



中国AIの父たち

中国AIを動かした十人の評伝

金京鎮弁護士

AI Library - kimkj.com - 2026

目次

1. 第1章 リ・カیف(李開復、Kai-Fu Lee)、世界の舞台から帰還した設計者、01.AIと主権AIへの道
2. 第2章 李彦宏(リー・イエンホン、Robin Li)、検索会社をフルスタックAI企業に変えたバイドゥ創業者
3. 第3章 唐傑(タン・ジエ、Jie Tang) - 清華大学知識工学からGLMへとつながった研究者の道
4. 第4章 梁文鋒(リャン・ウェンフォン、Liang Wenfeng)、ヘッジファンドの計算力を研究所に移したDeepSeekの創業者
5. 第5章 楊植麟(楊植麟、Yang Zhilin)、長いコンテキストでAGIを推し進めたMoonshotの若き創業者
6. 第6章 閔俊傑(イエン・ジュンジエ、Yan Junjie)、商湯からMiniMaxへ、消費者向けAIとLightning Attentionを推し進めた創業者
7. 第7章 王小川(ワン・シャオチュアン、Wang Xiaochuan)、検索の戦場を離れ、医療AIへと舵を切ったBaichuan創業者
8. 第8章 周靖人(Jingren Zhou)、Qwenのエコシステムを世界の開発者に開いたアリババクラウドの技術責任者
9. 第9章 姚順雨(ヤオ・シュンユー、Shunyu Yao)、ReActとTree of ThoughtsからTencent Hunyuanへ渡ったエージェント研究者
10. 第10章 吳永輝(吳永輝、Yonghui Wu)、Googleの翻訳研究からByteDance Seedの基盤モデル統括者へ

第1章 リ・カیف(李開復、Kai-Fu Lee)、世界の舞台から帰還した設計者、01.AIと主権AIへの道

成長期と学業

1979年、リ・カیفは米コロンビア・カレッジに入学しました。

コロンビア・カレッジはニューヨークのマンハッタンにある学部中心のカレッジで、リ・カیفはここでコンピュータサイエンスを専攻に選びました。1970年代末のアメリカにおけるコンピュータサイエンスは、今日のような独立した巨大な学問分野ではなく、数学科と電気工学科のあいだでようやく自分の居場所を見つけつつあった新興分野に近いものでした。彼が入学したのは、パーソナルコンピュータがちょうど大衆の前に姿を現し始めた頃のことでした。

人工知能という学問は、それ以上に馴染みの薄いものでした。1956年のダートマス会議でその名を得てから20年余りが過ぎていましたが、当時はまだ規則を一つひとつ手作業で書き込む記号論理方式が主流であり、ニューラルネットワークやディープラーニングといった発想は学界の主流から外れた位置に置かれていました。リ・カیفが学部生としてこの分野に初めて出会ったのは、まさにそうした停滞期の真ただ中でした。

1980年、大学2年生になったリ・カイクは、自らの進路を決定づける二人の教授に出会いました。一人は彼に自然言語処理を教え、もう一人はコンピュータビジョンを教えました。二つの講義は扱う問題こそ異なりましたが、共通する一つの問いを投げかけていました。機械は人の言葉を理解し、人の目が見るものを読み取ることができるのか、という問いです。

自然言語処理とは、人間の文章をコンピュータが分析し理解できるようにする分野であり、コンピュータビジョンとは、写真や映像の中の物体をコンピュータが認識できるようにする分野です。今でこそスマートフォン一つでこの二つの技術を毎日使っていますが、1980年当時はどちらもまだ始まったばかりの段階で、成功例より失敗例のほうがはるかに多いものでした。リ・カイクはその失敗を前にしても、むしろそこに可能性を見出しました。

二人の教授の講義を聴きながら、彼は人工知能という分野の大きさを実感したと伝えられています。人間の思考や認識を機械に移し替える作業は、ある一学科の狭い課題にとどまるのではなく、人間という存在そのものを理解する営みに通じているという気づきでした。彼はこの時から、残りの学業人生を人工知能研究に捧げると

心に決めました。

コロンビア・カレッジを卒業したリ・カイクは、1983年にカーネギーメロン大学のコンピュータサイエンス博士課程に進学しました。カーネギーメロンはピッツバーグにある工学系中心の大学で、1980年代初頭にはすでに人工知能とロボット工学の研究においてアメリカ国内でも屈指の地位を占めていました。この大学で博士号を取得したという事実は、その後、彼の経歴が紹介されるたびに欠かせない基本情報となりました。

博士課程に応募するにあたり、彼は研究計画書を提出しました。その計画書には、人工知能は人類が自らを認識し理解するための最後の一マイルである、自分はこの新しく花開こうとしている有望な科学分野に加わりたい、という一文が記されていました。この一文は、彼が人工知能を単なる基礎的な計算技術としてではなく、人間の精神を解き明かす学問として捉えていたことを示す記録として残っています。

「一マイル」という表現は、マラソンでゴール直前の最後の区間を指すときによく使われる言葉です。人類は宇宙や物質、生命の仕組みを一つずつ解き明かしてきましたが、肝心の自分自身の思考が

どのように生み出されるのかについては、長いあいだ解けないまま残してきました。リ・カイクは人工知能研究を通じて、まさにその最後の区間を越えようとしていたようです。

ただし、彼が博士課程で取り組んだ課題は、音声をコンピュータに認識させる研究でした。これは後年、彼が数々のメディアのインタビューで自身の学問的出発点として繰り返し語ることになるテーマとなりました。自然言語と音声、そして認識という三つの関心が、学部時代から博士課程まで一貫して続いていたようです。

リ・カイクが博士課程に在籍していた1980年代初めから半ばにかけては、人工知能研究全体が再び冬の時代を迎えた時期でした。政府や企業の研究資金が減り、華々しい約束を掲げていた初期の人工知能プロジェクトが次々と期待に届かない成果しか出せず、学界の外での関心は急速に冷めていきました。こうした空気の中で人工知能を生涯の研究テーマに選ぶことは、当時としてはありふれた道ではありませんでした。

それでも彼は揺らぎませんでした。人工知能研究が幾度となく関心の外へ押しやられては再び浮上することを繰り返すあいだも、いつか機械が人間の知能の限界を超える瞬間が訪れるという信念を手

放さなかったと伝えられています。この信念は、その後40年近くにわたる彼のキャリア全体を貫く一つの軸となりました。

博士課程で彼が没頭した研究は、人間の音声をコンピュータが文字に変換する音声認識技術でした。今日ではスマートフォンが数秒で話し言葉を文字に変えてくれますが、1980年代半ばの音声認識技術は、あらかじめ決められたわずかな単語をようやく聞き取れる程度の水準でした。彼がこの難題を博士論文の研究テーマに選んだという事実は、彼が最初から実用性と学問的挑戦の両方を追い求めていたことを物語っています。

コンピュータサイエンスという学問そのものが、彼の学生時代を通じて急速に規模を拡大していったことも付け加えておく価値があります。彼がコロンビアに入学した1979年からカーネギーメロンで博士課程を終えるまでのあいだに、パーソナルコンピュータは実験室の道具から家庭やオフィスの道具へと居場所を移しました。彼の学業期は、コンピュータサイエンスが独立した学科として完全に定着していく時期とちょうど重なっています。

博士号取得後、リ・カイクは学界と企業のあいだを行き来しながらキャリアを重ねました。アップルでは人工知能関連部門を率い、

シリコングラフィックスとマイクロソフトでは国際的な役職を務め、その後はグーグル・チャイナを率いました。

1998年、彼はマイクロソフト・リサーチ・アジアを設立し、再び教育者に近い役割を担うことになりました。当時、中国の大学におけるコンピュータサイエンス研究と教授陣の水準は、アメリカのトップクラスの大学と並べて比較するのが難しい状態だったと伝えられています。彼は聡明な中国人学部生たちを研究所に採用し、一からじっくり育て直す方式を選びました。

この方式は、彼がコロンビアとカーネギーメロンで自ら経験した教育過程をなぞったものに近いといえます。良い教授に出会って学問の大きさを実感し、難しい研究テーマに自ら取り組み格闘しながら実力を積み上げていく過程を、彼は今度は教わる側ではなく教える側の立場で繰り返したのです。このとき育てられた人材たちは、後に中国の人工知能産業全体にわたって研究者や起業家として広がっていきました。

彼の学業期の物語から繰り返し確認できる流れは明確です。まだ目立った成果のなかった学問分野に飛び込み、その分野が幾度も関心の外へ押しやられるあいだも踏みとどまり、学位を終えた後は学

んだことをそのまま次の世代へ返していったのです。この流れは、その後彼が歩んだすべてのキャリアの土台となりました。

2023年5月、62歳になった彼は北京で01.AIを設立し、自ら最高経営責任者の座に就きました。40年前の研究計画書に記したあの一文を、彼は自らの手で締めくくろうとしていたのです。

大学に入学した1979年当時のリ・カیفは、まだ名も定まらぬ学問にすぎる一人の青年にすぎませんでした。彼が学部と大学院で積み重ねた歳月は、華々しい成果の連続ではなく、幾度もの冬を耐え抜いた信念の記録に近いものでした。その信念こそが、後に彼をアップル、マイクロソフト、グーグルへと導き、ついには自らの会社を設立するところまで連れて行ったのです。

コロンビアで出会った二人の教授の講義、カーネギーメロンで書き記した研究計画書の一文は、今読み返しても素朴なものです。しかし、この素朴な出発点こそが、後にリ・カیفという名を人工知能の分野で語り継がれ続けさせる根となったという事実は、彼のその後の歩みをたどるほどに、いっそう鮮明になっていきます。

研究者、あるいは起業家としての土台を固めた時期

リ・カルフは1979年、米コロンビア・カレッジに入学して初めてコンピュータサイエンスに触れ、その場でこの学問が秘める可能性に引き寄せられ、生涯を懸けてみようと心を決めました。二十歳になる前に定まった方向が、その後四十年以上にわたって揺らがなかったという点で、この選択は彼のキャリア全体を貫く最初の結び目と呼ぶにふさわしいものです。

コロンビアを終えた後、彼は1983年にカーネギーメロン大学のコンピュータ工学博士課程に入りました。カーネギーメロンは当時すでに人工知能研究の中心地の一つであり、そこでリ・カルフは理論と工学の両方を扱う訓練を受けました。

博士課程に提出した研究計画書には、彼が生涯抱き続けることになる一文が記されていました。彼は人工知能を、人類が自らを認識し理解していく旅路の最後の一步であると記し、この新たに花開こうとしている有望な学問分野に加わりたいと述べました。二十歳そこそこの大学院生が人工知能を、単なる基礎的なプログラムや計算技術としてではなく、人間自身を理解する道具として捉えていたという事実は、後に人工知能研究が停滞期を迎えたときにも彼が信念を守り抜くことができた土台を示しています。

実際、人工知能研究はその後も幾度となく関心と資金が途絶える時期を経験しました。学界の中で人工知能がしばらく傍流として扱われていた時期にも、リー・カيفーは大学院時代に立てた原則を捨てず、いつかこの分野の時代が来るという確信を持ち続けました。この確信は、後に彼が六十歳を超えた年齢で再び起業の最前線に飛び込む決断へとつながっていきます。

博士号を取得した後、リー・カيفーは研究室にとどまらず、産業の現場へと足を踏み出しました。彼はアップルで人工知能関連部門を率い、シリコングラフィックスで役員を務め、その後マイクロソフトとグーグルで中核となる経営陣の一員となりました。学問から出発した人物がビッグテック企業の技術最前線を次々と経験することで、彼の感覚は理論から実務へ、さらに実務から組織運営へと広がっていきました。

このキャリアの中で最も明確な足跡を残したのが、1998年に設立したマイクロソフトアジア研究院、現在のMSRAです。当時、中国の大学における人工知能研究と教育の環境は、アメリカの上位大学と比べものにならないほど遅れていました。リー・カيفーはこの格差を埋めるため、中国各地から勤勉で優秀な若い学生を選抜し、

研究院の中で彼らを直接再教育する方式を選びました。

この再教育の方式は、単に一つの研究院の人材を育てるだけにとどまりませんでした。この時期にリー・カイフーの手を経た人材の多くが、後に中国の人工知能産業のあらゆる場面で中核的な役割を担うことになりました。その流れは、後に中国で大規模言語モデルを手がける多くの企業が、自国のエンジニアだけで組織を構成できる基盤へとつながっていきました。リー・カイフーが研究者から人材を育てる組織者へと移行した最初の地点が、まさにこの局面です。

MSRAの後、彼はグーグルチャイナの社長に就任し、中国国内のグーグルのエンジニアリング組織を率いました。大規模な検索サービスとそれを支える分散コンピューティングインフラを直接運営する中で、彼は理論的に優れた人材を集めることと、その人材たちを実際に稼働するサービスへと組織化することがまったく異なる技術であるという事実を身をもって学びました。この時期を経て、リー・カイフーは研究者と経営者の両方のアイデンティティを備えた人物として確立されていきました。

2009年、彼はグーグルチャイナを離れ、創新工場(イノベーション・ワークス)、英語名シノベーション・ベンチャーズを設立しました。この会社は中国において人工知能に焦点を当てた初期投資を切り開いた企業の一つに数えられます。リー・カイフーはこのファンドを通じて約30億ドル規模の資金を運用し、次世代のディープテックと人工知能スタートアップを陰で育てました。

投資家として過ごしたこの時期は、彼が現場の起業家から一步退き、産業全体を俯瞰する立場へと移った時間でした。彼はもはや自ら直接コードを書いたりサービスを運営したりすることはありませんでしたが、どの技術とどのチームが生き残るかを見極める眼力を、この時期に磨き上げました。

この安定した投資家としての日々は、トランスフォーマー構造の発展と、2022年末にチャットGPTがもたらした生成AIの波の前で揺らぎ始めました。リー・カイフーはかつて、自分が起業を始めるにはもう年を取りすぎていると考えたこともありましたが、目の前で繰り広げられる変化の規模を目にし、その渴望を抑えることができませんでした。

彼は、今回の生成AIの波があまりに大きく圧倒的であったため、自分が競技場の外に立って他人の挑戦に投資するだけの傍観者ではいられなかったと語りました。彼は自ら競技場の中へ、実際に事が起きている真ただ中へ飛び込まなければならなかったと述べています。

その結果生まれたのが、2023年5月に北京で設立した01.AI、中国語では「零一万物」です。六十二歳という年齢で、彼は投資家の立場から再び現場の最高経営責任者へと戻りました。彼はこの挑戦を、自らの学術的・産業的な旅路を締めくくる人生最後の駅であると同僚たちに語りましたが、それは1983年の博士研究計画書に記した「人類の自己認識に向けた最後の一步」という一文を、三十年余りを経て会社という実体に置き換えたものだったのかもしれませんが。

01.AIは設立から六か月で企業価値10億ドルを超え、ユニコーン企業の仲間入りを果たしました。同社は北京に本社を置き、産業用大規模モデルの学習と企業向けソリューション構築を軸とし、リー・カيفーは創新工場AI工程院の中核育成プロジェクトとしてこの会社を立ち上げました。

01.AIがその後どのような製品や組織構造で事業を広げていったかについては、この会社の事業と研究組織を扱う場で改めて見ていくことになります。リー・カイフーという人物が研究者から組織者へ、さらに投資家から現場の起業家へと移り変わっていった流れを押さえておくだけで十分でしょう。

振り返ってみると、リー・カイフーの経歴は、一人の人間が幾度もアイデンティティを変えながら積み上げてきた結果物です。大学院時代には人工知能という学問の哲学的な根を捉え、アップル、マイクロソフト、グーグルを経る中で大規模な組織を実際に動かす技術を身につけました。MSRAでは人材を育てる方法を、グーグルチャイナではその人材でサービスを動かす方法を、創新工場では産業全体を俯瞰する眼力を、それぞれ手にしました。

この四つの経験が一人の人間の中に積み重なっていたことが、彼が六十歳を超えてなお再び起業の最前線に立てた理由だと考えられます。若い起業家にはなかなか備わらない組織運営の経験と人材を見る眼力、そして投資家として産業の流れを読む感覚を、彼はすでに身につけた状態で01.AIを始めることができたのです。

そうした細部が欠けているからといって、この時期の意味が薄れるわけではありません。むしろ、大学院の研究計画書の一文から始まり、ビッグテック企業の役員の座を経て、投資家として一步退いた後に再び起業家として戻ってくる軌跡そのものが、彼がなぜ研究者と起業家という二つのアイデンティティの間に長く揺れることなく、結局その両方を自らのものにできたのかを物語っています。

ただ、一人の人間が研究者として出発し、組織者と投資家を経て再び起業家として完結するまでに、その結び目がどのように結ばれたのかを確認しておけば十分です。

会社と研究組織を築いた過程

2022年11月30日、オープンAIがチャットGPTを世に送り出した瞬間、リー・カイフーは四十年余り見守り続けてきた人工知能研究がついに実用の域を超えたと感じました。創業メンバーだった楊植麟の回顧によれば、チャットGPTの登場は業界全体の資本と人材が動く方向をまるごと変えてしまった出来事でした。投資家として競技場の外に立ち見守るだけの立場に、リー・カイフーはもはやとどまることができませんでした。

彼は2023年5月、北京に本社を置く01.AI、中国語では零一万物を設立しました。創業初日から彼が力を注いだのは、技術ではなく人でした。グー・シュエメイ、マー・ジエ、チー・ルイフォン、そしてアニタと呼ばれるホアン・ホイウェンまで、いわゆる「マイナス1階の中核経営陣」と呼ばれたこの四人は、それぞれの分野で最高水準と目される人物であり、会社が何度も方向転換する間もその座を守り、組織の基本構造としての役割を果たしました。リー・カイクーは彼らを通じて志を同じくする優秀な人材を次々と引き寄せることができた、後に語っています。強い人材密度とリー・カイクー一人の名声が重なり合い、01.AIは設立から六か月で企業価値10億ドルを超えるユニコーン企業へと躍り出ました。

会社を設立することとコンピューティング資源を確保することは、01.AIにおいて事実上一体のものでした。最も重要な投資家の一つがアリババクラウドでした。この会社が01.AIに提供したのは、グラフィックカードの実物ではなく、クラウドコンピューティングを利用できる代金券、いわば利用権クーポンでした。投資金が入ってくる瞬間、その資金はそのままアリババクラウドの演算資源を使う形へと還流する構造であり、01.AIは海外業者を通じて一部のグラフィックカードを別途借りることもありましたが、それはあくま

で補助的な手段にすぎず、会社の主力演算は最初から最後までアリババクラウドのインフラの上で稼働していました。2024年7月には追加資金調達を完了し、手元資金の流動性が頂点に達し、リー・カイフーは社内会議で、今後長期間使っても十分なだけの資金を確保したと述べました。

資金とコンピューティング資源を備えたからといって、モデル開発の方向まで自動的に決まるわけではありませんでした。2024年5月、01.AIはパラメータ規模が千億に達する大規模モデルYi-Largeを公開しました。この発表を終えた直後、リー・カイフーは自分たちが立っている位置を冷静に見つめ直しました。アメリカの先進的な研究機関は兆単位のドルを投じて超大規模モデルの学習に乗り出していました。中国のスタートアップがかき集められる資金は、それに比べれば億単位のドルにとどまっていた。パラメータをむやみに拡大する競争では勝てないという判断が下されたのです。

リー・カイフーが導き出した結論は、会社を守り抜く持久力がモデルの大きさではなく、インフラ効率と商業的成果の証明にあるというものでした。彼は2024年5月から6月にかけて、次世代の主力モデルとして準備していた超大規模モデルYi-X-Largeの学習計画を

取り下げました。代わりに研究組織の力を、より速く、より安価でありながら投資対効果の明確な混合エキスパート構造、すなわちMoEアーキテクチャへと移しました。いわゆる「小さくても美しい」方向へと舵を切った形でした。その結果生まれたのが、2024年10月に国際ブラインドテストプラットフォームLMSYSで世界6位、中国国内1位に輝いたYi-Lightningでした。

このモデルはグラフィックカード2000台を投入し、一か月半かけて学習させ、費用は300万ドルを上回りました。これは超大規模モデルの学習にかかる300万~400万ドルとさほど変わらない水準でありながら、はるかに実用的な性能を発揮しました。リー・カイフーは、スタートアップの初年度の戦略が翌年もそのまま通用するとは限らず、調整と転換は創業が経る当然の過程であると説明しました。

方向を転換した後は、組織をその方向に合わせて再編する必要がありました。2024年11月下旬から、社内ではアルゴリズムチームとクラウドコンピューティングチームを大きく異動させる議論が交わされました。そして2025年1月初め、市場には01.AIが資金難に耐えられずチーム全体をアリババに譲渡するという噂が急速に広まり

ました。リー・カيفーはWeChatモーメンツに即座に売却説を否定する投稿を掲載し、メディアのインタビューを通じて実際に何が変わったのかを自ら説明しました。

彼が明らかにした実際の変化はこうだった。2025年1月2日、01.AIとアリババクラウドは「産業向け大規模モデル共同研究所」を正式に設立した。事前学習アルゴリズムチームはアリババの通義大規模モデルチームへ、インフラエンジニアリングチームはアリババクラウド所属へと、それぞれ籍を移した。超大規模モデルのスケールアップを追い続けたかった研究者たちは、この移籍によってその道が続けることができ、01.AIは彼らがアリババの資本で作り上げていく超大規模モデルを、今後自社の応用モデルを教える「教師モデル」とする協力体制を築いた。

リー・カيفーはこの再編について、会社が事前学習を完全にやめるわけではないが、もはや超大規模モデルを追いかけないという意味だと整理した。人員をすべて出したわけでもなかった。01.AI内には中核アルゴリズム人材とインフラ最適化人材が一定規模で残り、彼らは汎用人工知能を狙う超大規模事前学習ではなく、微調整とプロンプトの精緻化、プラットフォームツールの開発、そして産

業別に安価で実用的なモデルを作る方向へと任務を切り替えた。

組織再編には人が去ることも伴った。2024年8月、技術副社長だった黄文昊とファン・シンが会社を去った。会社が汎用オープンソースエコシステムの拡大より企業向け事業に力を注ぐことを決めたことで、その年の12月前後には開発者リレーションとオープンソースコミュニティを担当していた部署の人員も大幅に整理された。リー・カイフーは、リンイーワンウーの中核共同創業者の陣容はこうした変化の中でも比較的安定して維持されていると明らかにした。

身軽になった後、01.AIが向かった先は企業向け、いわゆるToB事業だった。リー・カイフーは、損をしてまで受注する事業はしないという原則を明確にした。代わりに、エンジニアが直接顧客企業の現場に入りシステムを共に組み立てるFDE、すなわちフォワード・デプロイド・エンジニア制度を定着させた。この方式は基礎的なソフトウェア納品ではなく、人と物と手続きと意思決定の論理を計算可能な対象に変えるオントロジー作業に近く、そうして積み上げたビジネスマップをもとに企業が実際に使えるモデルを作る仕事だった。

成果は数字にも表れた。2024年の一年間で01.AIは1億元を超える実売上を上げ、そのうち約70パーセントが企業向け事業から生まれた。ゲームとエネルギー、自動車、金融の分野では1千萬元を超える契約が協議され、消費者向け事業では海外の有料サービスが売上の20から30パーセントほどを占め、損益分岐点に近づきつつあった。

具体的な事業事例も続いた。四川省内江高新区とは1億5千萬元規模の契約を結び、地域産業を押し上げる人工知能拠点を共に築いた。フォーチュン・グローバル500企業であるタイの正大集団とは、養鶏管理と農業労働を人工知能で支援するシステムを共同で作り上げた。カザフスタンとはQ.AIという合弁会社を設立し、元人工知能担当次官だったドミトリー・ムン氏が最高経営責任者を務めることにもなった。

企業向けプラットフォームも時間をかけて一つずつ整えていった。2025年3月には自社の大規模モデル技術をもとにした「万知」企業プラットフォームを発表し、企業がそれぞれの必要に合わせてモデルを調整し安全に配備できるようにした。2026年1月には国家級ハイテク企業と北京市専精特新中小企業に相次いで認定され、制度

的な信頼を積み上げた。そして2026年7月には「万策」という名の意思決定プラットフォームを披露した。ここには企業トップを助ける「老板AI」、投資判断を支える「投資官AI」、営業現場を管理する「小官AI」が共に組み込まれていた。

リー・カيفーはこうした変化を経る中で、人工知能転換の価値は試験的な事業をいくつか増やすことにはないと繰り返し強調した。戦略と組織、実際の業務シナリオ、データとモデルを一つに束ね、企業の意思決定の中枢にまで人工知能が入り込むようにしてこそ意味があるというのだ。コンサルティングで試行錯誤を減らす。オントロジーで業務マップを描き、FDEが現場で実行を前倒しし、そうして積み上がった経験を再びデータとモデルへ戻して磨き続ける循環構造が、01.AIの組織運営の土台として定着した。

振り返ってみれば、01.AIが歩んできた道は、最初から決められた青写真をそのままなぞった結果ではなかった。ChatGPTがもたらした衝撃の中で人材を集め、アリババクラウドの計算資源と資本を一つにまとめて初期成長の足場を作り、超大規模モデル競争での勝算が低いと悟るや、すぐさま方向を変えて組織を組み直した。そしてその組織をアリババとの協力実験室へ一部譲り渡してまで身軽にな

り、残った力を企業現場で直ちに価値を生み出す方向へと集めた。リー・カيفー自身が語ったように、これは勝てない戦いを避け、回収の見えない投資はしないという原則を組織設計そのものにそのまま映し出した結果だった。

中核モデルと開発プロセス

中国の大規模モデル産業が本格的な競争に入った時点で、リンイーワンウーは最初から二つの方向を同時に準備していた。世界の学界と開発者コミュニティから信頼を得るオープンウェイトモデルと、市場で実際に収益を生む閉鎖型商用モデルである。この二つの方向、いや二つの軸が互いを支えながら会社の技術路線を作り上げ、その結果がYiシリーズという名で残った。

最初に出たのはYi-34BとYi-6Bだった。パラメータ規模を大きく抑えたこれらのモデルは、安価なサーバー上でもすぐに動かしファインチューニングできるよう設計されていた。重みをそのまま公開し、海外の開発者たちが自由に使えるようにした。コード作業に特化したYi-Coderも同じ方式で公開され、プログラミングベンチマークで性能が検証される小型モデルの位置を占めた。

2024年5月、リンイーワンウーは約1千億パラメータ規模の閉鎖型モデルYi-Largeを世に送り出した。このモデルは公開直後、LMSYSの大規模モデルランキングで世界有数の企業の非公開モデルと並んで競い合い上位に入り、中国企業が作ったモデルの中では1位を占めた。リー・カイフーはこの結果を、リンイーワンウーが世界水準の事前学習能力を備えていることを証明した出来事として、複数のインタビューで繰り返し振り返っている。

だが、リー・カイフーはこの成功の裏側で、すでに別の計算をしていた。彼がLeanシャン・インサイトに明かしたところによれば、昨年5月にYi-Largeを出した瞬間から、この方向が時間とコストをあまりに食いすぎる道であることに気づいたという。超大規模モデルを育て続けるにはそれだけ多くのGPUと資源が必要で、スタートアップが背負うには重荷が大きかった。

そのため2024年5月から6月にかけて、会社は当初準備していたより大きなモデルYi-X-Largeの学習計画を取りやめた。代わりに、混合エキスパート構造、いわゆるMoEアーキテクチャを採用した新しいモデルの開発に着手し、その成果がその年の10月に発表されたYi-Lightningだった。この転換は社内で静かに行われたものではなく

、以後複数のインタビューでリー・カイフーが繰り返し説明する中核的な決定として位置づけられた。

Yi-Lightningは以前のモデルより速度が速く、推論コストも低かった。リー・カイフーは、事前学習一回にかかる費用が通常300万から400万ドルに上るところ、比較的安価な部類に入るYi-Lightningでさえ、GPU2千枚を一か月半動かして300万ドルを超える費用がかかったと明かした。この数字こそ、会社がなぜより大きなモデルを諦めたのかをそのまま示す根拠だった。

リー・カイフーはこの転換を、失敗の認定ではなく判断の結果として説明した。彼はメディアのインタビューで、スタートアップの一年目の戦略が二年目にも当てはまるとは限らず、調整と転換は起業の必然だと語った。2025年を商業化の勝負どころと定め、それに合わせて会社の事業優先順位も組み直したというのが彼の説明である。

技術の方向を変えた後も、超大規模モデル自体を完全に捨てたわけではなかった。リンイーワンウーは2025年1月2日、アリババクラウドと共に産業向け大規模モデル共同研究所を設立すると発表した。事前学習アルゴリズムチームとインフラチームの一部がこの研

究所に合流し、アリババクラウド所属へと移った。リンイーワンウー自体は、より速く安価で小さなモデルに集中する方向に残った。

この決定をめぐって、市場ではチーム全体をアリババに譲り渡すという噂が流れたが、リー・カイフーはWeChatモーメンツを通じて身売り説を直接否定した。彼がLateTalkのインタビューで語ったところでは、リンイーワンウーは買収を打診されたこともなく、今回の調整はスケーリング則が鈍化する局面で、会社が背負える戦いを選ぶ過程にすぎなかった。

超大規模モデルを完全に捨てなかった理由は別にあった。リー・カイフーは、アリババが超大規模モデルの学習を続ける理由を教師と生徒の関係にたとえた。体の大きなモデルがラベリング結果と合成データを作り出せば、それを体の小さな実用的なモデルが学び、性能を引き上げられるというのだ。彼はAnthropicがClaude 3.5 OpusでClaude 3.5 Sonnetを教えた方式を例に挙げた。スケーリング則を追い続けたかった一部の事前学習チームとインフラチームのメンバーが、アリババとの共同研究所を選びその旅を続けたのも、こうした文脈からだった。

モデルをどう作るかに劣らずリー・カイファーが強調したのは、モデルとインフラと応用を一つに結びつける方式だった。彼はZaijingのインタビューで、モデル、AIインフラ、アプリケーションの三要素がすべて揃わなければならないと語った。AIインフラを設計する段階からモデル学習に必要な条件を共に考慮してこそ全体の効率上がり、この三つを一つに配置する方式こそがリンイーワンウーの持つ競争力だというのだ。

こうしたアプローチが生まれた背景には、中国企業が置かれた資源条件があった。リー・カイファーはCapgeminiのインタビューで、中国企業はAnthropicやOpenAIをそのまま打ち負かそうとはしないと語った。超大規模モデルの学習に500億ドルをつぎ込むことはできず、確保できるGPU自体も限られているため、残る問いは持っている資源をどこまで絞り出せるかということだった。答えは工学的効率を極限まで引き上げることであり、DeepSeekがその代表的な事例として挙げられた。

リー・カイファーはこの条件を不利な条件とばかりは見なかった。資源が限られ突破口を用意しなければならないという圧力があるとき、むしろ制約の中から革新が生まれるというのが彼の説明である

。モデル構造を最適化しCUDAへの依存度を下げるようにして、中国企業は米国トップ企業の投資額の10分の1にも満たない費用で性能の90から95パーセントに達するモデルを作り出し、格差を6か月から9か月の水準にまで縮めてきたと彼は語る。

同じ文脈で彼は、中国企業がオープンソースを受け入れる理由を説明した。閉鎖型競争で勝てないのなら、公開し共有する道を選ぶというのだ。彼はこれを勉強会にたとえた。それぞれがノーベル賞を狙って中核技術を隠す米国企業とは違い、中国企業はそれぞれの成果を持ち寄り、集団全体が目標に近づく方式を取るという説明だった。

オープンソースに向けたこうした姿勢は、リー・カيفーにとって技術の問題を超えて国家単位の問題でもあった。彼は主権を幾つかの層に分けて説明した。技術そのものに対する統制権、データが外部に流出しない安全性、そしてある国で学習されたモデルはその国の価値観を帯びているため他の地域には合わないことがあるという適合性の問題である。モデルを直接手に握ってこそ組織の必要に合わせて調整できるというのが彼の要旨だった。

この時点で彼は、どの国も最初からオープンAIのような会社を新たに作る必要はないと述べた。代わりに彼が示した道は、先行するオープンソースモデルを半製品として取り入れ、自国の言語と規範に合わせて継続的に学習させることだった。彼はこれをピザに例えた。小麦を最初から栽培する必要がなく、冷凍ピザを買って自分の好みに合ったトッピングを載せて焼き直すように、オープンソースモデルを取り入れてインド風にも日本風にも作り変えられるというのだ。この方式のコストは最初からモデルを学習させるコストのわずか数パーセントに過ぎず、数億ドルではなく数百万ドル規模で済むと彼は説明した。

彼はこの第三の道を、米国の閉鎖型プラットフォームをそのまま受け入れてデータと統制権を渡してしまう方式や、無理に最初から自前のモデルを作る方式と区別した。検証済みのオープンソース基盤モデルの重みをダウンロードした後、その国の言語と法律、価値観に合わせて微調整する経路のほうが、大半の国にとってより現実的だというのが彼の結論だった。こうした考えは、01.AIがなぜYiシリーズの重みをオープンソースとして公開したのかを説明する背景でもある。

技術と戦略をこのように整理した後も、リ・カイクが繰り返し強調したのは結局のところ商業化だった。彼は事前学習に注いでいた力を引き上げる代わりに、実際の売上につながる応用と導入支援に組織の比重を移したと語った。ゲーム、エネルギー、自動車、金融といった分野で1000万元以上規模のソフトウェア契約を複数協議しているという説明も、この時期のインタビューで語られた。

その結果は数字にも表れた。2024年、01.AIの実際の売上は1億元を超え、このうち約70パーセントが企業向け事業から生まれたとリ・カイクは明らかにした。消費者向け売上のうち20から30パーセント程度は海外の有料課金製品によるもので、この分野はコストと売上が損益分岐点に近づきつつあるとも述べた。

リ・カイクは、01.AIが企業顧客の業界ノウハウと自社技術を組み合わせて合併会社を設立し、垂直産業別モデルを共同で作る案まで視野に入れていると語った。パートナーがコスト削減と効率向上において目に見える価値を確認して初めて、新技術に費用を支払う気になるというのが彼の判断だった。2024年8月以降、役員数名が会社を去ったことが明らかになったものの、中核となる共同創業者陣は比較的安定的に維持されたと伝えられる。

こうした一連の決定を貫くのは、リ・カイク自身がさまざまな場で語ってきた姿勢である。彼は勝てない戦いはせず、回収が見込みにくい大規模投資は避けると述べた。2025年にはアプリケーション市場が拡大するにつれ、01.AIが垂直的な細分産業で価値ある製品と市場の接点を見出すだろうという展望も併せて示した。

中国の大型モデルスタートアップが最終的にすべて崩れるだろうという見方について、リ・カイクはその可能性はないときっぱり否定した。自然言語で対話し、汎用的な推論と理解能力を備えた、AIなしには成立しないアプリケーション群が既存の構図を揺るがすほどの規模を持つからだというのがその理由だった。彼の見るところ、中国企業はアプリケーション開発と実際の適用においてより有利な立場にあり、これはインターネットとモバイルインターネットの時代にすでに一度証明された流れだという説明だった。

現在の立ち位置と中国AIにおける意味

2026年のリ・カイクは、もはや一社の創業者という言葉だけでは説明できない立場に立っている。彼はシノベーション・ベンチャーズ(創新工場)の会長として約30億ドル規模のディープテック・ベンチャーファンドを引き続き運用する一方、2023年5月に北京で自ら設立した01.AI(零一万物)の創業者兼CEOとして経営の最前線にとど

まっている。01.AIのほかにも複数のスタートアップの育成に関わり、タイム誌が選ぶ世界的なAIリーダーのリストに名を連ね、世界経済フォーラムの第四次産業革命センターAI委員会の共同議長を務めるなど、中国のテクノロジーが国際舞台で自らの立場を説明する際の代弁者の役割まで兼ねている。

01.AIは設立からわずか半年で企業価値10億ドルを超え、ユニコーン企業となった。短期間で技術的な飛躍を遂げた中国のフロンティアAI企業として位置づけられた。しかし2024年末から2026年現在に至るまで、この会社が歩んできた道は順調な上昇曲線ではなく、生存と商用化に向けた方向転換の連続だった。

2025年1月初め、中国のテクノロジー業界では、01.AIが事前学習に必要な資金を調達できず、チーム全体をアリババに譲渡するという噂が急速に広まった。リ・カイクは買収説を否定しつつも、実際に進行していた変化については隠さずに説明した。01.AIはアリババクラウドとともに産業向け大型モデルの共同研究所を設立し、超大規模モデルの学習とAIインフラを担っていた人員の多くがこの研究所に移り、アリババ所属となった。会社に残った組織は、速度が速くコストが低く規模の小さいモデル、つまりはるかに広い範囲に

応用できるモデルに集中する一方、独自のアプリケーション開発も並行することにした。

この再編の過程で、技術の方向性をめぐる意見の相違から中核人材の一部が会社を去り、オープンソースと開発者関係を専門に担っていた組織も2024年12月頃に解体された。01.AIがオープンソース・エコシステムの拡大よりも企業向け商用化に比重を置くと決めたことと連動した変化だった。リ・カイクはこの調整について、スケーリング則が鈍化する局面で商業化に関する根本的な問いに答えるべき時が来たと語った。

彼が挙げた例えは教師と生徒である。超大規模モデルはラベル付きデータと合成データを作り出して小型モデルの性能を引き上げる教師の役割を果たし、01.AI内部の研究陣はその結果を受け取って実際に使われる生徒モデルを磨き上げるという説明である。アンソロピックやオープンAIも似た方式を採っている点を根拠に挙げ、これは後退ではなく役割分担だという意味をはっきりさせた。

2024年5月以降、01.AIは最大規模で突出した性能を狙うモデルを作るコストがスタートアップの手に負える水準を超えると判断し、より速く安価なモデルに集中する方向へと舵を切った。その成果が

1000億パラメータ規模の混合エキスパート(MoE)構造を用いたYi-Lightningであり、2024年10月のLMSYSブラインドテストで世界6位、中国独自開発モデルの中では1位に上り、効率性を証明した。

リ・カیفは、自分たちは勝てない戦いはせず、回収の見込みにくい大規模投資も行わないという原則を繰り返し表明した。企業向け事業で手を出すべき仕事は、顧客の売上を実際に伸ばすこと、業界顧客と緊密に協力すること、そして複製可能性の高いソリューションの3つだけだと整理した。2024年、01.AIは1億元を超える実際の売上を上げ、2025年には大型モデル導入に適した業界企業と合併会社を設立し、そちらのデータとノウハウを自社技術と組み合わせる計画も打ち出した。

リ・カیفは、中国の大型モデルスタートアップが全滅する可能性はまったくないと言い切った。自然言語で対話し、汎用的な推論と理解能力を備え、AIなしには成立しない、いわゆるAIファーストアプリケーションがもたらす変化が既存の構図を揺るがすのに十分だという理由からである。彼は、中国企業がアプリケーション開発と実際の事業化においてより有利な立場にあると見ている。これはインターネットとモバイルインターネットの時代にすでに証明され

た流れだという説明である。

62歳でCEOの座に自ら就いた選択を後悔していないとも語った。2025年にはアプリケーション市場が本格的に開かれ、01.AIはさまざまな産業の細分野で実質的な価値を持つ製品と市場の相性を見出すだろうという展望も併せて示した。

リ・カイクが今もっとも力を込めて語るテーマは主権AIである。彼は主権が幾つもの層を持つ概念だと説明する。第一は技術そのものに対する統制権であり、第二はデータセキュリティである。どの企業も、自社の機微なデータが他国製の汎用モデルの学習に使われる状況を警戒するものである。ある国で学習されたモデルはその国の価値観とイデオロギーを内包しており、他の地域には合わない場合があるという点も指摘した。

北京とブリュッセルとワシントンが同じ言葉を使うからといって、同じことを語っているわけではないという診断も示した。どの国にも自国の言語と文化と法的規範を反映したモデルを望む層があり、その上に技術競争で地位を確保しようとする国家的野心の層が別に存在するというのである。日本、シンガポール、サウジアラビア、インドといった国々がこうした野心を明らかにしてきたが、米国

と中国を除く大半の国は最初から最先端の汎用モデルを学習させる資源を備えていないという現実も併せて指摘した。

だから彼は、すべての国がオープンAIを新たに作る必要はないが、すべての国が自ら統制できるAIは持つべきだと語る。現実的な道は、先行するオープンソースモデルを取り入れ、自国の言語と価値観と規制に合わせて継続的に学習させる方式である。彼はこれを冷凍ピザに例えた。小麦を自分で栽培する必要がなく、半製品を買ってきて自分の好みに合ったトッピングを載せて焼き直すように、オープンソースモデルを半製品として継続的に学習させれば、インド風であれ日本風であれ、その地域に合った結果物ができるという説明である。コストも最初からモデルを作る場合に比べてわずか数パーセントに過ぎず、数億ドルではなく数百万ドル程度で十分だと付け加えた。

米国のGPU輸出規制については、諸刃の剣だと表現した。中国企業はアンソロピックやオープンAIに勝とうとはしなかった。自分たちが超大規模モデルの学習に数百億ドルを投じることはできないと分かっているため、残された資源を極限まで使い切るエンジニアリング効率にこだわるというのである。DeepSeekがその代表例であ

る。中国企業の学習投資は米国トップ企業の10パーセントにも満たないにもかかわらず性能は90から95パーセント水準に達し、格差も6か月から9か月程度にまで縮まったと語った。

彼はこうした流れの背景に、異なる協働のあり方があると見ている。米国企業は独立した研究者のように、それぞれがノーベル賞を狙って中核技術を隠す一方、中国企業は一社だけではそうした成果を出せないため、各自の成果を持ち寄って集団で目標に近づく勉強会に近いやり方を取るというのである。オープンソースモデルは通常、閉鎖型に比べて6か月から9か月遅れてきており、中国モデルが米国モデルに遅れる時間差も同程度の幅だったという観察も付け加えた。

AGIを実現した側がすべての競争相手を圧倒するだろうという見通しには懐疑的である。彼はウィンドウズの成功が、他社が似たようなOSを作れなかったからではなく、その上に無数のアプリケーションが積み重なったからだという点を例に挙げ、本当の戦場はエコシステムだと語る。iPhoneとAndroidの関係も好んで用いる比喻である。アップルが高い利益を得る一方で、Androidは世界の利用者規模と普及率で上回っているように、AI分野でも閉鎖型が利益を

得てオープンソースが市場を得るという構図が生まれ得るというのである。彼は今後のAI地形が、LinuxとWindows、iPhoneとAndroidのように、少数の閉鎖型プラットフォームと幅広いオープンソース・エコシステムが長く共存する形に落ち着く可能性が高いと見通している。

ハードウェアの話に移ると、リ・カイクの診断はより具体的になる。iPhoneが最初に登場したとき、米国はデザインとソフトウェアと世界販売網をほぼすべて握っており、中国は製造だけを担っていたため、フォックスコンが稼いだ金額はアップルに比べればわずかなものだったというのである。今は事情が変わり、中国はスマートフォンとPCからロボット、スマートグラスやウェアラブル録音機器まで、ほぼすべてのハードウェアを作るサプライチェーンを備え、ファーウェイ、シャオミ、バイトダンスがすでにハードウェアへの大規模投資を始めている点を根拠に挙げているというのである。AI時代のiPhoneに相当する製品は、中国が設計し、作り、ブランドを付ける可能性が高いという見通しである。

彼は音声認識がほぼ完璧に近い水準に達している点、そして自身自身も毎日バイトダンスの録音機器を使い、会社の会議を大型モデ

ルに書き取らせて分析させているという経験も併せて紹介した。すべての会議に直接出席できなくても、AIの助けによって事実上すべての場に同席したのと似た効果を得られるという説明である。中国は人口規模とインターネット利用量においてすでにデータ優位を持っていた国である。大規模言語モデルの時代に一時的に英語圏データの品質に押されていたその優位が、音声と映像と現実世界の相互作用データが押し寄せることで再び戻りつつあるというのが彼の観察である。録音やデータ収集に対する文化的な抵抗感が相対的に薄い点、そして中国政府がAIを実体経済と結びつける政策方向をはっきりと打ち出している点も併せて指摘した。

こうした診断を貫く実務哲学は、損をしてまで仕事を引き受けないという原則です。李開復は、大規模モデルのスタートアップが技術そのものに酔いしれるロマンから抜け出し、実際の事業価値と投資回収を証明する冷静な経営者へと生まれ変わってこそ生き残れると、繰り返し強調してきました。01.AIが企業顧客に対応する際、リモートでAPIを提供するだけの方式を避け、現場に人員を直接派遣する方式を選んだのも、こうした判断の延長線上にあります。

事前学習とインフラ組織がアリババに移った件について、李開復はこれをスタートアップが背負いきれないコスト競争から抜け出し、実利のある方向へ転換した自発的な決定だと説明していますが、業界の一部では、中核人材と演算資源を大企業に譲り渡した事実上の部分的吸収と見る向きもあります。教師モデルと生徒モデルの関係によって自社開発能力を維持したという説明についても、巨大な演算インフラを丸ごとアリババのシステムに委ねた状態で、その関係がどれほど長く安定的に保たれるかは、今後さらに見極める必要がある点です。

李開復がカーネギーメロン大学でコンピュータ科学の博士課程を始めたのは1983年のことでした。当時、彼は研究計画書に、人工知能は人類が自らを理解するための最後の一步であり、この新たに開かれる学問分野に加わりたいと記しています。40年余りを経てようやく巡り会ったAIの時代を前に、挑戦してみなければ一生の後悔として残るだろうと考えたという彼の言葉は、今01.AIが歩んでいる実用主義的な行動と、オープンソースによる主権AI構想の両方を理解するための一つの手がかりとなります。

第2章 李彦宏(リー・イエーンホン、Robin Li)、検索会社をフルスタックAI企業に変えたバイドゥ創業者

生い立ちと学業

李彦宏は1968年、中国山西省陽泉市で生まれました。都市というより、炭鉱と火力発電で生計を立てる小さな工業都市であり、彼の家庭もその街でよく見られるごく普通の労働者家庭でした。両親は五人の子どもを授かり、李彦宏はその中で四番目、三人の姉に囲まれた唯一の息子として育ちました。姉たちの中で育った少年は、同年代の男の子に比べて口数が少なく、内気なほうだったと伝えられています。

彼の人生に方向を示したのは、ほかならぬ姉でした。李彦宏が十四歳になった1982年頃、姉が陽泉市全体でも数年に一人という快挙で北京大学に合格しました。小さな街から中国最高学府への進学は町中が知るほどの出来事であり、姉は荷物をまとめて旅立つ前、幼い弟にひと言を残しました。はっきりした夢もなく日々を過ごしていた少年にとって、その言葉は初めて勉強すべき理由を植え付けるきっかけとなりました。小さな街を出て、もっと広い世界へ出ていくという目標を抱かせたのです。

勉強ばかりの子どもではありませんでした。山西省陽泉地域は伝統演劇の文化が根強い土地であり、李彦宏も幼い頃から演劇やラジオの講談にすっかり夢中でした。毎日、河北放送から流れる講談に耳を傾け、一人で庭に出ては腰にタオルを巻き、おもちゃの武器を手し、昔話に登場する將軍たちの身振りをまねて遊びました。物静かな性格でありながら、体で表現する遊びを好む子どもだったようです。

その才能に目をつけた陽泉晋劇団の指導者たちが、彼を正式な団員として迎え入れようとしたこともありました。実際に彼は晋劇団の試験を受け、合格通知まで受け取っています。子どもの選択にめったに口を出さない父親は、その岐路を前に息子と向き合い、どちらの道を選んでも尊重すると言いました。しかし北京大学に進んだ姉の背中を見てきた少年は、結局、舞台ではなく机を選びました。俳優ではなく学生であり続けるというその選択が、彼の人生で初めて自ら下した決断でした。

学業に打ち込むと決めた李彦宏は、努力の末に陽泉市屈指の名門、陽泉第一中学校に首席で入学しました。1980年代初頭の陽泉はテレビさえ珍しい土地であり、コンピューターという物は概念その

ものが縁遠いものでした。そんな街で高校一年生になった年、学校に初めてコンピューター室が開設されました。

キーボードにいくつかの英単語や数式を打ち込むと、機械がたちまち計算を終え、画面に答えを映し出す光景は、少年の好奇心をまるごと揺さぶりました。彼はコンピューター室の担当教師に毎日のようにせがんで出入りの許可を得て、そうして同年代よりはるかに進んだ操作の腕を独学で積み上げていきました。学校を代表するほどの実力者と認められた彼は、山西省の省都・太原で開かれた全国中学生コンピューターコンテストに学校代表として出場しました。

小さな街では一番だったので大舞台でも好成績を収めると信じていましたが、結果は三位にも入らない完敗でした。大会を終えて立ち寄った太原市内の書店で、彼はもう一つの衝撃を受けました。書棚には、故郷では目にすることすらできなかった最新のコンピューター技術書やプログラミング教材がびっしりと並んでいました。太原の同年代の子どもたちはこうした本に日常的に触れながら育ったという事実を目の当たりにした瞬間、彼は自分が才能で劣っていたのではなく、情報に触れる機会そのものが違っていたのだと悟りました。情報への敷居を下げ、誰もが知識に近づけるようにしたいと

いう思いはこのとき初めて彼の中に芽生え、後に検索エンジンを作る出発点となりました。

1987年、李彦宏は陽泉市の高校首席として北京大学に入学しました。コンピューターを好む気持ちは変わりませんでしたが、専攻を選ぶ際には冷静に見極めました。コンピューター理論だけを掘り下げる計算機科学科よりも、コンピューターという道具を人類が積み重ねてきた膨大な知識データと結びつける応用学問のほうが今後より大きな価値を生むと判断し、情報管理・情報学の専攻を選びました。

北京大学での学業の負担は思ったより大きくありませんでした。同期生たちが文学や音楽、芸術に没頭している間、李彦宏は情報検索とコンピューターシステムという目標だけに集中しました。専攻科目を着実に修めるかたわら、情報管理学だけでは物足りないと感じ、計算機科学科の正規講義にも潜り込んでプログラミングの実力を磨きました。当時の中国の大学界では海外留学ブームが激しく巻き起こっており、李彦宏もより進んだコンピューター工学の環境を求めてアメリカの大学院へ進むという目標を固めていきました。

1991年、北京大学で学士号を修めた後、渡米の準備を進めましたが、その過程は決して平坦ではありませんでした。中国の大学の情報管理学に正確に対応するアメリカの大学院専攻がなく、出願書類には計算機科学へと方向を変えて記さねばならず、そのために奨学金を得るのは容易ではありませんでした。学校を選べる立場にないことを受け入れた彼は、奨学金を出してくれる大学を優先し、二十校を超える大学に出願書を送りました。結局、コンピューター科学分野で全米三十位前後に位置するニューヨーク州立大学バッファロー一校から入学許可と奨学金を得て、アメリカへと旅立ちました。

バッファローで迎えた留学生活は想像以上に苦しいものでした。アメリカの大学院生たちはコンピューターと高速通信網を日常的に使いこなしながら学んでいましたが、中国から渡ったばかりの李彦宏には言語の壁と専攻転換による知識の格差、その上、手続きの遅れで一学期を逃し、出遅れて合流したという不利まで重なっていました。最初の学期は見知らぬ土地での孤独な時間でしたが、北京大学時代から身についた聴講の習慣と真面目さで、眠る時間を除けばコンピューター室で過ごしました。二年目に差しかかる頃には、同僚の研究者たちの間でコンピューターアーキテクチャ設計に明るい学生として名を知られるようになっていました。

その頃経験したある出来事は、長く彼の心に残りました。学内のある研究チームに加わるために面接を受けた際のことです。彼の発表と受け答えを聞いた担当教授が彼を上から下まで眺め回し、中国にもコンピューターはあるのかという趣旨の質問を投げかけたことがありました。その無礼な問いは長く消えない傷として残り、後に検索や人工知能の分野で、中国の名で世界に示せる技術を作らねばという決意へとつながったと伝えられています。

バッファローで二年半を過ごした李彦宏は、優秀な成績で博士資格試験に合格しました。しかし資格試験を終えた後、彼は狭い理論の検証に時間を費やす学者の道に疑問を抱き始めました。自分の作ったソフトウェアが実際に多くの人々の日常を変える現場での成果を渴望していると気づいた彼は、博士課程を途中で切り上げ、コンピューター科学の修士号だけを取得して産業界へと身移しました。

学位を終えた李彦宏は、ダウ・ジョーンズに技術コンサルタントとして入社し、ウォール・ストリート・ジャーナル・オンライン版のリアルタイム金融情報システムの設計を任されました。その後、プリンストンのパナソニック研究所を経て、シリコンバレーの検索

企業インフォシークでシニアエンジニアとして働き、今日の検索エンジンの基盤となるハイパーリンク解析手法と関連特許を開発しました。この時期に積み上げた技術と特許が、彼が1999年末に北京へ戻ってバイドゥを創業する土台となりました。

李彦宏の学生時代を扱う資料は、このように山西省の小さな街から始まり、北京大学とニューヨーク州バッファローを経てシリコンバレーへと至る軌跡を比較的詳しく残しています。ただし両親の名前や具体的な職業については記録が乏しく、こうした空白は今後新たな史料が見つければ埋められる余地として、正直に残しておくのが適切でしょう。

彼が検索エンジンという事業を選び取った背景には、幼い頃に太原の書店で経験した情報格差の記憶が色濃く滲んでいます。小さな街で生まれ育った子どもが、情報に触れる機会の違いゆえに大舞台で後れを取ったという自覚は、後に彼が情報を選び出し、見つけ出す技術に人生を懸けさせた最初の動機だったと見ることができます。バイドゥという名の会社が設立される前、その種はすでに山西省のある少年の頭の中で育っていたようです。

