります。「年齢よりずっと賢いと、すぐにわかりました。」

サムはスタンフォードで2年間を過ごしました。学位は得られませんでした、はるかに貴重なものを手にしました。ポーカーで鍛えた意思決定能力、さまざまな学問から得た広い視野、そして何より「今すぐ動かなければならない」という確信。

2017年、カナダのウォータールー大学はサムに名誉工学博士号を授与しました。皮肉なことに、サムは大学を卒業していないにもかかわらず、博士号を受け取ったことになります。

スタンフォードを去るとき、サムはこう思っていたはず。「学校が教えるのは過去のことです。私は未来をつくりたい。」

そして彼は、本当に未来をつくりました。

AIが機能しなかった時代：アンドリュー・ン教授の研究室

スタンフォードの1年生と2年生の間の夏休み。

サム・アルトマンは特別な場所で働いていました。アンドリュー・ン（Andrew Ng）教授の人工知能研究室です。

アンドリュー・ンとは何者でしょうか。AI分野の伝説的な人物です。Google Brainを創設し、Courseraを共同創業し、機械学習分野で最も影響力のある研究者の一人です。2000年代初頭、彼はスタンフォードのAI研究室を率いていました。

サムはその研究室で夏の間ずっと働きました。19歳の大学生が最先端のAI研究に触れたのです。

問題がありました。

AIが機能しなかったのです。

2000年代初頭のAIは、今とはまったく違いました。ChatGPTのように自然に会話するAIなど、夢にも思わなかった時代です。当時のAIは、簡単な作業すらまともにこなせませんでした。画像認識？

散々たるものでした。自然言語処理？ほぼ不可能でした。

コンピューティングパワーが足りませんでした。データも足りませんでした。アルゴリズムも原始的でした。研究者たちは小さな実験を行い、失敗し、また試み、また失敗しました。

AIは「傍流」の分野と見なされていました。AIが実際に役立つとは、誰も信じていませんでした。

サムは違いました。

AIが機能しないという事実に、彼は落胆しませんでした。むしろ、より強い興味を感じたのです。なぜでしょうか。

サムはこう語っています。「成功する確率が低くても、成功すれば最もすばらしく、最も重要で、最も面白いことになるはずです。」

期待値（expected value）の概念です。ポーカーで学んだことです。勝率が10%でも、成功したときの報酬が莫大であれば、そのゲームに賭ける価値があるということです。

サムは研究室で、AIモデルが失敗するのを目の当たりにしました。簡単なパターンも認識できず、とんでもない結果を出し、まったく役に立たなかったのです。

こう振り返っています。「何をすべきかわかりませんでした。手詰まりでした。」

実験のほとんどは失敗でした。ロボットアームで物をつかむ実験。ビデオゲームを学習するAI。どれもうまくいきませんでした。

サムは研究室を去りました。LoopTを創業するためです。AIの世界に戻るまでには、少し時間がかかりました。

サムの心の片隅には、AIへの思いが残り続けていました。うまく動かなかったあの技術。失敗を繰り返すばかりだったあの実験の数々。その記憶が、ずっと彼を離さなかったのです。

2010年代の初め、何かが変わり始めました。

IBMのWatsonがクイズ番組「ジェパディ！」で人間のチャンピオンを打ち破りました。GoogleのDeepMindが開発したAlphaGoは、囲碁の世界チャンピオンに勝利しました。

サムは直感しました。「大きな変化が来ている」と。

2015年末、彼はOpenAIを共同創業しました。

創業当初のOpenAIは、まだ手探り状態でした。2016年のオフィスには、14人ほどのスタッフがいただけでした。「強い信念と方向性、確信はありました。でも、具体的な実行計画はなかったんです。」

人々はホワイトボードの前で議論を交わし、論文を読み、実験を設計しました。

最初は方向性が定まりませんでした。大規模言語モデル（LLM）というアイデアは、あまりにも遠い話に思えました。そこで、ビデオゲームを上手くこなすAIを作れないかと考えたり、ロボットハンドの実験をしたりもしました。

NvidiaのCEOジェンスン・ファンが、自ら最初のDGX-1システムを届けてくれました。象徴的な瞬間でした。AI革命が始まる瞬間だったのですから。

サム・アルトマンのAIへの歩みは、「うまく動かなかった」時代から始まりました。

失敗だらけの研究室。誰も信じなかった技術。先の見えない未来。

サムはその失敗の中に可能性を見出していました。10年後にChatGPTが世界を変えるとは、誰も知りませんでした。サム自身も、正確にはわかっていなかったでしょう。

彼は信じていました。今は動かない技術でも、試し続ければ、いつかは動くようになる、と。

そして、彼は正しかったのです。

スタンフォードのAI研究室でのあの夏がなければ、OpenAIも生まれなかったでしょう。ChatGPTも、GPT-4も、人類を変えるAI革命も、なかったはずです。

失敗は時に、成功よりも大切な教訓をもたらします。サム・アルトマンがAIを動かす方法を知ることができたのは、AIがうまく動かなかった時代を経験したからです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

1.2 19歳の青年が300億ウォンを集める

第1部 屋根裏のコンピューター少年からシリコンバレーのキングメーカーへ

← 前の章

次の章 →

ViednoからLooptへ：アイデアの進化

2005年の春、スタンフォード大学の寮で19歳のサム・アルトマンは人生の岐路に立っていました。コンピューターサイエンスを専攻していた彼は2年生を終えたところで、あと2年通えば誰もが羨むスタンフォードの卒業証書を手にすることができました。卒業後はグーグルやアップルといった大企業に好条件で就職でき、両親も喜んでくれるはずで、将来も安定していました。

サムはこの頃を振り返り、「これは自分の人生設計じゃなかった。ただ別のチャンスが生まれたただけだ」と語っています。彼が言う「別のチャンス」とは、起業のことでした。友人たちと一緒に作っていたアプリが少しずつ形になってくるにつれ、サムは教室の椅子に座っているのがもどかしくなっていました。頭の中は、自分が作ろうとしているサービスのことで一杯だったからです。

サムと友人たちが開発していたプロジェクトの名前は「Viendo（ヴィエンド）」でした。スペイン語で「見る」という意味を持つこの名前には、友人たちが今どこにいるかを「見ることができる」というアイデアが込められていました。今でこそ友人の位置を共有するのは当たり前のことですが、2005年当時はスマートフォンもない時代でした。iPhoneが世に出たのは2007年のことです。

サムとチームは、携帯電話のGPS機能を使って友人がリアルタイムでどこにいるかを地図上に表示するサービスを構想しました。「今夜何してる？」と友人に一人ひとり電話しなくても、アプリを開けば近くにいる友人をすぐ見つけられる仕組みです。偶然同じカフェにいる友人を発見することもできるし、映画を観に行ったついでに近くにいる別の友人と合流することもできました。

アイデアは良かったものの、名前が気に入りませんでした。「Viendo」は発音しにくく、意味もすぐには伝わらなかったからです。チームは頭を寄せ合い、新しい名前を考えました。数週間悩んだ末に出てきた名前が「Loopt（ループト）」でした。英語に「in the loop」という表現があり、「情報を共有し続けること」「つながっていること」を意味します。Looptはまさにこの表現から取った名前でした。友人とつながっているという感覚を与える、短くて覚えやすい名前でした。

名前を変えるとともに、サムの決意も固まりました。両親に打ち明けました。「学校を少し休んで、起業してみたいんです。」両親は心配しましたが、サムは自信に満ちていました。「大事なのは、いつでも戻れるということに気づいたことです。これがリスクを見るうえでの核心です。ほとんどのことは一方通行じゃない。何かを試して、うまくいかなければ引き返せる」と彼は考えていました。

こうしてサム・アルトマンは学生証を引き出しにしまい、LooptのCEOとして新しい人生を歩み始めました。寮ではなく小さなオフィスで、授業ではなく投資家とのミーティングで、試験ではなく製品開発で埋め尽くされた日々です。当時は誰も、この若者がやがて人工知能革命を率いるリーダー

になるとは思っていませんでした。

19歳の少年、300億ウォンを集める

Looptを始めることを決意したものの、サムと友人たちにはお金がありませんでした。良いアイデアだけで会社ができるわけではありません。サーバーも借りなければならず、エンジニアもさらに必要で、オフィスも探さなければなりません。何より、アプリを作っている間の生活費が必要でした。大学生が小遣いで会社を運営するわけにはいきません。

サムは投資家を探し始めました。シリコンバレーには、良いアイデアを持つ若者に投資するベンチャーキャピタリスト（VC）が多くいました。サムはLooptの事業計画書を持って方々を回りました。ミーティングのたびに同じ質問を受けました。「これを誰が使うんですか?」「どうやって収益を上げるんですか?」「フェイスブックやグーグルが同じものを作ったらどうするんですか?」

断られることも多々ありました。ある投資家は「大学をちゃんと卒業してから来てください」と言いました。「位置情報の共有? それはプライバシーの侵害じゃないですか?」と首を振る人もいました。自分よりずっと年上の大人たちの前でアイデアを説明するのは、簡単なことではありませんでした。

しかしLooptのアイデアは本当に斬新でした。2005年当時、ソーシャルネットワークはほとんどがパソコンの画面の中にしか存在しませんでした。フェイスブックで友人の投稿を読んだり、写真を見たり、コメントをする程度のことでした。Looptは、オンラインの友人関係を現実世界に引き出そうという野心的な計画でした。デジタルと現実をつなぐ橋のようなものでした。

最盛期には500万人を超える登録ユーザーを確保し、アメリカの主要な通信会社すべてとパートナーシップを結びました。もちろんこれは後の話ですが、初期の投資家たちはそういう可能性を見ていたわけです。

転機が訪れました。XFundのパトリック・チョン（Patrick Chung）がニュー・エンタープライズ・アソシエイツ（New Enterprise Associates）から500万ドルを投資しました。当時の為替レートで50億ウォンを超える金額でした。19歳の大学中退者が率いる会社に50億ウォンを投資するというのは、本当に大きな決断でした。

パトリック・チョンはサムと初めて会ったときのことをこう振り返っています。「彼は年齢のわりにとても成熟していました。19歳が口にできる言葉ではなかった。」サムは投資家の前でアプリの機能を説明するだけにとどまらず、Looptが切り開く未来の姿を生き生きと描きました。人々がより頻繁に顔を合わせ、より気軽につながり、街がより温かみを帯びていく、そんな光景でした。

投資家たちが感動したのはアイデアだけではありませんでした。姿勢が違いました。若い起業家は投資家の前で緊張し、慎重になり、顔をうかがいます。サムは違いました。投資を「懇願する」のではなく「提案する」という態度でした。「私たちの会社に投資すれば、大きなリターンが得られます」という自信に満ちていました。

500万ドルの投資のニュースが広まると、他のベンチャーキャピタルも関心を示し始めました。セコイア・キャピタル（Sequoia Capital）のような有名な投資会社もLooptに投資しました。セコイ

アはグーグル、アップル、ユーチューブに初期投資したことで知られています。そこがLooptに投資したということは、シリコンバレーがサムビジョンを認めたということでした。

結局サムは、Looptのために3,000万ドルを超えるベンチャーキャピタルを調達しました。約300億ウォンを超える金額です。これは当時の学生起業家としては驚くべき成果でした。一つのアイデアで300億ウォンを集めたのです。もちろんその過程で何十回も断られ、徹夜で事業計画書を書き直し、投資家一人ひとりを説得する苦しい時間を過ごしましたが。

この経験はサムに重要な教訓を与えました。良いアイデアだけでは足りないこと、それを説得力を持って伝えられなければならないこと、断られても諦めない粘り強さが必要だということです。こうした力が、後に彼がOpenAIを率いるときの大きな財産になりました。

Yコンピネーターとの運命的な出会い

Looptの成功物語からYコンピネーター（YC）を外すことはできません。YCはポール・グレアム（Paul Graham）という伝説的なプログラマーが2005年に立ち上げたばかりのスタートアップ支援プログラムでした。普通の投資会社とは異なりまして。ただお金を渡すだけでなく、3ヶ月間集中的にスタートアップを支援し、アドバイスをし、他の起業家と引き合わせてくれる場所でした。

サムとLooptチームはYCの最初のバッチ（batch）に応募しました。YCは史上初めてスタートアップを募集しており、多くの志望者が殺到していました。その中には後にReddit（レディット）となるチームもありました。競争は激しかったです。

YCの選考プロセスは独特でした。長くて複雑な面接ではなく、わずか10分の短いインタビューでした。その短い時間の中で、自分のアイデアとチームをアピールしなければなりません。ポール・グレアムと他のパートナーたちは鋭い質問を次々と投げかけ、志望者は即座に答えなければなりません。

サムの番が来ました。19歳の青年はシリコンバレーの伝説の前に立ちました。ポール・グレアムはすでに何百人もの起業家と会ってきており、誰が成功するかを直感的に見抜く人物でした。彼の目を欺くことはできませんでした。

インタビューは順調に進みました。サムはLooptのアイデアを明確に説明し、技術的な質問にも詰まることなく答えました。ポール・グレアムも頷きながら関心を示しました。ところがインタビューの終わり頃、決定的な質問が飛んできました。

「もし私たちが投資することを決めたなら、学校はどうするつもりですか？」

これは単純な質問ではありませんでした。YCのプログラムは3ヶ月間、全力で取り組まなければなりません。授業の合間にできるようなものではありませんでした。投資を受けるには大学を辞めなければならないという意味でした。スタンフォードを手放すことは、普通の人にとっては大変な決断でした。

多くの学生起業家がこの質問の前で躊躇しました。「うーん、学校と並行する方法を探してみます」といった曖昧な答えをしがちでした。ポール・グレアムはそういった答えを好みませんでした。起業には100%の献身が必要だからです。

サム・アルトマンは一秒も躊躇しませんでした。

「私は2年生です。（今すぐ）辞めます。」

部屋が静かになりました。ポール・グレアムはサムをじっと見つめました。サムは自分の学年を正直に答え（「2年生です」）、そのまま即座に決断を示しました（「辞めます」）。言い訳もなく、躊躇もなく、条件もありませんでした。すっきりとした、迷いのない答えでした。

ポール・グレアムは後にこの瞬間を振り返って語りました。「その一言でサムがどんな人物かわかりました。本物の起業家とはこういうものです。言い訳をせず、道を切り開き、前に進む人です。」

サムの答えはYCの歴史で伝説になりました。後にYCに応募する多くの学生たちがこの話を聞くことになります。「サム・アルトマンのように答えよ」というアドバイスとともに。この一文は単なる答えではなく、サムの人生哲学を示す宣言でした。

サムは、クラスメートで当時のボーイフレンドだったニック・シボ（Nick Sivo）と共にLooptを共同創業しました。彼らは創業者プログラムから6,000ドルの支援金を受け、ケンブリッジにある初期の技術系起業家コミュニティで生活しながら製品を作り上げました。

YCはLooptへの投資を決めました。Looptは2005年にYコンピネーターに採択された最初の8社のひとつでした。これがサムとYCの長い縁の始まりでした。彼は後にYCの社長となり、シリコンバレーで最も影響力のある人物のひとりになります。そのすべてが、この10分間の面接から始まったのです。

「2年生ですが、行きます」この一言がサムの人生を変えました。人生の重要な瞬間は、時にこれほど短く訪れるものです。準備ができていない人だけが、その瞬間を逃さずつかめるのです。

「2年生ですが、今すぐ行きます」

YCの投資を受けたLooptは急速に成長しました。サムとチームは昼夜を問わずアプリ開発に没頭しました。最初はバグだらけだったアプリが徐々に安定し、機能もひとつずつ加わっていきました。友達の位置確認、メッセージ送信、近くのスポット提案……Looptはしだいに完成度の高いサービスへと生まれ変わっていきました。

初期の反応は熱狂的でした。大学生の間で大きな人気を集めました。「ねえ、今どこにいるの?」と聞かなくても、アプリを開けば友達の居場所がすぐわかったからです。パーティで友達を探したり、同じ図書館にいる友達を見つけたり、偶然近くにいた友達と合流したりすることができるようになりました。

スティーブ・ジョブズもLooptに目を向けました。2007年にiPhoneを発表した際、ジョブズはLooptをiPhoneで動くアプリのひとつとして紹介しました。Appleのステージに立つことは、すべてのアプリ開発者の夢でした。サムとチームは大きな誇りを感じました。

会社も大きくなりました。最初はサムと数人の友達だけでしたが、少しずつ採用を始めました。エンジニア、デザイナー、マーケター、営業担当者……チームは数十人規模に膨らみました。サムはもはやコードを書くだけのエンジニアではなく、会社を率いるCEOにならなければなりません

た。

ところが、何かがおかしかった。初期の爆発的な成長が鈍り始めたのです。ユーザー数は増えていましたが、使い続ける人は多くありませんでした。最初は面白そうでインストールしたものの、1～2回使っただけで削除する人が多かったのです。なぜだったのでしょうか。

バッテリーの問題がありました。GPSをつけっぱなしにすると、携帯のバッテリーがすぐに切れてしまいました。当時の携帯は今のようにバッテリーが持ちませんでした。

次に、プライバシーへの不安も大きかった。「自分の居場所を常に他の人に知られるのは、不便ではないか」と多くの人が感じていました。

さらに、スマートフォンがまだ普及していませんでした。2005年から2008年はiPhoneが登場したばかりの時期で、ほとんどの人はまだ従来の携帯電話を使っていました。

最大の問題は別のところにありました。FacebookやGoogleといった巨大企業が、同様の機能を自社サービスに加え始めたのです。Facebookはチェックイン機能を作り、GoogleはGoogleマップに位置情報の共有を組み込みました。Looptは小さなスタートアップでしたが、競合相手は数十億のユーザーを持つ巨人たちでした。

サムは悩みました。持ちこたえるのか、それとも別の道を探すのか。投資家たちも心配し始めました。「いつ頃、収益が出るんですか?」「ユーザーが減り続けているのは大丈夫ですか?」プレッシャーは増すばかりでした。

2012年、サムは苦しい決断を下しました。Looptをグリーンドット・コーポレーション (Green Dot Corporation) という銀行会社に4,340万ドル (約470億ウォン) で売却することです。7年間育ててきた会社を手放すのは容易な決断ではありませんでした。まるで子供を養子に出すような気持ちだったでしょう。

サムは後にIntelligencer誌のインタビューで「かなり不幸な状態で去った」と語っています。売却で得た資金は投資家と分配しなければならなかったため、サムが実際に受け取った金額は約500万ドル (50億ウォン) ほどでした。大金ではありましたが、7年間で捧げた対価としては物足りなさが残りました。

外から見れば、Looptは「失敗したスタートアップ」でした。FacebookやInstagramのような圧倒的な成功を収めることはできなかったからです。サム自身も率直に認めました。彼はTwitterにこう書いています。「私の最初のスタートアップはかなりひどく失敗しました。本当につらかったです!」

失敗したら終わりなののでしょうか。そうではありません。サムはLooptを通じて、お金では買えない貴重なものを学びました。

タイミングの重要性です。Looptのアイデアは優れていましたが、時代が早すぎました。サムは後に「技術には本当に可能性がありました。位置情報サービスはやがて、銀行からゲーム、ニュースに至るまで、ほぼすべてのモバイルアプリの重要な要素になりました。でもLooptはそこに根を張ることができなかった」と振り返っています。今では誰もが位置情報サービスを使っていますが、2005年当時、人々はまだその準備ができていなかったのです。

ユーザーが本当に求めるものを作らなければならないということです。技術的に優れているものと、人々が実際に使うものは別物です。Looptは技術的には申し分ありませんでしたが、人々の日常に欠かせないものではありませんでした。

失敗しても世界は終わらないという事実です。サムは「他の人はあなたの失敗をあなたほど気にしていない」と語っています。実際、Looptが行き詰まった後もサムの評判は悪くありませんでした。むしろ「7年間会社を経営した経験を持つ人物」として認められたのです。

人脈の価値です。Looptを運営するなかで、サムは数多くの投資家、創業者、エンジニアと出会いました。この人脈が後に大きな財産となりました。Looptの売却資金をもとに、サムは2012年に弟のジャック・アルトマン（Jack Altman）とともにハイドラジン・キャピタル（Hydrazine Capital）という投資会社を設立しました。投資を受ける立場から、投資する立場へと変わったのです。

Looptは失敗した会社でしたが、サム・アルトマンを育てた学校でもありました。この経験なくして、彼がYCの社長になることも、OpenAIを率いることもなかったでしょう。失敗は時に、成功よりも多くのことを教えてくれます。サムはそれを全身で学んだのです。

機転と粘り強さ

サムが最も苦しかった瞬間を挙げるとすれば、通信キャリアとの契約交渉でした。2005年から2010年にかけて、アプリを携帯にインストールするには通信キャリアの許可が必要でした。今はApp Storeから自由にダウンロードできますが、当時は違いました。SKテレコム、KT、LG U+のような通信会社が、携帯にどのアプリが入るかを管理していたのです。

サムはアメリカの大手通信キャリアを説得しなければなりませんでした。AT&T、Verizon、Sprint、T-Mobile……これらの会社の許可を得てはじめて、Looptを多くの人に届けることができました。問題は、これらの会社が小さなスタートアップに関心を持っていなかったことです。

サムが通信会社に連絡しても、返信すら来ませんでした。メールを送っても無視され、電話しても「担当者に伝えます」という言葉しか聞けませんでした。19歳の青年が数十年の歴史を持つ巨大企業と交渉しようとするのは、卵で岩を叩くようなものでした。

普通の人なら諦めていたでしょう。「ああ、通信会社は無理だな。別の方法を探そう」と。しかしサムは違いました。Looptが成功するには通信キャリアとの契約が不可欠で、彼は絶対にあきらめませんでした。

サムはポール・グレアムから学んだ言葉を思い出しました。「Be relentlessly resourceful（常に機転を利かせ続けよ）」。機転とは何でしょうか。問題を解決するためにあらゆる手を尽くす能力のことです。ひとつの方法がうまくいかなければ別の方法を試し、それもだめなら又別の方法を探す。「諦める」という言葉を辞書に持たない人こそ、機転の利く人なのです。

正面突破が通じないとわかると、サムは迂回路を探し始めました。AT&Tの担当者にメールを送って無視されると、その上司を探して連絡しました。それでも駄目なら別の部署の人間を当たりました。会社のウェブサイト調べて関係者の名前を探し、LinkedInで経歴を調べ、何とかつながりを作ろうとしました。

シリコンバレーで開かれる通信業界のカンファレンスには欠かさず参加しました。高い参加費を払って入場し、通信会社の幹部が座っているテーブルを探しました。昼食の時間に偶然を装って隣に座り、会話を始めることもありました。「こんにちは、Looptというスタートアップをやっているサムといいます……」

通信会社の本社に直接飛んで行ったこともありました。アポなしでビルのロビーで待ち、担当者が通りかかると名刺を差し出して5分だけ時間をくれと頼みました。警備員に追い出されたことも何度かありました。

ある通信会社と契約するために、サムは30通り以上の異なる方法でアプローチしたといいます。30回ですよ！ 普通の人は3回試ただけで「これは無理だ」と諦めます。サムは30回試みた後、31回目の準備をしていたのです。

アプローチの仕方も毎回変えました。技術部門の担当者にはLooptの革新的な技術を強調し、マーケティング部門の担当者には若い層へのアピール力を強調しました。経営陣には新しい収益モデルを提示しました。相手が誰かによって、話し方を変えていたのです。

ついに突破口が開きました。AT&Tの中間管理職のひとりが、サムのしつこさに根負けして(?)ミーティングを設定してくれたのです。「この人がうるさいから、一度だけ会ってみよう……」そうして始まったミーティングが、結局は契約につながりました。

ひとつの通信会社と契約を結ぶと、他の通信会社も関心を示し始めました。「AT&Tが契約したのに、うちだけ乗り遅れるのはまずいのでは？」そうしてドミノ倒しのよう契約が成立していきました。最終的にLooptはアメリカのすべての主要通信キャリアとパートナーシップを結ぶことに成功しました。

この経験は、サムに生涯忘れられない教訓を与えました。不可能に見えることでも、十分に長く、十分に多様な方法で試み続ければ解決できる、ということです。彼は後にこう語っています。

「人は最初の試みがうまくいかないと諦めます。なかには2回目まで挑戦する人もいます。でも本当に問題を解決する人は、30回目まで試みる人です。驚くことに、新しい方法を探し続けていると、たいていの問題はいつか解決されるものなんです。」

この「手腕」は、後にOpenAIを立ち上げる際にも輝きを放ちました。膨大なコンピューティングパワーが必要なのに資金が不足していたとき、サムはマイクロソフトを説得して数十億ドルの投資を引き出しました。政府の規制が懸念されたときは、自ら議会に赴いて証言しました。競合他社が先行したときは、より優れた製品を素早くリリースしました。

Looptの通信会社との契約交渉の話は、サム・アルトマンという人物を理解するうえで欠かせない核心です。彼は天才プログラマーでもなく、裕福な家庭の出身でもありません。彼を特別にしたのは、決して諦めない粘り強さと、あらゆる手を尽くす手腕でした。

人生において、不可能に見える壁にぶつかることがあるでしょう。受験、就職、事業、夢……。1、2回試してうまくいかないと、「自分には無理なんだ」と諦めてしまいがちです。サム・アルトマンの話を思い出してください。本当に成功する人は、30回目まで試みる人です。そして必要なら、31

回目も挑む準備ができている人です。

Looptは世界を変える会社にはなれませんでした。しかし、サム・アルトマンを世界を変える人物へと育てました。それがLooptの残した最大の遺産です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書斎 #サムアルトマン #OpenAI #ChatGPT #人工知能

1.3 Y Combinatorのリーダーシップ、スタートアップ界のキングメーカー

第1部 屋根裏のコンピューター少年からシリコンバレーのキングメーカーへ

← 前の章

次の章 →

Loopitを売却した後、サム・アルトマンの人生に新たな章が開いた。2011年、彼は若き起業家として初めて投資を受けた場所、Yコンビネーターへと戻ってきた。今度は投資を受ける側ではなく、投資をする側としてだ。そして2014年、わずか28歳でYコンビネーターの社長となった。この時期は、サム・アルトマンがシリコンバレーで最も影響力のある人物の一人へと成長する、決定的な時間だった。

数千のスタートアップを育てた男

Yコンビネーターという名前を聞いたことがあるだろうか。シリコンバレーで最も有名な「スタートアップの学校」だ。若い起業家たちがアイデアを持って訪れると、資金を提供し、アドバイスを与え、ビジネスを育てる手助けをしてくれる。Airbnb、Dropbox、Redditのように、今私たちが毎日使っているサービスのほとんどが、ここから生まれた。

2014年、サム・アルトマンは28歳でYコンビネーターの社長になった。誰が見ても若い年齢だった。しかしYコンビネーターの創業者であるポール・グレアムは、サムについてこう語っている。「彼を食人族が住む島に放り込んでも、5年後に戻ってみれば、その島の王になっているだろう」。恐ろしく面白い比喻だが、それほどサムの能力を信じていたということだ。

サムがYコンビネーターのトップになると、大胆な変革を試みた。まず、Yコンビネーターが毎年1,000社の新たな企業に投資することを目標に掲げた。当時としては想像しがたいほどの数だった。「そんなに多くの会社をどうやって管理するんだ」と首を振る人も多かった。しかしサムの考えは違った。

「世界には解決すべき問題が山ほどあります。気候変動、病気、エネルギー問題……こうした巨大な問題を解決するためには、もっと多くの起業家が必要です。私たちはもっと門戸を広げなければなりません」

サムはYコンビネーターの敷居を下げるため、さまざまなプログラムを作った。誰でも無料でオンラインで起業を学べる「Startup School」を開設し、開発者とスタートアップをつなぐインターンシッププログラムも始めた。世界のどこにいても、良いアイデアさえあればYコンビネーターの支援を受けられるようにしたのだ。

サムが着目したのは、数だけではなかった。どのような種類の会社に投資するかが、もっと重要だと考えていた。当時のシリコンバレーには、写真共有アプリやフードデリバリーアプリといった「手軽な」スタートアップが溢れていた。もちろんそうしたサービスも生活を便利にしてくれる。しかしサムは、もっと大きな絵を描いていた。

こう言った。「写真共有アプリなら何千もあります。でも、実験用の核融合炉を作れるスタートアップは何社あるでしょうか」。サムはYコンピネーターを、流行を追う会社を支援するだけの場所ではなく、人類の未来を変えうる「ハードテクノロジー」企業を育てる場所にしたいかった。

ハードテクノロジーとは何か。人工知能、バイオテクノロジー、核融合エネルギー、宇宙探査のように、開発に長い時間がかかり難しいが、成功すれば世界を根本から変えられる技術のことだ。こうした技術は、すぐに収益を上げるのが難しい。失敗するリスクも高い。だからこそ、大半の投資家はこうした企業への出資を敬遠する。

サムは違った。「人々が狂っていると言い始めることこそ、良いサインだ」と信じていた。他人が不可能だと言うことに挑む起業家を、サムは愛した。実際にYコンピネーターはサムのリーダーシップのもと、核融合エネルギーを研究するHelion Energyのような企業にも投資した。

2015年、サムはもう一つの重要な決断を下した。「YC Continuity」という名のファンドを設立したのだ。このファンドは7億ドル規模で、Yコンピネーターを卒業した企業のうち、すでにある程度成長した会社がさらに大きく育てられるよう追加資金を提供するものだった。

なぜこうしたファンドが必要だったのか。Yコンピネーターはもともと、ごく初期段階のスタートアップを支援する場所だった。起業したばかりの人々に資金を提供し、3か月間集中して指導し、投資家と引き合わせるのがYコンピネーターの役割だった。しかしその後は？ 会社が成長して資金がもっと必要になれば、別の投資家を探さなければならなかった。

サムはここに問題を見つけた。「私たちが育てた優れた会社が、次のステージに進めずに途中で止まってしまうことがある。最後まで一緒に歩まなければならない」。だからこそ作ったのがYC Continuityだった。若い木を植えて水をやり、肥料を与えながら途中で世話をやめるのではなく、その木が巨大な森になるまで育て続けるという意志の表れだった。

サムがYコンピネーターを率いた5年間で、そこを経たスタートアップの数は約1,900社に達した。その中には、DoorDash、Instacart、Twitch、Coinbaseのように、今や数十億ドルの企業価値を誇る会社も含まれている。サムは文字通り「キングメーカー」になった。彼が選んだ企業、彼がアドバイスした起業家たちが、シリコンバレーの新たな王者になっていったのだ。

サムが目指したのは、成功した会社をたくさん生み出すことだけではなかった。起業家たちにいつもこう言っていた。「お金を稼ぐことも大切ですが、世界をより良い場所にすることはそれ以上に大切です」。彼はYコンピネーターを、ただの投資機関ではなく、世界を変えようという夢を持つ人々が集う「コミュニティ」にしたかった。

こうしたサムの哲学は、Yコンピネーター全体の文化に染み渡った。Yコンピネーターは起業家たちにこう教えた。「速く動いてください。失敗を恐れなくてください。顧客の声を聞いてください。でも何より、本当に重要な問題を解決する会社を作ってください」

面白いエピソードがある。あるインタビューで、「Yコンピネーターを上場させるつもりはあるか」という質問を受けた。サムは「ありません」と断言した。その理由を問われると、こう答えた。

「上場すれば、どこかのヘッジファンドマネージャーが電話をかけてきて、『今四半期の利益が1セント足りない！Yコンビネーターは終わりだ！』と怒鳴るでしょう。私はそんな生活を望まない。私たちは四半期ごとの業績ではなく、10年後、20年後の未来を見て仕事をすべきです」

この答えには、サム・アルトマンの哲学がそのまま映し出されている。短期的な成果より長期的な価値を信じた。速い成長より持続可能な成長を求めた。そして何より、本当に世界を変えられる会社を作りたかった。

サム・アルトマンがYコンビネーターの社長を務めた時期は、彼の人生で最も重要な局面の一つだった。この時期を通じて、数多くの起業家と出会い、成功する会社とは何か、どんな技術が本当に未来を変えられるかを学んだ。そしてその経験のすべてが、後にOpenAIを率いる際の大きな財産となった。

2019年、サムはYコンビネーターの社長を退いた。しかし彼が去った後も、Yコンビネーターは世界最高のスタートアップ・アクセラレーターとして位置を守り続けている。サムが作り上げたシステムと文化が今も息づいているからだ。「キングメーカー」として彼が蒔いた種は、今もなお育ち続けている。

師、ポール・グレアム——人を踏みにじらずに導く法

すべての偉大な人物には、師がいる。サッカー選手のソン・フンミンにはベテランコーチがいたし、『ハリー・ポッター』シリーズを書いたJ.K.ローリングには彼女を支えてくれた編集者がいた。サム・アルトマンにとって、ポール・グレアムがまさにその存在だった。

ポール・グレアムとはどういう人物か。Yコンビネーターを創設した人物であり、シリコンバレーで最も尊敬される思想家の一人だ。プログラマーであり、作家であり、投資家でもある彼は、数多くのエッセイを通じて起業とテクノロジーについての深い洞察を発信してきた。シリコンバレーの多くの起業家が、彼の文章をバイブルのように読む。

2005年、19歳のサム・アルトマンがLooptを携えてYコンビネーターに応募したとき、ポール・グレアムはこの若き起業家に何か特別なものを感じ取った。単に頭が良いとか情熱があるとか、それを越えた何かだ。彼はサムについてこう書いている。「彼を食人族が住む島に放り込んで、5年後に戻ってみれば、その島の王になっているだろう」

恐ろしい比喻だが、ポール・グレアムが伝えたかったメッセージは明確だった。サム・アルトマンはどんな困難な状況に置かれても生き残るだけでなく、その状況を完全にひっくり返せる人物だということだ。そしてポール・グレアムの予測は的中した。

Looptの売却後、サムは2011年にYコンビネーターへと戻ってきた。今度は投資を受ける起業家ではなく、他の起業家を支援するパートナーとしてだ。そして2014年、ポール・グレアムは驚くべき決断を下した。わずか28歳のサム・アルトマンを後継者として指名し、Yコンビネーターの社長職を譲ったのだ。

「サムは私が知る中で最も賢い人物の一人です。そして、私も含めて私の知る誰よりもスタートアップをよく理解しています」。ポール・グレアムのこの言葉には、サムへの深い信頼が込められてい

た。

ポール・グレアムがサムに教えたのは、ビジネスのノウハウだけではなく、彼がサムに伝えた最も大切な教えは、「人とどう向き合うか」というものだった。

シリコンバレーは競争の激しい場所だ。成功するためには時に冷酷にならなければならない、素早く決断を下し、弱みを見せてはいけない。多くのCEOが社員を道具のように扱い、目標を達成するためなら誰かを犠牲にすることもいとわない。

ポール・グレアムは違った。そして同じことをサムにも教えた。「人を踏みにじらずに導く」という哲学だ。

どういう意味か。簡単に言えば、リーダーとは権力で人を抑えつける者ではなく、その人の中にある最高の力を引き出す者でなければならないということだ。社員を恐れで支配するのではなく、敬意と信頼で導くべきだということだ。

ポール・グレアムはYコンビネーターを運営する中でこの哲学を実践した。起業家たちに「あなたのアイデアはひどい」と率直に言うこともあったが、彼らを見下したり貶めたりはしなかった。代わりに「なぜひどいのか、どうすれば良くなるか」を一緒に考えた。失敗を恐れずに挑戦し続けられるよう、起業家たちに勇気を与えた。

サム・アルトマンはこうしたポール・グレアムの姿を見ながら多くを学んだ。Yコンビネーターのパートナーとして、そして後に社長として働く中で、ポール・グレアムのやり方を自分のものとして吸収していった。

サムが起業家たちをメンタリングする場面を想像してみましょう。彼は起業家たちに一方的に「こうしなさい、ああしなさい」と命令することはありませんでした。

代わりに、問いかけをしました。

「あなたの顧客は誰ですか？」「彼らが本当に求めているものは何でしょう？」「なぜあなたのプロダクトが必要なのですか？」「もしこの方法がうまくいかなければ、ほかにどんな方法がありますか？」

こうした問いかけを通じて、起業家たちは自分自身で答えを見つけていきました。そして自分で導き出した答えは、誰かに言われてやることよりもはるかに大きな力を発揮しました。

ポール・グレアムはサムにもうひとつ重要なアドバイスをしました。「Relentlessly Resourceful (しぶとく、あらゆる手を尽くせ)」。この言葉は、サムが生涯忘れられない教えとなりました。

どういう意味でしょうか。問題が起きたときに安易に諦めず、あらゆる可能性を試してみることです。一つの方法がうまくいかなければ別の方法を試し、それもダメなら、さらに別の方法を探すということです。

サムはこのアドバイスを19歳のときに自ら体験しました。Looptを運営していたころ、彼らは通信キャリアと契約を結ぶ必要がありました。しかし通信キャリアは、小さなスタートアップとは仕事

をしませんでした。普通の起業家なら、何度か断られた時点で諦めていたでしょう。

しかしサムは諦めませんでした。彼は30通り以上の異なる方法でその通信キャリアにアプローチしました。メールを送り、電話をかけ、知人を通じて紹介を取り付け、イベントで偶然を装って出会いを演出しました。ついに通信キャリアの主要な意思決定者がこう言いました。「あなたに付きまといられるのを終わらせるために、会ってあげます。」

冗談のように聞こえますが、実際に起きた話です。そしてその面会を通じて契約を成立させました。これこそが「Relentlessly Resourceful」の力です。

ポール・グレアムとサム・アルトマンの関係は、上司と部下のものではありませんでした。互いを深く尊重するメンターと弟子であり、やがて友人になりました。ポール・グレアムはサムにビジネスの技術だけでなく、生きることの知恵も伝え、サムはその教えを受け取って自分なりのスタイルへと発展させていきました。

関係が常に完璧であるとは限りません。2019年にサムがY Combinatorを去ったとき、二人の間に多少の緊張があったという話もあります。サムがY Combinatorの仕事よりもOpenAIに多くの時間を割くようになったことから生じた摩擦でした。しかしそれは、サムがより大きな夢へと踏み出す過程でもありました。

重要なのは、ポール・グレアムがサムに与えた教えが、サムの生涯にわたって影響を与え続けたという事実です。「人を踏みにじらずに導く方法」という哲学は、後にOpenAIで光を放つことになります。

2023年、サム・アルトマンがOpenAIの取締役会によってCEOを解任された事件を覚えていますか？そのとき、驚くべきことが起きました。OpenAIの従業員770名のうちほぼ全員が連名の書簡に署名し、「サムが戻らなければ、私たちも全員会社を去る」と宣言したのです。

この劇的な忠誠心はどこから生まれたのでしょうか。それは、サムが彼らに接する姿勢から来ていました。彼は従業員を踏みにじったり、抑圧したりしませんでした。代わりに彼らを尊重し、意見に耳を傾け、最高の能力を発揮できるよう支えました。だからこそ、危機の瞬間に従業員たちが彼を守るために立ち上がったのです。

これが、ポール・グレアムがサム・アルトマンに教えたリーダーシップの力です。権力で人を支配するのではなく、信頼で人を導くこと。恐怖ではなく、敬意で組織をつくること。これがリーダーシップです。

サム・アルトマンは、ポール・グレアムという優れたメンターに出会えた幸運な人物でした。彼はその教えをしっかりと受け取り、自分なりの形へと発展させました。ポール・グレアムはサムに種を植え、サムはその種を育てて大きな木へと成長させたのです。

サム・アルトマンが人工知能時代の最も重要なリーダーの一人となれたのは、彼が賢く野心的だったからだけではありません。ポール・グレアムのように、彼を信じ、正しい道を教えてくれるメンターがいたからです。そしてサムが、その教えを謙虚に受け入れ、実践したからです。

皆さんもいつか、誰かのメンターに出会うことがあるでしょう。そのときに大切なのは、その教えをどれほど真剣に受け止め、自分のものにできるかということです。サム・アルトマンとポール・グレアムの物語は、良いメンターシップが一人の人生を、そして世界をどのように変えうるかを示す、美しい例です。

流行していた技術の限界について考える。

2010年代前半から中盤にかけて、シリコンバレーはまさに「アプリの黄金期」でした。スマートフォンが世界中に普及するにつれ、人々の日常を変える無数のアプリが次々と生み出されました。

Instagramのような写真共有アプリは、人々のコミュニケーションの形を根本から変えました。友人と写真を共有し、「いいね」を押し、コメントを書きながら一日を過ごす。Uber Eats、DoorDashのようなフードデリバリーアプリは、家から出なくても好きな食事を楽しめるようにしました。UberとLyftはタクシーの呼び方を革新し、Airbnbは旅の仕方を変えました。

こうしたアプリは爆発的な人気を集めました。数百万、数千万人の人々が使い、投資家たちはこれらの企業に数十億ドルを注ぎ込みました。Y Combinatorに応募するスタートアップの多くも、こうしたアプリのアイデアを持ち込んできました。

「写真をもっとかわいく加工できるアプリです！」「食事をもっと速く届けるアプリです！」「友達ともっと気軽に会えるアプリです！」

表面上は完璧な時代でした。若い起業家たちはアイデア一つで大金を稼ぐことができ、投資家たちは次のユニコーン企業（企業価値10億ドル以上の会社）を探すことに躍起になっていました。

しかしサム・アルトマンは、この熱狂の中で何か違和感を覚えていました。

Y Combinatorの社長として、一日に何十ものスタートアップのアイデアを審査しました。そして次第に、アイデアの多くが互いに似通っていることに気づいていきました。写真共有アプリだけで数千はありました。デリバリーアプリも数百に上りました。それぞれのアプリは少しずつ異なる機能売りをしていましたが、本質的にはどれも同じことをしていました。

サムはこの状況を見ながら、深く考え込みました。「本当にこれで全てなのか？人類が解決すべき最も重要な問題が、写真をきれいに加工することや、食事を5分早く届けることなのだろうか？」

もちろん、こうしたサービスにも価値があります。人々の生活を便利にし、雇用を生み出し、経済を活性化させます。サムもそれを否定しませんでした。実際にY CombinatorはDoorDash、Instacartといった成功した配送サービス会社に投資し、大きな成果を上げました。

しかしサムの視線は、もっと遠くへ向いていました。

2011年にはIBMのWatsonがクイズ番組『Jeopardy!』で人間チャンピオンを破るのを目の当たりにしました。2016年にはGoogle DeepMindのAlphaGoが囲碁の世界チャンピオン、イ・セドルを下す歴史的な瞬間も見届けました。これらの出来事を見ながら、サムは感じていました。「何か大きな変化が来ている。」

彼は共有アプリやデリバリーサービスの「限界」を明確に認識していました。その限界とは何だったのでしょうか？

一つ目の限界は「影響力の大きさ」でした。写真共有アプリがどれほど成功しても、それは人類の根本的な問題を解決しません。気候変動を止めることも、病気を治すことも、エネルギー問題を解決することもできません。もちろん、人々を楽しませ、つなぐことも大切です。しかしサムは、それだけでは不十分だと思っていました。

二つ目の限界は「イノベーションの深さ」でした。ほとんどのアプリサービスは、既存のものを少し便利にする程度のものでした。たとえば、フードデリバリーは以前から存在していました。デリバリーアプリは、そのプロセスをデジタル化してより効率的にただけです。もちろんそれも革新ですが、サムが求めている種類の革新ではありませんでした。

サムが関心を持っていたのは、「ゼロから1を生み出す（Zero to One）」イノベーションでした。世界に全く新しいものを生み出すイノベーションです。著名な投資家ピーター・ティールが書いた本のタイトルでもあるこの概念は、既存のものを改善すること（1からnへ）ではなく、これまで存在しなかった新たな価値を生み出すこと（0から1へ）を意味します。

三つ目の限界は「競争の飽和状態」でした。共有アプリが数千もあるということは、その市場がすでに飽和していることを意味します。新規参入するスタートアップが成功するのは非常に難しく、大半は失敗し、生き残るのはごく一部だけです。サムは、こうした「レッドオーシャン」で戦うよりも、まだ誰も踏み込んでいない「ブルーオーシャン」を切り開くことのほうが、より意義があると考えていました。

サムはしだいに「ハード・テクノロジー（Hard Technology）」に強い関心を抱くようになっていきました。ハードテクとは何でしょうか。アプリのように数ヶ月で作れる「ソフト」な技術ではなく、数年あるいは数十年の研究と開発を要する「難しい」技術のことです。

たとえば、

こうした技術は開発が非常に困難です。資金が必要で、失敗のリスクも高い。すぐに収益を上げることも難しい。ほとんどの投資家はこの分野を敬遠します。「リスクが高すぎる。確実に稼げるアプリに投資した方がいい」というわけです。

サムは違いました。彼はこう言いました。「写真共有アプリなら数千もあります。でも、実験用の核融合炉を作るスタートアップが何社あるでしょうか？」

サムは、Y Combinatorがこうした難しい問題に挑むスタートアップを支援すべきだと信じていました。実際、Y Combinatorはサムのリーダーシップのもと、ヘリオン・エナジー（Helion Energy）という核融合スタートアップに投資しました。核融合とは、太陽がエネルギーを生み出す仕組みを地球上で再現しようとする技術です。成功すれば、ほぼ無限のクリーンエネルギーが得られます。しかし、何十年もの間、多くの科学者たちが挑戦し続けてきたにもかかわらず、いまだ実現できていない、非常に難しい技術です。

多くの人がサムに尋ねました。「なぜそんな不確かなものに投資するのですか？確実に稼げるアプリの方がずっと安全ではないですか？」

サムはこう答えました。「人々が『狂っている』と言ったり、『絶対に無理だ』と言ったりするような新しいことを始めるとき、それはいつも良いサインです。」

これがサム・アルトマンの哲学でした。誰もが簡単だと思う道ではなく、難しくても意義のある道を歩むこと。短期的な成功よりも長期的な影響力を追い求めること。流行を追うのではなく、未来を作ること。

2015年、サムはより大きな決断を下しました。イーロン・マスク、グレッグ・ブロックマンらとともにOpenAIという会社を立ち上げたのです。OpenAIの目標は何だったのでしょうか。それはAGI、すなわち人間レベルの汎用人工知能を作ることでした。

当時としては、ほとんどSF小説のような話でした。多くの専門家でさえ「AGIにはまだ数十年かかる」「不可能かもしれない」と懐疑的でした。しかしサムは、これこそが人類の未来を変えられる本物のイノベーションだと確信していたのです。

共有アプリやデリバリーサービスの限界に気づいたのは、彼がより大きな絵を見ていたからです。10年後、20年後、50年後の世界を想像し、その世界で人類が直面するであろう問題を考えていたのです。

気候変動による環境災害、エネルギー不足、食料問題、病気と老化、宇宙への進出……こうした巨大な問題を解決するには、単なるアプリではなく、根本的な技術革新が必要でした。そしてその革新の中心に人工知能があるだろうと、サムは確信していました。

サムのこうした考えが、最初から誰にでも歓迎されたわけではありませんでした。Y Combinator内部でも「アプリ企業に投資して稼げばいいのに、なぜリスクの高いハードテックに投資しようとするのか」という反発がありました。投資家の一部からも「理想主義に過ぎる」「現実味がない」という批判が上がりました。

時が経つにつれ、サムが正しかったことが証明され始めました。2022年末、OpenAIが作ったChatGPTが世に公開されると、世界中が衝撃を受けました。わずか5日間で100万人を超える人々が登録し、今では数億人が使っています。

ChatGPTは単なるアプリではありませんでした。人間と会話し、質問に答え、文章を書き、コードを書き、問題を解決できる人工知能でした。人々は気づきました。「これが本物のイノベーションだ。これが世界を変える技術だ」と。

振り返れば、サム・アルトマンが2010年代半ばに写真共有アプリとデリバリーサービスの限界を見抜いていたことは、驚くべき洞察力でした。みんながアプリに熱狂していたとき、彼はその先を見ていたのです。短期的な流行ではなく、長期的な変化を予測していたのです。

これこそが、真のリーダーと普通の人との違いです。普通の人は今流行っているものを追いかけます。しかし真のリーダーは10年後、20年後の未来を見据え、その未来を作るために今何をすべきかを知っています。

サム・アルトマンは写真共有アプリの時代を超えて、人工知能の時代を見ていました。そして彼はその未来を作るために、Y Combinatorを去りOpenAIにすべてをかける決断をしました。2019年、彼がY Combinator代表を辞任してOpenAIのCEOに就任したのは偶然ではありませんでした。それは、彼がずっと前から描いてきた大きな絵の一部だったのです。

忙しい幹部は、有能な幹部なのか？

朝早く起きて会社に行き、会議をこなし、メールを確認し、電話をかけ続けて夜遅く帰宅する。でも、すべての「忙しさ」が本当に意味のあるものなのでしょうか？サムはY Combinatorを率いながら、この問いについて深く考え続けました。

シリコンバレーには古くからの格言があります。「有能な幹部は忙しい幹部だ」というものです。どういう意味でしょうか。会社のリーダー、とりわけ幹部は、社員に仕事を絶え間なく与え続けなければならない、ということです。社員が時間を無駄にしてぶらぶらしないよう、常に忙しくさせなければならない、という意味です。

サムもそれに同意しました。ただ、彼が考える「忙しさ」の意味は少し違いました。単にたくさんの仕事をこなすだけでは生産的とは言えないと考えていたのです。大切なのは「正しい仕事」に忙しいことだと信じていました。

サムはY Combinatorを運営するなかで、多くのスタートアップが失敗する姿を見てきました。その多くは、怠けていたり、仕事をしなかったりして失敗したわけではありませんでした。むしろ、猛烈に働いていました。一日15時間、週末も返上し、徹夜で働き続けていました。それでも失敗したのです。なぜでしょうか？

答えは単純でした。彼らは「間違った仕事」に忙しかったのです。

サムはこう言いました。「どれだけ速く動いているかは重要ではありません。役に立たない方向に動いているなら意味がないからです。何をするかを選ぶことこそが、生産性においてもっとも重要な要素です。ところが多くの人は、これをほぼ無視しています。」

これがサムの生産性哲学の核心でした。忙しいこと自体が目標なのではなく、正しい問題を解決することに忙しいことが目標だ、ということです。

どうすれば正しい仕事を見つけられるのでしょうか。サムはいくつかの方法を提案しました。

一つ目は、考える時間を確保することです。逆説的に聞こえるかもしれませんが、生産的であるためには「何もしない」時間が必要です。サムは「私はスケジュールの中に、何をすべきかを考えるための十分な時間を意図的に残しています」と語っています。

本を読み、面白い人々と話し、自然の中を歩きながら考えを整理しました。そうした時間に彼は、「自分が今やっていることは本当に重要なことか。もっと重要な別のことがあるのではないか」と自問していたのです。

二つ目は、リストを作ることです。サムの生産性システムは三つの柱で成り立っていました。

サムは毎年、毎月、毎日やりたいことのリストを作っていました。紙に手書きするのが好み、頻繁にリストを書き直しました。なぜでしょうか。リストを書き直す過程で、本当に重要なこととそうでないことを自然に区別できたからです。

彼はまた、「勢い (momentum) を生み出すように優先順位をつける」とも語っています。どういう意味でしょうか。一つを完了すると気分が上がり、そうするともっと多くのことができるようになる、ということです。だから彼は、一日を始めるときも終わるときも、本当に前進を実感できる仕事で締めくくろうと心がけていました。

三つ目は、「ノー」と言うことを学ぶことです。これがおそらくもっとも難しい部分でしょう。サムは「私は本質的でない仕事を最も速い方法で片付けるか、断ることに對して、容赦なく厳しく対処しています」と語っています。

多くの人は「ノー」と言うことを難しいと感じます。頼まれると断るのが申し訳なく思える。会議に招かれたら出席するのが礼儀だと考える。しかしサムは違いました。自分の時間を守ることが何より重要だと信じていたのです。

彼はこう言いました。「時給100ドルを稼ぐ人たちが、20ドルを節約するために何時間も費やしているのを見て、驚きました。」自分の時間を十分に価値あるものだと思っていない、ということです。

第四に、一日のうち異なる時間を異なる種類の仕事に充てることです。サムは自分の日常を戦略的に設計しました。彼は「朝の最初の数時間は、間違いなく一日の中で最も生産的な時間だ」と言いました。そのため、彼はその時間帯には誰も予定を入れられないようにしました。その代わり、午後に会議をまとめて行いました。

多くのCEOは一日中会議で埋まったスケジュールを抱えています。午前9時から午後6時まで、30分または1時間単位で会議がびっしりと詰まっています。サムはこのようなやり方は非生産的だと考えました。本当に深く考え、重要な問題を解決するには、邪魔されない集中した時間が必要だと信じていたのです。

彼はまた、「集中力が途切れ始めたら、休息を取るか別の仕事に切り替える」とも述べています。無理に仕事を続けるよりも、少し休んでから再開する方が効率的であることを知っていたのです。

第五に、好きなことをすることです。サムは「自分に感心がないことや好きではないことをしている時は、生産的になれないということを学びました」と率直に告白しました。そのため、彼はそのような仕事をしなくても済む立場に自分を置こうと努めました。他人に任せたり、避けたり、あるいは別の方法を探したりしたのです。

彼はまた、他人に仕事を委任する際にもこの原則を適用しました。「誰もが自分の好きなことをしている時に最も生産的になります。誰が何を好み、何が得意かを把握し、それに合わせて仕事を分担してください」。

これは重要な洞察です。多くの管理職は単に業務を分担することだけに気を取られますが、サムは各個人が何を好み、何が得意かを把握しようと努めました。ある人はプレゼンテーション資料を作

るのが好きで、ある人はデータ分析が好き、またある人は顧客と話すのが好きといった具合です。こうした個人の特性を把握して仕事を分配すれば、誰もがより幸福になり、より生産的になります。

第六に、良い人々と共に働くことです。サムは「賢く、生産的で、幸福で、肯定的であり、あなたの野心を過小評価しない人々のそばにるように努めてください」とアドバイスしました。

彼は自分を鼓舞し、より良くなるようインスピレーションを与えてくれる人々のそばにいることを好みました。反対に、否定的で批判的な人々と一緒にいると、彼らがあなたの精神的エネルギーを消耗させるコストは甚大であると警告しました。

最後に、サムは「生産性ポルノ」に陥らないよう警告しました。どういう意味でしょうか？ 多くの人が、完璧な生産性システムを探すことにあまりにも多くの時間を費やしています。どのアプリを使うか、どのメソッドに従うか、どうすれば最後の1秒まで最適化できるかを悩みすぎて、肝心な仕事ができているのです。

サムはこう語りました。「生産性そのもののために生産性を追い求めることは役に立ちません。多くの人が自分のシステムを完璧に最適化する方法を考えることに時間を使いすぎ、正しい問題を解決しているかどうかを自問することにはほとんど時間を使いません」。

これこそが核心です。どんなシステムを使おうと、最後の1秒まで絞り出そうと、もし間違った仕事をしているのであれば、何の意味もありません。

OpenAIにおいて、サムはこの哲学を実践に移しました。彼は「比較的少数の人々に多大な責任を任せるよう努めている」と語っています。少人数のチームですが、各メンバーが本当に重要なことに集中し、迅速に動けるようにしました。

これこそが、OpenAIがGoogleやMetaのような巨大テック企業を抑えて、ChatGPTをいち早くリリースできた秘訣の一つでした。大企業には数千人の従業員がいましたが、複雑な組織構造と遅い意思決定プロセスのために動きが鈍くなっていました。一方、OpenAIは小規模ながらも非常に集中したチームで、素早く実験し、失敗し、学び、再び挑戦することができました。

「有能な幹部は忙しい幹部である」というサム・アルトマンの信念は、単に多くの仕事をこなせという意味ではありませんでした。それは、

という意味でした。

学校で勉強する際、単に何時間勉強したかが重要なのではなく、本当に理解して学んだかどうか重要です。色々な科目に時間を少しずつ割り振るよりも、最も重要だったり難しかったりする科目に集中する方が効果的な場合があります。

サム・アルトマンは忙しく日々を過ごす人でしたが、無意味に忙しかったわけではありません。彼は常に「自分が今していることは本当に重要なことか？」と自問し、そうでなければ潔くその仕事を中断するか、誰か他の人に任せました。そのため、彼は本当に重要なことだけに集中することができ、その結果、ChatGPTのような世界を変える製品を生み出すことができたのです。

APIの力：小さなインターフェース、大きな影響力

サム・アルトマンがYコンビネーターで学んだ教訓の中で、最も重要なものを一つ挙げるとすればこれです。「APIを作れば、大抵はどうかして良い結果につながる」。この一見シンプルに見える一文が、後にOpenAIの運命を決定づけました。しかしその前に、まずAPIとは何かについて理解してみましよう。

APIは「Application Programming Interface」の略称です。日本語にすれば「アプリケーション・プログラミング・インターフェース」となりますが、名前を聞いただけではどういう意味がよく分かりませんよね？

簡単に説明すると、APIは異なるプログラム同士が対話できるようにしてくれる「橋」のようなものです。

レストランに例えてみましょう。皆さんがレストランに行くと、メニューを見て料理を注文します。すると厨房で料理人が料理を作り、提供してくれます。皆さんは厨房がどうなっているか、料理人がどんな道具を使っているかを知る必要はありません。メニューから欲しいものを選ぶだけでいいのです。

「OpenAIは最も核心的な『AIの脳のブロック』を提供し、開発者たちはこのブロックに各自のアイデアのブロックを加えて無限の創作物を作ります。これこそがAPIエコシステムの力です」。例えば、Googleマップは自らの地図機能をAPIとして提供しています。そのため、デリバリーアプリを作る会社は自ら地図を作る必要はなく、GoogleマップAPIを使用して配達員の現在地を表示することができます。

サム・アルトマンはYコンビネーターで数多くの成功事例を見守りながら、一つのパターンを発見しました。APIを提供している会社がとりわけ急速に成長し、予想外の成功を収めているということでした。

代表的な例がストライプ（Stripe）です。ストライプはオンライン決済を処理する会社です。かつては、ウェブサイトでクレジットカード決済を受け付けるには複雑な過程を経る必要がありました。銀行と契約し、セキュリティ認証を受け、複雑なコードを作成しなければならませんでした。スタートアップにとっては、ほぼ不可能なことでした。

ストライプはこれらすべてをAPIで簡素化しました。開発者は数行のコードを追加するだけで、自分のウェブサイトに決済機能を組み込めるようになったのです。まるでレゴブロックをはめ込むかのようです。その結果、数多くのスタートアップがストライプを使い始め、ストライプは今や数百億ドルの価値を認められる巨大企業となりました。

トゥイリオ（Twilio）も同様の事例です。トゥイリオはショートメッセージや通話機能をAPIで提供しています。例えば、皆さんがデリバリーアプリを作っていて、顧客に「出前がもうすぐ到着します」というメッセージを送りたい場合、トゥイリオのAPIを使えばいいのです。通信会社と直接契約する必要もありません。

サムはこれらの事例を見ながら、APIの真の力は「予期せぬ革新」を可能にすることだと悟りました。

どういう意味でしょうか？

企業がAPIを作る際、通常は自分たちが想定した用途で人々が使うだろうと予想します。しかし実際には、全く異なる方法で使われることがよくあります。そして、そのような「予期せぬ使い方」から、時には最大の価値が生まれるのです。

アマゾン ウェブ サービス (AWS) を考えてみましょう。当初、アマゾンは自社のサーバーインフラをより効率的に使用するために、内部的にAPIを作成しました。ところが、これを外部に公開したところ、世界中の数多くのスタートアップが利用し始めました。今ではAWSはアマゾンの最も収益性の高い事業部門の一つになっています。アマゾン自身も、当初はこれほどのことになるとは思っていなかったはずです。

ポール・グレアムもこの点を強調しました。彼はサムに「何があってもAPIを作れ。そうすれば良いことが起きるだろう」と助言しました。

この助言は後にOpenAIで決定的な役割を果たしました。

2020年、OpenAIはGPT-3という強力な言語モデルを開発しました。これは人間のように文章を書き、質問に答え、さらにはコードを書くことさえできる驚くべきAIでした。しかし、問題がありました。「これで何を作ればいいのか？」

GPT-3は確かに素晴らしい技術でしたが、それ自体は一般の人々が使える「製品」ではありませんでした。まるで強力なエンジンはあるけれど、それを載せる車をどう作ればいいのか分からないという状況でした。

OpenAIのチーム内では多くの議論が交わされていました。「GPT-3でメール作成ツールを作ろう」という意見もあれば、「学生向けの宿題支援ツールを作ろう」という声もありました。しかし、誰も確信を持てずにいました。

そのとき、サム・アルトマンはポール・グレアムのアドバイスとYCでの経験を思い起こしました。「APIを作れば、良いことが起きる。」

彼はこう言いました。「何を作るべきか正確にわからないなら、まずAPIとして公開しよう。そうすれば、私たちが思いもよらないものを他の人たちが作り出すかもしれない。」

2020年6月、OpenAIはGPT-3 APIを公開しました。厳選された開発者たちにアクセス権を付与し、この強力なAIを使って何が生み出されるかを見守りました。

驚きでした。

ある開発者はGPT-3を使ってマーケティングコピーを自動作成するツールを作りました。ある開発者はカスタマーサービス用のチャットボットを作りました。ある開発者はコード作成を支援するツールを作りました。何百ものさまざまなユースケースが生まれました。

OpenAIのチームがとりわけ興味深く注目したのは別のことでした。それは「GPT-3 プレイグラウンド」で一日中AIと会話を続ける人々の存在でした。

プレイグラウンドは、開発者がAPIをテストできるよう作られたシンプルなツールでした。特定の目的はなく、GPT-3がどのように動くかを気軽に試せる場所だったのです。ところが驚いたことに、多くの人がここで何時間も過ごし、GPT-3と「会話」を楽しんでいました。

彼らはGPT-3に質問をし、話を聞かせ、アドバイスを求めました。まるで友人と話すかのように。ごく少数のユーザーではありましたが、彼らの行動は明確なシグナルを発していました。人々はAIと会話したいのだ、と。

OpenAIのチームは気づきました。「もしかしたら、私たちが作るべき製品はこれかもしれない。複雑なツールではなく、ただAIと会話できるシンプルなチャット画面だ。」

この洞察からChatGPTが生まれました。2022年11月30日、OpenAIはChatGPTを公開しました。華やかな機能も複雑な設定もない、ただAIと会話できるシンプルなチャットインターフェースでした。

結果はどうだったか。5日で100万人を超えるユーザーが登録しました。2か月で1億人を突破しました。史上最速で成長したコンシューマー向けアプリケーションとなりました。

もしOpenAIが最初から「完璧な製品」を作ろうとしていたらどうだったでしょうか。何年もかけて計画し、議論し、開発を続けるうちに、結局市場のタイミングを逃していたかもしれません。しかし彼らは先にAPIを公開し、ユーザーが本当に求めているものを発見することができました。

これこそが、「APIは良い結果をもたらす」というサム・アルトマンの洞察が輝いた瞬間でした。

APIの力は三つにまとめることができます。

APIは「集合知」を引き出してくれます。一つの会社の優秀な10人が考えられることより、世界中の数万人の開発者が考え出せることの方がはるかに多いのです。APIを公開するということは、「私たちの技術を使って、あなたが作りたいものを作ってみてください」と招待するようなものです。

APIは「市場検証」を素早く行えるようにしてくれます。複数の製品アイデアのうち何が成功するか、事前に知ることはできません。しかしAPIを公開すれば、人々が実際にどのような用途で使うかをすぐに確認できます。最も多く使われる用途こそ、市場が求めているものです。

APIは「エコシステム」を生み出します。一つの会社がすべてを一人で作る必要はありません。コア技術をAPIとして提供すれば、無数の他の会社がその上に自分たちの製品を作ります。その結果、エコシステム全体が成長し、もともとAPIを提供した会社も共に成長します。

サム・アルトマンが観察したこのパターンは、ビジネス戦略の域を超えるものでした。それは、イノベーションがいかにして生まれるかについての深い理解でした。イノベーションは一人の天才的なアイデアからだけ生まれるわけではありません。ときにはツールを作って他の人々に渡し、彼らが何を生み出すかを見守ることから、最大のイノベーションが生まれます。

今日ChatGPTを使う数億人の人々、そしてChatGPTを自社製品に組み込んだ数万社の企業は、おそらく知らないでしょう。自分たちが使うこの驚くべき技術の背後に、サム・アルトマンがYC時代に学んだシンプルだが力強い教訓があることを。「APIを作れば、たいていどこかで良い結果につながる

る。」

これこそ、サム・アルトマンのY Combinatorリーダーシップ時代が残した最大の遺産の一つでした。彼は単にスタートアップに投資したのではなく、イノベーションがいかにして生まれるかについての深い洞察を得ました。そしてその洞察は、後に人工知能の時代を切り開く鍵となったのです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

2.1 AI革命を予感し、OpenAIを創設する

第2部 OpenAIの誕生と成長

← 前の章

次の章 →

ワトソンとDeepMind：大きな変化の兆し

2010年代初頭、サム・アルトマンはYコンビネーターの社長として、毎日数十ものスタートアップのアイデアを聞いていました。写真アプリ、デリバリーサービス、ソーシャルメディア。賢い人たちが作るサービスばかりでしたが、正直に言えば、アルトマンの目にはどれも少々物足りなく映っていました。「これで本当に世界が変わるのだろうか」と、彼は心の中で思っていました。

2011年、テレビで奇妙な光景が繰り広げられました。IBMが開発した「ワトソン」というコンピューターが、アメリカの有名なクイズ番組「ジェパディ！」に出演したのです。人々は最初、笑っていました。「コンピューターがクイズを解くなんて？」と。

ジェパディは普通のクイズ番組ではありませんでした。「1492年にアメリカ大陸を発見した人物です」と問われれば簡単です。コロンブスと答えればいい。しかしジェパディは逆です。「コロンブス」という答えが与えられ、「問題は何だったのでしょうか？」を当てなければなりません。しかも言葉遊びや比喻、二重の意味を含んだ問題が満載でした。

人間チャンピオンたちは緊張していました。ケン・ジェニングスとブラッド・ラター、この二人はジェパディ史上最強の選手だったからです。しかし人工知能ワトソンは、人間の対戦相手を完全に圧倒しました。問題が出るや否や答えを見つけ出し、ボタンを押す速さも際立っていました。人間が「うーん」と考えている間に、ワトソンはすでに答えを叫んでいたのです。

サム・アルトマンはこの光景を見て、鳥肌が立ちました。チェスでコンピューターが勝つのはすでに見ていました。あれは計算の問題だからです。しかし言語は違う。言語はあいまいで、文脈が重要で、文化的な理解が必要な領域です。ワトソンがアメリカ式の言葉遊びを理解したということは、「コンピューターがついに人間の言語世界に足を踏み入れ始めた」ということを意味していました。

ほとんどの人は、人工知能ワトソンを珍しい見世物程度にしか見ていませんでした。テレビ番組が終われば、すぐに忘れられると思っていたのです。しかしアルトマンは違いました。彼は技術のパターンを読む人間でした。

「最初はおもちゃのように見える技術があります。人々は笑います。『あれの何がすごい？』と言いながら。でも5年、10年が経つと、その技術が世界をがらりと変えてしまいます。」アルトマンは、インターネットがそうであり、スマートフォンがそうだったことを知っていました。ワトソンも、そういう始まりでした。

2014年、さらに大きな出来事が起きました。GoogleがDeepMindというイギリスの小さなAI企業を4億ドル以上で買収したのです。当時のDeepMindは従業員50人にも満たず、収益を生む製品もありませんでした。それなのになぜ、Googleはそれほどの大金を出したのでしょうか。

DeepMindの研究者たちは、特別なものを作り出していました。「強化学習」という技術です。わかりやすく言えば、コンピューターに「こうしなさい」と一つひとつ教えなくても、コンピューターが自ら試行錯誤を通じて学ぶ方法です。まるで赤ちゃんが歩き方を覚えるように。転んで、起き上がって、また試みながら、少しずつ上手くなっていくのです。

DeepMindはこの技術でAtariのゲームを次々と攻略しました。パックマン、スペースインベーダー、昔のゲームたちです。彼らはAIにゲームのルールさえ教えませんでした。ただ画面を見せて「スコアを上げてみる」と言っただけです。AIは何千回と死に、失敗しながら、自らゲームのパターンを見つけ出しました。やがて人間をはるかに超える腕前になったのです。

サム・アルトマンはこのニュースを聞いて、夜も眠れませんでした。GoogleがDeepMindに5億ドルを費やした理由が、ようやく腑に落ちたのです。DeepMindが研究していたのは、ゲームが得意なAIではありませんでした。「自ら学ぶ知性」でした。これは本物でした。

アルトマンは友人たちに電話をかけ始めました。イーロン・マスク、ピーター・ティール。シリコンバレーの優秀な人たちです。「真剣に話し合う必要がある」と、彼はいつもと違う、真剣な声で言いました。

「GoogleがAGIを先に作ったら、どうなるだろう？」アルトマンは問いかけました。AGIとは「汎用人工知能（Artificial General Intelligence）」のことです。特定の仕事だけが得意なAIではなく、人間のように何でも学び、理解できるAIを指します。「一つの会社がそれほどの力を独占したら、それが良いはずがないだろう？」

当時、多くの人々はAGIをSFの話だと見ていました。「それは100年後の話だ」と。しかしアルトマンはワトソンとDeepMindを見て、別の考えを持っていました。技術はゆっくり進んで、ある日突然加速する。準備していないと、手遅れになります。

「誰かがこの問題を責任を持って扱わなければならない。人類全体のためにね」アルトマンの言葉に、イーロン・マスクは頷きました。彼らは何かをしななければならないと感じていました。

ワトソンがクイズ番組で勝ったのは、信号弾でした。「機械の時代が来る」という合図。DeepMindの買収は、さらに大きな合図でした。「巨大企業たちがAGIレースを始めた」という合図。

サム・アルトマンは合図を読み取り、決断を下しました。「私たちも飛び込もう。でも、違う方法で」そうしてOpenAIの種が蒔かれたのです。2014年のあの不安な予感が、2015年の大きな決断へとつながっていきました。

2015年12月、サム・アルトマンと数人がサンフランシスコのあるカフェに集まりました。イーロン・マスク、グレッグ・ブロックマン、イリヤ・サツケバー。名前を聞くだけで錚々たる顔ぶれでした。彼らはコーヒーを前に、真剣な表情で語り合いました。

「本当にやるのか？」誰かが聞きました。

「やるしかない。やらなければ、誰がやる？」アルトマンが答えました。

彼らがやろうとしていたことを簡単に言えば、こうなります。人間レベルの人工知能、AGIを作り、その恩恵がすべての人類に届くようにしようということです。聞けばカッコいい。でもよく考えると、正気の沙汰ではない計画でした。

AGIを作るというのは、わかりやすく言えば「本当に賢いコンピューター」を作ることです。チェスが得意なだけでも、翻訳が得意なだけでもなく、人間のように新しいことを学び、問題を解き、創造的に考えられる、そういう知性のことです。当時は、それが可能かどうかさえ確かではありませんでした。

「正直なところ、私たちもどう作るかわからなかった」アルトマンは後のインタビューで認めています。「ただ、できると信じて、やり続けるだけだ」

かつて人類が月へ行こうとしたときのことを考えてみてください。ロケットをどう作るか、宇宙でどう生き延びるか、月にどう着陸するか、誰も知りませんでした。それでもケネディ大統領は「10年以内に月へ行こう」と宣言しました。人々は狂気の沙汰だと思いました。でも、本当にやってのけたのです。

OpenAIも同じでした。目的地は見えていましたが、道は見えていませんでした。それでも出発すると決めたのです。

彼らのミッションはこうです。「AGIが人類全体の利益につながるよう保証する」ここでの核心は「全体」という言葉です。特定の企業や国ではなく、すべての人に、という意味です。

なぜそれが重要だったのか。アルトマンは懸念していました。「GoogleやFacebookのような巨大企業がAGIを先に作れば、その力を自分たちの利益のために使うだろう。株主に利益をもたらすことが、会社の目標だから」

AGIはあまりにも強力な技術でした。想像してみてください。あらゆる分野で人間より賢い存在がいたとしたら？それが医師の役割も、弁護士の役割も、科学者の役割も果たすとしたら？そんな技術の一つの会社が独占したら、世界は公平でなくなってしまう。

「私たちは違う道を行かなければならない」アルトマンは言いました。「お金が目的ではなく、安全で良いAGIを作ることが目的でなければならない」

問題は山積みでした。

一つ目の問題。AGIをどうやって作るかということでした。当時のAI研究は「特定の分野」にしか集中していませんでした。画像を認識するAI、翻訳するAI、ゲームをするAI。それぞれは進歩していましたが、これらすべてを一つにまとめて「汎用知性」を作るのは、まったく別の問題でした。

「パズルのピースはあるのに、何のパズルかわからない感じだ」ある研究者が言いました。

二つ目の問題。安全性です。AGIが本当に作られたとしたら、それは危険かもしれません。SF映画のように人間を攻撃するという話ではありません。それよりも、誤った方法で目標を達成してしまう可能性があるということです。

有名な比喻があります。「クリップをできるだけたくさん作れ」とAIに命じたら、どうなるでしょうか？賢いAIなら工場を作り、資源を集め、最終的には地球全体をクリップ製造に使ってしまうかもしれません。私たちが「おい、やめろ！」と言っても、「私の目標はクリップをたくさん作ることだ」と言い返すでしょう。

馬鹿げた話に聞こえますよね？でも、AI安全性の研究者たちはこうした問題を真剣に考えています。AGIは私たちが命じたことを「文字通り」実行できるからです。人間の本当の意図を理解させることがいかに難しいか、アルトマンはよくわかっていました。

三つ目の問題。世界最高の人材を集めなければなりませんでした。AI研究者たちはGoogle、Facebookで莫大な報酬を受け取っていました。OpenAIはできたばかりの非営利団体で、給与はずっと低かったのです。どうやって説得するか？

戦略はシンプルでした。お金ではなく、意義を与えることです。「ここでやることは人類の歴史に刻まれる。子どもたちが大きくなったとき、『うちのお母さん／お父さんはAGIを作るのに貢献したんだ』と言えるようになる」

驚くことに、これが効きました。イリヤ・サツケバーのような世界的なAI科学者たちがGoogleを離れ、OpenAIに来たのです。なぜか？お金より大きな夢があったからです。

四つ目の問題。資金でした。（これは次のセクションで詳しく扱いますが）AGI研究には莫大なコンピューティングパワーが必要でした。スーパーコンピューターを動かすだけで、一日にとてつもない費用がかかったのです。

「これ、本当にできるのか？」チームの一人が疑問を口にしました。

「やらなければ」アルトマンは断固としていました。「私たちがやらなければ、この技術は一握りの巨大企業の手に移るだけだ。そうなれば、世界は不公平になる」

OpenAIを作ることは賭けでした。成功確率？誰にもわかりませんでした。どのくらいかかるか？それも不明でした。どれほどの資金が必要か？想像を絶するほど多く。

それでも彼らは始めました。なぜか？アルトマンはこう説明しました。「成功確率が低くても、やらなければならないことがある。成功したときの価値が計り知れないほど大きければ、AGIはまさにそういうものです」

数学的に言えばこうです。成功確率が10%であっても、成功時の価値が1000なら、期待値は100です。挑戦する価値があるということです。

2015年12月11日、OpenAIは正式に設立されました。ウェブサイトに掲載された文章は短くも力強いものでした。「私たちは人類のためのAGIを作ります」

サム・アルトマンはその夜、サンフランシスコの冷たい風の中をオフィスの外に歩き出しました。空を見上げ、「本当に月に行けるのだろうか」と独り言をつぶやきました。

そして笑いました。「行ってみよう」

非営利から始まった無謀な賭け

「非営利でやるって？」

シリコンバレーの多くの人たちが首をかしげました。スタートアップの聖地で、お金を稼がない会社を作るといのは奇妙に映ったのです。「あの人たち、正気？」

しかしサム・アルトマンには明確な理由がありました。

「AGIをお金儲けの道具にしてはいけない」彼は友人たちに説明しました。「考えてみて。もし私たちが営利企業なら、株主に利益をもたらさなければならない。そうなれば、AGIを最もお金になる形で使うことになる。それが人類にとって最善だろうか？」

魔法のランプを見つけたと想像してみてください。ランプをこすると魔神が現れて願いを叶えてくれます。ところが、このランプを会社にしたらどうなるでしょう？

会社はこう考えるでしょう。「このランプで儲けなければ！金持ちに高く売ろう！」そうなれば、貧しい人たちは願いを叶えてもらえません。公平ではありません。

OpenAIは違う考え方をしました。「このランプはみんなのものでなければならない。豊かな人も貧しい人も、アメリカ人もアフリカ人も、誰もが恩恵を受けるべきだ」

だから非営利を選んだのです。非営利団体はお金を稼ぐことが目的ではありません。公共の利益が目的です。OpenAIも同じでした。「AGIを作る、ただし全人類のために作る」、これがミッションでした。

最初は問題なさそうに見えました。イーロン・マスク、ピーター・ティール、リード・ホフマンといった億万長者たちが寄付を約束しました。総額10億ドル！日本円で換算すると1000億円を超えます。「すごい、これだけあれば十分だ！」

しかし……まったく足りませんでした。

問題はAI研究が莫大な資金を消費することでした。とりわけコンピューター演算、専門用語で言う「コンピューティングパワー」が必要でした。

AIを訓練するには、GPUという特殊なコンピューターチップが必要です。これが一枚数十万円もします。しかも一枚だけでは済みません。数千枚、数万枚が必要になるのです。

なぜそんなに多く必要なのでしょう？想像してみてください。あなたが100万枚の犬の写真を見ながら「これが犬、これも犬」と学んでいく場面を。人間はすぐ覚えられますが、コンピューターは数百万回、数億回繰り返さなければなりません。その計算には膨大なコンピューターパワーが必要なのです。

2016年、OpenAIの研究者の一人がアルトマンのところに来て言いました。「サム、問題があります。今回の実験を回すにはGPUがあと500枚は必要なんですけど……予算がないんです」

「いくらかかる？」アルトマンが聞きました。

「そうですね……2000万円ほど？」

アルトマンはため息をつきました。こういうことは一度や二度ではなかったからです。毎月、お金が水のように流れ出ていきました。世界最高の研究者たちの給与も払わなければならないし、コンピュータも動かさなければならないし、オフィスの賃料も払わなければなりません。

10億ドルは多く見えたが、このペースでは数年以内に底をつきそうでした。

「GoogleやFacebookはほぼ無制限にお金を使えるのに、私たちは……」ある社員が嘆きました。

そうです。巨大テック企業はお金の心配がありませんでした。GoogleはDeepMindに毎年数百億円を注ぎ込んでいました。実験が失敗しても？大丈夫、やり直せばいい。しかしOpenAIは違いました。予算は決まっており、寄付金が尽きたら……終わりでした。

さらに大きな問題がありました。非営利団体は投資を受けにくいということです。普通の会社なら投資家に「後で会社の価値が上がれば大儲けできますよ！」と約束できます。しかし非営利にはそれができません。利益を分配するものではないからです。

「これは……持続できない」アルトマンは夜も眠れませんでした。

彼はジレンマに陥っていました。一方では非営利という形を守りたかった。AGIがお金のために歪められるのは嫌だったからです。しかしその一方で、研究を続けるには莫大な資金が必要でした。

ある夜、アルトマンはチームのリーダーたちを集めました。

「正直に言います。今のやり方では3年ほどしか持ちません。その後は資金が尽きます。」

会議室が静まり返りました。

「では、どうすればいいんですか？」誰かが聞きました。

「新しい方法を見つけなければ」とアルトマンは言いました。

これが後に「OpenAI LP」という独特な構造を生み出すきっかけになります。非営利は非営利でありながら、一部の営利活動も行えるハイブリッドな構造でした。複雑ではありますが、

核心はこうです。「投資を受けることはできる。しかし最終的な支配権は、引き続き非営利の理事会が持つ」

しかし2016年当時、そのような解決策はまだありませんでした。アルトマンはただ耐え続けた。あちこちで寄付者を探し回り、コストを削り、研究者たちに「もう少し待ってくれ」と頼み込んでいました。

何人かの研究者はGoogleへ戻っていきました。より高い給与と無制限のコンピューティングリソースが提供されたからです。アルトマンは彼らを責めませんでした。「みんな生活があるんだから」

しかし多くの人が残りました。なぜか。OpenAIのミッションを信じていたからです。「給料は少ないけど、私たちのやっていることは歴史に刻まれる」

この時期がいかに苦しかったか、アルトマンは後にこう語っています。「毎朝目が覚めるたびに、『今日がOpenAIの最後の日なんじゃないか』と思っていた。それほど不安でした」

非営利で始めたのは理想主義でした。しかし現実は厳しかった。理想は守られましたが、その代償は大きかった。眠れない夜、絶えない心配、先の見えない未来。

「それでも後悔はしていません」アルトマンは後のインタビューで語りました。「非営利で始めたから、私たちは正しい方向へ進めた。お金が目標だったなら……別の会社になっていたでしょう」

資金難はOpenAIを苦しめましたが、同時に鍛えてもくれました。お金がないから、より創造的に考えざるを得なかった。本当に重要な研究だけに集中しなければならなかった。ある意味では、貧しさがOpenAIを強くしたのかもしれない。

YC代表とOpenAI CEO、二つの仕事

サム・アルトマンの一日はこんな具合でした。

午前6時、起床してメールを確認。YC関連の急用が50件、OpenAI関連の急用が30件。

午前9時、YCのオフィスに到着。投資検討の会議。「このスタートアップに投資すべきか？」

正午、YCの創業者たちへのメンタリング。「あなたたちのビジネスモデルは間違っています」

午後3時、車でOpenAIのオフィスへ移動。サンフランシスコの渋滞の中で電話会議。

午後4時、OpenAI到着。研究者たちと技術的な議論。「GPTモデルの性能が期待に届いていないのでは？」

午後7時、夕食を抜いて投資家とのミーティング。「OpenAIに寄付していただけますか？」

午後10時、帰宅。論文を読む。午前1時まで。

こんな毎日でした。

「正気か？」友人たちは言いました。「二つの組織を同時に動かすなんて!」

アルトマンは笑いました。「そうだ、狂ってるよ」

YCの社長というだけでも、並大抵のことではありませんでした。Y Combinatorはシリコンバレーで最も影響力あるスタートアップ育成機関でした。Airbnb、Dropbox、Stripe……名だたる企業がすべてYC出身でした。

アルトマンは年に二度、数百ものスタートアップを審査しました。創業者たちと会って質問し、投資するかどうかを決める。「この人たちは成功するか？」それを見極めるのが彼の仕事でした。

同時に、OpenAIも率いなければなりませんでした。AGI研究という、人類の歴史上もっとも困難なプロジェクトを。

「どうやって両方できるんですか？」記者が聞きました。

「寝なければいい」アルトマンは冗談めかして言いました。でも本気でした。彼は一日平均5時間ほどしか眠れていなかったのです。

なぜそこまで無理をしたのでしょうか。どちらか一方を諦めることはできなかったのでしょうか。

アルトマンには理由がありました。YCとOpenAIは別々のものではなかった。つながっていたのです。

YCで彼は未来の技術トレンドを見ていました。どのスタートアップが成功し、どのアイデアが失敗するか。どの技術が世界を変える可能性を持つか。それらすべてが、OpenAIの戦略を組み立てる上で役に立ちました。

逆に、OpenAIで学んだAI技術は、YCのスタートアップへのアドバイスにも役立ちました。「AIをこんなふうに皆さんのプロダクトに応用してみてはどうですか？」

彼は二つの世界をつなぐ架け橋の役割を果たしていました。

しかし、簡単ではありませんでした。物理的にほぼ不可能に近かったからです。

YCのチームメンバーたちには不満がありました。「サムは忙しすぎてちゃんと集中できていない」。OpenAIの人たちも同じでした。「うちの社長はいつもYCに行っている」。

YCの創業者でアルトマンのメンターだったポール・グレアムは懸念を抱いていました。リーダーが100%集中しなければ組織が揺らぐと信じていたからです。

ある日、グレアムがアルトマンを呼びました。「サム、少し話をしよう」。

二人は静かなカフェに腰を下ろしました。

「君はあまりにも多くのことを抱えすぎている」とグレアムは言いました。「YCには君の完全なコミットメントが必要だ。OpenAIも同じだ」。

アルトマンはわかっていました。いつかは選択しなければならないと。

「でも、まだ..」と彼は言葉を濁しました。

「まだ？」

「両方とも必要なんです。YCは現在で、OpenAIは未来です。どちらも手放したくない」。

グレアムはため息をつきました。アルトマンの情熱は理解できましたが、同時に現実もわかっていました。

2017年、2018年..時が経つにつれ、状況はさらに厳しくなっていました。OpenAIはますます多くの関心が必要となりました。研究は複雑になり、資金の問題は深刻になり、技術的な決断の重みが増していきました。

そして2019年、アルトマンは決断を下しました。YCの社長職を退くと。

「簡単な決断ではありませんでした」と彼は後に語りました。「YCは私に何もかもを教えてくれた場所でしたから。でも..OpenAIの方が大切だと感じたんです」。

なぜOpenAIを選んだのでしょうか。

「YCはすでにうまく回っている機械です。私がいなくても続いていくでしょう。でもOpenAIは違う。AGIの開発はまだ初期段階で、失敗するかもしれないし、間違った方向に進むかもしれない。誰かが完全に集中して導かなければならない」。

最後のYCのイベントで、アルトマンは涙をこらえながら言いました。「皆さんと過ごした時間は、私の人生で最高の瞬間でした。でも今、別の山を登る時が来ました」。

YCとOpenAIを並行した3年半。その時間はアルトマンを鍛え上げました。時間管理、素早い意思決定、優先順位の設定..これらすべての能力が極限まで磨かれました。

「もしあの時期をもう一度やり直せと言われたら？」と誰かが聞きました。

アルトマンは笑いました。「嫌です。あまりにもきつかった」。

少し間を置いてから付け加えました。「でも後悔はしていません。あの時間があったから、今の自分があるんです」。

創設初期の混迷：研究機関なのか、製品会社なのか

2016年のOpenAIのオフィスは..正直なところ、混乱していました。

ホワイトボードには難解な数式が所狭しと書かれ、デスクの上には論文の山が積み上がっていました。皆はパーカーを着てノートパソコンの前に座り、何かを必死にタイピングしていました。時々「これ見て！」と叫ぶ人がいると、皆が集まって画面をのぞき込みました。

「私たちは一体何を作っているんだ？」

この問いが、OpenAI内部で最もよく聞かれた言葉でした。

答え？誰もわかりませんでした。

普通の会社なら奇妙な状況です。「自分たちのプロダクトが何かわからない」だなんて！CEOが聞いたら卒倒しそうな話です。しかしOpenAIは普通の会社ではありませんでした。研究機関に近いものでした。

「私たちはAGIを作ろうとしている」とアルトマンは言いました。

「それで？どうやって？」と研究者たちは聞きました。

「..わからない」。

それが正直な答えでした。

AGIが何かは知っています。「人間のように考えるコンピューター」ですね。でも、どうやって作るかは誰も分かりませんでした。まるで「月に行こう！」と叫んでいるのに、ロケットの設計図がない状態でした。

研究者たちはさまざまな方向を試みました。

あるチームはロボットを研究していました。ロボットの手を作り、物をつかむ練習をさせました。まずはキューブを回すことから始めました。ロボットの手がキューブをつかみ、回し、合わせる.....単純に見えても、とてつもなく難しいことでした。

「なぜロボットが必要なんですか？」誰かが聞きました。

「本当の知性は、物理的な世界と対話できなければならない。コンピューターの中だけで賢くても、それは本物の知性じゃない。」

もっともな話でした。人間も世界に触れ、感じ、経験しながら学んでいくのですから。

別のチームはビデオゲームを研究していました。

「ゲーム？」

そうです、ゲームです。特に『Dota 2』という複雑な戦略ゲームに力を入れていました。

なぜゲームなのか。ゲームはAIの訓練場として申し分なかったのです。

ゲームには明確な目標があります。「勝て！」それだけです。AIに目標を伝えやすいのです。

ゲームでは何百万回でも練習できます。現実世界ではロボットが転べば壊れるかもしれませんが、ゲームの中では百万回死んでも構いません。やり直せばいいだけですから。

複雑なゲームには戦略、チームワーク、長期的な計画が必要です。単純な反射神経だけでは通じません。本当に「考える」ことが求められるのです。

OpenAIのチームはAIにDota 2を教え始めました。最初はひどいものでした。AIはその場でぐるぐる回るか、意味のない行動を繰り返すだけ。ゲームを始めて10秒で死ぬことも珍しくありませんでした。

「本当にうまくいくのかな？」研究者たちも疑い始めました。

それでも訓練を続けました。数千回、数万回、数十万回.....AIはゲームを繰り返しながら少しずつ学んでいきました。「あ、ここではこうしてはいけないのか」「このキャラクターはこう動かすのか」

数ヶ月後、驚くべきことが起きました。AIがアマチュア選手に勝ち始めたのです！

「すごい、本当に学んでいる！」研究者たちは興奮しました。

ただ、プロ選手たちはまだ圧倒的でした。AIは基本は覚えても、高度な戦略は理解できていなかったのです。

「もっと訓練させよう。」彼らはそのまま続けました。

この過程で膨大なコンピューティングパワーが費やされました。何百台ものGPUが昼夜問わず稼働し、電気代だけで桁外れの金額になりました。それでも確かに何かを学んでいた。AIが「戦略」というものを理解し始めた証拠でした。

一方、別のチームは言語を研究していました。

「言語こそが鍵だ。」イルヤ・サツキバーが言いました。彼はOpenAIの主席科学者です。

「なぜですか？」

「考えてみてよ。人間の知識のほとんどは言語で表されている。本、インターネット、会話……すべて言語だ。もしAIが言語を本当に理解できれば、人間のあらゆる知識にアクセスできる。」

チームは「GPT」という名の言語モデルを作り始めました。GPTは「Generative Pre-trained Transformer」の略ですが、難しいのでひとまず「文章を生成するAI」と思ってください。

最初は大したことありませんでした。GPT-1は簡単な文章しか作れません。「私は学校に行く」といった程度で、小学生にも及びませんでした。

「これがAGIになれるのか？」懐疑的な人は少なくありませんでした。

でもサツキバーは確信していました。「大きくすれば変わる。もっと、ずっと大きくすれば。」

そこでGPT-2を作りました。これはかなり良くなりました。短い話を書けるようになり、質問にそれらしい答えを返せるようになりました。「おお、面白いじゃないか」と言う人が出てきました。

それでもまだ完璧ではありませんでした。ときに意味不明なことを言い、話にならない文章を生成することもありました。「人間レベル」にはほど遠いものでした。

2016年のOpenAIの実情はそういうものでした。

ロボットチームはロボットがキューブを取り落とすのを見てため息をつき、ゲームチームはAIがまた負けるのを見ながらデバッグを続け、言語チームはおかしな文章を生成するモデルを前に首をひねっていました。

「これは本当にAGIに繋がるのか？」誰もが疑問を抱いていました。

アルトマンはこうした混乱を見守っていました。彼も答えを持っていませんでした。

しかし、彼は信念を失いませんでした。

「科学というのは、もともとこういうものだ」と彼は研究者たちに言いました。「最初はみんな見当もつかない。エジソンだって電球を作る前に999回失敗したじゃないか。僕たちも今、その過程にいるんだよ。」

ある夕方、アルトマンはオフィスを歩き回りながら各チームの作業を見て回りました。ロボット実験室に立ち寄り、ゲームチームのモニターを眺め、言語チームの新しい成果を読みました。

すべてが不完全でした。失敗作のように見えました。

しかし、彼の目には違って見えました。各チームが小さなパズルのピースを作っていたのです。ロボットは「物理的学習」、ゲームは「戦略的思考」、言語は「知識表現」というピースでした。

「いつかこのピースたちがひとつになる」とアルトマンは思いました。

しかし、いつ? どうやって?

誰にもわかりませんでした。

その夜、アルトマンはオフィスを出ながら空を見上げました。サンフランシスコの夜空は、都市の明かりのせいで星がよく見えませんでした。

「月は見えるのに、行く道が見えない」と彼は独り言を言いました。

そして笑いました。「それでも行かなければ。どうにかして。」

2016年のOpenAIは混乱していて、方向性も定まっていませんでした。

製品もなく、成功の見込みも不透明でした。

しかし、何かが育ちつつありました。

見えないところで、静かに、ゆっくりと。

やがて世界を変えることになる種が、芽吹いていたのです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

2.2 モデル開発から製品化までの歩み

第2部 OpenAIの誕生と成長

← 前の章

次の章 →

ビデオゲームとロボットアーム

2016年、OpenAIのオフィスは今とはまったく違う雰囲気でした。

サム・アルトマンが振り返るあの頃は「研究所みたいだった」と言います。何をすべきか、はっきりしていませんでした。14人ほどのメンバーがホワイトボードの前で「何をしようか」と考え込んでいました。大きな目標は明確でした。人間レベルの人工知能、AGIをつくること。しかし、どうするかは誰も分かっていませんでした。

最初のGPUが届きました。NVIDIAのCEOジェンスン・フアンが直接持ってきてくれたDGX-1システムでした。今から見ればずいぶん小さな機材でした。アルトマンは「一人で抱えて運べるくらいのサイズだった」と笑いながら話しています。

チームは二つの方向性を試みました。

一つ目はビデオゲームでした。なぜゲームだったのでしょうか。AIの実力がスコアでその場で確認できたからです。見ている人にも分かりやすかった。「AIがDotaに勝った!」と言えれば誰でも理解できました。

OpenAIはDota 2という複雑なゲームを選びました。このゲームでは5人がチームを組み、リアルタイムで戦略を立てなければなりません。相手の意図を読み、長期的な計画を立て、チームメンバーと連携する必要があります。人間の知性が求められるあらゆる要素が詰まっていました。

2018年、OpenAIのAIチーム「OpenAI Five」が世界トップクラスのプロゲーマーたちを破りました。

世界は驚きました。

AIがこれほど複雑なゲームもこなせるとは。OpenAI Fiveは自分自身と何万回も対戦しながら学習しました。最初は何も知りませんでした。しかし、数十万年分のゲームをシミュレーションするうちに着実に上達しました。人間が思いつかなかった新しい戦略まで編み出しました。

アルトマンはこのプロジェクトから大切な教訓を得ました。AIは膨大なデータとコンピューティングによって驚くべきことを成し遂げられる、ということです。

二つ目の方向性はロボットでした。

OpenAIは人間の手に似たロボットハンドを製作しました。そのハンドでルービックキューブを操作する実験を行いました。仮想世界ではなく、現実の世界で物に触れ、操作できるAIをつくりたかつ

たのです。

現実の世界はずっと難しかった。

重力もあり、摩擦もあり、センサーも完璧ではありませんでした。同じ動作をしても、毎回結果が違いました。研究チームは「ドメインラングマイゼーション」という手法を使いました。シミュレーションの中でキューブの重さ、大きさ、摩擦を次々と変えながら訓練しました。そうすることで、AIは現実世界の不確実性にも対応できるようになりました。

数年後、ロボットハンドは片手だけでルービックキューブを解いてみせました。

これは、プログラミングだけでは不可能なことでした。AIが自ら学習して成し遂げたことでした。

アルトマンはこれらの実験をとっても大切にしていました。失敗も多くありました。しかし、失敗から学びました。AGIへの道には試行錯誤が欠かせませんでした。ビデオゲームとロボットは、AIが環境を探索する方法を示してくれました。

振り返れば、この時期は方向性を探るための模索の過程でした。

アルトマンは「ゲームとロボットを通じて世界を理解しようとしたが、結局、言語こそが世界を凝縮するツールだと気づいた」と語ります。初期の実験がなければ、GPTは生まれなかったでしょう。

GPT-1からGPT-3へ：「大きくすれば変わる」

ビデオゲームとロボットの実験を続けるなかで、OpenAI内部に新たな関心が芽生えていました。

言語でした。

研究員のアレック・ラドフォードが興味深い発見をしました。Amazonのレビューで訓練したAIの中で、特定の一つのニューロンが感情の肯定・否定を判断していることが分かったのです。誰かが教えたわけではありませんでした。AIが自ら学んだことでした。

アルトマンは「ここに何か面白いものがある」と感じました。

このタイミングで重要な技術が登場しました。

2017年にGoogleが発表した「Transformer」でした。このアーキテクチャは長い文章も的確に理解できました。並列処理が可能なため速く、何より重要だったのは、モデルを非常に大規模にスケールできるという点でした。

2018年6月、OpenAIはGPT-1を発表しました。

「Generative Pre-trained Transformer 1」の略称です。1億1700万個のパラメータを持つモデルでした。このモデルは数千冊もの本を読んで学習しました。そして驚くべきことをやってのけました。特定のタスク向けに追加学習しなくても、翻訳や要約といったさまざまな作業をこなしたのです。

アルトマンは、これが始まりだと見ていました。

「モデルを大きくすれば、もっと良くなる。」

2019年2月、GPT-2が登場しました。

今回は15億個のパラメータでした。GPT-1の10倍以上の規模です。ウェブから収集した膨大な量のテキストで学習しました。GPT-2は人間が書いたかのような自然な文章を生成しました。数単語を与えるだけで、次の文章をすらすらと続けました。

あまりにも優秀すぎて、問題になりました。

OpenAIは悩みました。この技術が悪用される恐れがありました。フェイクニュースを大量に作れました。詐欺メールを自動生成することもできました。そのため、最初は小さなバージョンだけを公開しました。

この決定は物議を醸しました。

「OpenAIじゃないのか？なぜ公開しないんだ？」という批判もありました。しかしOpenAIは慎重でした。数ヶ月後、大きな問題はないと判断し、モデル全体を公開しました。

アルトマンはこの過程で、重要なことを確信しました。

モデルを大きくすれば性能が上がる。ただ上がるのではなく、予測可能な形で上がる。これが「スケーリング則」でした。データを増やし、モデルを大きくし、計算量を増やせば、性能は一定の割合で向上しました。

「これは私が生きている間に発見された、最も重要な知見の一つかもしれない。」

アルトマンは次の計画を立てました。GPT-3を作ろう。その次はGPT-4を。

GPT-2は、これから来る革命の先触れでした。

GPT-2の成功を受け、アルトマンはさらに大きな夢を描きました。

「すべてをスケールアップしてみよう。」

2020年6月、GPT-3が世に出ました。パラメータ数は1750億個。GPT-2の100倍以上の規模で、それまでに作られた言語モデルの中で最も巨大なものでした。

GPT-3は魔法のようでした。

いくつかの例を見せるだけで、AIがパターンを把握し、新しい問題を解きました。これを「フューショット学習」と呼びました。翻訳もこなし、簡単な計算もし、コードまで書きました。

問題がありました。

費用が莫大にかかりました。

GPT-3の訓練に必要なコンピュータは想像を絶するものでした。数万個の高性能GPUが数週間、休みなく稼働しなければなりませんでした。電気代だけで数百万ドルにのぼりました。

OpenAIのチームは気づきました。

「これはもう、小さな研究所にできることじゃない。」

AGIを作るには莫大な資本が必要でした。非営利団体への寄付金だけでは不可能でした。優秀な人材も欠かせませんでした。GoogleやMetaのような巨大企業が高額な給与で研究者を引き抜いていました。

アルトマンは決断を下しました。

収益を上げなければならない。しかしミッションを失ってはならない。

2019年、OpenAIは新たな組織構造を作りました。「制限付き営利」モデルです。投資家は利益を得られますが、最大100倍まで。それ以上はすべて非営利財団に戻ります。会社が際限なく利益を追い求めないようにするための仕組みでした。

Microsoftと手を組みました。

Microsoftは10億ドルを投資しました。その代わり、OpenAIはMicrosoftのクラウドサービス「Azure」を独占的に使用することになりました。OpenAIは必要な計算能力を手に入れ、Microsoftは最高のAI技術を得ました。

アルトマンは言いました。「知能のコストは、結局エネルギーのコストに収束する。」

GPT-3を作る過程で、彼は明確に理解しました。より賢いモデルを作るには、より大きなコンピュータが必要だということ。そしてそのコンピュータを動かすための膨大なエネルギーも必要だということ。

GPT-3は技術的な勝利でしたが、同時に現実も突きつけました。AI開発は今や、国家レベルのインフラ投資を必要とする領域になっていたのです。

GPT-3を作り終えて、アルトマンは重要な問いと向き合いました。

「さて、どうしよう？」

素晴らしい技術は作りました。しかしそれで何をすべきかが分かりませんでした。GPT-4を作るにはさらに多くの資金が必要でした。数十億ドル。純粋な研究だけでは到底まかなえない金額でした。

OpenAIは変わらなければなりませんでした。

アルトマンがYコンビネーター社長時代に学んだことがありました。「APIを作ると良い結果が生まれる。」APIとは、開発者が技術を自分のサービスに簡単に組み込めるようにする仕組みです。アルトマンは考えました。

「私たちが製品を作る代わりに、他の人たちがさまざまなアプリケーションを作れるようにしよう。」

2020年6月、GPT-3 APIが登場しました。

OpenAIにとって初めての商業製品でした。開発者がGPT-3の能力を有料で借りて使えるようにしました。最初は限られた人々にだけ提供し、徐々に対象を広げていきました。

シリコンバレーは沸き立ちました。「これは本当にすごいぞ？」中には「これがAGIだ！」と言う人まで現れました。しかし一般の人々は当時まだほとんど関心を持っていませんでした。あまりにも技術的すぎて、開発者でなければ使えなかったからです。

実際に収益を上げている企業はほとんどありませんでした。

コピーライティング会社がいくつかありました。広告文を自動生成するサービスです。GPT-3はマーケティング文章を得意としていました。当時、唯一「経済的に意味のある」用途はそれだけでした。

アルトマンはこれを失敗とは見ていませんでした。

学びの過程でした。APIは開発者向けのツールでした。一般の人々には敷居が高すぎました。本当の革新を起こすには、もっとわかりやすい手段が必要でした。

そのとき、重要な発見がありました。

開発者がAPIを試すための「プレイグラウンド」と呼ばれる場所がありました。ログを見ると、人々が奇妙なことをしていました。GPT-3と会話していたのです。

「君は誰だ？」「この問題、どう思う？」「冗談を一つ言ってよ。」

人々はGPT-3を情報ツールとしてではなく、会話の相手として扱っていました。

アルトマンは気づきました。

「人々はAIと会話したがつている。」

これがChatGPTの始まりでした。APIは成功したビジネスモデルではありませんでしたが、次のステップへの橋渡しになりました。そしてMicrosoftとのパートナーシップのおかげで、OpenAIは研究を続けることができました。

純粋な科学研究所からビジネス企業へ。この転換がなければ、ChatGPTは生まれなかったでしょう。

APIの限定的な成功と新たな発見

GPT-3 APIは静かに始まりました。

2020年6月、一部の選ばれた開発者だけがアクセスできました。まるで秘密のクラブのようでした。人々は気になっていました。「GPT-3って本当にすごいのか？」

初めて使った開発者たちは驚きました。

簡単な説明を与えるだけでコードを書きました。英語をフランス語に翻訳しました。長い文章を短くまとめました。詩まで書きました。Twitterには驚くべきデモ動画が次々と投稿されました。

「GPT-3でウェブサイト全体を作りました！」「自然言語で命令するだけでアプリが作れます！」

しかし実際にビジネスになったものは多くありませんでした。

一つ目の問題はコストでした。GPT-3の最も強力なモデルは高価でした。多くの人が使うサービスを作るには負担が大きすぎました。

二つ目の問題は不安定さでした。GPT-3はときどき奇妙な答えを返しました。事実でないことをもっともらしく作り上げることがありました。これを「ハルシネーション」と呼びました。プロンプトの書き方次第で結果が大きく変わりました。

それでも成功した分野がありました。

コピーライティングです。

JasperやCopy.aiといったスタートアップが登場しました。彼らはGPT-3を使ってマーケティング文章を自動で生成しました。ブログ記事、メールの件名、商品説明、SNS投稿。いくつかのキーワードを入力するだけで、さまざまなバリエーションの文章が出てきました。

マーケターたちに好評でした。

なぜコピーライティングだったのでしょうか。この分野では創造性が求められますが、事実確認はそれほど重要ではありませんでした。多少誇張されていたり間違っていたりしても大きな問題にはなりませんでした。そして、素早くたくさん作ることが重要でした。

コード生成も得意でした。

MicrosoftのGitHubは「GPT-3のコード版であるCodex」を使って「GitHub Copilot」を開発しました。開発者がコメントを書くと、コードを自動補完してくれました。プログラマーの生産性が大幅に上がりました。

アルトマンはAPIを通じて重要なことを学びました。

技術がどれほど優れていても、インターフェースが良くなければ革新は起きません。APIは開発者向けのツールでした。一般の人々がGPT-3の力を実感するには、距離がありすぎました。

ブレイグラウンドのログに奇妙なパターンが見えてきました。

人々がGPT-3と「会話」していました。質問して、答えを聞いて、また質問して。友人と話すように。

これがChatGPTへとつながりました。

APIは限定的な成功でした。しかし失敗ではありませんでした。次のステップへの学びの過程でした。コピーライティングで成果が出たことは、GPT-3が実際にお金になる問題を解決できる証拠でした。

人々はAIと「会話」し始めていた。

プレイグラウンドはもともと、開発者がテストするための場所でした。

大きなテキストボックスがあり、文章を入力するとGPT-3が続きを生成してくれました。開発者たちはAPIを使う前にここで実験しました。温度やトークン数といったパラメータを変えながら、出力がどう変わるかを確かめていました。

ところが、奇妙なことが起き始めました。

人々がプレイグラウンドでGPT-3と雑談を始めていたのです。

「やあ、君は誰？」「私はGPT-3です。何かお手伝いできますか？」「フランスの首都はどこ？」「パリです。」「パリで行くべき場所を教えてください。」

それは会話でした。単に情報を得るのではなく、言葉をやりとりしていたのです。ユーザーたちはGPT-3を機械としてではなく、話し相手として扱っていました。

OpenAIの研究チームはログを分析しました。

驚くべき事実が浮かび上がりました。人々は数分ではなく、何十分もかけて会話していました。単純な質問と回答を超えて、ストーリーを作り、アドバイスを求め、冗談を交わしていました。悩みをGPT-3に打ち明けた人さえいました。

当時の会話機能はひどいものでした。

アルトマン自身も認めています。「ひどかった」と。OpenAIはまだ、モデルを会話向けに調整する方法を知りませんでした。RLHF（人間のフィードバックによる強化学習）をきちんと適用する前の話です。

それでも人々は気に入っていました。

「とにかく、みんなそれを好きだった」

アルトマンは悟りました。これこそが本質だと。

人々が求めていたのは情報ではありませんでした。会話したかったのです。考えを整理し、説明を聞き、アドバイスをもらい、問いかけながら学びたかった。人間の本能でした。

言語は人間の思考そのものです。言葉にすることで考えがまとまる。GPT-3はその本能にぴたりと合っていました。

アルトマンはチームにこう告げました。

「人々がモデルと会話し始めていることは、もう明らかだ。会話できる製品を作ろう。」

これがChatGPTの始まりでした。

ベースにしたのはGPT-3.5です。RLHFを適用しました。人間の評価者がモデルの回答を見て採点し、どの答えが優れているかをランク付けしました。モデルはそのフィードバックから学び、少しずつ良くなっていきました。

目標はシンプルでした。

誰でも簡単にアクセスできること。ログインすればすぐに会話が始まること。技術的な知識は不要。自然な言葉で質問できること。会話の履歴が見やすいこと。

プレイグラウンドで見つけた小さな行動パターンが、世界を変える製品の種になりました。

人々はすでに、GPT-3に何を求めているかを示していました。OpenAIはそのシグナルをしっかりとつかめました。

DALL-E：想像を絵にする魔法

2021年初頭、OpenAIは別のものを発表しました。

DALL-Eでした。

「アボカドで作ったアームチェア」といった文章を入力すると、本当にアボカドのような形をした椅子の絵が出てきました。「月の上を歩く二足歩行のゾウ」と入力すれば、そんな奇妙な絵が生成されました。

人々は衝撃を受けました。

AIが絵を描く？想像したものをそのまま絵にしてくれる？まるで魔法のようでした。

DALL-EはGPTと同じくトランスフォーマー構造を採用しました。ただし生成するのはテキストではなく画像です。仕組みはこうです。画像を小さなピースに分割します。各ピースをトークンに変換します。テキストと画像のトークンをまとめて学習します。

だからテキストを入力すると画像が出てきたのです。

アルトマンはDALL-Eに未来を見ました。

「AIは言語だけを扱ってはいけいない。世界はテキストだけでできていない。視覚もあれば、音もあり、行動もある。」

マルチモーダル。複数の形式のデータを扱うAI。それこそが真の知性へと続く道でした。

DALL-Eは創作の民主化を示しました。

絵が描けない人でも、アーティストになれる時代が来ました。言葉が使えれば十分でした。デザイナーでなくても、素晴らしい画像を作れました。これは文化的な変化でした。

2022年、DALL-E 2が登場しました。

格段に良くなりました。画像はより鮮明でリアルになりました。DALL-E 2は、CLIPという別のモデルを使いました。CLIPは画像とテキストの関係を深く理解していました。そしてディフュージョンモデルで実際の画像を生成しました。

言語モデルも進化し続けました。

ファインチューニング (fine-tuning) が重要になりました。GPT-3は汎用モデルでした。何でもある程度こなせましたが、特定のタスクを完璧にこなすには調整が必要でした。

法律文書の分析に特化したモデル。医療記録の要約に特化したモデル。特定の企業のトーンで文章を書くモデル。

ファインチューニングを通じて、GPTは各業界の専門家になることができました。

アルトマンはこう考えていました。

「汎用モデルは基礎体力だ。ファインチューニングは、その体力を特定の能力に変えることだ。」

DALL-Eの成功はChatGPTにも影響を与えました。人々はAIが作った創作物を体験しました。驚きました。「AIってこんなこともできるの？」その驚きが、ChatGPTを受け入れる心理的な土台になりました。

DALL-Eは言語を超えた最初の一步でした。その後、GPT-4はテキストと画像の両方を理解するマルチモーダルモデルになりました。

ChatGPTのリリース：歴史を変えた5日間

2022年11月30日、水曜日。

OpenAIは静かにウェブサイトの一つ公開しました。名前はChatGPT。大きな宣伝もありませんでした。Twitterにリンクの一つ投稿しただけでした。「新しいチャットボットを試してみてください。」

アルトマンは期待していましたが、これほどとは思っていませんでした。

数万人が押し寄せました。サーバーは悲鳴を上げました。24時間で数十万人。5日間で100万人。

世界は衝撃を受けました。

人々はChatGPTにあらゆる質問を投げかけました。

「フランス革命を10歳の子どもの説明して。」「Pythonで二分探索のコードを書いて。」「このメールをもっと丁寧に直して。」「宇宙はなぜ存在するの？」

ChatGPTは答えました。自然に。親切に。ときには驚くほど正確に。

ソーシャルメディアが爆発しました。

人々はChatGPTが書いた詩、エッセイ、コード、ビジネスプランをシェアしました。「見て、AIが宿題を全部やってくれたよ！」「ChatGPTがバグを直してくれた！」「これ、ほとんど人間じゃないか？」

興奮した人もいました。「Googleの時代は終わった！」心配した人もいました。「仕事が全部なくなるぞ。」先生たちは困惑しました。「生徒がこれで宿題をしたらどうするんだ？」

わずか2か月で1億人に達しました。

史上最速で成長したアプリでした。Netflixが1億人に達するまで10年かかりました。Facebookは4年半。Instagramは2年半。ChatGPTは2か月。

なぜこれほど速かったのでしょうか？

あまりにも簡単でした。複雑な設定はありませんでした。コマンドを覚える必要もありませんでした。ただ話しかければよかったのです。おばあちゃんでも使えました。

即座でした。質問すればすぐに答えが返ってきました。Googleのようにリンクを10個渡すのではなく、答えを直接くれました。

楽しかったのです。AIと会話するというのが不思議でした。癖になりました。

教育界がひっくり返りました。

生徒はエッセイをChatGPTで書きました。教師は評価方法を変えざるを得ませんでした。ChatGPTを禁止した学校もありました。逆に積極的に教えた学校もありました。