研究者から創業者へと固まっていく時期

中国人留学生の李彦宏が、ニューヨーク州バッファローにあるニューヨーク州立大学大学院の計算機科学科に足を踏み入れたのは1991年のことでした。北京大学で情報管理学を専攻した彼でしたが、アメリカの大学にはその専攻に正確に対応する大学院課程がありませんでした。そこで彼は専攻そのものを計算機科学へと完全に切り替えて出願し、事務手続きが遅れたことで他の学生より一学期遅れて授業に加わることになりました。

専攻を乗り換えた代償は大きなものでした。基礎知識の格差と言葉の壁が同時に彼を苦しめました。彼は眠る時間を除いたほとんどすべての時間をコンピューター室で過ごし、コードを書いては書き直しました。留学二年目に差しかかる頃には、研究室内でコンピューターアーキテクチャを扱う腕前が際立って進んでいたと伝えられています。

この時期、彼はある研究チームに加わるために口頭面接を受けました。面接を担当したアメリカ人教授は、彼が中国出身であることを確かめると、中国から来た学生にこの研究がこなせるのかといった調子で見下すような質問を投げかけました。その瞬間の屈辱感は長く李彦宏の心に残り、いつか中国本土からアメリカの企業と肩を

並べる技術企業を築くという決意へと固まっていったと言われていきます。

彼は大学院に入ってから二年半で博士資格試験にただ一度で合格しました。計算機科学の博士課程で最も突破が難しいとされる関門を越えたわけですが、まさにそれ以降、彼は狭く理論的な研究に疑問を抱き始めました。論文一本のために何年も費やす学界のやり方が、自分の望む道とは違うという考えが彼をとらえました。

彼は博士課程を離れ、修士号だけを取得して学校を後にしました。理論を打ち立てる人間よりも、理論を世の中に応用する人間になりたいという判断でした。この選択は、後に彼が検索技術を実験室の外へ持ち出し、実際のサービスとして育て上げる道につながる最初の分岐点でした。

大学院に在籍していた間、彼はニュージャージー州プリンストンにあるパナソニック研究所で研究インターンとして働きました。時給二十五ドルという、当時の留学生としては高い報酬を得ながら、情報検索の基礎的な実務を身につけました。このインターンシップは、彼が学校で学んだ理論を実際のシステム設計へと移し替える最初の経験でした。

修士号を修めた後、李彦宏はウォール街のダウ・ジョーンズに上級技術コンサルタントとして入社しました。1994年から1997年まで続いたこの時期、彼が任されたのはウォール・ストリート・ジャーナル・オンライン版のリアルタイム金融情報システムを設計・構築する仕事でした。秒単位で流れ込む相場とニュースを安定して処理しなければならないこのシステムは、その後も複数の金融機関で長く使われたと伝えられています。

金融情報システムを扱う中で、彼は大量のデータを速く安定して処理する感覚を身につけました。秒単位で値が変わる株価情報をリアルタイムで選別し並べ替える作業は、のちにインターネット全体から必要なページを選別し並べ替える検索技術の問題と、根本において似通っていました。

1997年、彼はシリコンバレーに活躍の場を移し、インフォシーク(Infoseek)とその系列会社であるGO.COMでシニアエンジニアとして働き始めました。インフォシークは当時、黎明期のインターネット検索を牽引していた会社の一つで、彼はここで検索性能を引き上げる中核技術を次々と開発しました。

彼がインフォシークで作り上げたESP技術は検索速度を大きく向上させ、GO.COMでは画像検索エンジンを自ら完成させました。文字でできた文書を探すことと画像を探すことはまったく異なる問題でしたが、彼はその両方の領域で実務設計者としての実力を認められました。

この時期、彼の名を検索技術の歴史に残したのは、超リンク解析(ハイパーリンク解析)という発明でした。インターネットには無数のウェブページが散らばっており、その中のどのページが本当に信頼でき有用かは、文書内の単語を数えるだけではわかりません。彼はページとページの間に張られたリンク、すなわちあるページが別のページをどれほど頻繁に、どの程度の信頼度で引用しているかを数学的に計算する方法を考案しました。

多くのリンクを受けるページ、それもすでに信頼されているページからリンクを受けるページに、より高い点数を与えるこの方式は、文書内の単語照合にとどまっていた既存の検索技術の限界を超えようとする試みでした。この特許は、のちのグーグルのページランクに先立つ発明の一つに数えられ、今日の大規模言語モデルが情報を選別し順位付けする方式にもつながる原型になったと評価されて

います。

李彦宏がシリコンバレーで学んだのは技術だけではありませんでした。ネットスケープやヤフー、インフォシーク、マイクロソフトといった会社が資金を集め、特許を守り、人材をめぐって争う様子を、エンジニアという立場から間近で見っていました。技術がどれほど優れていても、資本と組織が支えなければ会社は生き残れないという事実を、彼は現場で学びました。

1999年、彼はシリコンバレーで経験し観察したことをまとめて『シリコンバレー商戦』という本を出しました。技術企業がどのように資金を調達し、競争し、崩れ、生き残るのかを扱ったこの本は、中国国内で少なくない反響を呼びました。コードを書くエンジニアから会社全体を設計する人間へと移っていく彼の心構えが、この本にそのまま込められていました。

その後あるインタビューで、彼は自分をシリコンバレーで長く働いた人間だと紹介しながらも、自身のアイデンティティは教育を受けたエンジニアという言葉でまとめました。中国に戻ると決めたとき、周囲の友人たちは英語の情報を中国語に訳す翻訳サービスを一緒に作ろうと勧めましたが、彼はその提案には従いませんでした。

今後、中国語のコンテンツそのものが爆発的に増えていくという自分の判断のほうを信じたからです。

その判断は、翻訳ではなく検索という方向へ彼を導きました。中国語のコンテンツが増えれば、それを見つけ出す技術が必要になるという予測であり、これが彼がシリコンバレーで身につけた検索技術を故国で展開する理由になりました。

1999年末、李彦宏はアメリカでの生活を整理し、北京へ戻りました。ダウ・ジョーンズとインフォシークで積んだ実務経験、超リンク解析という発明、そして本を書きながらまとめた経営についての考えが、彼の手の中にもともにありました。翌年である2000年1月、彼は北京でバイドゥを設立しました。

大学院で理論研究を離れ、実務へと方向を転じた選択。ダウ・ジョーンズで身につけた大容量リアルタイム処理の感覚、インフォシークで完成させた超リンク解析技術は、それぞれ異なる時期に得た断片でしたが、バイドゥという一つの中で再び出会いました。検索という技術的な課題を実際に動くサービスへと作り上げた力は、この八年にわたる実務経歴から生まれたものでした。

バイドゥを設立した後、李彦宏が歩んできた道は検索会社の創業者にとどまりませんでした。彼はのちにバイドゥ・ワールド大会の演説で、人工知能の能力が企業内に内在化されるとき、知能はコストではなく生産性になると語りました。AI産業の構造が、チップが価値を独占していた形から、モデルと応用サービスがはるかに大きな価値を生み出す形へと転換しつつあるというのが彼の見立てでした。

この見立ては、彼が若い頃シリコンバレーでリンク構造一つによって検索の勢力図を変えた方式から、そう遠くありません。バイドゥがその後打ち出した言語モデルERNIE(アーニー)シリーズは、テキストと画像、映像を併せて扱う構造を備えており、複数の専門家ニューラルネットワークを使い分ける方式で性能と効率を同時に狙っています。

検索窓に入力された単語を処理していた若きエンジニアは、三十年近い歳月が流れたのちには、画像と文章を併せて理解するモデルを設計する会社の創業者になっていたわけです。その間には、大学院での挫折と決断、ウォール街で身につけたリアルタイム処理技術、シリコンバレーで発明したリンク解析、そして帰国という選択が

、順に置かれていました。

会社と研究組織を作り上げた過程

帰国を決意した頃、シリコンバレーの友人たちは、英語で書かれた膨大なインターネット情報を中国人が読めるように訳す自動翻訳サービスをやってみてはどうかと勧めました。李彦宏はこの助言をそのまま受け入れませんでした。彼は今後、中国語のコンテンツそのものが爆発的に増えていくと見通し、翻訳ではなく、そのコンテンツを見つけ出す検索エンジンが必要だと判断しました。

上場後も李彦宏は、毎年総売上高の15パーセントほどを研究開発に固定的に投じるという原則を守りました。これは年間100億元、韓国ウォンにして1兆8000億ウォンを上回る規模の研究資金でした。これほどの比率を売上に対する研究開発予算として維持することは、中国の技術企業の中でも稀な事例に属します。

人材確保においても李彦宏は独特の方式を用いました。シリコンバレーですでに実力を証明していた研究者を北京の研究所に招く際には、現地の給与に15パーセント上乗せして支給しました。逆に北京のエンジニアがシリコンバレー支社への異動を望む場合は、基本給から15パーセント引き下げて調整しました。中国本土の研究能力

を守りながらも海外の人材を引き寄せようとするこの給与構造は、李彦宏がインフラの完成度は結局のところ人材密度から生まれると信じていたことを示しています。

ERNIE 4.5は、テキストと画像、動画を併せて受け取り、テキストを出力する構造として設計されました。テキストと視覚情報を異なる専門家グループに分けて処理する異種混合エキスパート構造を用いながらも、すべてのトークンが共有エキスパートと自己注意層を通過するようにし、二つのモダリティが互いを補強し合うようにしています。視覚入力には解像度が可変のビジョンエンコーダーを取り付けます。その出力をテキストと同じ埋め込み空間へ移すアダプターを設け、時間と縦横の位置を別々にエンコードする3次元位置埋め込みを適用しました。

このシリーズには、有効パラメータ47Bと3Bを用いる混合エキスパートモデル群と、総パラメータ424Bに達する最大モデル、そして0.3B規模の密モデルがともに含まれています。バイドゥの研究陣はPaddlePaddle(パドルパドル)フレームワークでこれらのモデルを学習させ、最大モデルの学習過程で47パーセントの演算使用率を達成したと明らかにしました。すべてのモデル重みと開発ツールを

Apache 2.0ライセンスで公開しました。

2026年2月に公開されたERNIE 5.0の技術報告書を見ると、この研究組織はいまや少数のスター研究者に依存する構造ではなく、数百人が水平的に協業する形へと定着しています。統括と諮問はワン・ハイフォン、ウー・ホア、ウー・ティエンが担い、アルゴリズムと学習を主導する役割はウェイ・スン、ジン・リウ、ディエン・ハイユー、イエン・ジュンマーらが担当します。彼らの下で400人を超える研究者が、データ構築、事前学習の最適化、評価体系の構築に分かれ、毎週モデルの性能を引き上げています。

2026年に公開されたERNIE-Imageの技術報告書は、バイドゥがテキストを超えて画像を生成する研究組織を別途立ち上げたという事実を示しています。このチームはジアシャン・リウが率い、視覚データを精製する人員としてジダ・フォンとポンユー・ジョウ、ジェンウェイ・チェンらがあり、ポストトレーニングとアラインメント作業はジェンウェイ・チェンとティエンルイ・ジュが担います。画像の質感や彩度、マーカーペンや水彩画のような画風を仕上げる美学専門グループも、ジアシャン・リウを中心に別途運営されています。

1999年、シリコンバレーを離れ、中関村の小さな部屋から始まったバイドゥは、こうした過程を経て検索エンジン会社からフルスタックの人工知能企業へと姿を変えました。その間には、15人で始まった「稲妻計画」から数百人規模のERNIE研究軍団に至るまで、李彦宏が資本と人材を絶えず研究組織に引き込んできた流れが続いています。

中核モデルと開発過程

中国のAI企業がそれぞれ言語モデル競争に乗り出していた時期。バイドゥの共同創業者であり最高経営責任者である李彦宏会長は、汎用人工知能へ至る道はテキストひとつだけでは開かれないと早くから判断していました。彼はテキストと画像、音声、映像を一つの知能体系の中に束ねなければならないと考え、そのためバイドゥは長らく先行的な資本投入を続けてきました。

リー・イエホン会長がとりわけ力を入れた応用分野は自動運転でした。彼は視覚大規模モデルの最大の価値が、もっともらしい映像を作り出すことにあるのではなく、機械が実際の道路上の世界を読み取り、計算し、その場で走行を判断することにあると考えていました。この考え方のもと、バイドゥの研究開発部門は自動運転インフラを世界的な水準まで引き上げました。

この方向性は数字にも表れています。バイドゥは毎年、総売上高の15パーセント前後、年間最低100億元以上を研究開発に固定的に配分し、事前学習と推論インフラを支えてきました。

バイドゥの人工知能研究の中心には、大規模言語モデル群であるERNIE、中国語で文心と呼ばれるシリーズがあります。バイドゥは2025年から2026年にかけてERNIE 4.5と5.0を相次いで発表し、大規模モデル市場で先頭の座を守りました。

2025年6月29日に公開されたERNIE 4.5は、十種類の異なるバリエーションを包括するモデル群として設計されました。推論時にはトークンごとに470億個、あるいは30億個のパラメータのみが起動する混合エキスパート構造を基本構造とし、最大のモデルは累積パラメータ数が4,240億個に達しました。低遅延のローカルサービングを狙った3億個規模の小型高密度モデルも併せて用意されました。

ERNIE 4.5の工学的な特徴は、多モーダル異種混合エキスパート構造にあります。異なる入力モダリティ同士がパラメータを共有しながらも、各モダリティ固有の表現力のために別途確保された専用エキスパートパラメータを併せて配置する方式です。こうすることで、言語モデルの文脈推論能力を維持しながら、視覚理解能力も同時

に引き上げることができました。

学習はバイドゥが独自に開発したオープンソースのディープラーニングフレームワーク、PaddlePaddle上で行われました。エンジニアたちは事前学習段階のハードウェア分散制御を緻密に磨き上げ、最大のERNIE 4.5言語モデルを学習させる過程でモデル演算使用率47パーセントを記録しました。

2026年2月4日には、次世代モデルであるERNIE 5.0の技術報告書が公開されました。

ERNIE 5.0は、テキスト、画像、動画、音声を最初から一つの自己回帰方式で共に学習するよう設計された基盤モデルです。従来のモデルが言語モデルにモダリティごとのデコーダーを付け加える後融合方式を採っていたのに対し、ERNIE 5.0はすべての入力を共有トークン空間に移して一つのシーケンスとして扱い、次のトークン群を予測する一つの目的関数のもとで共に学習させます。

このような統合方式を採用した理由は、後融合設計に現れるいわゆる能力シーソー現象にありました。一方のモダリティの性能を上げると、もう一方が下がるという問題が繰り返し起きていたのです。

。最初からすべてのモダリティを共に学習させれば、こうしたトレードオフなしに複数の能力を並行して伸ばせるという説明です。

構造の基本骨格は超希薄混合エキスパート方式です。ルーティングを決める基準は、そのトークンがどのモダリティに由来するかではなく、統合されたトークン表現そのものであり、そのため異なるモダリティのトークンが同じエキスパートプールへ共に送られます。活性化比率は3パーセントを下回り、演算負荷を大きく増やすことなく実効容量を拡大できました。

テキスト生成は、次のトークンを予測する標準方式に多重トークン予測技法を加え、品質と速度を両立させました。視覚と音声の生成は、トークン群を予測する問題として再定義されました。視覚側は次のフレームとスケールを予測する方式を、音声側は次のコーデックを予測する方式を採用しています。

学習過程で目を引くのは、弾力学習と呼ばれる新しい方式です。一つの事前学習過程の中で、深さと幅、ルーティングの希薄度がそれぞれ異なる複数のサブモデルを共に学習させ、メモリや時間の制約が異なる環境に合わせて選んで使えるモデルファミリーを一度に得る方式です。評価の結果、活性化パラメータの53.7パーセントと

全パラメータの35.8パーセントのみを使っても、完全なモデルに近い性能が維持されました。

事前学習の後には、教師あり微調整と統合マルチモーダル強化学習を併用する多段階の事後学習が続きました。複数のモダリティが混在するデータと超希薄構造が重なると学習が不安定になりやすいのですが、バイドゥの研究チームは偏りのないリプレイバッファと多段階重要度サンプリング、ヒントに基づく強化学習といった技法でこの問題を解決したと説明しています。

インフラ面では、4,000億個を超えるパラメータを扱うために、きめ細かなメモリ制御を伴うハイブリッド並列化を採用しました。トークナイザーを混合エキスパートバックボーンから切り離して別ノードに配置し、各構成要素がそれぞれに適した並列化方式を選べるようにしました。学習とロールアウト、環境とのインタラクションを共に調整する分散型強化学習インフラも新たに設計されました。

ERNIE 5.0の開発には、王海峰、呉華、呉甜をはじめとするバイドゥの最高科学者たちが率いる大規模な研究陣が参加し、于遜、景琉、田海宇、燕俊瑪らが中核アルゴリズムと学習過程を調整しました。学習データには多言語のウェブコーパスと数学推論ベンチマーク

、コードデータが共に含まれていました。

同時期、バイドゥは画像生成の分野でも別途の専門モデルを発表しました。2026年4月15日に公開されたERNIE-Imageは、テキストとレイアウト情報を受け取り、グラフィックデザイン水準の画像を生成する80億個パラメータのモデルです。ここから派生した高速版ERNIE-Image-Turboも同時に発表されました。

従来の拡散モデルは、精緻な画像一枚を得るために数十回のノイズ除去過程を経る必要がありました。バイドゥの研究チームは分布マッチング蒸留技法と強化学習を組み合わせ、ERNIE-Image-Turboがわずか八段階の推論だけで、細部の描写が崩れない画像を生成できるようにしました。

画像の品質を解像度ではなく人が感じる美的完成度まで引き上げるために、劉佳翔、周鵬宇、陳振宇らが率いる別の研究組織が組みられました。彼らは水彩画風の絵、色マーカーで描いた手描き画、蛍光ペンの質感、図表レイアウトといった細かいデータをモデルに組み込みました。

長い文章を画像の中に正確なタイポグラフィで配置する能力を測るLongText-Bench評価において、ERNIE-Imageは英語部門で満点、中国語部門で0.96、総合で0.89を記録し、高速版のERNIE-Image-Turboも英語満点、中国語0.96、総合0.87を出しました。これはByteDanceのSeedream 4.5やGoogleのNano Banana 2.0といった競合モデルと肩を並べられる水準でした。

文書処理の分野では、PaddlePaddleチームがPaddleOCR-VLを発表しました。低消費電力の組み込み機器にもそのまま搭載できるよう設計された9億個パラメータのビジョン言語モデルで、画面比率や解像度がまちまちでも歪みなく受け付けるNaViT方式の動的解像度エンコーダーと、バイドゥの小型言語モデルERNIE-4.5-0.3Bを一つにつなげたものです。

このモデルは、レイアウト分析と要素認識を分けた二段階学習を経ました。リアルタイム物体検出モデルRT-DETRにポインターネットワークを組み合わせたPP-DocLayoutV2により、文書内のテキストボックスや表、チャート、数式の位置と読む順序を把握しました。ビジョンエンコーダーはKeye-VLの重みで、言語モデル部分はERNIE-4.5-0.3Bの重みでそれぞれ初期化した後、再調整する方式が採

られました。

学習データの作成過程も複数の段階を経ました。まず軽量文書解析ツールPP-StructureV3で2,000万件に及ぶ画像とテキストのペアに仮ラベルを付けた後、これをERNIE-4.5-VLと競合モデルQwen2.5-VLに入力し、誤字や配置の誤りを再度修正するクロス検証を行いました。

表認識を磨く過程では、ERNIE-4.5-VLが表をHTMLに変換する仮ラベルを作成し、280億個パラメータの判別モデルERNIE-4.5-VL-28B-A3Bがその結果の整合性を自動的に選別しました。ここで選別された不良な表はDianJin-OCR-R1という専用ツールで再度追跡し、修正しました。これとは別に、学術論文や財務報告書、ウェブページから集めた80万件以上の中国語・英語のチャートデータも学習に用いられました。

配備性能を試験した結果では、資源の限られたRTX 3060でもvLLMアクセラレーションで每秒2,694.1件、SGLangアクセラレーションで每秒2,765.2件の処理量が出て、より高性能なRTX 4090Dではそれぞれ851.9件と932.4件を記録しました。実際のサーバー環境でも十分に耐えられるコストで運用できるということでした。

ERNIE 4.5の異種混合エキスパート構造、ERNIE 5.0の統合自己帰帰方式と弾力学習、ERNIE-Imageの高速拡散蒸留、PaddleOCR-VLの超軽量文書パースは、それぞれ異なる課題を解決するための設計でしたが、一つの方向に収斂します。言語と視覚、音声を別々に作るのではなく、一つのバックボーンと一つの学習体系にまとめ上げようとする試みです。独自フレームワークPaddlePaddleと独自チップインフラの上でこれらのモデルを共に動かす構造は、リー・イエンホン会長が描いてきたフルスタック人工知能企業の実際の姿に近いといえます。

現在の立ち位置と中国AIにおける意味

2026年の李彦宏会長は、今もバイドゥを作り、率いる人物です。創業者であり取締役会議長であり、最高経営責任者でもある彼は、会社の戦略と経営を一手に握っています。1990年代のシリコンバレーでハイパーリンクを解析して検索結果の順位を決める技術を独力で作り上げた経歴の持ち主らしく、今も技術そのものを手放さずに会社を率いるやり方を貫いています。

バイドゥが長く推し進めてきた戦略は、検索会社ではなく、人工知能のあらゆる階層を自ら作り上げる会社になることです。演算を高速化する自社チップから、そのチップを搭載して動かすクラウド

、モデルを訓練する自社フレームワーク、その上で動く言語モデル ERNIE、そして人々が実際に使うアプリケーションまでを一本の線でつなげた構造です。この構造は社内でしばしばフルスタックと呼ばれています。

このフルスタック戦略の最新段階は、自社設計の人工知能チップ事業を切り離し、香港市場に上場させる作業です。会社はこのチップ資産の価値を市場に直接評価してもらおうとしており、上場申請の手続きは計画通り進んでいるとされています。研究開発組織には、あらかじめ定められた枠内で演算資源を自由に使える権限を与え、決裁を待つ時間を無駄にしないよう組織を整えています。

李彦宏会長が守り続けてきた原則の一つは、売上のかなりの部分を研究開発に固定的に投じることです。彼は2018年、CGTNが設けた場で、会社が毎年売上の15パーセント前後を研究開発に投じていると明かしました。この比率は、年間最低100億元、日本円にして2000億円を超える規模の固定支出につながっています。

彼は社内だけで知られる人物ではありません。2015年に海南で開かれたボアオ・フォーラムでは、ビル・ゲイツとイーロン・マスクを直接招き、技術の未来をめぐる対談を主導しました。検索エンジ

ンを作ったエンジニアが世界的な起業家たちと肩を並べて質問を投げかける場でした。中国の技術業界を代表して舞台に立つ役割を、彼が長く担ってきたことをよく示す場面です。

彼が思い描く目標は、有名になることとは少し違います。彼はあるインタビューでこう語りました。世界中の多くの人が毎日自分たちのサービスを使っている、道で自分に気づかれなくてもかまわない。パリの市民も、東京の利用者も、サンフランシスコのエンジニアも、毎朝バイドゥの知能型サービスを自然に立ち上げる暮らしこそ、自分が望むものだと言っています。

ここ数年でバイドゥが送り出した成果は、このフルスタック構想が単なる言葉ではないことを示しています。言語モデルERNIEの新世代であるERNIE 5.0が発表されました。画像を生成するERNIE-Image、文書を読み取るPaddleOCR-VLも相次いで公開されました。三つとも、バイドゥが長年積み上げてきた自社フレームワーク「パドルパドル(PaddlePaddle)」の上で動くという共通点があります。

ERNIE-Imageは、美しい画像を生成する能力と、文字を正確に描き込む能力に重点を置いたモデルです。パラメータ数80億規模で、

人が見て自然な完成度と、画面内の文章をくっきりと配置する能力を同時に狙って設計されています。

このモデルを高速に動かすため、バイドゥの研究陣は分布マッチング蒸留という手法と強化学習を組み合わせで用いました。大きなモデルの描き方を短い工程に凝縮する方法といえるもので、その結果生まれたERNIE-Image-Turboは、わずか8ステップだけで細部が崩れない高精細な画像を生成します。

長い文章を画像内にどれだけ正確に描き込めるかを測るLongText-Benchという評価で、ERNIE-Imageは英語文で満点の1.00点を、中国語文で0.96点を獲得し、総合点は0.89点でした。バイトダンスのSeedream 4.5が総合0.988点、Zhipuの GLM-Imageが0.966点を記録したとと並べてみると、バイドゥが最上位圏で真っ向勝負を挑んでいることがうかがえます。

文書読み取りの分野では、PaddleOCR-VLが登場しました。パラメータ数9億規模とコンパクトに抑える一方、組み込み機器や低消費電力のモバイル端末でもそのまま動くよう設計されたモデルです。大きなモデルが入らない場所に文書認識機能を組み込もうという狙いがはっきりしています。

この小型モデルには、解像度を状況に応じて調整するNaViT方式の視覚認識器と、ERNIE 4.5系列のうち3億パラメータに縮小した言語モデルが一体化して組み込まれています。二つの機能を一つにまとめることで、109言語のオフライン文書まで読み取れるようにしています。

このモデルの訓練には、画像と文字を対にして自動で説明を付けた2000万組のデータと、表やグラフを含む80万件の高品質データが使われました。文書内の表やチャートまで読み取るには、通常の文字よりもはるかに精緻な学習データが必要だったためです。

実際に動かした結果も併せて公開されています。SGLangという高速化エンジン上で動かした場合、RTX 3060グラフィックカードでは毎秒2765.2件、RTX 4090Dでは毎秒932.4件の処理量を記録しました。高価なサーバー向け機器ではなく、一般的なグラフィックカードでも実用的な速度が出ることを数値で示した形です。

こうした発表は、バイドゥ開発者会議のような場を通じても継続的に行われてきました。2024年4月、北京で開かれた開発者会議では、バイドゥの人工知能の成果が2時間半以上にわたって紹介され

ました。ERNIE系列のモデルとその応用事例を幅広く取り上げた場
でした。

バイドゥが他の大規模モデル系スタートアップと異なる点は、アル
ゴリズム研究所の水準にとどまらないところにあります。チップ
、サーバー、フレームワーク、モデル、応用まで全階層を社内で直
接コントロールできる数少ない企業の一つです。この構造を築くま
で、李彦宏会長はトランスフォーマー理論が初めて登場してから生
成AIが本格的に定着するまで、かなり長い時間を耐え抜いてきまし
た。

この自立型の構造は、中国国内で政策的な意味も併せ持っていま
す。米国製チップやソフトウェアへのアクセスがさまざまな規制で
制限される状況下で、バイドゥが自社チップと自社フレームワーク
で国内の産業データを処理できるという事実は、国家レベルの技術
的自立を支える事例としてしばしば取り上げられます。他の企業が
外国製ツールを迂回して取り入れる一方、バイドゥは自社設計のチ
ップを大量生産する道を選びました。

もっとも、こうした成果が実際の売上にどれだけつながったかは
、別途見極めるべき問題です。バイドゥの株価は、一時1年前より4

0パーセント以上上昇する動きを見せました。この上昇分には、実際の売上増加と市場の期待の両方が混じっていました。

検索画面をAI中心に切り替える作業でも、バイドゥは自分たちが世界のどのポータルよりも先行していると主張してきました。トップ画面に写真や動画、リアルタイム情報といった多様な形式の結果を表示する割合を大きく引き上げたという主張です。しかし実際にどれだけの利用者が集まっているかを示す指標では、バイトダンスの豆包(ドウバオ)が依然として先行しており、アリババの通義千問(トンイーチェンウェン)がその後を追う形になっていて、検索画面の変化がそのまま利用者数の逆転にはつながっていません。

こうした差の背景には、中国の技術企業間のデータの壁もあります。テンセント、アリババ、バイトダンス、バイドゥといった大手企業は、それぞれ自社が集めた学習データを互いに共有せず、個別に蓄積する道を選んできました。バイドゥが自動運転データや膨大な書籍データを自前で集めてERNIEを訓練しているのも、こうした閉鎖的な環境の中で起きていることです。

バイドゥの財務を担う役員でさえ、この閉鎖的な構図について、異なる陣営が互いにデータを十分に開放していないと率直に認めた

ことがあります。こうした環境の中でERNIEが特定分野で他社よりデータが薄くなる問題をどう埋めていくのかは、今後も見守るべき点です。

結局、李彦宏会長が今立っている場所は、検索エンジンを作ったエンジニアが30年近く歩んできた道の延長線上にあります。チップを自ら作り、クラウドを自ら築き、モデルを自ら訓練し、その上に画像認識や文書認識といった実際に使われる道具までを載せた今のバイドゥは、彼が最初に検索アルゴリズムの特許を出願した時に描いていた方向が、大きく広がった姿だといえます。

これから残された問いは、このフルスタック構造が実際の商業的成果と利用者数の逆転につながるかどうかです。自社チップの上場が完了し、ERNIE系列のモデルが市場でより広く使われるようになったとき、バイドゥが今の差を実際に縮められるかどうかは、今後数年の実績と利用者指標を通じて明らかになる問題として残っています。

第3章 唐傑(タン・ジエ、Jie Tang) - 清華大学知識工学からGLMへとつながった研究者の道

成長期と学業

唐傑は1977年、四川省南充市で生まれた。家庭は決して裕福ではなく、幼い頃に抱いた目標もほかの子どもたちと変わらなかった。勉強に励んで良い学校に入り、良い職場を得ること、それがすべてだったと彼は後に振り返っている。

しかし彼は試験の点数だけを追う子どもには育たなかった。物理が好きで、電子回路や無線通信に心を奪われた。父親は決して豊かではない家計の中でも、息子の好奇心のために大金を投じてプログラム可能な学習機を買い与えた。唐傑はこの機械を分解しては組み立て直し、コードを書き換える過程でコンピューターという世界に初めて出会った。

1992年、彼は南充市の重点高校である蓬山中学に入学した。当時担任を務めた李中耀教師は、後に彼をこう振り返っている。唐傑はどんなことにも全身全霊で打ち込む生徒でした。試験の点数そのものよりも、知識の原理を最後まで突き詰めて完全に理解することに重きを置いていたという。

成績よりも理解を重んじる姿勢は、大学進学においても続いた。1995年、彼は中国の重機械分野の名門として知られる燕山大学に入学し、同校の看板学科である自動化を専攻した。機械制御という伝統的な工学を学ぶ一方で、彼の関心はほかのところにあった。真に興味を抱いたのは、ソフトウェアを扱うプログラミングだった。

結局、彼は修士課程で専攻をコンピューター分野へと移した。自動化という物理的な制御システムを扱っていた学生が、ソフトウェアとアルゴリズムの世界へと移った形である。この転換は、後に彼がハードウェアとソフトウェアの両方を理解する研究者へと成長する足がかりとなった。

2002年、唐傑は清華大学コンピューター科学技術専攻の博士課程に進んだ。この頃インターネット産業はすでに成熟しつつあったが、人工知能分野は依然として純粋理論の段階にとどまっており、今日のような大規模言語モデルという概念は存在しなかった。彼はこの時期、セマンティックウェブとデータマイニングを研究方向に定めた。

なぜよりによってデータマイニングだったのか。唐傑は、自然言語処理やセマンティックウェブといったほかの分野に比べ、データ

マイニングは新しいものに最も開かれ、受け入れる度量のある学風を持っていたと説明した。ちょうど中国国内でソーシャルネットワークサービスが芽生え始めた時期で、機会も多かった。2006年時点でこの分野を研究する人は国内にほとんどいなかったのだから、彼の眼力は時代を先取りしていたといえる。

2006年、唐傑は博士号を取得した。その年、フェイスブックは全世界での会員登録を開放し、ヤフーからの10億ドルの買収提案を拒否した。グーグルは設立からわずか1年のユーチューブを16億5000万ドルで買収した。中国国内でも電子商取引とオンラインソーシャルサービスが急速に規模を拡大していた。周鴻禕は奇虎360を設立し、ウイルス対策ソフト市場の競争に火をつけた。

清華大学コンピューター科学博士という肩書きがあれば、お金を稼ぐ道はいくらでもあった。シリコンバレーの大企業で高い年俵を得ることもできたし、国内大手企業でストックオプションを得て腰を据えることもできた。しかし唐傑は最も陰しい道を選んだ。清華大学に残り、研究を続ける道である。

この選択の背景には、ある人物の存在があった。漢字レーザー植字技術を生み出し、国家最高科学技術賞を受賞した北京大学方正グ

ループ創業者、王選である。唐傑は博士課程修了を控えた頃、学内行事で王選と出会った。1970年代から80年代にかけて、ソフトウェアとハードウェアの環境が海外に大きく後れを取っていた時代。王選が海外の40年の蓄積を一気に追いつき、漢字レーザー植字という難題を解き明かした話に、唐傑は深く感銘を受けたという。

唐傑が卒業する直前の2006年2月、王選はこの世を去った。唐傑が清華大学に残り研究を続けることを決めたのは、後から振り返れば王選が残した精神を受け継ぐことでもあった。ただし彼が受け継いだのは、レーザー植字ではなくソーシャルネットワーク分析とデータマイニングという新たな領域だった。

清華大学の教授として学問の道を歩むことを決めたものの、目の前の現実は厳しかった。人工知能研究分野には自然言語処理、セマンティックウェブ、データマイニングがあったが、唐傑が選んだデータマイニング、それもソーシャルネットワークを扱う研究は、当時の学界では誰も理解できない未知の領域だった。研究費を申請しても承認される見込みはなかった。

手立てのなかった唐傑は、自身が受け取った博士論文賞の賞金2万元を取り出し、研究プロジェクトの立ち上げ資金に充てた。成果

物を作って示さなければ、研究費が承認される可能性が生まれると判断したためである。この資金で彼は清華大学知識工学研究室、すなわちKEGのチームを組み、科学技術情報ビッグデータマイニング及びサービスシステムプラットフォームを作り上げた。名前はAMinerで、完全に独自の知的財産権を持つシステムだった。

AMinerは2006年12月に公開されるやいなや学界の注目を集めた。ある海外の名門大学は、正教授としての終身在職権と制限のない研究費予算を提示して彼を招こうとしたが、唐傑はこれを断った。当時の海外研究環境のほうが優れていたのは事実だが、国内で世界一流の研究を成し遂げたいという理由からだった。

その後もAMinerは着実に規模を拡大していった。220の国と地域の21万の企業及び機関がこのプラットフォームを利用し、登録された固有ユーザーだけで3000万人を超えた。研究データのダウンロード回数は230万回を超え、中国工程院や科学技術部はもちろん、グーグル、マイクロソフト、ハーバード、スタンフォード、MITといった機関もこのシステムを利用者名簿に載せた。

この時期、唐傑が受けた学術賞だけでも数多い。国家科学技術進歩賞2等賞、北京市科学技術進歩賞1等賞、北京市発明特許賞1等賞

、中国人工知能学会科学技術進歩賞1等賞を受賞し、最も権威ある国際学会で10年間の最優秀論文賞を受賞した国内初の人物となった。

2017年5月、アルファ碁が世界囲碁チャンピオンの柯潔を3対0で下し、その1か月後にトランスフォーマーモデルが発表された。今日のように世界を完全に一変させる出来事ではなかったが、学界を揺るがすには十分だった。翌2018年には、キングソフト元CEOの張宏江が北京智源人工知能研究院、すなわちBAAIを設立し、北京の人工知能人材を集め始めた。

2019年は唐傑の人生において一つの節目となった年だった。彼は清華大学コンピューター学科の教授に昇進した。BAAIの学者として名を連ね、清華大学とBAAIの支援を受けてAMinerを分社化し、智譜AIを立ち上げた。学者から起業家への転換であり、目標は汎用人工知能だった。

当時、人工知能分野で熱く注目を集めていた方向は、顔認識のようなコンピュータービジョンだった。曠視、商湯、依図、雲従といった企業が、商業化の速いこの分野で先を行っていた。ところが唐傑は自然言語処理と知識グラフに集中すると宣言した。学問には明

るくても事業には疎いのではないかという懸念がついて回った。

この判断は翌年、正しかったことが明らかになった。2020年、OpenAIが1万枚のカードで構築したGPT-3を発表し、言語の生成と理解の両方が可能な大規模モデルの時代が幕を開けた。唐傑が長年積み重ねてきた自然言語処理とデータマイニングの基盤は、すぐさま真価を発揮した。彼が率いたBAAIの悟道プロジェクトが、その最初の結実だった。

清華大学時代から積み重ねてきた研究の蓄積は、後にKEG研究室の名で続いたさまざまな学術成果においても確認できる。GLM系列モデルの初期論文、コード生成モデルCodeGeeX、画像生成モデルCogViewといった成果は、いずれもこの研究室から生まれ、唐傑はこれらの研究の中心にいた。

清華大学に残り研究を続けることを決めた2006年の決断、2万元から始まったAMiner。そしてそれに続く学術賞と国際的な評価は、一本の軌跡として結びついている。学者として生き残ることが難しかった時代の選択が積み重なり、後に智譜AIとGLMシリーズへとつながる土台を築いたと見ることができる。

研究者、あるいは起業家として定まっていく時期

1977年、唐傑は四川省南充で生まれました。家庭の事情は平凡で、目標もほかの子どもたちと変わりませんでした。勉強に励んで良い学校に進み、良い職場を得る道でした。ただ彼は幼い頃から物理が好きで、プログラミングや無線通信をいじることを楽しんでいました。父親は教育費を惜しまず、思い切ってプログラム可能な学習機を買い与えました。この品物は少年の手から長い間離れることがありませんでした。

1992年、彼は南充市の名門高校である蓬安中学に入学しました。担任を務めた李中耀教諭は、後にタン・ジエについて、何事にも全身全霊で打ち込む生徒だったと振り返っています。試験の点数よりも概念を完全に理解することに心を注いでいたといいます。点数にさほどこだわらなかったこの生徒は、1995年、中国の重機械分野の名門とされる燕山大学の自動化学科に入学しました。

機械と電気信号を扱う自動化の勉強は、彼の心を完全にはとらえませんでした。むしろ彼を引きつけたのは、目に見えないコード、すなわちプログラミングでした。彼は大学院進学過程で専攻をコンピュータに変え、2002年には清華大学コンピュータ科学技術専攻の博士課程に入り、セマンティックウェブとデータマイニングの

研究を始めました。この頃の博士課程の同期には、後にZhipu AIの最高経営責任者となる張鵬もいました。二人はここから長い縁を築いていくことになります。

2002年の人工知能研究は、今とはまったく事情が異なっていました。世界の情報技術産業はまさに膨張を続けていましたが、人工知能は依然として一部の学者による理論的関心事にとどまっていました。タン・ジエはこうした厳しい環境の中で、コンピュータがウェブ上の膨大な非定型テキストを自ら読み取り、その中の関係を掘り起こすデータマイニングを自身の研究方向に決めました。彼はこの分野が新しいものに最も開かれており、ちょうど中国国内でソーシャルネットワークサービスが芽生え始めた時期であったため、機会が多いと判断しました。

博士課程の修了を控えた2005年末から2006年初めにかけて、タン・ジエは学校の行事で王選と出会いました。王選は、1970年代から1980年代にかけて中国のソフトウェアとハードウェアの環境が海外に大きく遅れをとっていた時代に、独自の研究室の力だけで漢字レーザー植字技術を作り出した人物でした。北京大学の方正グループを創設した彼は、国家最高科学技術賞を受賞した科学者でもあ

りました。

王選は、海外の四十年の蓄積を一気に飛び越え、世界的な難題を解き明かした人物であり、タン・ジエはその話に深い刺激を受けました。タン・ジエが後に語った言葉を借りれば、技術は天を突くほど高くなければならず、それを市場にしっかりと根づかせなければならぬという考えが、このとき彼の中に刻まれました。王選は2006年2月、タン・ジエの卒業を数日後に控えた頃に世を去りました。清華大学に残って研究を続けるというタン・ジエの選択は、こうした意味で師の意志を継ぐことにほかなりませんでした。

清華大学のコンピュータ科学博士として、タン・ジエが金を稼ぐ道を探していたなら、それほど難しくはなかったはずです。シリコンバレーの大企業の高い年俸も、国内大企業のストックオプションも彼の前に用意されていました。ある海外の名門大学に至っては、終身准教授の地位と上限のない研究費予算を提示して彼を招こうとしました。しかし彼は、海外の研究環境の方が優れていると認めながらも、中国国内で世界一流の研究を成し遂げるという理由からその提案を断り、清華大学に残りました。

彼が受け継いだのは、レーザー植字ではなくソーシャルネットワーク分析とデータマイニングという方向でした。問題は、2006年当時、国内でこの分野を真剣に研究する人がほとんどいなかったことにありました。学界はこの研究を理解せず、研究費の申請はことごとく壁にぶつかりました。タン・ジエは結局、自身が博士論文優秀賞として受け取った2万元を取り出し、研究の元手としました。成果を作り出してこの研究が役に立つことを証明して初めて、その後正式に研究費を受け取る道が開かれるはずでした。

まさにこの2万元で、彼はチームを立ち上げ、科学技術情報ビッグデータ分析システムAMinerを作り上げました。清華大学知識工学研究室、すなわちKEGが生み出したこの完全な独自知的財産権プラットフォームは、世に出るとすぐに学界を驚かせました。2006年12月に公開されて以来、AMinerは220の国と地域、21万の企業・機関、3千万人を超える利用者を集め、研究データのダウンロードは230万件を超えました。中国工程院と科学技術部はもちろん、グーグル、マイクロソフト、ハーバード大学、スタンフォード大学、マサチューセッツ工科大学もこのプラットフォームを利用しました。

。

この成果によりタン・ジエは、国家科学技術進歩賞2等賞、北京市科学技術進歩賞1等賞、北京市発明特許賞1等賞、中国人工智能学会科学技術進歩賞1等賞をはじめとする数々の賞を受賞し、国際学術会議において過去十年で最も優れた論文に贈られる賞を国内で初めて受けた人物となりました。AMinerを率いながら彼は、膨大なウェブデータをふるいにかけて、知識グラフを精製する作業を十年以上手放しませんでした。この過程で大容量分散システムを扱う感覚と、非定型データを読み解く目を磨いていきました。

2017年5月、アルファ碁が世界の囲碁チャンピオン柯潔を3対0で破り、その一か月後にトランスフォーマーモデルが発表されました。今日ほど世間を揺るがすものではありませんでしたが、学界を揺さぶるには十分な出来事でした。翌年の2018年、キングソフトの元最高経営責任者である張宏江が中心となって北京智源人工智能研究院、すなわちBAAIを設立しました。営利を追わないこの新しい形の研究機関には北京の人工智能人材が集まり、タン・ジエもこの研究者として名を連ねました。

2019年は彼の人生において方向が定まった年でした。この一年で彼は三つの新しい立場を得ました。清華大学コンピュータ学科の教

授に昇進しました。智源、すなわちBAAIの首席研究員となり、清華大学とBAAIの支援のもとAMinerを分社化してZhipu AIを設立しました。学者から起業家への転換の目標は、汎用人工知能、すなわちAGIを実現することでした。当時としてはあまりに大きな目標であり、周囲には容易に受け入れられませんでした。

2019年、人工知能業界で最も注目されていた分野はコンピュータビジョンでした。Megvii、SenseTime、Yitu、Cloudwalkといういわゆる人工知能四小龍は、いずれもこの分野で事業を展開し、商用化への道筋もはっきりしていました。マンション団地の入り口ごとに顔認証装置を設置するだけでも、数千億元規模の市場が開けていました。しかしタン・ジエは、Zhipu AIが自然言語処理と知識グラフに集中すると表明しました。学問には明るくても起業には不慣れだろうという見方がついて回りました。

結果はそうした懸念を上回るものでした。2020年、OpenAIが1万枚の演算カードで学習させたGPT-3を発表し、言語を生成し理解する大規模言語モデルの時代が幕を開けると、中国の人工知能学界は焦りに包まれました。相手が1万枚規模のクラスターを使う一方、こちらはサーバー一台のカード8枚で実験してきた立場だったため

、その差は絶望的なほどでした。BAAIは直ちに「百人巨大モデル計画」を始動させました。研究所が清華大学の五道口に位置することにちなみ、このプロジェクトに悟道という名を冠しました。中国の人工知能が後れをとった立場を追い上げようとするこの計画の責任を、タン・ジエが担いました。

2021年初め、悟道1.0が世に出され、三か月後に悟道2.0がこれに続きました。パラメータ数は1兆7,500億に達し、GPT-3の十倍近く、当時としては世界最高記録でした。この成功はタン・ジエ個人にとどまらず、清華大学出身の研究者たちを中国人工知能業界の中心勢力に押し上げる契機となりました。清華大学コンピュータ学科の終身教授である黄民烈は聆心智能を、准教授の劉知遠は面壁智能を、清華大学人工知能研究院副院長の朱軍は瑞莎科技を、それぞれ創業しました。

若い世代の動きも負けてはいませんでした。90年代生まれの博士課程生である祁瞻超は、テンセントが初めて出資した大規模モデル企業である深言科技を設立しました。そしてタン・ジエの教え子であった楊植麟は、2023年にMoonshot AI、すなわちKimiの親会社を創業し、評価額を100億ドルを超える水準にまで育て上げました。

楊植麟は様々な場で恩師への感謝を述べ、タン・ジエが金をどれだけ稼げるかよりも、世界最高水準の成果を出せるかどうかを常に強調していたと語っています。

悟道という土台の上で、Zhipu AIのGLMシリーズのモデルは三、四か月ごとに新版を発表しながら急速に発展しました。言語、コード、複数の形態のデータを併せて扱うマルチモーダル、さらに自ら業務を処理するエージェントまでを包括するモデル群を備えることで、Zhipu AIは技術路線においてOpenAIやAnthropicのような世界トップクラスの企業と肩を並べられる数少ない中国企業として位置づけられるようになりました。タン・ジエはこの全過程について、学問研究とは山を登るようなもので、頂に立てなければ失敗したのも同然だと語りました。

頂に至る道は平坦ではありませんでした。タン・ジエはしばしばワーカホリックと呼ばれました。夜中の二時に起き、一日の大半を研究室で過ごしました。研究により集中するため、会社の会長職と最高経営責任者の座は他の人物に任せ、自身は主任研究員としてZhipu研究院を率いる役割だけを引き受けました。彼は運動も好みました。マラソンとトライアスロンに出場して幾度も入賞し、水泳と

自転車、ランニングを続けるこの競技が体力と意志を試すという点で研究に似ていると語っていました。

タン・ジエはまるで仕事と研究しか知らない人物のように描かれがちですが、彼にもささやかな楽しみがありました。コーヒーがそれです。清華大学時代から毎日コーヒーを飲み、自ら中毒レベルだと語るほどでした。彼は、研究をコーヒーのように中毒性のあるものにできるなら、うまくいかない理由がないという言葉を好んで口にしていました。

2006年に王選と出会った青年博士から、2019年にZhipu AIを設立した教授を経て、悟道とGLMを率いた研究責任者に至るまで、タン・ジエの歩みは学者と起業家の間を行き来する旅ではなく、両方の立場を共に守る方向へと固まっていきました。清華大学知識工学実験室から始まった彼の研究は、国家単位のプロジェクトを経てZhipu AIという会社へとつながり、この流れは次節で見ていく会社と研究組織の物語へとそのままつながっていきます。

会社と研究組織を作り上げた過程

2006年、清華大学コンピュータ学科の博士課程を修了したタン・ジエは、留学と学内残留という二つの分かれ道で学校に残る道を選

びました。彼が思い描いていた手本は、北京大学の方正グループを設立した王選でした。王選は北京大学教授という立場で研究室を築き、漢字レーザー植字システムを開発しました。その成果によって中国の印刷術に二度目の革命をもたらした学者でした。

同じ年の11月、王選が世を去りました。タン・ジエはその一か月後の12月、清華大学知識工学実験室KEGを率いて、科学技術情報を大規模に分析するシステムAMinerを世に送り出しました。ある学者が世を去った場所で別の学者が新たな組織を興したかのような形でした。このAMinerが、後にZhipu（智譜）創業の種となりました。

AMinerは最初から純粋な学術プロジェクトにとどまっていませんでした。共にプロジェクトに参加していた張鵬は、後のインタビューでAMinerが採算を度外視した研究用ツールではなかったと明かしています。国内外の企業から情報サービス費用を受け取り、早い時期から収益を上げる仕組みを備えていたという説明でした。

システムが積み上げたデータは年を追うごとに厚みを増していきました。国家863計画と973計画、国家自然科学基金の支援を受け、ファーウェイ、ソウグウ、テンセント、アリババといった企業と

共同プロジェクトを進めました。2010年頃には検索可能な研究者プロフィールが44万8千件余りに達し、科学技術部と中国工程院を含め20あまりの機関・企業がこのシステムを利用していました。

学術的にも成果が伴いました。2008年にデータマイニング学会KDDで発表した論文は、同学会が過去10年間に掲載した約1,600本の論文の中で上位5位以内に入りました。この結果は、AMinerを作り上げた研究チームに、自分たちの方向性が間違っていなかったという確信を与えました。

この組織を支えたもう一人の柱が張鵬だった。1998年に湖南省で大学入試を受けて清華大学コンピュータ学科に入った彼は、2002年に唐傑が同学科の博士課程に入学したことで、師弟であり同門でもある関係で結ばれることになった。二人がのちに揃って中国の大規模言語モデル産業を率いる人物になったことを思えば、この出会いは並のものではなかった。

張鵬は学部と修士を清華大学で終える間に、世界最高峰の学会に10本余りの論文を発表した。中国語と英語を併せて扱う知識グラフシステムを設計した。2006年前後に修士を終えた彼は、すぐにKEG研究室に入りAMinerプロジェクトの一員となった。以後、唐傑と

ともに研究室の実務を担っていった。

AMinerを会社として立ち上げようという議論は2013年から出ていたが、実際に動き出したのはそれよりずっとあとのことだった。2018年、政府部門が大学や国立研究所の技術資産を商業実体へ転換するよう促す指針を相次いで打ち出し、清華大学の研究陣に新たな道が開かれた。同じ年にOpenAIが第一世代のGPTを発表すると業界全体の研究意欲が大きく揺さぶられ、AMinerを会社として分社化する意志もそれだけ固まっていった。

2019年6月。清華大学教授の李涓子と唐傑をはじめとする学者たちが、AMinerを大学から切り離し北京智譜華章科技有限公司、すなわち智譜AIを設立した。清華大学コンピュータ学科教授で中国科学院院士の張鈔が首席顧問として名を連ね、中国科学院計算技術研究所出身で清華大学データサイエンス研究院副主任を務めた劉徳兵が理事長に就いた。

経営陣の構成で目を引くのは唐傑自身の立ち位置だった。彼は最高経営責任者や理事長といった座には就かず、最高科学者と智譜研究院院長という肩書だけを持った。王選が残した言葉、すなわち科学者がメディアに頻繁に顔を出すようになれば、その学術的な命は

ほぼ終わったも同然だという教えを、唐傑もそのまま守った形だった。

こうした役割分担のおかげで、智譜は学者出身の創業チームがしばしば陥る経営権争いを免れ、技術研究に集中することができた。初代CTOだった張鵬は会社が大きくなるにつれ次第に前面に出てCEOの座に就き、唐傑は裏でGLMアーキテクチャの研究開発を率い続けた。設立当初30人だった組織は最盛期に700人を超えた。

2020年は二つの出来事が重なった年だった。OpenAIがGPT-3を発表した日が奇しくも智譜の創立1周年記念日と重なり、その場に招待講演者として来ていた張鈇院士と張鵬が、発表されたばかりのGPT-3について深く語り合った。1750億個のパラメータを持つこのモデルに刺激を受けた智譜は、自前の大規模モデル開発に力を注ぎ始めた。

智譜は市場に出ていたBERT、GPT、T5といったフレームワークを検討したのち、それらをそのまま使うのではなく独自のアーキテクチャGLMを設計する道を選んだ。それぞれのモデルの長所を集めて多様な能力を兼ね備えたモデルを作りたい、という考えがその出発点だったと張鵬は説明した。

2021年、智譜の研究チームは1300億個のパラメータを持つGLM-130Bの学習に着手した。当時、中国国内でこの規模のモデルに挑む例は少なく、参考にできる海外事例も乏しかったため、張鵬でさえ成功を確信できずにいた。だが数々の検証でGLMが自然言語理解の分野でBERTやT5を上回る性能を示すと、チームは方向性への確信を強めていった。

モデル開発の問題が解けると、すぐに資金の問題が付いて回った。学習に必要な計算資源と人材、データを揃えるには莫大な費用がかかり、2020年から2022年にかけて資本市場は大規模言語モデルにさほど関心を示さなかった。のちに名を上げた他のモデル系スタートアップも、この時期は大きな投資を受けられないまま持ちこたえていた。

智譜の初期投資家である中科創星の内部でも、智譜をめぐる論争は激しかった。智譜が求めた初期投資額4千万元は、この投資会社が通常エンジェルラウンドに充てる額を大きく上回り、知識グラフ事業の収益経路もはっきりしないという指摘が出た。それでも投資会社の上層部は、唐傑と張鵬のチームがKEGとAMinerを10年以上率いて築いた信頼を見込んで投資を決めた。

しかし、この関係は長く続かなかった。2021年、中科創星はファンドの満期を理由に智譜の持ち分25%を回収して投資から手を引いた。より深い理由は人工知能市場全体に対する悲観的な見通しであり、実際にこの投資会社は同年から自前のAI投資チームを解散させた。智譜はこの一件で再び資金繰りに苦しむことになった。