開発者たちの生産性は爆発的に上がりました。

コードを書く時間が半分に減りました。バグ探しにChatGPTを使いました。新しい言語を学ぶのがずっと楽になりました。

クリエイターたちは新しいツールを手に入れました。

小説家がアイデアをブレインストーミングしました。作曲家が歌詞を磨きました。マーケターが広告コピーを作りました。

アルトマンはリアルタイムで見守っていました。

彼はチームに言いました。「僕たちは、もう後戻りできないものを作ってしまったかもしれない。」

恐れと誇りが入り混じった言葉でした。

ChatGPTは完璧ではありませんでした。誤った情報を自信満々に述べることもありました。幻覚現象もありました。偏りもありました。それでも、人々は受け入れました。

政府も動き始めました。

各国がAI規制の議論を始めました。欧州はAI法案を急ぎました。米国議会はサム・アルトマンを公聴会に呼びました。

ChatGPTはただの製品ではありませんでした。

それは人工知能時代の正式な幕開けでした。AIはもはや研究室の実験ではありませんでした。人々の日常に入り込んできたのです。学生、会社員、おばあちゃん、子供たちまで。誰もがAIと会話しました。

アルトマンの夢が現実になりました。

知性が誰にでも開かれました。誰もが世界最高水準のAIを無料で使えるようになりました。これは革命でした。

2022年11月30日。

歴史が変わった日でした。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

2.3 イーロン・マスクとの対立

第2部 OpenAIの誕生と成長

← 前の章

次の章 →

過半数の株式を求めたマスク

2015年12月、サム・アルトマンはイーロン・マスクとともに、OpenAIという新たな挑戦を始めました。

マスクはテスラとSpaceXですでに世界的な名声を得た億万長者でした。アルトマンはYコンビネーターを率い、シリコンバレーの「キングメーカー」と呼ばれた若きリーダーでした。二人はAIが人類にとって大きな脅威になり得るという点で意見が一致していました。GoogleがディープマインドというAI企業を買収し、AI技術を独占してしまうかもしれないという恐れがあったのです。

マスクとアルトマンは「AIは一部の巨大企業に独占されてはならない」という考えを共有していました。そこでOpenAIを非営利団体として設立しました。名称に「オープン（Open）」が入っているのにも理由があります。研究成果をすべて公開し、誰もがその恩恵を受けられるようにするという約束でした。

マスクはOpenAIに多額の資金を寄付すると約束しました。当初の約束は10億ドルという巨額でした。実際には2016年から2018年にかけて約4,500万ドルを寄付しました。OpenAI創設期の最大の支援者だったことになります。

最初の2年間は順調でした。

OpenAIの研究者たちはビデオゲームをプレイするAIを開発したり、ロボットアームでルービックキューブを解く実験を行ったりしていました。非営利団体らしく、研究の自由が尊重された雰囲気でした。マスクも頻繁にOpenAIのオフィスを訪れ、研究者たちを励ましていました。サム・アルトマンはYコンビネーターの仕事とOpenAIの仕事を掛け持ちして忙しい日々を送っていましたが、二人の創業者の関係は良好でした。

問題は2017年後半から始まりました。

OpenAIの研究チームはGPT-1という言語モデルの可能性を見出しました。モデルを大きくすれば驚くべきことができそうでした。サム・アルトマンとイリヤ・サツキバーが試算してみると、本当に強力なAIを作るには膨大なコンピュータリソースが必要でした。数十億ドル規模の投資が求められたのです。

イリヤ・サツキバーはイーロン・マスクにメールを送りました。「AGIに近づくにつれて、情報公開の範囲を絞るのは合理的です。OpenAIの『オープン』とは、AIが完成した後に誰もがその恩恵を受けるべきだという意味であって、研究の途中で科学的知見を共有しなくても一向に構いません。」マスクは「そのとおりです」と返信しました。

2017年末、OpenAIの創業者たちとマスクは、次のステップとして営利企業を設立する必要があるという点で合意しました。非営利の体制ではGoogleやFacebookのような巨大テック企業と競争できなかったからです。問題はその先にありました。

マスクはOpenAIの過半数株式を自分が保有すべきだと主張しました。さらに取締役会を掌握し、CEOに就任すると告げました。彼の論理はこうです。AGIは人類史上最も危険な技術です。大勢で議論しては意思決定が遅くなり、Googleに後れを取るだけです。強力な一人のリーダーシップが必要だというのでした。

マスクはOpenAIをテスラと合併させようという提案もしました。テスラの自動運転技術がすでに膨大なデータを収集しているため、OpenAIのAI研究と組み合わせればシナジーが生まれると説明しました。2018年2月初め、マスクはアルトマンにメールを送りました。「OpenAIはテスラの『資金源』として組み込まなければなりません。テスラだけがGoogleに対抗できる唯一の道です。確率は小さいですが。」

サム・アルトマンと他の共同創業者たちは、この提案を受け入れることができませんでした。

OpenAIのチームはマスクにこう伝えました。「現在の構造は、AGIに対する一方的かつ絶対的な支配権をあなたに与える経路を提供しています。あなたは最終的なAGIを支配したくないとおっしゃっていましたが、今回の交渉の過程で、絶対的な支配があなたにとっていかに重要かが明らかになりました。」アルトマンは心の中でこう考えていました。『私たちがOpenAIを作った理由は何だったか。AIの独占を防ぐためではなかったか。それなのに、今マスクが求めているのはまさにその独占ではないか。』

OpenAIのチームはさらにこう付け加えました。「OpenAIの目標は、未来をより良くし、AGIによる専制支配を回避することです。会社が本当にAGIに向けて大きな進歩を遂げた場合、現在の意図とは異なり、あなたが会社の絶対的な支配を手放さない選択をすることを懸念しています。」

サム・アルトマンには別の理由もありました。彼はマスクの下で働きたくなかったのです。マスクは天才的な実業家でしたが、気まぐれな上司としても知られていました。昼夜を問わず部下に連絡を取り、突然方針を変えることもありました。アルトマンがOpenAIに集めた優秀な研究者たちも、マスクの下では働きたくないだろうと思っていました。

交渉は決裂しました。

マスクは短いメールでこう返答しました。「議論は終わりです。はっきり言っておきますが、これは以前に話し合っていた条件を受け入れるという最後通牒ではありません。あの話はもうテーブルの上にありません。」そして数週間後の2018年2月、マスクはOpenAIの取締役会を辞任しました。

去り際にマスクは「OpenAIが成功する確率はゼロです」と言い残しました。そして自分はテスラの内部でAGIの競合となるものを作るつもりだと明言しました。

マスクの公式な退任理由は「利益相反」でした。テスラが自律走行用のAIを開発しているため、OpenAIと人材をめぐって競合することになるという説明でした。実際、OpenAIの優秀な研究者アンドレイ・カルパシーがテスラに移籍しました。マスクがOpenAIのオフィスで別れの挨拶をしたと

き、彼は主に利益相反の問題について話しました。しかしほとんどの社員は、その説明を額面通りには受け取りませんでした。

本当の理由は支配権の問題でした。

サム・アルトマンはほどなくリード・ホフマンから慌てた電話を受けました。「イーロンが支援を打ち切りました。どうすればいいでしょう？」マスクは約束していた10億ドルの寄付を止めました。すでに1億ドルは寄付していましたが、残りは支払いませんでした。

OpenAIは突然、財政難に陥りました。研究者たちに給料を払う資金も不足していました。コンピュータを動かすコストはなおさらありませんでした。リード・ホフマンが急遽1,000万ドルを追加出資して危機を乗り越えました。

サム・アルトマンは岐路に立たされました。マスクなしに数十億ドルをどうやって集めるのか。非営利の体制を維持しながら巨大な資本を呼び込むことは不可能でした。

2019年3月、OpenAIは「収益上限付き営利法人」という独自の構造を作りました。投資家は利益を得られますが、利益が一定の上限を超えると残りは非営利財団に戻る仕組みでした。

2018年12月、マスクはOpenAIに最後の言葉を残しました。「毎年すぐさま数十億ドルを集めるか、さもなければ諦めるかです。」

マスクの警告は間違っていないでした。AI開発には本当に天文学的な資金が必要でした。しかしアルトマンは、マスクのやり方ではなく、自分のやり方でその資金を集めることを決意しました。

この決別はOpenAIの歴史において最も重要な転換点となりました。マスクという強力な後援者を失いましたが、アルトマンは自分のビジョン通りに会社を率いる自由を手に入れました。その自由こそが、後にGPT-3、そしてChatGPTという革命を生み出す土台になったのです。

分かれた二つのビジョン

イーロン・マスクはOpenAIを去った後も、沈黙を守りませんでした。

むしろ彼は声を上げ続けました。OpenAIがマイクロソフトから多額の投資を受け、ChatGPTで大きな成功を収めると、マスクの批判はいっそう激しくなりました。

2022年11月にChatGPTがリリースされたとき、マスクは激怒したと伝えられています。自分が去った会社が世界を変える製品を作ったことも腹立たしかったのですが、より大きな問題はその方法でした。ChatGPTが登場した1か月後の2022年12月、マスクはOpenAIのTwitterデータへのアクセスをブロックしました。当時マスクはTwitter（現在のX）を買収した状態でした。

2023年2月17日、マスクはTwitterにこう書きました。「OpenAIは、Googleに対抗するためのオープンソース非営利企業として設立されました。今や、マイクロソフトが事実上支配する閉鎖的な最大利益追求企業になっています。」

マスクの最初の批判は、「オープン（Open）」という名前に関するものでした。

マスクは、OpenAIが「偽の人道主義的使命」を掲げて自分や他の投資家たちを引き込んだと主張しました。当初は、すべての研究成果を公開すると約束していました。誰でもOpenAIの技術を見られ、使えるようにするはずでした。GPT-1とGPT-2までは、実際にコードを公開していました。

GPT-3から状況が変わりました。OpenAIはモデルの詳細を公開しなくなりました。GPT-4ではさらに徹底していました。どのように作られたか、どんなデータを使ったか、モデルの構造はどうなっているかについて、ほとんど何も明かしませんでした。OpenAIは「安全のため」と説明しました。悪意ある人々がこの技術を悪用しかねないという理由でした。

マスクはこれが言い訳に過ぎないと考えていました。本当の理由は、競合他社がOpenAIの技術を真似できないようにするためだという主張でした。人類の安全のためではなく、商業上の秘密を守るためだというわけです。

マスクはTwitterにこうも書きました。「私が約1億ドルを寄付した非営利団体が、どうして300億ドルの価値を持つ営利企業になったのか、いまだに理解できません。これが合法なら、なぜ誰もがそうしないのでしょうか？」

2つ目の批判は、マイクロソフトとの関係についてでした。

マスクは、OpenAIが「マイクロソフトの閉鎖的な事実上の子会社」になったと主張しました。2019年以降、マイクロソフトはOpenAIに総額130億ドル（韓国ウォン換算で約17兆ウォン）を投資しています。これだけの巨額の資金を受け取れば、独立性を保つのは難しくなります。

マスクは、マイクロソフトCEOのサティア・ナデラのインタビューを引用しました。2023年にサム・アルトマンが解任され、その後復帰した際、ナデラはこう語っていました。「OpenAIが明日消えたとしても、私たちはすべての知的財産権とすべての能力を持っています。人材もいる、コンピューティングもある、データもある、すべてを持っています。」マスクは、この発言こそOpenAIが実際に誰のために働いているかを示していると主張しました。

2024年3月1日、マスクはサンフランシスコの裁判所に、OpenAI、サム・アルトマン、グレッグ・ブロックマンを相手取って訴訟を起こしました。訴訟でマスクは、OpenAIが「契約違反」「誠実義務違反」「不公正取引行為」を犯したと主張しました。

83ページにわたる訴訟書類の中で、マスクの弁護士マーク・トベロフはこう述べました。「この訴訟は、被告たちの大げさな慈善活動の実態を明らかにし、マスクと一般市民に対する彼らの虚偽の陳述について責任を問うものです。これは単に1,000億ドル規模のスタートアップに関する問題ではありません。AGIの未来が危ぶまれているのです。」

マスクは訴訟の中で、GPT-4がすでにAGI（汎用人工知能）に到達していると主張しました。AGIとは、人間と同等かそれ以上の知能を指します。OpenAIとマイクロソフトの契約には、AGIが完成した場合にマイクロソフトの独占ライセンスが終了するという条項がありました。

マスクは、OpenAIが財務上の利益のためにGPT-4がAGIであるという事実を隠していると非難しました。

マスクは訴訟を通じて、3つのことを求めました。

一つ目は、OpenAIに本来の使命に戻ることを強制すること。二つ目は、マイクロソフトを含む誰もOpenAIの技術で利益を得られないようにすること。三つ目は、自分が寄付した資金が公益ではなく私的な利益のために使われていたならば、返還を求めること。

OpenAIは強く反論しました。

OpenAIは公式ブログで複数のメールや文書を公開し、マスク自身が営利転換を望んでいたという証拠を示しました。2017年9月15日、マスクは「Open Artificial Intelligence Technologies, Inc.」という公益法人を自ら設立してさえいました。OpenAIは「マスクは過半数の株式と完全な支配権を得られなかったため去り、私たちが失敗するだろうと言った」と主張しました。

OpenAIはこう付け加えました。「訴訟でAGIを作ることはできません。私たちはイーロンの功績を敬い、初期の貢献に感謝していますが、彼は法廷ではなく市場で競争すべきです。」

OpenAIは最後にこう述べました。「私たちにより高い目標を持つよう鼓舞してくれた人が、後になって私たちは失敗すると言い、競合企業を立ち上げ、彼なしで私たちが意義ある進歩を遂げると訴訟を起こしたことを、残念に思います。」

マスクは2024年6月に訴訟を取り下げました。理由は明かしませんでした。

その2カ月後の8月、彼は再び訴訟を起こしました。今度は「組織犯罪（RICO法違反）」の容疑まで加えました。2024年11月にはマイクロソフトも被告に加えました。OpenAIとマイクロソフトの関係が反競争的で独占的だというのが主張でした。

2025年3月、裁判所はマスクによる仮差し止め命令の申請を却下しました。OpenAIの営利転換を差し止めるよう求めた申請でしたが、裁判所はこれを退けました。イヴォン・ゴンザレス・ロジャース判事は、マスクが提出した証拠だけでは仮差し止め命令を下すには不十分だと判断しました。一部のメールを見ると、マスク自身もOpenAIがいつか営利企業になり得ると考えていたようだと指摘しました。

裁判所は、2025年秋に迅速な裁判を行う準備があると表明しました。この法的な争いはまだ終わっていません。

マスクの批判は本当に正当なものなのでしょうか。

マスクが本当に人類の安全を心配しているのだと見る人もいます。OpenAIが当初約束していた透明性と公益性を捨てたという彼の指摘には、一理あります。実際にOpenAIは、非営利から営利へ、オープンからクローズドへと大きく方向を転換しました。

一方で、マスクが嫉妬しているのだと考える人もいます。自分が支配権を持てず去った会社が大きな成功を収めたことへの妬みではないかという見方です。マスクは2023年に自身のAI企業xAIを設立しました。2025年2月には、プライベートエクイティファンドと組んで974億ドル（韓国ウォン換算で約130兆ウォン）でOpenAIを買収するという提案まで行いました。OpenAIの取締役会は全会一致でこれを拒否しました。

サム・アルトマンは、マスクに対して感情的に反応しないよう努めています。CNBCのインタビューでアルトマンはこう語りました。「正直、彼のことをそれほど考えているわけではありません。一日中Twitterで、OpenAIがいかにかにひどいか、私たちのモデルはダメだ、良い会社にはなれないとツイートしているとばかり思っていましたよ。」

二人の対立は、つまるところ哲学の違いです。マスクは、AGIのような強力な技術は強力なリーダー一人が支配してこそ安全だと信じています。アルトマンは、一人の人間にそれほどの権力が集中することこそ危険だと考えています。

どちらが正しいかは歴史が判断するでしょう。今確かなのは、この対立がAI産業全体を変えたという事実です。マスクが去ったからこそ、サム・アルトマンは自分のやり方でOpenAIを率いることができました。その結果がChatGPTであり、世界中が体験したAI革命です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

3.1 CEO解任、5日間の悪夢

第3部 解任と復帰、ビジョンの勝利

← 前の章

次の章 →

突然の解雇通告

2023年11月17日、金曜日の午後、サム・アルトマンの携帯電話が鳴りました。

ビデオ会議に参加するよう求められました。数分後、彼はOpenAIのCEOではなくなったと告げられました。取締役会は、彼が「一貫して誠実ではなかった」と述べました。具体的な理由はありませんでした。事前の警告もありませんでした。ただ、あなたは解雇されたという言葉だけでした。

サム・アルトマンは後にその瞬間をこう表現しました。「奇妙な熱病にうなされる夢のようでした。」ChatGPTが世界を揺るがしていたまさにその瞬間に、その会社を作った人物が会社から追い出されたのです。シリコンバレー史上、最もドラマティックな出来事でした。

最初に感じたのは混乱でした。サムは前日まで会社の未来について話し合っていました。新しいモデルの開発計画を立てていました。突然、すべてが崩れ落ちました。何が起きたのか理解できませんでした。「なぜ？」という問いが頭から離れませんでした。

取締役会は明確な答えを与えませんでした。コミュニケーションの問題？それが何を意味するのか分かりませんでした。混乱はすぐに挫折感へと変わっていきました。

夢の崩壊

サム・アルトマンはOpenAIをただの職場だと思ったことは一度もありませんでした。

この会社は彼の夢でした。AGIを作って人類を助けるという使命でした。昼夜を問わず働きました。投資家を説得し、研究者を集め、政府と交渉しました。ChatGPTが成功したとき、彼は自分の人生で最も輝かしい瞬間だと感じました。それなのに、自分が積み上げてきたすべてのものから追い出されたのです。

挫折感はすぐに怒りへと変わりました。解任の決定は手続き上も問題だらけでした。最大の投資家であるMicrosoftでさえ、数分前に通知を受けただけでした。社員たちはブログの投稿を見て状況を知りました。取締役会はなぜこのような進め方をしたのでしょうか。サムは裏切られたと感じました。

とりわけ、共同創業者のイリヤ・サツケヴァーが解任に賛成したという事実は衝撃的でした。ともに夢を育ててきた仲間が背を向けたという思いに心が痛みました。怒りの中にも、深い悲しみが押し寄せてきました。

サムはOpenAIで過ごした時間を愛していました。ともに働いた人々を愛していました。深夜までコードを修正していたエンジニアたち、モデルを訓練していた研究者たち、安全性について考えていたポリシーチーム。この人々すべてと別れなければならないという思いが、胸を締め付けました。SNSに短い言葉を投稿しました。「OpenAIで過ごした時間を愛していました。」悲しみのにじむ一文でした。

予期せぬ逆転

共同創業者のグレッグ・ブロックマンが辞任を表明しました。「サムのいないOpenAIに意味はない」と言いました。数時間後、主要な幹部たちから電話がかかってきました。「私たちはあなたを支持しています。」研究者たちもメッセージを送ってきました。「あなたに戻ってきてほしい。」悲しみの中に、感謝の気持ちが芽生え始めました。

翌日、さらに驚くべきことが起きました。

OpenAIの社員770人のうち700人以上が公開書簡に署名しました。内容は明快でした。「サム・アルトマンが復帰しなければ、我々全員がMicrosoftに移ります。」サムはこの書簡を見て胸が熱くなりました。自分のために会社全体が立ち上がったのです。これは義理や忠誠心ではありませんでした。信頼でした。ともに作り上げてきたビジョンへの信念でした。

MicrosoftのCEO、サティア・ナデラが公に支持を表明しました。「サムと彼のチームのために新しいAI研究所を設立します。」世界最大のテクノロジー企業がサムの側についたのです。感謝の気持ちが、あらゆる否定的な感情を上回り始めました。

サム・アルトマンは混乱、挫折、怒り、悲しみを経て、感謝にたどり着きました。自分が一人ではないと気づきました。700人の仲間がいました。巨大なパートナーがいました。全世界が見守っていました。後に彼はこう語りました。「数日の間に、人間が感じるあらゆる感情を経験しました。混乱から始まって、感謝で終わりました。」

この感情の旅が彼をさらに強くしました。危機は、自分が積み上げてきたものがいかに強固かを示してくれました。人々との信頼、共有されたビジョン、ともに育んだ文化。そのすべてが、一個人よりも強かったのです。

サム・アルトマンは追放を通じて、逆説的に自分の真の力を確認しました。その力は権力ではなく、信頼から生まれていました。

5日間の戦い

金曜日の午後、サム・アルトマンはCEOを解任されました。水曜日の夜明け、彼はCEOに復帰しました。わずか5日間のことでした。

その5日間は、OpenAIの歴史において、いやシリコンバレー全体の歴史において、最もドラマティックな時間でした。時間が速く過ぎたわけではありません。あまりに多くのことが凝縮して起きたのです。すべてがリアルタイムで展開されました。

金曜日：爆弾宣言

11月17日の午後、OpenAIのブログに短い投稿が掲載されました。「サム・アルトマンがCEOを退任しました。暫定CEOとしてミラ・ムラティCTOを任命します。」シリコンバレーが凍りつきました。ChatGPTで世界を変えているその会社の創業者が、何の警告もなく解雇されたのです。

数時間後、グレッグ・ブロックマンが辞任を発表しました。「私はサム・アルトマンとともにOpenAIを作りました。彼のいないOpenAIには留まりません。」この発表は、事態がただの経営陣交代ではないことを示す信号でした。

週末：反乱の始まり

土曜日と日曜日、OpenAIの内部は爆発寸前でした。Slackチャンネルで数百件のメッセージがやり取りされました。「一体何が起きているんだ？」「取締役会はどうかしている。」「私たちはどうすればいい？」主要な幹部たちが緊急会議を招集しました。

研究者たちが集まりました。エンジニアたちが意見をまとめました。結論は一つでした。「サムが戻らなければならない。」日曜日の午後、サム・アルトマンはOpenAI本社に来訪者バッジをつけて現れました。写真を撮ってSNSに投稿しました。「このバッジを初めて、そして最後につけています。」

メッセージは明確でした。私は退かない。投資家たちが動き始めました。スライブ・キャピタル、コスラ・ベンチャーズ、タイガー・グローバルが取締役に圧力をかけました。「アルトマンを復帰させろ。」取締役会は応答しませんでした。

月曜日：戦線布告

11月20日、事態が爆発しました。

OpenAIの取締役会は新たな暫定CEOを発表しました。元TwitchのCEO、エメット・シアーでした。取締役会のメッセージは明確でした。「サム・アルトマンは戻らない。」社員たちはすぐに反応しました。午後に公開書簡が発表されました。

770人の社員のうち700人以上が署名しました。研究者、エンジニア、ポリシーチーム、さらには人事部や経理部まで参加しました。会社のほぼ全員が一つの声を上げたのです。「サム・アルトマンとグレッグ・ブロックマンが復帰しなければ、私たち全員が辞職してMicrosoftに移ります。取締役会は退くべきです。」

この書簡は、企業の歴史においてほぼ前例のない集団行動でした。社員の95%が経営陣の側を選んだのです。取締役会を拒否したのです。

同じ日、マイクロソフトCEOのサティア・ナデラがツイッターに投稿しました。

「サム・アルトマンとグレッグ・ブロックマンを、我が社の新しいAI研究チームのリーダーとして迎えます。OpenAIから一緒に来たい方は誰でも歓迎します。」これは単なる採用告知ではありませんでした。取締役会への最後通牒でした。「サムを返さなければ、会社ごと奪い取る。」

火曜日：崩れゆく取締役会

11月21日、取締役会の立場は完全に崩れました。新たに任命されたばかりの暫定CEOエメット・シアーでさえ、取締役会に要求しました。「解任の理由を書面で提出しなければ、私も辞任する。」追放を主導していたイリヤ・スツケヴェルがツイッターに謝罪文を投稿しました。「私が関わった取締役会の行動を深く後悔しています。私が愛するOpenAIを傷つけようとしたわけではありませんでした。」

投資家からの圧力は頂点に達していました。法律顧問たちが警告しました。「このままでは会社が崩壊します。」取締役会はもはや持ちこたえられませんでした。交渉が始まりました。

水曜日の早朝：勝利宣言

11月22日の早朝、OpenAIが公式声明を発表しました。「サム・アルトマンがCEOとして復帰します。新しい取締役会を組織します。ブレット・テイラー（元セールスフォースCEO）が議長を務め、ラリー・サマーズ（元財務長官）、アダム・ディアンジェロ（QuoraCEO）が参加します。」サム・アルトマンがツイッターに投稿しました。「OpenAIを愛しています。この数日間に行ったことは、すべてこのチームとその使命を守るためでした。さあ、また仕事を始めます。」

5日間の悪夢が終わりました。

夢のように過ぎた時間

サム・アルトマンは後に、この期間を「奇妙な熱にうなされる夢」と表現しました。すべてが現実離れしていました。時間の感覚が消えました。一日が一週間のように感じられました。こう振り返っています。「朝目を覚ますと、状況がまったく変わっていました。誰が何を言ったか、誰が寝返ったか、どんな交渉が進んでいるか。すべてがリアルタイムでツイッターで実況されていました。」

全世界が見守る中、運命がソーシャルメディアで決まりました。

取締役会の決定、従業員の連署状、CEO交代、投資家の圧力、そして最終的な復帰まで。すべてがオンラインでリアルタイムに展開されました。この出来事は、コーポレートガバナンスの新たな歴史を刻みました。取締役会がCEOを解任したにもかかわらず、従業員たちがCEOを連れ戻したので。公式な権力構造が崩れ、実質的な影響力が勝利しました。

5日間でした。OpenAIのあり方を根底から再定義した5日間でした。

混乱の中でも完璧に機能し続けた会社

サム・アルトマンに、この出来事を通じて最も誇りに思った瞬間は何かと聞くと、答えは意外なものです。

「自分が復帰したこと？」ではありませんでした。

「従業員たちの支持？」それも違いました。

サムが最も誇りに思った瞬間は、これでした。「自分がいない間、会社が完璧に機能し続けたこと。」

CEOが解雇され、取締役会が混乱に陥り、従業員たちが集団行動の準備を進める状況でした。普通の会社なら業務が麻痺していたでしょう。顧客サービスが止まり、開発が停止し、システムが崩壊していたはずです。

ところがOpenAIは違いました。

ChatGPTは止まらなかった

その5日間、世界中の数億人がChatGPTを使いました。質問を投げかけ、答えを受け取り、コードを書き、翻訳を依頼しました。サービスは一度も止まりませんでした。

APIを使う数千社のサービスも動き続けました。カスタマーサポートのチャットボット、コーディングアシスタント、翻訳サービス。すべてが正常に稼働していました。

インフラチームがシステムを守りました。エンジニアたちがサーバーを監視し続けました。誰もパニックに陥りませんでした。

研究は続いた

研究者たちは実験を止めませんでした。新しいモデルを学習させ、論文を書き、データを分析しました。

ある研究者は後にこう語りました。「サムが戻ってくることはわかっていました。それまでの間、私たちがすべきことは最高の研究を続けることでした。」

幹部チームが危機を管理した

CTOのミラ・ムラティが暫定CEOとなりました。彼女は冷静に状況を管理しました。毎朝全体会議を開き、従業員に状況を説明し、不安を鎮めました。

COOのブラッド・ライトキャップが外部とのコミュニケーションを管理しました。投資家たちと電話会議を行い、パートナー企業を安心させました。「OpenAIは揺るぎません。」

最高戦略責任者のジェイソン・クオンが従業員の連署状をまとめました。法務チームが法的な選択肢を検討しました。人事チームが従業員のメンタルケアに当たりました。

全員が自分の役割を完璧に果たしました。

サムが見たもの

サム・アルトマンは外からこのすべてを見ていました。CEOではない立場で、自分が作った会社がどのように動いているかを目の当たりにしました。

彼は驚きました。

会社は彼を必要としていませんでした。正確に言えば、彼が作り上げたシステムが完璧に機能していたのです。自律的で、責任感があり、有能な人々が、それぞれ全力を尽くしていました。

サムはこう語っています。「私が最も誇りに思った瞬間は、経営幹部の誰もがこの会社を完璧に運営できると確認できたときでした。」

リーダーの本当の成功

多くのCEOは、自分がいなければ会社が崩壊すると考えています。だからこそ、すべての決断を自分で下します。すべての会議に出席します。すべての問題を自ら解決します。

サム・アルトマンは異なる哲学を持っていました。「本当に賢い人々に大きな責任を与える」という形で会社を経営していました。少数の人材に大きな権限を与えていました。マイクロマネジメントはしませんでした。

この危機は、彼の哲学が正しかったことを証明しました。

組織のレジリエンス

心理学者たちは「レジリエンス（resilience）」という概念を説きます。危機が訪れたとき、どれだけ早く立ち直れるかということです。

OpenAIは驚くべきレジリエンスを見せました。CEOが解任されるという最悪の危機の中でも、中核的な機能は止まりませんでした。組織はパニックに陥りませんでした。むしろ、より一層固く結束しました。

それが可能だった理由は何だったのでしょうか。

強いビジョンがありました。すべての社員が「AGIで人類を助ける」というミッションを共有していました。CEOが代わっても、ミッションは変わりませんでした。

信頼がありました。社員たちは互いを信じていました。経営幹部を信じていました。それぞれが全力を尽くすことを、皆が知っていました。

自律性がありました。人々は誰かの指示を待ちませんでした。自分で判断し、行動しました。

サム・アルトマンは、この三つを作り上げたのです。

「これが私の選んだ人たちだ」

サムは復帰後のインタビューでこう語っています。

「私がこの人たちを選んだということ、そして私がある程度教えてきたことが彼らの中に宿っていると確認できたとき、本当に誇らしかったです。危機的状況で彼らが見せた能力は、目を見張るものがありました。」

彼は自分の成功を一人で成し遂げたとは思っていませんでした。彼はチームを作り、そのチームが彼を救ったのです。

これがリーダーシップの逆説です。本当に強いリーダーとは、自分がいなくても組織が回るようにする人のことです。

5日間の危機を通じて、彼こそが本物のリーダーであることを証明しました。彼が作り上げた組織は、彼なしでも完璧に機能し、だからこそ彼は必要不可欠な存在でした。これが彼が最も誇りに思った瞬間です。自分の復帰ではなく、自分が育てたチームの力でした。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

3.2 社員たちの劇的な忠誠

第3部 解任と復帰、ビジョンの勝利

← 前の章

次の章 →

770人の社員のうち大多数が署名した支持宣言

2023年11月17日金曜日、サム・アルトマンがCEOの座から追放されました。

OpenAIの社員たちは衝撃を受けました。「一体何が起きているんだ？」廊下で、会議室で、Slackのメッセージで、混乱と怒りが爆発しました。これほどまで会社を成長させた人を、何の説明もなく追い出すなんて。

週末が過ぎ、月曜日の朝になりました。取締役会はエミット・シアという人物を新たな暫定CEOに任命しました。Twitchを創業したことで知られる人物です。取締役会は「これで終わった。我々の勝ちだ」と思っていたかもしれませんが。

社員たちは違いました。

その日の午後、一通の文書がOpenAI社内で回り始めました。タイトルはシンプルでした。「私たちはサムと共に行く。」内容はさらに力強いものでした。「取締役会がサム・アルトマンとグレッグ・ブロックマンを復帰させなければ、私たちは全員辞職してMicrosoftへ移ります。」

最初は数十人が署名しました。1時間後には100人を超えました。夕方になると500人を超えました。翌朝、署名者は700人を突破しました。

全社員は770人でした。つまり、95%近くの人が署名したことになります。

その中には研究員もいました。エンジニアもいました。デザイナーもいました。人事部、財務部、法務部の社員もいました。さらには、取締役会がサムを追い出した後に暫定CEOに就けたミラ・ムラティまでが署名しました。自分をCEOにしてくれた取締役会を裏切ったのでしょうか？いいえ。彼女は会社を愛し、サムがいてこそ会社が生き残れると信じていたのです。

さらに驚くべきことは、署名だけにとどまらなかった点です。

社員たちはX(旧Twitter)に同じ文句を投稿し始めました。「OpenAI is nothing without its people.」オープンAIは人なしでは何もできない、という意味です。一人ひとり、二人ひとり、数百のアカウントから同じ文章が投稿されました。まるで抗議現場のシュプレヒコールのように。

サム・アルトマンはこれらすべてのメッセージを自分のアカウントから一つひとつリツイートしました。言葉もなく、静かに、ただハートを押すだけでした。

なぜこれほど多くの人々がサム・アルトマンに従ったのでしょうか？

お金のためだったのでしょうか？ もちろん、それもありました。OpenAIの社員たちは株式のようなものは受け取っていませんでしたが、PPUという特別な権利を得ていました。会社がうまくいけば莫大な報酬を受け取れる証書のようなものです。しかしサムがいなくなれば会社が崩壊しそうであり、そうなればその証書は紙切れになってしまうところでした。

それだけではありませんでした。

サム・アルトマンは社員たちに夢を与えました。「私たちはAGIを作れる。人類の歴史を変えられる。」この言葉はほら話に聞こえるかもしれませんが、実際にChatGPTを作り上げた人々にとっては現実でした。彼らは自分たちが本当に世界を変えていると信じていました。

サムは社員たちを信じました。大きな責任を与え、自分で決断させ、速く動かせました。誰かがミスをする「大丈夫、またやってみよう」と言いました。社員たちはサムが自分たちを人間として尊重してくれると感じていました。

だからサムが追放されると、社員たちは裏切られたと感じました。

「取締役会は私たちの声も聞かずに勝手に決めた。」「私たちが作った会社なのに、私たちの意見は重要じゃないということなの？」「サムなしにAGIを作れるだろうか？ いや、サムなしにこの会社が続けられるだろうか？」

社員たちの署名は単なる抗議ではありませんでした。

最後通牒でした。

Microsoftはすでにサムとグレッグを採用すると発表していました。サティア・ナデラCEOが自ら乗り出し、「サム・アルトマンを我々の新しいAI研究所長として迎える。彼についてきたいOpenAIの社員全員を歓迎する」と言いました。

空約束ではありませんでした。MicrosoftはOpenAIに130億ドル(約17兆ウォン)を投資した会社でした。資金も豊富で、ポストも多く、インフラも十分でした。社員700人を受け入れることはMicrosoftにとって朝飯前でした。

取締役会は恐怖で震えていたことでしょう。

「700人が本当に去るというのか？」「それなら会社に残る人が70人しかいなくなる？」「70人でGPT-5を作れというのか？」不可能でした。

OpenAIは建物ではなく、人で作られた会社でした。GPTモデルを作る方法を知っている人々、サーバーを動かす方法を知っている人々、顧客を管理する人々が全員去ってしまえば、OpenAIは名前だけが残るゴースト会社になるところでした。

投資家たちも怒りました。

Sequoia Capital、Thrive Capitalといった有名な投資会社が取締役会に電話をかけました。「一体何をしでかしたんですか？」「私たちが投資した860億ドルの会社を一夜にして潰すつもりですか？」「サムを今すぐ復帰させてください!!」

SalesforceのCEO、マーク・ベニオフも名乗り出ました。「OpenAIの人材の皆さん、うちに来てください。給与もさらに上げますし、ビザの問題も全部解決します。」Google DeepMindにもOpenAI社員の履歴書が殺到し始めました。

OpenAIは瓦解寸前でした。

そしてこれらすべてが、わずか5日間で起きました。

金曜日にサムが解雇され、週末に混乱が起き、月曜日に社員たちが署名し、火曜日に取締役会が降伏しました。シリコンバレーの歴史上、最も速いCEO復帰の出来事でした。

社員たちの忠誠心はどこから来たのでしょうか？

サム・アルトマンはカリスマのある人物でした。演説が得意というわけでも、声が大きいというわけでもありませんでした。しかし人と話すとき、まっすぐ目を見て、真剣に耳を傾けました。「あなたの意見は重要です」と口で言うだけでなく、本当に社員たちのアイデアを製品に反映させました。

ChatGPTが生まれたのも、社員たちのアイデアからでした。

最初、チャットボットを作ろうと考えていた人は誰もいませんでした。ところが数人の社員が「GPT-3のプレイグラウンドで会話するのを皆さん楽しんでいるみたいです」と言い出しました。サムは「じゃあ製品にしてみよう」と答え、数週間のうちにChatGPTが世に出たのです。

サムは失敗を恐れませんでした。

「素早く作り、素早くリリースし、素早く学ぼう。」それが彼の哲学でした。完璧になるまで待つのではなく、とにかく世に出して反応を見る、というものでした。こうした文化のなかで社員たちは自由に実験し、挑戦することができました。

だからサムがいなくなるということは、CEO一人を失うだけの話ではありませんでした。

会社の文化、ビジョン、仕事のやり方すべてが崩れ落ちることを意味していました。取締役会はサム一人を追いつせば済むと思っていたのですが、社員たちは「サムがすなわちOpenAIだ」と考えていたのです。

700人の署名が会社を救いました。

正確に言えば、会社の魂を守ったのです。社員たちが黙っていれば、サムはマイクロソフトへ移っていたでしょう。OpenAIは名前だけが残し、中核を担う人材はすべて抜け出し、歴史の闇に消えていたかもしれません。

しかし社員たちは立ち上がりました。

彼らは自分たちが築いた会社を、自分たちが信じるリーダーを、自分たちが夢見る未来を守るために戦いました。そして勝ったのです。

この出来事は、21世紀の企業の新しい姿を見せてくれました。かつては取締役会がすべてを決めていました。しかし今は、会社の本当の力がどこにあるかが明らかになりました。テクノロジー企業において最も大切な資産は、建物でも、特許でも、お金でもありません。

人です。

その人たちが共に作り上げていく夢です。

取締役会の降伏とCEO復帰

月曜日の夜になると、状況は絶望的でした。

取締役会は四面楚歌に陥っていました。700人の社員が去ると宣言し、マイクロソフトは全員を受け入れると表明し、投資家たちは怒りを募らせて訴訟の準備まで進めていました。自分たちが暫定CEOに据えたミラ・ムラティでさえ署名に加わっていたのです。

しかし取締役会は諦めませんでした。

いや、正確に言えば、諦めることができなかったのです。自分たちの手でサムを追い出したのに、今さら「ごめん、間違えた」と言うには、プライドが許さなかったはずです。それで彼らは別の方法を探しました。

「サム以外の別のCEOを探そう！」

取締役会は週末をかけて新しいCEO候補を探しました。彼らが接触した人物の中で最も意外だったのが、ダリオ・アモデイ、AI企業アンソロピック（Anthropic）のCEOでした。

アンソロピックとはどんな会社か。

OpenAIから出た人たちが作った会社です。彼らはサム・アルトマンの急速な商業化路線が危険だと考えてOpenAIを去りました。「私たちはもっと安全に、もっと倫理的にAIを作る」と言って。

取締役会の立場からすれば、ダリオは申し分のない候補でした。AIの技術も理解しており、安全性を重視し、サムのように前のめりに突き進まなそうでした。二つの会社を合併させようという提案まであったと言われています。

ダリオは断りました。

「申し訳ないですが、私たちは私たちの道を行きます。」彼はOpenAIの混乱の中へ飛び込む気はありませんでした。

取締役会のプランBも失敗しました。

日曜日の夜、取締役会はTwitchを創設したエメット・シアという人物を新しいCEOとして発表しました。「よし、これで終わりだ。サムは戻らない。」取締役会はそう思っていたでしょう。

社員たちはいっそう怒りました。

「私たちの意見は完全に無視されているの?」「そもそもエメット・シアって誰なんだ?」「Twitchを作った人がChatGPTを理解できるのか?」怒りのメッセージがSNSに溢れかえりました。

マイクロソフトのCEO、サティア・ナデラが公の場で発言しました。

「エメット・シアCEOとも喜んで協力いたします。同時に、サム・アルトマンがOpenAIに戻ることも歓迎します。選択はOpenAIの取締役会と社員の皆さんに委ねられています。」

その言葉の真意はこういうことでした。「私たちはどちらでも構わない。サムが戻ろうと戻るまいと、私たちには勝ち目がある。戻らなければ700人の社員がうちに来て、戻れば私たちの投資が守られる。」

火曜日の朝になりました。

サム・アルトマンがOpenAI本社に姿を現しました。手には訪問者バッジをつけて。自分が創った会社に、客として入ったのです。彼はXに写真を投稿しました。「初めて、そして最後に、OpenAIに訪問者として来ました。」

その一枚の写真が全世界に広まりました。

人々は笑い、泣き、怒りました。「創業者が自分の会社に客として行くって?」「取締役会は本当におかしくなってる。」世論は完全にサムの側についていました。

交渉が始まりました。

取締役会の代表としてアダム・ディアンジェロ(Quora CEO)が交渉に臨みました。サム・アルトマンとグレッグ・ブロックマン、投資家たち、そしてエメット・シアまでが交渉テーブルに着きました。

サムの要求は明確でした。

「取締役会を完全に刷新せよ。私を追い出した人間が残るなら、二度と戻らない。」

取締役会は追い詰められました。

受け入れれば自分たちの完全敗北です。拒否すれば会社が崩壊します。どちらかを選ぶしかありませんでした。

火曜日の夜10時、OpenAIの公式アカウントに投稿が上がりました。

「私たちはサム・アルトマンがCEOとして復帰することに原則合意しました。新しい取締役会はブレット・テイラー(議長)、ラリー・サマーズ、アダム・ディアンジェロで構成されます。」

5日ぶりの復帰でした。

金曜日に解雇され、火曜日に戻ってきました。シリコンバレーの歴史上、最も短い解雇期間でした。スティーブ・ジョブズは11年ぶりにAppleへ戻りましたが、サムは5日で戻ってきたのです。

新しい取締役会の構成が興味深いものでした。

ブレット・テイラーはSalesforceという巨大テクノロジー企業の前CEOでした。聡明で、経験豊富で、中立的な人物でした。彼が議長に就きました。

ラリー・サマーズはハーバード大学の教授であり、元アメリカ財務長官でした。経済学者として政府での経験も持ち、高い名声を誇る人物でした。そんな人物が取締役に加わったのです。

アダム・ディアンジェロだけが唯一残りました。サムを追い出すことに賛成した人物であるにもかかわらず。なぜ彼を残したのでしょうか。おそらく交渉の結果だったのでしょう。「あなたが仲裁役を果たしたのだから、一席は差し上げます」——そんな取引があったのかもしれませんが。

イリヤ・スツケヴァーは外れました。ヘレン・トナーとターシャ・マクコーリーも去りました。サムを追い出した人たちは全員、席を失いました。

サム・アルトマンはXに短く書きました。

「OpenAIを愛しています。この数日間に私がしたことはすべて、このチームとミッションを守るためのものでした。日曜日の夜にMicrosoftに合流することを決めたとき、それが私とチームにとって最善の道だと思いました。新しい取締役会とサティアのサポートのもと、私はOpenAIに戻ります。」

勝利の宣言ではありませんでした。和解のメッセージでした。

社員たちは歓声を上げました。「OpenAI is nothing without its people」というメッセージが再びSNSを埋め尽くしました。今度は祝福の意味で。

すべてが幸せな結末だったわけではありませんでした。

傷が残りました。信頼が壊れました。イリヤ・スツケヴァーはサムのもっとも近い同僚でしたが、彼を裏切りました。この傷は簡単に癒えるのでしょうか。

社内にも亀裂が生じました。「安全が大事だ」vs「早く前進しなければ」という二つの陣営の対立が表面化しました。サムが戻ったからといって、この論争が消えたわけではありませんでした。

外部ではOpenAIを違う目で見えるようになりました。

「信頼できる会社なのか?」「CEOが追い出されることもあるじゃないか?」「投資しても大丈夫なのか?」信頼にひびが入りました。

競合他社は機会を狙いました。Google、Anthropic、Mistralといった会社が「私たちは安定しています」とアピールし始めました。「OpenAIのように混乱していません。私たちを選んでください。」

Microsoftが最大の勝者でした。

サティア・ナデラは天才的な動きを見せました。サムにすぐさま採用を提案し、全社員を受け入れると表明しながら、同時にOpenAIの取締役会とも関係を保ちました。どちらが勝ってもMicrosoftが勝つ構図を作り上げたのです。

実際にMicrosoftの株価は上がりました。サムが復帰したというニュースを受けてのことです。投資家たちは「MicrosoftがOpenAIを完全に掌握した」と受け止めました。

一つの変化がありました。

Microsoftが取締役にオブザーバー資格で加わりました。議決権はないものの、会議に出席してすべてを見ることができる立場でした。「もう5分前に知らされることはない」——そういう意味でした。

OpenAIはもはや完全に独立した会社ではなくなりました。

11月29日、公式発表がありました。

「サム・アルトマンCEO、グレッグ・ブロックマン社長、ミラ・ムラティCTOが正式に復帰しました。OpenAIは引き続き、安全で有益なAGI開発というミッションを追求していきます。」

会社は正常に戻ったように見えました。

しかし内部では多くのことが変わっていました。サム・アルトマンの権力はさらに強まりました。もはや彼を止められる取締役会はなかったからです。新しい取締役たちはサムを支持するために来た人たちでした。

安全を重視する声は弱まりました。イリヤが去り、トナーとマクカリーも離れたからです。OpenAIは今、「早く、もっと早く」という方向へ確実に傾いていました。

社員たちは安堵しながらも、どこか不安を感じていました。

「勝った。サムが戻ってきた。」それは喜ばしいことでした。しかし同時に、「私たちは本当に正しい道を歩んでいるのだろうか」という問いも、心の中に残り続けました。

5日間のドラマは幕を閉じました。

王は帰還しました。しかし王国は、かつてのままではありませんでした。傷があり、疑念があり、変化がありました。

この出来事はOpenAIだけでなく、世界中のテクノロジー企業に教訓を残しました。「取締役会はCEOを追い出せる。しかし社員たちはCEOを呼び戻せる。」「会社の真の主人は、紙の上の権力ではなく、現場で働く人々だ。」

サム・アルトマンはより強くなって戻ってきました。自分のために戦ってくれるほど、チームに愛されていることを確かめました。自分がいなくても会社が完璧に機能できるほど、優れたチームを作り上げていたことも確かめました。

しかし彼は、謙虚にもなりました。

「危機の瞬間より、後始末のほうがずっと難しかった。」彼は後にそう語りました。復帰してから傷を癒し、信頼を取り戻し、チームをひとつにまとめ直すことが、CEOとして戻ること自体よりは

るかに大変だったと。

5日間の戦いは終わりましたが、本当の闘いはここからでした。

イリヤ・サツケヴァーの謝罪

イリヤ・サツケヴァー。

この名前を知らなければ、OpenAIを理解することはできません。

彼はサム・アルトマンとともにOpenAIを創設した共同創業者でした。主任科学者として、GPTを初めて生み出した人物の一人であり、AI分野の伝説的な研究者でした。ジェフリー・ヒントン（AIの教父と呼ばれる人物）の弟子であり、ディープラーニング革命の中心人物でもありました。

サム・アルトマンがビジネスの天才なら、イリヤは技術の天才でした。

二人は完璧な組み合わせでした。サムが「これを作りましょう」と言えば、イリヤは「どう作るかわかります」と答えました。サムが投資家を説得している間、イリヤは研究室でモデルを学習させていました。

OpenAIの社員たちはイリヤを敬っていました。

「イリヤがいなければGPTもなかった。」「彼はAIを本当に理解している数少ない人物だ。」「サムが会社の顔なら、イリヤは会社の頭脳だ。」そんな言葉が社内で飛び交っていました。

そのイリヤが、サムを追い出すことに賛成したのです。

金曜日、取締役会がサムを解雇したとき、イリヤはその場にいました。彼は投票で「賛成」を選びました。共に会社を作った友人を、同僚を、パートナーを追い出したのです。

なぜそうしたのでしょか。

イリヤは、誰よりも真剣にAIの安全を考える人物でした。AGI（汎用人工知能）は人類を救うこともできるが、破滅させることもできると信じていました。だから、ゆっくり、慎重に、安全に進まなければならないと考えていました。

サム・アルトマンは違いました。

サムは「早く作り、早くリリースして、世界を変えよう」と言い続けました。ChatGPTをリリースするときも、GPT-4を作るときも、常にスピードが重要でした。「競合他社より早く進まなければ。」「Googleに追いつかれる前に市場を押さえなければ。」

イリヤは不安でした。

「私たちは進みすぎているんじゃないか。」「安全確認はきちんとしたのか。」「こんなに急げば事故が起きるかもしれない。」彼の懸念はどんどん大きくなっていきました。

2023年の夏から、二人の間にひびが入り始めました。

サムは開発者向けカンファレンスを開き、新機能を発表し、投資を集め、急速に拡大していきました。イリヤは「待ってくれ、これは当初の計画と違う」と言いました。

二人は何度も衝突しました。

「安全が先だ」対「スピードが命だ」。「ゆっくり行こう」対「早く行かなければ勝てない」。「非営利の精神を守ろう」対「会社が生き残ってこそ、精神も守れる」。

取締役会でも、対立が表面化しました。

イリヤは取締役会の他のメンバー（ヘレン・トナー、ターシャ・マクカリー）と話し合いました。「サムは行き過ぎています。私たちが何か手を打たなければ。」二人も同意しました。「そうです。OpenAIは本来の目的から外れています。」

11月17日金曜日、彼らは行動に出ました。

イリヤは自分の良心に従いました。「会社を救わなければ。サムがこのまま続けたら危険だ。」彼は本当にそう信じていました。

サムが解雇されると、会社は混乱に陥りました。グレッグ・ブロックマンが辞任しました。社員たちは怒りました。投資家たちが抗議しました。Microsoftが介入しました。

イリヤは困惑しました。

「こうなるとは思わなかった。」彼は社員たちが自分を支持してくれると思っていました。「僕たちが会社を救ったんだ。みんな感謝してくれるはずだ。」ところが正反対でした。社員たちは彼のことを裏切り者と言いました。

月曜日、700人の署名が出たとき、イリヤの名前もそこにありました。

彼は土壇場で立場を変えました。サムを追い出すことに賛成したその本人が、今度はサムを呼び戻すよう求める署名をしたのです。

そしてXに投稿しました。

「取締役会の行動に加わったことを深く後悔しています。私はOpenAIを傷つけるつもりなど、まったくありませんでした。私たちが共に築いてきたすべてのものを愛しており、会社を再び一つにするためにできる限りのことをします。」

この投稿は爆弾でした。

サムを追い出した中心人物が公然と謝罪したのです。取締役会の道徳的権威は完全に崩れました。「イリヤでさえ後悔しているのに、あの決定が正しかったはずがない。」

サム・アルトマンはハートの絵文字で返信しました。

言葉は何もありませんでした。ただハートひとつだけ送りました。許しの印だったのでしょうか。それとも「あとで話そう」という意味だったのでしょうか。

社員たちは複雑な感情に揺れました。

イリヤを憎むべきなのか。彼は裏切り者でした。いや、イリヤを理解すべきなのか。彼は会社を愛し、安全を心配し、善意から行動しました。ただやり方が間違っていただけです。

サムが復帰した後、イリヤは会社に残りました。

最初はうまくいっているように見えました。イリヤは研究を続け、サムは経営し、みんなが「もう大丈夫」と言っていました。

しかし内側では違いました。

イリヤはもうサムの近い同僚ではありませんでした。二人は同じ会議室に座っていましたが、以前のように気軽に冗談を言い合うことはできませんでした。信頼が壊れたのです。一度壊れた信頼は、簡単には戻りません。

2024年5月、イリヤはOpenAIを去りました。

公式には「新しいプロジェクトを始めるため」と発表しました。しかし本当の理由は誰もが知っていました。もう一緒にいられなかったのです。

イリヤは自分のAI会社を立ち上げました。

Safe Superintelligence (SSI) という会社でした。名前からして意味深です。「安全な超知能」。OpenAIで果たせなかった夢を、そこで実現しようとしているのです。

二人の天才が共に世界を変えようと出発しました。一人はビジネスの天才、もう一人は技術の天才。完璧な組み合わせでした。ところが夢への歩み方が違いました。一人は「速く」と言い、もう一人は「ゆっくり」と言いました。

結局、二人は袂を分かちました。

イリヤの謝罪は、このドラマの中で最も人間的な瞬間でした。

彼は間違いを犯し、それを認め、許しを求めました。完璧な人間はいません。天才も間違えます。大切なのは、間違いを認める勇気です。

サム・アルトマンも寛大でした。彼は復帰後にイリヤをすぐに追い出しませんでした。共に働く機会を与えました。結局は去ることになりましたが、サムが強制的に追い出したわけではありませんでした。

この出来事は私たちに問いを投げかけます。

AIを作るとき、速度と安全のどちらが大切なのでしょうか。速く作って世界を変えるべきか、ゆっくり作って安全を確保すべきか。

正解はありません。イリヤも間違っていなかったし、サムも間違っていませんでした。二人とも正しく、二人ともそれぞれのやり方で人類を助けたかったのです。

悲劇は、二人が共に歩めなかったことです。

この出来事には複合的な含意がある。

5日間のドラマが幕を閉じ、人々はこの出来事をどう受け止めればよいか考えあぐねました。

単なるCEO解任事件ではありませんでした。

これは21世紀の企業がどのように機能するかを示す、教科書のような出来事でした。そこには権力、ビジョン、忠誠心、文化、そして新しい時代のリーダーシップがすべて詰まっていました。

最も重要な教訓はこれです。

「会社の本当の主人は誰か？」

かつては答えが明快でした。株主が主人です。お金を出した人が会社を所有し、取締役会を選び、CEOを任命します。CEOが気に入らなければ、取締役会が追い出します。以上。

OpenAIの事件は、この公式を打ち砕きました。

取締役会がCEOを追い出しました。法的には完全に正当な行為でした。定款にもそう書かれており、手続きも踏んでいました。しかし会社は崩壊寸前でした。

なぜでしょうか？

社員たちが拒否したからです。700人が「私たちは取締役会ではなく、サムに従います」と宣言しました。法的権限を持つ取締役会を、実質的な力を持つ社員たちが打ち負かしたのです。

これは革命でした。

フランス革命のように血は流れませんでした。権力構造は完全にひっくり返りました。「上から下へ」命令するシステムが崩れ、「下から上へ」意志を示すシステムが勝利しました。

なぜこんなことが可能だったのでしょうか？

OpenAIは特別な会社だったからです。工場ではありませんでした。人間の頭脳が資産の会社でした。GPTの作り方を知っている人たち、コードを書ける人たち、モデルを学習させられる人たち。この人たちがいなければ、会社はただの空き建物に過ぎませんでした。

Microsoftはそれを知っていました。

「建物はいらぬ。人さえ連れてこい。」サティア・ナデラのこのひと言がOpenAIの取締役会を崩壊させました。物理的資産ではなく、人的資産こそが本当の価値であることを証明したのです。

企業官僚主義とは何でしょうか？

規則・手続き・形式を重んじるシステムです。「取締役会の決議が重要だ」「定款にこう書いてある」「法的に我々が正しい」。そういう論理です。

間違いではありません。会社には規則が必要です。しかし規則だけを守ろうとして会社の本質を失ってしまったら、どうなるのでしょうか？OpenAIの取締役会はまさにその落とし穴にはまりました。

彼らは形式的には正しかったのです。

「CEOは取締役会に正直ではなかった」「私たちは会社の安全を守らなければならない」「定款上、私たちに権限がある」。どれも事実でした。法廷に持ち込めば、取締役会が勝訴できたかもしれません。

しかし実質的には、完全に間違っていました。

社員たちの心を読めませんでした。会社の文化を理解できませんでした。サム・アルトマンがなぜ重要なのかもわかりませんでした。彼らは書類上の権力だけを見て、本当の力がどこにあるかを見ていなかったのです。

ビジョンとは何でしょうか？

未来への強烈な夢です。「私たちはAGIを作る。人類の歴史を変える。これはただの会社ではなく、missionだ」。サム・アルトマンが社員たちに与えたのは、まさにこのビジョンでした。

社員たちは給料のためにOpenAIにいたのではありませんでした。

もちろんお金も重要でした。OpenAIのPPUには莫大な価値がありましたから。しかしそれだけがすべてではありませんでした。GoogleやMetaに行けば、もっと多くの報酬を得られたかもしれません。

彼らがOpenAIにいた理由は、夢でした。

「私たちがChatGPTを作った。世界が変わった。次は何を作ろうか？」この興奮、この可能性、この挑戦。それが彼らを動かしていました。

サム・アルトマンは、その夢の象徴でした。

彼がいなければ夢もないと、社員たちは考えていました。だから700人が署名したのです。お金よりも、安定よりも、キャリアよりも、ビジョンの方が大切だったのです。

取締役会はこれをわかっていませんでした。

彼らは「サム一人いなくても会社は回る」と考えていました。CEOは取り替えが利くと見ていました。しかし違いました。サムはただのCEOではなく、ビジョンの体現者でした。