突破口は思いがけない場所から現れた。大規模な計算資源を確保できず苦勞していた智譜は、ある国内クラウド企業が2020年に購入したまま使われずに倉庫に眠らせていたGPUチップセットを見つけ出した。相場よりはるかに安い値段でこの機材を借りられるようになり、学習を続ける道が開けた。

智譜はここにモデル圧縮と低精度学習、重み量子化といった技術を組み合わせ、計算コストを9割近く引き下げた。こうして完成したGLM-130Bは2022年11月、スタンフォード大学の大規模モデルセンターが実施した世界の主要モデル30件の評価で、アジア企業のモデルとして唯一名を連ねた。

この評価結果と、その年の年末に登場したChatGPTが重なったことで、智譜に対する投資心理も大きく変わった。2023年だけで25億元を超える投資を受け、2024年までの累計投資額は50億元を超

えて企業価値は200億元を上回った。2023年から2026年の上場前までの累計調達額は83億元を超えた。

唐傑は会社の外でも別の役職を持っていた。彼は北京智源人工知能研究院(BAAI)の学術副院長を2019年から2024年初めまで務め、大規模モデル「悟道」の研究を率いた。この時期に共に研究していた学者たちの中から、のちに自らの会社を立ち上げた者が何人も出ており、月之暗面を作った楊植麟や面壁智能を作った劉知遠がそうした例だった。

KEGと智譜研究院が生み出した成果は、言語モデル一つにとどまらなかった。ChatGLM-6Bは62億個のパラメータを備えながら低コストの量子化に対応し、オープンソースとして公開された。GitHubのスター5万個、Hugging Faceのダウンロード数1千万回を超え、これを土台にしたアプリケーションが600を超えて作られた。2023年5月28日、科学技術部が中関村フォーラムで発表した報告書はChatGLM-6Bをオープンソース影響力1位に挙げ、GLM-130BやCodeGeeX、CogVideoなど智譜系のモデル5つが上位10位以内に入った。

智譜研究院はCogViewやCogVLMといった画像・映像関連のモデルも併せて発表し、言語モデルにとどまらない産業化を進めていっ

た。CogViewは同時期にOpenAIが発表したDALL-Eと比べて複数の指標で上回るという評価を受け、CogVLMは視覚理解を優先する方式で設計され、初期のGPT-4Vに匹敵する水準だという評価を受けた。

2024年10月、智譜は対話型サービス中心からエージェント中心へと研究の重心を移し、AutoGLMを発表した。2025年8月には独自の記憶機能を備えたAutoGLM 2.0を披露した。会社の研究組織はこの時期からコーディングとエージェント分野へ重心を移し、組織を再編した。

智譜は2026年初め、香港証券取引所に上場し、大規模言語モデル企業として世界で初めて株式市場に名を連ねた。上場当日に鐘を鳴らしたのは理事長の劉徳兵とCEOの張鵬であり、唐傑はほかの役員たちの間に静かに立っていた。上場後株価が上がったことで、彼の持ち分の価値も大きく跳ね上がったと伝えられる。

唐傑は上場の頃に公開書簡を出し、会社が2026年から商業化より基礎モデル研究へと重心を戻すと明らかにした。彼が新たに率いることになった組織X-Labでは次世代モデルGLM-5を準備した。このモデルはしばらく正体を隠したまま評価プラットフォームのOpenR

outerで優れた性能を示し、そのあとになって初めて智譜のものと明らかになった。

中核モデルと開発の過程

中国の人工知能界で独自の基盤モデルについて語るとき、必ず名前が挙がるのが唐傑(唐杰、Jie Tang)教授である。清華大学コンピュータ工学科の知識工学研究室(KEG)を率い、人工知能基盤モデル研究センターにも籍を置く彼は、GoogleやOpenAIが作った枠組みをそのまま使うのではなく、GLMという名の事前学習方式を自らの手で設計した人物である。彼が打ち立てたこの方式は、その後数年をかけて百億、千億、兆の単位まで規模を膨らませてきており、その成長の流れをたどれば、中国の大規模言語モデル産業が歩んできた道がそのまま見えてくる。

話は言語モデルを作る二つの異なる方式から始まる。Googleが開発したBERTは文を両方向から同時に読み取るため理解力には優れていたが、文章を新たに生み出すことは苦手であり、OpenAIのGPTは前から順に文章をつないでいくことには強かったものの、文脈を幅広く見渡す力は相対的に弱かった。唐傑教授の研究陣はこの両者の長所を一つにまとめてみることにした。文中の特定の区間を空欄にしておき、その空欄を順番に埋めていくように学習させる方式

だった。

この発想は2021年、自然言語処理分野で権威ある学会ACLに発表された論文として世に出た。空欄埋めを自己回帰方式で学習させるというこの原理は、理解と生成という異なる目標を一つのモデルの中で同時に達成しようとする試みだった。この基本構造の上で最初に生まれた成果が、百億個のパラメータを持つGLM-10Bだった。これはのちに続くモデル系譜の最初の一步となった。

2022年8月には、これよりもさらに大きな規模へと飛躍します。1300億個のパラメータを持つ中国語・英語のバイリンガルモデルGLM-130Bが公開されたのです。このモデルは同年11月、スタンフォード大学の大規模モデルセンターが世界の主要モデル30個を横並びで評価した際、アジア発のモデルとして唯一名を連ねました。正確性と有害性の指標はGPT-3と肩を並べました。頑健性と校正誤差では評価対象全体の中で最も良い結果を出しました。

GLM-130Bの内部構造を見ると、いくつかの工学的な選択が目を引きまます。基本構造はトランスフォーマーのデコーダーに従いつつ、単語の位置情報を表す方式としては回転位置エンコーディング、すなわちRoPEを採用しました。層を1000を超えて深く積み重ねる

と学習が途中で崩壊する問題がよく起きるため、これを抑えるために層正規化の配置を見直したDeepNet方式と、RMSNormという正規化手法を併せて組み込みました。規模を拡大することと学習を安定させることを同時に解決する必要があったようです。

2023年初めには、この巨大なモデルをはるかに軽い体に移し替えたChatGLM-6Bがオープンソースとして公開されます。62億個のパラメータ規模にINT4量子化まで適用し、一般的なグラフィックカード1枚でも対話できるようにした成果物でした。1兆4000億字に及ぶ学習データを消化し、P-Tuning v2という効率的な微調整方式に対応していました。GitHubのスター数は5万個を超え、Hugging Faceのダウンロード数は1000万件を突破して、両プラットフォームともに人気ランキング1位に上がりました。

2023年5月28日、中国科学技術部が中関村フォーラムで発表した人工知能大規模モデル地図研究報告書は、ChatGLM-6Bをオープンソース影響力部門の1位に挙げました。同じ報告書では、GLM-130Bとコード生成モデルCodeGeeX、テキストを映像に変換するCogVideo、そして基礎アーキテクチャモデルGLMまで、5つのモデルが揃って影響力上位10位以内に入りました。清華大学の一つの研究室

から出たモデルがこれほど大量に上位を占めたことは、それ自体が異例の結果でした。

この頃からオープンソースのエコシステムは手がつけられないほど拡大します。ChatGLMを基本構造として作られたサードパーティ製アプリケーションだけで600個を超え、この流れは後に智譜(Zhipu AI)という会社の商業的な基盤となります。唐傑教授はこの会社の共同創業者であり、学術顧問として参加しました。会社が設立された後も、GLMの系譜を継ぐ新しいモデルを次々と世に送り出し続けました。

2024年に登場したGLM-4シリーズは、これまでのモデルが受けたフィードバックを集めて練り上げた成果物です。GLM-4、軽量版のGLM-4-Air、そしてGLM-4-9Bといった派生モデルがこのシリーズに含まれます。この世代でとりわけ力を入れたのは、モデル自らがコード実行機やウェブ検索、画像生成ツールを呼び出して使うGLM-4 All Tools機能でした。モデルが答えだけを出す代わりに、必要なツールを自ら選んで使う方向へ一歩進んだのです。

この世代を支えたデータ準備の過程も注目に値します。ウェブ文書や書籍、論文、コードリポジトリを集めて10兆個のトークンを確

保しましたが、ただ集めるだけにとどまらず、ウィキペディアや学術書籍のような信頼度の高い資料の比重を意図的に引き上げる再調整作業を経ました。データを整える過程は大きく三段階に分かれます。文章や文書がそのまま重複する場合を取り除く精密重複除去と、意味が似ている文書を見つけてふるいにかける曖昧重複除去を併せて行い、その後、有害コンテンツやスパム、壊れたエンコーディングを取り除くフィルタリングを経て、最後にコンピューターが読み取れる数字列に変換するトークン化を実施します。

トークン化の過程でも独自路線を選びました。他の中国企業が既存のオープンソーストークナイザーをそのまま使う一方で、THUD Mの研究陣はバイト単位のBPEアルゴリズムを独自に設計しました。中国語専用の語彙と多言語語彙を別々に学習させた後、OpenAIが使うcl100k_base語彙集と統合し、15万個の単語を収めた統合語彙集を完成させました。この語彙集のおかげで、トークン変換にかかるコストと処理速度を大きく削減できました。

2025年末に登場したGLM-4.5は、構造の面でこれまでのモデルとは趣が異なります。全体のパラメータは3550億個に及びますが、一つの単語を処理する際に実際に動くパラメータは320億個だけで

す。必要な部分だけを選んで使う混合エキスパート、すなわちMoE構造のおかげです。軽量の展開を狙ったGLM-4.5-Airは、全体パラメータ1060億個に活性化パラメータ120億個で、さらに体を絞ったバージョンです。

競合モデルと比べると、設計思想の違いがはっきりと表れます。DeepSeekのDeepSeek-V3やMoonshot AIのKimi K2は、隠れ層の横幅を7168まで広げ、層数は相対的に低く抑えました。智譜の研究陣はこれとは正反対の道を選び、横幅を5120に狭める代わりに層数を89個まで積み上げました。GLM-4.5-Airも隠れ層次元4096に45層を置く、薄く長い構造を採りました。

こうした選択には理由がありました。研究陣が何度も比較実験を重ねた結果、モデルを横に広げるよりも縦に深く積み重ねたときに、数学の問題のように複数の段階を踏む必要がある推論や長いコードを扱う作業で性能がより大きく伸びました。幅よりも深さのほう思考力に大きな影響を与えるという判断が、この設計の背景にあったようです。

デコーディング速度を高める仕組みも新たに組み込まれました。通常のモデルは次に来る単語を一つだけ予測しますが、GLM-4.5は

多重トークン予測、すなわちMTP層をもう一つ設けて、この先に来る複数の単語をまとめて先取りして推測します。こうすることで文章を生成する速度が速くなりながらも品質は保たれます。専門家同士に作業を均等に振り分けるルーティングの過程でも、補助損失なしで均衡を保つ方式とシグモイドゲートを併用し、トークンが一方に偏ったり遅延が生じたりする問題を減らしました。

GLM-4.5を学習させたデータは23兆個のトークンに及びます。NVIDIAが発表したNemotron-CCのデータ精製方式に着想を得て、クロールリングで集めた膨大なウェブ文書を品質スコアに応じて複数の等級に分け、スコアの高い文書ほど学習により頻繁に登場するようサンプルを増やしました。データの量を増やすことよりも、良いデータをより頻繁に見せる方式を選んだのです。

GLM-4.5が示すもう一つの特徴は、一つのモデルが二つの思考速度を行き来する点です。数学オリンピックの問題や複雑なシステム設計、コードのデバッグのように手間のかかる質問を受けると、モデルは最終的な答えを出す前に内部で計算時間を延ばしながら中間の論理を何度も点検します。これを思考モードと呼びます。反対に、日常的な質問や短い翻訳、要約のような軽い作業には、この点検

過程を飛ばしてすぐに答えを出す直接応答モードに切り替わります。

この思考モードを有効にした状態で、GLM-4.5はアメリカの数学オリンピックの過去問を集めたAIME 24評価で91.0パーセントの正答率を記録しました。OpenAIのo1-previewモデルを上回る結果でした。答えを急いで出す代わりに、必要なだけ考える時間を確保する構造が実際のスコアにつながった事例だと言えます。

同じ時期に登場したGLM-4.5VとGLM-4.1V-Thinkingは、この思考方式を視覚資料に広げた成果物です。図表や幾何図形、設計図のような画像を読み取る過程で起きやすい錯視や誤読を減らすため、画像の境界線を細かく分割して分析するエンコーダーを備えました。視覚情報と内部の思考過程を互いに一致させる段階では、人間のフィードバックを反映した強化学習とルールベースの報酬を併せて用い、画像理解能力と推論能力をともに引き上げました。

智譜がこれまで歩んできた道の中で目を引くもう一つの原則は、モデルを作るたびにこれを公開してきたという点です。ChatGLM-6Bの第1世代から第3世代、128Kと100万トークンまで長い文脈を扱うGLM-4-9B、視覚言語モデルGLM-4V-9B、コーディング専門モデ

ルCodeGeeX、ウェブ検索を組み合わせたWebGLMに至るまで、主要モデルのほとんどをHugging FaceとGitHubに無料で公開しました。その結果、2023年の1年間だけでこれらのモデルの累計ダウンロード数が1000万件を超えました。

こうして積み重ねられた技術は、今ではサービスとしても束ねられています。智譜はモデル層、ツール層、アプリケーション層の三層からなる構造を備えています。モデル層にはGLM-4やGLM-4.5のような基盤モデルが置かれます。ツール層には企業がAPI呼び出しと微調整を扱うコンソールがあり、アプリケーション層には一般利用者向けの対話サービスZ.aiとBigModel.cn、そしてモバイルサービスの智譜清言(ジープーチンイエン)が置かれています。研究室から生まれた一つのアーキテクチャが論文とオープンソースモデルを経て、日々使われるサービスへとつながっていく流れこそが、唐傑教授が描いたGLMの系譜の実際の姿だと言えます。

現在の立場と中国AIにおける意味

唐傑は清華大学コンピュータサイエンス学科の教授を務めながら、自らが長年率いてきた知識工学研究室KEGを通じて、ナレッジグラフと大規模言語モデルをつなぐ研究を続けている。大学の中では後進を育てる教授です。大学の外では智譜AI(Zhipu AI)の共同創業

者であり最高科学責任者として、会社が作るモデルの方向性を定める人物である。この二つの立場を行き来する姿は、彼が歩んできた道をそのまま映し出している。論文で積み上げた知識を実験室の外の製品へと移す仕事を、彼は長い時間をかけて行ってきた。

AMinerで培われたナレッジグラフ技術は、智譜AIへと移りながら大規模言語モデル研究の土台となった。智譜AIは2021年のGLM-130Bを皮切りに、独自開発したGLMアーキテクチャを育て続け、その流れは2024年のGLM-4へとつながった。GLM-4は前の世代よりも推論能力と言語アライメント性能を引き上げたモデルとして、技術報告書を通じて公開されました。智譜AIが運営するモデルサービス体系にそのまま組み込まれた。

GLM-4の技術報告書は、GLM-130BからGLM-4に至る系譜を一編の文書の中にまとめている。パラメータ規模をむやみに拡大する代わりに、学習方式とデータ構成を練り上げて性能を引き上げる方向に重きを置いた跡が、この報告書の随所に表れている。唐傑がKEGで長く扱ってきたデータ精製と知識表現の方法論が、この過程で実際に用いられたものと推測できる。

同じ年の8月にはCogVideoXが公開された。文章として書いたテキストを入力すると物理的にもっともらしい映像を作り出すこのモデルは、トランスフォーマー構造を映像生成向けに改良したエキスパート・トランスフォーマー方式を採用した。言語モデル中心だった智譜AIの研究範囲が映像生成という別の領域へと広がった事例です。言語と画像、映像を一つの会社の中で共に扱う体制を築こうとする意図がうかがえる。

2025年には視覚情報を扱うマルチモーダルモデルの分野で成果が続いた。GLM-4.5VとGLM-4.1V-Thinkingは、画像や図表を単に読み取るだけにとどまらなかった。その中に含まれる因果関係をモデル自らが検討するようにすることに重点を置いた。図表に記された数字を見て結論を出す水準を超え、なぜそのような結論になるのかをモデルが段階的にたどっていくよう学習させたのである。強化学習を組み合わせ、推論の過程そのものを学習対象とした方式は、テキスト専用モデルから始まったGLMシリーズが視覚推論にまで及ぶ方向へ広がったことを示している。

GLM-4.1V-Thinkingが図表を読み取り自ら検討するようにする方式は、正解一つだけを当てるよう学習させる従来の方法とは異なる

道である。モデルが中間の推論過程を文章として展開するよう誘導し、その過程が正しかったかどうかを改めて点数化して学習に反映する手順を踏む。こうした学習方式は計算コストが大きいですが、結果だけを暗記するモデルではなく理由をたどっていくモデルを作ろうとする智譜AIの方向性を表している。

同じ年、GLM-4.5の技術報告書が公開された。このモデルは総パラメータ数が3550億に達するが、実際に一つの質問に答える際に働くパラメータは320億にとどまる。混合専門家構造、いわゆるMoE方式を採用し、計算量を抑えながら知識の幅を広げる設計である。報告書によれば、このモデルは89のMoE層と1つのマルチトークン予測(MTP)層で構成されている。複数のベンチマークにおいて、オープンAIのo1-previewモデルを数学推論の指標で上回ったと述べている。

混合専門家構造を採用した理由は計算量の調整にある。すべてのパラメータを毎回使えば、それだけ計算コストが大きくなるからだ。質問の性質に応じて必要な専門家ネットワークだけを選んで使わせれば、知識全体の幅を広く保ちながらも、一回の応答にかかるコストを抑えることができる。GLM-4.5が355Bという大きな総量と3

2Bという小さな活性量を併せて掲げているのは、こうした計算が背景にあるからである。資源が潤沢でない企業が大型モデル競争を戦い抜くための方法でもある。

智譜AI(Zhipu AI)はこうして作ったモデルを研究室の外に出し、実際に使われるようにする経路を複数用意した。一般ユーザー向けの智譜清言(ジーチンイエ)とZ.aiは、人々が直接対話しながら使うサービスである。企業顧客向けのBigModel.cnは、開発者がAPIでモデルを呼び出して使う窓口だ。モデルを作る層と、それを道具として磨き上げる層、そして実際の応用へとつながる層を一つの会社の中に備えており、この構造は研究室から生まれた技術が売上へとつながる経路の役割を果たしている。

智譜清言、Z.ai、BigModel.cnという三つの窓口は、それぞれ狙うユーザーが異なる。智譜清言は中国国内で一般の人々が日常的に対話を交わす相手である。Z.aiは海外のユーザーと開発者に向けられており、BigModel.cnは他社が自社サービスの中にGLM系列のモデルを組み込んで使えるように開く扉である。この三つの窓口が並び立って存在するという事実は、智譜AIが研究成果を一つ的方式だけで送り出したのではなく、ユーザー層に合わせてサービスの形を分

けて設計したことを示している。

こうした垂直統合の構造は、唐傑が学者として過ごしてきた時間と無関係ではない。AMinerを作った頃から、彼は研究成果を論文としてだけ残すのではなく、実際に動くシステムとして完成させることに力を注いできた。智譜AIがモデル開発からサービス運営までを一つの会社の中で処理している今の姿は、その習慣が会社という単位に広がった結果と見ることができる。

智譜AIは会社の外にある開発者コミュニティにも門戸を開いている。オープンソース大規模言語モデル基金は、コミュニティの開発を支援するために演算用グラフィックカード1000枚を提供し、現金1000万元をプロジェクト支援金として拠出したほか、優秀な開発者にはAPIトークン1000億個を無償で提供する計画を盛り込んでいる。世界中の起業家を対象としたZプランは、エコシステムのパートナーとともに10億元規模の起業支援基金を作り、基盤技術の開発を後押しする。自社の技術だけを育てるのではなく、その技術を使う人々の裾野を共に広げようとする取り組みとして読み取れる。

唐傑の名前が論文の著者欄に載り続けていることも注目に値する。GLM-4技術報告書、GLM-4.5VとGLM-4.1V-Thinkingを扱ったマル

チモーダル報告書、GLM-4.5技術報告書はいずれも智譜AIと清華大学の研究陣が共同で作成した文書である。会社の最高科学者という肩書きが実際の研究活動から切り離された名誉職ではなく、論文作成とモデル設計に絶えず関わり続ける役職であることを、これらの報告書が示している。

GLM系列モデルの特徴の一つは、自社開発の事前学習構造にこだわり続けている点である。BERT、GPT、T5のように広く使われるフレームワークはいずれも海外企業が作ったものであり、智譜AIはこれらの長所を参考にしつつも、そのまま取り入れるのではなくGLMという独自の構造を一から設計した。この決定は2021年のGLM-130B開発当時から続いてきたもので、以後GLM-4、GLM-4.5へとつながるモデル系譜の根となった。

マルチモーダル領域へと広げていった流れも、同じ文脈で見ることができる。CogVideoXが映像生成を、GLM-4.5VとGLM-4.1V-Thinkingが画像推論を担うことで、智譜AIのモデル群はテキストだけにとどまらず、複数の形式のデータを併せて扱う方向へと領域を広げた。言語、画像、映像をそれぞれ別のチームが個別に作るのではなく、一つの技術系譜の中でつなげていくやり方は、KEG時代から知

識を複数の形式で結びつけて扱っていたアプローチと似ている。

CogVideoXが採用したエキスパート・トランスフォーマー構造も注目に値する。テキストと映像は性質がまったく異なるデータである。この二つを一つのトランスフォーマーの中で併せて扱うには、それぞれのデータの特性に合わせた別々の処理経路が必要になる。CogVideoXはテキストを扱う部分と映像を扱う部分を異なる専門家モジュールに分け、この二つが注意機構を通じて情報をやり取りするように設計した。言語モデルを作っていたチームが映像生成へと領域を広げながらも、トランスフォーマーという共通の構造原理はそのまま守り続けた。

GLM-4.5が掲げる性能比較は、数学推論の領域で際立っている。技術報告書は、特定のベンチマークでこのモデルがオープンAIのo1-previewを上回る結果を出したと述べている。ただし、こうした比較はどのベンチマークを、どのような条件で実施したかによって結果が変わりうるため、報告書に記された数値をそのままあらゆる状況に当てはまる絶対的な優位として受け止めるより、智譜AIが自社モデルの強みとして掲げる指標として理解するほうが正確である。

混合専門家構造を選んだ判断には、コストの問題も絡んでいる。大手インターネット企業と比べると、智譜AIが使える計算資源は潤沢ではない。すべてのパラメータを毎回起動させる方式の代わりに、必要な部分だけを選んで使う構造を選んだのは、限られた資源の中で大型モデルの性能を最大限引き出そうとした実用的な選択だったと見ることができる。

Z.aiという名前で整理された対外サービスの窓口は、智譜AIが自社モデルを海外のユーザーにも開いていることを示している。中国国内で使われる智譜清言とは別に国際的な接点を用意しているのは、会社が中国市場だけでなく海外の開発者やユーザーまで視野に入れていることを意味すると読み取れる。ブランドを二つに分けて運営するやり方は、言語と規制環境が異なる二つの市場を一つのサービスに無理に束ねるより、それぞれの市場に合った窓口を別々に設けるほうがよいという判断から出てきたものと見られる。

GLM-4.5のMoE構造において、89の層が専門家ネットワークを分担し、1つのMTP層が複数のトークンを一度に予測する役割を担っているという説明も心に留めておく価値がある。一度に一語ずつ予測する代わりに、複数の語をあらかじめ見通して予測させれば、同

じ学習量でも文章をより正確に仕上げることができる。パラメータの総量を増やす方法と学習手順を磨き上げる方法を併せて用いた結果がGLM-4.5の性能に表れており、これは智譜AIが大型モデルを作るいくつかの方法のうち、どれか一つにこだわっているわけではないことも意味している。

唐傑が歩んできた学者としての道と、智譜AIが歩んできた企業としての道は、こうして互いに重なり合っている。KEGで鍛え、AMinerで試した知識処理の方法がGLM系列モデルの土台となった。そのモデルは再び智譜清言、Z.ai、BigModel.cnというサービスに分かれ、人々の手に渡っていく。論文から始まった技術が製品になり、製品が再び論文としてまとめられて学界に戻っていく循環が、彼のキャリアの中で繰り返されてきた。

智譜AIがGLM系列を通じて成し遂げようとする目標は、大型モデルを作ることにとどまらない。テキスト、画像、映像を包括する複数のモデルを一つの技術系譜のもとに置き、それを個人ユーザー向けサービス、企業向けAPI、オープンソースコミュニティ支援という複数の経路へと送り出す構造を組み立てることに力を注いでいる。この構造を設計する場に唐傑が最高科学者として名を連ね続けて

いるという事実が、彼が今、中国の人工知能の中で占めている位置をはっきりと物語っている。

清華大学教授と智譜AI最高科学者という二つの立場を兼ねるやり方は、よくある組み合わせではない。ほとんどの研究者は学界に残るか、会社へ完全に移るかのどちらかを選ぶが、唐傑は二つの立場を同時に保ちながら、大学から生まれたアイデアが会社の製品へとつながる橋渡し役を自ら担ってきた。GLM-4からGLM-4.5に至る技術報告書に彼の名前が絶えず登場することも、この橋が今なお実際に機能している証拠である。

技術報告書のたびに繰り返し表れる姿勢も目を引く。パラメータ数をむやみに増やすのではなく、学習方式と構造を磨いて効率を高めようとする姿勢。テキストだけにとどまらず画像や映像までを包括しようとする拡張、作り上げたモデルを再び複数の形のサービスへと解き放とうとする試みが、GLM-4、CogVideoX、GLM-4.5V、GLM-4.1V-Thinking、GLM-4.5へと続く文書に共通して現れている。この一貫した流れが智譜AIという会社の技術の方向性を示すと同時に、その方向性を設計する人物として唐傑が立つ位置を教えてくれる。

唐傑は清華大学教授として次世代の研究者を育てる仕事と、智譜AI最高科学者としてGLM系列モデルの技術の方向を定める仕事を、共に続けてきた。二つの役割のどちらも手放さず並行して続けてきたこと自体が、彼がこれまで歩んできたやり方と変わらない。彼が身を置く研究室から生まれた知識グラフとデータ処理の方法論は、大規模言語モデルの土台となった。そのモデルは言語を超えて画像や映像を扱う方向へ、そして個人向けサービス、企業向けAPI、オープンソース支援にまで広がり続けている。GLM-4からGLM-4.5へと続く文書は、この拡張が一度の発表で終わったことではなく、年が変わるたびに少しずつ異なる方向で繰り返されてきた過程であったことを示している。学界と産業界を行き来しながら積み重ねてきたこの軌跡が、今、中国の人工知能エコシステムの中で彼が占める位置を説明してくれる。論文と製品、大学と会社を行き来する歩みが今後も続くかどうかは、次に出る技術報告書が教えてくれるだろう。

第4章 梁文鋒(リャン・ウェンフォン、Liang Wenfeng)、ヘッジファンドの計算力を研究所に移したDeepSeekの創業者

成長期と学業

梁文鋒は1980年代、広東省の無名の小さな町で育った。北京や上海、深圳といった大都市からは遠く離れた土地で、コンピューター教育や最新の技術書に触れることが難しい環境だった。しかし彼は、この片田舎を欠乏とは捉えなかった。むしろ大都市式の競争や型にはまった詰め込み教育から一步離れていたからこそ、自らの好奇心をじっくりと磨いていく余裕を得られたと、彼は後に振り返っている。他人が敷いたレールをそのままたどるのではなく、未知の領域に飛び込んでいく粘り強さを、彼はこの小さな町での日々から得たと語った。

彼の父親は小学校の教師だった。裕福ではなかったが学問を重んじる家庭で、梁文鋒は幼い頃から読書の習慣と学びを大切にする姿勢を身につけた。父親は目先の小さな利益を追う術ではなく、物事の本質と善悪を自ら見極める姿勢を息子に授けた。

彼が少年期を過ごした1990年代、広東省全域では鄧小平の南巡講話以降、資本主義的な商業化の熱風が激しく吹き荒れていた。金を

稼ぐ機会はいたるところに溢れ、彼の家を訪ねてくる近所の大人たちも例外ではなかった。父親に相談を持ちかけに来た隣人たちは口をそろえて、子供に本を読ませ学校に通わせるのは時間の無駄であり、いっそ早くから商売を教え込んだほうがよいと言った。「勉強など役に立たない」という考えが、この小さな町の大人たちの間に広がっていたのである。

幼い梁文鋒は、こうした会話をそばで聞きながら育った。そして、その俗物的な損得勘定に流されるところか、すぐには金にならないように見える数学や学問に、むしろより深く没頭する頑固さを培っていった。この時期の経験は、後に彼がたびたび語ることになるエピソードとなった。彼は、たやすく金を稼げた1990年代の豊かさは、実のところ一時の運にすぎなかったと語っている。その運が尽きた後には、本物の技術を手にできなかった人間や国家が厳しい代償を払うことになると指摘した。

この幼少期の気づきは、後にDeepSeekを立ち上げる時まで受け継がれた。彼は、米国の研究所が作ったLlamaのようなモデルをそのまま借りてきて、月20ドルの購読サービスで手軽に金を稼ぐ道を、ただ乗りする者になる道だと見なして軽蔑した。彼が常に強調し

ていたのは、ただ乗りする者ではなく、貢献する者になることだった。

一貫して中国本土で教育を受けた「国内派」という経歴は、彼を語る際にしばしば言及される点である。

海外の名門大学の学位やシリコンバレーのビッグテック経歴を前面に出す他の創業者たちと比べると、梁文鋒の経歴は素朴に見える。しかし、この素朴さこそが、むしろ彼のアイデンティティを形づくっている。彼は制度化された大学の研究室で指導教官の下につき、研究テーマを割り当てられるような典型的な学者ではなかった。商業の現場とアルゴリズムのコーディングに自ら体当たりしながら身につけた、独学者に近い存在だった。

こうした独学者的な姿勢は、彼がDeepSeekで人材を採用する際にもそのまま表れた。彼はアイビーリーグの肩書きを持つ者や、米国のビッグテックで働いて帰国したいいわゆる「海外帰り組」を特別扱いしなかった。代わりに、人工知能への純粋な情熱と粘り強い好奇心、そして自ら学ぶ能力を何より重視した。大学を出たばかりか、実務経験が1、2年しかない若い中国人エンジニアたちを抜擢し、自分自身に似たチームをゼロから作り上げていった。

梁文鋒が本格的に名を知られるようになったのは、人工知能企業ではなく金融業界だった。彼はDeepSeekを設立する前、中国屈指の規模を誇るクオンツ・ヘッジファンド、幻方量化(High-Flyer)を興した人物として知られていた。幻方量化は、数学モデルとアルゴリズムのコードだけで莫大な資産をリアルタイムに運用する、中国トップクラスの定量投資会社だった。

16年以上に及んだ幻方量化での日々の中で、梁文鋒は膨大な金融データを扱いながら自らの技術的な感覚を磨いていった。完璧に整理された数学的データ環境の中でトレーディングを設計する仕事こそが、コンピューターと知能を扱う最高の訓練の場だったと彼は振り返っている。金融データにとどまらず、人類が生み出した膨大な非構造化データ全体を扱い、機械に自ら思考させること、すなわち汎用人工知能の研究こそが自分の到達すべき地点であると、この時期に自覚するようになった。

2023年、彼は幻方量化で蓄えた資本と、1万台を超えるエヌビディア製GPUを含むハードウェア資産を新しい研究所へ移した。そうして立ち上げた会社がDeepSeek(深度求索)である。社名に込められた志向は明確だった。情報の壁と知能のコストを、その根本から

引き下げるということだった。

DeepSeekが初めて登場した時、中国の大規模モデル・スタートアップ7社の中で最も目立たない会社と評価されていた。1年前まで、こうした存在感の薄さはごく自然なことだった。その背後にクオンツ・ヘッジファンドの大手、幻方量化がいるという事実だけで説明がついていた時代だったからである。

しかし梁文鋒は、初めから別の前提に立っていた。他の中国企業が、数か月後には誰かが新しいモデルを出すのだから、ただそれに追随しながらアプリケーションサービスをうまく作ればよいと考えていた時。彼は自分たちの目的地を汎用人工知能(AGI)と定めた。それは、限られた資源の中でより強力なモデル性能を実現するために、新しいモデル構造そのものを研究しなければならないという意味だった。

彼は、本当の格差とは1年や2年の時間差ではなく、独自技術を生み出す能力とそれを模倣する能力との間の差だと見ていた。オープンAIがいくらソースコードを公開しなくても、その閉鎖性そのものが後発者の追い抜きを防いではくれないと判断していた。

この判断の根底には、中国国内の人材に対する彼なりの確信があった。彼は、最高水準の人材が中国では正当に評価されていないと考えていた。欧米ビッグテックの華やかな肩書きがなくても、純粹な好奇心と実力だけで世界水準の研究を成し遂げられる人々が、すでに中国国内にいと見えていたのである。

DeepSeekが初期に賭けた方向性も、こうした考えの延長線上にあった。梁文鋒は、数学とコード、マルチモーダル、そして自然言語そのものという三つの方向に会社の力を集中させた。この三つの方向は、その後DeepSeekが発表するモデル群の基本構造となった。

DeepSeekの初期の研究成果は、この方向設定が無駄ではなかったことを示した。推論コストを100万トークンあたり1元の水準まで引き下げたことは、梁文鋒が少年時代から抱いていた問題意識、すなわち技術の敷居を下げるという課題と、まさに重なり合うものだった。

DeepSeekが公開した技術報告書の中で、梁文鋒の名前は華やかな修飾語なしに登場する。DeepSeek-R1の論文もDeepSeek-V3の論文も、彼を研究とエンジニアリングに関わった共著者のリストに

、ただ名前だけで載せている。創業者であり最高経営責任者でありながら、自分を特別に前面に押し出さないこうした姿勢は、彼が少年時代に父親から受け継いだ態度と重なっている。

梁文鋒をめぐる話の多くが短く断片的である理由は、彼がメディアの前に姿を現すことが極めて稀だったからである。彼が対外的なメディアと直接向き合い、踏み込んだ話をした場合は、片手で数えられるほどしかない。

梁文鋒という人物の輪郭は、比較的是っきりと描き出せる。大都市ではなく、片田舎の小さな町で。商業化の熱風の真ただ中にあっても学問の価値を手放さなかった少年が、後にクオンツ金融というまったく異なる領域でデータを扱う感覚を培い、その感覚を今度は人工知能研究へと移していった、という流れである。

この流れの中で注目すべきなのは、彼がどの段階でも主流のやり方をそのままなぞらなかった点である。少年時代には、金稼ぎこそ最善だという時代の空気に逆らい、創業後には他人のモデルを真似て素早くサービスを売る手軽な道の代わりに、独自技術を研究する困難な道を選んだ。この選択の一貫性こそ、梁文鋒という人物を理解する手がかりだと言えるだろう。

DeepSeekがその後発表したV2、V3、R1といったモデル群と、オープンソース公開戦略。そして市場を揺るがした価格政策は、この成長期の問題意識が具体的な形で実現した結果である。

研究者あるいは創業者として定まっていく時期

梁文鋒は1980年代、広東省の小さな町に生まれた。大都市の熾烈な受験競争からは少し離れた場所で、彼はそこで、他人に追随するよりも自ら観察し、考え抜く習慣を身につけた。父親は小学校の教師であり、裕福ではない暮らしの中でも、利益よりも善悪を見極める姿勢を息子に受け継がせた。

彼が育った1990年代の広東省は、金儲けの算段で沸き立つ時代だった。家に出入りする近所の大人たちは、父親の前で子供を学校に通わせるのは時間の無駄だ、いっそ商売を教えたほうが良いと口をそろえた。少年梁文鋒は、そうした会話をそばで聞きながら育った。しかし彼は金の力に押し流されることなく、数学と知的探求の方へとより深く目を向けていった。

数十年後、彼はその頃を振り返り、当時人々が容易に金を稼げたのは、単に時代が与えた運にすぎなかったと語っている。その運が尽き、独創的な技術革新なしには生き残れない時代が来れば、タク

シー運転手すら居場所を失うだろうというのが彼の考えだった。この幼少期の気づきは、のちに彼が他人の作った技術を模倣して手軽に金を稼ぐ道を極度に嫌うようになる土台となった。

高校を終えた梁文鋒は浙江大学電子工学科に進学し、人工知能を専攻に選んだ。2000年代半ばは人工知能が今のように脚光を浴びていた時代ではなかったが、彼はこの分野に飛び込み、数学とアルゴリズムの基礎を固めた。

彼は2010年に情報工学の修士号を取得した。梁文鋒が大学院を終え社会に出ようとしていた2008年9月、米国のサブプライムローン問題に端を発した金融危機が世界を襲った。表面的には無秩序に見えるこの巨大な流れの底に隠れた因果関係を、機械学習で読み解けるのではないかという問いが彼を捉えて離さなかった。

この時期、彼は四川省成都の小さなアパートにこもり、世間と少し距離を置きながら機械学習アルゴリズムを独学した。彼のコーディング能力とデータを扱う感覚に目を留めた人物がいた。のちに世界的なドローン企業へと成長するDJIの創業者である。DJIの創業者は彼を直接引き抜こうとしたが、梁文鋒はこの誘いを断った。彼が望んでいたのは、空を飛ぶ機械を操るハードウェアではなく、デー

々に隠れたパターンを自ら見つけ出す智能そのものを、ゼロから作り上げることだった。

2013年から2016年にかけて、梁文鋒はいくつものデータ分析組織を相次いで立ち上げ、実戦で自らの実力を試した。ただし、これらすべての試みが向かう方向は一つだった。人間の目では追えないごく短い瞬間の価格変動パターンを、機械学習で捉えることである。

この経験は2016年2月、幻方量化(High-Flyer)の設立へとつながった。幻方は数式と機械学習アルゴリズムのみで資産を自動運用するクオンツ・ヘッジファンドだった。梁文鋒の指揮のもと幻方は急速に規模を拡大し、2021年末には運用資産が94億ドルに達した。市場平均指数を20から50ポイント上回る利回りを安定して出し続け、梁文鋒は30代半ばにして自力で財を成した億万長者となった。

梁文鋒が幻方を通じて築いたのは資本だけではなかった。彼は、モデルとデータだけでは人工知能の性能がある閾値を越えられないこと、そして計算資源も併せて積み上げてこそ知能が突如として飛躍することを、身をもって体得した。だからこそ彼は早くからグラフィックカードを買い集め始めた。2012年にはわずか1枚のGPUで実験を始め、2015年には100枚規模のクラスターを整えた。

2019年には1000枚の高性能GPUを確保し、2021年にはエヌビディアのA100とH100グラフィックカード1万枚以上を予約購入し、「萤火二号」と名付けたクラスターを構築した。これはアジア地域の企業の中でも屈指の規模だった。資金を運用するヘッジファンドの計算力が、いつしか人工知能研究を支えるインフラへと姿を変えていたわけである。

80億ドル規模のクオンツファンドを成功裏に運用し、十分な資本を手にした梁文鋒は、今や金融データを超えて、より広い知能の問題へと目を向けた。2023年4月、彼は幻方が積み上げてきた1万枚を超える計算資産と、相当額の個人資本を元手にDeepSeek(深度求索)を設立した。名前の通り深く探求するという意味を込めた会社であり、目標はただ一つ、汎用人工知能の事前学習技術を自らの力で切り開くことだった。

創業から半年後の2023年11月、DeepSeekは最初のモデルを世に送り出した。その後の展開は、中国の他の大型モデル・スタートアップ7社と比べても際立って静かだった。他社が華やかな発表とマーケティングに熱を上げる間、DeepSeekは特に目立った動きを見せず、ひたすら研究に没頭した。その沈黙を破ったのは、結局のと

ころ性能と価格だった。

2024年のあるインタビューで梁文鋒は、DeepSeekが目指す先は汎用人工知能であり、そのためには限られた資源の中でより強力なモデル能力を引き出す新しい構造を研究しなければならないと語った。彼は会社が確実に賭けた三つの方向性を明らかにした。第一に数学とコード、第二にマルチモーダル、第三に自然言語そのものである。この三つの方向はいずれも、人間の言語と論理、推論能力を機械へと移そうとする試みという点で、一つにつながっていた。

彼が繰り返し強調したのは、基盤技術と模倣の違いだった。真の格差は1年や2年という時間の問題ではなく、独創的な革新と模倣との間にある根本的な違いだというのが彼の考えだった。新しいモデルなど数か月経てば誰かが作り出すのだから、中国企業はただ後を追いかけてサービスさえうまく作ればいいという通念に、彼は同意しなかった。彼は、オープンAIが自らのソースコードを公開しないとしても、他者に追い越されるのを防ぐことはできないと指摘した。

こうした考えの根底には、中国国内の人材に対する彼自身の評価があった。彼は、最高水準の人材が中国では長らく過小評価されて

きたと見ていた。DeepSeekが世界各地から人材をかき集める代わりに自国内の研究者に機会を与えた背景には、こうした判断があった。梁文鋒自身も浙江大学で人工知能を学んだ国内派であり、その経歴が会社の採用方針にもそのまま反映された。

DeepSeekが世に広く知られるきっかけとなったのは価格だった。推論コストを100万トークン当たりわずか1元水準まで引き下げ、業界に価格競争の号砲を鳴らしたのである。この号砲がどう始まったのかという問いへの答えは、結局のところ梁文鋒が長年積み上げてきた計算インフラと、効率的なモデル構造の研究とが噛み合った結果だという点にあった。

梁文鋒はDeepSeekをただ乗りする存在ではなく、貢献者にしたいと考えていた。中国の他の大型モデル企業が外国で作られた技術を追いかけ、応用するだけにとどまるという印象を与えた一方で、彼は基盤技術そのものを世界に差し出す道を選んだ。この志向は、その後DeepSeekが自社モデルの重みと研究手法をオープンソースとして公開する形で具体化された。

DeepSeek公式ウェブサイトとHugging Faceのリポジトリには、こうしたモデル公開の成果がまとめられている。

広東の小さな町から始まり、浙江大学と四川省成都を経て、幻方のクオンツ・トレーディングルームに至る道は、振り返れば一つの問いにつながっていた。無秩序に見える世界の中で、機械に自らパターンを見つけさせることができるかという問いである。その答えを探し求める過程で、彼は研究者であると同時に創業者となり、その二つの役割はDeepSeekの中で互いに分かれることなく、今日まで続いている。

会社と研究組織を作り上げた過程

DeepSeek以前の梁文鋒の時間は、ヘッジファンド、ハイフライヤー(High-Flyer、幻方量化)のサーバールームで流れていった。浙江大学電子工学部で人工知能を専攻した彼は、資本市場の価格変動を機械学習で読み解くことに没頭し、その成果が2016年2月に正式に発足した定量投資ヘッジファンドだった。この会社がのちにDeepSeekを生み出す資金とハードウェアの源泉となったという点で、梁文鋒の経歴は金融工学から汎用人工知能研究へと渡る一本の橋を架けた形になっている。

その橋は、GPUカード1枚から始まった。ハイフライヤーの研究チームは2012年、わずか1枚のグラフィックカードで実験を回すところから始め、2015年には100枚規模のクラスターを整え、2019年

には1000枚にまで増やした。そして2021年、アジア太平洋地域でエヌビディアA100カードを1万枚以上確保した企業の一つとなった。大企業ではない組織の中で、これほどの物量を蓄えたところは他になかった。

このハードウェアの蓄積について、梁文鋒はのちのインタビューで淡々と説明している。自分たちがチャットGPT創業ブームを予見して先回りしたのではなく、モデルとデータと計算資源が一体となって結びついてこそ知能が飛躍するという信念のもと、早くから少しずつ積み上げてきた結果だというのである。こうして確保した計算資源は、のちに「萤火二号」という名の大規模コンピューティング拠点として整理され、梁文鋒はこの頃、自力で財を築いた富豪の仲間入りを果たした。

資金と設備を手にしたあと彼が選んだ道は、その資金をさらに定量投資に注ぎ込むことではなかった。2023年4月、梁文鋒はハイフライヤーとは別に、独立法人DeepSeek(深度求索)を設立した。社名の漢字の意味そのままに、深く探求する姿勢を掲げた組織であり、設立の趣旨は当初から明確だった。特定業種向けの応用製品を素早く売って稼ぐ道よりも、汎用人工知能という目的地に直接近づく

基礎研究を優先するというものだった。

こうした姿勢は、当時の中国人工知能業界の流れとは一線を画すものだった。新しいモデルなど数か月経てば誰かが作り出すのだから、中国企業はただ後を追いつながら応用サービスさえうまく作ればいいという考えが広く浸透していたが、梁文鋒はこうしたただ乗り路線を拒み、貢献者になる道を選んだ。創業初期、ベンチャー投資家たちが早期回収を求めてサブスクリプション型チャットボットサービスから出すよう圧力をかけたときも、彼は基盤研究を優先すると言ってこうした提案を退けた。

DeepSeekが最初に世に送り出した成果は、2023年11月に公開した大規模言語モデルだった。金融人工知能を扱う中で、機械の判断がいくらでも間違いうるという事実を身をもって経験していた梁文鋒は、この最初のモデルの限界を隠さず率直に明かす姿勢を組織文化とした。華やかな成果発表よりも謙虚な自己診断を前面に押し出したわけであり、これはのちにDeepSeekが築いていく研究中心の文化の最初の場面だった。

DeepSeekという名前が世に広く知られるきっかけとなったのは、2024年5月に発表したV2モデルでした。このモデルは100万トー

クンあたりの処理コストをわずか1元程度まで引き下げて登場しました。この価格は瞬く間に中国の大規模モデル業界全体を揺るがす価格競争の引き金となりました。DeepSeekの創業者にインタビューしたメディアは、この状況について「価格競争の最初の一発はどのように放たれたのか」と尋ねたところ、梁文峰は意図的に業界の秩序を乱す「なまず」役を演じようとしたわけではなく、成り行きでそうなってしまったと答えました。

彼の説明によれば、価格設定の過程はいたって基本的なものでした。原価を計算した上で、損をしない範囲でありながら暴利をむさぼらない程度の利益を上乗せしただけだということです。市場がこの価格にこれほど敏感に反応するとは予想していなかったとのことでした。V2発表の5日後には智譜(Zhipu)AIが値下げに踏み切り、続いてバイトダンス、アリババ、百度、テンセントといった大手企業が次々と追随しました。中でもバイトダンスは、自社のフラッグシップモデルの価格をDeepSeekと同水準まで引き下げた最初の手企業でした。

梁文峰は、この値下げがユーザーを奪うためのマーケティング施策ではなかったときっぱりと言い切りました。次世代モデルの構造

を研究する過程で原価そのものが下がり、人工知能という道具は誰もが安価に使えるべきだという考えがその根底にあったという説明でした。実際、DeepSeekはそれまで中国企業が慣例のように踏襲していたLlama系の構造をそのまま模倣する代わりに、最初から独自のモデル構造を打ち立てる道を選びました。

この選択の理由を梁文峰は、目指す先の違いとして説明しました。応用製品を素早く世に出すことが目標であればLlama構造をそのまま踏襲する方が合理的ですが、DeepSeekが目指す先は汎用人工知能(AGI)であり、これは限られた資源でもより強力なモデル性能を引き出す新たな構造の研究を必要とするということです。彼はLlama系構造が学習効率と推論コストの面で海外最高水準からおおよそ二世代ほど遅れていると診断しました。学習効率とデータ効率それぞれで生じる差を合わせると、同じ性能を出すために四倍もの演算能力を余計に使わなければならないと指摘しました。

こうした構造研究は2024年12月に公開されたV3モデルへとつながりました。マルチヘッド潜在アテンション(MLA)とDeepSeek独自のスパースな専門家混合構造を組み合わせたこのモデルは、Llama 3.1 405Bより学習コストを11分の1の水準まで抑えながらも、複数

の評価でオープンソース最高性能に到達しました。GPT-4oやClaude 3.5 Sonnetといった最上位のクローズド型モデルとも真正面から競える水準に達しました。

翌2025年1月には、強化学習で推論能力を高めたR1モデルが後に続きました。グループ相對方策最適化という強化学習手法により、人が一つひとつ正解を教えなくても自ら推論過程を磨いていくよう訓練されたこのモデルは、性能指標においてOpenAIの推論モデルと肩を並べる結果を出し、Hugging Faceで公開されて誰でもダウンロードして使えるようになりました。

DeepSeekがなぜあえてソースと論文を公開する道を選んだのかについて、梁文峰は明確な論理を持っていました。破壊的な技術の前でソースを閉じることで築いた堀は長続きせず、OpenAIがどれほどコードを非公開にしているとしても、他の者に追い抜かれる流れそのものを止めることはできないということです。だからこそ彼は堀をコードではなくチームに築きました。同僚たちが研究の過程で成長しながら積み重ねていくノウハウと、それを可能にする組織文化そのものこそが本当の競争力だと見たのです。

オープンソースを商業的な損益ではなく文化的な行為と見る視点も、同じ脈絡から生まれました。ソースを公開し論文を出したところで会社が実際に失うものはさほど多くなく、むしろ他の人々が自分たちの研究を追いかけてくることに研究陣が感じる達成感と、与える行為そのものが会社に積み重なっていく文化的な名声があると梁文峰は語りました。彼はこうした姿勢について、短期間で金を稼ぐやり方とは異なる方法で自分たちの地位を築いていっているのだと説明しました。

研究組織を満たす人材を選ぶ基準においても、DeepSeekは他の中国企業とは一線を画しました。DeepSeek V2を作ったチームには海外から戻ってきた人材が一人もおらず、全員が中国本土で教育を受けた人材で構成されていました。梁文峰は、世界最高水準の中核人材上位50人が今すぐ中国にいなくても、そうした人材は会社が自ら育て上げることができると思っていた。

こうした人材観は組織運営の方式にもそのまま浸透しました。DeepSeek内部には、あらかじめ定められた部署間の境界や上意下達の指示体系がはっきりと機能していませんでした。研究者それぞれが自ら興味を感じた問題を提案すると、そのテーマを中心に人が自

然に集まる形で協業が行われました。アイデアを検証するための演算資源の使用にも別途の決裁手続きや上限を設けず、研究者は学習クラスターのGPUを必要なだけすぐに引き出して使うことができました。

こうした水平的な構造が実際に成果を上げた事例としてよく挙げられるのが、V2とV3の中核技術であるマルチヘッド潜在アテンションの誕生過程です。アテンション構造の変遷を個人的に掘り下げていた若手研究者が代替案となりうる構造を提案すると、部署を問わず関心を持つ人々が集まってチームを組み、数か月でこのアイデアを実際に動くアルゴリズムへと完成させました。あらかじめ定められた研究ロードマップをたどった結果ではなく、組織が許容した自律性がそのまま技術へとつながった事例でした。

梁文峰はこうした試みについて、中国はこの数十年間、他国が成し遂げた技術革新を取り入れて応用し現金化することには慣れていたものの、肝心の革新そのものにはほとんど関与できていなかったと指摘しました。ムーアの法則がひとりで回るように十八か月ごとにより優れたハードウェアが登場する流れにただ乗るのに慣れていただけで、その流れを作った西洋の技術コミュニティが世代を継

いで重ねてきた創造の努力は目に入りにくかったというのです。彼は、中国経済の規模や大手企業の利益水準を見れば、革新に不足しているのは資本ではなく自信と、高密度に集まった人材を組織する方法だと診断しました。

彼は本当の格差は米中間の1年、2年という時差ではなく、独創性と模倣の違いにあると見ていました。この差が縮まらない限り中国は永遠に後を追う立場にとどまるほかなく、どのような形であれ自ら道を探る模索は避けられない課題だというのが彼の結論でした。DeepSeekが数学とコード、マルチモーダル、自然言語という三つの方向に並行して力を注いだのも、こうした判断から生まれた選択でした。

米国の先端半導体輸出規制が強化されるにつれ、DeepSeekの演算インフラ確保戦略も変化を余儀なくされました。大型化した次世代モデルの開発を準備する過程でNVIDIA系ハードウェアへのアクセスが次第に狭まると、DeepSeekはHuawei傘下のAscendプラットフォームへ中核演算インフラを移す作業に乗り出しました。異なるハードウェア構造に合わせて学習コードを組み直し最適化するこの作業には相当なコストがかかり、これは会社がそれまで守ってきた

資金自立路線に負担を加える要因となりました。

膨らむサーバー維持費と利用者急増に伴う運営コストを前に、梁文峰は創業以来初めて外部資本を受け入れる方向へと舵を切りました。テンセントと京東(JD.com)、電池メーカーCATL、そして国策色の強い人工知能産業基金が参加する大規模な資金調達がこの時期に行われました。ただし彼はこの資本に会社の意思決定権までは渡しませんでした。取締役会の議席を投資家に譲らず、研究方向に関する発言権も資本ではなく研究チームが握るという原則を明確にしました。金は受け取るがハンドルは渡さないというこの原則は、ヘッジファンドで築いた資本とGPUを自ら統制しながら研究を守ってきた梁文峰の長年の姿勢が、会社の統治構造へと移された結果だと言えるでしょう。

核心モデルと開発過程

梁文峰が描いていた目的地は、最初から応用サービスではなく汎用人工知能でした。彼は西側が作ったLlama構造をそのまま模倣してアプリケーションを素早く世に出す道を選びませんでした。代わりに、限られた演算資源の中でより強い知能を引き出すにはモデル構造そのものを新たに設計しなければならないと判断しました。その判断がDeepSeek(深度求索)の研究方向を最初から決定づけまし

た。

この判断を身をもって証明したモデルが、2024年12月に発表されたDeepSeek-V3です。トランスフォーマー層を61層積み重ね、隠れ次元を7168に設定した巨大な構造です。全体のパラメータは6710億個に達しながらも、実際にトークン一つを処理する際に動くパラメータはわずか370億個です。体は大きくても実際に力を使う筋肉は一部だけを選んで使うという意味であり、この設計が学習費用とサービス単価を同時に抑え込みました。

この圧縮の中核に、マルチヘッド潜在アテンション、略してMLAという仕組みがあります。通常のアテンション構造は、モデルが文章を読みながら蓄えるキーとバリューの情報をそのまま保存するため、モデルが大きくなるほどこの保存領域が雪だるま式に膨らみサービス原価を押し上げます。DeepSeekの研究陣はこのキーとバリューを512次元という狭い通路に圧縮して入れておき、必要なときだけ復元して使う方式を作り出し、クエリベクトルも1536次元に縮めて管理しています。

位置情報を運ぶ回転位置埋め込み、RoPEはこの圧縮と混ざると情報が潰れる恐れがあるため、64次元の別通路を専用に設け、位置情

報だけを独立して流し込むようにしました。こうして分けた結果、DeepSeekは推論サーバーで実際に使う記憶保存領域を従来方式の5パーセントから13パーセントの水準まで減らしながらも、モデルが元々持っていた判断力はそのまま維持しました。

フィードフォワード層でも同様の節約が起きています。DeepSeekは専門家一人ひとりを細かく分割したDeepSeekMoE構造を引き継ぎ、各層に共有専門家一つとルーティング専門家256個を配置しました。文章一つが入ってくると、その多くの専門家のうち8人だけが呼び出されて処理を行い、その8人が散らばる物理サーバーも4台を超えないよう通信経路まであらかじめ設計しておきました。

複数の専門家に仕事を均等に振り分けることは以前から悩みの種でした。人気のある専門家にばかり仕事が集まるのを防ぐため、従来のモデルは強制的にペナルティを科す補助損失関数を加えていました。このペナルティが智能そのものを損なう副作用を生んでいました。DeepSeekはこのペナルティを完全になくし、ルーティングの入口に設けたバイアス値を学習中にリアルタイムで少しずつ調整して仕事を均等に分配する方法を見つけ出しました。

次の単語一つだけを見据えて予測していく方式にも手を加えました。DeepSeekは本体の後ろに一層構造のマルチトークン予測層を付け加え、今の単語のすぐ次だけでなく、その先に来る単語まで併せて見据えるようにしました。こうすることで学習中の計算がより速く収束しました。実際のサービスで答えを出す際にも、一度に複数の単語をあらかじめ準備しておく方式で速度を引き上げることができました。

こうした構造を学習させるには材料となる文章も膨大でなければなりません。DeepSeekは世界中に散らばる文書や書籍、論文をかき集めて14兆8000億個のトークン規模の多言語コーパスを作り、その中で数学とプログラミング資料の比重を前世代よりはるかに高めました。Common Crawlという巨大なウェブ収集庫から良質なコードと数学の文章だけを選び出すため、Stack Overflowのようなプログラミングフォーラムとスタック・エクスチェンジの数学掲示板の投稿をシードとし、このシードで学習させた分類器を全ウェブページに適用して似た文章を探し出しました。この作業を三度繰り返して深く掘り下げた末に、GitHubのソースコード940億トークンと数学テキスト2210億トークンを含む1兆1700億トークン規模の専門資料を確保しました。

DeepSeek-V3全体を事前学習から文脈拡張、事後調整まで完走させるのにかった演算量は、H800グラフィックカード基準で278万8000時間でした。米国の大手テック企業が同程度の性能を出すために注ぎ込んでいた予算と比べると、桁が一つ丸ごと抜け落ちた規模です。この格差が2024年末以降DeepSeekを取り巻く騒動の根本的な背景となりました。

梁文鋒はこのV3の上に、もう一つの挑戦を重ねました。言語を扱う知能を超え、自ら思考を振り返り修正していく知能を作ろうとしたのです。その結実が、強化学習を極限まで押し進めたDeepSeek-R1シリーズです。

最初の一步は、人が事前に手を加えた学習データを一切排除する実験でした。V3のベースとなる重みの上にルールベースの報酬だけを載せて強化学習を回した成果がR1-Zeroです。ここでディープシークはGRPOというアルゴリズムを用いましたが、従来のPPO方式が方策モデルと同じくらい大きな価値モデルをもう一つ別に走らせる必要があり、メモリを大きく食っていたのに対し、GRPOは同じ質問に対して複数の答えを一度に生成し、その中で相対的に優れた答えに点数を付ける方式によって、価値モデルそのものをなくして

しました。

この訓練が積み重なる中で、驚くべきことが起きました。誰もコードとして組み込んではいませんでした。モデルが問題を解く途中で行き詰まると、自ら先の思考を振り返って再検討し、答えを出すまでにかかる思考の長さを自発的に伸ばし始めたのです。研究陣はこの瞬間を「アハ・モーメント」と呼びました。

ただし、R1-Zeroはそのまま世に出すには荒削りでした。思考を展開する途中で中国語と英語が入り混じり、段落すら分かれておらず、人が読むには不便でした。梁文鋒と研究陣はこの問題を直そうと、第二段階に進みました。読みやすい形に整えた少量の思考過程データを人の手で選び出し、モデルに教師あり学習で習得させたのです。

その次に、この中間モデルに比較的高い温度値を与えて多様な答えを生成させ、正解であり形式も整っている結果だけを選んで再びデータとしました。言語が入り混じった部分はV3本体に再び入力し、思考する言語を質問の言語に合わせるよう調整しました。こうして整えたデータの上に、役に立ち害にならない方向でもう一度強化学習を重ねた成果が、最終的なDeepSeek-R1です。

この訓練全体にかかった時間も注目に値します。R1-ZeroはH800グラフィックカード64枚を8基のクラスターに組んだ環境で198時間動かし、最終的なR1は同規模の設備で80時間、つまりわずか4日で訓練を終えました。これに加え、教師あり学習用データを作るための演算に、さらに5000GPU時間が投入されました。

R1が完成した後、ディープシークはこの知能を小型モデルにも分け与える作業に入りました。数学とコード。こうして作られた縮小版は15億から700億まで6種類が生まれ、元のモデルよりも数学とアルゴリズムの試験で目に見えて優れた点数を出しました。

こうした成果すべてを、ディープシークは商業利用や改変まで自由なMITライセンスでハギングフェイスにそのまま公開しました。V2とV3、R1と縮小版6種の重みはもちろん、推論とサービングに使うコードまでギットハブに合わせて上げておいたのです。これは梁文鋒が信念として語ってきた考えと、そのままつながっています。彼は、閉鎖的に築いた壁は長く持たず、オープンAIがソースを閉じても結局は誰かに追い越されるのを防げないと述べました。代わりに彼が本当の防壁だと考えたのは、研究陣がこの過程で積み上げたノウハウと組織文化そのものでした。

価格政策も同じ信念から生まれました。ディープシークは自社の最上位モデルのAPI利用料金を、100万トークン当たり入力1元。出力2元の水準に定めました。これは米国の閉鎖型モデルと比べると、短く見て10倍、長く見れば70倍も安い値です。梁文鋒は、この価格は暴利を取ることも損をすることもない、原価にわずかな利益だけを乗せた値だと説明しました。5日で智譜AIが追随し、バイトダンスが自社の旗艦モデルを同じ価格に下げたことで、中国の大型モデル業界全体に価格競争が広がりました。

梁文鋒はこの騒ぎを歓迎もせず、驚きもしませんでした。彼は、価格にこれほど人々が敏感に反応するとは思わなかったと淡々と語っただけで、利用者を奪うことが目的ではなかったと明かしました。彼が価格を下げた理由は二つありました。新しいモデル構造を探る過程で原価そのものが下がったこと、そしてAPIであれ知能であれ誰もが享受できるべきだという考えが、その根底にあったのです。

同じインタビューで彼は、ラマ構造をそのまま踏襲する道と自分たちが選んだ道との違いを、訓練効率とデータ効率という二つの方向から指摘しました。モデル構造と訓練方式で二倍。データを扱う

効率でさらに二倍の格差があるとすれば、同じ結果を出すのに海外最高水準より四倍多い演算を投じなければならないという意味です。彼がやろうとしたのは、まさにこの格差を一枚ずつ縮めていく作業でした。

梁文鋒はこの作業を、中国の人工知能産業全体の課題として広げて語りました。この三十年間、中国企業は他人が生み出した革新を持ってきて現金化することに慣れきっており、彼はこれが当たり前のことではないときっぱり言い切りました。本当の格差は一年や二年の時間差ではなく、そもそも作り出す力と後を追う力の違いです。この違いが縮まらない限り、中国は永遠に後を追う立場にとどまるほかないと彼は判断しました。

この判断を裏付ける根拠の一つが人材でした。彼はディープシークV2を作ったチームに海外から帰国した人材が一人もいなかったと明かし、世界最高水準の人材が今の中国に五十人いなくても、自前で育てることができるというのを語りました。最高水準の人材は中国において今なお過小評価されているというのが、彼の考えでした。

V3とR1という二つの柱を打ち立てた後、ディープシークが直面した次の課題は、この知能をどのように世に送り出すかということで

した。一般利用者はチャットアプリでR1のゆっくりとした思考を無料で体験し、企業はオープンAIの規格と噛み合うAPIを通じて同じ知能を利用できるようになりました。研究所から出発した計算力が消費者の指先に届くまで、梁文鋒が歩んできた道のりは、モデル構造を新たに組み立てる作業から始まり、価格表の一行で結ばれた形になっています。

現在の立ち位置と中国AIにおける意味

梁文鋒は現在、ディープシーク(深度求索)の創業者・最高経営責任者として、中国の人工知能業界全体を設計する立場に立っています。彼が以前に設立したクオンツ運用会社、幻方量化(ハイフラワー)は、数学モデルと機械学習で資産を運用していた会社であり、ディープシークはその中から分社化して独立した別法人です。梁文鋒自身も、大型モデルを作る仕事は定量投資や金融とは直接関係がないと述べたことがあります。この区分は、彼がディープシークを純粋な研究組織として位置づけてきた姿勢とつながっています。

ディープシークが世に知られたきっかけは、派手なマーケティングではなく価格表でした。2024年5月にディープシークV2が登場すると推論コストが100万トークン当たり1元の水準まで下がり、この数字一つが中国の大型モデル業界全体の価格競争に火をつけまし

た。梁文鋒は、この競争は自分たちが意図した結果ではないと語ります。原価を計算してわずかな利益を乗せただけなのに、皆がこれほど価格に敏感に反応するとは思わなかったというのが彼の説明です。

価格引き下げの余波は速いものでした。5日で智譜(ジープュー)AIが追随し、続いてバイトダンス、アリババ、百度(バイドゥ)、 Tencentといった大手企業が次々と価格を下げました。梁文鋒はこのうちバイトダンスが自社の旗艦モデルの価格をディープシークの水準に合わせた企業だと指摘し、大手企業のモデル原価が自分たちよりはるかに高いという点で、この流れが結局はインターネット時代の補助金競争と似た形に流れていったと評価します。

彼が価格を下げた理由として挙げたのは二つです。一つは、新しいモデル構造を研究する過程で原価そのものが下がったという事実であり、もう一つは、APIであれAIであれ誰もが使えるべきだという考えです。利用者を奪うことが目的ではなかったという彼の言葉は、この二つ目の理由につながっています。

ディープシークが他の中国企業と違っていた点は、モデル構造に対する姿勢にあります。当時、大多数の中国企業はラマ(Llama)構

造をそのまま取り入れて応用製品を素早く出す道を選びました。梁文鋒はこの道が悪くないと認めながらも、自分たちの目的地は汎用人工知能(AGI)であり、そのためには限られた資源の中でより強い性能を引き出す新しいモデル構造の研究が必要だと語ります。

彼は、ラマ構造は訓練効率と推論コストの面で海外最高水準よりおよそ二世代ほど遅れていると見ていました。訓練効率で二倍。データ効率でさらに二倍の格差が出れば、結局は同じ結果を出すために四倍多い演算力を使わなければならないという計算であり、ディープシークがすべきことはこの格差を一つずつ縮めていくことだとまとめます。

こうした説明の底には、より広い問題意識が横たわっています。梁文鋒は、この三十数年間、中国はムーアの法則が自然に回っていくかのように、何もしなくても十八か月ごとにより優れたハードウェアとソフトウェアが現れることに慣れきってきたと語ります。しかし、これは西側の技術コミュニティが世代を重ねて作り上げてきた成果であり、中国がその過程に参加してこなかったために、その存在自体に気づけなかっただけだということです。

そのため彼は、経済が成長した分だけ、中国もただ乗りする者ではなく貢献者にならなければならないという言葉を繰り返します。36クリプトン傘下のメディア「暗涌」との二度のインタビューで、彼はこの表現を核心的な一文として残しており、これはディープシークがオープンソース路線を守り続ける理由とも重なります。

梁文鋒が見る本当の格差は、時間ではなく性質の問題です。人々はよく、中国のAIは米国より一、二年遅れていると言いますが、彼はその格差がそもそも縮められる時間の問題ではなく、独創的な革新と模倣の間の違いだと指摘します。この違いが残っている限り、中国は永遠に後を追う立場にとどまるほかないというのが、彼の診断です。

革新のコストが低くないということも、彼は認めます。しかし今の中国の経済規模や、バイトダンス、テンセントのような大手企業の利益水準は、世界的に見て決して低い方ではありません。彼から見て不足しているのは資本ではなく自信、そして密度の高い人材を組織して実際に成果を出す方法です。

クローズド型モデル戦略についても、彼は明確な立場を示している。どれほどコードを閉ざしても他者の追随を止めることはできず

、だからこそDeepSeekは価値をチームの中に根付かせる道を選んだという。仲間たちが研究の過程で成長し、ノウハウを積み重ねながらイノベーションを続けられる組織と文化そのものが、自分たちの長く持ちこたえる力だという説明である。

オープンソースと論文発表で失うものは何もないという発言も、同じ文脈から出たものだ。技術陣の立場からすれば、他者が自分の後を追ってくるという事実自体から大きな達成感が得られ、オープンソースは商業的行為というよりも文化的行為に近く、与える行為そのものが付随的な名誉になるというのが彼の考えである。

投資家の朱嘯虎(ジュー・シャオフー)の短期収益論理については、その方式なりに一貫性はあるが、短期間で稼ぐ企業に向けたアプローチだと一線を画す。アメリカで最も稼いでいる企業を見れば、結局は長い年月をかけて内実を積み重ねた末に市場を主導する先端技術企業であるという点を根拠に挙げている。

DeepSeekが大規模モデル市場に参入した理由についても、彼は明確に一線を引く。アン・ロンの質問に対し彼は、この事業がクオンツ投資や金融と直接的な関係はなく、DeepSeekという独立法人を新たに設立したのもそのためだと答えた。目標は汎用人工知能で

す。大規模言語モデルはその道のりの途中ですでに初歩的な特徴を見せ始めているので、ここから出発し、映像資料の研究のような次の段階へ進むという構想である。

既存モデルを複製するだけの道と、本当の研究をする道との費用の違いについても彼は指摘する。公開された論文とオープンソースコードを基に数回の訓練や微調整を経るだけで済む複製は、費用が少なく済む。一方、あらゆる実験と対照を経なければならない研究は、計算資源と人的資源の要求がはるかに高く、費用が大きく膨らむ。

2021年に幻方(ハイフライヤー)がアジア太平洋地域で最も多くのA100グラフィックカードを1万枚単位で確保していたことも、こうした文脈で理解できる。ChatGPTブームが吹き荒れる前にすでになりの規模の計算資源を蓄えていた状況だが、梁文峰はこれを大げさな戦略というよりも好奇心から生まれた選択だと説明する。

V2モデルを作った人々についての質問にも、彼の答えは簡潔だ。元OpenAIポリシー責任者でAnthropic共同創業者のジャック・クラークが、DeepSeekについて正体不明の人材を雇っていると評したとき、梁文峰はV2開発チームに海外から帰国した人材は一人もいな

かったと明かした。世界トップ50に入る人材が中国にいないかもしれないが、それならば自分たちでそうした人材を育てればよいというのが彼の答えだった。

これらの発言を貫いているのは、梁文峰がDeepSeekを、中国AI産業の中でルールに従う会社ではなく、ルールそのものを書き換えようとする会社として描いているという点である。彼は創業初期のインタビューで、三つの方向に賭けたと明かしている。数学とコード、マルチモーダル、そして自然言語そのものがそれである。

技術的にDeepSeekが出した成果は、この方向性と結びついて現れた。2024年12月に公開されたV3は、複数の評価でオープンソース最高性能に達し、Llama 3.1 405Bを上回り、GPT-4oやClaude 3.5 Sonnetといった最上位のクローズド型モデルにも匹敵する水準に到達した。それでいて価格はClaude 3.5 Haikuよりも低く、Claude 3.5 Sonnetの9パーセント程度にとどまった。

訓練費用の面でも、V3はLlama 3.1 405Bの11分の1程度と推定されるという評価を受けた。

DeepSeekが大企業以外で唯一、A100チップを1万枚単位で備蓄していた会社だったという点も改めて見ておく価値がある。複数の事業を同時に手がける道を捨て、研究と技術にのみ集中しました。

こうした路線は、梁文峰が繰り返し語る文明史的観点とつながっている。彼は、アメリカで毎日起きている無数のイノベーションの中で、DeepSeekが特に話題になった理由は、この会社が中国企業でありながらイノベーションの貢献者という立場でそのゲームに参入したからだとしている。大半の中国企業がイノベーションよりも追従に慣れていたことと対照的な点である。

彼は過去三十年間、中国が金儲けばかりを重視しイノベーション自体をおろそかにしてきたと指摘しながら、イノベーションは商業的利害関係だけでは動かず、好奇心と作ってみたいという欲求が必要だと語る。彼が見るところ、今の惰性は永遠の条件ではなく、過ぎ去っていく一つの段階に過ぎない。

これらの発言がなされた時点では、DeepSeekはまだ世に広く知られていなかった。しかしV3、そして続いて公開された推論モデルが海外でも話題となったことで、梁文峰がインタビューで語っていた方向性と原則は、後になって裏付けを得た形になった。価格競争

の最初の号砲を鳴らした会社。モデル構造そのものを設計し直した会社、オープンソースを商業的損得ではなく文化的信念として扱った会社という評価がその後続いた。

梁文峰個人の立ち位置も、この流れの中で変わっていった。彼は華やかな脚光を求めるよりも研究とコーディングに集中する姿勢を保ち続け、メディア露出を極力控え、数少ない長文インタビューを通じてのみ自分の考えを明かしてきた。DeepSeekが引き起こした価格競争と技術論争が大きくなるほど、まさにその中心にいる創業者の言葉数はむしろ少なくなっていく点が目を引く。

第5章 楊植麟（杨植麟、Yang Zhilin）、長いコンテキストでAGIを推し進めたMoonshotの若き創業者

生い立ちと学業

楊植麟は1992年に生まれました。中国第二世代の生成AI大型モデル創業者の中では、年齢的に若い部類に入る人物です。彼が生まれた地域や幼少期を過ごした具体的な場所。彼の物語は、したがって工学的アイデンティティが実際に形作られ始めた大学入学の時点から始まります。

2011年、楊植麟は清華大学コンピュータサイエンス技術学科に入学しました。当時、清華大学のコンピュータ学科は中国のAI人材が集まる場所の一つでした。後に知譜AIの母体となる知識工学研究室も、まさにこの学科の中で少しずつ育ちつつありました。二人のその後の道を分けることになる学問的なルーツは、すでに同じ建物の中に置かれていたわけです。二十歳前後の楊植麟は、こうした環境の真ただ中で大学生活を始めました。

清華大学入学初期、彼がすぐに人工知能に打ち込んだわけではありません。コンピュータサイエンスはとても幅広い学問であって。新入生であれば誰でもソフトウェア工学やシステム、アルゴリズムといった様々な科目を一通り経験しながら自分の方向性を探ってい

くものです。楊植麟もまた、入学後しばらくはこの幅広いカリキュラムの中で基礎を固める時間を過ごしたものと見られます。

学部2年生だった2012年頃から、楊植麟は人工知能研究に没頭し始めました。最初から一つの分野だけを深く掘り下げたわけではありません。グラフ構造を扱うネットワーク演算。複数の形式のデータを同時に処理するマルチモーダルインターフェースのように、人工知能の中の様々な分野を幅広く見渡しながら自分に合った道を探していきました。二十代前半の学生にありがちなように、彼もまた広く見渡す時間を経てから一つの分野へと絞り込んでいきました。

学部卒業論文を準備する中で、楊植麟は清華大学コンピュータ学科の唐傑教授と出会います。唐傑教授は後に知譜AIを設立する人物でもあります。楊植麟は唐傑教授の指導のもと、自然言語処理とソーシャル知識グラフを結びつけるアルゴリズムを共に設計しました。この過程で、言語を扱う人工知能というテーマに本格的に足を踏み入れることになりました。彼は2015年、清華大学コンピュータサイエンス技術学科の学士号を取得しました。

一人の指導教授のもとで共に訓練を受けた二人が、数年後にそろって会社を設立し、それぞれ異なる道を歩むことになるという事実

は。当時としては誰も予想できなかったことでしょう。清華大学コンピュータ学科という一つのルーツから、知譜AIとMoonshot AIという二つの会社が共に育っていったのです。これは中国AI産業の人材地図を理解するうえでも参考になる点です。

唐傑教授が指導した学生の中で、後に独自に会社を設立した人物は楊植麟一人だけではなかったはずです。

清華大学コンピュータ学科在学中に彼が受けた訓練は、純粋な理論だけにとどまりませんでした。大手インターネット企業との産学協力。学内コンペティションや先輩研究者との勉強会といった経路を通じて、学生たちは論文を読む目とともに、実際にコードを書いて結果を確かめる習慣を身につけていきます。楊植麟もまた、こうした学内の雰囲気の中で研究者としての基礎体力を養ったものと見られます。

清華大学を卒業した後、楊植麟は米国カーネギーメロン大学コンピュータサイエンス学部の博士課程に進学しました。コンピュータサイエンス分野において世界屈指の大学院です。彼はここで二人の指導教授のもとで研究を続けましたが、一人は機械学習分野で名の知られたルスラン・サラフトディノフ、もう一人はグーグルAIでも

活動したウィリアム・W・コーエンでした。

二人の指導教授の組み合わせは、学生だった楊植麟にバランスの取れた訓練をもたらしたと言えます。サラフトディノフが機械学習理論の基礎を固める役割を担ったとすれば、コーエンは-googleという産業現場の感覚を共に伝えられる人物でした。学校の中の理論と学校の外の実践が、一人の指導体系の中に共に組み込まれていたのです。

カーネギーメロン大学自体が、コンピュータサイエンスと人工知能研究で長い歴史を積み重ねてきた場所です。こうした環境で博士課程を歩んだ学生は、自然と複数の実験室を行き来しながら異なるアプローチに触れることになります。楊植麟が自然言語処理という一分野にとどまらず、グラフ。推論、強化学習などへと研究の幅を広げることができた背景にも、こうした学校の雰囲気があったものと推察されます。

博士課程の間、楊植麟は教室の中だけにとどまりませんでした。彼は-google・ブレインでインターンとして働き、ディープラーニングのスケーリング研究で知られるクォック・V・。学校の外、産業研究所で得たこの経験は、理論研究が実際の大規模コンピューテ

イング環境とどのように噛み合うのかを、彼に直接示してくれました。

グーグル・ブレインでインターンとして過ごしていた時期。楊植麟は、人工知能モデルをより賢くする方法について一つの確信を得たと語ります。アーキテクチャを細部にわたって手直しし磨き上げる作業よりも、モデルとデータの規模を大きくする方が、より大きな飛躍を生むという感覚でした。これはいわゆるスケーリング則と呼ばれる考え方と重なり合うものであり、後に彼が設立した会社の研究方向にもそのまま受け継がれる発想でした。

博士課程の間、楊植麟はチューリング賞を受賞したヤン・ルカン、ヨシュア・ベンジオといった学者たちとも学術的に交流しました。彼らと意見を交わし、一部の研究を共に進める中で、彼は人工知能学界の最前線で繰り広げられる議論を間近で見守ることができました。カーネギーメロン時代の彼を知る人々は、彼がピンク・フロイドのようなロックバンドを好んで聴いていたとも伝えています。

この時期に楊植麟が名を連ねた研究の中で、今日でもしばしば言及されるのがXLNetです。当時グーグルが発表したBERTは、文章を双方向に読んで文脈を把握する方式で自然言語処理分野を牽引し

ていました。楊植麟はここに自己回帰方式の利点を組み合わせたXLNetを設計しました。この論文は2019年のNeurIPS学会で口頭発表論文に選ばれました。

XLNetが解決しようとした問題はこうです。BERTのように文章を双方向に一度に読めば文脈を幅広く把握できますが、文章を順序どおりに一つずつ生成していく作業にはうまく適合しない面がありました。楊植麟は、双方向に読む利点と順序どおりに生成する方式の利点を一つのモデルの中に共に取り込もうとしました。その結果は、様々な言語処理課題で良い成績につながりました。

同じ時期に発表されたTransformer-XLも、彼の名前とともに残る研究です。既存のトランスフォーマー構造は、一度に処理できる文章の長さが決まっています。長い文章を扱う際に前半の文脈を失ってしまう限界がありました。楊植麟は、区間をまたぐ再帰方式と相対的位置エンコーディングという手法でこの限界を解決しました。

Transformer-XLで扱った長文脈処理の問題は、後に彼が設立した会社が送り出すモデルの中核的な特徴とつながる点です。短い文章しか読めなかったモデルを。前後の文脈を途切れさせずに長く読み続けられるモデルへと変えようとする試みが、博士課程時代からす

でに彼の研究の中に根を下ろしていたわけです。

2018年に発表されたHotpotQAも彼が参加した研究です。このデータセットは、人工知能が一つの文書だけを読んで答える代わりに。複数の文書を行き来しながら段階的に推論して答えを見つけ、その根拠まで示すよう設計されています。今日でも、人工知能モデルが複雑な質問にどれだけうまく答えられるかを測る基準の一つとして使われています。

同じく2018年、楊植麟はヤン・ルカンと共にGLoMoという研究を発表しました。正解が定められたデータがなくてもデータ同士の関係をグラフの形で抽出し、その関係を他の作業にもそのまま応用できるようにする方法を扱った研究でした。2022年には唐傑教授の研究室と共にGLM系列の研究にも名を連ねましたが、これは後に知譜AIのモデル系列を支える設計へとつながりました。

この時期に彼が書いた論文を並べて見ると、一つの方向性が浮かび上がります。XLNetとTransformer-XLは言語をより長く正確に読み取る問題を扱い、HotpotQAとGLoMoは複数の断片的な情報をつなぎ合わせて推論する問題を扱いました。一見すると互いに異なるテーマのようですが、結局はすべて、言語モデルが世界をより広く

理解できるようにする作業へと収斂していきます。

清華大学とカーネギーメロン大学を行き来しながら積み重ねたこれらの論文は、単に学位を得るための手続きではありませんでした。各論文が扱った問題。すなわち双方向文脈と逐次生成の結合、長文脈処理、多段階推論、関係抽出は、後に一つの会社が送り出すモデルの中でそれぞれ異なる形で蘇ります。学生時代の研究テーマが創業後の技術の方向性どこまで重なる事例は、そう多くありません。

こうした論文が積み重なるにつれ、楊植麟の研究は学界の中で幅広く引用されるようになりました。彼の論文が受けた引用回数を合計すると、22,000回を超えていると言われています。三十歳にも満たない研究者としては稀な成果でした。この数字は、彼が後に会社を設立する際、投資家や同僚研究者たちに自らの実力を説明する裏付けとなったものと推測されます。

博士課程5年という歳月は決して短くありません。その間、楊植麟は清華大学で学んだ自然言語処理の基礎を土台に。学校と産業現場を行き来するこのような循環は、後に彼が会社を経営する際にも、学界と産業界を共に見渡す感覚として残ったものと見られます。

楊植麟は2019年、カーネギーメロン大学で博士号を取得した。清華大学の学部時代に自然言語処理に足を踏み入れてから4年目のことであり、学部入学を基準にすれば8年目のことだった。指導教授であったルスラン・サラフトディノフは、この卒業をオープンソース人工知能コミュニティ全体にとっても意義深い瞬間だという趣旨で語ったと伝えられている。

博士課程を経る中で、楊植麟の考えは少しずつ一つの方向へ収束していった。2017年頃から彼は、数ある人工知能分野の中でも言語モデルこそが汎用人工知能への確かな道であるという考えを固めていったという。人間が使う言語の中には世界の情報と理屈が凝縮されて詰まっており、次に来る単語を当てる訓練を極限まで押し進めれば、世界を理解するモデルに近づけるという考えだった。

この考えは、グーグル・ブレインでのインターン時代に得た感覚と自然に噛み合った。細部の技法を細かく磨き上げるよりも、モデルとデータの規模を大きくする方がより大きな成果を出すという確信が、言語モデルに対する信念と重なったのである。この二つの考えが合わさり、彼は言語モデルの規模を拡大することに自らの研究方向を賭けることになった。

博士課程を歩む学生の大半は、指導教授が定めた狭いテーマ一つに何年も取り組むものである。楊植麟はこれとは異なり、自然言語処理という大きな枠組みの中でも、文脈処理、推論、グラフ構造まで幅広く手を伸ばした。そうして複数の方向で積み重ねた経験が、後に言語モデル一つへと自然に収束していった。

清華大学で唐傑教授のもとで自然言語処理を学び、カーネギーメロン大学でサラフトディノフとコーエンの指導を受けながら産業界の研究室での経験まで積んだ時間は、楊植麟にとって確かな土台となった。論文を書く研究者にとどまらなかった。その研究を実際に大きく育ててみる機会を自ら作り出そうという方向へ彼の考えが移っていったのも、この頃からだったと見られる。

ただし、彼が創業以前まで自然言語処理研究に携わった期間が10年に及ぶと自ら明かしている点を見れば、博士号取得後もこの分野の研究を手放さずに続けていたという事実だけは明らかである。

2023年4月、自然言語処理研究に携わって10年目を迎えた年に、楊植麟はムーンショットAIを設立する。清華大学とカーネギーメロン大学で培った言語モデルへの確信、そして規模を拡大すればより大きな知能が生まれるという信念を、そのまま会社の方向性へと移

し替えたのだった。創業以降の具体的な歩みと製品にまつわる話は、後で改めて取り上げることにする。

こうして見ると、楊植麟の学業期は一人の人間が一つの道を長く掘り続けた物語に近い。19歳で清華大学に入り、23歳で自然言語処理というテーマを定め、27歳で博士号を取得するまで、彼が扱った問題は常に言語、文脈、推論という数本の軸から外れることがなかった。創業後に会社が世に送り出した製品群がこの軸の上でそのまま育っていったのも、そのためである。

清華大学時代の唐傑教授との出会い。学部から大学院まで8年。大学院を卒業してから創業までさらに4年という歳月の間、彼は自然言語処理という一つの分野を絶えず掘り下げ続け、その結果が後に彼の会社が世に送り出すモデル群の土台となった。

研究者、あるいは創業者として定まっていく時期

中国の人工知能業界において楊植麟という名前は、論文引用回数2万2千回を超える研究者としての経歴と、ムーンショットAI(月之暗面)という会社を率いる経営者としての経歴を同時に指す。1992年生まれの彼は、実験室の中で積み上げた成果を実験室の外の製品へと移すことに躊躇がなかった。その転換の根は、大学時代からす

でに培われていた。

楊植麟は2011年、清華大学コンピュータサイエンス学科に入学した。入学当初から人工知能を目標に定めていたわけではなく、学部2年生の頃、2012年あたりからこの分野に心を注ぎ始めた。初めのうちはグラフ演算、マルチモーダルインターフェースといった異なる方向のテーマを幅広く手がけながら、自分に合った場所を探っていた。

その模索の時期に彼に目をかけ、導いてくれたのが清華大学コンピュータ学科の唐傑教授だった。唐傑は後にジュー・エイ(智譜)を共に立ち上げることになる人物で、当時はデータマイニングとソーシャルネットワーク分析を教えながら、学部生だった楊植麟の研究を直接指導した。学生と教授として始まったこの縁は、その後、中国のAI学界で二人がそれぞれ会社を興してから度も度々交わる関係へとつながっていく。

唐傑のもとで生まれた最初の成果は国際学会で相次いで認められた。ネットワークデータの能動学習を扱った論文が2014年のWSDMで採択率5パーセントの口頭発表として採択された。同年、CIKMにもストリーミングネットワークデータを扱った論文が掲載された

。2015年には異種ソーシャルネットワークをつなぐ研究がKDDで口頭発表として採択された。この学会の全体採択率が19パーセントにとどまっていたことを考えると、学部生の成果としては異例だった。

唐傑と共に設計したソーシャル知識グラフモデルは2016年、IJCAIの論文としてまとめられた。清華大学の学術情報分析プラットフォームAMinerに実際に組み込まれ、研究者の関心事を抽出するために使われた。論文で終わらず人々が日々使うシステムへとつながったという点で、楊植麟は学部時代から理論と実装を切り離して考えない習慣を身につけていたと見られる。彼はこれらの成果を土台に2015年、清華大学の学士号を取得した。

卒業後、楊植麟は米国カーネギーメロン大学コンピュータサイエンス学院の博士課程に進んだ。指導教授は二人だった。深層学習研究で名を知られ、後にアップルのAI研究部門を率いることになるルスラン・サラフトディノフと、グーグルAI主任科学者のウィリアム・W・コーエンである。異なる背景を持つ二人の指導教授のもとで、彼は理論と応用の両方の視点を同時に鍛えられた。

博士課程の間、彼は大学の中だけにとどまらず、シリコンバレーの研究所を行き来しながら大規模な計算環境を直接扱う経験を積んだ。グーグル・ブレインでは、オートMLの研究で名を知られたクック・V・レと協働した。大学の実験室での論文作業と企業研究所での実戦感覚が、この時期に重なり始める。

この頃、彼が筆頭著者として参加した論文は今なお自然言語処理分野で頻繁に引用される部類に入る。2018年にEMNLPで発表したホットポットQAは、複数の文書を行き来しながら答えを見つける多段階推論を扱うためのベンチマークデータセットで、ヨシュア・ベンジオやクリストファー・マニングといった研究者たちと共に作った。機械が一つの文書だけを読んで答える方式を超えさせようという問題意識が、この作業の出発点だった。

2019年にACLで発表したトランスフォーマーXLは、その名の通り既存のトランスフォーマーの限界を扱った論文だった。当時のトランスフォーマーは一度に処理できる文章の長さが固定されており、それより長い文章を読ませるには切り分けて入力するしかなかった。トランスフォーマーXLは区間をまたぐ再帰構造と新しい位置エンコーディング方式を加え、この限界を解いた論文だった。長い文章

を切らずに読むモデルを作れる可能性を示した。

同年、NeurIPSで発表したXLNetも話題を集めた。グーグルが発表したBERTが双方向文脈予測で自然言語処理を牽引していた時期に、XLNetは自己回帰方式の利点と双方向文脈予測の利点を両方取り入れる構造を提案した。この論文はNeurIPS全体の投稿作のうち0.5パーセントにしか与えられない口頭発表の機会を得た。

こうした成果を土台に、楊植麟は2019年、CMUでコンピュータサイエンスの博士号を取得した。指導教授だったサラフトディノフは、彼の卒業に際してこの成果をオープンソース人工知能コミュニティ全体の勝利だと評価した。弟子一人の達成がコミュニティ全体の功績として語られるほど、トランスフォーマーXLとXLNetはその後、多くの研究者が踏み台とする土台となった。

博士課程とシリコンバレーを行き来する間、楊植麟の中で少しずつ固まっていたのは特定の論文一つではなく、研究に向き合う姿勢だった。彼は2017年頃から、言語モデルこそが人工知能全体が向かう通路であるという方向へ考えを絞り込んでいった。人間が使う言語の中には世界についての情報と論理がすでに溶け込んでおり、次に来る単語を正確に予測する能力を極限まで押し進めれば、世

界を理解するモデルに近づけるという考えだった。

グーグル・ブレインで過ごした時間は、この考えをより確かなものにした。決まった大きさのデータセットの上で細部の変数を少しずつ調整して小数点単位のスコアを上げる作業よりも、計算とデータの規模を拡大する方が性能をはるかに大きく押し上げるという事実を、彼は身をもって経験した。以後よく言われるスケーリング則への信念は、この時期の経験に由来すると見られる。

博士号を取得した後も、楊植麟は学界と完全に縁を切ったわけではなかった。2019年から2022年の間、清華大学KEG研究室との協力を続け、穴埋め方式を自己回帰学習に組み合わせたGLM構造を扱った論文を2022年ACLに発表し、複数言語に対応するコード生成モデルCodeGeeXの作業にも名を連ねた。会社を興す前の最後の数年間も、彼は依然として論文を書く研究者だった。

流れが変わった契機は、2022年11月30日にオープンAIがChatGPTを公開した出来事だった。言語モデルが実験室の外的一般ユーザーに直接届く製品になり得るという事実が目の前で証明されると、楊植麟はこの流れが人工知能分野の資本と人材が動く方向そのものを変えるだろうと判断した。彼はすぐに共に歩む人々を集め始めた。

2022年末から2023年初めにかけて、彼はアメリカに滞在しながら、事業計画を具体的な数字に落とし込む作業を進めていた。モデルの訓練と運用にかかる演算量、ユーザーが増えたときに必要となる推論コストといったものを、一つひとつ計算していった。論文を書いていた頃と同じように、勘ではなく数字で判断を下す習慣が、創業準備の過程にもそのまま受け継がれた形と見られる。

こうして固めた計算を携えて彼は中国に帰国し、2023年3月1日に Moonshot AIを正式に設立、初代最高経営責任者に就任した。創業当時、彼の自然言語処理の経歴はすでに10年近くに達していた。学部2年の頃に始まった関心が10年を超えた時点で、会社という形で再び生まれ変わったことになる。

Moonshot AIは当初から、法人顧客を相手にする受注型ビジネスや、短期間で模倣できる製品を選ばなかった。代わりに一般利用者向けの個人向けインテリジェントアシスタントサービスに会社の力を集中させ、そのサービスには彼の英語名にちなんでKimiという名がつけられた。2023年10月に発表されたKimiは、20万字に及ぶ長文を一度に読み込める初のインテリジェントアシスタントとして紹介された。

Kimiが長文を途切れさせずに読み込める能力を備えられた背景には、彼が博士課程で書いたTransformer-XLの問題意識がある。固定された長さの範囲内でしか文脈を理解できなかったモデルの限界を、学界で解いた人物が、数年後には同じ限界を製品の次元で再び解いたことになる。研究者時代の問いと創業者時代の答えが一本につながる場面である。

Moonshot AIはその後、自社開発の基盤モデルにこだわる路線を維持した。Zhipu AIをはじめとする他の中国AI企業と同様、汎用技術の限界を自ら突き破る道はより険しいが、楊植麟はこの路線を貫くことこそ企業が長く生き残る鍵だと、さまざまな場で語ってきた。創業前に学界で培った原則を、創業後もそのまま守り抜いた形と見られる。

この原則のもと、Moonshot AIは2025年7月、1兆個のパラメータ規模を持つオープンソースの基盤モデル「Kimi K2」を発表した。これは創業から3年に満たない時点での成果であり、社内ではこの短期間で強力な基盤モデルを作り上げられた理由として、人材と方向性の二つを挙げている。研究者として築いた人的ネットワークと、創業者として立てた方向性が共に働いた結果だと言える。

2026年1月には後継モデルの「Kimi K2.5」が公開された。評価では複数の項目で海外の主要なクローズド型モデルと肩を並べる水準に達したとされる。訓練インフラを再構築し、強化学習アルゴリズムを見直したことで、複雑な作業を処理する際の安定性と効率を高めた結果だと会社側は説明する。海外メディアも、中国から出たこのモデルがアメリカの技術業界に思いがけない波紋を広げたと伝えた。

清華大学の学部生時代、グラフとマルチモーダルの間を行き来していた模索。研究者と創業者という二つの立場を行き来しながら彼が守り続けてきたのは、立場を変えることなく、問題に向き合う姿勢そのものであったように見える。

会社と研究組織をつくった過程

楊植麟が清華大学コンピュータサイエンス・テクノロジー学部で唐傑教授の指導のもとデータマイニングを学んでいた頃、彼はすでに大規模モデル研究が結局は組織の問題であるという事実を目の当たりにしていた。2015年に学士号を取得した後、カーネギーメロン大学に渡り、ルスラン・サラクトディノフとウィリアム・コーエンの指導のもとで博士課程を歩み、2019年に学位を取得した。この時期に彼が参加したXLNetとTransformer-XLの研究は、後に彼が

作る会社の技術的なルーツとなった。

博士課程と重なる時期、彼は北京市と中国科学界が設立した非営利研究機関、北京智源人工知能研究院(BAAI)の国家級プロジェクト「悟道」に科学者として参加した。ここで彼は初期の大規模モデルのアルゴリズム設計を担当し、大規模分散演算環境でモデルを訓練しインフラを扱う経験を実務で積んだ。学界で身につけた理論を実際の規模のハードウェア上で検証したわけである。

この頃、彼が学問共同体の中で最も大きく受け継いだのは、恩師である唐傑教授の姿勢であった。唐傑は後にZhipu AIを設立する人物でもあるが、彼は弟子に対し、エンジニアが事業を起こす際に目先の商業的機会や収益よりも優先すべき価値は、自国の智能技術の世界最高水準に引き上げられるかどうかだと教えた。この教えは、楊植麟が後に短期的な売上を追うアプリケーション会社ではなく、基盤モデルそのものを研究する組織を立ち上げようと決意する土台となった。

2019年から2022年にかけて、彼は清華大学KEG研究室との協力を続け、GLMやコード生成モデルCodeGeeXといったアーキテクチャの設計に参加した。大学を離れてからも途切れなかったこの協力関

係は、彼が創業後も中国の学界とつながった研究方式を維持する背景となった。

2022年11月30日、OpenAIがChatGPTを発表したとき、彼はシリコンバレーに滞在していた。彼は大規模言語モデルの飛躍が資本と人材を吸い寄せる歴史的な転換点になると判断した。この流れに身を投じ、ゼロから新しく始める人工知能企業を立ち上げるべきだと決心した。

彼がシリコンバレーの滞在先で行ったのは、感傷ではなく計算だった。モデルの訓練にかかる演算量、学習コスト、ユーザーが増えたときに負担すべき推論コストを一つひとつシミュレーションした。その結果、商業的な妥協なしに基盤モデルをゼロから学習させる組織を作るには、創業初期の数カ月以内にベンチャーキャピタルから最低1億ドルに及ぶシード投資を受ける必要があるという結論に至った。彼はこの計算を根拠に帰国し、直ちに投資家の説得に乗り出した。

会社の公式な設立時期については、資料によって表記が異なる。楊植麟本人はインタビューで、法人設立日は2023年3月1日だと語った。一方、清華大学の公式紹介記事は2023年4月創業と記してい

る。両方の日付を並べてみると、3月1日の創業決断が4月に入って正式な組織体制の整備として結実したと見るのが自然である。

初期のオフィスは、北京の華やかなテック団地ではなく、あるビルの片隅の一室だった。会社がその後大きな投資を受けてからも、このオフィスの入り口には目立ったロゴは掲げられなかった。ドアの前には白いピアノが一台置かれ、その上にピンク・フロイドのレコードが置かれていたと伝えられる。派手な宣伝よりも静かな集中を前面に出した雰囲気だった。

人員構成も当初から少数精鋭を志向した。楊植麟は、モデルの内部構造を自らの手で改良した経験や、大規模なコンピューティングインフラの上で事前学習を完遂した経験を持つエンジニアを30人から40人ほど核となる研究開発人材として集めた。この狭く密度の高い組織は、創業からわずか6カ月で20万字分の長文を処理できるスマートアシスタント「Kimi」を完成させ、世に送り出す結果につながった。

Kimiは2023年10月に発表され、当時これほど長い文脈を扱えるサービスとしては初めてのものだった。この成果は、楊植麟が博士課程で扱ったTransformer-XLの発想、すなわち固定された長さの区

間に閉じ込められず文脈をつなげていく方式を、商用製品の規模にまで拡張した結果であった。学位論文にあったアイデアが、数年後に消費者が日々使うサービスへと移し替えられた過程であった。

資金調達の時期も会社の成長速度を物語る。2023年2月に最初の資金調達を始めた当時、中国の資本市場は新型コロナウイルスの封鎖解除直後の不安定さの中にあった。楊植麟はこの時期を逃していれば、会社が生き残るのに必要な初期資金を確保できなかっただろうと振り返っている。当初はアメリカの資本市場の門をたたいたが、結局は中国国内の資本の支持を受け、2023年下半期までに二度のラウンドを通じておよそ20億元を確保した。

2024年2月には、アリババが主導する投資ラウンドが成立した。リシキャピタルと小紅書(RED)が共に参加したこのラウンドで、会社の投資前評価額は15億ドル、投資後評価額は25億ドルと評価され、この時点までに累計で確保した資金は90億元を超えた。Kimiの利用者が急速に増えていた時期と重なる投資だった。

資金が増える一方で、人員は大きく増やさなかった。2024年初め時点で、会社の常駐人員全体はおよそ80人にとどまっていた。楊植麟は、基盤モデルを作る組織で重要なのは人数ではなく実力ある人

材の密度だと考え、製品企画や運営、システムリーダーシップといったポジションにも必要な分だけ人を加えた。

時間が経つにつれ、研究組織は徐々に専門化された構造へと落ち着いていった。自然言語処理を担当するチームは基盤モデルと長文脈処理の構造を研究し続けた。コンピュータビジョンを担当するチームは映像と物理世界を認識する問題を担った。強化学習チームは、モデルが自ら推論過程を点検しながら答えを探っていくゆっくりとした思考方式を研究した。インフラチームは、数万台規模のコンピュータを一つの装置のように扱う問題を担当した。

これらのチームは外部で作られたツールに頼るのではなく、自前のディープラーニングフレームワークをゼロから作り上げて使った。大規模なコンピューティングクラスターを運用し性能を最適化する経験は、この会社が他のスタートアップと異なる点だった。学習の成果を外部に公開する方向性も同時に維持した。これは彼の指導教授だったサラクトディノフが重んじてきた、オープンな研究の伝統と重なるものだった。

2025年7月、会社は1兆個に及ぶパラメータを持つ基盤モデル「Kimi K2」をオープンソースとして公開した。重みを無条件で世界中

の開発者に配布したこの決定は、会社が創業初期から守り続けてきた姿勢、すなわち商業的利益よりも技術そのものの前進を優先するという姿勢と重なるものだった。

2026年1月には後継モデルKimi K2.5が発表された。このモデルを作るにあたり、インフラチームは訓練過程で生じていたボトルネックを取り除く方向でパイプラインを組み直し、強化学習チームはモデルが複雑な作業を遂行している間も崩れることなく安定して学習を続けられるようアルゴリズムに手を加えた。この改善の結果、Kimi K2.5は複数の評価項目でGPT-5.2やClaude Opus 4.5のスコアに到達、あるいはそれを上回った。

同時期に発表されたKimi K3は、フロントエンドのコーディング能力を競う評価で上位に食い込んだ。米国の大手モデル企業の間でもこの結果は驚きをもって受け止められ、中国国外のメディアもこれを大きく取り上げた。

清華大学が2026年7月に発表した紹介記事は、Moonshot AIを北京を代表する人工知能企業として紹介し、楊植麟の言葉を直接引用した。彼は、基盤モデルを自前で研究開発することは汎用技術の限界を超えるためにはより険しい道だが、この方向を貫くことこそが

企業が長く成長していくための核心になると語った。彼が唐傑教授から学んだ姿勢がそのまま受け継がれている一節である。

同じ記事は、Moonshot AI一社だけでなく北京という都市全体の流れも合わせて説明している。2026年初め時点で北京市に登録された人工知能モデルは212件に上り全国で最も多く、北京市科学技術委員会と中関村管理委員会は今後も大規模モデルとエージェント技術の研究を支援し続ける方針を明らかにした。Moonshot AIの成長は、こうした政策的支援と歩調を合わせて進んできたと言える。

同社の企業価値は2025年から2026年にかけて100億ドルを超えたと伝えられている。創業当初30人余りで始まり、80人規模を長く維持していた組織が、世界的な規模の基盤モデルを次々と送り出す企業へと定着するまでにかかった時間は3年に満たない。

確かなのは、会社が歩んできた方向である。少数の人材を集め、基盤モデル研究に資源を集中させ、作り上げた成果を公開する方式を繰り返しながら、Moonshot AIは創業者の学問的背景とシリコンバレーで立てた計算を一つずつ実現していった。

中核モデルと開発の過程

2023年4月、彼はこの経歴を土台にMoonshot AI(月之暗面)を設立しました。創業者の年齢は若い部類に入りましたが、彼が掲げた方向性は実務的な妥協とはかけ離れたものでした。企業顧客の要求に合わせてシステムを構築する受注型事業の代わりに、汎用人工知能そのものに到達する基盤モデルを自ら作り上げるという路線を選んだのです。

この選択の根は、彼が博士課程で共著者として参加した一本の論文にまでさかのぼります。2019年に発表されたTransformer-XLは、既存のトランスフォーマー構造が定められた長さの文脈しか扱えず、文章が長くなると前半部分の情報を失ってしまうという限界を、区間をまたぐ再帰的な記憶方式で解決した研究でした。機械が人間のように膨大な知識を長く保持し続けるためには、記憶できる文脈の大きさそのものを広げる必要があるという発想が、この論文から始まりました。その発想は、Moonshot AI設立後も会社全体を導く原則として残り続けました。

Moonshot AIが世に初めて送り出した成果は、創業から半年後の2023年10月に登場した対話型アシスタントKimi智能助手でした。このサービスは一度に20万字に及ぶ長文を情報の欠落なく処理できま

した。当時の中国市場ではこれほどの長さの文脈を扱う商用サービスがなかったため、たちまち話題となりました。

2年に満たない2025年7月、Moonshot AIはKimi K2をオープンソースとして公開しました。総パラメータ数1兆400億に及ぶ混合エキスパート(MoE)構造を持つモデルでしたが、実際に一つの質問に答える際に稼働するパラメータはわずか320億にすぎませんでした。残りは眠った状態で必要なときだけ呼び出される仕組みです。こうした設計により、モデル全体の知識量を増やしながら計算量は抑えることができます。

この構造の中にはエキスパートニューラルネットワークが384個組み込まれています。同時期に登場した中国の別の大規模モデルDeepSeek-V3が256個のエキスパートを備えていたの比べると、はるかに細かく分割された形です。Moonshot AIの研究陣はエキスパートを細かく分けるほど、数学や論理といった難しい推論問題を解く性能が明確に向上することを実験で確認し、この方向を推し進めました。

その代わり、注意を向ける頭数、いわゆるアテンションヘッドは64個に減らしました。DeepSeek-V3の128個と比べると半分です。

エキスパート数を増やすことで生じる計算負担を、アテンション側で軽減しようとする均衡点でした。ここにメモリ使用量を減らすマルチヘッド潜在アテンション(MLA)方式まで組み合わせ、サービス単価を引き下げました。

2026年3月、GTCの舞台に立った楊植麟は、これをさらに一段発展させた構造を紹介しました。名称は「アテンション・レジデュアル」です。時系列を追うモデルが長く使ってきた残差学習方式を、トランスフォーマーの層が深くなる方向にそのまま応用したものです。層が深くなるほど前の層で学んだ情報がぼやけていく問題を、こうした形で抑え込んだと彼は説明しました。

画像や映像を理解する部分にも独自の設計が盛り込まれています。MoonViT-3Dという視覚エンコーダーは、SigLIP系統の重みで初期学習を始めた後、画像を無理に切り取ったり比率を歪めたりせず、原本のまま細かく分割した断片を一系列につなげて処理します。静止画であれ動画であれ同じ方式で扱うため、画面比率がどうであれ元の姿そのままに認識できるというのがこの設計の要点です。

モデルをこれほどの規模に育てるには、それに見合う学習素材が必要です。Kimi K2の訓練に投入されたデータは15兆5000億トーク

ンに及び、大きく四つの領域に分かれていました。インターネットから収集した低品質ページをふるいにかけてウェブ文書。

このデータを精製する方式は新たに組み立てず、先行するKimi k1.5モデルを作る際に確立した手順をそのまま引き継ぎました。一から作り直すより、検証済みのパイプラインを使い続ける方が安全だという判断だったのでしょう。

トークン一つひとつから最大限のことを学ばせるため、学習最適化の方式にも手が加えられました。よく使われるAdamオプティマイザーの代わりに、Muonとこれを改良したMuonClipを学習過程に組み込みました。これは大規模モデルを学習させる際にしばしば現れる非効率な収束速度を引き上げるための選択でした。

エキスパート構造を持つモデルを学習させる際には、勾配が急激に爆発して学習が崩れることがよく起こります。Moonshot AIはこの問題を防ぐためQK-Clipという技法を組み込みました。アテンションのクエリとキーの値が一定の限度を超えないよう事前に切り取る方式で、こうした安全装置を設けたことで学習曲線がはるかに安定して進みました。

モデルが学んだ知識を実際の指示に沿って反応させるよう調整する指示文データも、三つの経路で作りました。人の手で精巧に組み上げた複雑な指示文。AutoIFという独自の方法で自動的に拡張した指示文、そしてモデルがよく間違える弱点だけを選んで突き詰める専用モデルが作った指示文です。三つの経路を重ねて使えば、人の手だけで作るときより穴が少なくなります。