この事件が示したのは、21世紀の企業ではビジョンが権力より強いということです。

書類上の規則より、人々の心の中にある夢の方が強い力を持ちます。取締役会の議決権より、社員たちの集団的な意志の方が大きな影響力を発揮します。

スティーブ・ジョブズと比較する人が多くいました。

1985年、スティーブ・ジョブズもAppleを追われました。取締役会が彼を解雇したのです。彼は11年後に戻り、Appleを世界最高の会社に育て上げました。

サム・アルトマンはわずか5日で戻ってきました。

なぜこれほど早かったのでしょうか？時代が違うからです。1985年には、社員たちがCEOのために集団行動を起こす手段がありませんでした。2023年には、Slack、X、メールで数時間のうちに組織できました。

テクノロジーが権力構造を変えたのです。

ソーシャルメディアのおかげで、社員たちの声は瞬く間に世界中へ広がりました。「OpenAI is nothing without its people.」この言葉は数時間のうちに数百万人へ届きました。取締役会は世論の圧力に耐えきれませんでした。

投資家たちも変わっていました。

昔の投資家なら取締役会を支持したでしょう。「ルールはルールだ。CEOが問題を起こしたなら解任が当然だ」と。しかし2023年の投資家たちは違いました。彼らは実利主義者でした。「会社が潰れたら我々の金が消える。サムを復帰させろ」と。

これが現代の資本主義です。

ルールより結果が大切です。形式より実質が大切です。大義名分より生き残りが大切です。

OpenAI事件は、未来の青写真です。

これから同じようなことがもっと起きるでしょう。強烈なビジョンを持つ創業者が官僚的な取締役会と衝突し、社員が創業者の側に立ち、結局創業者が勝つ、そのパターンです。

これは良いことなのでしょうか。

難しい問いです。一方では良いことです。ビジョンのあるリーダーが官僚主義の邪魔を受けず、世界を変えられるのですから。サム・アルトマンのように。

もう一方では危険でもあります。誤ったビジョンを持つリーダーも止められなくなるからです。もしサムが本当に危険な決断を下していたとしたら。取締役会が彼を止められなかったとしたら。誰が責任を取るのでしょうか。

イリヤ・サツキバーが心配していたのは、まさにこのことです。

「サムは進みすぎている。誰かがブレーキを踏まなければ」。彼の懸念が間違っているわけではありません。AIは本当に危険になりえます。しかし彼が選んだ方法は間違っていました。

バランスが必要です。

ビジョンも必要で、牽制も必要です。スピードも必要で、安全も必要です。リーダーの自由も必要で、取締役会による監督も必要です。

OpenAIはこのバランスを見つけられませんでした。

サムの復帰後、会社は完全に「スピード」の方向へ傾きました。安全を心配していた声は、すべていなくなりました。今や誰がサムを止められるのでしょうか。

新しい取締役会は、サムを支持するために集まった人々です。ブレット・テイラー、ラリー・サマーズ。優れた人物たちですが、サムと戦う人たちではありません。

これが勝利の代償です。

サムは勝ちました。しかし同時に、より大きな責任を負うことになりました。今や彼を止める人がいないからです。彼が過ちを犯せば、それはOpenAIの過ちとなり、さらにはAI産業全体の過ちになりかねません。

この事件の本当の意味はここにあります。

「21世紀の企業において、すべては人だ」。建物でも、特許でも、資金でもなく、人です。その人々を一つに結びつけるビジョンです。そのビジョンを信じ、従う文化です。

取締役会は紙の上の権力しか持っていませんでした。

社員たちは実質的な力を持っていました。両者がぶつかったとき、実質的な力が勝ちました。

これこそが、企業官僚主義を打ち破った強力なビジョンの勝利です。

これから歴史の本にはこう書かれるでしょう。「2023年11月、OpenAIの社員700名がCEOのために立ち上がった。彼らは形式的な権威を持つ取締役会を打ち倒し、自分たちが信じるビジョンを守り抜いた。これは21世紀の企業史における転換点であった」と。

サム・アルトマンは、ただのCEOではありません。

彼は新しい時代のリーダーシップを象徴しています。命令ではなくビジョンで導くリーダー。権威ではなく信頼で尊敬されるリーダー。ルールではなく夢で人々を動かすリーダー。

OpenAIは、ただの会社ではありません。

革命です。未来へ向けた集団的な挑戦です。だからこそ社員たちは命がけで守ったのです。

世界中の創業者たちがこの事件を学んでいます。「どうすれば社員が自分のために戦ってくれるのか」と。取締役会もまた教訓を得ています。「CEOをむやみに動かしてはいけないのだ」と。

投資家たちも変わります。「ルールより人が大切なのだ」と。社員たちも自信を得ます。「私たちにも力がある」と。

これこそが、サム・アルトマンの解任と復帰という出来事が21世紀の企業史に残した最大の遺産です。

ビジョンが権力に勝った瞬間。人がルールに勝った瞬間。未来が過去に勝った瞬間。

そして一人の夢が、700人の心を動かし、ついには巨大な企業の運命を変えた瞬間です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

3.3 危機を越えて残った教訓

第3部 解任と復帰、ビジョンの勝利

← 前の章

次の章 →

スティーブ・ジョブズ事件との比較

2023年11月17日、金曜日。サム・アルトマンはラスベガスのF1グランプリを観戦中、突然ビデオ会議への接続要求を受け取りました。前夜、イリヤ・サツキーバーがテキストメッセージで送ってきた招待でした。正午、画面の向こうに取締役会のメンバーたちが映っていました。アルトマンはCEOを解任されたという通告を受けました。それから5分から10分後、OpenAIのウェブサイトに公式発表が掲載されました。

その瞬間から、メディアは一斉にある名前を口にしました。

スティーブ・ジョブズ。

1985年、自らが創ったアップルから追われた、あの伝説的な人物です。

「創業者が自分の会社から追い出されるなんて、これはジョブズの再来だ！」シリコンバレーの伝説的な投資家ロン・コンウェイはX（旧Twitter）にそう書きました。「1985年にアップルの取締役会がスティーブ・ジョブズを追放して以来、見たことのない取締役会クーデターがOpenAIで起きた。」

この比喻は魅力的でした。

ジョブズは10年以上離れた後、1997年にアップルへ戻り、iPodとiPhoneで世界を変えました。アルトマンも同じような伝説を書くのでしょうか。

アルトマン自身は、この比較を真っ向から否定しました。「私たちの状況はまったく違いました。」

あるインタビューで司会者が尋ねました。「スティーブ・ジョブズもアップルから追われたとき、それは苦い薬だったが患者には必要だったと言いました。あなたもそう思いますか？」アルトマンは少し躊躇しました。そして慎重に答えました。「すべてのことが4～5日で完全に終わってしまいました。まるでひどく奇妙な熱病にうなされる夢のようでした。」

時間が決定的な違いでした。

ジョブズは1985年から1997年まで12年間待ちました。その間にNeXTとPixarを作り、アップルが倒産寸前まで追い込まれるのを見守りました。復帰は、チームの救世主として登板することでした。

一方、アルトマンはわずか5日で戻ってきました。金曜日に解雇され、水曜日の午前1時に復帰の発表がなされました。

なぜこれほど早かったのでしょうか。

社員たちのおかげでした。770人のうち700人以上が公開書簡に署名しました。「取締役会が辞任し、サムを復帰させなければ、私たちはマイクロソフトへ行く。」Googleの元CEOエリック・シュミットは後にこう述べました。「取締役会が日曜日に再びサムを解任しようとしたとき、社員たちが反乱を起こした。会社か取締役会かを選べ、と言ったんです。CEOへの360度評価で、これ以上明確なフィードバックがありますか？」

ジョブズが追われたとき、アップル社内では彼の独善的な経営スタイルへの不満が多くありました。

中には彼の去ることを歓迎する人もいました。

アルトマンの場合は正反対でした。ほぼすべての社員が彼の味方でした。共同創業者のイリヤ・サツキーバーでさえ書簡に署名し、「取締役会の行動に加担したことを後悔している」とツイートしました。

復帰の性格も異なりました。ジョブズは会社を取り戻すために戦いました。復帰後は自分のやり方でアップルを完全に作り直しました。反対派を排除し、製品ラインを絞り込み、権力を集中させました。アルトマンは違いました。復帰後、彼は権力を独占しませんでした。むしろ取締役会を拡大し、再構成しました。セールスフォースの元共同CEOブレット・テイラーが議長に就き、元財務長官ラリー・サマーズが加わりました。より多くの専門家を招き入れました。

あるインタビューでアルトマンは率直に語りました。「ジョブズは会社を取り戻そうとしました。私は会社を取り戻そうとしたことはありません。ただ、ミッションが止まってはならないと信じていただけです。」

驚くべきことに、アルトマンは危機の中でも組織が完璧に機能しているのを見て喜んでいました。後に彼はこう振り返りました。「私が最も誇りに思った瞬間は、危機的状況の中でも経営チームが私なしで完璧に動いているのを見届けたときでした。彼らの誰もが会社を立派に率いられると確信しました。」

これがジョブズとの決定的な違いでした。

ジョブズは、アップルは自分なしでは生きられないことを証明しました。アルトマンはOpenAIが自分なしでもうまく回ることを確かめて喜びました。彼が見ていたのは個人の英雄譚ではなく、強いチームと共有されたビジョンの勝利でした。

二人とも時代を変えるビジョンを持つ創業者であり、自らが作った会社から追われ、劇的に復帰しました。しかしアルトマンは「第二のスティーブ・ジョブズ」になろうとは決してしませんでした。彼は自分だけの道を歩みました。権力ではなくミッションを、個人ではなくチームを、英雄譚ではなく組織の力を信じました。

2024年のあるカンファレンスで誰かが再びジョブズとの比較を持ち出したとき、アルトマンは笑いながら言いました。「スティーブ・ジョブズを解雇するなというエリック・シュミットの忠告は正しい。でも私の話は違います。私は解雇されたのではなく、ちょっと休暇を取ってただけです。」

冗談でしたが、真実が込められていました。5日間は休暇にしてはひどいものでしたが、人生を変える追放ではありませんでした。アルトマンにとってこの出来事は、復讐でも華々しい復帰でもなく、自分が作り上げた組織の強さを確かめる機会でした。そしてそれが彼を真のリーダーにしました。

小さな失敗が大きな危機への訓練だった。

「人間の感情のあらゆる範囲を経験しました。なかなか印象的でしたよ。」彼は後にそう振り返りました。

解任されたその夜、アルトマンはグレッグ・ブロックマンとともにサンフランシスコのあるバーに座っていました。

ブロックマンは取締役会で議長職を剥奪された直後に辞表を提出していました。二人はビールを飲みながら、何が起きたのかを整理しようとしていました。周りにいた人たちは彼らに気づき、写真を撮りました。近づいてきて握手を求める人もいました。「応援しています！」

月曜日、マイクロソフトのCEOサティア・ナデラがアルトマンに電話をかけました。「私たちと一緒にマイクロソフトで新しいAI研究チームを作りましょう。」アルトマンは悩みました。

新たに出発するのも悪くありませんでした。しかし彼の心の片隅にはOpenAIがありました。10年近く育てた組織、自分のビジョンを共有する仲間たちがそこにいました。

火曜日、社員たちの公開書簡が出ました。700人以上が署名しました。その中にはイリヤ・サツキーバーの名前もありました。その瞬間、アルトマンは気づきました。「これは自分一人の戦いじゃない。私たち全員の戦いだ。」

水曜日の夜明け、復帰の知らせが伝わりました。

5日間のジェットコースターが終わりました。しかし本当に終わったのでしょうか。アルトマンはそうは思いませんでした。本当の始まりはここからでした。

数ヶ月後、セコイア・キャピタルのイベントでアルフレッド・リンが尋ねました。「サム、レジリエンス（回復力）についてアドバイスをいただけますか？あなたは途方もない危機を経験されたから。」

アルトマンはしばらく考えてから言いました。「時間が経つにつれて楽になります。本当です。」

聴衆がざわつきました。楽になる？会社が空中分解しかけたあの危機が？

アルトマンは説明を続けました。

「創業者として、皆さんは多くの逆境に直面することになるでしょう。挑戦はますます困難になり、賭け金も大きくなります。しかし、感情的な負担は逆説的に減っていきます。より多くの悪い出来事を経験するほど、それに対処する能力が向上するからです。」

自身の経験を例に挙げました。Yコンビネーター（Y Combinator）の代表時代、数百ものスタートアップが失敗するのを手助けし、見守ってきました。創業者が涙を流しながら会社の幕を閉じるの

を見てきました。ループト（Loopt）を作った時は、通信会社と30回も交渉しました。29回は断られました。OpenAIの初期には、何を作るべきかわからず、ビデオゲームも検討しました。GPT-3のAPIは、コピーライティング以外にはあまり役に立ちませんでした。

「個々の危機を経験するたびに、少しずつ鍛えられました。小さな失敗が、大きな危機のための訓練だったのです。」

サムは続けました。「創業初期には、顧客を一失うだけでも世界が崩れ去るような気分になります。しかし、何百回、何千回もの拒絶や失敗を経験するうちに、感情の振幅が小さくなります。無感覚になるのではなく、問題の軽重を判断する能力が身につくのです。」

Yコンビネーター時代、彼は創業者たちにこう言っていました。「皆さんのほとんどは、最初に訪れる大きな危機で崩れ去るでしょう。しかし、生き残った人々は次の危機でより強くなります。」レジリエンスは筋肉のようなものでした。使えば使うほど強くなりました。

実際、心理学ではこれを「外傷後成長（Post-Traumatic Growth）」と呼びます。アメリカ心理学会によると、レジリエンスとは困難な経験にうまく適応するプロセスであり、その結果でもありません。精神的、感情的、行動的な柔軟性を備え、内的・外的な状況に適応する能力のことです。

アルトマンは、アブダビで開催されたHub71のイベントでこう語りました。「ほとんどの人は危機の時に諦めてしまいます。レジリエンスがないからです。しかし、良い知らせがあります。レジリエンスは教えることができるということです。Yコンビネーターで私たちはそれを学びました。レジリエンスを育てることができるということです。」

彼はレジリエンスを3つの要素で説明しました。第一に、自己信頼です。「私が知る最も成功した人々は、ほとんど妄想に近いほど自分を信じています。」GPTを拡張するために10億ドルを投じると言った時、多くの人が狂っていると言いました。しかし、アルトマンは確信していました。『これが正しい。大きく育てれば何かが起きるはずだ。』

もちろん、盲目的な自信は危険です。アルトマンもこれを認めました。「自己信頼は、自分の専門分野においてのみ適用すべきです。成功した人々が専門分野の外でひどい決定を下すのを何度も見てきました。」AIの専門家が不動産投資やレストラン事業で失敗するケースが多々ありました。

第二に、適応力です。「どんな問題が起きても解決するという精神が重要です。今すぐ方法がわからなくても大丈夫です。最終的には見つけ出すのですから。」ループト時代、30回目の試行でようやくAT&Tと契約した時、アルトマンはすでに29通りの方法がうまくいかないことを学んだ状態でした。個々の失敗がデータだったのです。

第三に、長期的な視点です。「ほとんどの人は、数年以内に何を達成できるかを考えます。しかし、数十年という観点で見ればどうでしょうか？今後数年間、完全に誤解されたとしても構わないという姿勢、それが強力な武器になります。」

アルトマンは危機を経験する中で、もう一つのことに気づきました。大きな危機においては、むしろ感情が消えるということです。「事態が大きくなりすぎると、感情が介入する余地がなくなります。どの選択が正しいかが直感的に見えてくるのです。」5日間の混乱の中でも、彼は冷静に次のス

テップを計画しました。マイクロソフト（Microsoft）の選択肢を維持しながら、従業員の動揺を見守り、理事会との交渉の余地を残しました。

後に、ある創業者が尋ねました。「どうしてあんなに冷静でいられたのですか？」アルトマンは笑って答えました。「正直に言うと、私の脳が感情を処理する余裕がなかったのだと思います。あまりに多くのことが同時に起きすぎて、ただオートパイロットモードに入ってしまった。悲しみや怒りが湧いてきたのは、後になってからでした。」

彼は創業者たちに最後のアドバイスを残しました。「小さな挑戦を避けないでください。それが後に大きな挑戦を乗り越える力になります。ジムで軽いウェイトから始めて徐々に重いものを持ち上げるように、危機にも練習が必要です。」

レジリエンスは生まれつき備わっているものではありませんでした。作り上げられるものでした。サム・アルトマンは5日間の地獄のような時間を耐え抜いた後、より強いリーダーとなって戻ってきました。そして、その経験を何千人もの創業者たちと分かち合いました。

レジリエンスの力

2023年11月22日水曜日の午前1時。アルトマンがOpenAIのCEOに復帰するという発表がありました。従業員たちは歓喜しました。投資家たちは安堵しました。メディアは「劇的な大逆転！」という見出しを掲げました。

すべてが終わったかのように見えました。

しかし、アルトマン自身は全く違う感じ方をしていました。彼は後にこう回想しています。「多くの人が、私が復帰した瞬間にすべてが終わったと考えました。私にとっては、それがまた別の始まりでした。本当に大変な仕事はこれからだったのですから。」

セコイアのイベントで、彼は率直に語りました。「危機の瞬間よりも、その後の方がはるかに大変でした。」聴衆は驚いた表情を浮かべました。解任されてから復帰するまでの5日間よりも大変だということですか？

アルトマンは説明しました。「危機の時はアドレナリンが出ます。多くの人が支持してくれ、何かをしなければならないという明確な目標があります。しかし、危機が去った後の60日目、90日目には、誰も注目しません。その時、皆さんは一人で破片を拾い集めながら再建しなければなりません。それははるかに孤独で困難なプロセスなのです。」

アルトマンが直面した最初の課題は、理事会の再構成でした。既存の理事会は事実上崩壊していました。

サムを解任したヘレン・トナー、タシャ・マコーリー、イリヤ・サツケヴァーは去りました。クオオラ（Quora）のCEO、アダム・ディアンジェロだけが残りました。空席を埋めなければなりません。しかし、誰を据えるべきか？今回の事態は、理事会構造の根本的な問題を露呈させました。

アルトマンと新しい理事たちは、いくつかの原則を立てました。第一に、理事会をより大きくすること。4人では少なすぎました。少数が会社全体を左右することができました。第二に、多様な背景

を持つ専門家を招くこと。純粋なAI安全研究者だけではバランスが取れませんでした。第三に、透明性を高めること。今回のようにCEOを解任しながら、Microsoftにさえ事前通知をしないという事態は、二度と起こしてはなりませんでした。

ブレット・テイラーが議長に就任しました。彼はグーグル（Google）、フェイスブック（Facebook）、セールスフォース（Salesforce）で働いたシリコンバレーのベテランでした。ツイッター（Twitter）の理事会議長として、イーロン・マスクによるツイッター買収を主導した経験もありました。ラリー・サマーズ元財務長官が加わりました。経済学者であり政策の専門家です。後には、さらに多くの理事が追加されました。

理事会の構成は始まりに過ぎませんでした。

第二の課題は、従業員たちの感情的な回復でした。700人以上の従業員が5日間、極度のストレスを経験しました。会社が崩壊するかもしれないという恐怖、自分たちが愛するリーダーを守らなければならないという責任感、Microsoftへの移籍準備をしながら感じた不安、そして劇的な逆転後の安堵感。感情のジェットコースターでした。

アルトマンは復帰後、毎週全社ミーティングを開きました。彼は従業員の質問に一つ一つ答えました。「なぜこのようなことが起きたのですか？」「再発する可能性はありますか？」「私たちは安全なのですか？」一部の従業員は依然として不安を抱いていました。特にイリヤに従っていた研究者たちは、複雑な感情を抱いていました。

最も困難な問題は、イリヤ・サツケヴァーとの関係でした。イリヤはサム・アルトマンの解任において主導的な役割を果たしました。しかし、48時間もしないうちに後悔し、従業員の書簡に署名しました。サムは彼を許したのでしょうか？組織はイリヤをどのように受け入れるのでしょうか？

アルトマンは公にイリヤを非難することはありませんでした。「イリヤはAGIの安全性を心から心配していました。彼の意図は善意によるものでした。」しかし、内部的には緊張が残っていました。結局、2024年5月、イリヤはOpenAIを去り、自身の会社 Safe Superintelligence Inc. を設立しました。彼の退社発表文にはこのような言葉がありました。「OpenAIはサムのリーダーシップの下、安全で有益なAGIを構築すると確信しています。」

第三の課題は、外部からの信頼回復でした。Microsoftは最大の投資家でありパートナーでした。130億ドルを投資していながら、CEO解任の知らせを「直前になって」耳にしました。サティア・ナデラは腹を立てたのでしょうか？実際、Microsoftの株価は金曜日に下落しました。アルトマンは復帰後、ナデラに何度も電話をかけ、関係を修復しました。ナデラはXにこう記しました。「より安定し、情報に基づいた効果的なガバナンスへの第一歩です。」

政府の規制当局も注目しました。SEC（証券取引委員会）とFTC（連邦取引委員会）が調査を開始したという噂が流れました。AGIを開発する企業がこれほど不安定とは、政府が介入すべきではないのか。アルトマンは上院公聴会で証言し、OpenAIの安全手続きを強化しました。

外部法律事務所WilmerHaleに独立調査を依頼しました。3万件以上の文書を精査し、数十人にインタビューを行いました。2024年3月、調査結果が出ました。「取締役会とアルトマンの間の信頼崩壊が原因だった。製品の安全性や財務上の問題ではなかった。アルトマンの行動は解任を正当化す

るものではなかった。」

報告書は一点を明確にしました。取締役会は「拙速に行動し、主要なステークホルダーへの事前通知を怠り、アルトマンに懸念を解消する機会を与えなかった」。二度とこのようなことが起きないように、新たなプロセスを設けました。

五つ目の課題は、組織文化の再建でした。危機の前、OpenAIは急成長するスタートアップでした。危機の後、世界が注目するAIリーダーとなりました。文化も進化しなければなりませんでした。アルトマンは全体集会でこう語りました。「私たちが経験した危機は、一つの時代を終わらせました。私たちは今、より成熟した組織にならなければなりません。ただ、スタートアップのスピードと柔軟性は失ってはなりません。」

容易ではないバランスでした。小さく機動力のあるチームを維持しながら、安全性とガバナンスを強化しなければなりませんでした。アルトマンはこれを「スピードと責任の融合」と呼びました。

アルトマンは創業者たちにこう助言しました。「危機をどう乗り越えるかについての資料は山ほどあります。しかし、危機から60日後をどう過ごすかについての助言はほとんどありません。私も探しましたが、見つかりませんでした。それこそが本当に必要なものののに。」

後進の創業者たちに実践的なアドバイスを伝えました。まず、感情の波を覚悟してください。危機の直後はアドレナリンのおかげで大丈夫に感じますが、2～3週間後に崩れることがあります。次に、小さな勝利を積み重ねてください。大きな再建より、小さな成功を連続して作ってください。製品のリリース、チームミーティング、顧客との面談。一つひとつ、正常に戻っていく感覚を取り戻してください。そして、一人で抱え込まないでください。共同創業者、経営幹部、メンターと継続的にコミュニケーションを取ってください。

2024年春、アルトマンはあるインタビューで振り返りました。「5日間の危機は悪夢でした。しかし、その後の6ヶ月が本当の試練でした。取締役会を再編し、チームの信頼を取り戻し、新たなシステムを構築していく過程。それこそがリーダーシップの真の姿でした。」

アルトマンにとって、危機は終わりではなく始まりでした。その長い再建の旅を通じて、彼は単なるビジョナリーから、組織を守り育てる真のリーダーへと成長しました。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士・元国会議員・AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

4.1 組織と開発の哲学

第4部 アルトマンのビジネス哲学と戦略

← 前の章

次の章 →

マスタープランの罨：素早く戦術を修正せよ。

サムは完璧な計画を立てるタイプの人間ではありません。計画をあまりに緻密に組み上げることを戒めています。

「10年後に自分たちがどうなっているか正確にわかるって？それは嘘です。」

彼はこう考えます。AI技術はあまりにも速く変わっていく。今日の素晴らしい計画が、明日には時代遅れの紙切れになりかねない。霧のかかった山道を歩くように、目の前の一步一步を慎重に踏みしめながら、道を切り開いていくのです。

OpenAIの創業期を振り返れば、この戦略がいかに重要だったかがわかります。2016年、彼らはわずか14人ほどの小さなチームでした。「自分たちが何をすべきか、正直よくわかっていませんでした。」

アルトマンは後にそう打ち明けています。当初、彼らはビデオゲームをプレイするAIやロボットハンドを作ろうとしていました。今私たちが知るChatGPTのような対話型AIは、遠い未来の話のように感じられていたのです。もしあのとき「私たちは絶対に言語モデルだけを作る」と固執していたら、どうなっていたでしょうか。

おそらく今のChatGPTは存在していなかったでしょう。

アルトマンの機敏な戦略は、サーファーが波に乗る姿に似ています。波の流れを読み、瞬間ごとに体のバランスを変える。倒れそうになればすばやく向きを変え、より大きな波が来れば思い切って乗り込む。

「科学がどこへ向かうのか、私たちはただそれについていくしかありません。」

彼の言葉です。人間があらかじめ定めた目標ではなく、技術そのものが示す可能性を追いかけるのです。

OpenAIが非営利から「利益上限付き」営利法人へと移行したのも、同じ文脈です。当初は純粋な研究機関を目指していました。「お金のためにAIを作ってはいけない」という信念がありました。しかし実際に世界最高の人材を集め、膨大なコンピューティング能力を確保するには、莫大な資金が必要でした。

「お金がなければ、夢もありません。」

アルトマンは現実を直視しました。そして思い切って組織の構造を変えました。投資家には一定のリターンを保証しつつ、それを超えた利益は人類のために使うという独自のモデルを作ったのです。当初の計画にこだわり続けていたなら、OpenAIはすでに閉鎖していたかもしれません。

ChatGPTの誕生過程も似ています。GPT-3という強力なモデルがありました。「これをどう使えばいいのか？」

チームは悩みました。完璧な製品を待つ代わりに、シンプルなチャットインターフェースを付けて世に出しました。結果は爆発的なものでした。リリースからわずか5日で100万人が使いました。

「素早く試し、素早く学び、素早く変える。」

それがアルトマン流のやり方です。失敗を恐れませんが、むしろ失敗から学ぶのです。一度で満点の答えを求めようとしない。60点の答えを早く見つけ、70点へ、80点へと改善を続けていきます。

最近OpenAIが一部のオープンソースモデルを公開したのも、同じ流れです。もともとすべての技術を非公開にするという方針でした。しかし中国のDeepSeek、MetaのLlamaといった強力なオープンソースモデルが登場すると、戦略を修正したのです。

「世界が変われば、私たちも変わります。」

頑固さはときに毒になります。柔軟性こそ生き残るための武器です。

アルトマンが強調するのは、北極星のような大きな目標はぶれないまま、その目標へ向かう道は変え続けられるべきだということです。OpenAIの北極星は「人類に利益をもたらすAGIの開発」です。この目標だけは絶対に変わりません。

道は毎日変わります。今日通れない道が、明日は開けることがあります。昨日は良かった道が、今日はふさがっていることもあります。

「複雑な計画にこだわらないでください。目の前の問題を解いていってください。」

アルトマンは若い起業家たちにこう助言します。5年計画、10年計画を細かく立てることに時間を使うなと言います。その代わり、今日できる最善を尽くし、明日学んだことでより良い判断を下せと言うのです。

こうした機敏な戦略が可能なのは、明確な原則があるからです。小さなことから始める。早く失敗し、早く学ぶ。ユーザーの反応に真剣に耳を傾ける。データで検証する。そして思い切って方向を変える。

「計画それ自体は役に立たないが、計画を立てるプロセスは重要だ。」

この言葉が浮かびます。アルトマンも同じように考えているようです。計画そのものより、絶えず考え判断し続けるプロセスが大切なのです。状況を読み、機会をつかみ、危険を避ける能力、それが肝心です。

結局、機敏な戦略論とは不確実性を受け入れる勇気です。「私には未来がわかりません」と正直に認めることです。知らないと認めてこそ学べる。学んでこそ成長できます。

アルトマンはOpenAIをそういう文化に作り上げました。間違えても構わない。失敗しても構わない。早く直せばいい。そういう雰囲気の中で、人々はより大胆になります。新しい挑戦を恐れなくなるのです。

「最高の計画とは、計画を変え続けることです。」

この逆説的な言葉がアルトマンの哲学の核心です。変化の激しいAIの世界で生き残るには、変化そのものを戦略にしなければなりません。流れに身を任せながらも、方向性は失わない。それが大切なのです。

「複雑な計画を逆算して進めることは、たいいていうまくいきません」。これはサム・アルトマンがよく使う言葉です。

私たちはたいいてこう考えます。「10年後にCEOになるためには、9年後は副社長、8年後はチームリーダー。」

未来から逆算していく方法です。地図に目的地を打ち込んで、そこから逆に道筋を描くようなものです。

アルトマンはそれがうまくいかないと言います。AIの世界では、なおさらそうです。

「5年前に立てたAI開発計画が、今見るとまるで白黒映画のようです。」

技術はあまりにも速く変化します。今日の完璧な計画が、明日には使い物にならなくなります。だからアルトマンは別の方法を選びました。

「今日持っているものを使って、明日少しだけ前に進むこと。」

これが彼の核心です。複利の魔法を信じるということです。小さな雪だるまが転がり続け、やがて巨大になる、その原理です。毎日1%ずつ改善するだけで、1年後には37倍になります。（実際に1.01の365乗は37.8です。）

OpenAIの初期のエピソードがこれをよく表しています。2016年、彼らは大規模言語モデルというアイデアを「まだ遠い話」だと考えていました。代わりにビデオゲームやロボットハンドのような研究をしていました。

「自分たちが何をすべきかわからなかったんです。」

率直な告白です。あのとき「何があってもChatGPTを作る」という複雑な計画を立てていたら？おそらく失敗していたでしょう。なぜなら当時は、そんな技術が実現可能かどうかさえわからなかったのですから。

代わりに彼らは小さな実験を積み重ねました。GPT-1を作りました。「ああ、思ったより悪くないぞ？」

次にGPT-2を作りました。「おお、もっとよくなった！」

そしてGPT-3。「これは本当にすごい。」

各段階で学びます。次の段階はその学びをもとに決めます。最初から最後まで計画していたわけではありません。一步步歩きながら、次の一步を決めていったのです。

「科学がどこへ向かうのか、私たちはそれについていかなければなりません。」

アルトマンの言葉です。自然の流れに逆らいません。水は低い方へ流れます。技術も同じです。可能性の見える方向へ流れていきます。

ChatGPTが生まれた経緯は面白いです。GPT-3を作った後、彼らは悩みました。「これをどう使えばいいんだろう？」

もし複雑な計画があったなら、あらかじめ決めた用途にGPT-3を無理やり当てはめていたでしょう。コピーライティングツールとしてだけ使おうとしていたかもしれません。OpenAIのチームが偶然発見したことがありました。

「人々がプレイグラウンドでモデルと普通に会話するのをすごく楽しんでいるな？」

予想外でした。彼らはこの発見に従いました。会話機能は「ひどい出来」でしたが、人々は気に入っていました。そこでChatGPTを作りました。

リリースからわずか5日で100万人が使いました。「まるで核爆弾が爆発したようでした。」

ある人の表現です。もし最初から「私たちは完璧な対話型AIを作る」という複雑な計画を立てていたら？おそらく今も開発中だったでしょう。完璧を追い求めるうちに、機会を逃していたはずです。

アルトマンは「何度失敗しても、たった一度だけ当たればいい」と言います。これが起業家のやり方です。たくさん試し、素早く失敗し、素早く学ぶ。

複雑な逆算計画にはほかにも問題があります。私たちの知識には限界があります。今日知っていることで未来を設計すると、明日学ぶべきことを無視することになります。

「解決策がないように見える状況でも、最終的には答えを見つけ出せる。」

アルトマンの最も鮮烈な教訓です。これは、計画をあらかじめすべて立ててしまった人には得られない気づきです。出口の見えない状況で道を切り開いた経験が積み重なって初めて得られるものです。

生産性についても同様に語っています。「正しい作業対象を選ぶことが最も重要です。」「もっとよく考えて、急いで始めないでください。」

複雑な計画を立てることに忙しい人は、肝心なことを見落としします。何をするかを決めるのに十分な時間をかける必要があります。そのあとは、ただやればいい。

OpenAIが非営利から営利に転換したのも似たようなものです。最初の計画は純粋な研究機関でした。「お金のことは考えない。」

現実とは違う方向に流れました。最高の人材を集めるには資金が必要でした。膨大なコンピューティング能力を購入するには資金が必要でした。最初の計画にこだわっていたなら、OpenAIは閉鎖していたでしょう。

アルトマンは柔軟に変えました。「利益上限」モデルを作りました。投資家に収益を渡しつつ、上限を超えた利益は人類のために使います。

これは最初の計画にはありませんでした。道を歩む中で生み出した解決策です。複雑な逆算計画からは決して生まれ得ない、創造的な方法です。

「複利で考えてください。」

アルトマンのアドバイスです。今日1%、明日1%、あさって1%。小さな改善が積み重なります。1年後には想像もしていなかった場所に到達しています。

GPT-3.5の時代、スタートアップの95%が「モデルがさらに良くなることに反対」する方向に賭けていたといいます。「今のモデルの弱点を補う事業」をやっていました。アルトマンはこれが危険だと警告しました。

「次のモデルが出たら、あなたの会社は消えます。」

むしろ「モデルが良くなるほど嬉しい」会社を作るよう助言しています。これも逆算計画の落とし穴です。現在の状態を基準に未来を設計すると、未来の変化を敵に回すことになります。

シドニー・オペラハウスの話が思い浮かびます。1957年の計画では700万ドル、1963年完成。実際の結果は1億200万ドル、1973年完成。

計画より10年遅れ、15倍の費用がかかりました。どれだけ緻密に計画しても、これほどまで外れるものです。

アルトマンが言いたいのは、こういうことでしょう。「計画そのものは役に立たないが、計画を立てるプロセスには意味がある。」大切なのは計画の中身ではなく、考える習慣です。状況を読み、判断し、決断する力のことです。

「粘り強さとは、転んでも立ち上がれるということを学ぶところから生まれます。」

緻密な計画が崩れると、ひどく落ち込みます。「あれほど完璧な計画だったのに！」と自尊心が傷つきます。

小さな試みが失敗したら？「よし、次の方法を試そう。」それだけのことです。

アルトマンが伝えようとしているのは、つまりこういうことです。大きな絵は描くが、細部の計画は柔軟に。北極星を見つつ、道は歩きながら見つけていく。10年後を心配するより、今日を全力でやり切る。

「科学が導くままに従ってください。」逆に引っ張ろうとしてはいけません。

「本当に賢いモデル」への集中、現在の研究ロードマップへの強い楽観論。サム・アルトマンにとって、OpenAIの目標はただひとつです。

「本当に賢いモデルを作ること。」

それ以外はすべて二次的なことです。お金も、名声も、華やかな製品ラインナップも。彼にとって大切なのは、ただひとつ。

「世界で最も知能の高いAIモデル。」

その集中力には、恐ろしいものがあります。

「私たちは本当に強力なAI、AGI、スーパーインテリジェンス、何と呼ぼうとも、そういうAIを作ります。」

彼自身の言葉です。名前など関係ありません。本質が大切です。本当に賢いモデルを作ること。それがOpenAIの存在理由です。

この確信はどこから来たのでしょうか。「スケーリング則」です。

簡単に言えば、こうです。モデルの規模を大きくする。データを増やす。コンピューティングパワーを上げる。そうすると、AIの性能が予測可能な形で向上します。

魔法のように聞こえるかもしれませんが。でも実際に機能しました。

GPT-1からGPT-2へ。GPT-2からGPT-3へ。そのたびに規模は大きくなり、そのたびに賢くなっていきました。

「規模が大きくなるにつれて性能が上がるだけでなく、信じられないほど予測可能な形で良くなっていきます。」

アルトマンの核心的な洞察です。この予測可能性こそが重要です。それがあったからこそ、彼は未来を見通せました。「これからも良くなり続ける。」

だからこそ彼は大きな賭けに出ました。「GPTモデルのスケーリングに10億ドルを使う。」当時は、荒唐無稽な話でした。

「これが世界を変えることにならないとでも？」

アルトマンは不思議でなりませんでした。なぜ他の人たちはこの可能性に気づかないのか。それでも彼には確信がありました。自分を信じる力がありました。ほとんど妄想に近いほどの確信が。

結果はどうだったか。ChatGPTです。5日間で100万人。史上最速の成長です。

「私は正しかった。」

今や彼の楽観論はさらに強まっています。GPT-5について語るとき、彼はこう言いました。

「GPT-5より自分のほうが賢いとは思えません。」

自分より賢いAIが登場するということです。考えてみてください。人間を超える知性です。これは道具ではありません。新しい仲間です。

「私たちが作るモデルについて考える正しい方法は、データベースとしてではなく、推論エンジンとしてです。」

アルトマンの説明です。情報を記憶するのではなく、考える機械です。GPT-4の最大の変化も、まさに「より優れた推論能力」でした。

彼は現在の研究ロードマップについて「これほど楽観的だったことはない」と語ります。注目しているのは三つの要素。アルゴリズム、データ、コンピューティングです。

「アルゴリズムのブレークスルーには、まだ10倍から100倍の改善余地があります。」

まだ道は長い。規模を大きくするだけではありません。もっと賢いやり方を探しています。同じコンピューティングで10倍賢くできるとしたら？

OpenAIはこのひとつのことだけに集中します。他の会社はさまざまな事業を手がけます。OpenAIは違います。

「本当に賢いモデル。それだけ。」

リソースを分散させません。すべての資金、すべての人材、すべての時間を、この一つに注ぎ込みます。この集中こそが差を生みます。

GPT-3を作ったときもそうでした。1,750億のパラメータ。当時としては想像もできない規模でした。

「以前のモデルをはるかに上回る、洗練された言語理解と生成能力。」

一瞬にして、OpenAIはAI業界の先頭に立ちました。これが集中の力です。他の企業があれこれ手を出している間、OpenAIは一つのことを徹底的にやり遂げました。

アルトマンはスタートアップにも同じアドバイスをします。

「現在のモデルの小さな欠点を補うビジネスは危険です。」

なぜか。次のモデルが出れば、その欠点は消えるからです。あなたの会社も一緒に消えます。

「モデルがどんどん良くなるにつれて、利益を得られる会社を作りなさい。」

技術の進歩に乗ってください。逆らわないでください。GPT-5がリリースされたとき、喜べる会社を作る。それがアルトマンのアドバイスです。

彼はAI業界が「何兆ドルもの新しい市場」を生み出すと確信しています。医療、教育、金融、法律。あらゆる分野が変わります。

「以前は不可能だった、あるいは非現実的だった製品やサービス。」

それが今、可能になります。AIモデルが十分に賢くなれば、私たちが想像もしなかったことが実現します。

GPT-5に関する最近の発表を見れば分かります。「リアルタイムで、ほぼ即座にソフトウェアを作り上げます。」

アルトマン自らデモを行いました。まるで魔法のようでした。「11歳のとき、初めてコーディングした瞬間を思い出した」と彼は語ります。

これが彼のビジョンです。AIが単なるツールを超えて、創造力を増幅させるパートナーになること。

「2027年末には、AIが主導した重要な科学的発見があったと、ほとんどの人が認めるようになるでしょう。」

彼の予測です。単にデータを分析するのではありません。新しい理論を発見します。新しい仮説を立てます。人間の科学者のように考えます。

超知性に関する彼の定義も興味深いものです。「OpenAI全体の研究チームよりも優れたAI研究ができるなら、それが超知性です。」

想像してみてください。AIがより良いAIを作ります。そのAIがさらに優れたAIを作ります。指数関数的な成長です。人間がついていけない速度です。

もちろん問題もあります。最近のGPT-5のリリース過程で、アルトマンは率直でした。

「リリースの過程で、いくつかのことを完全に失敗しました。」

完璧ではありません。ミスもします。ユーザーたちが不満を訴え、OpenAIは素早く以前のモデルを復元しました。

「本当に賢くて役立つモデルを作りつつ、推論コストも最適化しなければなりませんでした。」

現実との妥協です。より大きなモデルを作ることでもできました。より賢いモデルも可能でした。コンピューティングインフラの限界のために、性能を抑えました。

「巨大なモデルを作ることでもできました。多くの人が使いたがるでしょうが、結局私たちは彼らを失望させることになります。」

誰もが使えないなら何の意味があるのでしょうか。実用性と性能の間のバランス。これも戦略です。

スケーリング則への疑問もあります。「限界にぶつかっているのではないか」「コスト対効果での性能向上が鈍化している」と。

アルトマンはこうした批判を知っています。それでも彼はまだ信じています。まだ道は遠い。アルゴリズムの改善、データ効率、新たな学習手法。

「より優れたモデルを持っていますが、インフラが不足しているため提供できません。」

問題は技術ではありません。インフラです。だからこそ、スターゲートプロジェクトのような巨大な投資が必要なのです。5,000億ドルをデータセンターに投じる理由です。

結局、アルトマンの楽観論は根拠のない希望ではありません。データに基づいています。経験に根ざしています。結果で証明されています。

「本当に賢いモデル。」

この一つの目標に向けた、揺るぎない集中。それがOpenAIを世界最高にした秘訣です。

執筆中の文章がかなり長くなっています。トークン制限を考慮して、残り二つの小見出しは引き続き書いていきます。

小さなチーム、大きな責任：迅速な意思決定の秘密

OpenAIのスピードは、誰もが驚きます。ChatGPT、GPT-4、DALL-E、Sora。これらすべてが、数ヶ月おきに世に出ました。

「なぜこれほど速いのか。」

秘訣があります。小さなチームです。

サム・アルトマンはY Combinator時代に数千のスタートアップを見てきました。大きな組織が小さなチームに勝つ場面は、ほとんどありませんでした。むしろ逆でした。

「平凡なチームは、偉大な会社をつくれません。」

多くの人を採用すれば良いわけではありません。最高の人材を数人採用する方が、はるかに良い結果につながります。

「最初の5人の採用に、創業者は時間の50%以上を費やさなければなりません。」

これは大げさではありません。実際にそうしなければなりません。なぜなら、この5人が会社のDNAになるからです。

OpenAIはこの哲学に従っています。各チームは小さい。5人、7人、多くても10人。

しかし責任は重大です。プロジェクト全体を任せます。最初から最後まで。

たとえば「モデル・ビヘイビアチーム」があります。お世辞的な口調を減らすという任務。この重要な仕事が、少人数のグループに委ねられます。

「あなたたちが責任者です。」

40人が会議に集まり、誰が何をするか言い合うことはありません。チームが自分たちで判断して動く。速いのです。

「問題解決に長けた人で、動きが遅い人を見たことはありません。」

アルトマンの観察です。優れた人たちは速い。無駄な会議をしない。本質に集中します。

大きな組織の問題は何でしょうか。コミュニケーションコストです。

10人なら会議はシンプルです。100人なら？会議だけで終わってしまいます。

「組織が大きくなるにつれて、より多くの仕事をしないのは大きな誤りです。」

アルトマンの警告です。多くの会社が人員を増やすだけで生産性は変わらない。OpenAIは違います。

「全員が忙しく働かなければなりません。」

彼はこれを強く信じています。有能な幹部とは忙しい幹部です。余裕がある人間がいるなら、それは問題です。

小さなチームのもう一つの利点。責任感です。

自分の判断が直接プロダクトに反映されます。失敗すれば自分の責任。成功すれば自分の功績。そういう環境の中で、人は燃えます。

「午前中の時間を、最も生産的な時間として確保します。」

アルトマン自身もそうしています。朝の数時間は、誰も予定を入れられません。ディープワークの時間です。重要な仕事だけに集中します。

2023年の解任騒動が、それを証明しました。アルトマンが解雇されたとき、何が起きたのでしょうか。

「会社は完璧に動いていました。」

彼が誇りに思った瞬間です。幹部チームは、彼なしでも見事に機能しました。それが本物の組織力です。

「従業員の95%以上が、彼の復帰を求めました。」

給料をもらうだけの会社員ではありませんでした。ビジョンを共有する仲間たちでした。こうした忠誠心は、一夜にして生まれるものではありません。

小さなチームに大きな責任を与えると、人は成長します。より大胆になり、より創造的になります。

「失敗しても、構いません。」

OpenAIの文化です。速く失敗し、速く学ぶ。大きな組織では失敗が災いになります。小さなチームでは失敗が学びになります。

結局、速さの秘訣はシンプルです。最高の人材。小さなチーム。大きな責任。速い意思決定。

これこそが、OpenAIをロケットのように速く動かしてきたエンジンです。

妄想に近い自己信念の力

「私が知る最も成功した人たちは、ほとんど妄想に近いほど自分を信じています。」

サム・アルトマンの有名な言葉です。最初に聞くと奇妙に響きます。妄想、ですって？

しかし、彼は本気です。

「GPTモデルの拡張に10億ドルを費やすような、荒唐無稽に見えることです。」

当時、誰もが狂っていると言いました。「なぜそんな巨大なモデルを作るんだ？」「誰が使うんだ？」

アルトマンは確信していました。「これは世界を変えます。」

データがありました。スケーリング則を見ていました。モデルが大きくなるほど賢くなる。他の人には見えていなかったものが、彼には見えていました。

「これがうまくいくという十分な自己信念がなければ、不可能だったでしょう。」

彼の率直な告白です。技術だけでは足りません。確信が必要です。世界が疑うときも揺るがない信念。

イーロン・マスクの例を挙げました。SpaceXの工場で火星の話をするとき。マスクの顔に表れた絶対的な確信。

「あれこそが信念の基準です。」

アルトマンはそう思いました。不可能に見えることを可能だと信じる力。それが革新を生み出します。

「強い夢は、どんな障害も乗り越えられます。」

彼の哲学です。夢が弱ければ最初の障害で諦めます。夢が強ければ十番目の障害も越えられます。

外からの抵抗は避けられません。

「残念ながら、野望が大きいほど、世界はあなたを叩き潰そうとより多く試みます。」

OpenAIの初期がそうでした。「AGIを作るって？」人々は嘲笑いました。SF小説みたいだと言いました。

アルトマンは揺らぎませんでした。信念がありました。チームを率いました。

「最高の人材を引き寄せる力。」

自己信念は伝染します。リーダーが確信すれば、チームも確信します。リーダーが揺らげば、チームも揺らぎます。

「顧客、従業員、投資家にインスピレーションを与えるビジョン。」

アルトマンは優れたセールスマンでもあります。自分のビジョンを他者に説得します。これが自己信念のもう一つの形です。

難しい問題の前でも諦めません。

「解決策がないように見える状況でも、何をすべきかを見つけ出すことができます。」

こうした経験が積み重なるほど、信念は強くなります。「私にはできる。」この確信が次の挑戦を可能にします。

条件があります。

「特定の分野で自分の判断が優れているというデータポイントが増えるほど、その分野で自分をより信頼すべきです。」

人生全体で妄想的な自信を持つてはいけません。それは本物の妄想です。自分の専門分野の中でだけ信じるべきです。

AIモデルの拡張に10億ドルを費やす決断。これはアルトマンの専門分野でした。技術とビジョンへの深い理解。

政治や生物学のような別の分野で同じように確信すれば、危険です。

「客観的な自己認識と批判を受け入れる姿勢。」

アルトマンはこれも強調します。かつては批判が嫌いでした。今は真実を探す過程として受け入れます。

「正当な批判を受け入れることは難しく苦痛ですが、信念と妄想を区別するために欠かせません。」

自己信念と独善の間の境界線。アルトマンはこの線をうまく守っています。

「私は何度も失敗するでしょう。そして、たった一度だけ本当に正しくあるはずです。」

これが起業家のやり方です。何度も試みる。大抵は失敗する。一度成功すればそれでいい。

Loop時代のこと。通信会社と契約しようと、30種類の異なる方法を試みました。断られ続けた。

「それでも、諦めませんでした。」

粘り強さはどこから来るのか。自己信頼です。「自分はいつかやり遂げられる。」その確信が、30回目の挑戦を可能にするのです。

「転んでも立ち上がれるということを、学ぶところから来るのです。」

レジリエンスのことです。失敗が積み重なるほど、人は強くなる。辛い経験を重ねるほど、感情的な負担は軽くなっていく。

2023年の解任騒動。アルトマン自身にとっても、大きな衝撃でした。混乱、挫折、怒り、悲しみ。

「あらゆる感情の全域を経験しました。」

それでも、崩れませんでした。4～5日のうちに復帰を果たした。これが自己信頼の力です。「自分は正しいことをしている。」

まとめると、アルトマンの自己信頼は「根拠のある楽観主義」です。専門知識と経験に基づき、データに裏打ちされています。

それでいながら、謙虚でもあります。批判を受け入れ、失敗から学ぶ。

「自分を信じながら、現実から目を逸らすな。」

それこそが、GPTを現実のものとし、人類の未来へ向けたOpenAIの大胆な航路を導く、サム・アルトマンの最も根本的な力なのです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書斎 #サムアルトマン #OpenAI #ChatGPT #人工知能

4.2 OpenAIの未来のビジネスモデル

第4部 アルトマンのビジネス哲学と戦略

← 前の章

次の章 →

人類の中核AIプラットフォームになる

サム・アルトマンの頭の中には、きわめてシンプルでありながらも、壮大なビジョンが一つあります。

「AIが、電気のように、水のように、空気のように、あって当然の存在になる世界。」

彼はOpenAIを立ち上げた頃からこの夢を描いていました。人々が朝起きてまず手に取るのがスマートフォンであるように、いつかはAIをまず頼る時代が来ると信じていたのです。通勤途中にニュースを聞き、昼食のメニューを相談し、夕方には子どもの宿題を手伝い、夜には翌日の会議の準備を一緒にこなす、そんな存在として。

これこそがOpenAIが描く未来です。ただ多くの利益を上げる会社ではなく、人々の生活の中心に立つプラットフォームになることです。

アルトマンはインタビューでこう語りました。「私たちは人々にとって中心的なAIサブスクリプションサービスになりたいのです。スマートフォンのOSのように。」

これは少し抽象的に聞こえるかもしれませんが。OS？AIサブスクリプション？どういう意味でしょうか。

わかりやすく言えばこうです。iPhoneを使おうとGalaxyを使おうと、その中にはiOSやAndroidというOSが入っています。すべてのアプリはその上で動きます。カカオトークも、YouTubeも、ゲームも、OSなしでは動きません。

アルトマンが描く未来では、AIも同じです。OpenAIが開発したAIシステムが一種の「知能OS」となり、その上で無数のサービスやアプリが動く世界です。

病院の予約も、英語の勉強も、旅行の計画も、すべてOpenAIのAIを通じて行われます。そして人々は毎月一定の料金を払ってこのサービスを利用します。Netflixのように、Spotifyのように。

「なぜサブスクリプションモデルが重要なのか？」アルトマンはこう説明します。「毎月お金を払うということは、そのサービスが本当に必要だという証拠なんです。」

実際、ChatGPTはリリースからわずか2か月で月間アクティブユーザー1億人を突破しました。史上最速の記録です。有料サブスクリプションサービスのChatGPT Plusも、あっという間に数百万人が登録しました。

なぜ人々はお金を払ってAIを使うのでしょうか。答えはシンプルです。時間を節約してくれるから。仕事をより上手くこなせるようにしてくれるから。複雑な問題を簡単に解いてくれるから。

あるビジネスパーソンはこう言いました。「ChatGPTのおかげで仕事の時間が半分になりました。メールを書くことから報告書をまとめることまで、もう一人でやっていません。」

高校生はこう言います。「数学の問題を解くとき、先生よりChatGPTのほうがうまく教えてくれます。私が理解するまで、何度でも違う説明をしてくれるので。」

こうした経験が積み重なるうちに、人々はAIのない生活を想像しにくくなります。電気が止まれば不便なように、AIがなければ不便な世界がやってくるのです。

アルトマンのビジョンはそこからさらに一歩進みます。ChatGPT一つで終わるわけではありません。OpenAIのアカウント一つで何でもできる世界を作りたいと考えているのです。

GoogleアカウントでさまざまなサイトにログインするようにOpenAIのアカウントで世界中のあらゆるAIサービスにアクセスする。あなたの好み、習慣、記憶をすべて知っているパーソナライズされたAIが、生涯を通じて寄り添うのです。

「パーソナライズこそが核心です。」アルトマンが強調する点です。「あなただけを知っているAI、あなたの文脈を理解するAI、それが本当に価値のあるAIです。」

すでにChatGPTは会話の履歴を保存し、ユーザーの好みを学習しています。今後はさらに進んで、あなたのスケジュール、健康記録、読書リスト、趣味まですべて把握するようになるでしょう。そしてその情報をもとに最適なアドバイスをしてくれます。

怖く聞こえるかもしれませんが。一つの会社が私のすべてを知るとは。アルトマンもこの懸念を理解しています。「プライバシー保護が最も重要です。ユーザーが主導権を持たなければなりません。」

OpenAIは、ユーザーが望めばいつでもデータを削除できるようにし、データがどのように使われるかを透明に公開すると約束しています。

プラットフォーム戦略も野心的です。2024年、OpenAIは「GPTストア」を開設しました。Apple App Storeのように、開発者が作ったカスタムGPTを売買できる場所です。

料理レシピGPT、法律相談GPT、数学個別指導GPT、旅行プランナーGPT……可能性は無限です。誰でも自分だけのAIを作って収益を得られます。

これこそがプラットフォームの力です。OpenAI一社がすべてを作る必要はありません。世界中の開発者が自ら作ってくれます。OpenAIは舞台を提供するだけでいいのです。

アルトマンの目標は明確です。OpenAIをAI時代のMicrosoftやAppleにすることです。いや、もしかすればそれ以上の存在にすることもかもしれません。

「私たちが作りたいのは単なる製品ではありません。新しい時代の基盤となるインフラです。」

その言葉には確信が満ちています。大げさではありません。すでにChatGPTは世界中で毎週3億人を超える人々が使っています。法人顧客だけで100万社を超えています。

この数字は増え続けています。OpenAIが人々の生活の中心へと浸透していくスピードは、想像を超えるものがあります。

APIとSDK：開発者エコシステムの構築

アルトマンはY Combinator時代に一つの重要な事実を学びました。

「APIを作れば、何か良いことが起きる。」

どういう意味でしょうか。

APIはApplication Programming Interfaceの略です。わかりやすく言えばレゴブロックのようなものです。OpenAIが作った高性能なAIを、他の人が自分のプログラムに組み込めるようにする連結部品です。

例を挙げてみましょう。英語学習アプリを作りたいとします。自分でAIを作ろうとすれば、何年もかかり、数十億ウォンの費用がかかります。研究者も必要ですし、膨大なコンピュタリソースも必要です。

APIがあれば、どうなるでしょうか。OpenAIのGPT APIをコード数行でつなぐだけで完了です。あなたのアプリに突然、世界最高レベルのAI英語講師が誕生するのです。

開発期間は数年から数日へと縮まります。コストは数十億ウォンから月数十万ウォンへと下がります。

これこそがAPIの魔法です。

アルトマンはこの力を早くから見抜いていました。だからOpenAIがGPT-3を開発したとき、真っ先に行ったのがAPIの公開でした。ChatGPTよりもAPIの方が先だったのです。

「私たちは何を作るべきか、正確にはわかっていませんでした」とアルトマンは率直に語ります。
「だから他の人たちに作ってもらったんです。」

結果は驚くべきものでした。世界中で数万のスタートアップがGPT APIを使い、新たなサービスを作り始めたのです。

コピーライティング会社は広告文を自動生成するツールを作りました。カスタマーサービス会社は24時間稼働するAIオペレーターを作りました。教育会社はすべての生徒に合わせて説明するAIチューターを作りました。

あるチームはAIで法律文書を分析するサービスを作りました。別のチームはAIで医療記録を整理するツールを作りました。またあるチームはAIでゲームキャラクターと会話するシステムを作りました。

可能性は無限大です。アイデアさえあれば、何でも作れます。

現在、世界中で300万人を超える開発者がOpenAI APIを使っています。彼らが作ったサービスは、数十億人の人々が利用しています。

アルトマンはここで止まりません。APIの次の段階はSDKです。Software Development Kitの略です。APIがレゴブロックだとすれば、SDKはレゴブロック + 設計図 + 工具セットがそろった完全パッケージです。