文脈を処理する幅を広げる過程も段階を分けて進めました。文字と画像が混在した高品質データで3万2000トークンの長さの範囲内で教師あり学習を一度終えました。その後、同じ重みをそのまま維持しつつ入力長を12万8000トークンまで広げ、位置情報を補正する技法を加えて二度目の教師あり学習を行いました。この二段階を経ることで、モデルは画像混じりの非常に長い入力でも前半部分の情報を見失わないようになりました。

数学や図形のように複数の段階を経て解く問題は、別の方式で仕上げました。先行するKimi k1.5を教師モデルとして立て、このモデルが問題を計画し解答過程を評価し振り返り別の解法を探すという四段階を経て、長い解答過程を複数生成するようにしました。

こうして生まれた解答のうち、途中で計算が誤っていたり論理が食い違っていたりするものは、ルールベースの採点モデルがふるい落としました。正解にたどり着いた解答だけを残し、残りは捨てる拒否サンプリング方式でした。この過程を経てこそ、モデルは実際に正しく推論する習慣を身につけることができました。

こうして磨き上げられた成果が実際に市場で試されたのは2025年半ばのことでした。上海で開かれた世界人工知能大会の開幕直前に、Moonshot AIはKimi K3を発表し、このモデルは大規模モデルのランキングを競うArenaプラットフォームのフロントエンドコーディング部門で世界1位に輝きました。米国シリコンバレーの主要モデルをコーディングという具体的な領域で上回った点が、その中でもとりわけ注目を集めました。

年が明けて2026年1月、Moonshot AIはKimi K2.5を正式に発表しました。訓練インフラの基盤から組み直した改編でした。強化学習フレームワークとアルゴリズムに手を加えた結果、複雑な作業を処理する際の安定性と速度が目に見えて向上しました。評価結果では複数の項目で、GPT-5.2やClaude Opus 4.5といった米国を代表するクローズド型モデルに匹敵、あるいはそれを上回りました。

この成果について、楊植麟は「長期主義」という言葉を口にしました。基盤モデルを他社に頼らず自ら作る道は、汎用技術の壁を越えることがはるかに難しいものの、この道を貫くことこそ会社が長く生き残るための条件だという意味でした。同じ時期、北京に登録された大型モデルの数が212件に達し全国最多だったという統計も発表されましたが、これはMoonshot AIと同じ都市に根を下ろす智譜AI(GLMシリーズ)をはじめ、複数の企業がそれぞれ基盤モデル開発競争を続けている背景を示す数字でした。

Moonshot AIがこれまで守り続けてきた原則の一つは、複数の製品ラインに開発力を分散させるのではなく、Kimiというたった一つの消費者向けアプリにすべての開発力を集中させることでした。企業顧客向けのカスタムシステム構築事業には手を出しませんでした。人々が履歴書を検討し、分厚い本を翻訳し、長いコードをデバッグする過程で残す利用の痕跡が、次世代モデルを訓練する材料として還元される循環構造を作ることに力を注いだのです。

この循環は、会社設立からわずか3年に満たない期間で、Kimi智能助手からK2、K3を経てK2.5に至るモデル系譜を生み出しました。各モデルは、前のモデルが市場で確認した強みと弱みを次の学習

データとして吸収しながら磨き上げられ、その結果、アメリカでもこのモデルの性能に驚きを示す報道が出始めました。

Kimi K3の正確なパラメータ規模や学習に使われたトークン量は、関連報道には出ていません。こうした数字は、Moonshot AIがまだ外部に公開していない領域として残っています。

現在の立ち位置と中国AIにおける意味

2026年7月、北京市科学技術委員会と中関村管理委員会は、その年の初めに北京に登録された人工知能モデルが212件に達し、全国最高水準だったと発表しました。この統計発表の記事で一番先頭に置かれた名前が、楊植麟が率いるMoonshot

AIでした。記者は、今年1月にMoonshot

AIが新しいオープンソースモデルKimi

K2.5を正式に発表し、複数の評価項目でGPT-5.2やClaude Opus 4.5といった海外の主要なクローズド型モデルと同水準か、それを上回る成績を出したと伝えました。

同じ記事は、2月に智譜AIがGLM-5を発表したこと、そして同じ時期に智源のマルチモーダル大型モデルの研究成果が国際学術誌ネイチャーに掲載されたというニュースも合わせて取り上げました。Ki

mi K2.5とGLM-5が並んで言及されたのは偶然ではありません。北京という一つの都市の中で、汎用人工知能をめぐる競争が数社の実験室レベルを超え、地域経済指標として扱われ始めたことを意味していました。

楊植麟はこの記事の中で、人工知能の持続可能な成長について語りながら「長期主義」という言葉を口にしました。基盤モデルを自社で直接研究開発する道は汎用技術の限界を超えることがより難しいものの、その道を貫くことこそ企業が長く続くための方法だという趣旨でした。創業者個人の決意表明のようにも聞こえる発言ですが、新聞がこの発言を政策成果を説明する場に載せたという事実そのものが、彼の立ち位置を物語っています。

記事はMoonshot AIの短い歴史を改めて振り返りました。自然言語処理分野で10年働いてきた楊植麟が2023年4月に会社を設立し、その年の10月にはスマートアシスタントKimiが登場し、20万字に及ぶ長文を処理できる初のサービスとして地位を確立しました。2025年7月には1兆パラメータを持つオープンソースの基盤モデルKimi K2が世に出ましたが、このモデルこそKimi K2.5の前世代にあたるものでした。

3年に満たない期間でこの規模の基盤モデルを作り出せた秘訣を問われ、Moonshot AI側は人材と哲学が等しく重要だったと答えました。同社の研究開発人員は自然言語処理、コンピュータビジョン、強化学習、インフラといった分野にまたがっており、彼らが世界水準のディープラーニングフレームワークを構築し、超大規模な計算クラスターを運用・最適化してきた経験を積んでいると説明しました。

Kimi K2.5を作る過程では、訓練インフラそのものを組み直す作業がありました。強化学習フレームワークとアルゴリズム設計を手直した結果、複数の段階を経る複雑な作業を処理する際に、モデルがぶれることなく最後までやり遂げる能力と処理速度が目に見えて向上したと同社は説明しました。長文を読み解く力から出発した会社が、今では複数の段階を経て問題を解く力まで手にするようになったのです。

同じ記事は智譜AIを並べて紹介し、比較の枠組みを作りました。2019年に清華大学コンピュータ学科から派生した智譜は、2020年という早い段階で大型モデル分野に参入し、翌年に100億パラメータのモデルGLM-10Bを発表、2024年にはGLM-4で性能を60パーセン

ト引き上げました。智譜の最高経営責任者である張鵬は、会社が2021年からモデルサービス事業を準備していたと語りました。これは大型モデルが本格的に収益を生むようになる2年前のことでした。

智譜は今年1月に香港証券取引所へ上場し、汎用人工知能の基盤モデルを中核事業とする世界初の上場企業になりました。Moonshot AIがまだ上場手続きを踏んでいない間に、智譜は資本市場の敷居を越えたのです。両社とも基盤モデルを自ら作る道にこだわったという点では、同じ立場に立っています。

北京市科学技術委員会と中関村管理委員会の関係者は、これから続く第15次五カ年計画の期間においても大型モデルとエージェント技術の開発を支援し、技術を実際の産業へ移す作業を後押しし、人工知能分野の革新企業群を育てていくと明らかにしました。Moonshot AIと智譜は、この政策が示す成果であると同時に、今後数年間同じ支援を受けながら競い合う企業の手本として紹介されたのです。

Moonshot AIの歩みは2026年に入っても止まりませんでした。米AP通信は、中国発の新しい人工知能モデルが米国の技術業界を驚かせたという見出しの記事を掲載しました。上海で開かれる世界人

工知能大会の開幕を控えて、Moonshot AIがKimi K3を発表し、このモデルがフロントエンドのコーディング能力を競う国際評価ランキングで上位に入ったというニュースでした。

Kimi K1.5、K2、K2.5、K3へと続く流れを見ると、Moonshot AIがその都度異なる勝負どころを選んできたことが分かります。最初は長文を読み解く力でした。次は誰もが自由に使える1兆パラメータのオープンソース基盤モデルであり、その次は複雑な作業を最後まで安定して処理する力、そして今回は実際の画面を作り出すコーディング能力でした。毎回異なる項目で世界水準を上回ることで存在感を積み重ねてきたのです。

2026年3月、楊植麟はエヌビディアが米国で開催したGTC 2026の舞台に立ち、Kimi K2.5をどのようにしてその規模まで育て上げたのかを自ら説明しました。自国政府の発表の場で名前が挙がっていた創業者が、今度は米国の半導体企業が主催する世界最大規模のコンピューティングイベントの演壇に立ち、自らのモデルを説明したのです。

一人の人物が数か月の間に、北京市の統計発表記事と米国半導体企業の国際カンファレンスの演壇の両方に名前を連ねるというのは

、よくあることではありません。中国国内では人工知能の主権を支える代表企業の代表者として、海外では大規模モデルを自ら訓練し上げた技術者として、それぞれ異なる形で光を当てられているということであり、この二つの顔が今、楊植麟の立ち位置を形作っています。

同じ月に開かれたエヌビディアのGTC 2026本イベントでは、ジェンスン・フアンが基調講演を務め、コンピューティング産業全体の方向性を示しました。楊植麟のセッションはその大規模イベントの中に置かれた数ある発表の一つに過ぎませんでした。中国企業の創業者がこの舞台に招かれ、自らのモデル開発過程を直接説明したという事実そのものが、Moonshot AIを世界のコンピューティング産業が参考にすべき事例と見なしている表れとして読み取れます。

こうした行動が特別に映る理由は、Kimi K2をオープンソースとして公開した決断にもあります。1兆パラメータ級の基盤モデルを無償で公開する企業は多くありません。Moonshot AIはこのモデルを隠すことなく、世界中のどの研究者でも自由に使えるように開放しつつ、その上でKimi K2.5とK3を自社サービスとして開発し、商用市場で競争する戦略を同時に走らせています。

公開と商用化というこの二つの方式は互いにぶつかることなく、むしろ互いを後押ししています。公開と商用化を並行して進めるこの方式は、中国の大型モデル企業の間で一つの手本として定着しつつあります。無償で公開されたモデルの重みを利用した世界各地の研究者たちは、結果としてMoonshot AIという名前を自分たちの論文やコードの中に残し続けることになるのです。

北京という都市全体で見れば、Moonshot AIは一つの点ではなく、全体像の中の一片にすぎません。212件に及ぶモデル登録件数、智源の論文がネイチャーに掲載されたこと、智譜の香港上場、そしてMoonshot AIの相次ぐモデル発表が同じ一つの記事の中で合わせて扱われたという事実は、中国政府がこの分野を個別企業の成功物語としてではなく、都市全体の産業指標として管理していることを意味しています。

楊植麟という一個人の名前が、こうした統計や政策発表の場に繰り返し登場するという事実は、彼が今や一人の創業者という枠を超え、北京が自らの人工知能の実力を説明する際に持ち出す代表的な事例になったことをも意味しています。

こうした位置づけは自然に与えられるものではありません。北京に登録されたモデルが212件にも上る状況の中で新聞記事に名前を載せるには、その会社が出した成果が政策担当者と取材記者の双方を納得させるだけの根拠を備えていなければなりません。212件のうちわずか数社だけを選んで紹介する記事の中で、Moonshot AIが毎回その数社の中に入り続けてきたという事実こそが、楊植麟と彼が率いる会社が今立っている場所を作り上げた実質的な理由なのです。

智譜と並べたときに見えてくる違いもあります。智譜はコンピュータ学科という大学内の研究室から派生し、6年以上にわたってモデルサービス事業を準備した末に上場へとたどり着きましたが、Moonshot AIはそれよりはるかに短い期間のうちに、オープンソース公開と商用サービスという二つの経路を同時に切り開きました。出発点は異なるものの、両社とも外部で作られたモデルを借りてくるのではなく、最初から最後まで自ら作る道を選んだという点で、北京が後押しする産業の顔を共に形作っています。

世界市場におけるMoonshot AIの意味も、しだいに明確になってきています。AP通信が、中国発のモデルが米国のテック業界を驚

かせたと報じたことは、シリコンバレーを追いかけていた中国企業が、今やシリコンバレーから注視される存在へと変わったことを示す兆候として読み取ることができます。Kimi

K2.5はGPT-5.2やClaude Opus

4.5と並ぶ位置に置かれています。Kimi K3がコーディング能力の国際ランキングで上位に入ったことも、こうした評価を裏づけています。

ただし、こうした成果をそのまま完全な勝利として読み取るのは時期尚早です。確認できない部分は確認できないとそのまま残しておくほうが誠実でしょう。

Kimi K3がフロントエンド・コーディングのランキングで上位に入ったというニュースも同様です。一つの評価指標で先行したという事実が、モデル全般の優位を意味するわけではありません。AP通信がこのニュースを米国テック業界に衝撃を与えたという角度から取り上げたのは、コーディングという狭い領域から出た結果が、シリコンバレー全体に投げかけた心理的な衝撃が大きかったという意味に読み取るほうが正確です。

智譜(Zhipu)の事例が示すように、基盤モデルを自前で作る企業が上場にまで至るには、資本市場の信頼が必要です。Moonshot AIが今後も同じ道をたどるかどうかは、これから数年のうちに発表される財務指標と上場の有無が、この問いに答えを与えるでしょう。

GTCの舞台で語られた説明は、先の新聞記事が伝えた内容と同じ方向を指しています。訓練インフラを組み直し、強化学習フレームワークを手直しして複雑な作業における安定性と効率を上げたという説明が、北京の地方紙にも、エヌビディアが主催した国際イベントにも、同じように登場しました。国内向けの説明と海外向けの説明が互いに食い違わないという点は、Moonshot AIが示す技術的根拠が、少なくとも一貫性という面では信頼を得ていることを意味すると読み取れます。

それでも、これまで明らかになった事実だけを見れば、創業から3年余りの時点で楊植麟の立ち位置ははっきりしています。彼は、北京が人工知能産業政策を説明する際に名前を挙げる創業者です。清華大学と唐傑(タン・ジエ)研究室から出発し、米国で研鑽を積んで帰国した一人の研究者の歩みが、今や一つの都市の産業統計の中に刻まれる位置にまで至ったのです。

Moonshot AIが次に打ち出すモデルが何であれ、その発表はもはやスタートアップのニュースにとどまらない可能性が高いといえます。Kimi K1.5からK3まで続いてきた発表が、そのたびに北京の別の記事、別の政策文書、別の国際ニュースで引用されてきたように、楊植麟の次の一步もまた、中国の人工知能産業が自らを説明する言葉の一部として使われ続けるものと見られます。

会社を設立してわずか3年余りで、自国政府の統計発表と米国の半導体企業が主催する国際舞台の両方に名を連ねる立場にまで来たことは、今の中国の大規模モデル産業がどれほど速く動いているかを示す証拠でもあります。清華大学と唐傑研究室から始まった歩みがここまで至るのに要した時間は、それほどまでに短かったのです。

もちろん、この流れが今後もそのまま続くという保証はどこにもありません。智譜はすでに上場を終えており、他の大規模モデル企業もそれぞれ異なる速度で同じ目標を追いかけています。Moonshot AIと楊植麟が今後もこの位置を守り続けられるかどうかは、次に出るモデルの性能と、次に発表される財務指標、そして政府が続ける支援政策が合わさって決める問題です。

第6章 閆俊傑(イエン・ジュンジェ、Yan Junjie)、商湯からMiniMaxへ、消費者向けAIとLightning Attentionを推し進めた創業者

成長期と学業

閆俊傑は1989年、中国河南省のある小さな県城に生まれました。この土地は大都市が享受していた先進的なコンピュータ教育や体系的な先取り学習の環境とはかけ離れており、幼少期に訪れたことのある一番大きな都市といえば南京程度だったというほど、狭い地域社会の中で育ちました。

父親はその小都市の中学校教師で、母親は公務員でした。二人とも共働きで常に忙しくしていましたが、小都市が子どもに与えられる知的・道徳的な教育資源だけは惜しみなく用意してくれました。教師である父のそばで育った環境は、彼が学業に自然と没頭するようになった土台となりました。

祝日のたびに帰省して交わした祖父との会話は、彼にとって長く残る記憶となりました。祖父は生涯歩んできた話を自伝的な文章として残したいという願いを繰り返し口にしていましたが、コンピュータで文字を入力し文章をまとめる作業には明らかな限界がありました。

周りの親戚はこの願いを単なる酒の席での愚痴として聞き流していました。閆俊傑はこの場面で、ある老人の無力な挫折を目の当たりにし、深く印象に残ったと後に振り返っています。高度に発達した人工知能技術が北京や上海のような大都市の人々の娯楽の道具としてのみ使われてはならず、情報技術から取り残された田舎や小都市の普通の人々が抱える不便を解消する道具になるべきだという考えを、このとき抱くようになります。この信念は、彼が後に会社を率いる中で繰り返し語る原則、すなわち技術は結局のところ人の暮らしを楽にするために使われるべきだという考えにつながっていきます。

正規の教育インフラが十分でない環境の中でも、彼は数学の科目で際立った才能を示し、学校の教師たちから期待を寄せられました。

閆俊傑自身が自分の人生を貫く最も価値ある資産として挙げるのは、独学の能力です。良い教師に出会い、知識を楽に受け取れるという方式は一種の運に頼ったものであるという事実を、彼は早くから見抜いていました。図書館や教科書を自ら読み解き整理しながら積み上げた学習習慣は、後に彼が誰も歩んだことのないMoEとマル

チモーダル結合という未知の道を推し進める力につながったと彼は語ります。

青少年時代には戦略性の強いPCゲーム、Dota 2を好んで遊んでいました。このゲームに費やした時間は、思いがけない形で彼の進路とつながります。OpenAIが強化学習でDota 2のボットを訓練し、世界最高峰の選手たちを打ち破った論文を後に精読し、ゲームで感じていた興味が人工知能研究へと移っていくきっかけとなりました。

2006年、彼は故郷の河南省を離れ、南京にある東南大学数学学院に入学しました。この大学は中国内で名の知れた理工系大学であり、彼はここで数学という学問を本格的に掘り下げ始めます。

数学科在学中、閆俊傑は長く抱いてきた「一流の数学者」という夢を手放す時期を経験します。大学の優れた同期たちと自分を比較する中で、一部の人だけが持って生まれる数学的直観の壁を実感するようになります。彼は自分が数学界で一流の学者や研究者になるのは難しいという事実を、淡々と受け入れたと語っています。

数学者への道が閉ざされた後、彼は毎日図書館にこもり、数学の論理を大切にしながらも、世の中を実際に変えられる新しい学問は何かを探し求めました。その過程で、機械が人の視覚や言語のパターンを学習する学問であるパターン認識と人工知能に関する書籍に出会います。

彼は、宇宙の根本が数学であるならば、その数学をコンピュータの演算能力と結びつけ、機械自らに知性を生み出させることこそ最も価値のある仕事だと考えるようになりました。この気づきをきっかけに、彼は学問の方向を人工知能へと完全に転じます。

2010年6月、彼は東南大学で理学士の学位を取得して卒業しました。その後、中国科学院自動化研究所の大学院に進学し、工学博士の学位を取得します。この研究所は中国国内で人工知能研究を牽引する代表的な学術機関の一つでした。

一方、中国の一部メディアでは、彼が重慶郵電大学に通っていた、あるいは卒業したという報道がなされたことがあります。閆俊傑は放送人、羅永浩との公開対談で、この報道は事実と異なると自ら明らかにしました。彼は重慶郵電大学に入学したことも在学したこともないと明確に訂正しています。

博士課程を終えた後、彼は清華大学コンピュータ科学技術学科で博士研究員(ポスドク)課程を経ました。この時期、コンピュータビジョンと大規模ニューラルネットワーク構造に関する学問的基盤を一段と固めています。

中国の大規模モデル業界を率いる他の創業者の中には、カーネギーメロン大学で博士号を取得した楊植麟(ヤン・ジーリン)や、プリンストン大学で博士号を取得した姚順雨(ヤオ・シュンユー)のように、海外留学を経た人物が少なくありません。閻俊傑は初等教育から博士研究員課程に至るまで、海外留学の経験が一切なく、中国国内でのみ学業と研究を続けてきた学者です。

博士課程後半にあたる2014年、彼は中国の検索企業バイドゥ(百度)の研究所で短いインターンシップを経験します。この時初めて、数百枚のグラフィックカードが共に稼働する大規模GPUクラスターを直接扱い、大規模並列訓練を体で身につけました。彼はこの経験を通じて、学術論文の中のアルゴリズムと、膨大な演算資源を使う実際の産業インフラとの間には大きな隔たりがあるという事実を実感したと語っています。

バイドゥでの経験を足がかりに、彼は2014年11月、中国の学界からちょうど分かれ出て顔認識技術で名を上げ始めていたスタートアップ、SenseTime(商湯科技)に初期メンバーとして加わります。

商湯に入ってからすぐ、彼は物体検知と顔認識を研究する部署をゼロから立ち上げ率いました。数学的分析をハードウェア段階まで緻密に組み合わせる方式で顔認識アルゴリズムの精度を引き上げ、中国の公共安全とモバイルセキュリティサービスがすぐに使える99パーセント以上の水準にまで到達させました。

この成果が認められ、彼は商湯研究院副院長を経て、スマートシティ事業群の最高技術責任者(CTO)、そしてグループ副総裁の座まで急速に昇進します。大規模な公共安全網と企業向けスマートシティソリューションを陣頭指揮しながら、技術が実際の売上と事業規模につながっていく過程を現場で学びました。

商湯で積んだ経験は、後に彼が新しい会社を立ち上げ率いる礎となります。彼が商湯を離れ新たな道を準備し始めたという事実だけを短く記しておくにとどめます。

河南省の小さな県城に生まれ、教育資源の乏しい環境を自らの力で切り開いた経験。そして図書館でただ一人、人工知能という学問を発見した瞬間は、彼のその後の歩みを理解する上で重要な手がかりとなります。正規教育の助けを借りず独学で築き上げた姿勢は、彼が後に検証されていない技術路線を自ら選び取り推し進めるやり方と通じ合っています。

研究者、そして創業者として定まっていく時期

閻俊傑は河南省の小さな都市で育ちながら、自ら道を探す学び方を身につけました。2006年、南京の東南大学数学学院に入学し、2010年6月に理学士の学位を取得するまで、彼は純粋数学の世界の中で自分の居場所を見定める時間を過ごしました。大学で緻密な数学理論を掘り下げるほど、彼は自分が世界最高水準の数学者になるのは難しいという事実を淡々と受け入れるようになりました。その自覚は挫折ではなく、次の道を見つけるきっかけとなりました。

進路を定め直す作業は図書館で行われました。彼は毎日図書館にこもり、数学的思考を生かしながらも、世の中に実際に手を触れられる学問は何かを探し続けました。そうするうちに、コンピュータが散らばったデータから規則を読み取るパターン認識という分野に出会い、人の知性を機械に移し植える人工知能研究へと心を固めて

いきます。数学科から人工知能へと渡ったこの転換は、後に彼が大規模モデル企業を率いる際にも一貫して受け継がれる思考の根となりました。

東南大学を卒業した直後、彼は試験を経ずに大学院へ進む推薦入試枠で中国科学院自動化研究所に入り、修士と博士を続けて修める五年課程を始めました。中国科学院自動化研究所は、人工知能と画像処理の分野で中国を代表する研究機関でした。閆俊傑はここで工学博士号を取得したのち、清華大学コンピュータサイエンス学科で博士研究員の課程を続けました。ただしこの時期は常駐研究員としてキャンパスに留まる形ではなく、実務と並行する形だったため、彼が実際に清華大学キャンパスに滞在した時間は長くありませんでした。

彼の学歴をめぐっては、一時期、重慶郵電大学出身だとする報道が流れたこともありました。閆俊傑は羅永浩との対談でこの報道を直接正し、重慶郵電大学に通ったり在籍したりしたことは一切なく、これは外部メディアの誤報であるとはっきり述べています。

博士研究員の課程にあった2014年、彼はバイドウ研究所にインターン研究員として加わり、学問の世界から産業の世界へ足を踏み入

れました。当時のバイドゥ研究所はアンドリュー・ンが率いており、のちにアンソロピックを創業するダリオ・アモデイも中核研究者として在籍していて、世界的なAI人材が一堂に集まっていた時期でした。

バイドゥで彼が最初に向き合ったのは論文ではなくハードウェアでした。数百枚のグラフィックカードが同時に稼働する大規模演算クラスタを初めて目にした瞬間、彼は理論論文を数本仕上げるのとは次元の違う衝撃を受けました。インターンの身でありながら、彼は研究所が保有するグラフィックカード資源のかなりの部分を割り当てられ、大規模な画像学習実験を思う存分回すことができました。

この経験は彼にひとつの確信を残しました。ダリオ・アモデイが、大規模モデルの性能はデータとパラメータ規模に応じて予測可能な形で向上するというスケーリング則を学界に示す何年も前に、閔俊傑はすでにバイドゥで大規模な音声データをニューラルネットワークに投入すると性能が跳ね上がる現象をこの目で確かめていました。彼はこのとき、AIの飛躍は細かな数値調整からではなく、データと演算を大胆に増やすことから生まれるという考えを、身をもっ

て体得しました。

2014年11月、閆俊傑はセnstタイムに加わりました。セnstタイムは中国の学界から派生し、コンピュータビジョン分野の代表企業へと成長していた会社でした。彼はここで顔認識と物体検出を扱う研究組織を任され、自らエンジニアリングを率いました。

セnstタイムでの初めの一年半は順調ではありませんでした。彼が率いる顔認識アルゴリズムは業界のベンチマークで何度も下位に沈み、会社のかかなりの演算資源を割り当てられながら成果を出せないという重圧に耐えなければなりませんでした。彼はこの時期を、数学を学んでいた頃には味わったことのない挫折だったと振り返っています。しかし、この失敗の時期を経て、彼が率いた顔認識技術は精度99パーセントを超える水準にまで達しました。

この過程で閆俊傑は二つのことを学びました。ひとつは、目先のスコアを少しずつ手直しするやり方では根本的な性能差を埋めることはできず、時間がかかっても構造そのものを変える道を選ぶべきだということでした。もうひとつは、一人でコードを書き一人で実験を回していた従来のやり方では、この規模の問題は解けないという気づきでした。彼はかつて、なぜ研究にチームが必要なのか理解

できなかったと語るほど協業が不得手でしたが、セNSTタイムでの経験を通じて、複数の人の力量をひとつの組織へとまとめ上げることの価値を体得しました。

顔認識技術を軌道に乗せたのち、彼の立場は急速に上がっていきましました。セNSTタイム研究院副院長を経て、スマートシティ事業群の最高技術責任者、さらにグループ副総裁にまで上り詰め、彼は数万人規模のエンジニア組織を管理し、大規模な公共インフラ事業を率いる立場に立ちました。研究者として出発した彼が、組織と事業を併せて手がける経営者へと軸足を移した時期でした。

セNSTタイムが香港証券取引所への上場を目前にしていた2021年12月9日、閔俊傑は会社を去りました。上場を控え、持ち株の値上がり益が保証されていたその時点で、彼はすべての地位と持ち株を手放し、退社を宣言しました。彼がセNSTタイムを離れた理由は、技術の潮流が変わりつつあるという判断にありました。欧米発の大規模事前学習型生成モデルの登場を目の当たりにし、彼は顔認識や物体検出を中心とした従来型のAIが終わりを迎え、人が世界を認識するように言語と画像と音声を同時に扱う生成AIが次の時代を切り開くだろうと見通していました。

退社直後に彼が描いた構想は、ひとつのモデルの中で言語、映像、音声、音楽を同時に扱う全モダリティ基盤モデルでした。真の汎用人工知能であるならば、ある一種類の情報形式だけに閉じ込められてはならないという考えがその根底にありました。2022年初め、彼は上海でMiniMax、中国語名で「稀宇科技」を設立し、共同創業者兼CEOに就任しました。

MiniMaxがその後歩んだ道のりは、この転換の結果でした。生成AIブームの最中、他社が「中国版ChatGPT」を急いで打ち出す一方で、上海のこのチームはテキストと音声と映像を同時に扱うという、より難しい道を選びました。同社が自社開発した全モダリティモデルをもとに、複数の国でAIネイティブ製品を展開しています。設立からわずか4年後の2026年1月に香港証券取引所へ上場するに至った背景には、センスタイムで培った組織運営の経験と、バイドゥで得たスケール拡大への確信がともに存在していました。

彼が通った小学校と中学校。ただし彼は厳しい環境の中で独学の力を鍛え、東南大学、そして中国科学院へと進みました。

数学から人工知能へ、さらに産業現場の経営者へと移っていったこの経路は、彼がのちにMiniMaxで強調するようになる原則へとつ

ながっています。彼は羅永浩との対談で、技術の進むべき方向を信じるのであれば、目先の手軽な利益に妥協せず、上限の高い道に資源を注ぐべきだと語っています。バイドゥで目撃した規模の力と、セスタイムで経験した組織の力、この二つの経験がひとつに合わさり、彼が創業後に歩むことになる道の下絵をなしました。

会社と研究組織を築いた過程

2021年12月、閆俊傑はセスタイムの副総裁兼スマートシティ事業群CTOの座を退きました。当時のセスタイムは上場を控え、まさに莫大な資金を手にしようとしていたところであり、傍目には理解しがたい決断でした。彼はOpenAIが次々と公開していく事前学習型の大規模言語モデルを見つめながら、物体を区別し画像を認識することにとどまっていたAIの時代が終わりを迎え、自ら知識を推論する新しい時代が開かれると確信しました。

2021年12月9日、閆俊傑と志を同じくする初期のエンジニアたちは一室に集まり、九時間にわたって議論を重ねました。会社が進むべき五年の計画を、その場で文書にまとめました。これを土台に、2022年初め、上海にMiniMax、正式名称「上海稀宇科技有限公司」が設立されました。

社名MiniMaxは、数学者ジョン・フォン・ノイマンが練り上げたゲーム理論のミニマックス・アルゴリズムに由来します。このアルゴリズムは、相手が最悪の手を打つと仮定してもなお、数学的に最良の選択を導き出す方法を指します。数学を専攻した閆俊傑にとってはなじみ深い概念であり、創業初期の厳しい条件の中でも最善の道を探すという決意が、この名前にそのまま込められました。

創業初日から、MiniMaxは「みんなとともに知能をつくる」を意味する"Intelligence with Everyone"を会社の目指す姿として掲げました。閆俊傑は、AIが北京や上海の高学歴層だけのための道具にとどまってはならず、情報や技術から取り残された地域の人々にも届くべきだと、幾度となく語っています。

MiniMaxが組んだ研究組織は、決裁ラインが長く部署間の壁が厚いよくある大企業とは性格が異なりました。人員が400人を超えたのちも、閆俊傑自身とその下の二段階、この三層の指揮系統だけを維持しました。報告のための報告、決裁のための決裁をできる限り減らそうとする選択でした。

この組織の中では、アルゴリズムを扱うエンジニアとソフトウェアをつくる開発者、製品を企画する担当者との境界があいまいで

す。閔俊傑は、MiniMaxで製品企画を担うにはアルゴリズムがどう動くかを数学的に理解し、自らコードを書いて自分のアイデアを動く形で示せなければならないという原則を掲げました。

こうした人員密度を補うため、MiniMaxは社内で使うエージェント・インターンを自前でつくり配置しました。これらのエージェントは、エンジニアの依頼を受けてコードの一部を代わりに書き、採用の際に殺到する履歴書を選別して人事担当者に渡すところまで担いました。

法人顧客を相手にする営業組織も同じやり方で動いています。通常、企業向けソフトウェアの営業チームは華やかなプレゼン資料をつくり込むことに力を注ぎますが、MiniMaxのプリセールス・エンジニアたちはプレゼン資料の代わりに、顧客がその場で使える専用のAPIデモを自らコーディングして渡します。

創業初期、MiniMaxはモデル研究と一般向けサービスをひとつに結びつけて育てる道を選びました。この方向から生まれたのがTalkie、星野、Glow、ハイロー(海螺)AIといったサービスで、これらを通じて世界で累計2億人を超える利用者を集めました。

2024年末、DeepSeekがチャットボットアプリのダウンロード数にこだわらず、ひたすら基盤技術の水準を引き上げることだけに集中して世界の注目を浴びる過程を見つめながら、閆俊傑は会社の方角性を改めて見直しました。優れた基盤モデルは自然と優れた製品を生むが、優れた製品や多い利用者数が基盤モデルの知能を引き上げてくれるわけではない、というのが彼の得た結論でした。

この判断に基づき、MiniMaxは短期的な利用者数を増やすためのマーケティングと広告予算の執行を減らしました。資金と人員を基盤モデルの事前学習研究に集中的に投入する方向で組織を再編しました。閆俊傑自身が語ったように、巨大モデルの時代において製品の本質は結局モデルそのものであるという判断が、この転換の根底にありました。

資金調達においてもMiniMaxは急速に規模を拡大しました。アリババがシリーズA+ラウンドを単独で主導し6億ドルを投資し、この一度の資金調達でMiniMaxの企業価値は25億ドルを超えました。 Tencent、セコイア・チャイナ、ヒルハウス、IDG、さらにゲーム会社のミホヨまで名を連ね、株主構成をしっかりと固めました。

2026年1月9日、MiniMaxは設立からわずか4年で香港証券取引所に上場しました。生成AI基盤モデル企業としては世界最速の上場例とされました。上場初日の株価は公開価格の2倍以上に跳ね上がり、時価総額は1,000億香港ドルを超えました。

この上場により閻俊傑個人の資産も大きく増えました。ブルームバーグの億万長者指数は、上場当日時点で36歳だった彼の持株価値を約35億ドルと算定しました。彼は上場直後のインタビューで、調達した資金を短期的なアプリケーション加工や企業向けカスタム構築事業には使わないと述べました。基盤モデルの研究開発と計算インフラの拡充、優秀な人材の獲得に使うと明らかにしました。

MiniMaxの人材採用の方式もシリコンバレーでよく見られる方式とは異なりました。欧米の名門研究所での経歴や華々しい学位を最優先には置きませんでした。閻俊傑自身が海外留学の経験なく中国国内で数学を学んだ人物だったこともあり、彼は肩書きよりも自ら問題を掘り下げて解決する実務能力をより重く見ました。

創業初期、十数名程度だった中核人員の中には、数学やコンピュータ工学を正式に専攻していない人物も複数混じっていました。彼らを一つに結びつけたのは、知能の飛躍を自分の目で直接見たいと

いう信念でした。今もMiniMaxの中で基盤モデルの構造を変える中核エンジニアの多くは、卒業したばかりか実務経験が1、2年に過ぎない若い人材です。

閆俊傑は中核アルゴリズム職を採用する際、最終面接を自ら行います。この場で彼が最も注目するのは、コーディング能力よりも同僚と共に働ける姿勢です。一人で部屋にこもって完璧なコードを書く人より、自分の失敗を認め同僚の意見を受け入れて共にシステムを直していく人をより高く評価するというのが彼の説明です。

一部のスター研究者に株式を集中させる方式ではなく、MiniMaxは研究開発と運営、管理とマーケティングを含む全社員に株式を分配する体系を作りました。特定の人物に依存するより、組織全体が共に成長の分配を受けられるように設計しました。

計算インフラを固める過程は平坦ではありませんでした。中国業界で初めて大規模な専門家混合構造、すなわちMoE方式のabab6モデルを事前学習していた時期に、異なるグラフィックカード間の同期がずれて学習全体が崩壊することが3、4回起こりました。一度崩壊するごとに、最低2ヶ月の研究時間と1,500万ドルに及ぶ計算コストが失われました。

こうした失敗が繰り返されても、閆俊傑は安全性が確認された密構造に戻ることはありませんでした。彼は会社が使える計算資源の80パーセント以上をMoE経路に継続的に投入し、その結果2024年1月、中国初の大規模MoEモデルであるabab6を正常に完成させ発表することができました。

アリババの投資を契機に、MiniMaxはアリババクラウドが提供する計算資源と利用料支援を自社の分散学習体系に取り込みました。これによりハードウェアの利用コストを大きく下げつつ、大規模な事前学習を続けられる条件を確保しました。

エヌビディアのH20グラフィックカードの性能を最大限引き出すため、MiniMaxのエンジニアたちはLightning Attention演算に使われるCUDAカーネルを基礎から自ら書き上げ磨き上げました。その結果、分散学習クラスタ全体でモデルFLOPS利用率が75パーセントを超える水準まで上がりました。これはハードウェアの限界をソフトウェアで押し広げた結果でした。

このようなインフラ整備の成果は、2025年6月に公開されたMiniMax-M1で明確に表れました。ハイブリッドアテンション構造を採用したこのモデルは、オープンソースの重みとして全世界に公開さ

れました。10万トークン規模の推論性能をDeepSeek-R1と比較したところ、わずか25パーセントの計算量だけで同等水準の推論能力を示しました。

閻俊傑が創業して4年でMiniMaxを香港証券取引所の代表的なAI企業に育て上げる過程には、幾度もの研究の失敗と莫大なコストの損失が伴いました。それでも彼は、会社のすべての成果が結局は基盤モデルそのものの知能向上から生まれるという原則を揺るがずに守り抜き、その原則がMiniMaxという組織を今の位置まで導いた力であったと見ることができます。

中核モデルと開発過程

2022年の設立当時、MiniMax(稀宇科技)は中国国内でChatGPTをそのまま模倣した製品を急いで出す風潮とは異なる道を選びました。閻俊傑代表はテキストと音声と映像を一つのモデルの中で共に扱う全モーダル構造を築くことに会社の資源を賭けました。彼は本物の汎用人工知能であれば、最初から特定の入力形式に縛られてはならないと判断しました。この判断は後にさまざまなモーダルが混ざり合う流れが本格化する中で、会社が先を行く足がかりとなりました。

彼が繰り返し語った原則は、欧米の大型モデルの重みを持ち込んでマーケティング用の応用製品として包装する方式を拒み、インフラの基礎からアーキテクチャを自ら設計するということでした。彼は、巨大モデルの時代において製品の実体は結局モデルそのものであり、モデルの能力を高められないままマーケティングで埋め合わせようとすれば長続きしないと述べました。そのため会社のすべての商業的な反転はモデルの質的な飛躍から生まれるべきだという基準を立てました。

この原則が実際に体現された成果が、2025年1月15日に公開されたMiniMax-01シリーズ、すなわちMiniMax-Text-01とMiniMax-VL-01です。両モデルは混合専門家構造(MoE)とLightning Attentionと呼ばれる線形アテンション機構を共に用いる独自設計を採用し、この組み合わせはその後、世界のオープンソース陣営と商用API市場に少なからぬ波紋を広げました。

MiniMax-Text-01は全体パラメータが4,560億個に達する超大型MoEモデルですが、実際の推論過程で一度に有効になるパラメータは459億個程度に制限されます。このように設計すると、モデル全体の表現力は大きく保ちながらも計算コストは抑えることができま

す。層ごとに32個の専門家を置き、入力トークンごとにその中から2つを選んで使うTop-2ルーティング方式を採用しました。順伝播ニューラルネットワーク内部の隠れ次元は9216、モデルの基本隠れサイズは6144に設計されています。

既存のトランスフォーマー構造は、扱う文章が長くなるほど計算量が長さの2乗に比例して増える問題を抱えています。MiniMaxはこの問題をLightning Attentionで解決しました。線形アテンション方式を用いるトランスフォーマーブロック7個の後に、既存方式のソフトマックスアテンションブロック1個を続けて配置する構成を繰り返し、全体で80層を積み上げました。各アテンションモジュールは64個のヘッドから構成され、ヘッド1つの次元は128です。ソフトマックスアテンション区間にはキャッシュメモリを節約するため、グループサイズ8のグループクエリアテンションを組み込みました。位置情報を計算するロータリー位置埋め込みの基準周波数は10000に固定しました。

こうした構造のおかげで、MiniMax-Text-01は学習段階から最大100万トークンに達する入力を扱えるように作られ、実際のサービスでは速度が目立って遅くなることなく最大400万トークンまで文脈

を拡張して使うことができました。これはウェブ文書の規模をはるかに超える量の文章を一度に丸ごと読んで答えることが可能になったということを意味します。

モデルを学習させるデータを選ぶ過程でもMiniMaxは独自技術を動員しました。前世代に作っておいた、有効パラメータ50億個、総パラメータ600億個の小型MoEモデルをデータ採点者として立て、収集したウェブ文書や書籍、論文、ソースコードを知識の深さ、実際の有用性、主題の分布という3つの基準で採点させました。こうして選別した高得点データに重みを乗せてサンプリングした後、事前学習に投入しました。

視覚言語モデルであるMiniMax-VL-01は、Text-01の重みを土台にして、5,120億個規模の画像・言語トークンを追加で学習させて作りました。その結果、MMMU評価で68.5パーセント、MMMU-Proで52.7パーセント、チャート理解を測るChartQAで91.7パーセント、文書理解を測るDocVQAで96.4パーセントを記録し、米国のClaude 3.5 SonnetやGemini 1.5 Proと肩を並べる性能を示しました。

2025年6月16日には推論特化モデルであるMiniMax-M1が発表されました。OpenAIのo1やDeepSeekのR1のように答えを出す前に自

ら思考する過程を経るモデルで、Text-01の重みを基盤に数学推論とコーディング論理に重点を置いた7兆5,000億トークン規模のコーパスを追加で学習させた後、教師ありファインチューニングと大規模強化学習を順に経て完成させました。

推論モデルは思考を長く続けるほど計算量が増える構造的な問題を抱えていますが、既存のソフトマックスアテンションだけを用いるモデルはこのコストが特に急激に跳ね上がります。M1はLightning Attentionをそのまま引き継いだおかげで、10万トークンに達する長い思考過程を生成する際にも、DeepSeek-R1に比べて4分の1水準である25パーセントの計算量だけで同等またはそれ以上の結果を出しました。ここに100万トークンの文脈処理能力まで加わり、OpenAIのo3やClaude 4 Opusと比べても劣らない推論能力を示したと同社は発表しました。

M1を学習させる強化学習の段階で、MiniMax研究チームはCISPOという独自アルゴリズムを新たに作り出しました。名称を訳すと「クリップ重要度サンプリング方策最適化」という意味になります。従来広く使われていたPPO系アルゴリズムには、学習の途中でトークンの更新幅が急激に大きくなり、モデルがあらぬ方向へ発散して

しまうという問題がありました。CISPOは、これを抑えるために重要度サンプリングの重み自体を切り詰める方式を導入したものです。

このハイブリッド・アテンション構造とCISPOアルゴリズムが組み合わさった結果、MiniMaxはH800グラフィックカード512枚を借りただけで、3週間でM1の強化学習全過程を終えました。このとき実際にかかった演算コストは53万4700ドル程度で、韓国ウォンにして7億ウォンに満たない金額でした。大規模モデルの学習にかかる通常のコストに照らせば、かなり低い水準といえます。

同社は、思考トークンの上限を異なる設定にした二つのバージョンを同時に発表しました。4万トークンまで思考できるMiniMax-M1-40kは8万トークン版へ至る途中段階の成果物であり、8万トークンまで思考できるMiniMax-M1-80kが最終完成版に当たります。数学オリンピック形式の問題を競うAIME 2024で86.0パーセント、AIME 2025で76.9パーセント、ソフトウェアエンジニアリングの実務能力を測るSWE-bench Verifiedで56.0パーセントを記録しました。

ここに至る道のりは平坦ではありませんでした。2023年下半期、同業他社の大半がいまだ安定した密モデルの改良にとどまっていた

時期に、閆俊傑は会社の研究開発資源と演算資源の大部分を、当時は傍流とみなされていた混合エキスパート(MoE)構造に注ぎ込む決断を下しました。MoE方式ならより低い学習コストと推論コストで、密モデルに匹敵するか、それ以上の性能を発揮できるというのが彼の判断でした。

しかし実際の開発過程では、学習が三度から四度にわたって崩壊する事態に見舞われました。損失関数が発散し、事前学習した重みを使い物にならなくなるたびに、二か月近い研究開発期間がそのまま失われ、崩壊が起きるたびに最大1500万ドルに達するスーパーコンピューター使用料も同時に消えていきました。

投資家の間で懐疑的な声上がるなかでも、閆俊傑は退きませんでした。投資家の黄明明(ホアン・ミンミン)は後年、彼を評して「見た目は穏やかだが、内面は徹底している人物」と振り返り、資源が潤沢でなかった当時、閆俊傑が保有する演算資源の8割ほどをMoE開発一本に注ぎ込んでいたと語っています。こうした決断を下せる創業者だけが最後まで生き残れるというのが、黄明明の評価でした。

この戦いの末、2024年1月、MiniMaxは中国で初めて大規模MoEモデルabab6を正常に完成させ、商用化にこぎ着けました。閆俊傑は後日、M1が完成した日に自身のチャンネルで「大きな山は決して越えられない壁ではないと、生まれて初めて骨身にしみて感じた」との言葉を残していますが、これはこの時期の失敗と再挑戦とともに指した言葉でした。

2024年4月には、こうして作り上げたMoEモデルを発表し、複数のベンチマークで世界トップクラスの成績を収めました。この流れが続き、2025年6月にはMiniMax-M1のオープンソース公開へとつながりました。両者の間に横たわる一年余りの時間は、失敗を乗り越えて構造を磨き上げてきた開発過程そのものだったといえます。

MiniMaxは、こうして作り上げたモデルの重みをApache 2.0ライセンスの下でGitHub上に公開する方式を選びました。MiniMax-M1の事前学習済みの重みと、それをローカル環境で動かすための設定ファイル、モデリングコード一式を無償で配布したのです。閆俊傑は、モデルの価値は隠した瞬間から古びていくものであり、どうせ一年もすれば時代遅れになる重みなら、公開してより多くの人に使ってもらおうほうがよいという考えを明らかにしています。

こうして作られたモデルは、すぐさまコンシューマー向け製品へとつながりました。利用者が望む仮想の人格を作り出し、音声とテキストで会話できるTalkie(トーキー)と星野(シンイエ)は、世界で2億人を超える利用者を集め、高解像度の映像を生成する海螺AI(ハイルオAI)は、M1の推論能力とリアルタイム検索機能を組み合わせてサービスされています。

APIの料金設定にも、ライトニング・アテンションによるコスト削減効果がそのまま表れています。入力長が20万トークン以下の区間では、100万トークンあたり入力0.4ドル、出力2.2ドルという料金設定で、これはDeepSeek-R1よりも低い水準です。入力長が20万から100万トークンに及ぶ区間では入力1.3ドル、出力2.2ドルとなりますが、この区間はDeepSeekがアーキテクチャ上の限界からそもそも対応できていない領域であり、MiniMaxだけの市場として残っています。

閆俊傑はこの一連の開発過程を振り返り、創業以来守ってきた三つの原則を挙げました。技術主導の方向性を貫いたこと、オールモダリティの開発路線を選んだこと、創業初日から世界市場を見据えていたことです。彼は、モバイルインターネット時代のやり方、す

なわちトラフィックを買い集め、マーケティングと人員を投入して押し切るという手法を、そのままAIに持ち込もうとしても通用しないと気づいたと語っています。

彼はまた、スケーリング則に基づく効率向上への投資を続けています。モデルの性能が大きく飛躍し、推論コストが10分の1程度まで下がれば、事業機会はおのずとついてくると説明しました。彼の言葉を借りれば、長期的な技術価値を積み上げることこそが、短期的な収益性の課題を解く道だということです。

技術的なリスクよりも何を懸念しているかという問いに対し、閆俊傑は実証済みの道ばかりを追いかける「追随者心理」を挙げました。独自の領域を切り開かず他者の後ろを追うだけでは、中国のAI産業は常に二流にとどまらざるを得ないというのが彼の考えです。こうした考えが、ライトニング・アテンションやCISPOのような独自技術の開発を推し進めた土台になったといえます。

現在地と中国AIにおける意味

原稿の改訂を終えました。番号・表・検証メモ・小見出しはすべて散文の形に書き直し、6502字・26段落で、禁止表現チェックも通過しました。

2026年1月9日、閔俊傑が率いるMiniMax(上海稀宇科技)は香港証券取引所に正式上場しました。創業から上場までにかかった期間は4年余りで、これは世界中の生成AI大型モデル企業のなかでも屈指の短さでした。上場初日の株価は公募価格の2倍を超えて上昇し、会社の時価総額は1000億香港ドルを突破しました。当時36歳だった閔俊傑個人の資産は、ブルームバーグ・ビリオネア指数によれば35億ドル前後まで膨らみました。

10日後の1月19日、閔俊傑は北京で開かれた専門家・企業家・科学教育文化界代表者による座談会に出席し、発言しました。この座談会は、DeepSeek創業者の梁文鋒(リャン・ウェンフォン)に続き、大型モデル企業のトップとしては二人目の招待を受けての発言の場でした。上場という資本市場からの評価と、国家主催の座談会出席という政策面での評価が、同じ月のうちに重なって起きたことになります。閔俊傑とMiniMaxが、中国AI産業を代表する一角として確固たる地位を築いたことを示す場面です。

MiniMaxが歩んできた道は、最初から平坦だったわけではありません。2021年12月9日、チームを組んだ初日は、五つの会社を同時に立ち上げられるほどの自信があったわけでもなく、チームはただ

五か年計画を練るのに九時間を費やしただけでした。翌年初め、閆俊傑はセンスタイム(SenseTime)を離れ、MiniMaxを正式に設立しました。彼はかつて、創業を選んだ理由をこう説明したことがあります。同じくらい優秀な二人の卒業生のうち、一人は大企業に、もう一人はスタートアップに進んだ場合、2〜3年後には創業を選んだほうの成長速度がより速いということです。

2023年下半期、同業各社の多くが安定した密モデルの改良に追われていたころ、閆俊傑は会社の研究開発資源の大部分を、混合エキスパートシステム、すなわちMoE路線へ投入することを決めました。MoEは、より低い学習・推論コストで密モデルに匹敵するか、それ以上の効果を発揮できる構造でしたが、そこへ至る道は険しいものでした。MiniMaxは三、四度にわたる深刻な事前学習の失敗を経験し、モデルが崩壊するたびに二か月の開発期間と最大1500万ドルの演算コストが水泡に帰しました。

投資家の黄明明は後年、閆俊傑を評して「見た目は穏やかだが、内面は徹底している人物」と振り返っています。資源が限られていた状況のなかで、彼は当時手にしていた演算資源の8割をMoE開発に注ぎ込んでいたということです。そうした決断を下せる創業者だけ

が最後までやり抜けるという評価でした。2024年4月、MiniMaxはこの路線の成果である新たなMoEモデルを発表し、複数のベンチマークで世界最高水準の性能を記録しました。

2025年1月14日に発表されたMiniMax-01シリーズのモデル、すなわちMiniMax-Text-01とMiniMax-VL-01は、こうした積み重ねの上に生まれた成果です。このモデルは、標準的なトランスフォーマーのアテンション方式から離れ、線形アテンション機構を4000億パラメータ級のモデルに初めて適用しました。総パラメータ数は4560億個に達しましたが、トークン一つを処理する際に実際に働くパラメータは459億個程度に抑えられています。この設計のおかげで、演算コストを大きく増やすことなく400万トークンに及ぶ長い文脈を処理できる余地を確保しました。

2025年6月、MiniMaxは推論に特化したモデルMiniMax-M1をオープンソースとして公開しました。このモデルは、テスト時点での演算最適化と大規模強化学習を組み合わせたハイブリッド・アテンション構造を採用しています。7兆5000億トークン規模の推論重視コーパスによる継続事前学習を経たのち、独自開発の強化学習アルゴリズムCISPOを適用した成果でした。10万トークンに及ぶ長い推論

過程を生成する際にかかる演算量は、競合モデルであるDeepSeek-R1の4分の1程度とされています。

このモデルの学習に使われた資源も注目を集めました。H800 GPU 512枚を三週間稼働させて大規模強化学習を完了し、これにかかったインフラコストは53万4700ドル程度と見積もられています。思考の深さを調整できるよう、4万トークン版と8万トークン版の二種類に分けて発表しました。ただし、これらの数値はMiniMax側の発表に基づくものであり、独立した第三者のベンチマークで改めて確認された値ではない点は明記しておく必要があります。

MiniMaxはテキストにとどまらず、音声、映像、音楽までを網羅するオールモダリティ製品群を広げ続けてきました。映像生成モデルの海螺AI(Hailuo AI、ハイルオAI)や音声モデルのSpeech、音楽モデルなどを相次いで発表し、これらの製品は企業向けのオーブンプラットフォームと並行して、一般消費者向けアプリケーションとしてもサービスされています。閻俊傑は、真の汎用人工知能であるならモダリティを選ばないはずだと繰り返し語っており、この信念が三つのモダリティを同時に押さえる戦略の土台になっています。

閆俊傑が繰り返し強調する原則は、巨大モデル時代における製品とはモデルそのものであるということです。モデルの論理力と推論能力を引き上げられないままマーケティングで事業を推し進めても長続きしないと彼は語ります。そのため、会社のあらゆる反転はモデル性能の質的飛躍から生まれなければならないという原則を打ち立てています。この原則は、創業初期のモバイルインターネット時代のやり方、すなわちトラフィックを買ってマーケティングで押し切るやり方をそのままAIに持ち込もうとして試行錯誤を重ねた末に固まったものです。