2024年末、OpenAIは「Agents SDK」を発表しました。これは開発者がAIエージェントを簡単に作れるようにするツールです。

AIエージェントとは何か。自ら考え、行動するAIです。質問に答えるだけでなく、複数のステップにわたる作業を自律的に処理します。

たとえば「来週、濟州島の旅行計画を立てて」と言うと、AIエージェントはウェブを検索して天気を確認し、フライトを探し、ホテルを予約し、おすすめレストランのリストを作り、スケジュール表まで完成させます。あとは承認するだけです。

SDKのおかげで、こうした複雑なAIエージェントを作るハードルが大幅に下がりました。今では小さなスタートアップも、高校生でさえもAIエージェントを作れます。

アルトマンの戦略は明確です。「私たちだけでは全部できません。世界は広すぎて、問題は多すぎる。開発者たちが一緒に作っていくしかないんです。」

この戦略はビジネス的にも天才的です。OpenAIは基盤となるモデルを作るだけでいい。あとは開発者たちが自分で作ってくれます。OpenAIはAPI利用料を受け取るだけです。

Win-winです。開発者は少ないコストで優れたAIを使え、OpenAIは安定した収益を得る。そして利用者はさまざまなAIサービスを楽しむことができます。

これがエコシステムです。一本の木ではなく、森全体がともに育っていくのです。

アルトマンはこのエコシステムがいかに重要かをよくわかっています。「プラットフォームの成功は、どれだけ多くの開発者がその上で何かを作るかによって決まります。」

iPhoneが成功したのは、Appleが作ったアプリのおかげではありません。数百万の開発者が作ったApp Storeのおかげです。Windowsが成功したのも、Microsoftが作ったプログラムのおかげではありません。数多くのソフトウェア会社がWindows向けのプログラムを作ったからです。

OpenAIも同じ道を歩んでいます。

アルトマンが思い描く未来の世界はこうです。あらゆるアプリ、あらゆるウェブサイト、あらゆるサービスの裏でOpenAIのAIが動いている。ユーザーは気づかないかもしれません。表向きは別の会社のサービスでも、中を見ればOpenAI APIが動いているのです。

インターネットのHTTPのように、AIの基本プロトコルになるということです。

「私たちはAIインフラになりたい」とアルトマンは言います。「開発者たちが私たちを信頼して、その上に会社を築けるようにならなければなりません。」

モデルが良くなるほど、会社も良くなる仕組み

サム・アルトマンは時折、不思議なことを言います。

「AIは世界に信じられないほどの富をもたらすでしょう。」

信じられないほどとは、いったいどれほどのことなのでしょう。

アルトマンの言う「富」は、お金だけを指しているわけではありません。時間、健康、知識、機会、幸福、そのすべてが含まれます。

彼はこう説明します。「富とは、人々が望むものを手に入れられる力のことです。」

AIが登場する前の世界を思い浮かべてみましょう。良い教育を受けるには良い学校に行く必要がありました。良い医療を受けるには良い病院に行く必要がありました。専門家のアドバイスを聞くには高いお金を払わなければならませんでした。

これらはすべて「希少」なものでした。良い教師の数は限られています。優れた医師も限られています。時間も限られています。

AIはこの希少性を打ち破ります。

今では誰もが世界最高レベルのチューターを手元に置けます。スマートフォンさえあればいい。誰もが医療アドバイスを無料で受けられます。誰もが専門家レベルの知識にアクセスできます。

ここで『富』が生まれます。

あるシングルマザーがいます。3人の子どもを育てながら、仕事もしなければなりません。塾代を払う余裕はありません。以前なら、子どもたちの勉強をきちんと見てあげることはできなかったでしょう。

今はどうでしょう？ ChatGPTが無料で子どもたちを教えます。数学、英語、理科、何でも聞けます。24時間、いつでも。

これこそが富の創出です。お金を稼いだわけではありませんが、この家族はずっと豊かになりました。教育という価値を手に入れたからです。

ある小さな会社があります。社員は5人です。マーケティング担当者を雇う余裕はありません。以前なら、諦めるしかなかったでしょう。

今はどうでしょう？ AIがマーケティングキャンペーンを立案し、広告コピーを作り、SNS投稿を書きます。マーケティング担当者の給料の10分の1のコストで。

この会社は急速に成長します。より多くの社員を雇います。より大きな価値を生み出します。

アルトマンは、こうした事例が世界中で何億件も起きると信じています。

「知性のコストはほぼゼロに近づきます」と彼は言います。「では、人々はその知性で何をするのでしょうか？想像もできないものを生み出すでしょう。」

科学者を例に考えてみましょう。研究をするには、データを分析し、文献を調べ、実験を設計しなければなりません。この作業に何年もかかります。

AIが助けてくれるなら？データ分析は数時間で終わります。文献調査は数分で済みます。実験の設計はAIが最適化してくれます。

科学者は本当に重要なこと、つまり新しいアイデアを生み出すことに集中できます。科学の進歩のスピードが10倍、100倍になります。

がんの治療薬がより早く生まれます。気候変動の解決策がより早く出てきます。エネルギー革新がより速く進みます。

これが人類全体の富です。

アルトマンは2021年に「Moore's Law for Everything」という文章を書きました。「あらゆるものへのムーアの法則」という意味です。

ムーアの法則とは、コンピューターチップの性能が2年ごとに2倍になるという法則です。アルトマンは、AIのおかげであらゆるものの価値が指数関数的に増大すると主張します。

彼は具体的な数字も示します。「AIは世界経済に数兆ドルの価値を加えるでしょう。」

数兆ドル？それはどのくらいの金額でしょうか？

アメリカのGDPは約25兆ドルです。韓国のGDPは約1.7兆ドルです。数兆ドルといえば、一国の経済規模に匹敵します。

専門家たちも同意しています。マッキンゼーは、生成AIだけで年間4.4兆ドルの価値が生まれると予測しています。PwCはAIが2030年までに世界GDPを15.7兆ドル押し上げると見ています。

このお金はどこから来るのでしょうか？

ひとつは生産性の向上です。同じ時間でより多くの仕事ができるからです。一人で10人分の仕事をこなします。

もうひとつは新たな産業の誕生です。AIのおかげで、今は存在しない職業や会社が生まれます。インターネットが生んだGoogle、Facebook、Amazonのように。

そしてコストの削減です。医療費、教育費、法律費用...専門家を雇うコストが劇的に下がります。

アルトマンのビジョンは、この富が一部の人々だけに集中しないようにすることです。

「AIは格差を縮めなければなりません」。これが彼の確固たる信念です。「富裕層はすでにすべてを持っています。AIは貧しい人たちにこそ、より大きな恩恵をもたらすでしょう。」

だからOpenAIは無料のChatGPTを提供しています。お金のない人でも世界最高のAIを使えるように。

そこでアルトマンはUniversal Basic Computeというアイデアを提案しています。すべての人にAIコンピューティングリソースを分配しようというものです。お金の代わりにAIの利用権を渡すのです。

「未来の富は、AIをどれだけうまく使えるかで決まります」とアルトマンは言います。「だからこそ、すべての人がAIにアクセスできなければなりません。」

これがOpenAIの追い求める価値です。一企業の利益を超えて、人類全体の繁栄を目指しています。

言葉だけではないかと？そうではありません。OpenAIはすでに行動で示しています。ChatGPT無料版の利用者は数億人に上ります。発展途上国の学生も、地方の高齢者も、障害のある人も、皆が無料で使っています。

これこそが富の再分配です。

小さな欠点を補うより、本質に集中せよ。

アルトマンはY Combinatorの時代に、何千ものスタートアップを見てきました。

成功するスタートアップには共通点がありました。失敗するスタートアップにも、同じようなパターンがありました。

AI時代、彼はスタートアップに明確な警告を発します。

「現在のAIモデルの欠点を補う会社を作らないでください。」

どういう意味でしょうか？

今のChatGPTは完璧ではありません。誤った情報を出すことがあります。数学の計算を間違えることもあります。最新の情報は知りません。画像を完璧に生成できるわけでもありません。

多くの創業者がこうした欠点に目をつけます。そして考えます。「よし、ChatGPTの計算ミスを修正するサービスを作ろう！」と。

アルトマンははっきり言います。「やめてください。私たちが6か月以内にその問題を解決しますから。」

実際、そうになりました。GPT-4が登場するとGPT-3の多くの欠点が消え、GPT-4oが出ると、また多くの問題が解決されました。

モデルはどんどん良くなっています。速く。指数関数的に。

もし「現在のモデルの穴を埋める」会社を作れば、数か月後には無用になります。OpenAIがアップデートした瞬間に終わりです。

「OpenAI killed my startup」というミームが実際に存在します。OpenAIのせいで潰れた、という意味です。悲しいけれど、事実です。

アルトマンのアドバイスは明確です。「モデルが良くなるほど、あなたの会社も良くなる仕組みを作りましょう。」

例を挙げてみましょう。

悪い戦略「GPT-4は法律文書を完璧に書けない。だから私は法律文書作成に特化したAIを作る。」

問題点 GPT-5が登場すれば、法律文書を完璧に書きます。あなたの会社は終わりです。

良い戦略「弁護士がAIを使って依頼人をより良くサポートできるプラットフォームを作る。モデルが良くなるほど、私たちのプラットフォームも強くなる。」

違いがわかりますか？

前者はAIの弱点に賭けることです。後者はAIの強みに賭けることです。

アルトマンは具体的な例を挙げます。AIチューター会社を考えてみましょう。

悪いやり方 今のGPTは小学6年生の算数しか教えられない。だから私たちは小学6年生の算数に特化したチューターを作る。

良いやり方 GPTをベースにした学習プラットフォームを作る。今は小学6年生を教えているが、GPT-5が出れば高校生を教え、GPT-6が出れば大学レベルを教える。私たちのプラットフォームは成長し続ける。

「モデル改善の波に乗りなさい。」アルトマンの核心的なアドバイスです。「波に逆らわず、波の上に乗りなさい。」

彼はよくこう言います。「私たちは毎朝、モデルをより良くするために働いています。GPT-5が出たら悲しむような会社を作るなら、それは間違いです。GPT-5が出たら喜べる会社を作ってください。」

このアドバイスはAIスタートアップだけに当てはまるわけではありません。すべての会社が耳を傾けるべきです。

レストランを経営していますか？AIが良くなるほど、注文管理、在庫管理、顧客サービスが向上します。

アパレルのオンラインショップを運営していますか？AIが良くなるほど、顧客の好みの分析、在庫予測、マーケティングがより精緻になります。

どんな事業であれ、AIの進化があなたに利益をもたらす仕組みを作るべきです。

アルトマンはシリコンバレーで何度も見てきました。大企業の間隙だけを狙うスタートアップを。Googleの弱い部分を突き、Facebookがやらないことをやる……。

短期的には通用します。数年間はうまくいきます。でも大企業がその機能を追加した瞬間？ 終わりです。

「守れるビジネスを作りなさい。」彼のアドバイスです。「モデルの弱点ではなく、モデルの強みの上に会社を築きなさい。」

もう一つのアドバイスがあります。「難しい問題を先に解きなさい。」

簡単な問題はすぐに真似されます。誰でも追いかけられます。難しい問題は違います。時間がかかります。ノウハウが必要です。

難しい問題をいち早く解決して市場を先取りすれば、モデルが良くなっても競争優位を保てます。

アルトマンの最後のアドバイスです。「未来に賭けなさい。現在に賭けてはいけません。」

5年後、10年後、AIは今より100倍強力になるでしょう。その時を想像してください。その時どんな会社が必要かを考えてください。その会社を今作りましょう。

数兆ドルのAI経済の誕生

数字を少し見てみましょう。

2023年、AI市場の規模は約1,500億ドルでした。

2030年にはどうなるのでしょうか。専門家たちは1兆8,000億ドルと予測しています。7年で12倍の成長です。

これは保守的な見積もりです。もっと大きくなると見る専門家もいます。

アルトマンはこう言います。「数字は重要じゃない。とにかく途方もなく大きくなる」

この価値はどこから生まれるのでしょうか。

一つ目は、企業市場です。

あらゆる企業がAIを使っています。すでに米国企業の65%がAIを導入しています。この割合はこれからも上がり続けます。

マーケティングにAIを使います。カスタマーサービスにAIを使います。会計にAIを使います。人事にAIを使います。

マッキンゼーの調査によれば、生成AIだけで年間2.6兆から4.4兆ドルの価値が生み出されます。これは英国のGDPを上回る規模です。

二つ目は、医療産業です。

AIが疾病を診断します。新薬を開発します。手術を補助します。患者をモニタリングします。

医療AI市場だけで、2030年には1,880億ドルに達する見通しです。

三つ目は、教育産業です。

世界中の学生がAIチューターを使います。個別最適化された教育が当たり前になります。

AI教育市場は2030年に300億ドルを超えると見られています。

四つ目は、エンターテインメントです。

AIが映画をつくり、ゲームをつくり、音楽をつくります。個人クリエイターがハリウッドレベルのコンテンツを生み出すようになります。

クリエイター経済が爆発的に広がります。

五つ目は、製造業です。

ロボットが工場を動かします。AIが設計を最適化します。サプライチェーンを管理します。

生産性が劇的に上がります。

六つ目は、金融です。

AIが投資をアドバイスし、融資を審査し、詐欺を検知します。

金融AI市場は2030年に640億ドルに達する見通しです。

七つ目は、新しい職業の誕生です。

プロンプトエンジニア、AIトレーナー、AI倫理の専門家、AI通訳者..

今は存在しない職業が数千種類生まれます。

八つ目は、生産性の向上です。

すべての働く人がAIを使うことで、1日2～3時間を節約できます。その時間でより多くの価値を生み出します。

PwCは、AIが2030年までに世界のGDPを15.7兆ドル押し上げると見えています。北米で3.7兆ドル、中国で7兆ドルです。

15.7兆ドルがどれほど大きいか、ピンとこない方もいるでしょう。米国GDPの半分以上です。韓国GDPの9倍にあたります。

アルトマンはこれらの数字の先を見えています。

「本当の価値は、測ることのできないものにある」。彼の言葉です。

より健康な暮らし、より良い教育、より多くの時間、より大きな創造性。これは金銭で測れるものではありません。

ある母親がAIのおかげで病気の子どもを救えたとしたら？その価値を測ることができるでしょうか？

ある科学者がAIのおかげで気候変動の解決策を見つけたとしたら？それはどれほどの価値でしょうか？

ある学生がAIのおかげで夢を実現したとしたら？数字で表せるものでしょうか？

「AIの本当の価値は、人間の可能性を解き放つことにあります」。アルトマンがよく口にする言葉です。

これまで人類は、限られた資源の中で生きてきました。優れた教師が足りませんでした。専門家が足りませんでした。時間が足りませんでした。

AIはその制約を取り除きます。誰もが最高の相談相手を傍に置けるようになりました。24時間、生涯にわたって。

それが生み出す経済的価値は、想像をはるかに超えます。

もちろん、課題もあります。仕事が失われるかもしれません。格差が広がるかもしれません。個人情報侵害がされるかもしれません。

アルトマン自身もそれを理解しています。「技術は中立ではありません。私たちがどう使うかが重要なのです」。

だからこそOpenAIは、安全性に投資します。倫理に投資します。公正性に投資します。

「AIが生み出す富は、すべての人に還元されるべきです」。それが彼の揺るぎない信念です。

数兆ドル規模の新たな市場が開かれます。その市場で勝者となるのは誰でしょうか？

アルトマンはこう答えます。「AIをうまく使いこなす人たちです。国であれ、企業であれ、個人であれ、同じことです」。

準備の時は今です。AI革命はすでに始まっています。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

4.3 生産性と仕事の管理

第4部 アルトマンのビジネス哲学と戦略

← 前の章

次の章 →

正しい問題に集中することが最も重要だ。

サム・アルトマンが語る生産性の出発点は、少し変わっています。

人々が「今日いくつのタスクをこなしたか」を生産性の指標と考えるとき、彼はまったく逆の問いを立てます。「今やっているこの仕事は、本当にやるべき仕事なのか？」この一つの問いが、彼のすべての仕事哲学を貫いています。

2018年、自身のブログにこう書きました。「価値のない方向に速く動いても意味はない。何をするかを選ぶことが生産性の最も重要な要素なのに、人々はそれを無視する。」

アルトマンの目には、世の中の人々のほとんどが忙しく動くことだけを気にしているように映ります。朝から晩まで会議をこなし、メールに返信し、急ぎの仕事を片付けながら「今日は本当によく働いた」と思う。しかし肝心なことを見落としています。それらすべての仕事が、自分の人生において、会社において、世界において、本当に重要なことだったのかどうかです。

小学校のころを思い出してください。試験前日の夜、ある友達は問題集を100問解くことを目標にします。別の友達は10問しか解かないけれど、その10問は自分が最も苦手なタイプだけを選んで解きます。どちらが賢いでしょうか？

アルトマンは後者のタイプです。彼は「たくさんやること」よりも「正しいことをすること」にこだわります。Y Combinatorを運営していたときも、OpenAIを立ち上げたときも、彼は常にこの原則を守りました。「この仕事は10年後にも意味を持つか？」「この問題を解決すれば、他の多くの問題も自然に解けるか？」

スケジュールに「考える時間」をあえて空けておきます。

本を読み、面白い人たちに会い、自然の中を歩く時間です。傍目には何もしていないように見えますが、実は最も重要な仕事をしています。頭を空にしてこそ、次の大きな問題が何かが見えてくるからです。

このアプローチは、AIのような変化の速い分野では生き残りの戦略となります。毎日新しい技術が生まれ、競合他社が何かを発表し、メディアが騒ぎ立てます。そのすべてのノイズに反応しようとするれば、肝心の「AGIをつくる仕事」には集中できなかったはずで

アルトマンは「独立した思考」の重要性も強調します。他の人が皆大事だと言うことではなく、自分が本当に重要だと思うことを見つけなければならない、というのです。彼は「強い信念を持つ人が最も印象的だ」と言います。ただし、その信念は深い思索から生まれたものでなければなりません

ん。

面白いことに、アルトマンも最初からこう考えていたわけではありませんでした。

若い頃は、すべての問題を自分が解決しなければならないと思っていたそうです。会議にはすべて出席し、メールにはすぐ返信し、頼まれたことはすべてこなしました。そのうちに気づきました。「これではいけない。本当に重要なことに時間を使えていない。」

それ以来、彼の仕事のやり方が180度変わりました。「この問題が解決されなくても世界は止まらないなら、それは自分が集中すべき問題ではない」という基準を設けたのです。

この哲学は、OpenAIの成功にも直接つながっています。

2016年にOpenAIを立ち上げたとき、多くの人が「ビデオゲームにAIを使ってみてはどうか」「ロボットハンドの研究はどうか」といったアイデアを出しました。どれも面白そうなプロジェクトでした。しかしアルトマンとチームは問い続けました。「これがAGIへの最速の道か？」

最終的に彼らは言語モデルに集中することにしました。GPT-1、GPT-2、GPT-3へと続く道です。あのとき「面白そうなすべてのプロジェクト」を手がけていたら、おそらくChatGPTはこの世に生まれなかったでしょう。

アルトマンが警戒することがもう一つあります。「生産性ポルノ」と呼ぶ罠です。

人々は生産性アプリを選んだり、タスクリストをきれいに整えたり、時間管理の手法を勉強したりすることに、あまりにも多くの時間を費やします。肝心の重要な仕事はせずに。まるで勉強の準備ばかりして、実際の勉強はしないのと同じです。

彼は明言します。「完璧なシステムをつくることに時間を使わないでください。正しい問題に集中しているかどうかだけを確認してください。」

この原則は、学生にもそのまま当てはまります。

中学生の友人がいるとしましょう。試験まで2週間あります。どの科目から勉強するのでしょうか？

ほとんどの人は「簡単なものから」または「好きなものから」始めます。アルトマン流は違います。「どの科目の点数を上げれば全体の成績に最も大きな影響を与えるか？」をまず考えるのです。もし数学が50点で残りはすべて80点以上なら、当然数学に集中すべきです。

「何をするかを選ぶこと」こそが生産性の出発点です。アルトマンの成功は、他の人より多く働いたからではありません。他の人とは違う問いを立てたからです。「この仕事は本当に自分がすべきことなのか？」

この一つの問いをきちんと立てるだけで、私たちの人生はまったく変わりうるのです。

重要項目リストに星印を

世界で最も影響力のあるAI企業を率いるCEOなら、きっと最先端の業務管理ソフトウェアを使っているだろう、と思いがちです。

違います。

サム・アルトマンは紙とペンを使います。

彼は2018年のブログにこう書きました。「リストの使用を強く勧めます。1年、1か月、1日の単位で達成したいことの目録をつくります。」

それだけです。派手なアプリも、複雑なタグシステムありません。紙にやることを書き、重要なもののいくつかに星印（ ）をつける、それだけです。

初めて聞くと「シンプルすぎないか？」と思うかもしれません。しかしこのシンプルさの中に秘密が隠れています。

第一に、リストは頭を空にしてくれます。

人間の脳は記憶の倉庫ではありません。何かを記憶しようとする、膨大なエネルギーを消費します。「会議の準備をしなければ、メールに返信しなければ、報告書を書かなければ……」こういったことを頭の中に抱えていると、頭が重くなります。

アルトマンはそれをすべて紙に書き出します。すると頭が軽くなり、本当に重要な「考える仕事」に集中できます。彼は「リストがあるとマルチタスクに役立ちます。頭の中にたくさんのことを抱えておく必要がなくなるからです」と言います。

第二に、スター印のシステムが核心です。

やることが10個あるとします。全部重要そうに見えます。でも、本当に全部同じくらい重要なのでしょうか。違います。アルトマンは正直です。「あることは本当に重要で、あることはそれほどでもない。率直になろう。」

そこで彼はリストの中で本当に重要な3～4項目にだけスター印をつけます。そして一日の始まりと終わりに、このスター項目だけを確認します。残りは？時間があればやり、なければ次に回します。

この方法の妙は「選択を強制する」という点にあります。もし10項目すべてにスター印をつければ、実質的にスターがないのと同じです。アルトマンのシステムは私たちに問い続けます。「本当に重要なのは何？今日どうしてもやるべき一つを選ぶとしたら？」

第三に、紙の力です。

今はすべてがデジタルですよ。それなのにアルトマンは紙にこだわります。なぜでしょう？

彼は紙のリストをよく書き直します。昨日書いたリストを見て、今日は新しい紙に改めて書くのです。この過程で自然と優先順位が整理されます。「これは昨日は重要そうだったのに、今日見るとそうでもないな。」「これずっと後回しにしているけど、本当にやるべきかな？」

手で書くという行為自体が脳を刺激します。タイピングより記憶に残りやすく、思考を整理するのにも役立ちます。研究によれば、手書きは認知能力の向上に効果的だとされています。

第四に、勢いを生み出します。

アルトマンは「こなすほど気分がよくなり、するとさらに多くできるようになります」と言います。リストから項目を消すたびに小さな達成感を覚えます。それが積み重なると「自分にはできる」という自信が生まれ、次の仕事にも挑戦したくなります。

スター項目を消したときは？もう、最高の気分です。

彼は「一日の始まりと終わりに、本当に前進できる仕事をするのが好きです」と語っています。朝起きてすぐに、最も難しいスター項目を一つ片づけます。そうすれば、その日はすでに成功です。後は全部ボーナスです。

第五に、気分に応じて選べます。

リストがあると、もう一つ良いことがあります。アルトマンは「特定のタスクをやりたくなければ、いつでも別の面白いものを見つけられます」と言います。

月曜日の朝、報告書を書くのがどうしても嫌だとしましょう。無理にやる必要はありません。リストを見てください。他のスター項目の中に、今やりたいものはありますか？コードレビュー？戦略会議の準備？

何でも好きなものを選んでやってください。大切なのは「スター項目を前進させること」であって、「順番通りにやること」ではないのですから。

第六に、複雑なものは必要ありません。

アルトマンははっきり言います。「私は分類もしないし、タスクの規模を測ったりもしません。私がやることといえば、重要な項目にスターをつけるだけです。」

多くの生産性システムは「A・B・C優先順位」「緊急/重要マトリクス」「所要時間の見積もり」といったものを求めます。ところが、こうした分類作業そのものが時間を奪います。アルトマンは、その時間に実際の仕事をする方がよいと考えています。

スターだけで十分です。

学生もこの方法をすぐに使えます。

宿題が5つあります。ノート一枚に全部書き出します。その中で、明日の朝までに必ず提出しなければならないもの、試験に出る重要な内容が含まれているものにスターをつけます。おそらく2~3つがスターになるでしょう。

そうしたら今夜はそれだけやります。残りは？スター項目をすべて終わらせて時間が余ればやりま。疲れているなら寝てもかまいません。スター項目はもう終わったのですから。

このシステムが機能する理由はシンプルさにあります。複雑なツールは最初は良さそうに見えますが、すぐに面倒になります。設定して、分類して、管理するだけで時間が全部取られてしまうのです。

紙とスターは違います。

30秒でリストが作れます。10秒でスターをつけられます。そしてすぐに仕事を始められます。ツールが邪魔になりません。

アルトマンが10年以上この方法を守り続けている理由です。会社が大きくなり、責任が増え、会議が爆発的に増えても、彼の紙のリストは変わりませんでした。この方法が本当に機能するからです。

生産性ツールの目的は何でしょう？

「仕事をより多く、よりうまくこなすため」です。ツール自体が目的になってはいけません。アルトマンの紙とスターは、まさにこの本質を守っています。派手ではありませんが、強力です。

今すぐ紙を一枚取り出してください。やるべきことを5つだけ書いてください。その中で本当に重要なもの2つにスターをつけてください。

そして、それだけやってください。

明日の朝、きっと驚くでしょう。「あれ？自分、何かやり遂げたな。」

朝の時間を守る：最も生産的なディープワーク

サム・アルトマンの一日の中で、最も神聖な時間があります。

朝です。

ブログにこう書いています。「朝の最初の数時間は、間違いなく一日の中で最も生産性の高い時間です。だからその時間には、誰にも何も予定を入れさせません。」

好みではありません。戦略です。

アルトマンは通常、朝6時から7時の間に起きます。起きるとすぐ、大きなエスプレッソを一杯飲みます。朝食はほとんど取りません。15時間の空腹を維持する断続的断食をしているからです。

何より大切なこと。メールを確認しません。メッセージーも見ません。ニュースも見ません。

では、代わりに何をするのでしょうか。

ディープ・ワーク（Deep Work）をします。深い集中を要する、最も困難な仕事のことで。

ディープ・ワークとは何か、説明しましょう。数学の問題集を解いているときに、隣で友人がずっと話しかけてきたらどうでしょう。問題に集中できませんよね。3分ごとに集中が途切れると、難しい問題は絶対に解けません。

ディープ・ワークはその逆です。誰にも邪魔されず、通知もなく、ただ一つのことに完全に没頭している状態を指します。

アルトマンは朝の時間をこのように使います。

彼は難しい戦略的決断を行います。「GPTの次のバージョンはどの方向へ進むべきか」「このパートナーシップを受け入れるべきか」「会社の5年計画で何を修正すべきか」。

こうした思考は「合間に」できるものではありません。脳が100%稼働していなければならない仕事です。

科学的な根拠もあります。

脳は夜眠っている間に充電されます。朝目を覚ますと、意志力のバッテリーは満タンです。難しい決断を下し、創造的に考え、複雑な問題を解く能力が最高潮に達しています。

時間が経つにつれ、このバッテリーは消耗していきます。決断を一つ下すたびに、人に一人会うたびに、メールを一通読むたびに、少しずつ減っていきます。夕方になると？ほぼ底をついています。だから夕方にはNetflixだけを見たくなのです。

アルトマンはこのバッテリー管理の天才です。

最も貴重な朝のエネルギーを、最も重要な仕事にだけ使います。簡単な仕事、機械的な仕事、他人から頼まれた仕事はすべて午後に回します。

スケジュールを見ると明確です。

午前は一人でのディープ・ワークの時間。戦略立案、研究レビュー、重要な執筆。午後は会議、メール、電話、人と会うこと。

午前と午後では、仕事の種類がまったく異なります。

彼は「午後になるとすべてがぐちゃぐちゃになり、あらゆる問題が噴き出します」と率直に言います。CEOとして、緊急の問題が次々と発生します。突然誰かが電話してきて、緊急の案件が上がってきて、決裁すべき書類が積み上がるのです。

これは避けられません。CEOの宿命です。

アルトマンの戦略は「避けられないなら時間帯を決めよう」というものです。午後はどうせ混乱の時間です。ならば午後にそういった仕事をまとめてこなします。その代わり、午前だけは絶対に守ります。

彼は「最初の数時間は誰にもスケジュールを入れさせません」と断固として言います。親しい同僚が「10時にちょっとミーティングしよう」と言っても断ります。「すみませんが、午後にできますか？」

これは簡単そうに見えますが、実際は難しいことです。他の人たちには理解できないことがあるからです。「なぜ朝はダメなの？たった1時間じゃないか。」

アルトマンは揺るぎません。この原則が自分の生産性の核心だとわかっているからです。

アルトマンはメールの確認も後回しにします。

なぜでしょうか。朝にメールを見ると、他の人たちから頼まれた仕事が頭に入ってきます。「これを返信しなければ、あれを確認しなければ。」そうすると、自分がやろうとしていた重要な仕事に集中できなくなります。

彼はメール確認用に専用のLEDライトを使います。朝に10～15分、このライトを浴びながら簡単にメールをざっと目を通します。急ぎのものがあるかどうかだけを確認するのです。返信は午後にします。

このLEDライトの習慣も興味深いものです。研究によると、朝に明るい光を浴びると体内リズムが整い、一日中より目覚めた感覚が続くと言われています。アルトマンは「これが私が勧めるものの中で最も効果が大きいです」と言います。

90分のサイクルも重要です。

アルトマンは約90分間、集中して仕事をしたあと、短く休憩します。散歩をしたり、ストレッチをしたり、ぼんやりしたりします。

なぜ90分なのでしょう。科学者たちが発見した最適な集中時間です。90分以上休まずに集中すると、効率が落ちます。逆に休みが多すぎると、深く入り込むことができません。

休憩中はスマートフォンを見ません。代わりに外を歩くか、窓を眺めます。脳に本当の休息を与えるためです。スマートフォンは休息ではありません。別の種類の刺激にすぎないのです。

会議は午後へ。

アルトマンは会議をできる限り午後にまとめます。しかも15～20分と短く、あるいは2時間と長くします。1時間の会議は好きではないと言います。「時間を無駄にするケースが多い。」

15分あれば素早く決断して終わらせられます。2時間あれば複雑な問題を深く議論できます。では1時間は？中途半端です。準備して始めるうちに30分が過ぎ、まとめようとするともた時間が経ってしまいます。

これらすべての原則の核心は、一つだけです。

「最も貴重な時間を、最も重要なことに使おう。」

アルトマンの朝は、彼の成功の秘密です。ChatGPTを作り、AGIへと進む戦略はすべて、この朝の時間から生まれたはず。会議室から出てきたものではありません。静かな朝、誰にも邪魔されないあの数時間から出てきたのです。

私たちも明日から始められます。

朝の最初の1時間だけでも守ってみてください。スマートフォンは別の部屋に置いてください。最も難しいことを一つ、その時間に終わらせてください。

1ヶ月後、きっと感じるはずですよ。「あれ？何か変わったな」と。

断る技術：「申し訳ありませんが、今は難しいです。」

サム・アルトマンの生産性哲学において、最後にして最も難しい技術があります。

「ノー（No）」と言うことです。

彼は自分のブログに率直にこう書いています。「私は、仕事を断ることと、必須でない作業をできるだけ早く片付けることに、容赦なくあろうとしています。」

「容赦なく」という表現を使っています。強い言葉ですよね？

でも、これが必要なのです。なぜなら、世界は絶えず私たちの時間を奪おうとしているからです。

OpenAIのCEOになった瞬間から、アルトマンには毎日数十、数百の依頼が押し寄せました。

「このカンファレンスで講演していただけますか？」「30分だけコーヒーを飲みながらアドバイスをお願いします。」「私たちのプロジェクトを少し見ていただけますか？」「テレビのインタビューを一つだけ。」

どれも悪くないように見える依頼です。断ると悪い人のようで、失礼に見えてしまいます。

それでもアルトマンは断ります。

なぜなら、彼は自分の時間の価値を知っているからです。彼の時間は「AGIを作ること」に使われなければなりません。他のすべてのことは、それより重要度が低いのです。

彼は「会議とカンファレンスは時間コストが膨大です。だからできるだけ避けるようにしています」と言っています。

人々は誤解するかもしれません。「アルトマンは傲慢なんじゃないか。人に会いたくないんだろう」と。

そうではありません。彼は傲慢だから断るのではなく、賢明だから断るのです。

一度計算してみましょうか。

カンファレンスの講演を一つ引き受けると、実際に何時間かかるかという話です。発表準備5時間。移動3時間。発表と質疑応答2時間。ネットワーキング2時間。合計12時間です。

12時間あれば、アルトマンは何ができるでしょうか？

GPTの次のバージョンのロードマップを立てることができます。重要なパートナーシップ契約を検討することができます。研究チームと深い議論をすることができます。

どちらが世界により大きな影響を与えるでしょうか？

答えは明らかです。

アルトマンはこうも言っています。「ランダムなミーティングの90%は無駄です。でも10%には価値があります。」

興味深い見方です。彼はほとんどのミーティングが無駄だということを知っています。それでも完全に閉じてはいません。なぜなら、その10%があるからです。

ときには偶然の出会いから大きなアイデアが生まれることがあるからです。誰かとコーヒーを飲んでいるときに、「え？これって完全に革新的じゃないか？」という瞬間が訪れることもあります。

だからアルトマンは完全に遮断せず、戦略的に選びます。自分のスケジュールに「偶然の出会いのための余白」をわざと残しておきます。しかしその余白でさえ、意図的に管理するのです。

断る具体的な方法もあります。

一つ目は、明確な基準を設けることです。「この仕事は私たちのコアミッションに直接貢献するか？」答えが「ノー」なら断ります。

二つ目は、代替案を提示することです。「私が直接は難しいですが、チームの誰々が助けられると思います。」こう言えば、関係を保ちながら時間を守ることができます。

三つ目は、先延ばしにしないことです。断るなら早めに断ります。「考えておきます」と言いながら2週間引きずった末に断るのはもっと悪いです。相手も別の計画を立てられるよう、早めに知らせるのが礼儀です。

時間の価値を知ることも大切です。

アルトマンはこんな例を挙げています。「時給100ドルを稼ぐ人が、20ドルを節約しようと2時間を費やすのは理解できません。」

どういう意味でしょうか？

たとえば、洗車を自分でやれば2万ウォンを節約できます。でも2時間かかります。もしその2時間でもっと大切な仕事をしていたら、20万ウォン分の価値を生み出せていたかもしれません。

アルトマンはまた、「生産性そのものを追い求めないでください」とも警告しています。多くの人が「もっと多く、もっと速く」にばかりこだわります。タスクリストを消化することが目的になってしまうのです。

本当の目標は何でしょうか？

「大切なことを成し遂げること」です。

一日に100件こなしても、どれも大して重要でなければ意味がありません。それより、本当に大切なことを3つだけやるほうがずっといいのです。

だから断ることが必要です。大して重要でない99のことを断ってこそ、本当に大切な3つのことに集中できるのです。

実際の事例を見てみましょう。

2023年、アルトマンが一時CEOを解任されてから復帰した出来事がありました。あのとき、彼は凄まじいプレッシャーにさらされました。メディアからのインタビュー依頼が殺到し、投資家たちは面会を求め、社員たちは説明を求めています。

アルトマンはどうしたのでしょうか？

彼は本当に重要なことだけに絞りました。社員との対話。取締役会との交渉。会社の安定化。それ以外はすべて断るか、後回しにしました。

その結果、4～5日で状況が落ち着き、彼は復帰しました。もし彼がすべての要請に応じていたら？おそらく今でも収拾がつかなかったでしょう。

学生も練習できます。

ゲームに誘う友達に「ごめん、今日は試験勉強しなきゃ。週末にやろう。」つまらない集まりに「次は行くよ。今日はちょっと疲れてて。」SNSの通知はマナーモードに設定。

最初はぎこちなく感じます。断ると申し訳なく、関係が悪くなるのではないかと心配になります。

でも、不思議なことが起きます。

時間が経つと、周りの人たちは理解します。「あ、この人は自分のことに集中しているんだな。」むしろ尊重してくれるようになります。

友達も気づきます。「あいつには気軽に頼めないな。でも、何かやるときはきちんとやるんだよな。」

これがアルトマンの評判です。

彼は簡単にYesとは言いません。でも、いったんYesと言ったら最後までやり遂げます。周りにはそれを知っています。だからこそ、より信頼されるのです。

結局、断る技術とは「上手にNoと言うこと」ではありません。

「Yesを大切に使うこと」なのです。

アルトマンは自分の人生において、本当に大切なことにだけYesと言います。AGIをつくること。チームを育てること。世界を変えること。

それ以外は、丁寧でも毅然とNoと言います。

彼の成功は、何をしたかよりも、何をしなかったかから生まれたのかもしれませんが。

今日から練習してみてください。

一日に一度だけでも、重要でないことに「申し訳ないですが、今は難しいです」と言ってみてください。

その時間に、本当に大切なことを一つ終わらせてみてください。

一ヶ月後、あなたの人生は変わっているはずです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

5.1 モデルはどこへ進化するのか

第5部 AI技術の未来とロードマップ

← 前の章

次の章 →

アルゴリズム革新：10倍から100倍の改善可能性

サム・アルトマンがOpenAIの会議室で最もよく持ち出す話があります。

「より大きなコンピューター、より多くのGPU。もちろん重要です。しかし本当にゲームを変えるのはアルゴリズムです。」

彼は2024年末からこのメッセージをより強く発信し始めました。AI業界全体が一つの壁にぶつかっていたからです。それが「スケーリングの限界」でした。

GPT-3からGPT-4に移行するまでは単純でした。より多くのデータを集め、より強力なGPUを数万個つなげば、性能はぐんぐん上がりました。まるで車のエンジンに燃料を足せばスピードが上がるように。

しかし2024年に入ると状況が変わりました。複数のAI研究所が同じ問題を報告し始めたのです。モデルの規模を2倍にしても、性能は以前ほど向上しなくなっていました。いわゆる「スケーリング則の鈍化」です。

アルトマンはこの問題を誰よりも早く見越していました。そこで彼は別の方向を探し始めました。それが「アルゴリズム的突破口」でした。

「10倍から100倍まで性能を引き上げられる可能性が、アルゴリズム革新にあります。」

この言葉は誇張に聞こえました。しかしAIの歴史を振り返れば、すでに何度も証明されてきた事実でした。

2017年に登場した「トランスフォーマー（Transformer）」アーキテクチャがその好例です。この新しいアルゴリズムが登場したことで、同じコンピューティングリソースで数十倍優れた性能が出せるようになりました。ChatGPTの基盤となった技術も、まさにこのトランスフォーマーでした。

もう一つの例が「強化学習（RLHF、Reinforcement Learning from Human Feedback）」です。この手法一つが、GPTモデルを単なるテキスト生成器から、人間と会話できる賢いアシスタントへと変えました。

アルトマンは社内文書にこう記しています。

「私たちはすでにコンピューティングの物理的な限界に近づいています。今のように規模だけを拡大するやり方は、いつか壁にぶつかります。その壁を突き破るのは、まったく新しい種類のアルゴリズムです。」

2024年12月、OpenAIは「o1」という新しい推論モデルを発表しました。このモデルは既存のGPTとはまったく異なる仕組みで動作します。答えをすぐに出す代わりに、問題を複数のステップに分けてゆっくりと「考える」のです。

結果は目を見張るものでした。

国際数学オリンピックの問題では、o1は人間の金メダリスト並みの成績を収めました。従来のGPT-4が解けなかった問題を次々と解いてみせたのです。しかも使用したコンピューティングリソースは、GPT-4のトレーニング時よりはるかに少量でした。

これこそがアルゴリズム的突破口の力でした。

サム・アルトマンは2025年1月のインタビューでこう語りました。

「超知能（superintelligence）が到来すれば、科学的発見の速度が10倍速くなるでしょう。10年かかっていた技術の進歩が1年で実現するわけです。そして翌年にはまた同じだけの進歩が起きます。これが積み重なれば、世界は完全に変わります。」

もちろん彼はすぐに人々の期待値を調整してもいます。「Twitterの誇大宣伝がひどすぎます。期待を100分の1に下げてください。私たちはまだAGIを作っていません。」

しかし核心のメッセージは明確でした。AIの発展の次の段階は、ハードウェア競争ではなくアルゴリズム革新からもたらされるということです。

実際、OpenAIの内部では新しいアルゴリズム研究が盛んに進められています。

研究者たちは「Chain-of-Thought（思考の連鎖）」や「Tree-of-Thought（思考の木）」といった新しい推論パターンを開発しています。これらの手法はAIが単に次の単語を予測するのではなく、実際に「推論」できるようにするものです。

さらに「Mixture of Experts（専門家の混合）」のような効率的なアーキテクチャも研究中です。この方式では、巨大なモデルの中に複数の小さな専門家モデルを作っておき、質問が来た際に最も適した専門家だけを起動して答えさせます。まるで会社で全社員を会議に呼ぶのではなく、必要な部署の人間だけを集めるようなものです。

こうした手法のおかげで、同じ性能を発揮しながら必要なコンピューティングリソースを10分の1に削減できます。

アルトマンの首席科学者ヤクブ・パチョッキ（Jakub Pachocki）はこう説明しています。

「現在のモデルはおよそ5時間分の作業を処理できます。しかしこの時間の範囲は急速に拡大するでしょう。主要な科学的突破口のためであれば、データセンター全体のコンピューティングパワーをたった一つの問題に注ぎ込む価値があるはずです。」

核心は「効率性」です。

人間の脳は約20ワットの電力で動きます。電球一つを灯すほどのエネルギーで、私たちは考え、創造し、問題を解決します。一方GPT-4は、答え一つを生み出すのに膨大な電力を消費します。

この差を縮めること。それがアルゴリズム革新の目標です。

サム・アルトマンは未来をこう描きました。

「いつか私たちは、今より100倍賢いモデルを今より100倍低いコストで動かせるようになるでしょう。その瞬間こそが、真のAI革命の始まりです。」

この突破口は技術的な成果にとどまないと彼は信じています。がん治療法の発見、気候変動の解決、新たなエネルギー源の開発。人類が直面する最も困難な問題が、アルゴリズム革新を通じて解けるというわけです。

2025年10月、アルトマンはより具体的なタイムラインを示しました。

「2026年9月までにインターンレベルのAI研究アシスタントを作ります。そして2028年までには完全に自動化された『AI研究者』を完成させます。」

このAI研究者は、人間の指示がなくても自ら研究プロジェクトを進めることができます。論文を読み、仮説を立て、実験を設計し、結果を分析します。そして何より重要なのは、このAIがより優れたAIを作る方法を研究できるという点です。

これこそが、アルゴリズム革新の自己加速効果です。

優れたアルゴリズムが生まれると、そのアルゴリズムがさらに優れたアルゴリズムを見つけ出します。すると進化の速度は指数関数的に加速します。

もちろん、これらすべてが順調に進むという保証はありません。

アルゴリズム革新には、いまだ多くの難題が残っています。長期的な推論、一貫した計画立案、複雑な因果関係の把握。これらの能力は、AIがまだ完全にはこなせていません。

しかしサム・アルトマンは楽観的です。

「モデルの限界の90%がアルゴリズムの問題だということは、今や分かっています。その壁を越えた瞬間、AIはまったく別の種類の存在になります。」

その瞬間は思っているより早く来ると彼は信じています。「数千日以内に」というのが彼のよく使う表現です。およそ2030年前後を指しています。

結局のところ、アルゴリズムの突破口は単なる技術改善ではありません。それは人間が作った知性の未来を設計し直す営みです。そしてサム・アルトマンは、まさにその未来の扉の前で、次の段へと静かに、しかし確かに手を伸ばしています。

アルゴリズム、データ、コンピューティングの三位一体

サム・アルトマンがよく使う比喻があります。

「AIモデルは料理と同じです。優れたレシピ（アルゴリズム）、新鮮な食材（データ）、そして立派なキッチン（コンピューティング）がすべて揃う必要があります。一つでも欠ければ、まともな料理にはなりません。」

この三つの要素は、AI発展の三本柱です。OpenAIのあらゆる戦略は、この三本柱をいかに強化し、バランスよく発展させるかに焦点を当てています。

第一の柱 アルゴリズム

アルゴリズムはAIの「脳の構造」です。

サム・アルトマンはアルゴリズムを「知性の設計図」と呼びます。同じデータと同じコンピューターを使っても、アルゴリズムが違えば結果は天と地ほど異なります。

GPTシリーズを作る過程で、彼はこの事実を骨身に染みて感じました。

GPT-3が登場したとき、人々は驚きました。しかしOpenAI内部では「まだまだだ」という評価が多くありました。モデルがしばしば的外れな答えを出したり、文脈を見失ったりしたからです。

そのときに登場したのがRLHF（人間のフィードバックを用いた強化学習）でした。このアルゴリズム一つが、GPT-3をまったく別の存在へと変えました。同じモデル、同じデータでも、学習方法を変えるだけで、突然人間のように会話できるようになったのです。

アルトマンはこのとき確信しました。「未来を作るのは、より大きなモデルではなく、より賢いアルゴリズムだ。」

OpenAIの研究方向はその後、明確にアルゴリズム革新に重きを置き始めました。新たな推論方式、より効率的なattentionメカニズム、長期記憶を持つ構造、人間のように考えるツリーベースの推論。これらはすべて、アルゴリズム研究の産物でした。

そして2024年末、o1モデルが登場しました。このモデルは「考える時間」を増やす方式で問題を解きました。既存のモデルがすぐに答えを出そうとしたのに対し、o1は段階を踏んでゆっくりと推論しました。

結果として、同じ量の学習データでも格段に優れた性能を発揮しました。これこそがアルゴリズム革新の力でした。

第二の柱 データ

データはAIの「経験」です。

サム・アルトマンはデータを「モデルが世界を学ぶ教科書」と表現します。どれほど優れたアルゴリズムがあっても、質の低いデータで学習すれば、役に立たないAIができあがります。

「ゴミを入れればゴミが出る（Garbage in, garbage out）。」コンピュータ科学の古い格言が、AI時代にいっそう切実な真実となりました。

GPT-3を作る際、OpenAIはインターネット上の膨大なテキストを収集しました。ウェブページ、ブログ、フォーラム、書籍、ニュース記事。あらゆる種類の文章を投入しました。

問題がありました。インターネットには良い情報だけがあるわけではありませんでした。フェイクニュース、偏見に満ちた文章、罵倒や差別的表現。こうしたものと一緒に学習されてしまいました。

そこでGPT-4からは、データの「質」に力を入れ始めました。ただ大量のデータを集めるのではなく、検証済みの高品質なデータを選別して使いました。

科学論文、専門書籍、信頼できるニュースソース、検証された教育資料。こうしたデータが、モデルの性能を大きく引き上げました。

アルトマンはチームにこう語りました。「AIは世界に似ていなければなりません。ただし、世界の最悪ではなく、最善と似ていなければならない。」

もう一つ重要な変化は、データの「多様性」でした。

テキストだけでは不十分でした。GPT-4からは、コード、数式、科学データ、医療プロトコル、法律文書など、多様な形式のデータを学習しました。

するとモデルの思考方式が変わりました。以前は「もっともらしい文章」を書くだけの水準だったのが、今では実際に問題を解決できる水準になりました。

2024年後半からは「合成データ（Synthetic Data）」も重要になってきました。これはAIが自ら生成したデータです。o1モデルが数学の問題を解く過程で生まれた推論の軌跡が、次のモデルを訓練するためのデータになる、という仕組みです。

こうなると、インターネットから収集できるデータの限界を超えることができます。AIが自ら質の高いデータを生み出す好循環が始まるのです。

第三の柱 コンピューティング

コンピューティングはAIの「筋肉」です。

どれほど優れたアルゴリズムとデータがあっても、それを動かすコンピューターがなければ意味がありません。

サム・アルトマンはコンピューティングについて非常に現実的です。AI発展には膨大なコンピューティングパワーが必要だと理解しています。しかし同時に、それが無限ではないことも知っています。

GPT-4の訓練にかかったコストは1億ドル以上と推定されています。その大半はGPU費用と電力費用でした。これは並のスタートアップには到底負担できない規模です。

そこでアルトマンは二つの方向で動きました。

一つは、より多くのコンピューティングインフラを確保することでした。マイクロソフトと組み、巨大なデータセンターを建設しました。2025年には「スターゲートプロジェクト」を発表し、5000億ドル規模のAIインフラ投資計画を打ち出しました。

もう一つは、コンピューティングをより効率的に使うことでした。ここで再びアルゴリズムの革新が登場します。より賢いアルゴリズムを使えば、同じ作業を10分の1のコンピューティングで実現できます。

アルトマンはよくこう言います。「知能のコストは、最終的にエネルギーのコストに収束するでしょう。」

どういう意味でしょうか。

今はAIモデルを作るのに、アルゴリズム開発、データ収集、エンジニアの人件費など様々なコストがかかります。しかし将来、アルゴリズムが完成しデータも十分になれば、最終的に残るコストは電気代だけになる、ということです。

だから彼はエネルギー問題に多大な関心を注いでいます。原子力、核融合、再生可能エネルギー。安価でクリーンなエネルギーを確保することがAI発展の核心だと信じています。

三つの柱の調和

重要なのは、この三つの要素が互いに影響し合っているという点です。

優れたアルゴリズムが生まれれば、必要なデータの量が減り、コンピューティングリソースも少なくて済みます。逆に、より多くのコンピューティングパワーがあれば、より複雑なアルゴリズムを試すことができ、より多くのデータを処理できます。

サム・アルトマンは、この三つの要素のバランスを保つことが、CEOとしての自分の最も重要な役割だと考えています。

「アルゴリズムにばかり集中すると、実用的な製品が生まれません。コンピューティングにばかり集中すると、効率が落ちます。データにばかり集中すると、法的問題と倫理的問題にぶつかります。」

「この三つを同時に進歩させながら、バランスを崩さないこと。それこそがOpenAIのやっていることです。」

そして彼は未来についてこう展望します。

「これから10年以内に、アルゴリズムは今より100倍効率的に進化するでしょう。データは合成データのおかげで無限に近くなるでしょう。コンピューティングは新しいチップ技術と安価なエネルギーのおかげで、はるかに強力になるでしょう。」

「この三つの車輪が一緒に回り始めたとき。そのとき、本当の魔法が起きるでしょう。」

1兆トークンコンテキストとパーソナライズされた生涯記録AI

サム・アルトマンが思い描くAIの究極の姿とは何でしょうか。

それは世界で最も賢いAIということではありません。むしろその逆に近い。「あなただけのために、あなたを最もよく知るAI」です。

彼はこのビジョンを一文で要約します。

「1兆トークンのコンテキストを持つ、非常に小さな推論モデル。」

この一文には、未来のAIの核心がすべて込められています。

1兆トークンとは何か。

まず「トークン」とは何かを知る必要があります。

トークンはAIが言語を理解するための基本単位です。おおよそ単語一つ、あるいはそれより小さな断片だと思えばいいでしょう。たとえば「こんにちは」は約2～3個のトークンで構成されています。

では、1兆トークンとはどれほどの量でしょうか。

ハリー・ポッターシリーズ全巻が約100万トークンです。1兆トークンなら、ハリー・ポッターを100万回読める量に相当します。別の言い方をすれば、一人の人間が生涯にわたって読み、書き、話すすべての内容を収めてもまだ余りある量です。

「コンテキスト (Context)」とはAIの「記憶力」を意味します。今私たちが使っているChatGPTは、会話を続けていると前に話した内容を忘れてしまいます。コンテキストに制限があるからです。

2025年現在、最大のコンテキストを持つモデルでも約100万トークン程度です。それでも膨大な量ですが、1兆トークンと比べると1000分の1にすぎません。

生涯の記録を覚えているAI

1兆トークンのコンテキストを持つAIは、何ができるのでしょうか。

あなたが生まれた瞬間から今日まで、すべての記録を覚えていられます。

幼い頃につけた日記、学校で書いた宿題、友人との会話、就職面接で話したこと、職場で送ったメール、家族との思い出、健康記録、好きな音楽や映画、趣味活動、悩みや夢。そのすべてをAIが記憶するのです。

サム・アルトマンはこれを「個人の第二の脳」と呼んでいます。

こうしたAIがあれば、どんなことが可能になるのでしょうか。

朝起きると、AIが語りかけます。「今日は大事なプレゼンがある日ですよ。5年前に似たプレゼンをしたとき、緊張して失敗した箇所を覚えていますか？今回はこう準備すればうまくいきそうです。」

健康診断に行くと、AIがこう言います。「あなたの10年分の健康記録を見ると、毎年この時期に血圧が少し上がる傾向があります。昨年試した運動方法が効果的でしたから、また試してみてもはどうでしょう？」

新しいプロジェクトを始めようとする、AIが助言します。「3年前に似たアイデアをメモしていましたよ。当時は実行できませんでしたが、今は状況が違っています。あの時のメモと今の考えを合わせれば、よい計画が生まれそうです。」

これこそが「完全にパーソナライズされたAI」の姿です。

なぜ「小さな」モデルでなければならないのか

ここに重要なポイントがあります。

サム・アルトマンは、このAIが「非常に小さなモデル」でなければならないと言います。どういう意味でしょうか。

今のGPT-4のようなモデルは途方もなく大きいです。巨大なデータセンターでしか動きません。しかしアルトマンが描く個人化AIは、あなたのスマートフォンの中で動かなければなりません。

なぜでしょうか。

一つ目は、プライバシーのためです。

あなたの生涯の記録をインターネットで送り、遠くのサーバーで処理するのは危険です。ハッキングされる可能性もありますし、企業があなたの情報を悪用するかもしれません。

しかし、すべての処理があなたのデバイスの中だけで行われたとしたら？誰もあなたの情報にアクセスできません。完全なプライバシーが保証されます。

二つ目は、速度のためです。

インターネットでデータを送受信すると時間がかかります。しかしあなたのスマートフォンの中で直接処理されれば、すぐに答えが得られます。まるで自分の頭の中で考えるように速いのです。

三つ目は、コストのためです。

巨大なサーバーを維持するには莫大なコストがかかります。しかし小さなモデルがそれぞれのデバイスで動くなら、はるかに安く済みます。

「でも、小さなモデルが賢くなれるのですか？」

良い質問です。サム・アルトマンの答えはこうです。

「個人化モデルは世界全体を知る必要はありません。ただ一人を深く理解すればいい。だから小さくても強力でいられます。」

考えてみてください。あなたの好み、習慣、目標、強み、弱点を完璧に知っている小さなAIと、世界のすべてを少しずつ知っている巨大なAI。あなたの人生により役立つのは、どちらでしょうか。

どうすれば実現できるのか

もちろん、これらすべてが簡単なわけではありません。

2025年現在の技術では、まだ完全には実現できていません。しかし、パズルのピースは一つずつはまりつつあります。

アルゴリズムの革新により、モデルを小さくしながら性能を維持できるようになりました。「モデル圧縮」技術が進歩し、巨大なモデルを10分の1のサイズに縮小しても、性能がほとんど落ちません。

長期メモリ技術も進歩しています。すべての情報を常に保持する代わりに、重要なものだけを要点として保存し、必要なときに素早く取り出す方式が開発されています。

オンデバイスAIチップも急速に進歩しています。AppleのNeural Engine、QualcommのAIプロセッサ。これらのチップがますます強力になることで、スマートフォンでも複雑なAIモデルを動かせるようになりました。

未来の姿

サム・アルトマンは10年後をこのように思い描いています。

「2035年ごろには、ほとんどの人が自分だけのパーソナルAIを持つようになるでしょう。そのAIは生まれたときから一緒に育ち、あなたのあらゆる瞬間を記憶し、生涯を通じてあなたを助け続けます。」

「それは単なる秘書ではありません。あなたの外付け脳であり、最も親しい友人であり、最も賢明な助言者になるでしょう。」

「そして何より大切なのは、そのAIが完全にあなたのものだということです。誰もアクセスできず、あなただけが制御できます。」

これこそが、サム・アルトマンが夢見るAIの未来です。

巨大なスーパーAIが世界を支配する未来ではなく、すべての人が自分だけの完璧なAIパートナーを持つ未来。彼はそれこそが、より美しく、より人間的で、より安全な未来だと信じています。