競争に対する姿勢にも彼の性格が表れています。2024年に入り大型モデル市場の競争が一段と激化し、大手企業の挟撃によってスタートアップの居場所が狭まると、閆俊傑は、電気自動車やスマートフォン、モバイルインターネットのように中国で大きく成長した産業はいずれも複数の企業が長期間にわたり激しく争った末に世界市場を主導する製品を生み出してきたと語りました。競争は避けられず、避けられないのであれば強くなれる方向を最大限伸ばすしかないというのが彼の判断でした。

彼が挙げる二つの核心課題は、技術をいかに高めるか、そしてユーザーとともにいかにより良いものを作り出すかです。どちらも長い時間をかけた蓄積が必要だと彼は説明します。こうした見方は、スタートアップが大手企業と対峙した際に資源の規模では勝てない以上、技術と製品という狭い道で勝負するしかないという現実認識から生まれたものと見られます。

国内市場の価格競争をめぐる見方も注目に値します。複数の大型モデル企業が価格を引き下げると、大型モデルは高いとばかり考えていた伝統企業が大挙して利用に踏み出しました。コストが下がったことでエラーが出て再度呼び出せばよいという心理が広がり、モデルの呼び出し量そのものが大きく増えたと閆俊傑は評価します。彼は、こうした競争のおかげで中国のモデルが非英語圏の言語においてGPTに匹敵する水準まで引き上げられたと見ています。

海外市場に進出する戦略は創業当初から定まっていた。閆俊傑は、規制や国内競争を避ける目的ではなく、汎用人工知能を人類全体が使うインフラと捉えるなら一地域にとどまることはできないという理由を挙げています。海外市場には大手企業のトラフィック支援や既存の流通網がないため技術と製品力だけで勝負しなければ

ならず、そうして得たユーザーの反応こそがマーケティング費用をかけて得たデータよりも価値があるというのが彼の説明です。

こうした戦略は数字にも表れています。MiniMaxの製品は世界で2億人を超えるユーザーを集め、同社の発表によると2025年第3四半期の累計売上高は前年同期比で170パーセント超増加し、そのうち海外売上高の比率は70パーセントを上回りました。これらの数値は同社の発表に基づくものであり、外部監査や第三者集計によって再確認された値ではない点も付け加えておきます。

収益化の構造について閔俊傑は大きく二つの軸で説明しています。一つは企業顧客と開発者を対象としたMiniMaxオープンプラットフォームであり、もう一つはコンシューマー向け製品に組み込まれた広告です。オープンプラットフォームには3万を超える企業顧客と開発者が集まっています。自社でモデルを構築するのが難しい企業がMiniMaxの音声・視覚の能力を借りて使う仕組みです。ただし彼は、現段階で優先すべきは収益化ではなく、技術が広く使われる水準に到達することだと線を引いています。

大型モデルスタートアップの生存可能性をめぐる懐疑的な見方に対しても、閔俊傑は淡々と答えます。投資家の朱嘯虎(Zhu Xiaohu)

が、大型モデルの開発コストがあまりに大きいため収益化だけでは自立が難しく、こうしたスタートアップの最善の結末は大手企業に買収されることだと述べたのに対し、閆俊傑は、製品を使う人がいない、あるいは収益を上げられないのはユーザーのせいではなく、たいてい自社の技術と製品が不足しているせいだと答えました。

彼はこの状況を業界が通過する一種の試験台と見ています。QQも2000年頃には収益化の方法が分からずさまざまな収益モデルを試みては失敗を重ね、最終的にモバイル付加サービスとゲームにたどり着いたという事例を挙げながら、試験に合格すれば生き残り、合格できなければ淘汰されるのが道理だと語りました。ユーザーやエコシステムのせいにせず、自ら足りない点を埋めようとする姿勢が彼の語り口に繰り返し表れています。

マルチモーダルへ拡張する理由についても閆俊傑は明確な論理を持っています。人々が日々接するコンテンツの大半は文章ではなく写真と動画であり、テキストはそのうちの洗練された一部にすぎないということです。幅広いユーザー層と高い利用頻度を同時に獲得するにはテキストのみを出力するモデルでは足りず、複数のモダリティを併せ持つ動的なコンテンツ出力が必要だというのが彼の核心的

な判断です。

動画生成がテキスト生成より難しい理由についても彼は具体的に指摘します。動画一本には数千万単位の入力と出力が伴うため処理そのものが厄介で、データ容量もテキストとは比較にならないほど大きくなります。5秒間の動画が数メガバイトに達するのに対し、同じ時間分のテキストは1キロバイトにも満たず、保存容量の差は数千倍に及びます。テキストを前提に組まれた既存インフラをデータ処理・整形・ラベリングの段階から作り直さなければならない理由はここにあります。

世界市場に向けたオープン戦略も彼の事業観を示しています。MiniMax-M1の重みはApache 2.0ライセンスの下でHugging FaceとGitHubに公開されました。閆俊傑は、AIの基盤モデルはどのみち1年もすれば大きく後れを取るものなので、抱え込んでいても得るものは少なく、むしろ世界中に自分たちの構造を追従させるほうがアーキテクチャのエコシステムで足場を広げるのに有利だと説明しました。

2026年を控えて彼が明かした懸念事項は、技術的な難関よりもむしろ心構えにあります。検証済みの道をひたすら追従する追従者心

理にとどまれば、中国のAI産業は永遠に二流にとどまらざるを得ないというのが彼の診断です。シリコンバレーが切り開いた経路を少しずつ手直しして世界2位の水準に安住しようとする姿勢こそ、どのような技術的ボトルネックよりも危険だと彼は語ります。

こうした発言はMiniMaxが歩んできた実際の選択と結びついています。他社が安定した高密度モデルを磨いている間に、線形アテンションという馴染みのない構造に会社の資源の大半を賭けた決定。三つのモダリティを同時に手がける困難な道を選んだ決定、国内競争ではなく海外市場で正面から検証を受けようとする決定は、いずれも同じ姿勢から生まれたものと見られます。安住せず自ら困難な道を選ぶ姿勢が、彼の経営に一貫して表れています。

MiniMax-M1がDeepSeek-R1に比べて4分の1程度の演算量しか使わないという数値や、訓練コストが53万ドルだという数値も、同社の発表に基づく値です。

上場と国家座談会への出席という二つの出来事が同じ月に重なったことで、MiniMaxは商業的成功と政策的な評価を同時に得た企業として浮かび上がりましたが、その成長の財務的な持続可能性は、今後発表される実績によって改めて検証されるべき課題として残っ

ています。

閆俊傑という人物を現時点でまとめるなら、数学を専攻し商湯科技(SenseTime)で実務経験を積んだのち三十歳前後で会社を設立し、わずか4年で香港証券取引所に上場させた創業者であり、線形アテンションという非主流の技術路線に会社の資源を賭け、その成果をオープンソースとして世界に送り出したエンジニアです。彼の次の一手は、MiniMaxが今の成長速度をどれだけ長く維持できるか、そして彼が警戒し続けてきた追随者心理を避け、どこまで独自の道を切り開いていけるかにかかっていると見られます。

第7章 王小川（ワン・シャオチュアン、Wang Xiaochuan） 、検索の戦場を離れ、医療AIへと舵を切ったBaichuan創業者

成長期と学業

王小川は中国四川省成都市で生まれた。彼が後に設立した会社、百川智能(Baichuan)という社名には、あらゆる川の流れが集まって大きな川になるという意味が込められているが、その最初の流れは成都の教室から始まった。数学の問題を解いていた少年は、それからほどなくして全国的な舞台に名を連ねることになる。

十四歳になった年、王小川は中国全国高校数学コンクールで一等賞を受賞した。通常は高校の上級学年の生徒たちが競い合う場で、はるかに若い年齢で頂点に立ったわけで、同世代よりも数歩先を歩んでいたと言えるだろう。

数学から始まった才能は、やがてコンピュータプログラミングへと広がっていった。十七歳になった彼は国際情報オリンピックに中国代表として出場し、金メダルを獲得した。数学コンクールが数字と論理を扱う場だったとすれば、情報オリンピックはその論理を実際に動くコードへと落とし込む力を競う場だった。王小川はこの二つの関門を十代のうちにすでに突破していたのである。

こうした実績を土台に、彼は中国屈指の名門校である清華大学のコンピュータサイエンス学科に進学した。

清華大学在学中の王小川の歩みで目を引くのは、彼が学部生の身でありながら、すでに自身の最初のベンチャー企業を立ち上げていたという事実である。大学の教室で学んだ理論を卒業後に先延ばしにせず、学生の身分のまま会社を興して市場に飛び込んでいたようだ。

このように理論と実践を同時に手にする姿勢は、その後の彼のキャリア全体を通じて繰り返し現れる。搜狗(Sogou. 搜狗)検索エンジンで二十七歳にして副総裁の座に就いたことも、後に百川智能を設立して自ら会社を率いたことも、すべて学生時代に学んだことをすぐさま実行に移す習慣の延長線上にあると見ることができる。

清華大学の学部課程を終えた後、王小川は同大学の大学院課程に進学した。ところがこの時期、彼が没頭していたテーマは、人工知能や検索アルゴリズムのような、後の彼のキャリアを象徴する分野ではなかった。彼は生命体の設計図を解読する問題に取り組んでいたのである。

2000年、王小川は大学院の卒業論文で遺伝子塩基配列解析の組み立てアルゴリズムを扱った。人の遺伝子を構成する長い塩基配列の断片をコンピュータで再び組み合わせる問題であり、生物学とコンピュータ工学が交わる地点に位置するテーマだった。検索エンジンもチャットボットもまだこの世になかった時代に、彼はすでに生命のコードをコンピュータの言語で解き明かす作業をしていたのである。

この論文作業の根底には、一つの問いが横たわっていた。王小川は生命を支配する数学的原理とは一体何なのか、自らに問い続けていた。この問いは大学院時代の一時的な関心事に終わらず、その後数十年にわたって彼が抱き続ける命題となった。

彼はこの問いに答えを見出そうとする過程で、西洋式の数学と物理学だけでは限界があると感じた。方程式とアルゴリズムで精巧に設計されたコンピュータプログラムとは異なり、人の体は数え切れないほど多くの変数が互いに絡み合った複雑な体系だからである。一つの定まった公式では説明しきれない対象を前に、彼は別の道を探し始めた。

その別の道が中医学だった。王小川は中医学を迷信や代替療法としてではなく、生命という複雑な体系を全体として捉える一つの哲学であり、もう一つの科学の形として受け止めた。西洋医学が体を部位ごとに分けて分析するとすれば、中医学は体全体の均衡と流れを見る形で近づいていく。王小川はこの二つのアプローチを対立させるのではなく、生命を理解するための異なる二つの枠組みとして受け入れている。

こうした関心は机上の勉強だけにとどまらなかった。王小川はその後、実際に中医学関連企業に投資し、この関心を具体的な行動へと移した。大学院時代に抱いた一つの問いが、その後の彼の投資判断にまで痕跡を残したようだ。

清華大学を後にした王小川は搜狗に加わり、検索と入力ソフト事業でキャリアを積んでいき、三十歳前後の若さですでに副総裁の座に就いた。彼はその後、搜狗の最高経営責任者を務めたが、2021年にその座を退いた。検索エンジン会社を率いていた時代の話は次の節で改めて扱うが、大学院時代に抱いた生命に関する問いが、この時期にも彼の心の片隅に残り続けていたという点は記しておく価値がある。

搜狗を離れた後、王小川は2023年4月に百川智能を設立した。社名を決めチームを編成していく過程は後の節で詳しく扱うが、その出発点には、二十三年ほど前の大学院時代に書き上げた遺伝子組み立てアルゴリズムの論文と、生命の数学的原理を知りたいという、あの古くからの問いが依然として横たわっていた。検索エンジン事業で名を上げた彼が、結局医療人工知能へと舵を切った話は、つまり新たな決心というよりも、学生時代から抱き続けてきた問いへと立ち返った旅路に近い。

彼の物語は成都の数学教室から始まり、清華大学のキャンパスを経て生まれたものであり、家族史を通じて伝えられる物語ではない。

一つは数学と情報学において若くして全国と世界の舞台の頂点を極めた才能の線である。もう一つは大学院時代から手放さなかった、生命に関する問いの線である。

この二つの線は互いに異なる方向を向いているように見えるが、実際には一人の人間の中でずっと重なり合ってきた。コンクールとオリンピックで示した論理的能力は、後に検索エンジンと人工知能モデルを作る技術的な土台となった。生命の数学的原理を問い続け

た問いは、結局彼が医療人工知能へと会社の方向を転じる糧となった。

もう一つ触れておくべきは時代的な位置づけである。王小川は1990年代生まれの若い人工知能起業家たちよりも上の世代に属する人物として紹介される。検索エンジンの時代を経て人工知能の時代へと移り変わる中でも現役で活動し続けているという点で、彼の学業期の物語は中国情報技術産業の一つの世代の流れとも重なっている。

結局、王小川という人物を理解するうえで、成都の数学教室と清華大学の大学院研究室は欠かせない二つの場所である。前者の場所で彼は数字とアルゴリズムを扱う力を世界水準にまで引き上げ、後者の場所では、その力を生命というはるかに大きく複雑な対象へと向け始めた。

彼が検索エンジン会社の経営者から、再び人工知能会社の創業者へと立場を移していく間も、この二つの場所で得たものは消えることなく、常に共にあった。次の節で扱う搜狗時代の検索と入力ソフトをめぐる競争、そしてその後の百川智能創業の物語は、結局のところこの学業期の延長線上にあるようだ。

研究者、あるいは創業者として定まっていく時期

中国四川省成都で生まれた王小川は、十四歳で全国高校数学コンクールで一位を獲得し、十七歳では国際情報オリンピックで金メダルを受賞した。幼い頃から数字と論理で頭角を現したこの経歴は、後に彼が検索エンジンと人工知能というデータ産業に飛び込む土台となった。

1996年、彼は清華大学コンピュータサイエンス学科に入学した。同じ時期、この学科には、後に智譜AIを設立するチャン・ポン、百川を共に作ることになるルー・リーユンのような、後に中国人工知能産業を率いることになる人々が共に学んでいた。チャン・ポンは学生時代の王小川を、学科内ですでに名の知れた人物として記憶している。

大学を卒業する前から、彼は中国初期のポータルサイトであるチャイナレン(ChinaRen)の設計に参加した。その実力に目をつけた搜狐(Sohu)創業者の張朝陽が彼を直接見出し、王小川は学位を終える前から商業技術の現場に足を踏み入れた。

Sohu(搜狐)に入社した後、検索とポータルのインフラを担当して成果を上げ、27歳でSogou(搜狗)の副総裁に就いた。大規模なトラ

フィックを扱い、検索アルゴリズムを運用しながら培った感覚は、後に言語モデルを扱う上でも受け継がれる基盤となった。

清華大学大学院時代に彼が書いた修士論文は、検索や人工知能ではなく、遺伝子塩基配列を組み立てるアルゴリズムだった。2000年に完成したこの論文で、彼は生命体という複雑な系を数学で解き明かす問題に没頭した。

彼はこの時期から、生命が働く原理を数式で説明できるのかという問いを長く抱いてきたと語る。西洋式の還元主義科学だけでは生命という対象を十分に説明できないと感じ、その限界を埋める別の視点を探し求めた。

その過程で彼は中医学に関心を持つようになった。彼にとって中医学は迷信ではなく、人体を一つの統合された系として見る哲学に近く、この関心は後に中医学関連企業へ個人資金を投じるころにまでつながった。

Sogouを率いていた時期、彼は検索という事業の本質が結局のところ人の言語を理解し、適切な情報へとつなげることにあるという点を繰り返し確認した。膨大なウェブデータを扱いながら、言語こ

そが知能の中心軸であるという考えを育てていった。

2021年、彼はSogouの最高経営責任者の座を退いた。すぐに新しい事業に飛び込むことはせず、しばらく学術的な探索の時間を過ごしたが、この時期に彼が長く抱いてきた生命と言語に関する問いが再び浮かび上がった。

2022年末、アメリカでChatGPTが登場すると、彼は大きく心を動かされた。ChatGPTが示した言語理解と推論の能力は、彼が長い間問い続けてきた問い、すなわち人間の言語知能を機械で実現できるのかという問いに、実際の答えを示した出来事だった。

彼はこの変化を、単なる基礎技術の発展としてではなく、自分がずっと前から描いていた生命モデルと医師人工知能という構想を現実に移す道具がついに現れた瞬間として受け止めた。この判断が彼を再び起業の場へと導いた。

2023年4月、彼はBaichuan(百川智能)を設立し、代表を務めた。共同創業者として、清華大学の同窓生でSogouの最高執行責任者を務めた茹立雲(ルー・リーユン)を迎えた。Sogou、百度、Huawei、Microsoft、Tencent、ByteDance出身の人材を集めて研究組織を

編成した。

Baichuanはシード資金5000万ドルで出発し、設立から半年後の2023年10月にAlibaba、Tencent、Xiaomiなどから総額3億ドル規模のシリーズA投資を受けた。累計調達額は50億元を超え、広告や価格競争に資金を浪費しなかったという評価を受けた。

同年、Baichuanの研究陣は2兆6000億トークンで学習させたBaichuan 2-7BとBaichuan 2-13Bをオープンソースとして公開した。これらのモデルはMMLU、GSM8K、C-Evalといった評価において、同規模の他のモデルに匹敵する性能を示した。

こうして汎用言語モデルを世に送り出した後も、王小川の関心は最初から医師人工知能を作ることにあったと彼は語る。ChatGPTの登場は、その目標へ向かう道のりを早めてくれたきっかけにすぎず、方向そのものを変えたわけではなかった。

2025年4月、会社が設立2周年を迎えようとしていた時点で、彼は大幅な調整を断行した。汎用モデルの研究組織を大きく縮小し、金融をはじめとする複数の産業分野の事業ラインを整理したうえで、会社の資源を医療大規模モデル一つに集中させた。

彼はこの決定が必ずしも正しかったと断定はしないと語る。ただ、上場という成果は目の前の華やかさをもたらしてはしても、会社がどんな価値を生み出しているのかを自ら答えられない状態のほうが、彼にとってはより耐え難い問題だったと明かした。

この調整の後、会社の人員は300人以下に圧縮され、彼に直接報告する人員も10人余りの水準まで減った。多くの社員と一部の投資家がこの方向転換に反対し、その時期に会社を去った人も少なくなかったと彼は振り返る。

医療分野へ方向を転じて以降、彼が強調した原則は、特定の医師の診療記録を学習させ、その人物の判断方式を真似るやり方とは距離を置くことだった。彼はこうしたアプローチを診療事例をコピーする模倣学習と呼び、質問する能力が欠けたモデルだと指摘する。

その代わりBaichuanは、医師たちが正解を一つひとつ教える方式ではなく、どのような質問の流れが良い診断につながるのかを判断させ、強化学習の報酬基準を打ち立てる方式を選んだ。この作業は医療スタッフを説得し、評価体系を新たに作り上げなければならない、費用よりも時間のかかる仕事だったと彼は語る。

こうして作られた医療大規模モデルM4と対話型サービスBaixiaoyi(百小意)は、2026年5月22日にそろって公開された。M4は患者の診療歴全体を記憶する能力とエージェント構造を備え、OpenAIが作成したHealthBench評価において、難易度の高い領域と専門家領域の両方で最も優れた成績を収めたと彼は明らかにした。

Baixiaoyiはアプリとウィーチャットの二方向で運用されており、アプリでは診療記録と処方とを照らし合わせる真剣な判断を扱い、ウィーチャットのボットでは数日後に症状が良くなったか尋ねたり、薬を飲む時間を知らせたりする日常的な伴走役を務める。北京の児童病院ではこのモデルの小児科版が多職種診療に実際に投入され、専門家の診断とかなり一致する結果を出した。河北省の複数の県級病院にも拡大されている。

清華大学時代にデータを扱う感覚を身につけ、大学院で生命という対象を数学で解き明かそうとした王小川の関心は、検索エンジンの経営者を経て、再び研究者の問いへと戻ってきた。彼は自分自身を、商業的成果よりも一つの問いに長く答え続ける人間として描き、その問いの先には人間の病を診る人工知能があると語る。

会社と研究組織を作った過程

王小川は2021年にSogouの最高経営責任者の座を退いた後、すぐに新しい事業を興すことはせず、学術研究と知識の探索に時間を費やしました。検索エンジンを長く率いて培った技術的感覚を人工知能へと移していくための準備だった状況ですが、その成果は2023年初めに具体的な形を取り始めました。

法人登記簿に残る記録を見ると、会社の正式名称である北京百川智能科技有限公司は、2023年3月24日に正式に設立登録されました。初期株主名簿には、王小川が全体持分の99パーセント、出資額495万元を担い、残り1パーセントに当たる5万元は長年の事業仲間である茹立雲(ルー・リーユン)が担ったと記されています。出資期限は2043年3月15日となっていますが、これは中国の法人登記でよく用いられる長期履行条件であり、そのまま会社の財務状態を意味する数字ではありません。

茹立雲は、Sogou時代に最高執行責任者として検索アルゴリズムとトラフィック運用を長く指揮していた人物です。王小川がBaichuanを設立するにあたって真っ先に呼んだのが、まさにこの仲間であり、そうして二人が共同で会社の基本構造を築き上げました。

報道と会社の資料がともに示す創業時点は2023年4月です。このとき会社が掲げた目標は明確でした。中国独自の基礎的な大規模言語モデルを学習させ、その上で人々が実際に使える応用サービスを作り出すというものでした。OpenAIがアメリカで行っていたことを、中国の環境で改めてやってみようという抱負だったとまとめることができます。

研究開発組織を埋めた人材の顔ぶれも注目に値します。搜狗出身者をはじめ、百度、ファーウェイ、マイクロソフト、バイトダンス、テンセントといった企業で働いていた人材数十人が百川智能に加わり、モデルの事前学習をともに進めました。検索と人工知能の分野ですでに実力が証明された人々を素早く集められた背景には、王小川と茹立雲の搜狗時代の人脈が大きく作用したものと見られます。

大学との協力も早い段階から定着していました。北京大学と清華大学の二校が百川智能の初期モデルを研究ツールとして採用し、関連研究を進めたほか、2023年5月には北京市経済信息化局が指定する汎用人工知能産業イノベーションパートナー計画の第1期リストにも名を連ねました。

同社と学界の関係は論文にもそのまま表れています。2023年9月19日に公開された論文「Baichuan 2: Open Large-scale Language Models」の著者リストには、王小川本人を含め数十人の名前が並んでいますが、その中でも紀嘉明、張博榕、範学海、劉ミケル、戴俊涛、孫瑞陽、楊耀東など複数名が北京大学所属と記されています。社内外の研究人材が一本の論文のもとに名を連ねた形です。

この論文が示した学習規模もかなりのものです。百川智能の研究陣は2兆6000億トークンにのぼる多言語データでモデルをゼロから学習させ、その成果であるBaichuan 2-7BとBaichuan 2-13Bの基本モデルおよび対話モデルをApache 2.0などのライセンスで外部に公開しました。多額の資本を投じて作ったモデルを即座に開発者コミュニティに開放する形で、会社の名を広めた形です。

資金調達のプロセスも速いものでした。創業時点の2023年4月には投資家を明かさない5000万ドル規模の初期資金が入り、半年後の2023年10月には3億ドル規模のシリーズA投資を受けました。この投資にはアリババ、テンセント・インベストメント、シャオミ・グループをはじめとする複数の投資機関が参加しました。

ある投資家は同社について、累計調達額が50億元を超えながらも資金をそれほど使い込んでいない、数少ない安定した大規模言語モデルのユニコーン企業の一つだと評しました。

王小川は2026年1月、新浪科技とのインタビューで、会社が2027年に新規株式公開(IPO)の手続きを開始する見込みだと述べました。ただし、これはあくまで創業者本人の見通しであり、確定した上場スケジュールではない点は留意しておく必要があります。

こうして順調に見えていた会社の方向性が大きく変わった契機は、医療分野への集中でした。王小川は、大規模言語モデル業界の多くが依然として汎用モデルと応用サービスに力を注いでいた2024年半ば、百川智能の力を医療分野に集中させると表明しました。この決定により、しばらくの間「変わった」「遅くなった」と評され、業界の関心の外に押しやられることもありました。

王小川本人は、この関心が突然のものではないと説明しています。彼は2000年、大学院の卒業論文のテーマとして遺伝子シーケンシングを組み立てるアルゴリズムを研究しました。その頃から、生命を支配する数学的原理とは何かに興味を持っていたと語ります。数学や物理学といった既存の科学体系だけでは生命現象をすべて説

明できないと感じ、漢方医学を学び、関連企業に投資したこともあります。彼は漢方医学を一つの哲学体系として理解していると述べています。

王小川の論理は、医療を特定の一産業分野として捉えていない点にあります。彼は、医療分野に飛び込むことがいわゆる特定市場への方向転換ではないと言います。どんな人工知能技術であれ、医療現場で使われる場面があるため、他の応用分野が流行に乗って消えていったとしても、医療はその流れに巻き込まれないだろうというのが彼の判断です。

この方向転換は組織にも痕跡を残しました。百川智能で共同創業パートナーとして活動していた鄧江(トウ・コウ)は、その後会社を離れて独立起業しました。百川智能が言語モデルと一般消費者向けサービスに集中する間、彼はマルチモーダル技術と企業向けサービスという別の道を選びました。二人が同じ会社からスタートし、異なる方向に分かれていった事例と言えます。

医療分野の成果はモデル公開へとつながりました。百川智能は医療能力を高めたオープンソースモデルBaichuan-M2を発表し、医療評価基準であるHealthBenchで60.1点を記録して、320億パラメー

タ規模ながらOpenAIが公開した1200億パラメータ級モデルgpt-oss-120Bのスコアを上回ったと発表しました。

北京児童病院との協力もこの頃に進められました。王小川は対談の中で、小児科から協力を始めた理由を説明し、2025年第1四半期以内に海淀区の住民一人ひとりにAI医師アシスタントを提供する計画に言及しました。

医療モデルの開発にかかるコストは一般的な汎用モデルよりもはるかに大きいと王小川は語ります。それでもそれだけの価値があるとする根拠として、彼は中国で年間の外来診療件数が84億回を超える規模であることを挙げました。医療産業が汎用人工知能に吸収されず独自の地位を保つだろうという彼の見通しも、データと診療現場の特殊性に基づいています。

病院代や診療費を誰が支払うのかという問題についても、彼は三つの経路に言及しました。政府、病院、商業保険がそれぞれ費用を分担して支払う構造が可能であるということで、これに加えて海外市場への進出も一つの機会と見ていと述べました。大規模言語モデル事業を始めてからは地方政府関係者と会う機会も増え、入札手続きを経験する場面も積み重なったと彼は語りました。

王小川は百川智能を設立する前に率いていた搜狗と、現在の百川智能を比較して語ることもありました。検索というすでに確立された市場で競争していた搜狗と、まだルールが定まっていない人工知能医療分野で新しい道を切り開かねばならない百川智能は、創業者の立場から見てまったく異なる戦いだったと彼は振り返りました。彼はまた、ともに働く若い人材がもっと必要だという願いも合わせて語りました。

王小川は検索エンジン経営者として培った組織運営の経験と、清華大学時代から続く人脈をもとに百川智能を素早く軌道に乗せ、創業初期の汎用モデル路線から医療という具体的な方向へ力を結集する決断を下しました。その決断の根底には、生命現象を数学で解き明かしたいという、彼が大学院時代から抱き続けてきた長年の関心があります。

中核モデルと開発プロセス

王小川が搜狗で培った能力は、自然言語を扱う感覚と大規模トラフィックを制御する工学でした。百川智能を設立した後、彼はこの二つをそのまま持ち込むのではなく、米国発の公開重みを取り入れてアプリケーションだけを乗せる道を避けました。事前学習の段階からデータを整え、構造を作り直す道を選び、その成果が2023年

下半期に公開したBaichuan 2シリーズでした。

Baichuan 2は7Bと13Bの二つのサイズで登場し、いずれもゼロから新たに学習したモデルでした。学習に用いたトークン数は2兆6000億で、先行するBaichuan 1の2倍を超える規模でした。これほどの量の多言語データを集めて精製する作業は、量よりも選別の方法が重要になりますが、百川智能は独自に構築した重複除去システムをこれに組み込みました。

このシステムは、局所性鋭敏型ハッシュ系の手法と、文をベクトルに変換する密なエンベディングを併用して、類似文書を見つけ出し重複データをふるい落とししました。数兆字規模の原文コーパスを数時間以内に整理できたのは、このクラスタリング基盤のおかげでした。そうしてこそ、次の段階である学習に清潔なデータをすぐに投入できたのです。

トークナイザーにも手が加えられました。語彙集のサイズをBaichuan 1の64,000語から125,696語へとほぼ倍増させましたが、これは圧縮効率と学習密度の間でより良いバランス点を探るための選択でした。数値を扱う際には整数と小数点を一桁ずつ個別に分割するよう設計し、モデルが計算を行う際に桁数を取り違えないようにし

ました。

プログラムコードに頻出する連続した空白にも個別に配慮しました。空白専用の特殊トークンを入れることで、コードのインデントが崩れないようにしました。中国語のような長い文脈を扱う必要がある状況を考慮し、トークン処理長の上限をあらかじめ設定しました。細かな選択ではありますが、実際の文書を扱う際の精度の差を生む部分です。

位置情報をエンコードする方式は、二つのモデルサイズによって異なる方法が採られました。7Bモデルには回転変換を用いるRoPEを組み込み、フラッシュアテンションとうまく噛み合うようにし、13Bモデルには文脈長を超える拡張に強いALiBiを組み込みました。モデルサイズが大きくなるほど、より長い文脈を扱う必要がある場面が増える点を踏まえた使い分けでした。

ALiBiを用いるとグラフィックメモリの負担が増える問題が伴います。百川智能はこれを軽減するためxFormersライブラリのメモリ効率的アテンションを併せて組み込み、トランスフォーマーブロックの入力側にレイヤー正規化を配置することで、大規模分散学習の途中で損失値が跳ね上がり学習が崩壊することを防ぎました。こう

した調整は派手ではありませんが、実際に大規模モデルを最後まで学習させるために必要な安全装置でした。

性能評価では、MMLU、CMMLU、C-Eval、GSM8K、MATH、HumanEvalといった国際学術基準の試験で、同規模のモデルと競い合い良好な成績を収めた。中でも医学問題を扱うMedQAや法学問題を扱うJEC-QAのように、複数の段階を経て答えを導く専門領域で、他国のモデルより優れた結果を示したと記録されている。

ライセンス方針も注目に値する。7Bと13Bのベースモデルおよび対話型モデルの重みを、学術目的では無条件で公開した。一般ユーザーや企業もメールで申請し承認を受ければ、商用でも無償で利用できるよう門戸を開いた。4ビットに量子化したバージョンも併せて配布し、大規模サーバーを持たない環境でもモデルを動かせるようにした。

テキストの次に百川智能(バイチュアン・インテリジェンス)が手をつけた領域は音声だった。Baichuan-Audioは、人の音声をそのまま受け取って理解し、再び音声で応答する end-to-end 構造を目指して作られたモデルである。音声をテキストに変換し、再びテキストを音声合成する従来の方式とは異なり、中間変換の過程でイン

トネーションや感情といった情報が失われる問題を減らそうとする試みだった。

このために作られたのがBaichuan-Audio-Tokenizerである。Whisperエンコーダーで音声から意味情報と音色情報を同時に抽出し、8層からなる残差ベクトル量子化を経てこれを離散的なトークンに変換する。この過程を毎秒12.5回のフレームに圧縮することで、リアルタイム通話のように遅延が重要な状況でも負担なく処理できるようにした。

モデル内部では、音声専用のヘッドを別途設けた。テキストトークンと音声トークンが混在して互いに干渉する問題を防ぐための設計であり、テキストと音声ペアになったデータを学習させることで、会話中のどの時点で発話の受け渡しをすべきかをモデル自らが習得できるようにした。

推論された音声トークンは直接音として出力されるのではなく、フローマッチング方式のデコーダーを経てメルスペクトログラムという中間形式に復元される。その後、ボコーダーを通過することで、実際に人が耳にする音声波形へと変わる。複数の段階を経る理由は、圧縮されたトークンだけでは自然な音色やイントネーションを

直接復元することが難しいためである。

音声データを大量に投入して学習させると、モデルが本来持っていた推論力や計算能力がかえって低下する副作用が生じる。百川智能はこれを防ぐため、テキスト資料と音声-テキストのペア資料を交互に組み合わせる2段階の事前学習方式を用いた。その結果、中国語と英語のいずれにおいても途切れのないリアルタイム音声対話性能を確保したとしている。

音声モデルの次に百川智能が力を注いだ領域は医療だった。王小川は2025年初めのインタビューで、多くの大規模モデル企業がまだ明確な方向性を定めていなかった2024年半ばの時点で、会社が医療に集中するという決定を明確に下したと明かした。その年の旧正月前まで約8か月間、大きなモデル更新がなかったのも、実はこの方向転換を準備する期間だったという説明である。

この決定の背景には、彼が大学院時代から抱いていた関心があった。2000年の修士論文で遺伝子配列を接合するアルゴリズムを扱った当時から、生命を数学で説明できるのかという問いを持ち続けてきたと彼は語る。数学と物理学だけでは生命をすべて説明できないという考えから中医学にも目を向け、これを一種の哲学として理

解していると述べた。この関心が同社の医療モデル開発の方向性へとつながっているというのが彼の説明である。

王小川は、医療を特定業種に限定された狭い領域とは見なさない
と強調する。彼の論理によれば、どのようなAI技術であれ医療現場
でその活用を見出せるため、医療は今日の汎用AIの潮流に埋も
れて消えていく分野ではないという。彼は実際、同業他社がモデル
APIの価格競争やマーケティングによるトラフィック獲得競争に乗
り出す中でも、これに加わらなかったと語る。

こうした動きをめぐり、一時期ソーシャルメディアでは百川智能
が出遅れているとの見方も出た。ただしある投資家は、百川智能に
ついて累積調達額が50億元を超えながらも資金を大きく消耗してい
ない、数少ない安定した大規模モデルのユニコーン企業だと評価し
たと伝えられる。

Mシリーズのうち世に出たのは推論モデルであるBaichuan-M1-Pr
eviewで、旧正月連休の直前に公開された。王小川はここでマルチ
モーダル機能を組み込んだ理由について、ユーザーとのコミュニケ
ーション手段を広げるためのものにすぎず、マルチモーダル自体が
新しい技術領域を切り開くわけではないと説明した。彼は、OpenA

lのo1モデルが言語こそ知能の中心軸であるという事実を改めて確認させてくれたと見ている。

競合モデルに素早く追いつく過程については、知識蒸留の手法を用いることは特に隠すべきことではないという立場を示した。ディープシークが技術的には半歩遅れて動きながらも予想を上回る成果を出したという評価も併せて示した。これは、大規模モデル開発が必ずしも最前線でゼロから作り上げる方式のみで成り立つわけではないという彼の認識を表している。

医療モデルを実際の現場に導入する作業も並行して進められた。北京儿童医院と提携し、小児科から着手したのがその一例である。王小川は2025年第1四半期中に、北京市海淀区の住民一人ひとりがAI医師秘書を持てるようにするという目標を掲げた。彼はこのサービスの費用を負担する方法として政府、病院、商業保険の三つを挙げ、海外進出も別の機会として言及した。

医療専用モデルの学習にかかる費用は一般的なモデルよりはるかに高いが、それに見合う価値があるというのが彼の判断である。中国国内で年間の外来診療件数が84億回を超える規模であることを根拠に、彼は医療産業がデータや現場状況の特殊性ゆえに汎用人工知

能に容易には吸収されないだろうと見通した。オープンソースとして公開する方針を続けているのも、医療現場で各機関が自ら微調整して使えるようにするためだと説明した。

こうした方向性のもとで生まれた成果が、オープンソースの医療モデルBaichuan-M2である。国際的な医療ベンチマークであるHealthBenchで60.1点を獲得した。これは、OpenAIが発表した1,200億パラメータ級のオープンソースモデルgpt-oss-120Bの医療関連スコアを、32Bというはるかに小さな規模で上回る結果だった。王小川は、会社が30億元規模の資金を確保し医療AIに集中投資していることも明らかにした。

百川智能という会社そのものは2023年3月24日に法人として設立され、王小川と元搜狗最高執行責任者(COO)の茹立雲が共同で立ち上げた。資金調達には2023年4月に5,000万ドルで始まり、同年10月にはアリババ、テンセント、シャオミなどが参加したAラウンドで3億ドルを調達した。王小川は、会社が2027年に新規株式公開を開始する見通しだと述べたことがある。こうした資金と時間軸は、彼が検索時代に培った技術を医療という新しい領域へと移し込むための元手となっている。

現在の立ち位置と中国AIにおける意味

中国の大規模言語モデル業界の中で王小川は、検索会社を去った創業者ではなく、医療人工知能という明確な方向を自ら選び取って歩んできた人物として位置づけられている。彼が2021年に搜狗CEOの座を退き、2023年4月に北京で立ち上げた百川智能(Baichuan Intelligent)は、汎用人工知能サービスを研究・提供する会社として出発した。法人登記上の設立日は2023年3月24日で、王小川が99%の持分を保有し法定代表者を務めている。残る1%は元搜狗COOの茹立雲が保有しており、二人が検索時代の同僚関係をそのまま新会社に持ち込んだことを示している。

創業初期から投資家の反応は早かった。2023年4月に5,000万ドルを調達したのを皮切りに、同年10月にはアリババ、テンセント、シャオミグループをはじめとする複数の投資機関が参加したAラウンドで3億ドルを調達した。ソーシャルメディアで会社に懐疑的な声が上がった際にも、ある投資家は百川智能について、累積調達額が50億元を超えながらも資金をむやみに使わない、数少ない安定した大規模モデルのユニコーン企業だと評価した。同社は2023年5月、北京市経済信息化局が指定した第1次「北京市汎用人工知能産業革新パートナー計画」の一員となった。北京大学と清華大学が百川の

大規模モデルを導入し研究に活用する提携関係も結んでいる。

王小川自身の言葉を借りれば、医療への関心は創業よりもはるかに古いものだった。彼は2000年、大学院の修士論文で遺伝子シーケンシングの組み立てアルゴリズムを扱った当時から、生命を動かす数学的原理とは何かに関心を抱いていたと語る。数学や物理学といった既存の科学体系だけでは生命をすべて説明できないと悟った後には中医学を学び、中医学関連企業に投資したこともある。彼は中医学を一つの哲学として理解していると述べており、この関心が後に百川の方向性を定める土台となった。

2016年にアルファ碁が世界に衝撃を与え、魏則西事件が中国社会に医療情報の信頼性という問題を投げかけた時期の経験も、彼の判断に影響を与えたとインタビューで振り返っている。百川を立ち上げた当初から医療を念頭に置いていたが、会社設立初期には汎用モデルと応用サービスを前面に押し出して語っていた。そのため、後に医療へと方向を固めた際、外部からは会社が変わった、あるいは動きが鈍くなったといった評価も出た。

転換の契機がはっきりと表れたのは2024年半ばだった。複数の先頭集団の大規模モデル企業がそれぞれ異なる道を選んでいたその時

期、百川は医療集中という明確な選択を下し、以後はモデルAPIの価格競争にも加わらず、トラフィックを買い集めるマーケティング競争にも参加しなかった。王小川の説明によれば、医療は狭い意味での垂直領域ではない。自然言語理解であれ推論であれ、いかなるAI技術にも医療の中で活かせる場があるため、医療は今日の汎用人工知能の潮流に埋もれて消える分野ではないという論理である。

こうした選択を裏付けるように、百川は約8か月間新しい大規模モデルのバージョンを発表しなかったが、旧正月連休を控えて推論モデルであるBaichuan-M1-Previewを発表した。王小川は、このモデルが内モンゴルの重篤な脳梗塞患者の症例を扱った結果、モデルが示した診断方向が実際の協和病院の合同診療結果とかなり一致していたと紹介した。彼は、オープンソースとして公開する方針は当初から決まっていた計画だったと述べ、医療業界が直接取り入れて微調整しやすくするための選択だと説明した。

技術面では、OpenAIのo1モデルを素早く追いつける過程で知識蒸留の手法を用いることをあえて隠さなかった。ディープシークが技術的には半歩遅れて追随しながらも予想を上回る成果を出した点も彼は認めた。実際、百川のモデルをオープンソースとして配布す

る手順は、Baichuan2のGitHubリポジトリを通じて公開されている。これは、モデルを配布し利用する方法を文書として整理した基礎資料である。

王小川は、ChatGPTの登場後に得た判断についてもはっきりと述べています。言語こそが知能の中軸であり、人工知能はいわゆる第四次産業革命のようなものではないという考えです。この言語中心の見方は、オープンAIのo1が言語こそ知能の主軸であることを証明したという彼の評価につながり、次の段階の転換は人工知能が自ら道具を作って使う方向になるだろうという展望にも結びついています。ヤン・ルカンが、言語だけでは人工知能が現実世界の法則を理解できないと主張したことについてどう思うかと尋ねられた際にも、彼は自らの言語中心の見方を曲げずに答えました。

実際に医療サービスを展開する過程で、百川(バイチュアン)は北京児童病院と協力し、小児科から始めたと明らかにしました。王小川は、今年第1四半期のうちに海淀区のすべての住民にAI医師秘書ができるだろうと述べました。彼は費用を賄う方式として政府、病院、商業保険という三つの支払い構造を挙げ、これに加えて海外進出も一つの機会と見ていると明らかにしました。

大規模モデル事業を始めてから、以前の検索エンジンCEO時代よりも多くの省長や市委書記に会うようになったという彼の言葉は、医療人工知能が地方政府の入札や政策決定の過程といかに密接に絡み合っているかをうかがわせます。彼は実際の入札過程を経験した話も合わせて語りました。医療専用の超大規模モデルを訓練するのにかかる費用は一般的なモデルよりはるかに大きいものの、それに見合う価値があるというのが彼の判断であり、中国で一年間に行われる外来診療が84億件を超える規模であることを、彼はその根拠として挙げました。

医療産業が汎用人工知能に吸収されて消えることはないだろうという彼の展望は、医療データと診療現場が持つ特殊性に根拠を置いています。M1にマルチモーダル機能を組み込んだ理由については、これは主にユーザーとの相互作用を滑らかにするための目的にすぎず、マルチモーダル自体が独自の技術路線を新たに切り開くものではないと一線を引きました。より遠い未来を問う質問には、身体を人工的に補完する義体化や、漫画『攻殻機動隊』が描き出した世界、そして人工知能が人類文明を引き継いでいく姿までも挙げながら答えました。

搜狗(ソウゴウ)で検索と入力法の競争を繰り広げた経験と、百川での二度目の創業を比較してほしいという質問にも彼は答えを出し、今の会社にはより多くの若い人材が必要だという点を強調しました。この一年間何を作り、健康で幸せだったかを問う質問には、常温超伝導体のニュースが与えてくれた驚きと、人工知能という巨大な潮流に参加できたことを幸運だと思うと答えました。彼は、99.99パーセントの人々が今起きている変化を過小評価していると述べ、大規模モデルが結局は世界をより平らにしてくれるだろうという期待を語りながら対談を締めくくりました。

医療へと方向を定めた後、百川が打ち出した成果の一つが音声モデルBaichuan-Audioです。2025年2月に公開されたこのモデルは、テキストを介さずに音声を直接やり取りするエンドツーエンド方式を採用しており、8層から成る残差ベクトル量子化オーディオトークナイザーと専用のオーディオヘッドを組み合わせ、人とリアルタイムで対話できるよう設計されています。中国語と英語を行き来するバイリンガル環境でも自然な意思疎通ができるようにすることが、この技術報告書が掲げた目標でした。

医療分野での成果としてはBaichuan-M2が際立ちます。このモデルは国際的に通用する医療ベンチマークであるHealthBenchで60.1点を獲得し、320億個規模のパラメータのみで、オープンAIがオープンソースとして公開した1200億個規模のgpt-oss-120Bが記録した点数を上回りました。大きなモデルを小さく絞り込みながらも性能を維持したこの結果は、百川が汎用競争から手を引き、医療に資源を集中させた選択が無駄ではなかったことを示す根拠として言及されています。

もちろん、この方向転換が社内で全員の同意を得て静かに進められたわけではありませんでした。百川のパートナーであった鄧江(ダン・ジャン)は、言語モデルと個人消費者向けサービスに集中する王小川の路線とは異なる道を選び、会社を去ってマルチモーダルと企業向けサービスを志向する自らのAI医療会社を新たに設立しました。メディアは王小川について、最も孤独な人工知能創業者の一人であり、結局はすべてを一人で背負っていると評しました。会社が競争で後れを取っており、彼が掲げるAI医師という目標を支えるには厳しいという批判的な見方も合わせて報じられました。

こうした評価の分かれる中でも、王小川は2027年の新規株式公開(IPO)を目標に据えていると明らかにしました。百川智能(バイチュアン・インテリジェンス)の創業者兼CEOとして、彼は会社がその年に上場手続きを開始すると見込んでいると述べました。これはまだ確定した日程ではなく、経営陣が示した見通しとして扱うべき内容です。資金規模についても、30億元を確保して医療人工知能に全力を注いでいるという報道がありましたが、この数値もメディア記事に登場した表現であるため、そのまま確定した財務数値と断定するのはではなく、慎重に扱う必要があります。

検索エンジンと入力法の競争で経験を積んだ技術経営者が、医療という、失敗の代償が人の命に直結する領域へと会社全体の方向を移した事例は、中国の大規模モデル産業の中でも珍しいものです。王小川の選択は、単一モデルの性能競争というよりも、人工知能技術をどこに使うのかという問いに自ら答えを出した事例に近いといえます。彼が対談で繰り返し述べたように、言語を扱う能力から出発し、物理現象を数学に置き換え、さらに生命現象を数学に置き換えていく過程が彼の描く人工知能の次の段階であり、百川の医療モデル群はその考えを実際の製品へと移そうとする試みとして読み取ることができます。

第8章 周靖人(Jingren Zhou)、Qwenのエコシステムを世界の開発者に開いたアリババクラウドの技術責任者

成長期と学業

周靖人という名前は、今日人工知能を語るとき、会社紹介資料や発表会の壇上よりも先に、彼が歩んできた役職の名で浮かび上がってきます。彼のLinkedInのプロフィールには、アリババでシニア・バイス・プレジデント兼チーフサイエンティストを務めていると記されています。この短い肩書きの中に、彼が歩んできた道のかなりの部分が凝縮されています。アリババのような大規模企業においてチーフサイエンティストという席は、誰もが座れる席ではありません。会社全体が人工知能研究の方向性をどう定め、どこに資源を投じるかを判断する役割を担う席です。

彼が公の場に頻繁に登場した時期は、アリババクラウドが通義千問(Tongyi Qianwen)を世に送り出していた頃です。当時の彼の肩書きは、アリババクラウド・インテリジェンス最高技術責任者(CTO)でした。チーフサイエンティストと最高技術責任者、この二つの肩書きを並べてみると、彼が担ってきた役割の幅がうかがえます。研究の方向を定めることと、その研究を実際のサービスに仕立て上げること。この二つを併せて担う席だったと見ることができます。

2023年4月11日、彼は北京で開かれたアリババクラウド・サミットの壇上に立ちました。その日彼が残した言葉は、以後アリババの人工知能事業全体を貫く一文になりました。あらゆるソフトウェアは、大規模言語モデルを導入してアップグレードする価値があるという言葉でした。この一文は単なる製品宣伝文句ではなく、彼が長らく抱いてきた判断が凝縮された表現だったと見ることができます。

発表会が終わったあと、彼は南方都市報の記者と向き合い、もう少し長い話を聞かせました。このインタビューで彼は、アリババの大規模モデル研究が始まった時点を正確に示しました。アリババは中国で早くから大規模モデル研究に着手した企業の一つであり、2019年からさまざまな事前学習型大規模モデルを開発してきたと明かしました。通義千問が世に出る4年前から、すでに研究が始まっていたということです。

この発言は、彼が歩んできた研究者としての時間を押し量らせてくれます。2023年に突然現れた成果ではなく、それよりはるかに以前から積み重ねてきた作業の結果だという話です。彼は通義千問について、自分たちが歩んできた技術の道筋の上にある製品であり

、ここ数年の努力を示す結果物だと語りました。自分が作ったものを一瞬の成果としてではなく、長年の蓄積の産物として説明する姿勢が、この言葉の中に込められています。

彼が明かした2019年という時点は、世界的に見ても意味のある時期です。大規模言語モデルという概念がまだ広く知られていなかった頃でした。その時点ですでに事前学習型の大規模モデルを開発していたという彼の言葉は、アリババクラウドの中で彼が率いた研究組織が、かなり早い段階から方向を定めていたことを示しています。

記者がアリババの大規模モデルとOpenAIのChatGPTを比較するよう求めたとき、彼は言葉を濁さず答えました。GPTがさまざまな面で先んじているのは、科学技術発展の過程における当然の段階だと述べました。先行する技術を認めることにためらいはありませんでした。今は互いに追いつき追い越されつの局面であり、足りない部分を絶えず埋めていかなければならないと付け加えました。自分が率いる研究を誇張しない姿勢が、この答えにそのまま表れています。

それでいて彼は、次の段階では自分たちのモデルが技術革新に貢献できることを願っていると語りました。後れを取っているという事実を認めることと、これから前進していくという意志を示すこと。この二つが一つの文の中に共に込められていました。研究者が自らの立ち位置を冷静に見極めながらも、方向を見失わない姿がこの発言から読み取れます。

同じインタビューで彼は、通義千問がChatGPTを狙って作ったモデルではないという点も明確にしました。アリババは大規模モデル研究において自分たちなりの道を歩んできたと述べ、マルチモーダル、視覚、自然言語を順に研究してきたと説明しました。通義千問はその複数の研究領域が会う一つの地点にすぎず、完成形ではないとも語りました。自分の研究を一つの到達点としてではなく、今も進行中の過程として見つめる視線です。

彼が説明した研究の幅は、言語だけにとどまりませんでした。マルチモーダル、視覚、自然言語をひとつおき挙げたくだりから、彼が長年にわたり複数の分野を横断しながら研究を進めてきたことがうかがえます。一つの技術だけを深く掘り下げたのではなく、異なる領域の研究を並行して進め、それらを一つにまとめる作業をして

きたと見ることができます。

記者がモデルの学習・運用コストを10分の1、あるいは100分の1に減らす方法を尋ねたとき、彼の答えはかなり具体的でした。彼は、数千枚のグラフィックカードを使っても、特定の時点で一部しか稼働せず残りが遊んでいるなら、それは資源の無駄遣いだと指摘しました。この説明は、彼が単にモデルを上手に作る人ではなく、そのモデルを動かすコンピューティングインフラ全体を深く理解している人物であることを示しています。

彼は分散コンピューティングの実行グラフをどう動的に最適化するか、ネットワークの輻輳をどう避けるかについても触れました。こうした細部にわたる技術は、言語モデルを扱う研究者よりも、大規模システムを設計するエンジニアにとってなじみの深い領域です。彼がアリババクラウドというクラウドインフラ企業で長く地位を築いてきた理由が、ここに表れています。

サービス段階でコストを下げる方法についても、彼は具体的に説明しました。モデルを圧縮したり間引いたりして知識を蒸留する過程が必要だと述べ、企業専用モデルを配備する際には、必要な性能は維持しつつ不要な機能を削ぎ落とさなければならないと指摘しま

した。この説明の中には、大きなモデルを作る能力だけでなく、それを実際に使えるよう仕上げる能力までも併せ持たなければならないという彼の判断が込められています。

彼は、こうした技術的な細部の作業こそがアリババクラウドの持つ競争力だと語りました。華やかなモデル性能の数値よりも、目立ちにくいインフラ最適化の作業を自分たちの強みとして挙げたようです。長年クラウドシステムを扱ってきた者だけが出せる答えだったと見ることができます。

企業専用モデルと汎用モデルの違いを説明する際にも、彼は実務的な言葉を用いました。通義千問は特定の業務に合わせて作ったモデルではなく、全般的な知識体系を一つにまとめた汎用モデルだと説明し、企業ごとが持つ固有の知識や経験をモデルに落とし込むことが企業専用モデルの目的だと述べました。この説明において、彼は研究者の言葉と経営者の言葉を併せ持って語っていました。

彼は、企業が自社データを安全に隔離された専用空間に入れ、事前処理なしに文書や画像、動画をそのまま登録できるようにすると明かしました。ワンクリックでモデルを作る機能を支援するとともに付け加えました。複雑な技術を使いやすい形に変える作業に、彼がど

れほど気を配っているかが、この発言に表れています。

彼はまた別の方法として、APIを通じてモデルの機能を開放することにも言及しました。多くの開発者がAPIやSDKを使ってモデルの利用範囲を広げ、自分たちなりの価値を生み出すことを期待していると語りました。自分が作ったモデルを自社の中だけに閉じ込めませんでした。外部の開発者たちに開いていこうとする姿勢が、このくだりにはっきりと表れています。

こうした開放の姿勢は、クラウドサミットでの発言にも受け継がれていました。彼は通義千問の機能を開放し、個々の企業が自社専用モデルを構築できるよう支援すると語りました。アリババクラウドを通じて誰もが自社モデルを開発することを歓迎すると付け加えました。スタートアップを含む多くの企業や機関が、より容易にイノベーションを実現できるよう支援するという言葉も残しました。

記者がアリババ社内の全製品に通義千問を導入する日程を尋ねたとき、彼は決まった日程はないと率直に答えました。モデルをどう上手に使いこなせるかは、アリババを含むすべての企業が同じスタートラインに立っているのと同じだと述べ、一朝一夕に成し遂げられることではないと指摘しました。性急に日程を約束するのではな