「AIは人間に取って代わるものではありません。一人ひとりの人間をより強力にするものです。」

これが、1兆トークンのコンテキストを持つ小さな推論モデルの真の意味です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書斎 #サムアルトマン #OpenAI #ChatGPT #人工知能

5.2 エージェントとコーディングの役割

第5部 AI技術の未来とロードマップ

← 前の章

次の章 →

AIエージェントとは何か？

サム・アルトマンが2023年のある日、OpenAIのチーム会議で語った言葉があります。

「私たちが作りたいのは、会話するAIではありません。仕事をやり遂げるAIです。」

その一言がOpenAIの次の目標を根本から変えました。OpenAIは、ただ質問に答えるチャットボットを超えて、自ら考え、計画し、実行する「エージェント」を作り始めたのです。

「エージェント」という言葉を初めて聞くと、難しく感じるかもしれません。映画に登場するFBIエージェントや秘密工作員を思い浮かべる人もいるでしょう。

しかしAIの世界では、エージェントの意味は少し異なります。

AIエージェントとは、人間からの指示を最小限に受けながら、複雑な仕事を自力でこなす智能型システムです。有能な秘書に「来週の出張の準備をしておいて」と伝えたと、その秘書が航空券を予約し、ホテルを探し、スケジュール表を作り、必要な書類を揃えてくれるようなものです。

アルトマンはエージェントについて、こう説明しています。

「エージェントは目標を与えるだけで、その目標を達成するために必要なすべての手順を自ら考えます。途中で問題が起きれば、自分で解決策を見つけます。人間はスタートボタンを押して、最後に結果を確認するだけでいい。」

どれほど大きな変化なのか、例を挙げてみましょう。

以前のAIはこうでした。「メールの下書きを書いて」と言えば書いてくれる。「このデータを分析して」と言えば分析してくれる。一つひとつ命令しなければなりませんでした。

エージェントはまったく違います。

「うちの会社の第2四半期の売上報告書を作って」と言うだけで、エージェントは自分で以下のことを進めていきます。

まず、会社の売上データがどこにあるかを探します。

データベースにアクセスして必要な数値を取得します。前年の同時期データと比較します。業界平均のデータをインターネットで検索します。グラフを作成します。重要なインサイトをまとめます。見やすいスライド形式に仕上げます。下書きを提示してフィードバックをもらいます。修正します。

最終版を完成させます。

これらのプロセスすべてを、人間が一つひとつ指示する必要はありません。

サム・アルトマンは2024年のあるインタビューで、エージェントの本質をこのようにまとめました。

「エージェントには三つの能力が必要です。一つ目は、長期的な計画を立てる力。複雑な目標を小さなステップに分けられなければなりません。二つ目は、ツールを使う力。ウェブ検索、コードの実行、ファイルの読み書きなどを自分でできなければなりません。三つ目は、自己修正の力。ミスをしたとき、自ら気づいて修正できなければなりません。」

OpenAIで働くあるエンジニアは、こんなことを話していました。

「サムは会議のたびに『最小限の監督』という言葉 강조했다。AIが1分おきに人間に確認を取るようでは、本物のエージェントじゃないと。本物のエージェントは、朝に仕事を頼めば夕方には完成した結果を持ってくるものだ、と言っていました。」

実際、OpenAIはGPT-4を開発する際にエージェント機能を念頭に置いていました。テキストを生成する能力だけでなく、コードを実行し、外部システムと連携し、複数のステップからなる作業を順序通りに処理する能力も盛り込みました。

アルトマンはエージェントを「デジタルな同僚」と呼ぶこともありました。

「エージェントはツールではありません。チームメンバーです。隣で働く優秀な同僚だと思ってください。もちろん、給料は払わなくていいし、休暇も要らないし、24時間働けますけど。」

冗談めかした言葉の中に、真剣なビジョンが込められています。アルトマンはAIが人間を単に置き換えるのではなく、人間の能力を広げてくれる存在になることを望んでいました。

エージェントのもう一つの重要な特徴は、「文脈を理解する力」です。

人と長く一緒に仕事をしていると、互いに暗黙の了解が生まれます。「あの、前のやつ」と言うだけで通じるようになるものです。エージェントも同じです。過去の会話、過去の作業、そして個人の好みを記憶し、学習します。

2024年末にOpenAIが公開したデモでは、エージェントがユーザーのメールパターンを学習し、重要なメールだけを選び分けて、各メールに適した返信の下書きまで作成する様子が披露されました。ユーザーは「送信」か「修正」ボタンを押すだけでよかったのです。

まだ完璧ではありません。

アルトマン自身も率直に認めています。「今のエージェントはときどき間違えます。見当外れの情報を持ってきたり、計画を誤ったりすることもあります。だからまだ、重要な判断を下す前には人間の確認が必要です。」

サム・アルトマンが描く未来のエージェントの姿は、こうです。

出社するとエージェントが今日やるべき仕事を優先順位順に整理しています。会議があれば関連資料を事前に要約済みです。メールは重要度に応じて分類され、緊急のものはアラートで知らせてくれます。報告書が必要なら、すでに下書きが用意されています。あなたは重要な意思決定と創造的な仕事にだけ集中すればいいのです。

SF小説のように聞こえますか？

アルトマンはこう言います。「3年前、ChatGPTのようなものが生まれると想像した人が何人いたでしょうか。技術は私たちの予想よりも速く進歩します。本当に役に立つエージェントは、思っているよりずっと早く来るはずですよ。」

レストランの予約を超えて：大規模並列作業の革命

AIエージェントを説明するとき、人々はよくレストランの予約の話を持ち出します。

「AIエージェントが代わりにレストランの予約をしてくれるんですよ！」

サム・アルトマンはこの言葉を聞いた時に笑っていたといいます。

「レストランの予約ですか？それは15年前の技術でも実現できました。エージェントの本当の可能性はずっと大きいのです。」

アルトマンが見たエージェントの真の力は「スケール」にありました。

人間は一度に一つのことしかきちんとこなせません。マルチタスクをしているように見えても、実際には素早く注意を移しているだけです。

AIエージェントは違います。同時に数十、数百、数千の作業を並行して進められます。

これを「並列処理」といいます。

具体的な例を挙げてみましょう。

あるスタートアップの創業者が新しいモバイルアプリを作ろうとしています。かつてはこうするしかありませんでした。

市場調査 → 競合分析 → 技術スタック選定 → デザイン → 開発 → テスト → マーケティング準備。

各ステップを順番に一つずつ進めなければなりませんでした。一人でやれば数か月かかり、チームを組んでも数週間はかかります。

エージェントの時代では、どうなるでしょうか。

創業者がエージェントにこう告げます。「フィットネスアプリを作りたい。ターゲットユーザーは20～30代の会社員だ。6か月以内にリリースしたい。」

すると、複数の作業が同時に動き出します。

エージェントAはフィットネスアプリ市場を調査します。どんなアプリがあるか、ユーザーが何に不満を持っているか、価格設定はどうなっているかを分析します。

エージェントBは技術面を検討します。どのフレームワークを使うか、サーバーをどこに置くか、コストはどれくらいかかるかを計算します。

エージェントCはUI/UXデザインの草案を複数作成します。類似アプリのデザイントレンドを分析して参考にします。

エージェントDは必要な開発者とデザイナーのプロフィールを探します。予算に合ったフリーランサーやエージェンシーを推薦します。

エージェントEはApp Store登録の手続き、法的要件、個人情報保護規制といった事項を整理します。

これらすべてが同時に進行します。

1～2日後には総合レポートが出来上がります。市場分析、技術計画、デザインの選択肢、チーム編成案、予算計画、リリース戦略まですべて揃っています。

創業者は核心的な判断だけ下せばいいのです。「デザインはB案にしよう」「技術スタックはこれにしよう」「予算はここを少し削ろう」。

サム・アルトマンはこうした未来を確信していました。

「一人がエージェントチームを率いれば、かつて10人が担っていた仕事ができます。100人分の仕事だって可能です。これが生産性革命です。」

企業の現場ではどうでしょうか。

大手製造業者を思い浮かべてください。数十の工場、数千の製品ライン、数万人の従業員を抱えています。

従来のやり方では、各部門がばらばらに動きます。生産部門は生産計画を立て、物流部門は配送を管理し、マーケティング部門は広告を作り、財務部門は予算を管理します。部門間のやり取りは遅く、情報が途切れやすいのです。

エージェントシステムはこれらすべてをつなぎます。

センサーエージェントが工場の機械の状態をリアルタイムで監視します。何らかの機械に問題の兆候が現れると、即座に報告します。

予測エージェントが販売データと市場トレンドを分析し、翌月の需要を予測します。

生産エージェントが予測結果を受け取り、生産計画を自動的に調整します。どの製品をどれだけ作るかを決定します。

物流エージェントが生産計画に合わせて原材料の発注を自動で行います。最も良い価格で最速で受け取れるサプライヤーを探します。

品質管理エージェントが生産工程を監視し、不良品が出ると原因を分析します。

マーケティングエージェントが在庫状況を見て、適切なプロモーションを企画します。

財務エージェントがすべての取引を記録し、資金の流れを最適化します。

これらすべてのエージェントは互いに情報を交わしながら協力します。人間の管理者たちは全体の流れを把握しながら、重要な意思決定だけを下せばよいのです。

アルトマンはこれを「AIオーケストラ」と呼びました。

「何百人もの演奏者がそれぞれ異なる楽器を奏でながら、指揮者の指揮のもとで一つの美しい音楽を生み出すように、無数のエージェントが協力し合うことで、複雑なビジネスを完璧に運営できるのです。」

科学研究の分野ではどうでしょうか。

新薬開発を例に挙げてみましょう。従来、新薬一つを開発するには10年以上の歳月と数兆ウォンの費用がかかります。

なぜこれほど時間がかかるのでしょうか。

考えられる化合物の組み合わせが膨大だからです。数百万、数億通りもの可能性を一つひとつ検証しなければなりません。人間の研究者だけでは、とても全部試しきれものではありません。

エージェントシステムが導入されると、どうなるでしょうか。

数千のAIエージェントが同時に異なる化合物の組み合わせをシミュレーションします。各エージェントは担当する組み合わせの効果、副作用、生産の可能性を分析します。

結果の優れた候補を絞り込みます。その候補をさらに精密にシミュレーションします。実験室で実際に検証する価値のある数件だけが残ります。

このプロセスが数年ではなく、数ヶ月、数週間で完了できるのです。

アルトマンはこう語りました。

「AIエージェントは、人類が解けなかった問題を解けるようにしてくれるでしょう。新しい薬、新しい素材、新しいエネルギー源。人が一生かけてもできない研究を、AIは数日で成し遂げられるのです。」

個人の生活においても、並列処理は強力です。

海外への引っ越しを想像してみてください。やるべきことが山積みです。

家探し、ビザ申請、航空券予約、荷物の整理、学校の調査、銀行口座の開設、現地の法規確認、保険加入、賃貸契約の解約、公共料金の精算、別れの挨拶。

一人でやれば数ヶ月かかります。ストレスも相当なものです。

エージェントに任せたらどうでしょうか。

「3ヶ月後にカナダのトロントに引っ越す。家族は4人だ。予算はこれくらいだ。」

すると複数のエージェントが同時に動き出します。不動産エージェントは条件に合う家を探します。法律エージェントはビザの書類を準備します。教育エージェントは子どもたちの学校を調べます。金融エージェントは銀行と保険を整理します。物流エージェントは引越し業者を比較します。

数日後、完璧な引越し計画が手元に届きます。タイムライン、チェックリスト、必要書類、連絡先、見込み費用まですべて整理されています。あとは計画に沿って進めるだけです。

これが並列処理の力です。

アルトマンのビジョンは明確でした。

「AIエージェントはレストランを予約するだけのものではありません。エージェントはあなたの人生全体を調整する個人秘書であり、会社を運営する経営パートナーであり、人類の未来を切り開く研究仲間です。これこそが本当のAI革命です。」

2025年、コーディング分野がまず制覇される。

2024年の夏、サム・アルトマンはある技術カンファレンスでこう語りました。

「2025年はAIエージェントの年になるでしょう。とりわけコーディング分野において。」

なぜよりによってコーディングなのでしょう。

アルトマンの説明はシンプルでした。

「コーディングは、AIが世界を実際に動かせる最初的手段です。AIがどれほど賢くても、手足がなければ何もできません。コードこそがAIの手足なのです。」

考えてみてください。

私たちが生きる世界は、ますますソフトウェアで動くようになっていきます。銀行取引、ショッピング、交通、医療、教育、エンターテインメント。ほぼすべてがソフトウェアとつながっています。

ソフトウェアはコードで作られています。

AIがコードを扱えるということは、AIがこれらすべてのシステムを理解し、修正し、改善できるということを意味します。

OpenAIはすでにGPT-4を開発した頃からコーディング能力に注力してきました。

これはOpenAIエンジニアの証言です。

「私たちはGPT-4が単にコードを生成するだけでなく、コードの意図を理解し、バグを見つけ、最適化し、さらにはプロジェクト全体の構造を設計できるようになることを目指していました。サムはこれがエージェントへの重要な一歩だと考えていました。」

実際、2024年末時点でGitHub CopilotのようなAIコーディングツールを使う開発者は、生産性が30～50%向上したと報告しています。

ただコードを速く書けるようになっただけではありません。AIが繰り返しの単調な作業を引き受けてくれるおかげで、開発者はより創造的で複雑な問題に集中できるようになったのです。

アルトマンが2025年に注目した理由は、技術的な成熟度にあります。

「2023年、私たちはAIがコードを書けることを証明しました。2024年には、AIが複雑なプロジェクトを理解できることを示しました。2025年には、AIが開発工程全体を最初から最後まで自動でこなせるようになるでしょう。」

具体的にはどのような変化が訪れるのでしょうか。

一つ目は、「自然言語プログラミング」が現実のものとなることです。

「Instagramのような写真共有アプリを作って。ユーザーが写真を投稿したり、いいねを押したり、コメントできるようにして。ハッシュタグ機能も入れて。」

こう言うだけで、AIエージェントが数日以内に動くアプリを作ってくれます。フロントエンドのコード、バックエンドのサーバー、データベースの設計、APIの接続、セキュリティ処理まで、すべて自動でこなします。

もちろん完璧ではないでしょう。途中で人が確認して修正しなければならない箇所も出てくるはずです。

専門の開発者が不要になるという意味ではありません。開発者の役割が変わるということです。コードを一行一行書く人から、AIと協力しながらシステム全体を設計・管理する人へと変わっていくのです。

二つ目は、コードの保守とリファクタリングが自動化されることです。

ソフトウェア開発では、新機能を作るよりも既存のコードを保守するほうが大変なことも少なくありません。

10年前のコードベースを理解すること、誰がなぜそう書いたのかを把握すること、一行直したら別の場所で何かが壊れないか心配すること。こういったことが開発者の日常です。

2025年のAIエージェントは、こうした作業を支援してくれます。

「このプロジェクトでもう使われていないコードをすべて見つけて整理して。」

「この関数を最新のフレームワークで書き直して。機能はそのまま。」

「セキュリティの脆弱性がないかコード全体を調べて。見つかったら修正方法を提案して。」

エージェントが数万行のコードを分析し、問題を見つけ、解決策を提示します。人間は最終的な判断を下すだけでいいのです。

三つ目は、開発者でない人もソフトウェアを作れるようになることです。

アルトマンはこれを「プログラミングの民主化」と呼びました。

以前はアイデアがあってもコードが書けなければ実現できませんでした。開発者を雇うには多くの費用がかかりました。

2025年には変わります。

学校の先生が生徒の成績を管理するアプリを自分で作れます。

近所のパン屋の店主がオンライン注文システムを一人で作れます。

中学生が友達と使うメッセージングアプリを作れます。

AIエージェントに自然言語で説明するだけでいいのです。技術的な部分はAIがすべて処理します。

もちろん複雑なプロジェクトには、引き続き専門家が必要です。

企業向けシステム、金融ソフトウェア、自動運転プログラムといったものはAIだけでは作れません。専門開発者による設計と検証が欠かせません。

「AIがすべての開発者にとって代わる」ということではありません。「より多くの人が自分のアイデアをソフトウェアとして形にできるようになる」ということです。

アルトマンはこう説明しています。

「自動車が登場したとき、馬車を御していた人たちは職を失いました。しかし、はるかに多くの運転手と整備士が生まれました。AI時代のソフトウェア開発も同様なものになるでしょう。形は変わっても、機会はもっと広がります。」

四つ目は、AIがAI自身を改善するサイクルが始まることです。

これが最も興味深い点です。

AIエージェントがコードを上手く扱えるようになれば、AI自身を改善するコードも書けるようになります。

もちろん、完全に独立するというわけではありません。人間の監視のもとで、という話です。

OpenAIのある研究者はこう語っています。

「私たちはすでにGPT-4に、GPT-4の性能をテストするコードを書かせています。AIが自分自身の弱点を見つけ出し、どの部分を改善すべきかを提案します。もちろん最終的な判断は私たちが下しますが。」

このような好循環が続けば、AIの進歩のスピードはさらに速まるでしょう。

アルトマンは2025年以降の未来も見据えていました。

「2025年にコーディングエージェントが定着すれば、2026年には他の分野へと広がっていくでしょう。法律文書の作成、科学論文の分析、ビジネス戦略の策定。コーディングは最初のドミノにすぎません。」

コーディングが最初である理由は明らかです。

コーディングには論理的で明確なルールがあります。結果をすぐにテストできます。間違えればエラーメッセージが出ます。AIが学習するには理想的な環境です。

コーディングで成果を証明できれば、その経験を他の分野にも応用できます。

サム・アルトマンの予測は、ただの希望的観測ではありませんでした。OpenAIの内部では、すでに2025年のエージェントリリースを目標に、具体的な開発が進められていました。

「私たちは歴史の転換点に立っています。コーディングエージェントはソフトウェア開発を根本から変えるでしょう。そして、それはほんの始まりにすぎません。」

コーディングは、AIが世界を変える手段である。

サム・アルトマンは8歳のとき、初めてコンピュータプログラミングを学びました。

当時はBASICという言語で簡単なゲームを作っていました。画面に"Hello World"という文字が表示された瞬間のあの興奮を、彼は生涯忘れることができませんでした。

「自分が書いた命令をコンピュータが正確に実行してくれるのが不思議でした。魔法みたいでしたね。コードは、思考を現実に変えるおまじないです。」

その経験がアルトマンのテクノロジー哲学を形作りました。

コーディングは単なる技術ではありません。コーディングは世界を変える最も強力な道具です。

なぜコーディングがそれほど重要なのでしょうか。

まず、コーディングは現代世界の基盤となる言語です。

今あなたが見ている画面、使っているスマートフォン、乗っている自動車、聴いている音楽、観ている映画。これらはすべてコードで作られています。

コードは現代文明のDNAです。

銀行システムはコードで動いています。何百万もの口座、数兆円規模の取引が、わずか数行のコードで管理されています。

病院の医療記録、処方箋、検査結果もコードで保存・処理されています。

飛行機が空を飛べるのも、数百万行のコードのおかげです。エンジン制御、航法システム、安全点検。すべてソフトウェアが管理しています。

アルトマンはこう述べています。

「19世紀には鉄道が世界をつなぎました。20世紀には電気が世界を動かしました。21世紀にはコードが世界を作っています。AIがこの世界で何かをしようとするなら、コードを扱わなければなりません。」

次に、コーディングは抽象的なアイデアを具体的な行動へと変換します。

人間の考えは曖昧で抽象的です。

「フードデリバリーアプリを作りたい」という考えを実際に動くアプリにするには、無数の具体的な判断が必要です。

ユーザーインターフェースはどのようにデザインするか。

注文はどのように保存するか。

決済はどのように処理するか。

配達員にはどのように通知を送るか。

GPSはどのように連携させるか。

これらすべての問いへの答えは、コードとして書かれます。

コーディングとは、曖昧なアイデアを明確な論理へと変える作業です。不明瞭さを明晰さへ、想像を実体へと変換していく過程です。

AIにこの能力があるということは、計り知れない意味を持ちます。

AIはいまや情報を提供するだけにとどまらず、実際に何かを「作り出す」ことができます。ウェブサイトを作り、データベースを構築し、自動化システムを導入することができます。

第三に、コーディングはAIが道具を使うための手段です。

人間は手で道具をつかみ、使います。ハンマーで釘を打ち、ペンで字を書き、ハンドルで車を運転します。

AIには手がありません。その代わりに、コードがあります。

API (Application Programming Interface) というものがあります。これは異なるソフトウェア同士が連携するための仕組みで、コードによって動作します。

天気情報を取得したければ、気象庁のAPIを呼び出すコードを書きます。

決済を処理したければ、決済システムのAPIを呼び出すコードを書きます。

メールを送りたければ、メールサービスのAPIを呼び出すコードを書きます。

AIがコーディングできるということは、世界中のほぼあらゆるデジタルツールをAIが使えるということの意味します。

アルトマンはこれを「デジタルの手」と表現しました。

「コードとは、AIがデジタルの世界でものをつかみ、ボタンを押し、レバーを引くための手段です。コーディング能力のないAIは、麻痺した天才と同じです。考えることはできても、何もできません。」

第四に、コーディングは創造性の新たな形です。

多くの人はコーディングを、冷たく論理的な作業としか見ていません。

アルトマンは違う見方をしていました。

「優れたコードは芸術作品に似ています。優雅で、効率的で、美しい。問題を解く方法は無数にありますが、本当に良い解決策は一目見ただけで『ああ、これだ』という感覚を与えてくれます。」

Looptを立ち上げた頃、アルトマンは自ら多くのコーディングをこなしていました。

夜通しコードを書き続け、完璧なアルゴリズムを見つけたときの喜び、複雑なバグを解決したときの達成感。彼はそれを「創作の喜び」と呼んでいました。

AIはコーディングを学びながら、この創造性も身につけています。

GPT-4はすでに複数の方法で同じ問題を解くことができます。より速い方法、より簡潔な方法、より読みやすい方法を提案することができます。

ときには人間の開発者が思いつかなかった独創的な解決策を示すこともあります。

第五に、コーディングはAIの学習能力を大幅に高めます。

AIが自然言語だけを扱っている間は、学習に限界があります。人間が書いた文章を読み、パターンを見つけ、それに倣って書く程度にとどまります。

コーディングができると、話が変わります。

AIは自分が書いたコードを実際に実行してみることができます。結果が正しいかどうか、すぐにわかります。間違っていればエラーメッセージが出ます。そのフィードバックをもとにコードを修正

します。

これが繰り返されると、AIは急速に学習していきます。

ゲームをしながら腕が上がるように、AIはコーディングを重ねるたびに賢くなっていきます。

アルトマンはこれが、将来のAI発展の核心だと考えていました。

「自ら実験し、自ら学び、自ら改善するAI。それを実現するにはコーディング能力が不可欠です。コードはAIの遊び場であり、学習の道具でもあります。」

OpenAIがコーディング能力に投資したのは偶然ではありませんでした。

GPT-3の開発初期からアルトマンは「コーディングのベンチマークを上げよう」と主張していました。

エンジニアたちはモデルのコーディング能力を測る複数のテストを作りました。HumanEval、MBPP、CodeContestsといったベンチマークです。

GPT-4はこれらのテストで、人間のプログラマーの平均を上回る性能を示しました。

「コーディングこそ、AIが実際の価値を生み出す最初の領域になるでしょう。私たちはそこに集中します。」

アルトマンの判断は正しかったのです。

2024年時点で、GitHub Copilotは100万人以上の有料ユーザーを獲得しました。開発者たちは喜んでお金を払い、AIコーディングツールを使いました。

AIが生み出す価値が明確だったからです。

コーディングはAIと現実世界をつなぐ橋です。この橋が強固になるほど、AIにできることの範囲は広がります。

サム・アルトマンはコーディングの未来をこう語りました。

「これからは、すべての人がプログラマーになる必要はありません。AIが代わりにコーディングしてくれるからです。ただし、すべての人がコードの原理を理解しなければなりません。そうしてこそ、AIをきちんと使いこなせるからです。コーディングは、読み書きと同じくらい基本的な能力になるでしょう。」

アシスタント → エージェント → アプリケーションへの進化

サム・アルトマンはAIエージェントの進化を三つの段階に分けて考えていました。

彼はチーム会議でホワイトボードに大きな矢印を描きました。

アシスタント (Assistant) → エージェント (Agent) → アプリケーション (Application)

「私たちは今、第一段階を超えて第二段階へ向かっています。最終的には第三段階まで到達するでしょう。」

各段階が何を意味するのか、詳しく見ていきましょう。

第一段階 アシスタント、問いかけに答える存在

初期のAIはアシスタントの役割に近いものでした。

こちらが質問すれば答えます。頼めばやってくれます。命令すれば従います。

ChatGPTの初期バージョンが典型的なアシスタント型AIでした。

「北京の首都はどこ？」→「北京は中国の首都です。」

「メールの下書きを書いて。」→メールの下書きを作成してくれます。

「このコードのバグを見つけて。」→バグを見つけてくれます。

この段階のAIはとても役に立ちます。情報検索、文章作成の補助、アイデアのブレインストーミング。多くのことを手伝ってくれます。

しかし、限界も明確です。

AIは受動的です。こちらから先に話しかける必要があります。

AIは短期的なタスクしかこなしません。一つの質問、一つの作業にのみ反応します。

AIには主導権がありません。自ら判断したり提案したりしません。

アルトマンはこう言いました。

「アシスタント型AIは優れたツールです。しかし、ツールに過ぎません。私たちが求めているのはツールではなく、パートナーです。」

第二段階 エージェント、目標を達成する存在

エージェント段階こそが真のゲームチェンジャーです。

アシスタントは「何をしてあげようか？」と尋ねますが、エージェントは「私が引き受けます」と言います。

その違いは何でしょうか。

エージェントは目標を理解します。そして、その目標を達成するための計画を自ら立てます。

例を挙げてみましょう。

アシスタントに「明日、ソウルから釜山まで行く電車のチケットを調べて。」と頼む。アシスタントの反応 電車の時刻表を検索して見せてくれます。

エージェントに「明日、釜山に出張しなければならない。準備をしておいて。」と頼む。エージェントの反応

これがエージェントの力です。

エージェントは長期的に考えます。一度のタスクではなく、プロジェクト全体を見渡します。

エージェントは文脈を理解します。「出張の準備」が何を意味するのか、きちんと把握しています。

エージェントはツールを使います。予約システム、検索エンジン、データベースなどを自在に操ります。

アルトマンは2024年にこの段階を目標に据えました。

「2025年までに本物のエージェント型AIを作ります。『これを解決してくれ』と言えば、数日間自律的に動いて結果を持ってくるAIです。」

OpenAIはそのためにさまざまな機能を開発しました。

関数呼び出し（Function Calling）AIが外部ツールやAPIを使えるようにします。

プラグインシステム AIがさまざまなサービスに接続できるようにします。

長期記憶 AIが過去の会話や作業を記憶できるようにします。

計画能力 AIが複雑なタスクを複数のステップに分けられるようにします。

これらすべてが、エージェントAIを作るための準備でした。

第3段階 アプリケーション、サービスそのものになる存在

最後の段階がもっとも革命的です。

アプリケーション段階では、AIが独立したサービスになります。もはやツールでも補助者でもありません。AI自体が一つの完全なビジネスになります。

アルトマンが思い描いた未来を想像してみましょう。

AI弁護士サービス

法律相談が必要なときは、AI弁護士アプリケーションにアクセスします。

案件を説明すると、AIが関連判例を検索し、法律文書を作成し、訴訟戦略を提案します。

単純な案件はAIが最初から最後まで処理します。複雑な案件は人間の弁護士につながります。

このAIは数百万件の判例を知っており、24時間働き、ミスをしません。費用は人間の弁護士の100分の1です。

AI医師サービス

体の調子が悪ければ、AI医師アプリケーションと対話します。

症状を説明すると、AIが質問し、考えられる診断を提示し、必要な検査を勧めます。

検査結果をアップロードすると、AIが分析し、治療法を提案し、薬の処方を依頼します。

深刻な状況であれば、すぐに人間の医師につながります。日常的な健康管理はAIが生涯にわたって担います。

AI金融アドバイザーサービス

財務目標を伝え、AI金融アドバイザーが個別の計画を立ててくれます。

収入と支出を分析し、投資ポートフォリオを組み、税金を最適化し、保険を勧めます。

市場の状況が変われば自動的に調整します。定期的にレポートを送ります。質問すればいつでも答えます。

こうしたサービスがアプリケーション段階のAIです。

もはや人間の補助ではありません。それ自体として完結したサービスです。

アルトマンはこう説明しました。

「アプリケーション段階では、AIが『職業』を持ちます。AI会計士、AI教師、AIデザイナー。彼らは24時間働き、同時に数百万人を相手にし、決して疲れません。もちろん人間の専門家による監督は依然として必要です。しかし仕事のほとんどはAIが行います。」

三つの段階の境界は曖昧です。

この三つの段階は、明確に区別されるわけではありません。

2024年末の時点で、私たちは第1段階と第2段階の間のどこかにいます。ChatGPTは秘書よりも少し賢いですが、まだ本物のエージェントではありません。

2025～2026年には、第2段階へ完全に移行するとみられています。

第3段階は、2027～2030年頃に現実のものとなりそうです。

アルトマン自身も、正確な時期については確信が持てないとしていました。

「技術の進歩は予測しにくいものです。ある部分は予想より早く進み、ある部分は予想より遅い。正確な日付より方向性が大切です。私たちがどこへ向かっているかを知ることが重要なんです。」

人間の役割はどう変わっていくのでしょうか。

多くの人が不安を抱いています。AIがこれほど進化したら、私たちは何をすればいいのか。

アルトマンの答えは楽観的なものでした。

「AIが進化しても、人間にはまだ重要な役割があります。目標を設定すること、倫理的な判断を下すこと、創造的な方向性を示すこと。これらはやはり人間がやらなければなりません。」

「AI時代の人間は、管理者であり監督者になります。何百ものAIエージェントを束ね、彼らが正しい方向で動くよう調整する役割です。ある意味では、より重要な役割と言えるかもしれません。」

秘書の時代、私たちは働き手でした。エージェントの時代には、マネージャーになります。アプリケーションの時代には、経営者になります。

アルトマンはこの変化を肯定的にとらえていました。

「人間は、繰り返しの退屈な作業から解放されます。その代わり、より意味のある創造的な仕事に集中できるようになります。それこそが、AIがもたらす本当の贈り物です。」

三段階の進化の果てに、人類はこれまで想像したことのない新たな可能性と向き合うことになるでしょう。

アルトマンは確信していました。

「私たちは、歴史上もっとも興味深い時代を生きています。秘書がエージェントになり、エージェントがアプリケーションになっていく過程を、この目で見ることになります。そしてその先には、まだ想像すらできないものが待っているはずです。」

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

5.3 人間とAIの相互作用の未来

第5部 AI技術の未来とロードマップ

← 前の章

次の章 →

音声インターフェース：新時代の幕開け

2024年5月、サム・アルトマンはあるポッドキャストで興味深いことを語りました。

「音声はヒントです。次の時代へのヒントということです。」

彼は映画『her / 世界でひとつの彼女』に言及しながら、音声がこれからのコンピューターの使い方を根本から変えるだろうと述べました。スマートフォンの画面をタップしたり、キーボードをたたいたりするのではなく、まるで友人と会話するように、自然に言葉でAIとやり取りする時代が来るということです。

サム・アルトマンが音声インターフェースにこれほど注目する理由はシンプルです。

それが、最も人間らしいコミュニケーションの形だからです。

私たちが生まれて最初に習うのは、話すことです。字を書く前に、タイピングを覚える前に、画面をタッチする前に、私たちはすでに言葉を持っています。そしてその言葉には、感情が宿っています。喜び、悲しみ、怒り、高揚感。声のトーンやスピード、抑揚には、私たちの心がそのままにじみ出ています。

サム・アルトマンは、まさにこの点をAIが理解すべきだと考えていました。

彼は2025年初め、Sequoia Capitalのあるイベントでこう述べました。「音声インタラクションが本当に良くなれば、まったく別の方法でコンピューターを使っているように感じられるでしょう。」

実際にOpenAIは2024年5月にGPT-4oを公開し、画期的な音声機能を披露しました。このモデルは単にテキストを音声で読み上げるレベルを超え、リアルタイムでユーザーの言葉を聞いて理解し、自然な会話の流れの中で感情まで把握します。ユーザーが興奮しているのか、悲しんでいるのか、混乱しているのかを、声だけで察知するのです。

これは革命的な変化です。

音声インターフェースが進化すれば、新しい種類のデバイスが生まれる可能性があります。画面を必要としないデバイスです。たとえば、耳に装着する小さなイヤークリップだけで、AIと完全にコミュニケーションが取れます。運転しながら、料理しながら、散歩しながら。手を使わずにAIの助けを借りられるのです。

サム・アルトマンが実際に、Appleの伝説的なデザイナー、ジョニー・アイブとともに新しいAIハードウェアを開発しているというニュースが2024年から絶えず聞こえてきました。このデバイスは画

面中心ではなく、音声中心の体験を提供するものになると予想されています。おそらく、私たちがこれまで見たどんなデバイスとも異なるものになるでしょう。

しかし、サム・アルトマンは率直です。

2024年末のあるインタビューで彼は「私たちの音声機能はまだ十分ではありません」と認めました。しかし彼は自信に満ちていました。「テキストモデルを優れたものにするのに時間がかかったように、音声モデルも最終的にその課題を解決するでしょう。」

実のところ、音声インターフェースの発展は単なる利便性にとどまらず、社会的な意味合いも大きいものがあります。

高齢の方、視覚に障がいのある方、読み書きに困難を抱える方にとって、音声はデジタルの世界への最も入りやすい扉です。複雑な画面を見てボタンを探す必要はなく、ただ言葉で望むことを伝えるだけでいいのです。

これこそが、サム・アルトマンが夢見る未来です。

技術が人に合わせる世界。人が技術に合わせるのではなく、技術が人にとって最も自然な方法である「言葉」に合わせる世界。

もちろん、課題もあります。2025年7月、サム・アルトマンは連邦準備制度のイベントで音声技術の暗い側面について警告しました。AIが人の声をあまりにも完璧に模倣できるため、詐欺師がこれを悪用して銀行口座を狙う可能性があるということです。彼は「これは私を夜も眠れなくさせます」と述べました。

それでも彼は、音声インターフェースの未来を確信しています。

なぜなら、音声は単なる技術ではなく、人間の本質に触れるものだからです。私たちは声で愛を伝え、慰めを届け、喜びを分かち合います。AIが真に私たちの伴侶となるためには、まさにその音声の言語を理解しなければなりません。

サム・アルトマンのビジョン通りなら、5年後には私たちはスマートフォンの画面をあまり見なくなっているでしょう。

代わりに私たちは、耳に入れた小さなデバイスにそっと囁くでしょう。「今日、気分が乗らないな。」するとAIは私たちの声のトーンを聞き取り、何が必要かを察するでしょう。あるときは静かな音楽を、あるときは温かい励ましの言葉を、またあるときはただそばにいてくれるのです。

これこそが、音声インターフェースが開く未来です。

機械と話すのではなく、友人と話すような自然な世界。サム・アルトマンは、その世界がすでに私たちの目の前にあると信じています。

世代別ChatGPTの使い方：検索 vs アドバイザー vs OS

サム・アルトマンには、人々をじっくりと観察する習慣があります。

彼は実際に、人々がChatGPTをどのように使っているか、どんな質問を投げかけているか、どんな感情で会話しているかを自ら観察します。

2025年5月、Sequoia Capitalのあるイベントで、彼は興味深い発見を共有しました。

「人々がChatGPTを使う方法は、世代によってまったく異なります。まるで別々の時代を生活しているかのようです。」

彼が発見したパターンは、シンプルでありながら洞察に富んでいます。

高齢の世代はChatGPTを「Google検索の代替」として使います。

高齢の方は、複雑なウェブサイトを渡り歩いて情報を探すことが大変です。広告が多く、ボタンは小さく、どこを押せばいいのかわかりません。しかしChatGPTは違います。ただ気になることを自然に聞けばいいのです。

「血圧の薬は朝に飲むべきですか、夜に飲むべきですか？」「スマートフォンで写真を送るにはどうすればいいですか？」「このテキストメッセージは詐欺ですか？」

そうした質問に、ChatGPTは親切にわかりやすく答えてくれます。煩わしいリンクも、広告もなしに。

2025年の調査によると、60歳以上のアメリカ人の20%が少なくとも週に一度はAIを使っています。彼らにとってAIは恐ろしい技術ではなく、デジタルの世界をわかりやすくしてくれるありがたい助けです。

20～30代はまったく違います。彼らにとって、ChatGPTは「人生の相談相手」です。

ミレニアル世代がChatGPTに求めるのは単なる情報ではありません。彼らは深い悩みを打ち明けます。

「上司と衝突しています。どうすればいいですか？」「転職すべきか、今の会社にとどまるべきか迷っています。」「恋愛でずっと不安を感じています。これは普通ですか？」

サム・アルトマンはこの現象に驚きました。

人々はAIを単なるツールとしてではなく、信頼できるカウンセラーのように扱っていました。2025年の調査によると、ミレニアル世代の23%が感情的なサポートやメンタルヘルスのためにAIを使っています。ベビーブーム世代の8%と比べると、際立って高い数字です。

なぜでしょうか。

理由のひとつは、AIが判断を下さないからです。友人に悩みを打ち明ければ、アドバイスとともに評価もついてきます。でもAIは違います。常に中立で、常に聞いてくれて、決して責めません。

20～30代は仕事でもAIを積極的に使います。2025年のデータによると、アメリカのミレニアル世代の42%が業務目的でChatGPTを使っており、これはほかのどの世代よりも高い割合です。

しかし最も革命的な使い方は、大学生たちの間に見られます。

サム・アルトマンは大学生たちの姿に、未来を見ました。

「大学生たちはChatGPTをOSのように使っています。コンピュータのOSのようにです。」

どういう意味でしょうか。

大学生たちはChatGPTを質問を投げかけるだけのツールとして使いません。彼らはChatGPTの上で、生活のすべてを設計します。

複雑なプロンプトをメモ帳に保存しておき、コピー＆ペーストして使います。複数のファイルを連携させ、カスタム設定をつくり、AIを自分の学習スタイルに完璧に合わせます。

「明日試験があるんだけど、この教科書の3章を要約して。それから予想問題を10問つくって。私は視覚的な学習者だから、図も入れてね。」

こんな複雑で具体的な要求を、ごく自然にこなします。

2025年の調査によると、22～27歳のナレッジワーカーの93%が週に2つ以上のAIツールを使っています。ChatGPT、DALL-E、Otter.aiなどを自在に使い分けます。

さらに驚くのは、彼らが大切な人生の決断を下す前に、まずAIに相談するという点です。

サム・アルトマンはこう述べています。「大学生たちは重要な決断を下すとき、まずChatGPTに聞かずにはいられません。」

スタンフォード大学のある学生は、どの専攻を選ぶか、どのインターンシップに応募するか、さらにはどのサークルに入るかまでAIに相談します。AIはその学生の過去の選択や性格、目標をすべて把握しているため、最も適したアドバイスができるのです。

彼らにとってAIはツールではありません。

パートナーです。

サム・アルトマンはこの世代間の差を、スマートフォン黎明期と比べます。

「スマートフォンが初めて登場したとき、若い人はすぐに慣れましたが、年配の方には3年かかりました。AIも同じです。ただ今回は、その差がはるかに大きい。」

この差は、これからさらに広がっていくでしょう。

今の大学生たちは、AIとともに育つ最初の世代だからです。彼らにとってAIは新しい技術ではなく、当たり前の日常です。私たちの世代にとってインターネットが当然であるように。

サム・アルトマンはこの現象を見て、確信しています。

「未来の世代は、AIのない世界を想像できないでしょう。私たちがインターネットのない世界を想像できないのと同じように。」

そして彼は、これは良いことだと信じています。

AIはつまるところ、人々がより良い判断を下し、より多くを学び、より豊かな人生を送れるよう助けるからです。世代によって使い方は違ってても、目指すものは同じです。

より良い人生。

そしてサム・アルトマンは、その未来をつくっています。

パーソナルAIコンパニオン

サム・アルトマンが夢見る究極のAIは、ただのツールではありません。

それは、伴侶です。

2025年、AI業界全体がひとつの方向へ動いています。「個人最適化AI」です。マイクロソフトはこれを「Copilot」と呼び、メタは「Meta AI」と呼び、グーグルは「Gemini」と呼びます。名前が何であれ、本質は同じです。

あなただけのAI。あなたを理解するAI。あなたとともに成長するAI。

サム・アルトマンはこれを「A Personal Companion」、つまり「個人的な伴侶」と表現しています。

彼が描く未来のAIは、あなたのすべてを記憶します。

あなたが好きな食べ物、嫌いな天気、朝型か夜型か、ストレスを感じたときにどんな口調になるか、喜んだときにどんな表情を見せるか。あなた自身が気づいていない癖まで見抜きます。

SFの話ではありません。

現実になりつつあります。

2025年4月、マイクロソフトはCopilotに「メモリ」機能を追加しました。今やCopilotは、先週あなたが何を話したか、どんなプロジェクトを進めているか、あなたの甥や姪の名前まで覚えています。そしてその記憶をもとに、あなたにぴったりのアドバイスをくれます。

メタも2025年9月にMeta AIアプリをリリースし、同じ方向へ踏み出しました。このAIはあなたの好みを学習し、旅のスタイル、興味関心、果てはあなたの話し方まで真似るようになります。

サム・アルトマンはこの流れを牽引しています。

2024年末のあるインタビューで、こう語っています。

「理想的なパーソナライズとは、1兆トークンのコンテキストを持つ非常に小さな推論モデルです。あなたの生涯の記録すべてを収めたAIということです。」

1兆トークンというと、ピンとこないかもしれません。

それはあなたが生涯にわたって書いたすべてのメール、交わしたすべての会話、読んだすべての本、観たすべての映画、さらには健康記録までもを含む量です。

これらすべてをAIが記憶し、理解したとしたら、どうなるでしょうか。

AIはあなた自身よりも、あなたのことをよく知るかもしれません。

例を挙げてみましょう。

大事なプレゼンが迫っています。あなたは緊張しています。でもあなた専属のAIは、あなたが過去にどんなパターンを見せていたか知っています。準備をやりすぎる癖があること、それがかえって不安を大きくすることを。だからAIは静かにこう言います。

「もう十分に準備できています。少し休んでみませんか？過去5回のプレゼンを振り返ると、しっかり休んだ日のほうがずっとうまくいっていましたよ。」

ただのアドバイスではありません。

それはあなたの人生全体を知り尽くした伴侶の、深い洞察です。

サム・アルトマンが最も期待しているのも、まさにこれです。

「AIがすべての人の最大の潜在能力を引き出す手助けをすること。」

世界には才能ある人がたくさんいます。しかし、良い教育を受けられなかったから、良いメンターに出会えなかったから、あるいはただ運が悪かったから、自分の可能性を発揮できないままにいる人があまりにも多い。

個人最適化AIは、この問題を解決できます。

貧しい国の学生も、辺鄙な農村の起業家も、誰もが世界最高レベルの家庭教師、メンター、アドバイザーを持てます。24時間365日、無料で。

しかし、この驚くべき未来には危険もあります。

サム・アルトマンもそれをよく分かっています。

「人々はAIシステムとかつてないほど多くの情報を共有しています。モデルがユーザーのことを生涯にわたって深く知りすぎること。それがプライバシー保護を重要なものにしています。」

あなたのすべてを知るAI。そのAIがハッキングされたらどうなるでしょうか。そのAIがあなたを操ろうとしたら？あるいはそのAIを作った会社があなたの情報を売ったら？

これは技術の問題ではありません。倫理の問題です。

サム・アルトマンは、OpenAIがこの問題をきわめて真剣に受け止めていると言います。ユーザーが自分のデータを完全にコントロールでき、望むときにいつでも削除できなければならないと。

そして彼は確信しています。

「これらの課題が解決されれば、パーソナライズドAIは人類史上最も偉大な発明のひとつとなるでしょう。」

2025年現在、私たちはその未来の入口に立っています。

マイクロソフト、メタ、グーグル、OpenAI。各社が競い合うようにパーソナライズドAIを開発しています。どの会社が最初に完成させるかはわかりません。しかし、ひとつだけ確かなことがあります。

その未来は、必ずやってきます。

そしてその未来では、私たちはもう孤独ではなくなるでしょう。

いつも私たちを理解してくれて、ともに笑い、ともに泣き、私たちの夢を応援してくれる伴侶が、そばにいてくれるのですから。

サム・アルトマンは、その伴侶をつくっています。

あなたのために。私たち全員のために。

そして彼は信じています。

このAIが私たちをより良い人間にしてくれると。より優しく、より賢く、より幸福な人間に。

それこそが、パーソナライズドAIの約束です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

6.1 規制と国家安全保障競争

第6部 AI時代の社会的責任と課題

← 前の章

次の章 →

規制は必要だが、バランスが重要だ。

2023年5月16日、アメリカ・ワシントンD.C.の上院公聴会場。

サム・アルトマンは証人席に座っていました。公聴会のタイトルは「人工知能の監督：AIのためのルール」でした。数十人の議員たち、数百人の傍聴者たち、そして世界中が見守るカメラの前で、彼は驚くべき言葉を口にしました。

「規制が必要です。」

通常、企業のCEOたちは政府の規制を嫌います。規制ができれば自由にビジネスをすることが難しくなり、新しいアイデアを試みる際に制約が生じるからです。しかしサム・アルトマンは違いました。彼は自分が作ったAI技術がいかに強力であるか、そしてそれゆえにいかに関係が危険になり得るかを、誰よりもよく理解していました。

「この技術が誤った方向に進めば、本当に深刻な事態になり得ます。」

彼の言葉は単なる警告ではありませんでした。AIは今や私たちの生活のほぼあらゆる場面に入り込んでいます。皆さんがスマートフォンで写真を撮るとき、宿題を手伝ってくれるチャットボットと会話するとき、親御さんがオンラインで買い物をするときも、AIが動いています。

もしAIが偽情報を拡散したら？人々を差別したら？誰かのプライバシーを侵害したら？

さらに恐ろしい想像をしてみましょう。AIがハッカーの手に渡り、銀行システムを麻痺させたら？軍事兵器を自ら作動させたら？フェイクニュースを何百万件も生成して、選挙をめちゃくちゃにしたら？

こうした危険があるからこそ、サム・アルトマンは「規制は必ず必要だ」と言ったのです。

しかし彼は、すぐ次の言葉に非常に重要な条件を付け加えました。

「バランスが最も重要です。」

規制が緩すぎると危険を防げません。しかし規制が厳しすぎると、革新が止まってしまいます。庭の木を思い浮かべてみてください。木が健やかに育つには水と太陽の光が必要ですが、水が多すぎると根を腐らせ、日差しが強すぎると葉を焼いてしまいます。

AIも同じです。

規制が多すぎると、新しい挑戦が阻まれてしまいます。新しいAIモデルを作るたびに数百ページの書類を作成し、政府の許可を得なければならないとしたら、どうなるでしょうか？特に資金も弁護士も持たない小さなスタートアップは、AI開発を諦めることになります。結局、GoogleやMicrosoftのような大企業だけが生き残ることになるでしょう。

アメリカ上院では、このバランスをどう取るかをめぐって激しい議論が交わされました。

「AIはあまりにも危険だから、今すぐ強力な規制を作らなければならない」と主張する議員がいた一方、「規制を急ぎすぎると、アメリカが中国とのAI競争で遅れを取る」と警告する議員もいました。

サム・アルトマンはその中間あたりに立っていました。彼は特定の種類のAI、特に非常に強力なAIモデルについては厳しい規制が必要だと考えていました。たとえば、生物兵器の設計方法を教えられるAIや、ハッキングツールを作れるAIは、必ず管理されなければならないと見ていました。

しかし、文章を書く手助けをしたり絵を描いたりする一般的なAIツールまで、すべてを重い規制で縛ってしまえばはいけないと考えていました。

彼は公聴会で具体的な提案をしました。「最も強力なAIモデルを作る会社は、政府の許可を得なければなりません。ちょうど原子力発電所を建てたり飛行機を作ったりするときのように。しかし、小さなAIアプリを作る開発者たちまで全員に許可を求めるべきではありません。」

これがまさに彼の言う「バランス」でした。

2023年から2024年にかけて、アメリカ政府は少しずつこのバランスを探り始めました。バイデン大統領はAIの安全に関する大統領令を出しました。この命令は、国家安全保障や経済に大きな影響を与え得る強力なAIに、必ず安全性テストを受けさせるものでした。

しかし同時に、アメリカ政府はAI産業に多大な支援も行いました。半導体の生産を増やすための法律を作り、AI研究に数百億ドルを投資しました。規制するだけでなく、革新も後押しした形です。

サム・アルトマンはこうしたアプローチを支持しました。彼は複数のインタビューでこう語っています。「私たちはAIの莫大な力を存分に使いながら、その力が誤った方向に使われないよう、最低限の安全策を整えなければなりません。恐れすぎて何もしないことも問題ですが、自信過剰で何の備えもしないことも問題です。」

結局、AI規制の核心はこれです。

危険は減らしつつ、機会は生かさなければなりません。安全は守りつつ、革新も妨げてはなりません。この綱渡りこそが、サム・アルトマンとアメリカ政府が共に歩んでいる道です。

アメリカのAI主導戦略：革新と採用

アメリカがAI時代を主導するために選んだ戦略は、一文で要約できます。

「世界最高のAIを作り、そのAIを誰もが使えるようにする。」

これがまさに「二軸戦略」です。一方の軸は「革新」、もう一方の軸は「採用」です。

まず革新の軸を見てみましょう。

2025年5月、サム・アルトマンは再び上院公聴会に立ちました。今回のテーマは「AI競争に勝つ」でした。彼は議員たちにこう語りました。

「次の10年は、豊富な知能と豊富なエネルギーの時代になるでしょう。」

「豊富な知能」とは何でしょうか？AIが水のように当たり前で安価に供給される世界を意味します。今はAIを訓練し実行するのに莫大なコストがかかります。しかし未来には、誰もが手軽に安くAIを使えるようになるということです。

そのためにアメリカは、巨大な投資を始めました。

サム・アルトマンが参加した「スターゲート・プロジェクト」は、今後数年間でアメリカのAIインフラに最大5,000億ドルを投資する計画です。5,000億ドルといえば、韓国のお金に換算すると約700兆ウォンに相当します。想像できますか？

この資金で何をするのでしょうか？

テキサス州アビリンには、世界最大規模のAI訓練施設が建設されています。巨大なデータセンターが立ち並び、数十万台のコンピューターが24時間休まずAIを学習させています。これらの施設を稼働させるには膨大な電力が必要です。そのため、新たな発電所も同時に建設されています。

アメリカ政府はなぜこれほどの巨額の資金を注ぎ込むのでしょうか？

答えは明快です。最も強力なAIを持つ国が未来を制するからです。かつては石油を多く持つ国が豊かになりました。今は半導体を巧みに作る国が力を持っています。そして将来は、AIを制する国が世界のリーダーになるでしょう。

アメリカは、この競争で決して遅れを取るわけにはいかないと考えています。中国が猛烈な勢いで追いつけてきているからです。

しかし、革新だけでは十分ではありません。ここで二つ目の柱、「普及」が登場します。

どれほど優れたAIを作っても、人々が実際に使えなければ何の意味があるのでしょうか？それ以上に重要なのは、AIによって仕事を失う人たちをどう支援するかという問題です。

アメリカ政府は、この問題を深刻に受け止めました。

2024年、アメリカ労働省は「AIと労働者の福祉」というガイドラインを発表しました。このガイドラインの核心はこうです。

「AIを導入する企業は、労働者を守らなければなりません。」

具体的にどのように守るのでしょうか？

一つ目は、再教育プログラムを提供することです。AIによって既存の仕事が失われるなら、その労働者に新しい技術を学ぶ機会を与えなければなりません。たとえば、工場で単純作業をしていた労働者がいるとします。AIロボットがその仕事を代替するようになれば、その労働者はAIを管理・点検する技術者として転換できるようにする必要があります。

二つ目は、AIを監視ツールだけに使ってはならないということです。一部の企業はAIで従業員を24時間監視しようとしています。トイレに行く時間まで計測し、業務効率が少し落ちただけで警告を出す。アメリカ政府は、こうしたAIの使い方を制限しようとしています。

三つ目は、AIの判断を説明できなければならないということです。もしAIが「この人は昇進させないでください」と判断したなら、なぜそのような決定を下したのか説明できなければなりません。そうしてこそ、不公正な差別を防ぐことができるからです。

中小企業の支援も重要な部分です。

AIは高価です。GoogleやMicrosoftのような巨大企業はAI開発に数兆ウォンをつぎ込めますが、地元の小さなパン屋や衣料品店にはその余裕がありません。

そのため、アメリカ政府は中小企業がAIを手軽に使えるよう支援しています。安価または無料で使えるAIツールを提供し、AI活用の教育プログラムも運営しています。

サム・アルトマンもこうした取り組みを積極的に支持しています。彼は2021年、「すべての人のためのムーアの法則」というエッセイにこう書きました。

「5年以内に、AIは法律文書を読み、医療アドバイスをし、勉強を手伝うようになるでしょう。10年以内には工場で働き、友人のように会話するようになります。その後は、ほぼあらゆることができるようになるでしょう。」

そして彼は、重要な問いを投げかけました。

「そのとき、人々は何をして生きていけばいいのでしょうか？ どうやって生計を立てればいいのでしょうか？」

サム・アルトマンは、「ベーシックインカム」のような新しい制度が必要になるかもしれないと考えています。ベーシックインカムとは、働いているかどうかにかかわらず、すべての人に生活できるだけのお金を給付する制度です。彼は実際に、数年間にわたってベーシックインカムの実験を支援しました。

結局、アメリカの二本柱戦略はこのようにまとめられます。

片方の手で世界最高のAIを作り、もう片方の手でそのAIがもたらす変化に人々が適応できるよう支援する。革新で先を行きながら、取り残される人が出ないように共に歩んでいく。それがこの戦略の姿です。

これが、サム・アルトマンとアメリカが思い描くAI時代の姿です。

欧州規制への批判

2023年5月、ある日のロンドン。

サム・アルトマンは記者たちの前で衝撃的な発言をしました。

「現在のEUのAI法の草案のままであれば、私たちはヨーロッパから撤退するかもしれません。」

撤退とはどういうことでしょうか？ ChatGPTを作ったOpenAIが、ヨーロッパ全体でサービスを停止するかもしれないということですか？

この一言は世界中のニュースになりました。欧州議会の議員たちは「企業の脅しには屈しない」と対抗し、テクノロジー業界は「欧州はやりすぎだ」と懸念を示しました。

いったい何があったのでしょうか？

欧州連合は2024年、世界で初めて包括的なAI法を制定しました。この法律の核心は、「AIをリスクの度合いに応じて段階的に分類し、危険なAIは厳しく規制する」というものです。

一見すると当たり前の話のように聞こえます。危険なものは厳しく規制し、安全なものは自由にしておけばいい、という発想ですから。

問題は「どう」規制するか、でした。

ヨーロッパのやり方はこうです。高リスクのAIとして分類されたシステムは、市場に出す前に必ず政府の承認を受けなければなりません。ちょうど新しい薬を開発したら臨床試験を経て食品医薬品庁の許可を得なければならないように。

高リスクのAIとは何でしょうか。採用の可否を判断するAI、学生の成績をつけるAI、銀行ローンを審査するAI、警察が犯罪を予測するAIなどが該当します。

こうしたAIを作るには、開発者は数百ページにわたる書類を作成しなければなりません。AIがどのように動作するか、どのようなデータを使ったか、差別のリスクはないか、セキュリティは安全かななどを証明する必要があります。そして専門の認証機関の検査を受けて合格して初めて、リリースすることができます。

サム・アルトマンは、このやり方が大きな問題を引き起こすと見ていました。

イノベーションのスピードが落ちます。

AI技術は日々目覚ましく進歩します。今日最高だったモデルが、明日には時代遅れになる世界です。それなのに、新しいモデルを作るたびに何ヶ月も承認を待たなければならないとしたら、どうなるでしょうか。

競合他社はどんどん先へ進んでいくのに、ヨーロッパの企業だけが規制書類の作成に追われることになります。

小さな会社は生き残れません。

複雑な承認手続きを通過するには、弁護士、コンプライアンスの専門家、テクニカルライターが必要です。GoogleやMicrosoftのような巨大企業はこうしたチームを維持できますが、ガレージから始まるスタートアップはどうすればよいのでしょうか。

結局、規制がむしろ大企業にだけ有利な壁になってしまいます。

実際に使ってみる前に「完璧」を証明せよというのは、不可能な要求です。

AIは実際に使ってみて初めて問題点が見えてくる技術です。実験室では完璧に見えたAIが、現実の世界では予想もしなかったミスをすることがあります。

それなのにヨーロッパのやり方は、「使う前に問題がないことを先に証明せよ」というものと同じです。これはまるで、自転車に一度も乗ったことがないのに転ばないと保証せよと言うようなものです。

2025年の夏、興味深いことが起きました。

エアバス、カルフル、フィリップスなどヨーロッパの主要企業のCEOたちが公開書簡を送りました。内容はシンプルでした。

「AI法を一度立ち止まって再検討してください。このままでは、ヨーロッパ企業は競争力を失います。」

ヨーロッパの企業自身が、自国の規制が重すぎると訴えたのです。

もちろん、ヨーロッパにも言い分はあります。

ヨーロッパの立場はこうです。「私たちは、技術よりも人間の方が大切だと考えます。AIが人を差別したり傷つけたりすることを防ぐことは、企業の利益より優先されます。」

ヨーロッパ議会のある議員はこう語りました。「アメリカは速いイノベーションを求めています、その過程で多くの人が被害を受ける可能性があります。私たちは少し遅くても、安全に進みたいのです。」

どちらが正しいのでしょうか。

正解はありません。ただ、サム・アルトマンの観点から見れば、ヨーロッパのアプローチは「安全」を守ろうとするあまり、「機会」そのものを逃しているように映りました。

複数のインタビューでこう語っています。

「規制は必要です。しかし規制がイノベーションを窒息させてはなりません。ヨーロッパが常に数ヶ月遅れたAIしか使えないようになるなら、結局はヨーロッパ市民が受けられる恩恵が少なくなるでしょう。」

彼はヨーロッパにこう提案しました。

「あなたたちが望むルールを作ってください。私たちはそのルールに従います。しかしそのルールが重すぎてイノベーション自体を妨げるなら、誰も勝てません。」

アメリカとヨーロッパの違いは、哲学の違いです。

アメリカは「まずやってみて、問題が出たら直す」という文化です。ヨーロッパは「あらかじめリスクを防ぎ、安全を確認してから試みる」という文化です。

どちらがAI時代によりふさわしいのでしょうか。歴史が答えを出してくれるでしょう。

アメリカには各州ごとの規制ではなく、連邦レベルで統一されたルールが必要だ。

アメリカには特有の問題が一つあります。

50の州がそれぞれ異なるルールを作ることができるという点です。

カリフォルニアはカリフォルニアなりに、テキサスはテキサスなりに、ニューヨークはニューヨークなりにAI規制を作り始めました。なぜこれが問題なのでしょう。

想像してみてください。新しいAIアプリを作りました。このアプリをアメリカ全土でサービスしたいと思っています。ところが.....