く、ありのままの状況を伝える姿勢が、この答えから読み取れます。

この頃彼が残した発言の数々を一つにまとめてみると、研究者としての彼の気質がおぼろげに浮かび上がってきます。先行する技術を認めることに出し惜しみをしませんでした。自分が作ったものを完成形ではなく過程として説明し、華やかな数値よりもインフラの細部の作業を強みとして挙げる姿勢です。こうした姿勢は、長年にわたり大規模システムを扱いながら失敗と改善を繰り返してきた人物によく見られるものです。

2019年から続く事前学習型大規模モデルの研究は、その後いくつかの世代のモデルへと引き継がれていきました。アリババクラウドは後に、大規模モデルの推論と配備をワンストップで行えるサービスを整えるに至りました。その研究の延長線上から生まれた技術報告書の一つが、2025年に公開されたQwen3技術報告書です。2019年という出発点とこの報告書との間に横たわる時間が、彼が明かした研究の始まりが決して短くなかったという事実を物語っています。

彼の学歴や成長の背景を具体的に確認できるくぐりには多くありません。その代わり、彼がどのような席に座り、どのような言葉を残したか、その言葉がどれほど具体的で慎重であったかを通じて、彼の人となりを推し量ることができます。チーフサイエンティストと最高技術責任者という二つの肩書き、2019年という出発点、分散コンピューティングの効率化についての詳細な説明。この三つが、彼が歩んできた道を理解する手がかりとなります。

彼が残した発言の数々は、アリババという巨大な組織の中で人工知能研究がどのように位置づけられてきたかを示す窓口でもあります。一人の人物の肩書きが変わり、彼が立つ舞台が移り変わっていく間にも、その背後には2019年から積み重ねられてきた研究の時間がありました。そしてその時間は、後にアリババクラウドが世界の開発者たちへ門戸を開くという、より大きな物語へとつながっていくことになります。

研究者、あるいは創業者として固まっていく時期

周靖人博士の技術的な土台は、コンピュータサイエンスの中でもとりわけ手間のかかる分野である大規模データベースと分散システムで培われた。米国コロンビア大学大学院でコンピュータサイエンスを専攻していた頃、彼はデータをどう分割し、どう複数のコンピ

ュータに分散して格納し、それをどうずれなく統合するかという課題に取り組んでいた。この時期に扱ったテーマは分散コンピューティング、大規模データ処理アルゴリズム、大容量データベース管理システム(DBMS)の構造であり、この三つはのちに彼がアリババクラウドで超大規模AIインフラを率いる際に、そのまま生きる下地となった。

一人のコンピュータ科学者がどれほど評価されているかは、学界が与える肩書きで測ることができる。周靖人は世界のコンピュータ学界で最高権威とされる二つの団体、ACMとIEEEの双方でフェロー(Fellow)の地位を得た。この二団体はそれぞれ米国計算機学会と米国電気電子学会であり、そこでのフェローは論文を数本発表しただけで得られる肩書きではなく、ある分野に長年はっきりとした足跡を残した者にのみ与えられる称号である。データベースと分散システムという、目立たないがあらゆるサービスの土台を支える領域でこうした評価を受けたという事実は、彼が華やかな応用よりも基盤技術そのものを掘り下げる研究者として育ってきたことを示している。

研究者は一人では育たない。周靖人の学問的な感覚は、分散システムとトランザクション処理の分野を長く牽引してきた複数の学者たちとの交流のなかで磨かれた。IEEEコンピュータ学会の理事を務めたデイビッド・ロメット博士。IBM研究所に三十八年以上在籍し、現代的な分散トランザクション理論の基礎を築いたC・モハン(C. Mohan)博士。分散ディープラーニングエンジンであるApache SINGAプロジェクトを設計したシンガポール国立大学のベンチン・オイ教授が、その中にいる。彼らはそれぞれ異なる角度から大規模システムに取り組んできた人物たちだった。周靖人は彼らと成果を分かち合いながら、数万個の演算コアが遅延なくかみ合って動く原理を身につけていった。

ここで確認しておくべき点がある。周靖人がアリババに加わる前にマイクロソフトなど他社でどのような役職に就いていたか。大学院の学位論文の正確なタイトルと提出年が何であるか。この部分は未確認のまま残しておき、確認された事実を中心に彼の経歴をたどるのが正確である。

データベース研究者がクラウド企業の最高技術責任者へと転じる例は多くないが、周靖人の場合、その移行は自然なものだった。デ

一タを分割し、分散して格納し、再び統合する感覚は、そのまま超大規模AIモデルを学習させる感覚へとつながっていったからである。アリババクラウドの最高技術責任者として、彼は通義千問、すなわちQwenシリーズを率いる立場に立った。データベース最適化で培った目で、AIコンピューティングパイプライン全体を見渡すようになったのである。

Qwen3モデルを学習させる過程で彼が直面した課題は、規模そのものだった。このモデルは三十六兆個に及ぶトークンを事前学習データとして用いた。これだけの量を安定的に処理するには、データをどの順序で流し、どのハードウェアにどう分散して格納するかを細かく調整しなければならない。大学院時代に大规模データ処理アルゴリズムを扱った経験が、ここで再び生きることになる。

Qwen3シリーズが打ち出した数々の工夫の中で目を引くのが、「思考予算」と呼ばれる制御方式である。これはモデルが答えを出す前にどれだけ長く考えるか、すなわちデコード過程でどれだけの演算資源を使うかを、ユーザーが直接調整できるようにする仕組みだ。簡単な質問には資源を節約し、複雑な問題にはより長く考える時間を割り当てることで、同じハードウェアでより多くのリクエスト

を処理できるようにした方式である。

文脈をどれだけ長く記憶できるかも、この時期のQwenシリーズにとって重要な課題だった。デュアルチャンクアテンション(Dual Chunk Attention、DCA)とYARNという二つの技術によってモデルが扱える文脈の長さを伸ばしたが、これは長い文書や長い対話をまるごと読んで答えなければならない実際の利用環境で、モデルの限界を押し広げる作業だった。ハードウェアはそのまま扱える情報量を増やさなければならないこの課題もまた、分散システムで資源を無駄なく使う感覚と重なり合っている。

Qwenシリーズの前世代にあたるQwen2.5の技術報告書を見ると、このシリーズのモデルが最初から多様なサイズで設計されていたことが確認できる。小型モデルは軽量の機器でも動くように、大型モデルは複雑な推論まで担えるように分けて作られた構造は、一つの巨大なモデルだけを追求するのではなく、用途に応じた複数のサイズを同時に運用する戦略を示している。

同じシリーズで視覚と言語を併せて扱うQwen2.5-VLの技術報告書も、この時期の方向性をよく示している。画像や文書を読み取る視覚言語モデルを別系統として育てたことは、周靖人が率いるチー

ムがテキストにとどまらず、画面や画像、文書をまるごと理解するAIへと領域を広げていったことを意味する。

こうした拡張はその後も続いた。Qwenシリーズの最新世代としては、画面を理解しアプリケーションを操作するマルチモーダルエージェントモデルであるQwen3.7-Plus、プログラミングや事務自動化、長期タスクの遂行までを担うQwen3.7-Maxが登場した。これらは画面を見て、考え、文章を書き、実行し、検証するという一連の流れを一つのモデルの中で処理するように設計されたエージェント型モデルである。

テキストと画像の枠を超える拡張も顕著である。画像生成を専門に担うQwen-Image系列は、二百億個のパラメータ規模で作られたQwenシリーズ初の画像生成モデルだった。動画生成を担うWan系列は、自然言語の指示だけでオンライン上で動画の画面やストーリーを書き換えられる水準にまで進んだ。画像生成と編集を一つに統合したQwen-Image-2.0-Proも、同じ流れの上にある。

音声の領域も避けては通らなかった。十七種類の表現力豊かな声を備えたQwen3-TTS-Flashは、低遅延かつ高い安定性で音声を合成する。Qwen3.5-LiveTranslateは六十の言語を行き来しながら二

- ・五秒という短い遅延で同時通訳を行い、元の話者の声までリアルタイムで再現する。テキスト、画像、音声、動画を一つのモデルの中で理解しやり取りする全モーダル(Omni)モデルであるQwen3.5-Omni-Flashも、この系列の最近の成果である。

こうして複数の感覚へと広がっていったQwen系列の土台には、常にインフラの課題が横たわっている。アリババクラウドはQwenを実際に企業が使えるようにするため百煉(バイリエン)というプラットフォームを運営しており、モデル開発ツールと呼び出しサービスを一括で提供し、オープンソースフレームワークに対応したエージェント開発ツールまで備えている。研究室で磨き上げたモデルを企業の現場で実際に使えるサービスへと移すこの段階が、周靖人が最高技術責任者として担う役割のかなりの部分を占めている。

百煉を通じてQwenを導入した企業の事例は、このインフラが実際に機能していることを示している。大学のシステムがあちこちに分散しており、AI導入のハードルが高いという課題を抱えていた超星智創グループは、百煉と手を組みAI能力センターを設立した。ベスパンは自社開発のエージェントプラットフォーム「ChatOmni」で企業現場の課題に対応した。

乳幼児向けEコマース分野の唯亿科技(ウェイイー・テクノロジー)は、Qwen大規模モデルをもとに顧客対応を支援し、意図の精密な識別に成功したと明らかにした。ハローは、Qwenの複数のモデルを組み合わせて作ったエージェントによって、取引ロボットの全工程を新たに構築した。自動車メーカーの零跑(リープモーター)はアリババクラウドと手を組み、C10車種の車内に音声大規模モデルを搭載し、質疑応答と描画機能を実現した。

採用分野では、アリババクラウドが獵聘(リエピン)と提携し、Qwenをもとにプロンプトを磨き上げ、求人情報をインテリジェントに生成することで、人材と求人を結びつける速度を高めた。これらの事例は個別に見ればそれぞれ異なる産業に属しているが、いずれも同じ土台の上で動いている。周靖人が大学院時代から培ってきた大規模システム運用の感覚が、Qwenという形でさまざまな産業現場に浸透した結果である。

Qwen3が世に出た時期をめぐるインタビューで、周靖人はこのシリーズが目指す方向性を自ら語っている。彼が強調したのは、一つの巨大なモデルですべての問題を解決するというアプローチではなく、用途に応じてサイズと方式を変えた複数のモデルを併用する戦

略だった。これは、彼がデータベース研究者だった頃から持ち続けてきた、資源を無駄なく配分する感覚と変わらない。

Qwenシリーズがオープンソースという形で公開されている点も注目に値する。アリババクラウドはQwenモデルの多くを外部の開発者が利用し改変できるよう公開してきており、これは世界中の開発者がQwenをもとに自らのサービスを作れるようにする通路となった。研究者として周靖人が積み重ねてきた学界との交流のあり方、すなわち成果を分かち合い互いの研究の上に積み上げていく習慣が、会社としてのオープンソース戦略へとつながっているようである。

認証や規範を整える作業もこの時期の役割だった。Qwenシリーズのモデルは国内で初めて大規模モデル事前学習モデルテストを通過し、国家標準の要求を満たした。国家インターネット情報弁公室の大規模モデル登録手続きも国内で初めて完了したと伝えられている。こうした手続きは技術そのものと同じくらい、規模の拡大したAIサービスが社会の中に根を下ろすために経なければならない通過儀礼に近い。

AIマネジメント体制についても、国際認証機関のIQNetから世界初の認証書を受けたという記録が残っている。

振り返ってみると、周靖人の経歴で注目すべき点は、彼がデータベースと分散システムという目立たない基盤技術から出発し、その土台の上にAIという応用を積み上げたという順序である。華やかなモデルを一つ作り、あとからインフラを整えたのではなく、数万個の演算コアを扱う感覚を身につけたうえで、その上にQwenというモデルシリーズを築き上げたのである。

この順序は、Qwenシリーズが他のAIモデルと一線を画す点でもある。テキストモデルから始まり視覚言語モデルへ、さらに画像・動画・音声を含む全モーダルモデルへと広がっていく過程で、段階ごとにハードウェア資源をどう配分するかという課題が常に伴っていた。これは研究者として培われた習慣が、会社の技術戦略へと引き継がれた事例と見ることができる。

アリババクラウド最高技術責任者という役職は、結局のところ研究者として積み重ねてきた感覚を、会社規模のインフラへと移す役職だった。コロンビア大学院で培った分散システムの知識、学界の同僚たちと分かち合った大規模コンピューティングの原理に関する

洞察が、Qwenという名のもとに一つに結実した。この流れは今も、Qwenシリーズが新しいモデルを発表するたびに続いている。

会社と研究組織を作り上げた過程

周靖人はアリババクラウドの最高技術責任者(CTO)に就くと、すぐに目に見える製品づくりよりも人材と組織を整えることに取りかかった。分散システムと大規模データベースを扱いながら培ってきた感覚は、数万枚のグラフィックカードを一つにまとめて動かし、36兆個にのぼるトークンを遅延なく学習に投入するインフラを設計するうえでそのまま生かされた。通義千問、すなわちQwenという名のチームは、このインフラの上で生まれた。

Qwenチームは最初から完成した青写真を持って出発したわけではなかった。周靖人と現場の開発者たちは会議と実験を重ねるうちに、モデルをまるごと公開して世界中の開発者に使わせるほうが、かえってエコシステム全体を先に押さえる道になるという結論に至った。計画を立てて実行したというより、ぶつかりながら道を見つけたのちに、その方向を正式な戦略として定めた形だった。

こうして固まったオープンソース路線は、中国の中でも最も徹底した部類に入る。Qwenチームは数十億パラメータ規模の密なモデ

ルから、複数の専門家(エキスパート)を使い分ける混合構造のモデルまで、重み全体を惜しみなくコミュニティに公開した。モデル一つを誇示するよりも、使えるツール一式をまるごと引き渡す道を選んだのである。

スピードもまたこのチームの武器だった。Qwenチームは一、二か月のうちに対抗する改良版を出すやり方で応じてきた。中国の複数の大型モデル研究所の中でも、こうしたスピード勝負を続けて維持してきたのはQwenがほぼ唯一だと言われている。

この組織を実際に動かしてきた人々のうち、中核を担うのが林俊暘だ。彼はQwen、Qwen2、Qwen2.5、Qwen3へと続くシリーズ全体のアーキテクチャ設計を任されてきた技術リーダーであり、周靖人が描いた大きな絵図のもとで、実際のモデルの基本構造を組み上げる役割を果たしてきた。二人の関係は、インフラを設計する人とその上でモデルを組み立てる人が分かれていながらも密接に噛み合っている構図を示している。

Qwen3の技術報告書には、このチームに参加した名前が長々と並んでいる。An Yang、An Fengli、Bao Songyang、Bei Chenzhang、Bin Yuanhui、Bo Zheng、Bowen Yu、Chang

Gao、Chen Genhuang、Chen Xulu、Chu Jiazheng、Dai Hengliu、Fan Zhou、Fei Huang、Feng Hu、Hao Ge、Hao Ranwei、Huan Lin、Jia Rongtang、Jian Yang、Jianhong Tu、Jianwei Zhang、Jianxin Yang、Jiaxi Yang、Jing Zhu、Jingren Zhou、Junyang Lin、Kai Dang、Ke Qinbao、Ke Xinyang、Le Yu、Liang Haodeng、Mei Li、Ming Fengxue、Mingjie Li、Fei Zhang、Peng Wang、Qin Zhu、Rui Men、Rui Zhaogao、Shi Xuanliu、Shuang Luo、Tianhao Li、Tianyi Tang、Wenbiao Yin、Xingzhang Ren、Xinyu Wang、Xinyu Zhang、Xuancheng Ren、Yang Fan、Yang Shu、Yichang Zhang、Ying Ezhang、Yu Wan、Yuchong Liu、Zekun Wang、Zeyu Cui、Zhenru Zhang、Zhifeng Zhou、Zihan Quにいたる、この長い名簿こそが、一つのモデルを生み出すのにどれほど多くの手が必要かをそのまま物語っている。

これらの名前はテキストモデルにとどまらなかった。同じ人材が、画像と動画を併せて扱うQwen2.5-VL、音声までリアルタイムでやり取りするQwen2.5-Omniの開発にも受け継がれた。一つのチームが、テキスト、視覚、音声という異なる感覚の知能を順に組み込んでいった形である。

組織を立ち上げる作業と並行して、周靖人はアリババがすでに握っていた資本とクラウドの力を活用する戦略も併せて練った。アリババは中国国内で独自モデルを開発する複数のスタートアップに投資し、その結果、半数を超えるスタートアップがアリババクラウド上でサービスを動かすことになった。スタートアップに投じられた資金のかなりの部分が、やがてアリババクラウドのサーバー賃料とグラフィックカード使用料として戻ってくる構造が定着したのである。

こうして生まれた循環は、アリババクラウドを中国の大型モデル計算の中心に据えた。スタートアップの側から見れば、投資を受けると同時に計算資源を確保する形だった。アリババの側から見れば、投じた資金が自社クラウドの売上として戻ってくる構造だった。双方に利があったからこそ、この関係は容易には揺らがなかった。

法人顧客を相手にする側では、別のプラットフォームが必要だった。周靖人は、企業がわずかなコードで複数のモデルを呼び出して使える「アリババ百煉(Bailian)」、すなわちモデルスタジオを立ち上げ、商用化した。このプラットフォームを通じて、Qwen2.5-TurboやQwen2.5-Plusといった商用モデルが、低い単価で世界各地に

サービス提供された。

モデルが積み重なってきた過程を時系列で追うと、この組織がどのように育ってきたかがはっきりと見えてくる。2023年下半期に初めて登場したQwenは、次に来る単語を予測するトランスフォーマー構造をもとに、ウェブテキストとコード、数学データを併せて学習した成果だった。この時点ですでに、複数種のデータを一つのモデルに混ぜ込む方式が定着していた。

2024年上半期にはQwen1.5とQwen2シリーズが登場した。Qwen 2-72B-Instructを含む複数の密なモデルとともに、全体で570億個のパラメータのうち140億個だけを選んで使う混合エキスパート構造のモデルも、この時期に併せて学習を終えている。モデルの規模をやみくもに大きくするよりも、必要な部分だけを選んで使う方式を、この時点から並行して試していたのである。

2025年1月にはQwen2.5のフラッグシップ・ラインナップが登場した。事前学習データは18兆トークンまで増え、数学問題を専門に解くQwen2.5-Mathと、コードを扱うQwen2.5-Coderのための独自のデータ合成体系もこの時期に整えられた。一つのモデルですべてをこなすよりも、用途に合わせた方向性を増やしていく流れが明確

になった時期である。

同じ年の2月には、視覚知能を備えたQwen2.5-VLが公開された。画像や動画の中のピクセルを読み取る能力を磨く一方で、文書中の表や段落構造を読み取るPDFレイアウト解読技術は、次のモデルの学習データを準備する作業にそのまま使われた。一つのモデルが次のモデルを作る材料になる流れが、ここで確認できる。

3月には音声を実タイムでやり取りするQwen2.5-Omniが登場した。音を離散的なトークンに変換して処理する方式と、リアルタイム対話が可能な構造を備え、テキストと画像に続いて音声までもが一つのモデル系列の中に組み込まれることになった。

2025年5月15日には、次世代の統合推論モデルQwen3が発表された。36兆トークンにのぼる多言語データを学習に用いた。問題の難易度に応じて、どれだけ長く考えるかを自ら調整する「思考モード」と「思考バジェット」という仕組みが初めて導入された。ここにデュアルチャック・アテンションとYaRNという手法を組み合わせ、一度に読み込める文脈の長さを四倍に伸ばした。

こうしてモデルを増やしていく一方で、肝心の、目に見えやすい消費者向けチャットサービスでは苦戦を強いられていた。バイトダンスの豆包(Doubao)が大々的なマーケティングで首位の座を守る間、通義千問単独のアプリはダウンロード数、利用者数のいずれでもほとんど存在感を示せない水準にとどまった。モデルの性能とサービスの人気が必ずしも一致しないという事実を、この一件が示している。

この問題を解決するため、周靖人とアリババの首脳陣は、既存のやり方をそのまま押し通す代わりに組織を組み直す道を選んだ。通義千問という単独のアプリを育てる代わりに、すでに数億人が使っているブラウザやクラウドドライブ、検索ツールの「夸克(Quark)」事業部の中に、Qwenの製品とサービスを移す作業を進めているとされる。

この再編は、単なる組織改編というより、知能をどこに植え込むかという問いへの答えだった。新しいアプリをもう一つ作って利用者を呼び込むよりも、すでに人々の手の中にあるツールの中に知能を溶け込ませるほうがはるかに早いという判断があった。モデルをうまく作ることと、そのモデルを実際に人々のもとへ届けることが

、それぞれ別の課題であることを、この決定はよく示している。

振り返ってみると、周靖人が築いたのは一つのモデルではなく、モデルを絶えず生み出し続けられる仕組みだった。インフラを敷き、オープンソースで開発者を集め、資本とクラウドでスタートアップを引き込み、夸克のような既存サービスに知能を植え込む一連の動きは、それぞれ単独で見れば平凡に見えるが、一つにまとめれば完結した循環構造をなす。

米国の研究所が出したモデルを後追いして真似するやり方では、こうした循環は生み出せない。周靖人が選んだ道は、巨大な計算インフラを備え、その上で最も開放的なオープンソース・エコシステムを育てたのち、モデルをサービスとして販売する形で市場を広げていくというものだった。この順序が入れ替わらなかった点が、この組織の特徴だと言えるだろう。

中核モデルと開発の経緯

中国アリババクラウドの最高技術責任者、周靖人博士が率いるQwen(通義千問)チームは、一つのモデルを磨き上げる水準にとどまらなかった。大規模分散コンピューティング・インフラとデータ精製パイプラインを最初から一体で設計するやり方で研究を導いてき

た。言語モデルと視覚-言語モデル、音声・リアルタイム対話モデルまでを一つのチームの中で順次送り出してきた背景には、こうした全体構造を軸とするアプローチがある。

その成果は、Qwen2、Qwen2.5、Qwen3へと続く言語モデル系列と、視覚と言語を併せて扱うQwen2.5-VL、音声とリアルタイム対話を包含するQwen2.5-Omniへと大きく広がっていった。米国発の構造をそのまま持ち込むのではなく、ハイブリッド・アテンションや動的解像度ビジョン・トランスフォーマー、絶対時間エンコーディングといった独自に設計した要素をあちこちに組み込んだ点も、このチームの特徴である。

大きなモデルが学んだ知識を、小さく軽いモデルに移し替える手法も絶えず磨き上げてきた。いわゆる「強弱蒸留」と呼ばれるこの方式は、巨大モデルを学習させたのちにその知識を小型モデルへ伝え、小型モデルを最初から個別に学習させるよりかはるかに少ない費用で使い物になる性能を引き出す方法である。複数の規模のモデルを同時に世に出せた理由の一つが、ここにある。

2024年7月に公開されたQwen2シリーズは、以降のモデルの基本構造となる骨格を固めました。パラメータ数5億、15億、70億、72

0億の高密度モデル4種に加えて、複数のエキスパートニューラルネットワークを用意し、必要なものだけを選んで使う専門家混合(MoE)構造の57億-A14億モデルも同時に発表しました。このMoEモデルは総パラメータ数が570億に達しますが、トークン1つを処理する際に実際に動くのは140億のみで、300億規模の高密度モデルに匹敵する性能を、サービング(推論提供)コストの3分の1程度に抑えて示しました。

Qwen2の高密度モデル群は、SwiGLU活性化関数と回転位置埋め込み(RoPE)、RMSNormベースの事前正規化を組み合わせで用い、学習を安定させました。大量にサービスを提供する際にメモリを圧迫するキー・バリュー(KV)キャッシュの問題を減らすため、グループ化クエリアテンション(GQA)も導入しました。この手法により、キャッシュサイズを削減しながらも性能を大きく落とさずに済みました。

学習に用いたデータ量も目を引きます。5億パラメータモデルを除く高密度モデルは、7兆を超えるトークンで事前学習を行いました。端末に組み込んで使う超軽量版のQwen2-0.5Bは、知識密度を高めるためにむしろ12兆に及ぶ多言語データを学習に投入しました

。扱える文脈の範囲を広げるため、位置補間技法であるYaRNを適用し、基本のコンテキストウィンドウを12万8000トークンまで拡張して事前学習しました。

2025年1月に登場したQwen2.5は、事前学習データの規模を従来より2倍以上に拡大した系列です。Qwen2が7兆トークンを学習したのに対し、Qwen2.5は18兆トークン規模の高品質な多言語データをゼロから学習し、常識的知識と科学技術分野の知識、そして複数段階を経る推論能力を同時に引き上げました。

モデルを仕上げる後段の学習では、100万件を超える教師ありファインチューニングデータを投入した後、人間の選好を反映する方法として、オフライン方式の直接選好最適化(DPO)とオンライン強化学習方式のグループ相對方策最適化(GRPO)を併せて適用しました。二つの方式を段階的に組み合わせて使った理由は、人間が残したデータを安定的に反映しつつ、リアルタイムの強化学習で細部の性能を継続的に引き上げるためです。

Qwen2.5の基礎重みは、アリババ社内の別の研究組織へと引き継がれ、数学専用モデルのQwen2.5-Mathとコード専用モデルのQwen2.5-Coderへと分岐していきました。一つの基礎モデルを複数分野

の専門モデルの出発点とするこの方式は、後のQwen3のデータ準備過程でも再び用いられます。

視覚と言語を同時に扱うQwen2.5-VLは、2025年2月に技術レポートとして公開されました。リソースの少ない環境から大規模な演算クラスタまで対応できるよう、30億、70億、720億パラメータの3段階で発表しました。従来、画像を決まったサイズに切り出して入力していた前処理方式が画面比率を歪めてしまう問題を避けるため、元の画像の縦横比をそのまま保つ動的解像度ビジョントランスフォーマーを一から新たに設計して学習させました。

演算量を抑えるため、全レイヤーのうち一部の区間にのみウィンドウアテンションを適用する方式も併用しました。動画のように時間が流れる資料を扱う際に時間変化を見失わないよう、従来の1次元回転位置エンコーディングに手を加え、絶対時間に合わせたマルチモーダル回転位置埋め込みを新たに導入しました。この方式によって、数時間に及ぶ長尺の映像でも流れを見失わずに認識できるようになりました。

文書を読み取る能力を高めるため、書式や領収書、図面、数式のように決まった型を持たない画像を、HTMLタグとレイアウト座標

情報に変換した大規模な合成文書データセットを作成し、モデルに学習させました。この学習過程を経て、Qwen2.5-VLは多様な形式の文書を構造化された形で読み取る能力を備えるようになりました。

2025年5月15日に公開されたQwen3は、周靖人が主導してきた研究成果が結集した系列で、推論能力の扱い方において従来のモデルとは明確な違いを見せています。事前学習に用いたトークン数は36兆に達し、これは前作Qwen2.5の18兆トークンから再び2倍に増えたものです。

この膨大なデータを収集する過程で、チームは以前に発表したモデルを道具として活用しました。PDFのような非定型文書から文字を正確に抽出するため、自社の視覚モデルであるQwen2.5-VLをデータのラベリング作業に投入しました。数学とコード関連の合成データを増やすため、Qwen2.5-MathとQwen2.5-Coderが生成したデータも併せて用いました。対応言語の範囲も29言語から119の言語と方言へと大きく広げました。

事前学習は3段階に分けて進めました。第1段階では30兆トークン規模で一般的な知識と言語理解を固め、第2段階では6兆トークン分

の科学技術・コーディングデータを集中的に投入して推論関連の知識を補いました。最終段階では回転位置エンコーディングの基底周波数を1万から100万へと大幅に引き上げ、YaRNとデュアルチャックアテンションを併せて適用し、学習段階から最大3万2768トークンまで文脈を処理できるようにしました。

モデル構成は高密度モデルと専門家混合(MoE)モデルの2方向に分かれます。高密度モデルは0.6億から32億まで6段階の規模を揃え、MoEモデルは30億-A3億と、フラッグシップである235億-A22億の2種類を用意しました。235億-A22億モデルは総パラメータ数が2350億に達しますが、トークン1つを処理する際に実際に起動するのは20億のみで、128個の細分化されたエキスパートニューラルネットワークのうち、トークンごとに8個だけが呼び出される構造を備えています。

このモデルは、常時オンになっている共有エキスパートを別途置かない代わりに、エキスパート間の業務配分が一方に偏らないよう、バッチ単位でバランスを取る損失関数を適用しました。エキスパートごとの役割をはっきりと分けながらも、特定のエキスパートに仕事が集まらないよう調整しています。こうした設計は、大規模サ

ービスにおいて応答速度を一定に保つのに役立ちます。

Qwen3で際立つもう一つの特徴は、一つのモデルの中に、じっくり考えるモードとすぐに答えるモードを同時に組み込んだ点です。複雑な問題を解く際には複数の段階を経て順を追って推論する方式を用い、単純な質問には即座に答えを出す方式を用いるよう、一つのモデル構造の中で切り替えられるようにしました。従来はチャットボット用モデルと推論専用モデルを別々に用意して行き来しなければならなかった手間を、この方式によってなくすことができました。

利用者は、答えを出す前にモデルが生成する思考トークンの数をあらかじめ設定できる「思考予算(シンキングバジェット)」という仕組みも併せて使うことができます。回答速度と精度のどちらに重きを置くかを状況に合わせて調整できるようにした仕組みで、周靖人はインタビューの中で、この機能が実際のサービス環境でコストと品質を同時に管理するために使われていると説明しました。

アリババクラウドは、こうして作り上げたQwen3を、自社のモデルサービスプラットフォームであるアリババクラウド・モデルスタジオ(通称「百煉」)を通じて、世界各国の開発者に低単価で開放し

ました。オープンウェイトとして公開したモデルの重みと、有料APIサービスを併せて運営するこの方式は、開発者には無料でダウンロードして使える道を開きつつ、大規模にサービスを運営しようとする企業には別途の有料アクセス経路を提供する構造です。

言語モデルの外に目を向けると、Qwen2.5-Omniは視覚と音声、文字を一つのモデルの中で同時に扱い、リアルタイムの会話まで対応するよう設計された系列です。文字モデルと視覚モデルを別々に用意してつなぎ合わせるのではなく、一つのモデルが複数の形態の入力を同時に処理できるようにしたこの流れは、先に見た言語モデルと視覚言語モデルの設計経験が引き継がれた結果と見るができます。

Qwen2からQwen3に至るまでの間に、事前学習データは7兆トークンから36兆トークンへと5倍以上に増え、その間モデルは高密度構造中心のものから、専門家混合構造と思考モードを併せ持つ形へと変化してきました。この流れは、単一モデルの規模をやみくもに拡大するのではなく、データの作り方とモデル構造、後段の学習方法を同時に見直しながら効率を高めてきた、周靖人とQwenチームのアプローチを示しています。

現在の立ち位置と中国AIにおける意味

2026年現在、周靖人(ジョウ・ジンレン、Jingren Zhou)博士は、アリババグループの中核事業であるアリババクラウド(阿里雲)で最高技術責任者(CTO)を務めています。彼が指揮する組織は、2019年からアリババが積み重ねてきた大規模自然言語処理と視覚・音声知能の研究を一つにまとめた組織で、周靖人はこの組織全体の研究方針と資源配分を決める立場にあります。

彼が指揮する組織こそが通義千問(Tongyi Qianwen)、すなわちQwenチームです。このチームはアリババ内部で自生的に育ってきた研究組織で、四半期ごとに新しいファウンデーションモデルを発表し、事前学習からアライメント工程まで全過程を周靖人の承認のもとで進めています。

学術界において周靖人は、ACMフェローとIEEEフェローの両方に名を連ねる人物です。彼の学問的な基盤は米国コロンビア大学大学院のコンピュータサイエンス課程で築かれ、この時期の訓練が、後の大規模分散システムを設計する感覚へとつながったと見ることができます。

彼が長年交流してきた人物の顔ぶれも、彼の工学的な背景をうかがわせます。分散データベース分野の重鎮であるモハン・C.(Mohan C.)博士、分散ディープラーニングエンジンApache SINGAを設計したシンガポール国立大学のベン・チン・オイ(Beng Chin Ooi)教授、IEEEコンピュータ協会のデイビッド・ロメット(David Lomet)博士らとの学術的な交流がその例です。

こうしたシステム設計の経験は、数万個のGPUを遅延なく同期させて稼働させる、アリババクラウドの大規模事前学習パイプラインを構築する土台となりました。モデルの知能を高めることと同じくらい、その知能を学習させるハードウェアを効率的に動かすこともまた、周靖人の長年の関心事だったようです。

2025年5月15日、Qwenチームは次世代のフラッグシップモデルQwen3を公開しました。密構造(Dense)と専門家混合(MoE)構造を併せ持つこのモデル群は、パラメータ規模を0.6Bから235Bまで細かく分けて設計されています。

Qwen3の事前学習には、36兆個にのぼる多言語トークンが投入されました。この膨大なデータを精製する過程で、Qwenチームは自

社の視覚言語モデルであるQwen2.5-VLをレイアウト解析器として動員し、形式がまちまちなPDF文書からテキストを抽出する作業を自動化しました。

Qwen3には、二つの応答方式が一つの重みの中に統合されています。複数の段階を経て答えを推論する思考モード(Thinking Mode)と、即座に答える非思考モード(Non-thinking Mode)がそれぞれ、ユーザーはこの二つを行き来しながら、必要に応じて応答方式を選ぶことができます。

この思考モードには、思考予算(Thinking Budget)という調整装置が付いています。ユーザーが推論過程で生成する思考トークンの数の上限を自ら設定できるようにしたもので、応答速度と正答率の間で望ましいバランスを見つける助けとなる装置です。

対応言語の範囲も大きく広がりました。それまで29言語に対応していたQwenシリーズは、Qwen3に至って119の言語と方言まで扱うようになりました。これはアリババが中国国外の市場までも見据えているというシグナルと読み取れます。

2026年7月14日には、音響知能分野へ領域を広げたQwen-Musicが公開されました。テキストから音楽を作り、曲のスタイルを別の曲に移すという作業を一つのモデルで処理するこのシステムは、500万時間を超える多言語音楽データを学習しています。

Qwen-Musicで目を引く点は、メロディを自ら構想する過程を推論の一段階として組み込んだことです。歌唱を始める前に曲の基本構造となるメロディをあらかじめ思考として組み立てる、いわゆるメロディ思考連鎖(Melody-CoT)方式と、高解像度のステレオレンダラーを併せ持っています。

2026年上半期には、エージェント型モデルの拡張も続きました。汎用エージェントモデルであるQwen3.7-Max、視覚と動作を同時に処理するQwen3.7-Plus、60言語を2.5秒以内の遅延でリアルタイム通訳するQwen3.5-LiveTranslateが相次いで登場しました。

動画生成の分野では、200億パラメータ規模の万相(Wan)2.7シリーズが登場しました。画像生成用のWan2.7-Image、テキストを動画に変換するWan2.7-T2V、参照動画をもとに生成するWan2.7-R2Vまでを備え、テキストと画像、音声、動画を横断するモデル群を完成させました。

こうして積み上げられた技術は、アリババクラウドが企業顧客に販売するサービスの基盤となります。周靖人が立ち上げたアリババ百煉(阿里云百炼、Model Studio)プラットフォームは、企業がわずかに数行のコードでQwenシリーズのモデルを呼び出して使い、自社のデータに合わせて調整できるようにしたサービス型モデル(MaaS)ハブです。

アリババクラウドは、中国国内で独自に大規模モデルを開発しようとする複数のスタートアップに出資する形でも影響力を広げました。こうした投資を通じて、新興企業の研究インフラをアリババクラウド上に載せる流れが生まれ、結果としてアリババは中国の大規模モデル演算インフラ市場で大きな地位を占めるようになりました。

米国の半導体輸出規制という制約下でも、Qwenチームは独自の訓練インフラとアルゴリズム最適化で対応してきました。YARN手法で文脈長を延ばし、デュアルチャンクアテンション(Dual Chunk Attention、DCA)と事前正規化(pre-normalization)といった手法を組み合わせ、36兆トークン規模の事前学習を安定的に収束させたことがその一例です。

こうした技術の蓄積は、中国の大規模モデル市場の構図にも影響を与えました。従来型のポータル企業が消費者向けサービスで停滞を経験する間、Qwenは中国発の大規模モデルの中でも実質的な競争力を備えた軸として位置づけられていきました。

周靖人が推し進めたもう一つの方法は、完全なオープンソース配布です。Qwenチームは主要モデルの重みをApache 2.0ライセンスでHugging FaceとGitHubに無償公開しました。これは技術を隠すよりも広く行き渡らせてエコシステムを育てる道を選んだ決断でした。

この開放戦略の結果、世界各地の中小開発企業や研究機関がQwenアーキテクチャをもとに自社のサービスを構築するようになりました。モデルを無料で公開しても、その上に築かれるエコシステムがアリババクラウドへと還元されるという計算が、この戦略の根底にあります。

性能の面でも、Qwenシリーズは国際舞台での存在感を高めています。LMSYS Chatbot Arenaのようなベンチマークでは、Qwen3-235B-A22BとQwen2.5-MaxがGPTシリーズ、Geminiシリーズ、Claudeシリーズのモデルと並んで上位に位置し、順位を競い合ってい

ます。

Qwenチームが歩んできた道を振り返ると、周靖人の役割は目先の消費者指標を追いかけるよりも、演算資源とデータ規模を年々拡大していく長期的な設計に近いといえます。分散システムを手がけていた彼の初期のキャリアが、大規模モデルの学習インフラを築く仕事へと自然につながっていったようです。

Qwen2シリーズの技術報告書に盛り込まれた基礎アルゴリズムの最適化、Qwen2.5-Coderが示したコード特化学習、Qwen2.5-Omn iのオムニモーダル統合まで、周靖人が率いるチームの仕事は一つのモデルをうまく作ることにとどまらず、言語と視覚、音声、コードを横断する一つの技術の軸を築くことへとつながりました。

こうした流れすべての中心には、アリババクラウドというインフラとQwenというオープンソースエコシステムを共に握る周靖人の立場があります。彼はモデルを作る研究者であると同時に、そのモデルが動くクラウドを設計する技術者でもあり、二つの役割を一人で兼ねているという点が、彼の位置を他の研究責任者たちと区別する特徴です。

第9章 姚順雨(ヤオ・シュンユー、Shunyu Yao)、ReActとTree of ThoughtsからTencent Hunyuanへ渡ったエージェント研究者

成長期と学業

姚順雨は1998年、中国安徽省合肥に生まれた。幼い頃から数学とコンピュータアルゴリズムに才能を見せ、その才能は地元の学校を経て全国規模の大会へとつながっていった。合肥第45中学を終えて合肥第一中学に進学した彼は、中学・高校時代を通じて理系科目で頭角を現しながらも、決して一方に偏った生徒ではなかった。担任教師の回想によれば、彼は自分で計画を立てて実行する性向がはっきりしていたという。文章を書くことや音楽、バスケットボールのような活動にも関心を絶やさなかった。

2014年、高校生だった姚順雨は全国情報学オリンピックに出場し、銀メダルを獲得した。この大会は中国でコンピュータサイエンスに才能のある生徒を選び出す代表的な登竜門であり、アルゴリズムと離散数学を扱う能力を競う場である。翌2015年には、中国の大学入試である高考で安徽省理系全体3位の成績を収めた。この二つの成果が重なり、彼は清華大学交叉信息研究院のコンピュータサイエンス専攻、通称姚班と呼ばれる課程に入学する資格を得た。

姚班はチューリング賞受賞者である姚期智教授が創設した課程で、コンピュータサイエンスの基礎を極めて緻密に教えることで知られている。姚順雨は2015年から2019年までここで学部課程を歩み、アルゴリズムと数学の基本を鍛えた。彼は学業だけが続けたのではなく、姚班の学生会長を務め、同級生同士の学術交流と行政的な連絡を調整する役割も併せて担った。後に彼が数々の研究プロジェクトを組織し、他者と協働する方式には、この時期の経験が土台になったと推測できる。

2019年に清華大学を卒業した姚順雨は、米国プリンストン大学コンピュータサイエンス学科の博士課程に進学した。彼が向かう研究テーマを定めるまでには時間が必要だった。すでに合格通知を受け取っていたにもかかわらず、5年を捧げるに値する研究の方向をなかなか定められずにいた。そんな中、彼はプリンストン大学の助教授だったカーティク・ナラシムハンに連絡を取った。当時オープンAIが公開したGPT-2のような事前学習言語モデルが示した文脈理解能力に注目し、このモデルをテキスト生成だけに閉じ込めず、テキストベースのゲーム環境の中で自ら判断し行動させることができるかを尋ねた。

ナラシムハン教授はこの提案を受け入れた。二人はその後5年間、言語モデルを行動する主体へと変える一連の研究を共に築いていた。ナラシムハン教授は2017年から2018年にかけてオープンAIに客員研究員として滞在し、初期のGPT論文の共著者として参加した経歴があった。彼が積み重ねた言語モデル研究の経験は、姚順雨のプリンストン博士課程全体に影響を及ぼした。

プリンストンで姚順雨が没頭した問題は、言語モデルを次の単語を予測する道具から脱却させることだった。彼はコンピュータ画面やウェブページ、コードリポジトリのような実際の環境の中で、モデルが目標を立てて行動を取るよう設計する研究を続けた。この流れから生まれた最初の成果がReActである。ReActは、モデルがある行動を取る前に、なぜその行動を選んだのかを人間が読める言葉で記述させ、その記述をもとに実際の行動へとつなげる方式である。

推論と行動を一つの循環の輪の中に結びつけるという発想は、当時としては新しいアプローチだった。モデルが即座に答えを出す代わりに中間過程を言葉で示せば、その過程を人間が点検でき、モデル自身も自分の判断を振り返ることができる。ReActはその後、多

くの研究者や企業が言語モデルを自律的なアシスタントに仕立てる際に参照する基本的な枠組みの一つとして定着した。

姚順雨はここで止まらず、モデルの推論方式そのものをさらに洗練させる研究へと進んだ。その成果が2023年に発表したTree of Thoughtsである。従来の思考連鎖方式は一本道の推論経路しか辿らないため、途中で行き詰まると引き返す方法がなかった。Tree of Thoughtsはグラフ探索の技法を借用し、モデルが複数種類の推論候補を同時に作り出し、その中から有望なものを選んで進み、行き詰まった経路は前の段階に戻って別の候補を試みるよう設計した。

この方式は、数字パズルや創作的な文章のように一度の計算では答えが出ず、複数の段階を経なければならない課題でとりわけ効果を発揮した。モデルが自ら複数の案を立てて比較した上で選ぶことで、単一の経路だけを押し進めていた従来の方式よりも誤答から抜け出す余地を広げたようだ。この論文は、姚順雨の名を言語モデル推論研究の分野にはっきりと刻む契機となった。

プリンストン博士課程の間、姚順雨は言語エージェントの能力をどう測定するかについても共に考えた。既存の評価方式が概ね仮想のゲーム環境にとどまっていたのに対し、彼は人間が実際に直面す

る業務に近い評価の枠組みをいくつも打ち立てた。WebShopは、モデルがウェブショッピングモールで商品を探しカートに入れる課題を扱う、インターネットベースのウェブエージェント評価としては初期の事例に属する。InterCodeは、モデルがプログラムコードを書き実行結果を確認しながら問題を解決するよう設計した、初期のコーディングエージェント課題だった。

このほかにも彼は、モデルが実際のGitHubの 이슈を追跡して大規模なコードリポジトリ内の欠陥を見つけ修正するSWE-bench、ユーザーとツール、システムが絡み合う相互作用を長時間にわたって観察するtau-benchのような評価体系を打ち立てる作業に参加した。こうした評価の枠組みは、言語モデルが試験問題をうまく解けるかどうかではなく、実際の業務環境でどれほど役立つ形で動けるかを測る尺度として使われた。

2024年、プリンストンで博士号を取得する際に提出した論文の題目は《Language Agent: from next token prediction to digital automation》だった。次の単語を予測する言語モデルから出発し、デジタル環境の中で自ら業務を処理する自動化の主体へと進んでいく過程をまとめた論文であり、彼が5年間続けてきた研究の方向を一

文で要約してくれる題目だと言える。

工学論文を書く合間に、姚順雨は多分野の本を幅広く読んできた人物としても知られている。彼の個人ホームページに掲載された読書リストには、ジェームズ・ワトソンの《二重らせん》、デイビッド・トンの一般相対性理論講義録、エリザベス・スペルキの《赤ちゃんは何を知っているか》、リチャード・ハミングの《科学と工学をするアート》のような本が並んでいる。生物学と物理学、認知心理学を横断するこのリストは、彼の関心がコンピュータコードの中だけにとどまらなかったことを示している。

プリンストンで積み重ねた研究は、卒業と同時に彼をオープンAIへと導いた。彼はオープンAIで、コンピュータを代わりに操作するエージェント製品であるオペレーターと、科学や金融のように情報量の多い分野で深みのある調査を行うディープリサーチの開発に参加した。ReActとTree of Thoughtsで鍛えた推論と行動の結合、そして実際の業務に近い評価方式を打ち立てた経験が、これらの製品の設計にそのまま受け継がれた。

姚順雨の学生時代と大学院課程を振り返ると、一つの流れがはっきりと見える。彼は理論を打ち立てるだけにとどまらず、その理論

が実際の環境で機能するかを常に併せて確認しようとした。合肥の学校でオリンピックの問題を解いていたやり方も、清華大学姚班で学生会を率いていたやり方も、プリンストンで言語モデルに行動を任せる研究を続けていたやり方も、いずれも結果を目で確認できる形にすることに重きを置いていた。

2025年末、テンセントが彼を最高AI科学者として迎え入れたことで、プリンストンで鍛えた言語エージェント研究はそのまま産業の現場へと移された。

研究者、あるいは起業家として定まる時期

姚順雨(ヤオ・シュンユー)という名前が中国AI業界に広く知られるきっかけとなったのは、2026年6月、テンセントクラウドAI産業応用大会の舞台だった。この場で彼はテンセントの首席AI科学者としての肩書きで初めて大衆の前に立った。テンセントがAI競争で後れを取っているという世間の指摘に直接答えた。だが、その答えの重みを理解するには、彼がその場に立つまでに歩んできた道、すなわちプリンストンの大学院生からオープンAIの研究員を経てテンセントの首席科学者になる過程を見ておく必要がある。

姚順雨は米国プリンストン大学コンピュータサイエンス学科で博士課程を歩みながら、言語モデルを扱う自分なりの視点を築き上げた。彼が身を置いた時期は、巨大言語モデルがテキストをそれらしくつなげる道具から、自ら判断し行動するエージェントへと移り変わる転換点だった。彼はこの転換の真ただ中で論文を書き、その論文群は後に業界全体が従う方法論となった。

彼の博士学位論文の題目は《Language Agent: from next token prediction to digital automation》、すなわち言語エージェントが次のトークン予測からデジタル自動化へと進んでいく過程を扱ったものだった。題目そのものが彼の研究人生を要約している。コンピュータが次に来る単語を一つ当てるという、一見素朴な計算を繰り返させるだけで、ウェブブラウザを操作しコードを修正することまでやり遂げられるという確信がその根底に流れている。

この確信を実際の方法論として初めて世に送り出した論文が、2022年に発表されたReActである。ReActはモデルがすぐさま行動に出ることをしなかった。行動に先立ち、なぜそう判断したのかを人間が読める言葉で説明させる方式である。推論(Reasoning)と行動(Acting)を交互に織り込んだことから、名前もReActとなった。

この方式が重要な理由は単純である。モデルが一人で内心で判断してすぐさま行動すれば、その判断がどこで狂ったのかを人間が辿るのは難しい。ところがモデルが段階ごとに自分の考えを言葉で示せば、人間もそれを見守ることができ、モデル自身も辻褄の合わない考えを取り除く余地が生まれる。その結果、でたらめを作り出す誤りが減り、行動を人間が制御できる余地も増えた。

ReActは発表されてほどなく、エージェントを作る者なら誰もが通らねばならない出発点となった。今、ウェブを検索し、ツールを呼び出し、複数の段階を経て問題を解くAIエージェントの大半がこの枠組みの上で動いている。姚順雨が大学院生時代に打ち立てたこの方法論一つが、その後数年にわたって続く彼の研究方向全体を規定したと見ることができる。

ReActの次に出てきたのが2023年のTree of Thoughts、略してToTである。ReActが思考と行動を一本の線で交互につなぐ方式だったとすれば、ToTはその思考の筋を木のように複数の種類へと枝分かれさせる方式である。モデルが一つの解法だけを押し通す代わりに、複数種類の解法を同時に広げて見て、行き詰まった道が出れば引き返してまた別の道を探すようにする。

この後戻りする動作をバックトラッキングと呼びます。人が難しい問題を解くとき、ひとつの方法が行き詰まれば最初に戻って別の方法を試すのと似ています。ToTはこの自然な思考の習慣をコンピュータに模倣させるため、探索アルゴリズムを言語モデルの文章生成プロセスに組み込んで実装しました。おかげでモデルは、数学オリンピックの問題のように複数の段階を経て解く難しい課題において、一度に答えを出す代わりに、複数の候補となる解法を比較しながら答えに近づけるようになりました。

ReActとToTがモデルの思考と行動の仕方を変えた研究だったとすれば、ヤオ・シュンユーはその成果を測る物差しにも手をつけました。既存の定型的な問題解決データセットだけでは、エージェントが実際に役立つかどうかを見極めるのは難しいと考えた彼は、人が実際に働く環境をそのまま移した評価ツールをいくつも作りました。

その中のWebShopは、実際のオンラインショッピングサイトのよう
に商品を検索しカートに入れて決済まで進む一連のやり取りを評価する枠組みです。ヤオ・シュンユーはこれをインターネット上のウェブエージェントを試す最初の課題として紹介しました。人がも

のを買おうとして検索し、比較し、選ぶというごく当たり前の過程を、エージェントがどれだけ人間らしくこなせるかを測る物差しにしたようです。

もうひとつは、コーディングエージェント向けの課題として紹介されたInterCodeです。ヤオ・シュンユーはこれをコーディングエージェント評価の初期形態として言及しました。実際の開発者がコードを直し、実行結果を確認してはまた直すという反復作業を、言語モデルがどこまで真似できるかを試すツールでした。この流れは後に、実際のGitHubのissueを読み、膨大なコードリポジトリの中から自らエラーを見つけて修正するソフトウェアエンジニアリングエージェントの評価へとつながっていきました。

ヤオ・シュンユーはこうした評価を作るたびに、同じ原則を繰り返し口にしました。モデルの性能は、決められた試験問題をどれだけうまく解けるかではなく、人が実際に直面する仕事を代わりにこなせるかどうかで判断すべきだという原則です。ベンチマークのスコアを上げることと、実際に役立つ仕事をこなすことは違うという区別を、彼は大学院時代から一貫して強調してきました。

プリンストンで培われたこの研究の感覚は、その後彼が場所を移してからそのまま引き継がれました。彼は博士課程を終えた後、OpenAIで中核となる研究員として働き、エージェントの方法論を実際のサービスに仕立て上げる経験を積みました。コンピュータの画面を人間のように操作するエージェントや、インターネットを調べて調査報告書を自律的に作成するツールといった製品が、この時期に生まれました。

大学院で論文として証明した方法論が、実際の利用者が対価を払って使う製品へとつながっていく過程を経験しながら、ヤオ・シュンユーは研究者でありながら製品を見る目も同時に養いました。この経歴が、後に彼がテンセントに移り、研究組織と製品組織を共に率いる立場に就く背景となります。

2025年末、ヤオ・シュンユーはテンセントに加わり、最高経営責任者直属のチーフAIサイエンティストに任命されました。AIインフラ部門と大規模言語モデル部門を共に率いる立場を任されました。テンセントはこの人事を大規模言語モデルの研究開発組織を根本から再編するシグナルと位置づけ、ヤオ・シュンユーはテンセントHunyuan大規模モデルの研究開発を統括する中心的責任者となります。

した。

テンセントに加わった後、彼が投げかけたテーマは大学院時代から培ってきた問題意識と大きく変わりません。彼はAI競争の前半戦、すなわちより大きなモデルを学習させランキング表を押し上げる戦いはすでに幕を閉じ、コーディングエージェントとマルチモーダル、身体性を持つ人工知能が競い合う後半戦がいままさに始まったと語りました。

ヤオ・シュンユーは、この後半戦では訓練そのものよりも評価と企画がより重要になると見えています。どの任務をAIに任せるかを正確に定める作業、そしてその任務をうまくこなせたかどうかを測る評価体系を築く作業こそが、次の段階の核心だという判断です。これは彼がプリンストンでWebShopやInterCodeといった評価ツールを自ら作っていた頃の問題意識とつながる考えです。

彼はテンセントでHunyuan 3を開発する過程でインフラを大幅に再構築し、データと評価体系にも大きな変更を加えたと明かしました。あわせて事後学習を担う精鋭人員をテンセントの大衆向けサービスであるYuanbaoチームに優先的に投入し、モデルを磨き上げる作業と実際の利用者の反応をつなぐ回路をしっかりと固めたと説明

しました。

ヤオ・シュンユーはコストパフォーマンスという言葉で性能とコストの二つに分けて説明しました。まず性能を見て、その後にコストを検討するという順序です。本当の強みは、大きくないモデルでも価値ある任務をきちんとこなせる点にあると語りました。資本と電力が潤沢ではない環境で規模を競う代わりに、小さなモデルを精緻に磨き上げる道を選びました。

テンセントがAI競争で後れを取っているのではないかという質問に対し、彼はGPTやClaudeだけが唯一の答えになるわけではないと答えました。テンセントは指標よりも信頼を軸に動く組織だという説明も付け加えました。これは彼が大学院とOpenAIを通じて身につけた評価中心の思考と、大きな組織の中で合意を形成していくやり方が交わる地点を示しています。