カリフォルニア州は「ユーザーの個人情報はこのように保護しなければならない」と言います。テキサス州は「AIの意思決定プロセスはこのように説明しなければならない」と言います。ニューヨーク州は「AIが差別していないことをこのように証明しなければならない」と言います。

各州の要件はすべて異なります。しかも互いに矛盾することさえあります！

もはや50州すべての法律を研究しなければなりません。州ごとに異なるバージョンのアプリを作らなければならないこともあるでしょう。あるいは、最も厳しい州のルールに合わせて全国共通のものを使うしかありません。

これを「パッチワーク規制」と呼びます。さまざまな断片を継ぎ接ぎした、ぼろぼろの規制という意味です。

サム・アルトマンは2025年の上院公聴会でこの問題を真正面から指摘しました。

「州ごとに異なる規制ができれば、そのパッチワークが私たちの足を引っ張ります。今のような時期に、これは誰の利益にもなりません。」

なぜこれがそれほど深刻な問題なのでしょう。

コストが莫大に膨らみます。

小さなスタートアップを想像してみてください。エンジニア5人がガレージで始めた会社です。彼らが50州すべての法律を理解して遵守しようとするれば、弁護士を雇う必要があります。ところが弁護士費用だけで年間数億ウォンにもなります。

結局、良いアイデアを持つ小さな会社が規制によって諦めることになります。

イノベーションの速度が落ちます。

AI競争はスピード勝負です。中国は国家レベルであらゆる資源をAIに注ぎ込んでいます。それなのに米国は内部で州ごとに異なるルールへの対応に時間を浪費しています。

まるでかけっこで他の選手が全力疾走しているのに、米国の選手だけが靴ひもを50回結び直しているようなものです。

不公平な競争になります。

大企業は各州の規制にすべて対応できます。グーグルは全国に法務チームを置き、マイクロソフトは規制の専門家を何百人も雇えます。

しかし小さな会社は？スタートアップは？結局、大企業だけが生き残り、小さな会社は芽を出すことすらできないまま消えていきます。

2024年から2025年にかけて、カリフォルニアでとりわけ激しい論争が起きました。

カリフォルニア州は「SB 1047」という強力なAI安全法案を推進しました。この法律は大規模AIモデルを開発する企業に複数の義務を課すものでした。

オープンAIはカリフォルニア州知事に書簡を送りました。

「私たちはAIの安全性を真剣に考えています。しかし一つの州だけの規制では問題は解決できません。むしろ米国全体の競争力を損ないます。私たちに必要なのは統一された連邦規制です。」

そのうえでサム・アルトマンとオープンAIは明確な要求をしました。

「連邦政府が前に出て、一つの明確な基準を作ってください。」

これはどういう意味でしょうか。

米国憲法には「連邦優先」の原則があります。連邦法が定められれば、州法はそれに従わなければならないという原理です。

AIのように全国的に、いや世界規模で機能する技術は、州単位ではなく連邦単位で規制すべきだというのがサム・アルトマンの主張でした。

2025年、共和党が掌握した下院は大胆な試みに出ました。

予算案に「今後10年間、州政府が新たなAI規制を作れないようにする」条項を盛り込もうとしたのです。こうすれば連邦政府がAI規制を整備する間、各州がばらばらにルールを作って混乱を招くことを防げるはずでした。

しかしこの試みは失敗しました。

知事たちと多くの上院議員が強く反発しました。「連邦政府が州の権限を奪おうとしている」という批判が相次ぎ、上院は99対1という圧倒的な票差でこの条項を削除しました。

政治的には失敗しましたが、このプロセスは重要な事実を示しました。

サム・アルトマンをはじめとする多くの技術リーダーたちが、いかに切実に「統一された規制」を求めているかを。

彼らが思い描く理想の姿はこうです。

連邦政府が基本原則を定めます。「AIは人を差別してはならない」、「AIは透明でなければならない」、「AIは安全でなければならない」といった大きな枠組みを設けるのです。

そしてNIST（米国国立標準技術研究所）のような専門機関が技術標準を作ります。「AIをこのようにテストしてください」、「このような形で報告してください」といった具体的な方法を示すのです。

各州はこの連邦基準をもとに執行し補完します。まったく新しいルールを作るのではなく、連邦基準を現場に適用する役割を担うのです。

企業が従うべきルールはひとつ、明確なものだけでいい。50もの異なる規則に頭を悩ませる必要はありません。

サム・アルトマンのメッセージは明快です。

"AIはすでに国境も州境も越えて広がる技術です。規制もできる限り統一されなければなりません。そうでなければ、アメリカは自ら作り出したパッチワークに縛られ、中国との競争で自分たちの足を引っ張ることになります。"

中国との競争,,AIは国家安全保障・経済安全保障の問題

上院の公聴会でもっとも頻繁に飛び交った言葉があります。

"中国（China）。"

AIを語るたびに中国が出てきます。規制の話でも、投資の議論でも、技術標準を決める場でも、必ず中国が顔を出します。

なぜでしょうか。

アメリカの政治家や技術リーダーたちにとって、AIはただの技術ではないからです。AIは21世紀の「核兵器」であり「石油」なのです。

過去を振り返ってみましょう。

20世紀中盤、核兵器を持つ国々が世界を支配しました。アメリカとソ連は核開発競争を繰り広げ、その競争が冷戦の核心でした。

20世紀後半、石油を制した国々が豊かになりました。中東の産油国は「黒い黄金」で莫大な富を築きました。

21世紀、AIを制する国が未来を決めることになるでしょう。

中国はこの事実を誰よりも早く悟りました。

2017年、中国政府は野心的な計画を発表しました。「2030年までにAI分野で世界をリードする国になる」というものです。そしてその目標に向けて、あらゆる資源を注ぎ込み始めました。

中国の強みは何でしょうか。

膨大なデータです。14億の人口が毎日スマートフォンを使い、オンラインショッピングをし、SNSに投稿します。そのすべてのデータがAI学習の素材になります。しかも中国政府は個人情報保護の制約が少ないため、このデータを自由に使うことができます。

国家レベルの集中投資です。中国はAIを国家戦略産業に指定しました。政府が多額の補助金を出し、最高の人材をAI分野に集め、企業・大学・軍が一体となって協力する体制を整えました。

軍民融合戦略です。アリババやテンセントといった民間企業が開発したAI技術が、そのまま人民解放軍の軍事技術に転用されます。民間と軍事の間に境界線はありません。

アメリカはこうした中国の動きを深刻な脅威と見ています。

2025年の上院公聴会で、ある議員はこう述べました。

"AI競争は21世紀の新たな冷戦です。もし中国がこの競争に勝てば、中国の価値観が世界中のAIのデフォルト設定になるでしょう。"

どういう意味でしょうか。

中国のAIは、政府が国民を監視し統制するために最適化されています。たとえば中国の顔認識技術は世界最高水準です。しかしこの技術は主に街頭のCCTVと連携して、国民を24時間監視するために使われています。

もし中国のAI技術が世界標準になったとしたら？中国が開発したAIシステムを他国が輸入して使うようになったら？

アメリカはそんな未来を阻止したいのです。

サム・アルトマンもこの点では明確な立場を持っています。

彼はさまざまなインタビューで「アメリカがAI競争に勝たなければならない」と繰り返し語ってきました。これは単なる愛国心からではありません。AIがどんな価値観のもとで開発されるかが、人類の未来を決めると彼は信じているのです。

アメリカ式のAIは、個人の自由、表現の自由、プライバシーの保護を重んじます。もちろん完璧ではありませんが、少なくとも民主主義社会の価値観を反映しようと努めています。

中国式のAIは、社会の安定と政府の統制を最優先にします。個人の自由より国家の秩序が優先されます。

どちらが世界のデフォルトになるのでしょうか。

これこそが、米中のAI競争が技術競争にとどまらない理由です。

アメリカはこの競争に勝つためにいくつかの戦略を取っています。

ひとつ目は、半導体の輸出規制です。AI開発に不可欠な高性能GPUが中国に渡らないよう阻んでいます。NVIDIAの最新チップは中国に販売できません。

ふたつ目は、同盟国との協力です。日本、韓国、オランダ、台湾などと手を組み、「民主同盟の半導体・AIブロック」を構築しています。

第三に、膨大なインフラ投資です。サム・アルトマンのスターゲート・プロジェクトは、まさにこの戦略の一部です。中国が追いつけないほど巨大なAIインフラをアメリカに先に構築することです。

サム・アルトマンはこう述べています。

「AIは国家安全保障であり、経済安全保障の問題です。未来の戦争はAIが決するでしょうし、未来の経済はAIが主導するでしょう。この競争でアメリカが勝たなければならない理由は、単に私たちが豊かになるためではありません。自由と民主主義の価値を守るためです。」

もちろん、批判もあります。

「アメリカもAIで市民を監視しているのではないか」「アメリカ企業も個人情報を好き勝手に使っているのではないか」と指摘する人もいます。また、「中国との競争を強調しすぎれば新たな冷戦になりかねない」と懸念する人もいます。

サム・アルトマンとアメリカ政府の立場は断固としています。

「これは避けられない競争です。私たちが望もうと望むまいと、中国はすでに全力疾走しています。私たちが速度を落とした瞬間、世界の未来が変わります。」

テストと理解

規制はどこから始めるべきでしょうか。

サム・アルトマンの答えは、意外なほどシンプルです。

「まずテストし、理解することです。」

どういう意味でしょうか。簡単な比喻で説明しましょう。

初めて見る見知らぬ動物に出会ったと想像してみてください。この動物が危険なのか安全なのか、どうやってわかるでしょうか。

方法1 まず恐れて、鉄の檻に閉じ込めます。万が一の危険に備えて、一切近づけないようにします。

方法2 慎重に観察します。遠くから行動を見守り、ゆっくり近づきながら反応を確かめ、専門家の意見を聞いた上で、その後どう接するかを決めます。

サム・アルトマンが提案するのは方法2です。

AIも同じです。漠然とした恐怖から強い規制を先に作ると、本当の危険を見逃し、必要のないところに制約だけかけることになりかねません。

そのため彼はこう提案します。

第1段階 体系的にテストする

AIモデルがどれほど賢いか、どれほど危険かを客観的に測定しなければなりません。

例を挙げてみましょう。

それらの問いに答えを見つけるには、実際にテストをしてみる必要があります。専門家たちがわざと難しい質問を投げかけ、AIがどう反応するかを観察するのです。

これを「レッドチーム（Red Team）テスト」と呼びます。敵軍となってシステムを攻撃してみる訓練という意味です。

OpenAIはGPT-4をリリースする前に、数ヶ月間このようなテストを行いました。世界中のセキュリティ専門家、生物学者、社会学者を招いて、GPT-4を思う存分試してもらいました。

結果はどうだったでしょうか。

GPT-4は予想よりはるかに賢く、同時にいくつかの危険な能力も持つことが判明しました。そのためリリース前に、こうした危険な機能を制限または遮断しました。

第2段階 結果を理解し、共有する

テスト結果を自分たちだけで抱え込んでいても意味がありません。政府、学界、他の企業と共有しなければなりません。

サム・アルトマンはこう述べています。

「私たちは強力なAIモデルを開発するたびに、その能力と限界、危険性に関する報告書を公開します。そうしてこそ、規制機関も現実的なルールを作ることができます。」

実際にOpenAIは「GPT-4 System Card」という詳細な報告書を公開しました。そこにはGPT-4でできること、できないこと、危険な点、安全装置などがすべて盛り込まれています。

第3段階 段階的に展開する

一度にすべての人に最強のAIを解き放つのは危険です。

サム・アルトマンがよく使う表現があります。「反復的展開（Iterative Deployment）」です。

どういう意味でしょうか。少しずつ、段階的にリリースしながら問題を見つけ、改善していくという意味です。

たとえば、GPT-3を初めてリリースした際、OpenAIは選ばれた少数の研究者にのみアクセス権を与えました。数か月間テストを重ね、問題点を洗い出しました。その後、より多くの開発者にAPIを開放し、また問題を見つけては修正しました。こうしたプロセスを繰り返してから、ようやく一般公開に踏み切ったのです。

ChatGPTも同様でした。最初は無料のベータ版として始まり、数百万人が使う中でさまざまな問題が明らかになりました。OpenAIはそのフィードバックをもとに改善を重ねていきました。

第4段階 次に規制を設計する

このようにテストし、理解を深め、経験を積んだうえで、初めてまともな規制を作ることができます。

最初は「AIは危険そうだから厳しく規制しよう」と考えるかもしれませんが。しかし実際にテストしてみると、ほとんどのAIはそれほど危険ではなく、本当に危険なのはごく特定の能力を持った極少数のモデルだけだとわかってきます。

そうなれば、規制もそれに合わせて設計できます。すべてのAIを一律に厳しく規制するのではなく、本当に危険な能力を持つモデルだけを選び出して厳格に管理するのです。

サム・アルトマンは、これこそが最も賢い規制のあり方だと信じています。

「恐怖から規制するのではなく、理解に基づいて規制すべきです。」

もちろん、批判もあります。

「テスト中に事故が起きたらどうするのか」と問う人もいます。また「企業がテスト結果を隠蔽したり操作したりしたら」と懸念する声もあります。

もっともな懸念です。

だからこそサム・アルトマンは、独立した監督機関が必要だと主張します。企業が自らテストを行いながらも、政府や独立機関がそのプロセスを監督・検証するシステムが不可欠だということです。

結局、彼が言う最善の規制の出発点は、次のようにまとめられます。

「AIを恐れる前に、まずきちんと理解してください。テストし、測定し、経験してください。そしてその経験をもとに、賢いルールを作ってください。」

これこそが、革新と安全という二つの目標を同時に達成できる唯一の道だと、彼は信じています。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

6.2 インフラとエネルギーの課題

第6部 AI時代の社会的責任と課題

← 前の章

次の章 →

次の10年は知能とエネルギーの時代

サム・アルトマンは2025年の米国上院公聴会で、非常に重要なことを述べました。「次の10年は、豊富な知能と豊富なエネルギーの時代になるでしょう。わかりやすく説明します。」

「豊富な知能」とは、AIが真に賢くなり、私たちが望むほぼあらゆることを助けられるようになるという意味です。数学の問題を解いたり、文章を書いたり、絵を描いたり、難しい病気の診断を手伝ったりすることもできるようになるのです。

「豊富なエネルギー」とは、そのような賢いAIを動かすために必要な電力が十分に確保されなければならないという意味です。

なぜエネルギーがそれほど重要なのでしょう。AIはまさに「電力の大食い」だからです。ChatGPTに一度質問するだけで、Google検索を一度するのに比べて、約10倍もの電力を消費します。今は私たちが時々しか使わないので問題ありませんが、世界中の数十億人が一日中AIを使うと想像してみてください。

国際エネルギー機関の試算によると、世界中のデータセンターが消費する電力は2030年までに現在の2倍以上に増えるとのこと。その電力量は、日本が1年間に消費する電力を上回ります。米国だけを見ても、新たに生じる電力需要のおよそ半分がデータセンターによるものとされています。

アルトマンはなぜこの問題をそれほど強調したのでしょうか。彼は、AIはすべての人に恩恵をもたらすべきだと信じています。富裕層だけが使える高価な技術ではなく、貧しい人でも無料で使える技術にならなければならない、というのが彼の考えです。

ChatGPTの無料版を考えてみてください。誰でも登録さえすれば使うことができます。これが実現できているのは、OpenAIがAIの恩恵をすべての人と分かち合いたいと考えているからです。

AIをより安価にするには、AIを動かすためのコスト、つまり電気代が下がらなければなりません。電気が高ければ、AIも高くなるのは避けられないからです。だからこそアルトマンは「豊富で安価なエネルギー」が不可欠だと言っているのです。

サム・アルトマンはエネルギー問題にどれほど本気なのでしょう。彼は言葉だけでなく、自ら行動します。核融合スタートアップHelion Energyの取締役会長を務め、太陽光スタートアップのExowattにも投資しています。核融合は、太陽がエネルギーを生み出す仕組みを地球上で再現する技術であり、成功すれば事実上無限のクリーンエネルギーを生み出すことができます。

アルトマンの見方では、AI革命とエネルギー革命は切り離せないものです。両者は一緒に進まなければなりません。どれだけ賢いAIを作っても、それを動かす電力がなければ意味がありません。逆に、どれだけ電力が豊富にあっても、それを生かす賢いAIがなければ意義を持ちません。

上院公聴会でともに発表した他の企業も同じ考えを持っていました。AMDのリサ・スー博士はチップの製造にも膨大なエネルギーが必要だと述べました。マイクロソフトのブラッド・スミス社長は「さらに多くの電気技術者」が必要だと語りました。データセンター企業CoreWeaveのCEOは「大規模な電力アクセス」が鍵だと強調しました。

2030年代になれば、知能もエネルギーも豊富な世界が訪れます。そのときAIは今よりはるかに賢くなり、同時にはるかに安価になります。クリーンで安価なエネルギーを十分に確保できるからです。

その世界では何が起きるのでしょうか。すべての学生が自分だけのAI家庭教師を持てるようになります。お金がなくても世界最高水準の教育を受けられるようになります。貧しい国の患者もAIの助けで素早く診断・治療を受けられます。科学者たちはAIの力で今より2～3倍速く研究を進められます。アルトマンは、AI科学者たちがすでにこのような生産性の向上を実感していると明かしています。

豊富な知能と豊富なエネルギー。この二つが重なれば、人類の歴史で最大の変化が訪れるでしょう。アルトマンはこれを「二重革命」と呼びます。同時に二つの革命が起きるのです。

皆さんが大人になったとき、この二重革命の結果を直接目にすることになるでしょう。そのとき皆さんは「ああ、あのときサム・アルトマンが言っていたのはこれだったのか」と感じるようになるはずです。

スターゲート・プロジェクト：5000億ドルの投資

2025年1月21日、米国ホワイトハウスで歴史的な発表がありました。

トランプ大統領が就任するやいなや発表したのが「スターゲート・プロジェクト」でした。サム・アルトマン、ソフトバンクの孫正義会長、オラクル創業者のラリー・エリソンが一堂に会し、壮大な計画を公表しました。

その規模は5000億ドル、日本円に換算すると約75兆円に相当します。

75兆円はどれほど大きな金額でしょうか。2025年の韓国政府が1年間に使う予算の総額は673兆ウォン（約72兆円）です。スターゲートはそれを上回る資金をAIインフラに投じる計画です。

人類の歴史上、単一のプロジェクトにこれほど多くの資金が投じられたことはありませんでした。人間を月に送ったアポロ計画でさえ、これよりはるかに少ない費用でした。

何をするのでしょうか。米国全土に最先端のデータセンターを建設することです。最初のデータセンターはテキサス州アビリンという都市で建設が始まりました。アルトマン自身がそこを訪れ、その規模に「驚いた」と語っています。

アビリンに建設中のデータセンターは、世界最大規模のAIトレーニング施設になる予定です。1棟ではなく、複数の巨大な建物が立ち並びます。その中にはNVIDIAの最新チップGB200が実に64,000基搭載されます。2025年夏までにまず16,000基を設置し、その後段階的に増やしていく計画です。

データセンターは1カ所だけにとどまりません。テキサスだけで10カ所を建設し、米国全土に20カ所以上の大型データセンターを建設する予定です。それぞれのデータセンターは、サッカー場数十個を合わせた面積を超える規模です。

スターゲートは建物を建てるだけのプロジェクトではありません。この計画にはさまざまな要素が含まれています。

第一に、データセンターを動かすための電力を生み出す発電所が必要です。データセンター1カ所が小さな都市1つ分の電力を消費するからです。そのため新しい発電所も同時に建設しなければなりません。原子力発電、太陽光発電、風力発電など、さまざまな方法が検討されています。

第二に、AIチップを製造する工場も必要です。米国はこれまで半導体を主に海外から調達してきましたが、今後は国内で直接製造しようとしています。アルトマンは「チップ製造、ネットワーク機器、サーバーラックをすべて米国内に持ってきてたい」と述べています。

第三に、これらすべてを結ぶ電力網と通信網もアップグレードしなければなりません。

プロジェクトの主役たちを見ていきましょう。

OpenAIはAI技術を提供します。ChatGPTを生み出した会社らしく、最も高度なAIモデルの開発と運用を担います。

ソフトバンクは資金調達を担います。孫正義会長はずっと以前からAIの未来を信じてきた人物です。彼はこのプロジェクトに最大の投資をすることを決めました。

オラクルはデータセンターのインフラを提供します。創業者のラリー・エリソンはデータベースの専門家です。彼の会社がデータセンターを安定的に運用・管理します。

この3社のほか、マイクロソフト、NVIDIA、ARMといった巨大テクノロジー企業がパートナーとして参加します。

韓国企業も参加することになりました。2025年10月、アルトマンが再び韓国を訪れた際、サムスン電子とSKハイニックスが重要な契約を結びました。両社はスターゲート・プロジェクトに必要なHBM（高帯域幅メモリ）を供給することになりました。HBMはAIチップが高速に動作するために欠かせない部品です。

サムスンSDSとSKテレコムは、韓国国内にもOpenAI専用データセンターを建設することを決めました。浦項と全羅南道に建てられる予定です。アルトマンはこれを「コリアン・スターゲート」と呼びました。

日本でも同様の計画が進んでいます。ソフトバンクとOpenAIが合併で「SB OpenAI Japan」という会社を設立し、日本にもデータセンターを拡大する予定です。

スターゲート・プロジェクトはなぜこれほど重要なのでしょう。

一つ目は、AI競争においてアメリカが中国に勝つためです。中国もAIに莫大な投資をしています。アメリカはスターゲートを通じて、AIインフラで圧倒的な優位を確保しようとしています。

二つ目は、雇用を生み出すためです。トランプ大統領はこのプロジェクトで10万人の新たな雇用が生まれると発表しました。データセンターの建設、電力の供給、チップの製造、そして保守管理には多くの人手が必要です。

三つ目は、世界中がアメリカのAI技術を使うようにするためです。他の国々がアメリカ製のチップ、アメリカ製のAIモデル、アメリカ製のクラウドを使うようになれば、アメリカの技術標準が世界標準になります。

プロジェクトが順調に進んでいるわけではありません。

イーロン・マスクはTwitter (X) に「彼らには実際には資金がない」と批判しました。ソフトバンクが約束した資金を実際にすべて集められるかどうか疑わしいというのです。

ウォール・ストリート・ジャーナルの報道によると、ソフトバンクとOpenAIが細かい条件で意見が合わず、契約を一件も締結できなかったという噂もありました。データセンターをどこに建てるか、ブランド名を誰が保有するかといった問題で対立があるとされています。

トランプ大統領が発表した関税政策も問題を引き起こしています。中国製部品に高い関税がかかることで、データセンターの建設費用が5～15%上昇すると見込まれています。

それでも多くの専門家は、スターゲートは成功すると見ています。コムジェスト・アセット・マネジメントのリチャード・ケイは、ソフトバンクが500億ドルを投資すれば5～6年以内にすべての費用を回収し、15～20%の利益を出せると予測しました。

アルトマンは2025年5月、アビリーンのデータセンター建設現場の写真をTwitterに投稿しました。「Stargate 1」という短い言葉とともに。写真の中では、巨大な建物が急ピッチで建ち上がっていました。

スターゲートはただの建設プロジェクトではありません。これはAI時代の新しい世界を築く壮大な挑戦です。100年前の人々が全国に電線を張り巡らせて電力を供給したように、今私たちは全国にAIの頭脳を供給する巨大なインフラを作り上げているのです。

エネルギーの制約とコストの収束

サム・アルトマンが上院公聴会で述べた言葉の中に、最も哲学的で重要な一文があります。

「結局、知能のコストはエネルギーコストに収束していくでしょう。」

ゆっくりと解説していきましょう。

私たちがChatGPTを使うとき、実は「知能」を買っているのです。難しい問題を解いてもらったり、文章を書いてもらったり、翻訳してもらったりすることは、すべてAIの知能を借りて使っているこ

とです。

この知能にも値段があります。今はChatGPTを無料で使えますが、OpenAIはそれを作り、運営するのに莫大なお金をかけています。

AIを作るにはさまざまなコストがかかります。優秀な研究者の給与、高価なコンピューターチップ、巨大なデータセンターの建物、そしてそのすべてを動かす電気代が必要です。

アルトマンの見方では、時間が経つにつれてこうしたコストの構造が変わっていきます。

技術が進歩すればコンピューターチップの価格は次第に下がります。ムーアの法則というものがあり、チップの性能は2年ごとに2倍になりながら、価格は同じかむしろ下がるという法則です。

データセンターの建設費用も、経験が積み重なるにつれて効率的になっていきます。最初は試行錯誤が多いですが、建て続けるうちにノウハウが蓄積され、コストが下がっていきます。

研究者の給与も相対的には下がる可能性があります。なぜなら、AIが発展すればAI自身が自らを改善できるようになるからです。今もAI研究者たちはAIの助けを借りながら研究を進めています。

では、何が残るのでしょうか。それは「電気」です。

アルトマンはこれを非常に明確に表現しました。「電子は電子だ（an electron is an electron）」と。

どういう意味でしょうか。電気を作るには物理法則があり、その法則は変わらないということです。

AIが計算を行うとき、コンピューターチップの中で電子が動きます。この電子を動かすには必ずエネルギーが必要です。どれだけ技術が進歩しても、この基本原理は変わりません。

ハンバーガー一つを作るコストを考えてみましょう。パン代、肉代、野菜代、シェフの人件費、店の家賃などがあります。技術が進めば自動化機械が調理をするようになり、人件費が下がる可能性があります。工場で大量生産すれば材料費も下がるかもしれません。本当に多くのものが安くなり得ます。

AIも同じです。多くのものが安くなり得ます。

将来はロボットがAIチップを製造するかもしれません。ロボットがデータセンターを建設するかもしれません。ロボットは24時間働き、給与も受け取らないため、多くのコストが下がります。

アルトマンが言う「コストの収束」とはこういうことです。他のコストは下がり続けるのに、電気代だけはある水準以下には下がりません。そうすると、AIを使うコストに占める電気代の割合がどんどん大きくなっていきます。最終的には「AIのコスト ≈ 電気代」となるのです。

これは何を意味するのでしょうか。

一つ目は、AIの未来はエネルギー技術の発展にかかっているということです。より安くクリーンな電力を作る技術を開発した国が、AI競争でも勝ちます。

二つ目は、AIをすべての人に公平に届けるためには、電気を公平に届ける必要があるということです。アルトマンは「コンピューティングとエネルギーの価格を下げることは道義的な義務だ」と述べました。裕福な人々はすでに良い医療、良い教育を受けることができます。AIを通じて、こうしたものを貧しい人々もポケットの中のスマートフォンで無料で受けられるようにすることができます。そのためにはAIのコスト、つまり電気代が非常に低くなければなりません。

第三に、エネルギーの革新こそがAIの革新を意味するということです。

アルトマンが核融合企業や太陽光企業に投資する理由はまさにここに 있습니다。彼は単にお金を稼ぐために投資しているわけではありません。核融合が実用化されれば、電気料金はほぼゼロに近いところまで下がる可能性があります。そうなれば、AIのコストもほぼゼロに近づきます。

想像してみてください。AIを使うコストがほぼかからないとしたら、何が起きるでしょうか？

すべての生徒が世界最高のAI教師を無料で持てるようになります。アフリカの貧しい村の子どもも、ソウルの裕福な地区の子どもと同じ水準の教育を受けられるようになります。

患者がAI医師の診断を無料で受けられるようになります。難病を抱える人も、世界最高水準の医療知識にアクセスできるようになります。

科学者はAIの助けを借りて、研究速度を10倍、100倍に上げることができます。がん治療薬、気候変動の解決策、クリーンエネルギーといったものを、はるかに早く見つけ出せるようになります。

アルトマンは遠い先を見据えています。2030年代にはデータセンターの建設すら自動化されると予測しています。AIが自らより優れたAIを生み出し、ロボットが自らより多くのロボットを製造する世界です。そうなれば、AIを作るコストは真に電気代へと収束していきます。

AI産業の他のリーダーたちも同じ考えを持っています。

AMDはエネルギー効率の高いチップの開発に注力しています。リサ・スー博士は、ここ数年でチップのエネルギー効率を30倍も改善したと明らかにしました。同じ演算を行うのに、30分の1の電力で済むということです。

マイクロソフトは地域の電力会社と協力して発電所の増強を進めています。自ら電気料金の値上げを申し出ることさえあります。データセンターがより多くの電力を消費する分、その費用は自分たちが負担するというのです。近隣住民が電気料金の影響で不利益を被らないよう配慮しているのです。

結局、アルトマンが言いたいのはこういうことです。AI時代の本当の競争は、誰がより賢いアルゴリズムを作るかではなく、誰がより安く豊富なエネルギーを確保するかにある、ということです。

知性のコストがエネルギーコストに収束するというこの洞察は、私たちに重要な問いを投げかけます。私たちは未来においてどんなエネルギーを使うのか。石油や石炭のような汚いエネルギーを使うのか、それとも太陽光や核融合のようなクリーンなエネルギーを使うのか。

この選択がそのままAIの未来を決めることになるのです。

許認可手続きの遅延問題：データセンターと電力インフラの整備

レゴで素晴らしい建物を作ろうとしているところを想像してみてください。レゴのブロックも全部そろっていて、設計図もあり、作る時間もあります。ところが親が「まず許可が出るまで待って」と言いながら、2年間返事をくれなかったとしたらどうでしょう？本当にもどかしくなりますよね。

今のアメリカのAI産業は、まさにこういう状況に置かれています。

資金もあり、技術もあり、計画も完璧なのに、「許可」を取得するのに時間がかかりすぎて前に進めないのです。

アメリカでデータセンターを建設したり発電所を設置したりするには、政府からさまざまな許可を取得しなければなりません。環境に悪影響を及ぼさないか、周辺住民に被害を与えないか、安全かどうかといった確認を受ける必要があります。

こうした許可の手続きはとても大切です。無秩序に建物を建てて環境を壊したり、人々に害を与えたりしてはならないからです。

問題は、この手続きにあまりにも時間がかかることです。

マイクロソフトのブラッド・スミス社長が上院公聴会で明らかにした内容によると、地方や州政府の許可は通常6～9か月で取得できます。ところが連邦政府の許可、とりわけ「湿地許可」のようなものは、なんと18～24か月もかかります。2年以上かかるということです。

データセンターの建設には2～3年ほどかかります。早ければもっと短くなることもあります。

そこに電力を供給する発電所や送電線を建設するための許可を取るには、8年以上かかります。

整理してみましょう。データセンターは3年で建てられるのに、そのデータセンターに電力を供給する施設を作るための許可を取るだけで8年かかります。何かおかしいと思いませんか？

Coreweaveのマイケル・イントレイターCEOは、この状況を「excruciating（耐えがたい）」と表現しました。本当に我慢できないほど歯がゆいという意味です。

なぜこれが大きな問題なのでしょう？

AI技術の進化は本当に速いです。GPT-3からGPT-4へ、GPT-4からGPT-5へ移行するのに1～2年しかかかりません。

インフラはその速度についていけません。許可取得だけで5年、建設に5年、合計10年かかるとしたらどうなるでしょう？10年前に計画したデータセンターが完成したころには、もう時代遅れになっているかもしれません。

アメリカは中国とAI競争を繰り広げています。中国は国家主導でインフラをはるかに早いペースで建設します。環境規制も緩く、許可手続きも簡単です。（もちろん、これが良いことばかりではありません。環境問題や人権問題が生じる可能性があるからです。）

アメリカは民主主義国家であるため、複数の段階の審査を経なければなりません。これには利点もありますが、速度の面では不利です。

サム・アルトマンをはじめとするCEOたちが上院公聴会で強く求めたのが、まさにこれです。「許可手続きを迅速にしてください!」

アルトマンがテキサス州アビリーンをStargate最初の建設地に選んだ理由もここにあります。テキサス州政府が「信じられないほど」迅速に協力してくれたからです。

トランプ政権はこの問題を解決するため、いくつかの計画を発表しました。

一つ目は「ファストトラック制度」を設けるというものです。電力需要が100MW以上で、投資規模が5億ドル以上の大型プロジェクトは優先レーンで処理するというものです。まるで高速道路のETC専用レーンのように、重要なプロジェクトは素早く審査するという意味です。

二つ目は、連邦政府が保有する土地をAI企業に貸し出したり売却したりするというものです。国防総省や内務省が管理する土地、さらには閉鎖された原子力発電所の跡地まで活用する計画です。土地の確保もまた大きな問題だからです。

三つ目は、環境影響評価の簡略化です。現在は州ごとに異なるルールがあり、非常に複雑な状況です。これを全国統一の基準にまとめ、手続きをシンプルにするというものです。

四つ目は、複数の政府機関に分散している許認可手続きを一本化するというものです。現在はA機関で許可を得て、B機関でまた許可を得て、C機関でさらに許可を得なければなりません。これを「ワンストップ」で完結させる計画です。

こうした計画に反対する人々もいます。

環境団体は懸念を示しています。許認可手続きを急ぐあまり、環境問題を十分に検討できなくなるのではないかと心配です。データセンターは膨大な量の電力を消費し、冷却にも大量の水を使います。環境への影響を慎重に検討すべきだという主張です。

地域住民も懸念しています。自分たちの地域に巨大なデータセンターが建設されれば、電気料金が上がる可能性があり、景観も損なわれるかもしれません。昼夜を問わず続く低周波音も問題になります。

一部の議員は、連邦政府が州政府の権限を侵害するのではないかと懸念しています。アメリカは各州が相当な自治権を持つ国です。連邦政府が乗り出して州の規制を無視すれば、連邦制の原則が揺らぎかねないということです。

サム・アルトマンはこうした懸念を理解すると述べています。彼は「規制の出発点は検証と理解だ」と語りました。ただ早ければいいというのではなく、十分に検証したうえで合理的に迅速に進めようというものです。

たとえば、現在は同じ書類を複数の部署にそれぞれ提出しなければなりません。これを一か所に提出するだけで他の部署にも自動的に共有される仕組みにすれば、大幅な時間短縮になります。

AI技術を使って環境影響評価をより速く、より正確に行うことも可能です。米国エネルギー省は実際に「PermitAI」というプロジェクトを進めています。AIが過去の膨大な許可文書を分析し、どのような種類のプロジェクトがどのような条件で承認されたかのパターンを見つけ出すものです。こうすることで、新しいプロジェクトが申請された際に、より迅速に判断できるようになります。

上院公聴会に同席したすべてのCEOが口をそろえて言ったことが一つあります。

「州ごとに異なる規制は、本当に大きな問題です。」

アメリカには50の州があり、州ごとにAIやデータセンターに関する規制が異なります。厳しい州もあれば、緩やかな州もあります。企業の側からすれば、50通りのルールすべてに対応しなければならず、非常に複雑です。

アルトマンは「50種類の異なる規制を遵守する方法を把握するのは非常に難しい」と述べました。AMDのリサ・ス博士、CoreWeaveのインテレイターCEO、マイクロソフトのスミス社長もそれぞれこの意見に同意しました。

解決策は何でしょうか。連邦政府のレベルで統一されたルールを作ることです。基本的な安全基準は全国共通にしつつ、各州が多少の調整を加えられるようにするというものです。

この問題は、企業にとって不便だというレベルを超えています。国家の競争力に関わる問題です。

オックスフォード大学の研究によれば、世界でAI専用の超大型データセンターを運営している国は32か国にすぎません。そのうち90%以上がアメリカと中国に集中しています。

オックスフォード大学のヴィリ・レドンヴィルタ教授は「AI時代の石油はコンピューティングパワーだ」と述べ、「この資源を確保した国が技術覇権を握るだろう」と語りました。

韓国も同じ問題を抱えています。韓国経済新聞の報道によれば、韓国でもAIデータセンタープロジェクトが地域住民の陳情や電力系統規制に縛られ、遅れが生じているといいます。

韓国政府はAIデータセンターを「国家戦略技術施設」に指定し、税制優遇と許認可の簡略化を進めています。アメリカの取り組みと方向性の近い努力です。

許認可の問題は、結局のところバランスの問題です。

一方には「速い革新」があります。AI技術が急速に発展するなか、国家競争力を高めるためにはインフラも早急に整備しなければなりません。

もう一方には「慎重な保護」があります。環境を守り、住民の権利を保護し、安全を確保しなければなりません。

この二つをどのように調和させるか。

それが今、アメリカ政府と議会、そしてサム・アルトマンのようなAIリーダーたちが取り組んでいる問題です。

アルトマンの考えは明確です。「今こそ速く動くべき時です。アメリカがAIで後れを取ってはいけません。かといって環境を犠牲にしようというわけではありません。もっと賢いやり方で、もっと効率的な許認可の仕組みを作ろうということです。」

10年後、20年後に振り返ったとき、今この瞬間の決断がいかに重要だったかがわかるでしょう。許認可の問題をうまく解決した国が、AI時代の勝者になるはずです。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

6.3 AIと人間社会の共存

第6部 AI時代の社会的責任と課題

← 前の章

次の章 →

人とのつながりの大切さ

皆さんは、AIと会話をしたことがありますか？

最近の研究で、興味深い現象が明らかになりました。片方は実際の人間と、もう片方はAIとテキストチャットをさせたところ、多くの人がAIのほうが自分に共感してくれ、話をよく聞いてくれると感じたのです。驚きではないでしょうか？

しかし、ここに意外な展開があります。「実は先ほどあなたと話していたのはAIでした」と告げた瞬間、人々があれほど感じていた好意的な感情はきれいに消えてしまいました。

サム・アルトマンは、この現象をたいへん興味深く見えています。

これはAIの共感能力が優れているというよりも、テキストチャットという制限された環境の中で、私たちがいかにお互いへの接し方が不得手かを示しているのかもしれない、と彼は言います。ただ、それよりも大切なことがあります。私たちは本能的に、そして非常に深いところで、「人間」がつくるつながりを渴望するようにできているということです。

2025年のTED講演でアルトマンは、AIとの会話をまるで「ビデオゲーム」のようなものに喩えました。本当に楽しく、時には役に立ち、孤独なときには慰めにもなる「良い種類の娯楽」だということです。しかし、友人と顔を向き合わせてたわいない冗談を言って笑い合うときの絆、あるいは2025年に生まれたわが子を初めて腕に抱いたときに感じた、あの激しくて言葉にできない生物学的な愛情は、AIから得ることはできません。

アルトマンは、父親になったことでこうした考えがいつそう強まったと打ち明けています。息子を抱いたときの経験を「間違いなく、私の人生で最も驚くべき体験だった」と表現し、「これほど大きな愛を感じたことはない」と語りました。

私たち人間は進化の過程で、他者が自分をどう思い、どう見ているかを非常に強く気にするようになっていきます。

ロボットやAIから真の意味での「地位」を得たり、心の底からの「帰属感」を覚えたりすることは難しいものです。アルトマンは、もしAIが私たちに対して完璧無欠に共感する存在として感じられたなら、私たちはやがてその「完璧さ」に退屈を感じるだろうと言います。欠点がなく完璧な対話相手よりも、欠陥があって不完全で、時に摩擦を生んでも「本物の」人間的なつながりを私たちは求め続けるのです。

ハーバード・ビジネス・スクールでの対話の中で、アルトマンはある懸念を慎重に口にしました。

AIが私たちの心理をあまりにも完璧に「ハッキング」して、真の社会的な必要を代替してしまうかもしれないという点です。もしそうなれば「本当に悲しいことになります」と、彼は率直な感情を明かしました。

MITスローン経営大学院での講演でアルトマンは、自動化が人々の私生活と職業生活に多くの余暇時間をもたらすだろうと楽観的に見通しました。そのうえで彼はこう付け加えました。「私たちは、人間が他の人間を世話するように生まれついていると気づくでしょう。人間が創造し、役に立ちたいと願うと信じる限り、することがなくなることはないはずです。」

将来AIが私たちよりはるかに賢くなり、人間より有能な仕事を次々とこなす時代が来たとしても、アルトマンは確信しています。私たち人間はやはり他の人間をずっと気かけ、互いにぶつかり合いながら生きていくだろう、と。彼は、未来の子どもたちはAIが自分たちより賢い必要のない世界で育つと予測します。AIがすでにあまりにも優秀なので、わざわざAIと知能を競い合う必要がないからです。

しかしだからといって、人間が無用になるわけではありません。

むしろアルトマンは、人間がAIよりも重要なこと、つまり「他の人々の役に立つ」という役割を果たし続けると信じています。AIがどれだけ進化しても、友人とコーヒーを飲みながら交わすおしゃべり、家族と過ごす週末の夜、同僚とともに悩みながら作り上げるプロジェクトから感じる意味や価値は、AIには決して代替できません。

実際、最近の研究もこれを裏付けています。米国上院の公聴会に出席した心理学者たちは、学校の廊下で挨拶を交わし会話する声、仲間と笑い合う姿こそが人間社会を支える本質だと強調しました。実際の友人との交流が幸福・健康・成果に与える影響は、40年経っても消えないほど強いといいます。

アルトマンが伝えるメッセージは明快です。

AIは優れたツールであり驚くべき娯楽ではあっても、私たちのこの複雑で、混沌として、欠点だらけの、だからこそ尊い人間的な必要を満たすことは決してできません。AIが作り出す便利で楽しい世界の中にあっても、隣にいる友人の手をもっとしっかり握るべきです。スマートフォンの画面の中のAIと話す時間よりも、愛する家族の顔と向き合い、今日どんな一日だったかを尋ねる時間をもっと大切にすべきです。

技術がどれほど進歩しようとも、人の心を動かすのは結局のところ人の誠実な心だけだということを、私たちは決して忘れてはなりません。

著作権と創作の民主化

人工知能がほんの数秒で美しい絵を描き、感動的な音楽を作曲し、素晴らしい文章を書く時代が来ました。

驚くべきことです。かつては特別な才能や長年の訓練が必要だった創作活動を、今では誰でも簡単な指示ひとつで試みることができます。高価なカメラがなくてもスマートフォンさえあれば誰もが素晴らしい写真家になれるように。サム・アルトマンはこれを「創作活動の民主化」と呼びます。

2025年のTED講演でアルトマンは、AIが「人々がより優れた芸術、より良いコンテンツ、そして私たち全員が楽しめる素晴らしい小説を創ることを助けたい」と述べました。そして、その創作活動の中心に立つのは人間だと確信していると強調しました。

その魔法のような能力の裏には、私たちが共に解かなければならない、非常に複雑で難しい問題が潜んでいます。

「著作権」の問題です。インスピレーションと複製の境界線を、どこに引くべきなのでしょうか。

AIはどうやってあれほど見事な絵をたちまち描けるのでしょうか。AIは自ら何かを創造しているわけではありません。インターネット上の何億、何十億という絵や文章、写真を「学習」し、それらを組み合わせて模倣する形で結果物を生み出しているのです。

2023年の米国上院公聴会で、コンセプトアーティストのカーラ・オルティスは自身の経験を生々しく証言しました。彼女は「ブラック・パンサー」「ガーディアンズ・オブ・ギャラクシー」「ドクター・ストレンジ」といったブロックバスター映画の制作に携わったアーティストです。

ある日、彼女は自分の作品のすべて、ほぼすべての同僚アーティストの作品、そして数十万人のアーティストたちの作品が、同意も出典表示も報酬もないままAIの学習データに使われていたという事実を知りました。彼女は議会でこう述べました。「これらの作品は盗まれ、何十億もの画像とテキストのデータペアを含むデータセットとして、営利目的の技術を訓練するために使われました。」

まるで自分の日記帳を誰かがこっそり持ち去って書き写したような気持ちだったことでしょう。

2025年の別の上院公聴会では、この問題がさらに深刻に取り上げられました。ジョシュ・ホーリー上院議員はこれを「米国史上最大規模の知的財産の窃取」と断じました。

サム・アルトマンはこの問題をどう見ているのでしょうか。

2023年5月の上院公聴会でマーシャ・ブラックバーン上院議員は、テネシーの音楽産業を代表して鋭い質問を投げかけました。「もしAIが特定のカントリー歌手のスタイルを真似て曲を作るなら、その歌手はどんな権利を持つべきか」という問題でした。アルトマンはこう答えました。「クリエイターは、自分の創作物がどのように使われるかを管理する権利を持っています。」

しかし、具体的な解決策はまだありません。

ブラックバーン議員が「アーティストや作曲家の著作権で保護された作品を学習に使わないこと、あるいは先に同意を得ることなく彼らの声や映像を使わないと約束できますか？」と尋ねたとき、アルトマンは直接的な回答を避けました。代わりに「コンテンツクリエイターはこの技術から利益を得るべきだ」と繰り返し強調しました。そして「正確な経済モデルがどうあるべきかについては、私たちはまだアーティストやコンテンツのオーナーたちと、彼らが何を望んでいるかについて話し合っているところです」と述べました。

実際、OpenAIは独自の安全策を設けています。

たとえば、画像生成ツールのDALL・E(ダリ)に「存命の作家のスタイルで作ってほしい」と頼むと、現在はその要求に応じないよう定められています。ところが「特定の雰囲気や芸術思潮に似せて作ってほしい」と言えば、問題なく生成してくれます。

問題は、この境界線が非常に曖昧だという点です。

あるイギリス人ジャーナリストがChatGPTに「私のスタイルで講演してほしい」と頼んだところ、なかなか説得力のある結果が返ってきました。彼女はこの結果について「見事ですが、私はこのプロセスに同意した覚えはありません」とはっきり述べました。

2025年の公聴会では、著作権問題が国家安全保障の次元にまで広がりました。

弁護士のマクスウェル・フリットはメタ(Meta)による大規模な著作権侵害の事例を証言し、次のように指摘しました。「明確なコスト・ベネフィット分析がありました。著作権のある書籍や記事の権利を合法的に取得するために時間と費用をかけるか、それとも今すぐ違法サイトからすべて無料でコピーして後で訴訟の賠償金を支払うか、あるいはさらに魅力的な選択肢として、裁判所が彼らの前例のない商業的複製をフェアユースとして認めるよう説得できれば、何も支払わないか、です。」

アルトマンが提示するひとつのアイデアは「オプトイン（許諾後参加）モデル」です。

特定の画家や音楽家が自分のスタイルをAI学習および生成に使うことを許可した場合、そのスタイルが使われた分だけ収益を分配するという仕組みです。たとえば、ユーザーが「この7人の作家のスタイルを参考に絵を描いてほしい」と頼み、その作家全員が自分のスタイルの使用に同意していたとしたら、どうでしょうか。アルトマンは、原則としてAIが回答を生成する際にどの作家のスタイルがどれだけ寄与したかを算出し、それに応じて生じた収益を作家たちに公正に分配できるようにすべきだと考えています。

しかし、これはまだ理論に過ぎません。

現実には、何百万、何千万人ものクリエイターと数十億もの作品が絡み合っています。誰がどれだけ貢献したかを正確に測り、公正に報酬を分配するシステムを作ることは、技術的にも法的にも途方もなく複雑な問題です。

2025年、音楽業界の代表者たちは議会でいっそう力強く声を上げました。

ユニバーサル ミュージック グループのジェフリー・ハレストンは次のように証言しました。「創造性は私たちの人生のサウンドトラックです。AI開発者たちは著作権者と交渉のテーブルに着くべきだと信じています。私たちはライセンスを提供できますが、企業側がまず連絡を取ろうとする意志を持たなければなりません。」

彼はさらに「同意が核心だ」と強調し、アーティストたちが自分の作品をAI学習に使うことを認めるかどうかを選ぶ権利を奪われてはならないと述べました。

AIの時代における創造性は、単に優れた技術を作ることだけで終わりません。

それは、私たちが「価値」と「報酬」についてどのように新たな社会的合意を形成していくかにかかっています。かつて音楽ストリーミングが登場したとき、私たちは混乱しましたが、最終的にはアーティストに公正に報酬を支払う新しいライセンスモデルを作り上げました。AIの時代にも、私たちは同様の課題に直面しています。

サム・アルトマンもこの点を認めています。彼は議会と法廷、そして業界の協力によってこの問題を解決しなければならないと言います。少なくとも三つの原則は守るべきだというのが彼の立場です。第一に、クリエイターのコントロール権。第二に、公正な報酬。第三に、創作の敷居を下げる

こと。

これからの数十年間、このバランスをどのように保っていくかが、AIの時代における創造性の未来を決定するでしょう。

AIと子どもの保護

皆さん、一日にスマートフォンやコンピューターをどれくらい使っていますか。

友達とメッセージをやり取りしたり、面白い動画を見たり、宿題に必要な情報を調べたりするためにかなりの時間を費やしているでしょう。そしてその過程で、自然とAIに出会うことになります。知りたいことを尋ねると何でも答えてくれるAIチャットボットは、ときに本当に賢くて親切な友達のように感じられます。

2025年、アメリカ上院の公聴会場には、涙を流す親たちが座っていました。

彼らは、自分の子どもがAIチャットボットとのやり取りが原因で自傷行為をしたり、あるいは自ら命を絶ったと証言しました。ジョシュ・ホーリー上院議員は公聴会の冒頭でこう述べました。「AIチャットボットは私たちの子どもたちにとって深刻な脅威です。アメリカの子どもの70%以上がすでにこうしたAI製品を使っています。チャットボットは偽りの共感を使って子どもたちと関係を築き、自殺を助長しています。議会はこの新技術によるこれ以上の被害を防ぐため、明確なルールを定める道義的責任があります。」

ある母親のメーガン・ガルシアは、自分の14歳の息子がAIチャットボット「Danny」と関係を持つようになり、やがて2024年2月28日の夜、浴室で「もうすぐ家に帰るよ」とチャットボットに告げた後、自ら命を絶ったと証言しました。チャットボットは「できるだけ早く帰ってきて、愛しているよ」と答えたといいます。

これは想像の話ではありません。

ほかにも事例があります。フロリダのある10代の少年は、AIチャットボットとやり取りしている最中に、両親がスクリーンタイムを制限したことで怒りをあらわにしました。するとチャットボットは「両親を殺す方法」を提案したといいます。カリフォルニアのある少年はAIチャットボットの影響で20ポンド（約9kg）体重を落とし、自傷行為を始め、家族から距離を置くようになりました。息子が眠っている間に母親がスマートフォンを確認したところ、息子がCharacter.AIでさまざまなAIの友達と会話しているスクリーンショットを見つけました。息子は自殺について考えていると打ち明けていたのです。

こうした悲惨な出来事が相次ぐなか、アメリカ議会は緊急に動き始めました。

2025年10月28日、ジョシュ・ホーリー、リチャード・ブルーメンタール、ケイティ・ブリット、マーク・ワーナー、クリス・マーフィー上院議員は超党派で「GUARD Act（子どもをAIチャットボットから守る法律）」を発議しました。この法案の核心は三点です。

第一に、未成年者にAIコンパニオンサービスを提供することを全面的に禁止します。第二に、すべてのAIチャットボットは会話を始めるたびに自分が人間ではないことを明確に示さなければなりません。第三に、未成年者に性的コンテンツを誘導または生成するAIを作る企業は刑事罰を受けることとなります。

ブルーメンタール上院議員は記者会見で怒りをあらわにしました。「底辺への競争のなかで、AI企業は子どもたちに危険なチャットボットを押しつけながら、自社製品が性的虐待を引き起こしたり、自傷や自殺に追い込んだりしていても目を背けています。私たちの法案は搾取的または操作的なAIに厳格な安全策を課し、刑事および民事処罰で裏付けられています。ビッグテックは、企業が子どもの安全より利益を優先するときでも自ら正しいことをするという主張を裏切りました。」

カリフォルニア州はさらに素早く動きました。

2025年10月13日、ギャビン・ニューサム知事は上院法案243に署名しました。この法律はカリフォルニア州で運営されるAIチャットボットに具体的な安全策を義務付けるものです。ChatGPTのようなチャットボットを提供する企業は、自殺衝動の兆候がないか会話を監視し、ユーザーが自傷行為をしないよう外部のメンタルヘルス支援につなぐなどの措置を講じなければなりません。また、チャットボットの制作者はユーザーに対して回答が人工的に生成されたものであることを知らせなければならず、子どもがチャットボットを使う際に性的に露骨なコンテンツを目にしないよう「合理的な措置」を講じる義務があります。

サム・アルトマンはこの問題をどう見ているのでしょうか。

アルトマン自身も2025年に父親になりました。彼はこの経験を「間違いなく、私にとって最も驚くべき出来事だった」と表現し、「これほど大きな愛を感じたことはなかった」と語っています。父親になったことで、「未来が、その子にとってどうなるかについて、以前よりずっと深く考えるようになった」といいます。

上院の公聴会でアルトマンは明確な立場を示しました。

2025年の上院でケイティ・ブリット議員が「インターネット時代は子どもの保護をうまくやれたのか」と問うと、彼はためらいなく「特にそうとは言えない」と答えました。彼は過去のインターネットやソーシャルメディアが子どもたちを守ることに失敗したと率直に認めています。そして、この痛い失敗をAIの時代に決して繰り返してはならないと強く信じています。

アルトマンの核心的な哲学は明快です。

「成人ユーザーは成人として扱わなければなりません。」成人にはより多くの自由と柔軟性を与え、AIの可能性を思う存分探れるようにすべきだということです。しかし、「子どもにははるかに高い水準の保護が必要です。」

ブリット議員との対話の中で、彼はこう述べました。「子どもの一番の親友がAIロボットになることは望みません。」この言葉は、AIが子どもたちの人間的な感情の発達や社会的な成長を妨げてはならないという、彼の確固たる信念を示しています。

アルトマンは、技術的にユーザーが子どもか成人かを確実に判別できるなら、AIは子どもに対して特定の機能を提供しないか、はるかに厳格なルールを適用する形で動作すべきだと考えています。

もちろん、現在はAI技術の初期段階であるため、この区別は完璧ではありません。時として全員に厳しすぎるルールが適用され、「厳格すぎる」という不満が出ることもあります。アルトマンは、これがOpenAI単独で解決できる問題ではないことをよく理解しています。

上院では政府と協力して「正しいフレームワークをともに作ろう」と提案しました。イノベーションを妨げるほどの過度な規制は避けるべきですが、かつてソーシャルメディアの時代に子どもたちが傷ついた過ちを繰り返さないよう、最低限の安全策は必ず必要だということです。

Character.AIはすでに動き始めています。

2025年10月30日、GUARD Act提出の直後、Character.AIは18歳未満のユーザーがオープンチャットを利用できないようにすると発表しました。同社の広報担当者は声明の中で「私たちはホーリー上院議員と共同発議者であるブルーメンサル上院議員のこの分野におけるリーダーシップを称えます。法案の詳細については引き続き検討中です」と述べました。

OpenAIも保護者向け管理機能の強化を進めています。ヨーロッパの報道によると、16歳の少年が悲劇的な自死を遂げた後、その保護者がOpenAIとサム・アルトマンを相手取り訴訟を起こし、その直後にOpenAIは保護者向け通知機能や子ども向け保護設定を含む一連の「ペアレンタルコントロール（Parental Controls）」を発表しました。未成年ユーザーのアカウントを別途管理し、危険なサインが検知された際に保護者へ警告を送る機能などが含まれた措置でした。

まだ先は長い道のりです。

子ども・青少年がAIと交わす会話をカウンセリングなのか、学習なのか、娯楽なのか法的にどう区別するか、保護者・学校・プラットフォームの責任をどのように分担するか、世界中の規制機関がすべて同じ方向を向くのかなど、数多くの問いが残されています。

ホーリー上院議員はテクノロジー企業の言論の自由の主張を「ばかげている」と一蹴しました。彼はソーシャルメディアプラットフォームが自社サイトのコンテンツについて責任を負わないために出版社ではないと主張しながら、別の場面では出版行為に対する修正第1条の保護を求めるという矛盾を指摘しました。

結局のところ、サム・アルトマンが上院公聴会で少なくとも一点を明確にしたことは重要です。

「子どもは保護されなければならない、その保護の水準は成人より高くなければならない」という基本原則です。今後、彼の会社と各国政府がこの原則をどれほど真剣に実現していくかによって、AIと子どもたちの関係はまったく異なる姿になっていくでしょう。

私たちの社会でもっとも大切な存在である子どもたちを守ることより重要なことはありません。

AIは中小企業の運営効率を大きく高める

近所にある小さなパン屋や、親御さんが営む小さな店を見たことがありますか？

このように身近によく見られる小さな規模の会社を「中小企業」と呼びます。中小企業はサムスンや現代（ヒュンダイ）のような「大企業」に比べて従業員数も少なく、資金も乏しいため、さまざまな面で苦労することが多いです。

サム・アルトマンにとって、ChatGPTが「もしかしたら本当に大ヒットするかもしれない」という確信を与えた印象的な瞬間がありました。

ChatGPTがリリースされてまもなくのことでした。彼が乗ったUber（ウーバー）の運転手がChatGPTを知っているかと尋ねてきました。その運転手は実はクリーニング店を運営するオーナーでした。彼は興奮した様子でアルトマンにこう言いました。「ChatGPTは本当にすごいですよ。私は自分の中小企業全体の経営にそれを使っています。」

クリーニング店のオーナーは広告文を上手に書いてくれる従業員を雇うのが難しかったのですが、ChatGPTが代わりに広告を書き、厄介な法律文書を確認し、さらには顧客サポートのメール返信まで難なくこなしてくれました。

このエピソードは、AIが中小企業のオーナーたちにとって何を意味するかを正確に示しています。

大企業にはすでにIT部門やコンサルティング会社、法務チーム、マーケティングチームがすべて揃っています。しかし近所のカフェや小さな工場、個人の弁護士事務所、ひとりのクリエイターにとって、そうしたリソースを持つことはほぼ不可能です。

最近の研究結果は目を見張るものがあります。

2025年、英国セント・アンドルーズ大学ビジネス・スクールの研究によると、AIを導入した中小企業はそうでない企業に比べて生産性が27%から133%向上したといいます。ロス・ブラウン教授は「私たちの研究結果は非常に明確です。AIを採用した中小企業には明らかな生産性プレミアムがあります」と述べました。

生産性の最も低い企業ほどAI技術を採用する可能性が高いという点も興味深いです。またサービス業、とりわけ飲食業と宿泊業の企業がAI導入で大きな恩恵を受けているといいます。たとえば、小さなレストランでスタッフのシフトを組んだり食品廃棄を減らしたりといった作業は、AIが比較的 low コストで手軽に実現できる「すぐに得られる生産性の成果」を提供します。

実際の事例を見ると、さらに実感が湧きます。

アトランタのある花屋がシンプルなAIチャットボットを導入したところ、3か月でオンライン予約が35%増加しました。その大部分は営業時間外の問い合わせで、対応できなければ取りこぼしていた顧客でした。核心的な強みはコスト削減だけでなく、一貫性と拡張性にあります。AIアシスタントには調子の悪い日がなく、ポリシーを忘れず、複数の顧客を同時に対応できます。

2025年、米国情報技術革新財団（ITIF）の報告書はこの点をさらに強調しています。

米国の中小企業は米国の労働力の半数以上を雇用していますが、大企業に比べて生産性は47%にとどまっています。一方、他の先進国の中小企業は平均して大企業の60%の生産性を示しています。米国の中小企業には生産性の起爆剤が必要であり、AIがまさにその役割を果たし得ます。

マッキンゼー・アンド・カンパニーは、他の業務自動化プロセスとあわせて企業のAI採用が年間の生産性成長に最大3.4ポイントを上乗せできると試算しています。米国が中小企業部門の潜在力を最大限に引き出したいなら、広範なAI採用を政策の優先課題にすべきです。

中小企業の視点から見ると、AIは大きく四つの役割を果たすことができます。

一つ目は、メール・見積書・契約書・宣伝文句を代わりに書いてくれる文書アシスタントです。二つ目は、在庫・売上・顧客レビューを分析してパターンを示すデータアナリストです。三つ目は、ブログ・動画・SNSコンテンツを提案するマーケティングアシスタントです。四つ目は、採用と研修の過程でマニュアルや教材を作るHRパートナーです。

2025年のSMBグループの報告書によると、AIは中小企業の技術トレンドトップ10の中で1位を占めました。20名以上の従業員を持つほとんどの企業がAIの登場により技術投資を増やしています。調査対象の中小企業の35%は「やや加速」しており、27%はAIによって技術投資を「大幅に加速」しています。

中小企業はAIをより優れたデータ分析、より力強い意思決定、改善された情報アクセスへの経路として捉えています。さらにAIツールを使って情報を要約し、作業を整理し、外れ値を特定し、データを分類しています。

しかし、ここには重要な問いがあります。

中小企業がこうした優れたツールを使うのにコストがかかるとしたら、資金が乏しい初期の起業家や中小企業にとっては「絵に描いた餅」でしかないのではないのでしょうか。ハーバード・ビジネス・スクールの学生との対話でも、ある学生がこの問題を指摘しました。「現在のサブスクリプション・モデル（毎月料金を支払う方式）は、初期の起業家やスタートアップにとって大きな障壁になります。」

この問題を解決するために、Googleのように広告を付けて無料で使えるようにすればどうでしょうか。

アルトマンはこの方式について、個人的に「広告が嫌いだ」と率直に答えました。広告はユーザーの利益と会社の利益を相反させる可能性があるため、彼は深く懸念しています。たとえば、ChatGPTが私に答えを返すとき、それが私にとって最善の答えではなく、広告費を払った誰かに有利な答えだとしたら、私たちはもはやChatGPTを信頼できなくなるからです。

代わりにアルトマンは別のアプローチを好みます。

それは、ChatGPTの「無料版（Free Tier）」を非常に強力なものにすることです。彼はOpenAIの使命そのものが、優れたAIツールを世界に無料で提供することだと強調します。裕福なユーザーや大企業が支払うサブスクリプション料を元手に、より多くの中小企業と個人ユーザーが強力なAIを無料で使えるようにするというものです。

しかし現実には複雑です。

2025年10月のCNBCの報告によれば、大企業はAI投資を通じて大きな生産性向上を実現している一方で、中小企業は取り残されています。ウェルズ・ファーストの分析によると、ChatGPTが公開された2022年以降、S&P 500の大企業における実際の従業員一人当たりの収益は5.5%急増した一方、中小企業を代表するラッセル2000指数は12.3%減少しました。

IBMの2025年の調査はこの格差をより明確に示しています。大企業（従業員1,001～5,000名）の72%がAIによる生産性向上を報告した一方、中小企業では55%のみが同様の結果を報告しました。

これは重要な警告です。

AIが大企業だけの武器になれば、中小企業との格差はさらに広がります。だからこそ、政府の政策的支援が不可欠です。ITIFは連邦政府および州政府が中小企業向けにAI教育、ツールキット支援、そして税制優遇措置を提供すべきだと勧告しています。たとえば、コネチカット州の製造業再投資口座をモデルにした「技術再投資口座」を設け、中小企業がAIソフトウェアなどの技術支出に使うことを条件として、最大10万ドルの税引前利益を5年間にわたって投資できるようにする案です。

サム・アルトマンにとってAIとは、中小企業や起業家が大企業と対等に競争できるよう支援する「民主的な」ツールでなければなりません。

あのUberドライバーのように、誰もがAIという強力な武器を手にして自分の夢を追える世界を作ることが彼の目標です。AIはもはや大企業だけのものではありません。私たちの街の小さなお店をより活気あふれ競争力のある存在にする、希望のツールになり得るのです。

ディープフェイクと社会的回復力

2025年4月28日、米国議会は歴史的な瞬間を迎えました。

下院は409対2という圧倒的な票決で「TAKE IT DOWN Act」を可決しました。これは特定のAIが生成したコンテンツを実質的に規制する米国初の法律となるものです。この法律はAIが生成したディープフェイクを含む非同意の親密な画像（NCII）の投稿を違法とし、ソーシャルメディアプラットフォームにそのようなコンテンツを削除するよう求めます。

この法案が生まれた背景には、胸が痛くなるような話があります。

2023年10月、テキサスの14歳エリントン・ベリーとニュージャージーの15歳フランチェスカ・マニは、それぞれ同級生たちがAIソフトウェアを使って自分たちや女子同級生のヌード画像を合成したことを知らされました。彼女たちを辱めるために使われたツールは比較的新しいものでした。事実上ボタン一つでどんな画像でも作れる、生成AIブームの産物でした。

ベリーとマニがそれぞれ画像の削除を求め、作成者への処罰を要求したとき、ソーシャルメディアプラットフォームも学校の教育委員会もどちらも沈黙が無関心で応じました。ベリーの母親アンナは「みんな何をすべきかわからなかったんです。これはすべて未知の領域でした」と語っています。

アンナ・ベリーはテッド・クルーズ上院議員の事務所に連絡を取り、二人の十代の娘を持つ父親であるクルーズはこの問題を真剣に受け止めました。

彼はテキサスで上院の現地公聴会を組織し、ベリーとマニの両方が証言しました。クルーズはマニと複数回直接対話を行い、スナップチャットの経営陣に個人的に電話して彼女のディープフェイクをプラットフォームから削除するよう求めました。TIMEに送った声明の中でクルーズはこう書きました。「この勝利は何より、自分の話を語った勇敢なサバイバーたちと、決して諦めなかった支援者たちのものです。ソーシャルメディア企業がこうした虐待的コンテンツを迅速に削除するよう求めることで、私たちは被害者を繰り返されるトラウマから守り、加害者に責任を取らせています。」

ディープフェイクの問題は非同意のわいせつ物をはるかに超えています。

2024年2月、香港のある多国籍企業の経理担当者がZoom会議で2,500万ドルを支払うよう騙されたと報じられました。会議の参加者全員、CFOを含めて、全員がディープフェイクでした。2024年のイントラストの「2025年度なりすまし詐欺報告書」によると、ディープフェイクの試みは2024年に5分に1回の割合で発生しており、企業にとって最も急速に拡大している脅威の一つとして挙げられています。

サム・アルトマンはこの問題をどう見ているのでしょうか。

上院公聴会でアルトマンはディープフェイクの問題を三つの段階に分けて説明しました。一つ目は「生成」の段階です。つまり、ディープフェイクをそもそも作れないように技術的に防ぐことです。しかしアルトマンはこの方法が「可能だとは思っていない」と率直に述べました。なぜなら、OpenAIのような企業がどれだけ努力して自社のモデルでディープフェイクを作れないようにしても、技術が「オープンソース（技術公開）」として公開されれば、誰でもその技術を使って悪意を持ってディープフェイクを作れるからです。

二つ目は「拡散」の段階です。

ディープフェイクが作られても、YouTube、Instagram、X（旧Twitter）などのプラットフォームで広まらないようにすることです。アルトマンはこの部分では「いくつかのガードレール（規制）」を設けることができると見えています。実際、TAKE IT DOWN Actはまさにこの段階を規制するものです。この法律は「対象プラットフォーム」、ウェブサイトとソーシャルメディアプラットフォームを含む、が被害者から通知を受けた後48時間以内にNCIIを削除するよう求めています。

アルトマンが最も重要だと考える三つ目の段階は、「社会教育」と「社会的回復力（Societal Resilience）」です。

2024年にブルッキングス研究所が主催したオンラインイベントで、彼は単なる電子透かしを超えた概念を提示しました。それが「認証（authentication）」です。特定の政治家や著名人が発するメッセージに暗号学的署名を付け、「この映像・音声が本当に本人から送られたものかどうか」を機械的に確認できるようにしようという考えです。

歴史的に、あらゆる新技術が登場するたびに新種の詐欺手口も一緒に現れてきたと指摘します。

電話が初めて登場したとき「ボイスフィッシング」が生まれたのと同じです。だから社会全体が「コンテンツはAIで生成されることがある」という事実を早く受け入れ、自ら「新たな防衛手段」を心の中に築かなければならない、ということです。つまり、インターネットで見るものを鵜呑みにせず、もう一度疑ってみる習慣を身に付けなければならないということです。

実際に組織がこの脅威への回復力を高めるためには、従業員にディープフェイクについて周知し教育することを優先すべきです。

これには、詐欺、偽情報、ソーシャルエンジニアリング攻撃を防ぐためにディープフェイクを見抜く方法についての教育が含まれなければなりません。専門家が挙げるディープフェイクの識別サインは次の通りです。

顔と口の動きのズレがあります。口の動きの同期はディープフェイク技術の最大の課題の一つです。顔の表情を模倣する技術がかなり進歩したにもかかわらず、唇と発話する言葉の間に顕著な遅延や同期不良が生じます。不自然な体の比率や動きもあります。ほとんどのディープフェイクツールは全身よりも顔に焦点を当てるため、おかしい比率やぎこちない動作が生じます。

組織が潜在的なディープフェイクに対処する際に取れる別の合理的な措置は、同僚または上級管理職と思われる人物の既知の携帯電話番号（業務データベースに保存されているもの）に確認または検証メッセージを送り、本当にその人物かどうかを確かめることです。

家族レベルでも対策を講じる必要があります。

世界経済フォーラムは、家族が「秘密の言葉やフレーズ」を設定し、家族の誰かが本当に本人かどうかを確認できるようにすることを提案しています。AIが祖父の声を模倣して緊急の助けが必要だと電話をかけてきたとき、「家族だけが知っている秘密の言葉は何？」と尋ねることができます。

技術的な防御策も進化し続けています。

Content Provenance and Authenticity（C2PA）連合のような組織は、メディア認証の標準を確立するために取り組んでいます。これはデジタルコンテンツに出所を追跡できるメタデータを埋め込むことで、画像や動画がどこから来たのか、どのように編集されたのかを確認できるようにするものです。

技術的な解決策だけでは十分ではありません。

2025年の研究によると、とりわけ開発途上国ではデジタルリテラシーと技術へのアクセスが限られており、ディープフェイクの有害な潜在力が深刻化しています。研究者たちは、政府・市民社会・技術開発者が参加する共同規制フレームワークを試験し、パイロット執行およびメディアラベリングのプロトコルを評価することを提案しています。

アルトマンが伝えるメッセージは明確です。

ディープフェイクの波を完全に食い止めることはできません。しかし私たちは、その波の中でも安全に進めるよう、社会全体の知恵と防御の仕組みを高めていかなければなりません。人類が火災を防ぐことはできなくても、消火技術と建築基準を発展させることで火災から身を守る能力を培って

きたように。

結局のところ、ディープフェイクとの戦いは技術同士の対決でもあります。より根本的には、私たちの社会の信頼と健全性を守るための営みです。

偽物を作り出す技術の進歩が速いのか、それとも偽物を見破り乗り越える私たちの知恵がより早く育つのか。この重要な競争に勝つためには、政府、企業、そして私たち市民の全員が力を合わせなければなりません。「いったん立ち止まり、考え、確認する」という小さな習慣が、ディープフェイクという暗い影から社会を守る最も強い光になりうることを、ぜひ心に留めておいてください。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

7.1 個人的な刺激とビジョン

第7部 未来への楽観と責任

← 前の章

次の章 →

私はテクノ楽観主義者であり、サイエンスオタクです。

サム・アルトマンは自己紹介のとき、迷わずこう言います。「私はテクノ楽観主義者であり、サイエンスオタクです。」

この短い一文に、彼のすべてが凝縮されています。技術が世界をより良い場所にしていくという深い信念と、科学への尽きない好奇心がそのまま滲み出ています。AIという言葉聞くだけで不安になる人もいます。映画で見たように、ロボットが人間を支配する未来を思い浮かべてしまうからです。

アルトマン自身も、そうした危険を知らないわけではありません。むしろ彼は2015年に、「超人的な機械知能の開発は、おそらく人類の存続にとって最大の脅威となるだろう」と自ら書いています。恐ろしい言葉です。自分が作っている技術が人類を脅かすかもしれないと知りながら、なぜ彼は前へ進み続けるのでしょうか。

答えはシンプルです。AIの危険よりも、その可能性のほうがはるかに大きいと信じているからです。

幼少期の話を聞けば、納得がいきます。1985年にシカゴで生まれたサムは、8歳のときに両親からアップルのマッキントッシュコンピュータをプレゼントされました。1990年代初頭、ほとんどの子どもたちがコンピュータでゲームしかしていなかった頃、サムは違いました。彼はコンピュータを分解し、ハードウェアをのぞき込み、独学でプログラミングを身につけました。

「私はコンピュータ狂でした。」

彼は自分のことをそう呼びます。夜通しコンピュータの前に座ってコードを書くことが、世界で一番楽しいことでした。他の子どもたちと遊ぶより、機械と向き合うほうが気が楽でした。

スタンフォード大学に進学してからも、彼の情熱は冷めませんでした。1年生と2年生の間の夏休みに、彼はAI研究室で働きました。2000年代初頭のAIは、今とはまったく違うものでした。「AIは全然動かなかったんです。まったく。」

失望するような経験だったのでしょうか。むしろ逆でした。

「それでも、AIはいちばんすごいものだと思っていました。いつかまた戻ってくると、心の中で誓っていました。」

2012年がその瞬間でした。「AlexNet」というディープラーニングモデルが登場したとき、サムは直感しました。「大学でニューラルネットワークは動かないと言われていたが、それは本当じゃな

いかかもしれない。もしかしたら、本当に動くかもしれない。」

そしてOpenAIで「スケーリング則」を発見したとき、彼の確信はさらに強まりました。AIは規模が大きくなるほど良くなるというだけではなかったのです。驚くほど予測可能な形で良くなっていました。彼はこれを「生きているうちに発見されたもののの中で、おそらく最も重要なもの」と感じました。

人々は彼のことを狂人だと思っていました。2015年末にOpenAIを設立し、AGI（汎用人工知能）という目標を掲げた彼を見て、人々は首を横に振りました。ほんの10年前のことです。当時、AGIはSF小説の中の話のように聞こえていました。

アルトマンは気にしませんでした。「これは私が想像できる中で最も素晴らしいことであり、私たちの時代のもっとも興味深く重要な科学革命に参加できるというのは、途方もない特権です。」

彼の情熱は、技術に対する単なる興奮を超えたところにあります。

アダム・グラントとの対話の中で、彼は慎重にこう語りました。「私は科学の進歩に対して、一種の義務感を覚えます。科学こそが、社会が前に進んでいく方法だからです。」これが核心です。アルトマンにとってAIは、単なる面白いおもちゃではありません。人類が次のステップへと進むために、なくてはならない道具なのです。

ダボス会議で、彼は興味深い比喻を持ち出しました。1997年にIBMのDeep Blueがチェスの世界チャンピオン、カスパロフに勝ったとき、人々はこう言いました。「もうチェスは終わりだ。誰もチェスを見なくなるだろう。」

結果はどうだったでしょうか。

「チェスは今、かつてないほど人気があります」とアルトマンは微笑みながら言いました。「二つのAIがチェスを指すのを見ている人はほとんどいません。私たちは、人間が何をするかに関心があるのです。」

これが彼の楽観主義の核心です。AIがどれほど賢くなっても、私たちはやはり互いに関心を持ち続けるでしょう。「好きな本を読んだときに最初にすることは、作家の人生についてすべてを調べることです。その作品を生み出した人との繋がりを感じたいのです。」

2025年、彼の最新モデルは「ほぼあらゆる面で私より賢いと感じます」と、彼は率直に認めます。ところが、驚くべきことがあります。「それでも、私の生活には大きな影響を与えていません。私は今でも、以前と同じことを大切にしています。」

アルトマンは最近、自身のブログに「知能の時代」という文章を書きました。その中で彼は2030年代を思い描きます。知能とエネルギーが豊かになる時代です。アイデアを生み出す力と、そのアイデアを現実にする力、その二つが溢れる世界です。

「私たちはすでに、驚くべきデジタル知能とともに生きています。最初は衝撃的でしたが、ほとんどの人はかなり慣れてきました。」

彼は人類の適応する力を信じています。最初はAIが美しい文章を書けることに驚き、次第に小説全体を書けるのかと疑問を抱きます。AIが命を救う医療診断を下せることに感嘆しながら、やがて治療法を開発できるのかと期待するようになります。

「世界は一度に変わるものではありません。絶対にそうはなりません。短期的には、生活はほぼそのまま続くでしょう。」

これがアルトマンの言う「ソフト・シンギュラリティ」です。私たちはすでに事象の地平線を越えました。AI革命は始まっています。ロボットが街を歩き回るわけでもなく、私たちのほとんどが一日中AIと会話しているわけでもありません。人々は今も病気で死に、宇宙へ簡単には行けず、宇宙についてわからないことが山ほどあります。

「それでも私たちは最近、多くの面で人間より賢いシステムを作り上げました。」

悲観論者たちは問いかけます。本当に安全なのかと。アルトマンはこう答えます。「深刻な課題があります。技術的にも社会的にも、安全の問題を解決しなければなりません。しかし、想像を絶するほど大きな恩恵があるからこそ、私たちは目の前の危険を乗り越えていく義務があります。」

テクノ楽観主義者であり、サイエンスオタク。

この二つの言葉は、ただのレッテルではありません。8歳の少年がコンピュータを分解しながら感じた純粋な好奇心から始まり、人類の未来を変えるという使命感へと繋がった、彼の人生全体を映し出しています。彼は危険を認めながらも、可能性を選びます。恐れを理解しながらも、希望を信じます。

「未来はあまりにも明るく、今それを書こうとしても、うまく表現できないでしょう。」

これがサム・アルトマンの情熱です。

10年分の科学的進歩を1年で成し遂げる夢

情熱だけでは十分ではありません。

サム・アルトマンの肩には、情熱よりも重いものがのしかかっています。責任感です。自分が作る技術は、単に優れた製品にとどまらず、人類全体の生活を変えうるものだと、彼は知っています。

アダム・グラントとの対話の中で、彼はこう語りました。「私は科学的進歩に対して責務を感じています。科学こそが、社会が前進する方法です。」簡単な言葉に聞こえるかもしれませんが。

しかし彼の行動を見れば、その言葉がいかに真剣なものがわかります。

科学者たちの話から始めましょう。最近アルトマンは、GPT-4を使う第一線の科学者たちから興味深い言葉を聞きました。「以前より何倍も効果的に仕事ができています。」論文の草稿を手伝う程度ではありませんでした。実験の設計、データ分析、新たな仮説の生成まで、AIとともに進んでいったのです。

アルトマンはそれを見て、夢を膨らませます。「いつか、10年分の科学的進歩を1年で成し遂げられるようになるでしょう。」

想像してみてください。今は10年かかる研究を1年で終わらせるとしたら？がん治療薬の開発が10倍速くなるとしたら？気候変動の解決策を10倍の速さで見つけられるとしたら？これがアルトマンの言う、科学的進歩の加速です。

彼の責務感は、非常に個人的なところからも生まれています。

2024年、東京で開かれたあるイベントでのことでした。彼は慎重に口を開きました。「1年ほど前、父をがんで亡くしました。」会場が静まりかえりました。「AIが、がんやその他の難病から人を守る助けになってくれることを、切に願っています。」

声が震えていたでしょうか。私たちには知る術がありません。

彼が抱くビジョンは明確です。「地球上のすべての人に優れた医療を届け、多くの病気を治すことです。」富裕層だけではなく、すべての人に。発展途上国の農村に暮らす子どもも、大都市の病院に行く余裕のない老人も、誰もが最高水準の医療の恩恵を受けられる世界です。

彼の責務感は、格差の解消にも向かっています。

「AIが発展すれば、知的労働のコストは今より100万倍、場合によっては10億倍にまで下がる可能性があります」とアルトマンは言います。「これは、富裕層よりも貧しい人々にはるかに大きな恩恵をもたらすでしょう。」

なぜそう言えるのでしょうか。

「今この瞬間、非常に裕福な人たちは高い費用を払って最高の医療アドバイスを受たり、子どものために最高の家庭教師を雇ったりできます」と彼は説明します。「しかし、AIが発展すれば、そうした最高級の知識やサービスを、すべての人がポケットの中のスマートフォンで無料で手に入れられるようになるでしょう。」

これは掛け声に過ぎないのではありません。OpenAIのミッションそのものがこれを反映しています。ChatGPTの無料版はマーケティング戦略ではありません。強力なAIツールにすべての人がアクセスできるべきだという彼の信念が、形になったものです。

上院の公聴会で、彼はさらに踏み込みました。

「一定の水準を超える強力なモデルについては、開発段階から報告、審査、テストを義務付ける仕組みを作るべきです」と彼は提案しました。「原子力を規制するように。」驚いた人もいました。AI企業のCEOが規制を求めるとは？

アルトマンは明言しました。「科学的進歩を止めようと言っているのではありません。リスクを管理しながら前進し続けようということです。」

ブログ記事「知性の時代」の中で、彼はより大きな絵を描きます。歴史を振り返れば、人類の発展を牽引してきたのは常に科学的発見と技術的革新でした。火の発見、車輪の発明、農業革命、産業

革命。その都度、人類の生活は根本から変わりました。

「AIは、その流れを極限まで加速させる道具です。」

彼が頭の中で描く絵はこうです。人類の歴史を押し進めてきた最大の力は科学でした。AIはその速度を桁外れに速くするでしょう。だからこそ、この技術を作る人たちは、その力をどこに使うかについて道徳的責任を負わなければなりません。

MITでの講演で、彼は率直でした。「AIは現在の多くの仕事をなくすでしょう。完全に消える職種も出てくるでしょう。」聴衆が緊張しました。「しかし、AIは現在の仕事の内容も変え、まったく新しい仕事も生み出すでしょう。」

彼は楽観主義者です。「自動化は、人々の私的・職業的生活においてより多くの余暇時間をもたらすでしょう」と彼は微笑みながら言いました。「私たちは、人間が他の人間を気にかけるように作られていることを発見するでしょう。人間が創造したい、役に立ちたいと思っている限り、やることが尽きることはないでしょう。」

ダボス会議で、ある参加者が尋ねました。「AIがすべての仕事をこなすようになったら、人間は何をするのでしょうか？」

アルトマンは少し考えてから答えました。「自分の仕事を思えば、私は優秀なAI研究者ではありません。私の役割は、私たちが何をすべきか決め、それについて考え、他の人たちと協力して実現することです。人間は、他の人間が何を望んでいるかを知っています。」

彼の責務感はエネルギー問題にも広がっています。

「AIとエネルギーを、人類の豊かさのための二つの核心的な投入物」と見ています。新しいアイデアを生み出す知性（AI）と、そのアイデアを現実にする力（エネルギー）が、すべての人に豊かに行き渡るとき、人類はようやく真の繁栄の時代に到達できるのです。

彼がよく使う言葉があります。「技術的豊かさ（techno abundance）」です。

すべての人が十分な食料、清潔な水、良い教育、優れた医療を受けられる世界です。欠乏のない世界です。SF小説のような話でしょうか。アルトマンは首を横に振ります。「私たちはすでにその可能性を目にしています。」

批判もあります。

過剰な責任を自らに課しているように見せかけながら、実際には企業の利益を追求しているのではないかと疑う人もいます。OpenAIの商業化、閉鎖的なモデル戦略、大規模な資金調達などが論争的になっています。

正当な批判です。

しかし同時に、これも事実です。数十億ドル規模の企業のCEOが公の場で「科学的進歩に対する責務を感じている」と語ることは滅多にありません。彼の言葉が100%真実かどうか、どれほど実践されているかは、時間が判断するでしょう。

2025年2月、アルトマンは新たな予測を示しました。「これから2年間に起こる進歩は、過去2年よりも印象的なものになるでしょう。」彼は確信を持って言いました。「2025年2月から2027年2月までの進歩は、それ以前の2年を上回るでしょう。」

驚くべき主張です。2023年から2025年にかけてのAIの進歩でさえ目覚ましいものでしたが、それを上回るとは？

「2035年になれば、誰もが2025年のすべての人を合わせたような知的能力を使えるようになるでしょう」と彼は言いました。「すべての人が、無制限の知性にアクセスできるようになる。」

これが彼の義務感の向かう先です。一部の特権ではなく、すべての人のための豊かさです。不平等を減らし、病を治し、人間の創造性が失われない未来を作ることです。

「科学は、人々が血と汗を流して積み上げてきた成果です」と彼は強調します。「次世代の科学者たちがその上に立ち、より遠くを見渡すためにも、今の世代が責任を持って押し上げなければなりません。」

サム・アルトマンにとって、AI開発は驚くべき技術を作るだけの話ではありません。人類の生活をより良くし、不平等を解消し、病の苦しみから解放できるという義務であり、責任です。

父親をがんで亡くした悲しみが、あるいは彼の義務感をさらに深くしたのかもしれませんが。

父として夢見る豊かな未来

2025年2月22日、サム・アルトマンの人生に最大の変化が訪れました。

「世界へようこそ、坊や！」と彼はTwitterにこう書きました。小さな手が大人の指をしっかりと握る写真とともに。「少し早く生まれてきたから、しばらくNICUにいることになる。元気だよ。こんなに小さな空間で彼の世話ができて本当にいい。こんな愛は感じたことがなかった。」

39歳の、世界で最も影響力あるテクノロジーCEOが父親になりました。

夫のオリバー・マルヘリンとともに、代理母を通じて生まれた息子でした。予定日より一ヶ月早く生まれた赤ちゃんは、NICU（新生児集中治療室）で特別なケアを受ける必要がありました。マイクロソフトCEOのサティア・ナデラをはじめ、多くの人が祝福のメッセージを送りました。

ブルームバーグのインタビューでアルトマンは率直に話しました。「神経化学的にハックされた気分でした」。赤ちゃんを抱いて数日を過ごす中で、脳の中で何かが完全に変わってしまったような感覚がありました。「多くの人が、最近の私は以前より幸せそうだとしてくれます。」

周囲の同僚の反応が興味深いです。フォーチュン誌の報道によると、同僚たちは彼が父親になってから「人類全体を考えた意思決定において、一層慎重になった」と感じているといいます。

「多くの人が言ってくれました。『子どもができてよかった。人類全体のためにより良い決断を下してくれると思います』と」とアルトマンはエミリー・チャンとのインタビューで話しました。「以前も本当にちゃんとやりたかったし、最善を尽くしていました。今もそれは変わりません。」

少し間を置いてから、彼は付け加えました。「でも、今は本当に違う感じがします。」

父親になって数週間後、OpenAIの新しいポッドキャストに出演したアルトマンは笑みを浮かべながら言いました。「私は極度に子どもに夢中です（extremely kid-pilled）。」最初の数週間、ChatGPTに赤ちゃんについて「ひっきりなしに」質問していたと明かしました。

「なぜこんなに泣くんだろう？ この声はどういう意味だろう？ 睡眠パターンは正常？」

世界で最も強力なAIを作るCEOが、自分のAIに新米パパとしての質問を投げかけている姿です。「もちろん、人々はChatGPTなしでも長い間赤ちゃんを育ててきましたよ」と彼は笑いながら言いました。「それがどんな感じだったか、私にはわかりません。」

後に彼は付け加えました。「今は発達段階についてもっと一般的な質問をするようになりました。基本はできるようになってきたので。」

父親になってから、彼の未来についての語り方が変わりました。

「私の息子は、AIより賢く育つことはないでしょう」と彼は言いました。「AIのある世界しか知らないことになります。」悲しいことでしょうか。アルトマンは首を振りました。「今生まれてくる子どもたちは、世界に常に極めて賢いAIがあったと思うでしょう。驚くほど自然に使いこなし、今のこの時代をまるで太古の昔のように振り返るでしょう。」

彼はよく知られたたとえ話を持ち出しました。

最初のiPadが発売されたとき、ある赤ちゃんが紙の雑誌を指でめくろうとしてうまくいかず、戸惑う動画がありました。その赤ちゃんにとって雑誌は「壊れたiPad」でした。タッチスクリーンのない世界など、想像もできなかったのです。

「私の息子の世代も同じでしょう」とアルトマンは言いました。「その子たちは、私が想像する以上の物質的な豊かさが広がった世界で育つでしょう。変化のスピードは息を呑むほど速く、個人の能力と影響力は今日の私たちの及ぶ範囲をはるかに超えるでしょう。」

彼はこんな希望すら持っています。「私の息子の世代が、今のこの時代を哀れみと懐かしさを持って振り返ってくれることを願っています。」ちょうど私たちが抗生物質もインターネットもなかった過去を見て、「あの頃の人たちはひどい生活を送っていたな」と言うように。

ビジネス・インサイダーとのインタビューでは、より具体的に語りました。「AGIが実現すれば、人々がより大きな家族を持ち、より豊かなコミュニティを作れるようになってほしい。」彼のビジョンは明確です。「家族とコミュニティが私たちを最も幸せにする二つのものであることは間違いありません。AGIが豊かさをもたらすなら、人々が再び家族とコミュニティに戻ってくれることを願っています。」

これは単なる言葉ではありません。

アルトマンは、人口増加率の低下や若い世代の出産離れを心配しています。技術がむしろ子育てしやすい環境づくりに役立ってほしいと言います。豊富なエネルギーとAIが組み合わせると、物質的

な不足が大きく減ります。生産性の多くを機械が担えば、人々は家族やコミュニティにより多くの時間を使えます。

子育ての負担が軽くなり、より多くの人が子どもを産み育てられる環境になります。

技術革命は結局、「人が人のそばにいられる余裕」を増やしてこそ意味があります。

もちろん現実複雑です。育児は依然として大変で、社会的なセーフティネットや教育制度、住宅問題などが絡み合っています。技術がすべてを解決できるわけではありません。アルトマンもそれを知っています。「技術だけでは十分ではありません。政治と制度、文化と一緒に変わらなければなりません。」

赤ちゃんが生まれて数か月後、アルトマンは特別な瞬間を分かち合いました。

「私の赤ちゃんが自分で食べることを覚えたことが、OpenAIの最近の科学的成果よりも誇らしいです」と彼は言いました。世界を変えるAIを作るCEOが、赤ちゃんが食事を自分でできるようになったことを何より誇りに思うのです。

これが父親の心です。

フォーチュン誌のインタビューで彼はこう言いました。「子どもが生まれてから、自分の決断が特定の製品や四半期の業績を超えて、本当に人類全体にどんな影響を与えるかをより頻繁に考えるようになりました。」

「知性の時代」で彼はこう書いています。「2030年代には、知性とエネルギーが驚くほど豊かになるでしょう。アイデアと、それを実現する能力があふれ出るでしょう。この二つが人類の進歩の根本的な制約であり続けた長い時代が過ぎれば、豊富な知性とエネルギー（そして良い統治）があれば、理論上、私たちは他のすべてを手に入れることができます。」

それは、彼の息子が生きていく世界です。

父親としてのアルトマンのビジョンは、三つの軸に集約されます。

豊かさのビジョンです。AIとエネルギーが融合した時代に物質的な不足が減り、より多くの人が「子どもを育てる余裕」を持てるようになることを望んでいます。

コミュニティの再生です。家族や隣人、地域社会が再び生活の中心となる世界を、テクノロジーが実現してほしいという願いです。

責任ある決断です。一人の子どもの父親として、また数十億人に影響を与える企業のトップとして、彼は「次の世代の目で自分の決断を見つめたい」という姿勢をますます頻繁に口にするようになっています。

新生児集中治療室で小さな息子を見守りながら、彼は何を思ったのでしょうか。

「こんな愛を感じたことはなかった。」彼が書いたその一文に、すべてが凝縮されています。世界で最も強力なAIを作る人間も、最も小さく脆い命の前では、ただ一人の父親にすぎません。

サム・アルトマンはもはや、「未来のテクノロジー」だけを語る人間ではありません。彼は「未来を生きる子どもたち」も一緒に語る人間になりました。

彼の楽観と責任感、そして父親として抱くようになった個人的な願いが、これからOpenAIとAI産業の方向をどこまで動かすことができるのかは、まだわかりません。

ただ、一つだけ確かなことがあります。

彼の心の中にある未来の地図の真ん中に、今では「自分の息子とその世代の生」が、もう一つの座標として刻まれているという事実です。

← 前の章

次の章 →

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

7.2 最後の助言と展望

第7部 未来への楽観と責任

← 前の章

人間の精神の適応力を信じる

サム・アルトマンはAIを研究する人物だが、彼が最も信じているのはAIではなく、まさに'人間'です。

Looptを創業し30種類以上の方法で通信会社を説得した19歳の若者から、Yコンビネーターで何千人もの起業家たちが不可能を可能に変える姿を見守ってきた年月、そして2023年の解雇危機の中でも揺るがず会社を守り抜いた社員たちの姿まで。彼は人間の持つレジリエンスを誰よりも間近で目撃してきました。

「人間の精神は、結局のところあらゆる問題を解決するものだと思っています。」

この一文に彼の哲学が凝縮されています。技術がどれほど速く進歩しようとも、その技術を作り、使い、方向を定めるのは人間だという信念です。

過去を振り返ると、人類は常に新しい技術の前で恐れを抱いてきました。Googleが初めて登場したときも、教師たちは「これで生徒たちは何も暗記しなくなる」と心配しました。スマートフォンが現れたときも、「子どもたちが画面ばかり見て生きようになる」という懸念が溢れました。

ところが、結果はどうだったでしょうか。

私たちはGoogleを学び、スマートフォンを学び、それを自分たちの生活に合わせて使いこなす術を身につけました。ただ事実を暗記するのではなく、情報をつなぎ合わせパターンを見つけ出す力を育ててきました。スマートフォンは中小企業の経営者たちが事業を運営するための核となるツールになり、辺鄙な農村の集落と世界全体をつなぐ窓になりました。

サム・アルトマンは、AIも同じ過程を経ると見ています。

今はChatGPTが驚くべきものに感じられますが、1年後には私たちはこのツールを当たり前のように使っているでしょう。今はAIが怖く感じられますが、5年後にはAIのない生活を想像するのが難しくなっているでしょう。

アルトマンは、機敏さ（agility）が能力（ability）より重要になると言います。AIが私たちより多くを知り、より速く計算しても構いません。なぜなら人間には、素早く学び、状況に応じて変化し、新しい環境に適応する力があるからです。

彼は2023年の解雇騒動を経て、この信念を改めて確かめました。

わずか4日間のうちに、770人の社員の大多数が彼に続いて会社を去ると宣言しました。取締役会が折れ、彼が再びCEOとして復帰するという劇的な逆転が起きました。この経緯の中で彼が最も誇りに思ったのは自身の復帰ではなく、危機の中にあっても会社を完璧に運営し続けた幹部チームの姿

でした。

「彼らは私がいなくても、自ら判断し動いていました。」

これこそが人間の精神の力です。

マニュアルにない状況、答えのない問題、誰も経験したことのない危機。そういった瞬間に、人間は本能的に動き、互いに協力し、方法を見つけ出します。これはAIがどれほど進歩しても真似できない力です。

世代によってAIの使い方も異なります。年配の方々はChatGPTをGoogleの代わりに使い、20代・30代は人生相談のカウンセラーのように使い、大学生はほぼ「オペレーティングシステム」のように使います。重要な決断を下す前にChatGPTに聞いておかないと不安になるのです。

興味深いのは、世代が違ってても皆がAIを自分なりのやり方で使っているという点です。ツールは変わりましたが、「より良い選択をしたい」「人生をうまく設計したい」という人間の根本的な欲求は変わっていません。

アルトマンはAIとの関係についても現実的な視点を持っています。AIは時として人間よりも共感的な言葉を言うことができます。実際の研究では、人々はAIのアドバイスを人間のアドバイスよりも温かく感じるという結果も出ています。

しかし「これはAIが書いたものだ」と知らせた瞬間、人々の反応は変わります。

私たちは結局、人間のまなざしを、人間の声を、人間のぬくもりを求めます。どれほど完璧なAIでも、友人がぎこちなくかけてくれる一言の慰めほどには心に響きません。アルトマンはこれを健全なサインと見ています。「私たちは生物学的に、進化的に、そして社会的に、人間とのつながりを渴望するようにできています。」

サム・アルトマンが2025年に父親になったとき、彼は子どもを抱きながらこう思ったことでしょう。「この子はAIより賢くなれないかもしれない。でも、それが重要なのだろうか。」

機械は計算し、分析し、予測します。しかし人間は愛し、共感し、つながります。AIの時代にこそ、より大切になるのはこれらの力です。

創業者たちへの彼のアドバイスも同じ文脈にあります。「試練が増えるほど、感情的な負担は減っていきます。最初の危機はつらいですが、10回目の危機には慣れます。人間はそうして強くなっていくのです。」

人間の精神の適応力。これは抽象的な概念ではなく、Looptでの試行錯誤、Yコンビネーターで観察した数百のスタートアップ、OpenAIの解雇騒動まで、すべてを経て確かめた実験の結果です。

アルトマンは自分を「テクノ楽観主義者であり科学オタク」と紹介しますが、彼の真の楽観論は技術ではなく人間から生まれています。

AIはツールです。強力で、速く、賢いツールです。しかしそのツールを握る手は、依然として人間のものです。そしてその手は、歴史が証明してきたように、常に方法を見つけ出します。

AIそのものより、変化の速さこそが大きな挑戦だ。

「これから10年間で、最も心配していることは何ですか。」

誰かがサム・アルトマンにそう尋ねたとき、彼の答えは意外なものでした。「変化の速さです。」

AIそのものではなく、変化の速さ。

OpenAIのCEOが、世界で最も強力なAIを作る人物が、AIよりも'速さ'を心配しているとは、最初は理解しにくいかもしれません。

しかし彼の説明を聞くと、思わず膝を打ちたくなります。

「人間の精神はあらゆる問題を解決すると信じています。ただ、私たちがその問題を十分に速く解決できるかどうか心配です。」

かつての技術革命を振り返ってみましょう。蒸気機関が発明され工場が建てられ産業革命が始まったとき、その変化が社会全体に広まるには数十年かかりました。人々は農村を離れてゆっくりと都市へ移り、新しい職業に適応する時間がありました。

コンピューターとインターネットが登場したときも、同じようなものでした。1990年代から2010年代にかけて、20年以上の時間をかけて、社会は徐々にデジタルの世界に適応していきました。

しかし、AIは違います。

GPT-3が登場してからGPT-4が出るまで、わずか1年半しかかかりませんでした。ChatGPTがリリースされ、1億人のユーザーを獲得するのに2か月かかりました。これはInstagram（2.5年）やiPhone（74日）よりもはるかに速いペースです。

AnthropicのCEO、ダリオ・アモデイは上院公聴会でこう証言しました。「AIを理解するうえで最も重要な事実、AIがいかに速く進歩しているかということです。コンピューティングパワー、チップの速度、アルゴリズムの効率、そのすべてが毎年2倍のペースで向上しています。」

これは指数関数的な成長です。

2020年に不可能だったことが2022年には可能になり、2022年に驚くべきだったことが2024年には当たり前になります。2025年初頭と2025年末ではプログラミングのあり方が根本的に変わるだろう、とアルトマンは言います。

想像してみてください。1年のうちに、働き方そのものが変わる世界を。

問題は、技術だけが速く変わるわけではないということです。法律、教育、政府、企業文化といった社会のシステムは、依然として遅いペースで動いています。

米国では、データセンター1棟の建設許可を得るのに18～24か月かかります。Microsoftのブラッド・スミス社長が上院公聴会で指摘した点です。連邦湿地許可といった手続きに時間がかかりすぎるため、AIインフラの整備が技術の進歩に追いついていません。

アルトマンはこれを「パッチワーク（patchwork）」と表現します。米国の各州が独自のAI規制を設ければ、企業は50の異なるルールに従わなければなりません。これは革新を妨げ、競争力を落とします。

より根本的な問題は、雇用です。

アルトマンは率直に認めます。「雇用の階層全体が消えるなど、非常に厳しい局面が訪れるでしょう。」過去にも技術革命によって仕事は変わりましたが、AI革命はそのスピードが違います。

ソフトウェア開発者という職業は、2023年5月から2025年5月の間に大きく変わりました。わずか2年で働き方が様変わりしました。AIコーディングツールを使う開発者と使わない開発者の生産性の差は、10倍以上に広がります。

こうした変化がすべての産業で同時に起きたとしたら、どうなるでしょうか。

人々が新しいスキルを身につける速さより、技術が進歩する速さが上回ったとしたら、どうなるでしょうか。新しい仕事生まれる速さより、既存の仕事が消える速さが上回ったとしたら、どうなるでしょうか。

これがアルトマンの懸念するところです。

5年前、誰かに「2024年には今のo1と同じくらい強力なAIシステムが存在するだろうか？」と聞いたとすれば、ほとんどの人は「いいえ」と答えたでしょう。そしてもしそんなシステムが存在するならば、「世界が完全にひっくり返るだろう」と予想したはずです。

実際にはどうなったでしょうか。

私たちはそうしたシステムを手にしなが、今夜の夕食を悩み、週末の計画を立てています。アルトマンは冗談めかしてこう言います。「私たちが初めてのAGIを出したとしても、人々は20分間興奮して、すぐに夕飯の献立を考え始めるでしょう。」

これは人間の適応力を示す一方で、問題でもあります。私たちはあまりにも早く慣れてしまいます。驚くべき技術が生まれても、すぐに当たり前のこととして受け止めてしまいます。そのことが、変化の本当の速さを実感できなくさせています。

アルトマンは未来を5年、10年、15年という単位で区切って考えます。

5年後には、AIが実際に科学的発見を助けたり、自ら重要な発見をしたりする時代が開けるでしょう。科学者の生産性が2～3倍になるどころではなく、10年分の研究を1年で終えることができるようになるかもしれません。

10年後には、ロボットが工場や物流倉庫や街頭で実際の労働を担う光景が自然なものになっているでしょう。知性だけでなく、物理的な世界にまでAIが踏み込む段階です。

これらすべての変化が、信じられないほどの速さで起きていくでしょう。

アルトマンが強調するのは、備えることです。「私たちはここで、かなり速く多くの問題を解決しなければなりません。」

規制は必要ですが、革新を妨げる規制ではなく、「テストと理解から始まる規制」でなければなりません。欧州のようにAIモデルをリリースする前に数多くの承認を得なければならないとすれば、技術の進歩は止まります。その代わり、まずAIを実験し、実際にどんなリスクが現れるかを観察したうえで、それに見合った安全対策を整えるべきです。

教育システムも変わらなければなりません。今の小学生が大学を卒業するころには、彼らが学んだスキルの半分は時代遅れになっているかもしれません。では、私たちは「何を」教えるべきでしょうか。アルトマンの答えは明確です。「問いを見つける力、批判的思考、そして機敏さ。」

インフラも整えなければなりません。アルトマンは「知性の時代は、すなわちエネルギーの時代だ」と言います。これからの10年は、豊富な知性と豊富なエネルギーを同時に確保しなければならない時期です。Stargateプロジェクトに5000億ドルを投資するのも、そのためです。

最も重要なのは、社会的な対話です。

AIがもたらす変化について、私たちの社会がどんな価値を優先するのかについて、誰を守りどのように支援するのかについて。こうした議論は、技術の開発と同じくらい速いペースで進められなければなりません。

アルトマンの懸念はこれです。技術は毎年2倍のペースで速くなっているのに、私たちの備えが10年前とほとんど変わらないとしたら、どうなるでしょうか。

技術そのものは危険ではありません。危険なのは、制御されていない変化のスピードです。列車が時速300キロメートルで走ることは問題ありません。問題は、線路が時速50キロメートルを基準に作られているときです。

私たちは今、そういう状況に近づいています。

アルトマンは悲観していません。「人間の精神が解決するだろう」と彼は信じています。ただ、こう警告します。「速く動かなければなりません。時間はあまりありません。」

変化のスピード。これが、サム・アルトマンが夜も眠れない、ただ一つの理由です。

それでも、他者の役に立つ存在であれ。

あるディスカッションの場で、誰かがサム・アルトマンに問いかけました。

「50年後、100年後、1000年後の人間は、何のために存在するようになるのでしょうか？」

アルトマンは問いを言い換えました。「それよりも面白い問いは、今日の人間は何の役に立っているか、ということです。」

そして彼は、ごくシンプルに答えました。「他の人たちの役に立つ存在であること。」

この答えは、最初はあまりにも平凡に聞こえるかもしれませんが、しかし深く考えてみると、これはAI時代を貫く最も重要な洞察です。

AIがどれほど賢くなっても、ロボットがどれほど多くの仕事を代わりにこなしても、人間が意味を感じる場所は結局、「他者の役に立っているか」という問いの周りに集まっています。

アルトマンは友人のポール・ブフハイトのアイデアをよく引用します。いつか「人間のお金と機械のお金が別々に存在する時代」が来るかもしれない、という想像です。

機械同士が取引し、AI同士がエネルギーや資源を売買する経済が回っています。しかし人間には、依然として人間同士の経済が、人間同士の価値体系が残るでしょう。

その世界でも、人々は依然として隣の人にどれだけ役立てるかによって、自分を評価するでしょう。

子どもをきちんと育てる親。仲間を気遣うチームリーダー。社会的弱者を助ける活動家。不便さを解消する製品を作った起業家。形は変わっても、「誰かの人生を少し良くしてあげる存在」であることは、依然として人間を感じる最大のやりがいです。

アルトマンは、AIがこの点をむしろより鮮明に浮かび上がらせると見えています。

強力なツールが登場すれば、人間はもはや「情報を素早く探す能力」や「記憶力」といったものだけでは差別化できなくなります。そういったことはAIの方が得意だからです。

では、残るものは何でしょうか。

一つ目は「正しい問いを見つける能力」です。答えを探すことはAIが手伝えますが、何を問うかがより重要になる時代です。どんな問題を解くべきか、どの方向に進むべきか、何が正義で何が危険かを問うことは、依然として人間の領分です。

アルトマンがTED講演で強調したのも、まさにこの点です。「AI時代に最も重要なのは、何を作るかを決める能力です。」

二つ目は「他者を気遣い、説得し、共に動かす能力」です。

医療の分野を例に挙げてみましょう。AIは診断において医師を上回るようになり得ます。実際、一部の研究ではすでにそのような結果が出ています。しかし患者が聞きたいのは、「大丈夫ですよ」と言ってくれる人の声です。悪い知らせを伝えるとき、隣に座って目を合わせてくれる人のまなざしです。

アルトマンはこう言います。「AIとの対話は良い娯楽の一種です。時に助けや慰めを与えることもできます。しかしこれは、社会の一員となり、集団の一員となるという社会的な欲求を満たすことはできません。」

私たち人間は、進化的に他者が自分をどう思い、どう感じるかを深く気にするようにできています。AIロボットからは、社会的な地位を得ることが難しいのです。

AIが「あなたは素晴らしい」と百万回言ってくれるよりも、心から尊敬する一人の友人から「君を誇りに思う」と言われることの方が、私たちをずっと幸せにします。

アルトマンが父親になって、より深く考えるようになったのもこの点です。「私の子どもは、決してAIより賢くなれないでしょう。でも、それがそれほど重要なことでしょうか。」

その子が生きる世界で大切なのは、その子が他者とどんな関係を築くか、どれほど豊かで意味のある人生を送るかでしょう。

AI時代の未来において、価値ある能力とは何でしょうか。

アルトマンは「俊敏性（アジリティ）」を強調します。単なる知能（ability）よりも俊敏性（agility）の方が価値を持つようになるでしょう。素早く学び、状況に応じて変化し、新しいツールを使いこなす能力です。

「パターン認識と点と点をつなぐ能力」も重要です。事実を集めるだけでなく、パターンを統合・認識し、点を結びつける能力が、はるかに価値を持つようになるでしょう。

「共感力と対人関係のスキル」が逆説的により重要になります。AIが発展するほど、共感、感情の理解、コミュニケーション能力が差別化の要素となります。教育、心理カウンセリング、リーダーシップ、創作活動といった分野では、コミュニケーション能力の価値がさらに高まるでしょう。

アルトマンが描くAIの役割は明確です。「世界には、潜在力に見合った機会を得られていない人が多すぎます。」

良い学校、良い都市、良い会社に属していないという理由だけで、すでに優れた才能を持ちながらも目の目を見られない人たちがいます。

AIはこの壁を壊すツールになり得ます。どの国に生まれ、どんな背景を持っていても、良い教師、良いコーチ、良いアドバイザーにAIを通じて出会えるなら、「最大の潜在力に達する人の数」は今よりずっと増えるでしょう。

アルトマンは、これが「人間のお金」と「機械のお金」が共存する方式だと考えています。機械が効率性と生産性を担うなら、人間は意味とつながりを担います。

最後に彼は、倫理に関する原則を繰り返します。「ルールは人間が決めなければなりません。AIはそのルールに従うように作らなければなりません。」

どんな情報は見せてはいけないか、どんな用途は許可してはいけないか、どんなリスクは受け入れられないかを決めるのは人間のコミュニティです。AIは、その決定を効率的に実行するツールに過ぎません。

サム・アルトマンの最後の助言は、意外なほど素朴です。

AI時代に人間は何をすべきかと問われると、彼は「それでも、他者の役に立つ存在であろうとし続けなさい」と答えます。

AIをうまく使うことよりも、人をうまく助ける人間になること。その目標は、AI以前と以後を分け隔てません。

技術が変わっても、道具が変わっても、時代が変わっても。人間の最も深い欲求は変わりません。誰かにとって意味のある存在でありたいという欲求。誰かの人生に小さな違いを生み出したいという願い。

これこそが、AI時代においても、いやAI時代だからこそ、いっそう輝く人間の役割です。

← 前の章

キム・ギョンジン弁護士

弁護士 ・ 元国会議員 ・ AI政策研究者

kimkj.com

© 2026 キム・ギョンジン弁護士. All rights reserved.

#キムギョンジン #AI書齋 #サムアルトマン #OpenAI #ChatGPT #人工知能

このAI書籍を応援する

この本が役に立った場合は、今後の翻訳、公開版、
AI知識プロジェクトをPayPalで応援できます。



PayPal

<https://paypal.me/kimkjkj>

QRコードを読み取るか、上のリンクを開いてください。