プリンストンでナラシマン教授とともに言語エージェントというアイデアを論文として練り上げた大学院生が、数年のうちに中国を代表するテクノロジー企業の研究開発の方向を定める立場に就いたということのようです。ReActとToTという二つの方法論、そしてWebShopのような評価ツールは、その間に彼が積み重ねてきた成

果の基本構造をなしています。

ヤオ・シュンユーの経歴を振り返ると、彼は創業者と呼べるような会社を自ら立ち上げたことはありません。それでも彼が大学院時代に築いた方法論と評価哲学は、テンセントという巨大組織のAI研究の方向性を決める立場へと彼を導き、その意味で彼は研究者でありながら同時に組織を率いるリーダーとしての地位を固めたと言えるでしょう。

会社と研究組織を作り上げた過程

ヤオ・シュンユー(姚顺雨)が会社を興したやり方は、Moonshot AIを立ち上げたヤン・ジーリンやBaichuanを立ち上げたワン・シャオチュアンとは異なります。彼は自分の名前を冠したスタートアップを興しませんでした。彼はプリンストン大学コンピュータサイエンス学科でカーティック・ナラシマン(Karthik Narasimhan)教授のもとに籍を置き、言語エージェントの研究を磨きました。2024年8月に学位を修了した後はOpenAIの研究チームに加わり、OperatorやDeep Researchといった製品開発に参加しました。会社と組織を実際に作り上げる舞台は、その後、すなわち2025年12月にテンセントへ加わったことで開かれました。

プリンストン時代とOpenAI時代がこの節の中心ではありません。ただ、彼がテンセントで即座に組織設計の権限を与えられた背景だけは押さえておく必要があります。言語エージェントの理論を生み出した研究者でありながら、同時に大規模製品の実戦開発に参加した経歴を併せ持つ人物は珍しいものでした。テンセントはその稀な組み合わせに組織再編の重みを託しました。プリンストンでは実験室規模のアイデアを扱っていたとすれば、テンセントではそのアイデアを数万台のサーバーと数億人の利用者の前に立たせる仕事が彼を待っていました。

テンセントがヤオ・シュンユーを迎え入れた時点は、社内外で静かに見過ごせる出来事ではありませんでした。証券時報の報道によれば、テンセントは彼の加入を機にAI戦略組織の構造を大幅に手直ししました。人材採用にも拍車をかけました。表面的にはひとりの人物を招いた人事でしたが、実際には組織図全体を描き直す作業でした。テンセントのようにすでに膨大な組織を抱える企業でこれほどの規模の改編が行われるということは、大規模モデル競争がもはや一部門の実験的な課題ではなく、会社全体の存亡がかかった課題へと格上げされたことを意味してもいます。

2025年12月、テンセントは大規模モデルの研究開発アーキテクチャをアップグレードすると発表しました。このとき新設した部門がAIインフラ(AI Infra)部、AIデータ(AI Data)部、データコンピューティングプラットフォーム部の三つです。名前だけを見れば技術部門の再編のように聞こえますが、三部門を一挙に新設したという事実は、テンセントが大規模モデル競争を会社の最優先課題に引き上げたことの表れと読めます。部門をひとつ増やすことと三つ同時に立ち上げることとは性格が違います。後者は既存の組織に手を加える程度ではなく、大規模モデルの研究開発全体をゼロから設計し直すという意志の表明に近いものです。

AIインフラ部は大規模分散訓練と低遅延の推論サービスを支えるハードウェア・ソフトウェア基盤を担います。AIデータ部はモデル学習に使うデータを整え、評価基準を築く役割を担います。データコンピューティングプラットフォーム部はこの両者をつなぐ演算インフラを管理します。三つに分けた理由は、ひとつのモデルを訓練するのに必要な作業がそれだけ多層に分かれていたからです。モデルが良い答えを出すには良いデータが必要であり、良いデータを大量に処理するにはそれ相応の演算インフラが必要であり、そのインフラを実際に動かすには安定したハードウェアとソフトウェアの基

盤が必要です。テンセントはこの三つの層をそれぞれ独立した部門として切り分けることで、初めて互いにかみ合って動くようにしようとした。

ヤオ・シュンユーはこの改編とともに「CEO/総裁室」所属のチーフAIサイエンティストに任命されました。彼はテンセント総裁のマーティン・ラウ(刘炽平)に直接報告する立場を与えられました。会社を興したことの無い元OpenAI研究員が、中国の大手プラットフォーム企業の最高意思決定ラインに即座に連なった事例です。通常、外部から招かれた科学者は研究所の所属として入り、いくつかの段階を経てようやく経営陣と向き合う場合が多いのですが、ヤオ・シュンユーは加入と同時に総裁直轄のラインに立ちました。

同時に彼はAIインフラ部と大規模言語モデル部の責任者の地位も兼任しました。この地位はテンセントの技術工程事業群(TEG)を率いるルー・シャン(卢山)に報告する構造でした。ヤオ・シュンユーは一方では総裁室に、もう一方ではTEG組織にまたがる二重の報告ラインを持つことになりました。ひとつの肩書きが二つの系統に同時にまたがる構造は珍しい配置です。テンセントがそれだけこの組織に力を入れたことの表れと見ることができます。

こうした二重構造は、研究の方向性を定める権限と、実際にサーバーを動かしモデルを訓練する実行権限をひとりの人物に集約する効果を生みます。研究者がアイデアだけを出し、インフラチームは別に動くという方式とは異なります。ヤオ・シュンユーがプリンストンの研究室で身につけたのは、理論と実行コードを共に扱う習慣でした。テンセントはその習慣をそのまま組織設計に落とし込みました。

これまでテンセントの大規模言語モデル開発を率いてきたワン・ディ(王迪)大規模言語モデル部副総経理は、ヤオ・シュンユーの下に配置が変わりました。長く実務を担ってきた管理職が新しく加わった科学者の下部ラインに再配置されたという事実は、今回の組織改編が形式的な肩書きの調整にとどまらなかったことを示しています。長年在籍してきた管理職を新規加入の外部人材の下に移す決定は、社内で容易に下せるものではありません。それだけテンセントがヤオ・シュンユーの判断を信頼していたことの表れと読み取れます。

証券時報はこの人事をテンセントのAI戦略が加速している兆候だと解説した。姚順雨をテンセント混元(Hunyuan)大規模モデル研究

開発の中核責任者と位置づけたのである。混元はテンセントが長年育ててきた自社製大規模モデル系列であり、その指揮権が新たに加わった研究者に委ねられたことになる。中国メディアが人事のニュースを組織戦略転換の兆候として読み解いたという事実そのものが、この招聘が業界内でどれほど重く受け止められたかを物語っている。

テンセント混元モデルの技術的な土台を押さえておくと、2024年11月に公開されたHunyuan-Largeは、3890億のパラメータのうち520億のみを活性化させるオープン型の専門家混合(MoE)モデルであった。入力ごとに全パラメータを使うのではなく、必要な専門家サブネットワークだけを選んで使う構造であるため、大規模なモデルを動かしながらも演算コストを抑えられる。Hunyuan-Largeの技術報告書は、大規模な合成データ、専門家ルーティング戦略、キー・バリューキャッシュの圧縮技法、専門家ごとの学習率調整を核心技術として挙げていた。姚順雨が加わる以前から、テンセントはこの規模のモデルをすでに自社技術で構築し終えていたということである。

姚順雨が引き継いだのは白紙ではなく、すでに稼働していた大規模モデル基盤であった。彼がテンセントで手がけたのは新しいモデルをゼロから作り上げる作業ではなく、既存の研究開発体制を組み直し、その上に新しい指揮系統を据える作業に近かった。すでに性能が実証されたモデルとインフラが存在していたからこそ、彼は技術を一から証明する時間を費やす代わりに、すぐさま組織構造の手直しに取りかかることができたのである。

テンセント創業者でCEOの馬化騰(マー・ホアテン)は2026年1月26日の公式発言で、姚順雨の合流以降、会社が人材招聘のスピードを引き上げたと明らかにした。彼はこの一年でテンセント混元大規模モデルが「深い再編」を経たとも語った。最高経営責任者自らが組織改編に言及したという事実は、この人事が取締役会レベルでも重く扱われたことを示している。

馬化騰はテンセント混元の研究開発組織が社内的にコデザイン(Co-design)、つまりインフラ設計とアルゴリズム設計を同時並行で進める方式を加速させているとも説明した。モデルを作るチームと、そのモデルが動くハードウェア・ソフトウェア基盤を作るチームが別々に動くのではなく、最初から同じ目標を見据えて設計をすり

合わせるという意味である。これは先に述べた姚順雨の二重報告構造とそのままつながっている。

混元大規模モデルとテンセントの看板チャットサービスである元宝(ユエンバオ)との相乗効果を高める作業も、この改編の目標の一つとして挙げられた。研究室から生まれたモデルが実際のユーザーサービスとどれだけうまく連携して動くかが、テンセント経営陣の関心事であったということだ。研究成果を論文だけにとどめず、製品へと直結させようとする意図がこの発言からうかがえる。

姚順雨自身の発言も、彼がテンセントで担う役割が研究顧問にとどまらないことを裏づけている。2026年6月5日、テンセントクラウドAI産業応用サミットにおいて、彼はテンセントクラウド・インテリジェント産業グループCEOのドウソン・トン(唐道生)とともに壇上に立った。この場で彼の肩書きは、テンセント混元の大規模言語モデルとAIインフラを統括する首席AI科学者として紹介された。

この場で姚順雨は、テンセントがAI競争で後れを取っているという懸念に対して明確に否定した。彼はAI競争がようやく後半戦に入ったばかりだと述べ、いまの状況を1970年代にパーソナルコンピューターが普及し始めた時期になぞらえた。競争で数歩遅れている

という指摘が出るたびに焦って反応するのではなく、ゲームはまだ序盤だという大きな枠組みで話を展開したようである。

彼はChatGPTとClaude Codeだけが唯一のスーパーアプリとして残るとは考えていないとも語った。そのような世界はひどく息苦しい世界であり、新たな機会は絶えず生まれてくるはずだという趣旨であった。コーディングエージェントや生産性ツールが重要性を増していく流れは確かだが、まだ満たされていない兆単位の市場が残っており、マルチモーダルAIや身体を持つ知能(embodied intelligence)といった新領域も次々と開かれているという説明も付け加えた。

姚順雨は、大規模言語モデルとAI製品を作る初期段階では試行錯誤と後退が自然なことだという点も認めた。ただし後半戦で重要なのは、足りない部分を正直に認め、調整するにあたって忍耐強くあることだと語った。研究者時代に失敗した実験をそのまま論文に残していた態度が、会社の経営陣としての語り口に移し替えられた姿である。

こうした発言は、彼がテンセントで作り上げた組織の性格をうかがわせる。成果を性急に急ぎ立てる組織ではなく、失敗を認めて次

の実験へ移る研究室の習慣を、大企業の内部にそのまま植え付けようとする組織である。プリンストンの研究室で培われたやり方が、テンセントというはるかに大きな舞台でも維持されているように見える。論文を書いていた人物が会社を率いると、実験ノートに記していた正直な失敗の記録が経営の言葉に姿を変えて現れるようだ。

姚順雨がテンセントクラウドのイベントにCEO級の人物と並んで登場したという事実も注目に値する。彼は研究所の内側にとどまる科学者ではなく、会社の対外メッセージをともに発信する立場に立っていたのである。この点は、彼が担っている位置が研究顧問ではなく事業の方向性まで包括する経営ラインであることを改めて裏づけている。テンセントクラウドCEOと並んで壇上に立つということは、彼が混元モデルの技術責任者を超えて、テンセントのAI事業全体を対外的に代表する顔の一つとして位置づけられたということでもある。

こうして見ていくと、姚順雨がテンセントで作り上げたのは一つの製品ではなく組織そのものである。AIインフラ部門、AIデータ部門、データコンピューティングプラットフォーム部門という三つの部署、総裁室とTEGを同時につなぐ二重報告ライン、既存の実務責

任者を再配置した人事、インフラとアルゴリズムを一体として設計するコデザイン方式まで、彼が合流してから数か月のうちにテンセント混元の研究開発組織の基本構造が新たに組み直された。

この改編が目を引く理由は、姚順雨が創業者ではないという事実にある。彼は自己資本で会社を立ち上げる代わりに、すでに存在する巨大企業の内部で研究権限とインフラ実行権限を同時に握る立場を手に入れた。テンセントという既存の資本とクラウド基盤を梃子にして、スタートアップであれば何年もかかったはずの組織構築を短期間で成し遂げたようである。

他の大規模モデル系スタートアップが毎月膨大な資金を燃やしながらインフラコストに苦しめられているの比べると、テンセントという枠組みの内側で組織を作る方式は資金調達の圧力からかなり自由である。姚順雨は投資家を説得する代わりに、テンセント取締役会の決定一つで大規模なクラウド演算力と組織再編の権限を同時に手にしたのである。毎月投資家に進捗を報告し次のラウンドに備えなければならないスタートアップ創業者と、取締役会がすでに承認した予算の枠内で組織を組み立てる彼の立場は、根本的に異なる。

こうした構造は利点とともに負担も伴う。テンセントの内部では資金の心配なくインフラを組み直せる一方で、総裁室とTEGの双方に成果を報告しなければならない立場でもある。姚順雨が率いる組織が成果を出せなければ、その責任もまた創業者個人ではなくテンセント取締役会のレベルで評価される構造である。資本の安定性と成果への圧力が同時について回るポジションだということである。

むろん、この過程で確認できていない部分もある。テンセントが姚順雨の招聘に正確にいくら支払ったのかである。

プリンストンで培ったReActやTree of Thoughtsといった言語エージェント研究、OpenAIで磨いた実践的な製品感覚は、結局のところテンセントという舞台で組織を設計する能力へとつながった。姚順雨の経歴において、会社と研究組織を作り上げる過程は、創業ではなく既存の大企業内部に新たな指揮系統を植え込む作業として現れた。

その結果が、いまのテンセント混元組織の姿として残っている。AIインフラ部門とAIデータ部門、データコンピューティングプラットフォーム部門がそれぞれ担う階層で動き、その上に総裁室とTEGへとつながる二重報告ラインがかかっており、その中心に姚順雨と

いう一人の名前が置かれている。創業をせずとも組織を築き上げるとはどういうことか、彼のテンセント合流以後の数か月がそれを示しているようである。

核心モデルと開発過程

姚順雨博士が学界にその名を知らしめたきっかけは、大規模言語モデルを、テキストを吐き出すだけの道具から、自ら判断し行動する存在へと変えようとした試みであった。彼はプリンストン大学で研究を続け、OpenAIの研究者として働くなかで、人工知能が次の単語を当てることにとどまらず、デジタル環境のなかで自ら行動する方向へ進むべきだという問題意識を論文として具体化した。その成果がReActとTree of Thoughtsである。この二つは今なお言語エージェント研究を語るうえで欠かせず引用され続ける枠組みとして残っている。

ReActは「ReAct: Synergizing Reasoning and Acting in Language Models」という題名で発表された。ICLR 2023において上位5パーセントに入る口頭発表に選ばれ、エージェント研究の基本教材のような位置を占めるに至った。この論文が解こうとした問題は明確であった。それまで広く使われていた思考の連鎖(Chain-of-Thought)方式は、モデルに頭の中だけで思考を続けさせるものであった。反

対に、道具を直接操作する行動中心のモデルは、なぜそうした行動を取るのかを説明しなかった。考えはするが世界と接点を持たない、あるいは行動はするが理由を語らない、この二つの方式が別々に動いていたのである。

ReActはこの二つを一つの循環の輪として結び合わせた。モデルがまず思考(Thought)を文章として書き出し、その思考に基づいて行動(Action)を実行し、その結果として返ってきた観察(Observation)を再び次の思考の材料とする方式である。この三段階を繰り返させると、モデルは自分がなぜこの検索を行い、なぜこのクリックをするのかを、段階ごとに自然言語として書き残すようになり、人間はその経路をたどりながらどこで誤りが始まったのかを突き止められるようになる。

この構造の効果を確かめるため、姚順雨博士と共同研究者たちは実際の環境を二つ用意した。一つはウィキペディアの文書を行き来しながら答えを見つけなければならない質問応答課題、HotpotQAである。もう一つはウェブサイトを巡って商品を検索し、購入まで完了させなければならないWebShopであった。いずれの環境も、モデルがただ答えを吐き出すのではなく、複数の段階を経ながら情

報を集め、判断を修正していかなければならない課題であった。ReAct方式を適用したモデルは、思考を外に表さない既存モデルに比べ、誤りから回復する能力が目に見えて向上した。

同時期に姚順雨博士が打ち出したもう一つの柱がTree of Thoughts、略してToTである。論文題名は「Tree of Thoughts: Deliberate Problem Solving with Large Language Models」であった。NeurIPS 2023で口頭発表として採択された。ReActが思考と行動を一つに結びつけることに主眼を置いたのに対し、ToTは思考そのものをどう幅広く展開させるかという問いに答えた研究であった。

従来の言語モデルは、左から右へ文章を一行に押し出しながら答えにたどり着いていました。一度誤った方向に踏み出すと、後戻りする方法がなかったのです。ToTはこの方式を、人が難しい問題を解くときの行動に近づけました。人は一つの解法だけを押し進めることはしません。複数の種類のアプローチを同時に思い描き、行き詰まった道だと感じれば前の段階に戻って別の道を試みます。ToTはこの過程を、コンピュータ科学で長く使われてきた木構造探索法である幅優先探索と深さ優先探索で実装しました。

具体的には、問題解決の各段階を複数の思考の枝に分けて広げ、モデル自身に各枝が正解に近いかどうかを評価させます。この自己評価スコアをもとに、可能性の低い枝は切り捨て、有望な枝だけを残して次の段階に進みます。もしある枝が行き止まりに達したと判断されれば、モデルは演算を以前の状態に戻すバックトラッキングを実行し、別の枝を再び探索します。こうした構造のおかげでToTは、計算量の多い数学パズルや幾何学の問題のように、一度では答えが出ない課題で強みを見せました。このコードはプリンストン公式リポジトリを通じて完全に公開されており、誰でも自由に利用できます。

姚順雨博士の研究歴を見ると、彼は自身の個人ホームページでも学業と論文の経歴を詳しく整理しています。このホームページによれば、彼はReActやToTのような学術研究にとどまりませんでした。OpenAIで働きながら、これらの概念を実際の製品へと落とし込む作業にも関わりました。代表的なものがComputer-Using Agent、略してCUAと呼ばれる技術です。これは人がキーボードとマウスでコンピュータを操作するように、人工知能が自らオペレーティングシステムとウェブブラウザの画面を認識し操作できるよう助ける技術の原型でした。

もう一つはDeep Researchという製品でした。ユーザーが複雑な調査テーマを入力すると、エージェントがインターネット上の複数の資料を自ら探し、情報を集め、異なる出典を比較しながら一つの報告書としてまとめる方式です。ReActで築かれた思考と行動の循環構造が、ウェブ検索と文書読解を行き来する実際の調査過程へと拡張された形でした。

エージェントを作ることと同じくらい重要なのが、そのエージェントが実際にどれほど使い物になるかを測る物差しを作ることでした。姚順雨博士はこの課題にも取り組み、その成果がtau-benchと呼ばれる評価体系です。このベンチマークは、ツールとエージェント、そして実際の人間ユーザーの間で何度もやり取りが交わされる商業的業務環境を模し、エージェントが一度きりの回答ではなく、長い対話と交渉の過程でも仕事をきちんと処理できるかを測定するよう設計されています。tau-benchは2025年のICLRで正式に採択されました。

姚順雨博士が歩んできた道は、一つの問いによって貫かれています。言語モデルが考えるだけの存在から抜け出し、実際に行動するためには何が必要かという問いです。ReActは思考と行動を結ぶ最

小単位の答えでした。ToTはその思考を広く深く展開する方法でした。CUAとDeep Researchは、その思考と行動を実際のコンピュータ画面とインターネット上で具現化した成果物であり、tau-benchはその成果物が実際にどれほど信頼できるかを測る道具でした。

こうした経歴を積み重ねてきた研究者が中国市場へと舞台を移した時点は、2025年12月でした。証券時報の報道によれば、テンセントはこの時期に大規模モデル研究開発組織を大幅に見直し、AIインフラ部、AIデータ部、データコンピューティングプラットフォーム部を新たに設けました。そしてこの改編の中心に、元OpenAIシニア研究員であった姚順雨を据えました。

彼が任された職務は一つではなく複数でした。まず、CEO兼総裁室所属の首席AI科学者として、テンセント総裁の劉熾平に直接報告する立場に立ちました。同時に、AIインフラ部と大規模言語モデル部を率いる責任者として盧山に報告する役割も兼ねました。一つの肩書きではなく、組織構造全体を横断する複数の職責を一人に重ねて任せた形でした。これはテンセントがこの人事を単なる採用ではなく、組織改編の中心軸に据えたことを意味していました。

証券時報の報道は馬化騰の発言も併せて伝えました。馬化騰は、テンセントHunyuanがこの一年間、人材採用と組織構造の両面で大きな変化を経て、その過程でネイティブAI人材をより多く取り込んだと明らかにしました。この発言でHunyuanという名前が出てくるのは偶然ではありません。テンセントHunyuanはこの会社の大規模モデルブランドです。姚順雨の合流は、まさにこのHunyuanモデル開発組織の再編と結びついていたからです。

報道はさらに、この一年間でテンセントHunyuan大規模モデルが大きな再構築を経たと伝え、姚順雨が合流して以降、会社が人材採用を加速させ、研究開発チームを再編したと述べました。社内的には共同設計を加速させ、Hunyuan大規模モデルと対話型サービス「元宝(Yuanbao)」の間の相乗効果を高めているとも記しました。この部分は、彼の役割が一つのモデルを手直しするだけにとどまらず、モデルとそのモデルを使うサービスとの間の設計を共に調整するところまで広がっていることを示しています。

こうした措置は、テンセントのAI戦略が加速しているという兆候として解釈されました。姚順雨はこの報道の中で、テンセントHunyuan大規模モデル研究開発の中核責任者として位置づけられまし

た。プリンストンで思考と行動をつなぐ枠組みを作り、OpenAIでその枠組みを製品へと移した研究者が、今や中国最大級の技術企業の一つの大規模モデル組織全体を設計する立場に立つことになったのです。

ここで押さえておくべき点があります。テンセントHunyuan研究チームが発表した代表モデルであるHunyuan-LargeとHunyuan-TurboSの技術報告書は、それぞれ2024年と2025年に公開されましたが、両報告書の主要貢献者名簿に姚順雨の名前は挙がっていません。彼がテンセントに合流した時期と組織改編報道の時期を踏まえると、この二つのモデルは彼が着任する以前から進められてきたテンセントHunyuan研究チームの成果であった可能性が高く、彼がその後担うことになった役割は、こうした既存のモデル群を包括する研究開発組織全体を指揮する立場に近いと見るべきでしょう。

それでも、この二つのモデルがどのような方向で設計されたかを押さえておくことは、彼が今何を引き継いで率いることになったのかを理解する助けになります。Hunyuan-Largeは2024年11月に公開されたモデルで、全体のパラメータ数は3890億個に達しますが、実際に一つの質問に答える際に活性化されるパラメータは520億

個程度に抑える専門家混合構造を採用しました。この設計は、モデル全体の知識は膨大に蓄えつつも、その都度実際に動く演算量は減らし、サービスコストを下げる狙いを込めています。

Hunyuan-TurboSはそこからさらに一歩進んだモデルです。このモデルは2025年に公開されました。既存のトランスフォーマー構造にMambaと呼ばれる系列の構造を組み合わせたハイブリッド設計を採っています。トランスフォーマーは文脈を幅広く理解するのに強い一方、計算量が文章の長さに比例して増えるという欠点があり、Mamba系列の構造はその負担を大きく減らす代わりに文脈の扱い方が異なります。二つの構造を組み合わせたのは、長い文書を扱いながらもサービスコストを抑えられるようにするための折衷案でした。

Hunyuan-TurboSで注目すべきもう一つの設計は、質問の難易度に応じて思考の深さを変える方式でした。天気を尋ねるような簡単な質問には、モデルがあえて長い思考過程を経ずに素早く答えるようにし、数学の問題や複雑なコードの誤りのように複数の段階を経る必要がある質問には、自ら計画を立て、途中で自分の判断を振り返る過程を経るように分けたのです。これは、先にReActとToTが

扱った問題、すなわち言語モデルがいつ思考を展開し、いつ即座に行動すべきかという問いと同じ根を持つ設計と見ることができます。

姚順雨博士がテンセントで任された立場は、こうした既存モデルの設計思想を引き継ぎながら、同時に彼がプリンストンとOpenAIで培ってきた言語エージェントに対する視点を、テンセントの研究開発体系の中に取り込む位置にあると言えます。組織図の上ではCEO兼総裁室所属の首席AI科学者でありながら、AIインフラ部と大規模言語モデル部を共に率いる立場です。実質的には、テンセントが自社の大規模モデルとそれを使うサービスを一つにまとめて設計しようとする試みの中心に立つ人物だと言えるでしょう。

プリンストンの研究室で次の単語を当てるモデルを超えようとした問いは、今やテンセントという巨大な組織の研究開発体系を作り直す作業へとつながっています。ReActが打ち立てた思考と行動の循環の輪、ToTが打ち立てた複数の種類を同時に見渡し後戻りする探索方式は、一本の論文の中にとどまりませんでした。彼が今身を置く組織が次のモデルを設計する際に参照する下絵として残っているものと見られます。

現在の立場と中国AIにおける意味

姚順雨は2025年12月にテンセントに合流した後、「CEO/総裁室」所属の首席AI科学者に任命され、テンセント総裁の劉熾平に直接報告する職責を受けました。同時に、AIインフラ部と大規模言語モデル部の責任者を兼任し、技術工程事業群を率いる盧山にも報告します。彼が組織の中でどのような職責を受けたかは前節ですでに触れたので、この職責が彼自身に、そして中国AI業界全体にどのような意味を持つのかを見ていきます。

1998年生まれの彼がこうした職責を受けるに至った背景には、清華大学姚班を経てプリンストン大学でコンピュータ科学の博士号を取得し、その後OpenAIの研究者として働いた経歴があります。27歳という年齢で中国最大のプラットフォーム企業の一つの大規模モデル研究開発の方向を定める立場に立つことになったのは、彼が大学院時代に書いた論文が実際に業界標準として定着したという事実なしには説明が難しいでしょう。

彼がテンセントに合流する数か月前の2025年4月、姚順雨は「AIの後半戦(The Second Half)」という文章を発表し、業界に少なからぬ波紋を広げました。彼はこの文章の中で、人工知能が囲碁、SAT、数学オリンピックのような公認のベンチマークですでに人間

の成績を上回っているにもかかわらず、それに見合う生産性の変化がなぜ現実世界では起きていないのかを問いました。

彼はこの問いを、人工知能が直面する重い課題として挙げました。AGIに近づいたとベンチマーク突破を祝うよりも、すでに高得点を得たAIが実際にどのような価値を生み出しているのかを問い直すべきだという趣旨でした。得点を上げることと、人に役立つことを成し遂げることは別問題だという指摘です。この文章以降、業界ではこれをしばしば「効用問題」と呼んでいます。

この指摘が向けられたのは、研究者たちの古い習慣でした。事前学習と事後学習を終えたモデルをどこに適用し、どのような価値に変えるかを考える作業は、論文を書きランキング表に名前を載せる作業とは異なる筋肉を要求します。姚順雨は、研究者自身が問題定義と指標設定、そして反復的な改善を導くプロダクトマネージャーの姿勢を身につけるべきだと考えました。

この考えは、彼がテンセントで担う実務とぴったり重なります。彼はテンセントで、事後学習の精鋭人材を「元宝(Yuanbao)」チームに投入し、コストパフォーマンスを性能とコストの二段階に分けて説明しました。ランキング表で一位を取れなくとも、ユーザーが

実際に使いやすいモデルが市場で勝つという彼の説明は、大学院時代から培ってきた問題意識が企業経営の言葉へと移った姿です。

こうした問題意識を裏づけるのは、彼がプリンストンとOpenAIで残した実績です。ReActは推論と行動を一つの循環の輪として結びつけ、言語エージェントを作る方法論の出発点となりました。Tree of Thoughtsは、その思考過程を複数の方向へ広げて探索させる枠組みを提示しました。二つの論文は、今やウェブを検索し、ツールを呼び出し、複数の段階を経て問題を解くエージェントの大半が拠って立つ土台となっています。これは中国企業一社だけでなく、世界中の研究者と企業が共通して従う規格に近いものです。

MITテクノロジーレビューはこうした成果を認め、彼を「35歳以下のイノベーター」リストに選出しました。言語エージェントのパラダイムを打ち立てたReAct、そしてWebShopやSWE-benchのような実践的な評価環境を作り上げた功績が選出理由として挙げられました。審査評では、彼が打ち立てたこうした成果がOpenAIのOperatorやDeep Researchのような実際の製品へとつながり、人が使えるエージェントへ近づく道を切り開いたと記されています。

彼が個人ホームページにまとめている研究リストを見ると、言語エージェントをどのような基準で評価するかという問いが、ReActやTree of Thoughtsに劣らず一貫して続いていたことが分かります。ツールを呼び出し、ユーザーと対話しながら任務を完遂する能力を測るtau-bench、そして言語エージェントを記憶・計画・実行という要素に分けて整理した「言語エージェントのための認知アーキテクチャ」論文がその例です。ベンチマークを作ること自体が、彼にとって常に研究の半分を占める作業でした。

目を引く事実の一つは、2025年2月を最後に更新された彼の個人ホームページが、いまだに自身をOpenAI所属の研究者として紹介している点です。テンセント入りを報じた記事も、彼のホームページがまだ新しい役職を反映していないという事実を別途指摘しました。世界各地で引用される研究者の公開履歴と、中国大手企業の組織図の中に新たに描かれた彼の実際の立場との間に時差があるということです。この時差そのものが、彼の転職がいかに速く進んだかを示しています。

学界と産業界の評価が国籍を問わず彼に向けられたという事実は、彼が打ち立てた方法論が特定企業の資産ではなく、業界全体が共

有するツールとなったことを意味します。テンセントが彼を迎え入れたことは、一人の人材を招いたことであると同時に、世界がすでに標準として受け入れた方法論を作った人物を自社の組織内に取り込んだことでもあります。

こうした国際的な評価を得た研究者が中国企業へ転職した時期も注目に値します。米国が先端半導体や演算機器の輸出規制を強め続ける状況の中、シリコンバレーの研究所の中核研究員が自国の大手プラットフォーム企業へ移った事例は、中国AI業界内で人材が回帰する流れを示す事例として受け止められています。

こうした流れを、性急に国家間競争の物語としてのみ読み取る必要はありません。ただし確認できる事実ははっきりしています。プリンストンで学位を取得し、OpenAIでキャリアを積んだ27歳の研究員が、わずか数か月の間に自国最大のプラットフォーム企業の研究開発の方向性を決める立場へと移ってきたという事実そのものが、人材の移動経路が一方向にのみ流れるわけではないことを示しています。

ヤオ・シュンユー以前にも、シリコンバレーで名を上げた研究者が中国企業へ転職した事例はありましたが、彼のように若い年齢で

一気に大企業の研究開発組織全体を任される役職に就いた例は珍しいものです。彼の移籍が複数のメディアで繰り返し報じられた理由も、この人事が一個人の転職を超えて、中国の大規模モデル業界が海外で実績を認められた研究者をどのように処遇するかを示す事例として読み取られたからです。

プリンストン時代に彼が取り組んだもう一つの軸は、言語エージェントをどう測るかという問いでした。この問いは、テンセントで彼が担う実務に直結しています。彼はツールを呼び出し、ユーザーと対話しながら任務を最後まで完遂する能力を測るtau-benchを共同で設計しました。言語エージェントを記憶・計画・実行に分けて整理した論文も共同執筆しています。ベンチマークを新たに組み立てること自体を、彼は研究の半分だと考えてきました。

このような経歴を持つ人物がAIインフラ部と大規模言語モデル部を兼任する役職に就いたということは、テンセント内でモデルをどう評価し、どこに配置するかを決める権限が、評価方法論を長年研究してきた人物の手に握られたことをも意味します。彼がテンセントクラウドAI産業応用大会でテストも前倒しで設計すべきだと述べたことも、こうした経歴が実務の言葉に置き換えられた例と見るこ

とができます。

テンセントが彼の入社に合わせてAIインフラ部、AIデータ部、データコンピューティングプラットフォーム部を新設し、既存の大規模言語モデル副総経理を彼の下に再配置した経緯は、これまでに触れました。ここで指摘しておきたいのは、この改編が一個人の招聘にとどまらなかったという点です。テンセントが大規模モデル競争を会社の最優先課題に据えたというシグナルとして、業界で受け止められたということです。

ただし、この役職が持つ重みを正確に見極めるには、彼が何を新たに作り、何を引き継いだのかを区別する必要があります。テンセントが2024年11月に発表したHunyuan-Largeと、2025年5月に発表したHunyuan-TurboSの技術報告書に記載された主要貢献者リストには、ヤオ・シュンユーの名前は載っていません。

この事実は、彼がこの二つのモデルを直接設計したり訓練させたりしたわけではないことを意味します。彼がテンセントに入社した時期は2025年12月です。二つのモデルはそれより前にすでに開発・公開された状態でした。彼が引き継いだのは白紙ではなく、すでに稼働していた大規模モデル研究開発体制でした。彼が担った実務

は、その体制を再編成し、新たな指揮系統を組み込むことにありました。

こうした区別は、彼の役割を貶めるためのものではなく、彼がテンセントで実際に担っている仕事があるかを正確に描き出すためのものです。彼はHunyuanという名を最初に生み出した人物ではなく、すでに存在していたHunyuan研究開発組織に新たな原則と指揮体制を据え、次の段階へ押し進める役割を担いました。

テンセントクラウドAI産業応用大会で彼が示した発言も、こうした役割と重なります。彼はテンセントがAI競争で後れを取っているという指摘に一線を画し、競争はようやく後半戦に入ったばかりだと述べました。ChatGPTとClaudeだけが唯一の答えにはならないだろうという彼の説明は、自身が移ってきた組織の立ち位置を彼がどう理解しているかを表しています。

彼はコーディングエージェントやマルチモーダル、身体性を持つ知能のように、まだ満たされていない市場が残っていると見ており、こうした市場こそがテンセントのような後発企業にも扉が開かれている理由だと説明しました。現在の状況を1970年代、パーソナルコンピューターが普及し始めたばかりの時期になぞらえた彼の言

葉は、彼が競争の速さよりも競争が今後どれほど長く続くかをより重く見ていることを示すものと読み取れます。

テンセントが今回の組織改編と人材招聘に実際いくらを支払ったのか。こうした数値は、彼の役割が向かう方向を判断するうえで必ずしも必要な情報ではないため、

ヤオ・シュンユーがプリンストンで身につけた学問的な習慣、すなわちどんな研究であれ論文として残し、失敗した実験さえも記録に整理しておく習慣は、彼が今身を置く企業組織においても大きく変わらず続いています。彼はテンセントクラウドAI産業応用大会で、大規模モデルやAI製品を作る初期段階では試行錯誤と後退が自然なことだと認めました。不足している点を率直に明かし、調整する際に忍耐力を持つ姿勢が次の段階の要だと述べました。

研究室で失敗を隠さず次の実験へと進んでいた姿勢が、大企業の経営陣の話法へと移し替えられた様子は、彼がテンセントで作ろうとしている組織の性格をうかがわせます。成果を性急に急き立てる組織ではなく、失敗を認めて再挑戦する研究室のやり方を大企業の中に移植しようとする試みと見ることができます。

中国AI産業の中で彼が持つ意味は、いくつかに分けて見ることができます。一つは、彼が打ち立てたReActとTree of Thoughtsが、国境を問わない業界標準として定着したという事実です。もう一つは、彼が米国で訓練を受けた研究者が中国企業へ移る流れの目立つ事例となったという事実です。

三つ目の意味は、彼がリーダーボード競争から実際の価値を生み出す競争へと重心を移そうと語った効用の問題が、テンセントという大組織の研究開発の方向性を決める原則として、すでに定着し始めているという事実にあります。彼はその問題を論文ではなく、組織改編と人事配置によって現実に移そうとしています。

論文で方法論を打ち立てた大学院生が、数年も経たないうちに中国を代表する技術企業の研究開発の方向性を決める立場に就いた経緯は、一個人の経歴を超えて、中国のビッグテックが海外で実績を認められた研究者をどのような立場に迎え入れるかを示す事例でもあります。ヤオ・シュンユーがテンセントで担う役職は、そうした流れの上に置かれています。

彼が大学院時代から持ち続けてきた効用という問いは、今や一つの会社の研究開発全体を測る物差しとして使われています。ベンチ

マークのスコアではなく、実際に人が使う価値を基準にしようという彼の問題意識が、中国を代表するプラットフォーム企業の一つの大規模モデル戦略を導く原則となりました。彼がテンセントで今後どのような成果を出すにせよ、この出発点を振り返れば、彼が歩んできた道がなぜこの役職へとつながったのかを推し量ることができるでしょう。

彼の現在の立ち位置は、三つの層が重なり合ったものです。大学院生時代に打ち立てた方法論が、国境のない業界標準となっている状態。米国で名を上げた研究者が、自国企業へと移ってきた状態。そして、リーダーボードよりも価値を優先しようという長年の問題意識が、テンセントの大規模モデル戦略の原則として再び現れた状態です。この三つの層はそれぞれ異なる時点に始まりましたが、今はヤオ・シュンユーという一人の人物の役職の中で重なり合いながら動いています。

ただし、彼が打ち立てた研究方法論がすでに世界的に実証されているという事実。

第10章 呉永輝(吴永辉、Yonghui Wu)、Googleの翻訳研究からByteDance Seedの基盤モデル統括者へ

生い立ちと学業

呉永輝という名前は、中国語では吳永輝、英語表記ではYonghui Wuと書きます。彼はシステムエンジニアリングとディープラーニングの両方を手がける数少ない研究者の一人として知られ、検索エンジンの基盤インフラを作っていた人物が、のちに翻訳や画像生成モデルまで設計するようになった珍しい経歴を持っています。

大学で何を専攻していたかについても同様です。実際に記録が残っている時点から話を始めるのが妥当でしょう。

記録がはっきりと始まる時点は2008年です。呉永輝はこの年に博士号を取得し、すぐにGoogleにソフトウェアエンジニアとして入社しました。

Google入社後の最初の7年間、彼は検索順位を決定するエンジニアリング組織で働きました。人々がGoogleの検索窓に語句を入力したとき、どの結果をどれほど速く、どれほど正確な順序で表示するかを計算する仕事でした。世界中から押し寄せる検索リクエスト

を遅延なく処理し、順位付けの計算を大規模サーバー上で並列に走らせる作業でした。

この時期の経験は、のちに彼がByteDanceで巨大言語モデルを訓練する際に再び生かされることになります。数千億、さらには1兆に達するパラメータを持つモデルを安定して訓練するには、結局のところ大規模分散コンピューティングを扱う感覚が必要であり、その感覚の根はこの検索エンジンを作っていた7年間にあったと見ることができます。

Googleでの彼の経歴が大きく変わる契機となったのは、ディープラーニングの台頭でした。彼は検索ランキングのエンジニアリングを離れ、Google Brainへと異動しました。Google Brainは当時Google内で人工知能研究を牽引していた組織であり、ここで彼はエンジニアから研究者に近い役割へと移っていきました。

Google Brain時代に彼が残した最大の成果は機械翻訳の分野にあります。彼はGoogleのニューラル機械翻訳システム、通称GNMTと呼ばれるプロジェクトを率いた研究者の一人でした。このシステムは2016年に発表され、翻訳を単語ごとに置き換える方式ではなく、文全体の意味を一つのベクトルに圧縮してから再び展開する方式

で処理しました。

それまで機械翻訳は、句単位で対応させる統計的手法が中心でした。文全体を丸ごと読み取って意味を把握し、再び文として展開するニューラルネットワーク方式は、翻訳の自然さを目に見えて向上させ、Google翻訳サービスに実際に適用されて世界中のユーザーがその変化を実感しました。以降、世界の自然言語処理研究はこの方式を標準として発展していきました。

翻訳の次に彼が名を連ねた分野は音声でした。彼はConformerという名のニューラルネットワーク構造を共同設計した研究者の一人として名を連ねました。Conformerは、Transformerが持つ文全体を広く見渡す能力と、畳み込みニューラルネットワークが持つ近接した音を細かく捉える能力を一つに統合した構造でした。この構造は2020年頃に発表されたのち、音声認識分野で広く使われる基本的な枠組みの一つとして定着しました。

翻訳と音声を経て、彼の関心は画像と言語を共に扱う方向へと移っていきました。彼はCoCaという名のマルチモーダルモデルの研究にも参加しました。CoCaは、写真とその写真を説明する文をペアとして置き、両者の関係を一つのモデルの中で同時に学習させる

構造でした。このモデルは2022年頃に公開されました。

この画像と言語を共に扱う経験は、のちに彼がByteDanceでSeed reamという画像生成モデル群を統括することになった際に再び役立ったものと見られます。写真を説明するモデルと写真を生成するモデルは方向こそ逆ですが、画像と言語がどのように対応し合うかを理解する感覚は、両方の作業に共通して必要とされるからです。

翻訳と音声と画像を横断したこの経歴のおかげで、彼はGoogle内で特定の一分野の専門家としてではなく、複数の分野を横断してインフラとアルゴリズムを同時に扱うことのできる数少ない研究者として認められるようになりました。

こうした成果が積み重なり、彼は2023年9月にGoogle Fellowに任命されました。Google Fellowは、Google内で技術職として到達できる最高位のポストの一つであり、社内でもごく少数の人員しか授与されない称号です。同時期、彼はGoogle DeepMindの研究担当バイスプレジデントのポストも兼任することになりました。

Google BrainとGoogle DeepMindが一つの組織に統合される過程で、彼はGeminiという名の巨大モデルを開発する作業にも参加しました。検索エンジニアとして出発し、翻訳、音声、画像を経て、Googleが作った最大のモデルの開発にまでつながったわけですから、17年という歳月の間に彼が歩んだ道は、Googleの人工知能研究が歩んだ道とかなりの部分で重なっていると言えるでしょう。

こうしてGoogleで築いたキャリアの頂点にあった彼が、2025年2月にByteDanceへと移籍しました。彼はByteDanceの巨大モデル研究組織であるSeedで、基盤研究を統括する責任者として着任しました。最高経営責任者である梁汝波に直接報告する立場を与えられました。

彼がSeedに着任した当時、Seedは厳しい状況に置かれていました。数千人規模の研究組織が数百億元を投じて2年近く追いつけた末、中国国内で上位に入る基盤モデルを作り上げたものの、はるかに少ない人員と資源で動いた他のチームにほどなく追い抜かれる事態が起きていました。組織の責任者はこの問題を認め、会社の最高経営責任者も全社会議で、もっとうまくできたはずだという趣旨の指摘をしたことがあります。

Seedの複数の関係者は、彼に寄せた期待を二つにまとめています。モデルの能力を中国国内で第一位に押し上げること、そして世界最上位圏のモデル企業と肩を並べて競える水準まで引き上げることでした。同時に、彼と接した人々が共通して挙げる印象は「落ち着き」でした。

彼は着任後、100人余りに及ぶ中核研究員たちと密度の濃い一対一の面談を行ったと伝えられています。長年組織を率いてきた人物らしく、新しい組織を任された後にまず行ったのは、その組織の中の人々を一人ひとり会ってみることだったようです。

彼を間近で見てきた人々は、彼がGoogle出身というよりも、むしろByteDanceの中で長く育った人物に近いという印象を受けたと語ります。実用的でありながらロマンチストでもあるという評価も付きます。技術への理解が深く、どの方向から成果が出せるかを素早く見抜くという評もあります。

彼がSeedで立てた目標は大きく二つにまとめられます。一つは、基盤モデルの能力を引き上げ、研究効率を高めて決められた期限内に成果物を出すことです。もう一つは、研究を重視する雰囲気を作り、最高水準の研究組織へと成長させることでした。この二つの目

標の間でバランスを取ることは、彼が着任して以来ずっと抱えている課題でもあります。

ここで押さえておくべきは、検索エンジニアとして出発し、翻訳、音声、画像を一通り経てGoogle Fellowという地位にまで上り詰めたのち、新しい組織へと移っていったこの軌跡そのものです。

一方で、2008年の博士号取得以降の経歴です。

この17年間のGoogleでの経歴を丸ごと眺めてみると、彼は特定の一本の論文で名を知られた研究者ではなく、検索という最も基礎的なインフラから翻訳、音声、画像という複数の応用分野を順に経て実力を積み上げてきた人物だと言えます。こうした幅広い経験が、彼がByteDanceで複数方向のモデルを同時に統括する立場を任されるようになった背景と見られます。

研究者、あるいは起業家として定まっていく時期

中国の人工知能を率いる人物たちの中でも、呉永輝ほど静かに、しかし着実に自らの地歩を広げてきた人物は珍しい。彼は2008年に計算機科学の博士号を取得した直後、Googleにソフトウェアエンジニアとして入社した。華やかな起業ストーリーも、劇的な抜擢

のドラマもなく、彼はGoogleという巨大な組織の中で最も基礎的な仕事から始めたのである。

入社後の最初の7年間、彼が担当したのは検索順位を決めるエンジニアリング、いわゆるサーチランキングだった。表向きは地味な業務だったが、Googleという会社の心臓部にあたる仕事だった。一日に何十億回も発生する検索結果の順序を決めるシステムを扱う仕事であり、この時期に大規模分散システムを扱う感覚が身についたものと見られる。

7年という歳月は短くない。一人の人間が同じ場所でそれほど長く働くということは、華やかな転職ではなく、一つの問題を最後まで掘り下げる性向を示していると言える。後に彼の同僚たちが彼について最も強く思い出す印象が「落ち着き」だったという事実も、この時期の習慣と無関係ではなさそうである。

その後、彼はGoogle Brainへと異動する。Google Brainは人工知能研究を専門に担う組織であり、ディープラーニングという新しい方法論が実験室の外に出て実際のサービスへと浸透していった時期の真ただ中であつた。呉永輝はここでディープラーニングを翻訳分野に応用する研究を率い、その成果は検索順位アルゴリズムにも

反映された。

翻訳という問題は、人工知能研究において長らく厄介な課題だった。一つの文を別の言語に移す作業は、単語を一つずつ置き換える算数ではなく、文脈と語順とニュアンスをまるごと理解しなければならない作業だからである。ディープラーニング以前の翻訳プログラムは、規則と統計を幾重にも積み重ねて作られた精巧な装置に近く、文章が少しでも長くなると不自然になりがちだった。

呉永輝が携わった仕事は、こうした翻訳方式を一つのニューラルネットワークにまとめて学習させる方向へと変える作業だった。文をまるごと読み、まるごと訳す方式が定着するにつれ、翻訳の品質はそれまでとは異なる次元へと引き上げられた。この時期の経験は、彼が後に扱うことになる大規模言語モデルの土台となる感覚、すなわち膨大なデータを一つのニューラルネットワークに消化させる感覚を育んだものと見られる。

検索から翻訳へ。そして翻訳から再び検索アルゴリズムへと成果が還元されていく循環は、彼が一つの技術を狭く掘り下げる人間ではなく、異なる問題の間から共通の原理を見いだす人間であったことを示している。Googleという会社が検索と翻訳という、一見異

なる二つの事業を同じ人工知能技術で結びつけることができた背景には、こうした研究者たちの仕事があった。

このように十数年をGoogle社内で過ごした末、2023年9月、彼は「Google Fellow」という地位に就く。Google Fellowは社内でも数えるほどしかいない研究者にのみ与えられる地位であり、長年積み重ねてきた技術的成果を会社自らが公認する手続きに近い。同時に彼はGoogle DeepMindの研究担当副社長を務め、Geminiの開発と初期の追撃戦に加わった。

Geminiは、Googleが競合他社の大規模言語モデルを追いかけて作り上げた統合人工知能モデルである。検索順位エンジニアとして出発し、翻訳研究を経た彼が、今や会社全体の人工知能戦略がかかったモデルの研究を指揮する立場に立ったのである。17年近くにわたって一つの会社の中で歩んできた軌跡を見れば、この地位は突然与えられたものではなく、着実に積み上げてきた結果に近い。

ところが2025年2月、彼はこの地位を離れ、ByteDanceへと移る。ByteDanceの大規模モデル研究組織であるSeedの基礎研究責任者として合流したのである。欧米メディアはこの移籍を「Googleのベテラン人工知能研究者がByteDanceに合流した」というニュー

スとして短く扱ったが、中国国内ではそれよりもはるかに重い期待が彼を待ち受けていた。

彼が合流した当時、Seedの状況は厳しかった。数千人規模の研究チームと数百億元に及ぶ投資を投じて2年間追いかけた末、中国国内で第一級と評される基盤モデルを作り上げたものの、それよりもはるかに少ない人員と資源で動いた他チームのモデルにあっという間に追い抜かれた状態だった。部門責任者たちはこの過ちを認め、最高経営責任者は会議の場で直接悔しさを表明した。

こうした状況の中、呉永輝に与えられた任務は基礎的な技術改善ではなかった。Seedの内外の関係者は、彼がモデルの能力を国内最高水準に引き上げ、世界トップクラスの企業と肩を並べる組織を作らねばならないという期待を背負っていたと語る。長年Googleという安定した組織で働いてきた研究者が、一夜にして業績プレッシャーがはるかに切迫した立場へと移った格好である。

合流直後、彼が選んだ方法は華やかな宣言ではなく、人と直接会うことだった。彼は中核研究員100人余りと一人ひとり面談を行い、この過程を経てモデル構造研究を率いる人材を新たに抜擢した。検索順位エンジニア時代から身についた落ち着いた態度が、見知ら

ぬ組織に初めて足を踏み入れた瞬間にもそのまま表れたようである。

彼を間近で見てきた人々は、彼がGoogle出身というよりまるでByteDanceの中で育った人物のように見えると語った。技術への理解が深く、どちらの方向から成果が出るかを素早く見極めるという評価も続いた。長年にわたり様々な問題を横断しながら培った感覚が、新しい組織への適応の速さにもそのまま表れたものと見られる。

ByteDanceに移って以降、彼が実際に手がけた組織再編や技術的負債の整理、Seed1.5やSeeDreamといったモデルの具体的な成果については、この章の別の節で詳しく扱う内容である。ただ、ここでははっきりと押さえておくべき事実は、彼がこの地位に至るまでたどった道が、検索エンジニアから翻訳研究者へ、そして再び大規模モデル研究担当副社長へと続く一つの長い流れだったという点である。

振り返ってみれば、彼の経歴はある一瞬の飛躍ではなく、一つの問題を最後まで捕まえて次の問題へと移っていく反復の連続だった。検索順位という問題から翻訳という問題へ。翻訳という問題から再び統合人工知能モデルという問題へと移りながらも、彼は常にそ

の問題の底に横たわる原理を見いだす方法で仕事をしてきた。

創業者として華やかに脚光を浴びる他の人物たちと比べると、呉永輝の経歴は際立って地味である。しかし組織の中で研究者としての信頼を積み重ねていく道もまた、会社を興す道に劣らず重要な道であるという事実を、彼の経歴は示している。彼は会社を興した人間ではなく、会社が困難にぶつかるたびに呼ばれて問題を整理する人間として、自らの地位を固めてきたようである。

彼がGoogleで積み重ねた17年という歳月は、単に長く勤めたという事実以上の意味を持つ。検索と翻訳という異なる問題を行き来しながら積んだ経験。そしてその経験を再び大規模モデルというより大きな問題に適用してみせた能力は、彼がSeedに合流した後に起きた組織再編と技術的成果の土台となった。

呉永輝という名前が馴染みのない読者であっても、彼の経歴を一つずつたどっていくと、見覚えのあるパターンに気づくだろう。検索という基礎技術から出発し、翻訳という応用問題を経て、最終的には会社全体の未来がかかった大規模モデル研究にまで到達する道である。こうした道を歩んできた研究者が中国の人工知能産業へと身移したという事実は、それ自体、この産業がどのような人材を

求めているのかを物語っている。

会社と研究組織を作り上げた過程

2025年初め、ByteDanceの大規模モデル研究部門Seedは苦境に立たされていた。数千人の研究人員と数百億元を投じて2年間で中国国内でも指折りの基盤モデルを作り上げたものの、それよりはるかに少ない人員と資源で動いていたチームにあっという間に追い抜かれた。部門責任者は会議の場で過ちを認め、ByteDance CEOの梁汝波は全社会議でもっとうまくやれたはずだと直接指摘した。呉永輝がSeedを任されたのは、まさにこうした雰囲気の中でのことだった。

呉永輝は見知らぬ人物ではなかった。2008年に博士号を取得した後Googleにソフトウェアエンジニアとして入社し、最初の7年間は検索ランキングエンジニアリングを担当した。その後Google Brainに移り、ディープラーニングを翻訳と検索ランキングに応用する研究を率いた。2023年にはGoogle DeepMindの研究担当副社長となり、Geminiの開発初期の過程に加わった。17年をGoogleで過ごした彼が2025年2月にByteDanceへと移った事実は、業界でも大きなニュースとして扱われた。

Seedに合流した後、彼はただちに中核研究員100人余りと密度の高い一対一の面談を行った。この過程でモデルアーキテクチャを扱う研究員数名を新たに抜擢した。Seedの関係者が彼に対して抱いた第一印象は、華やかな弁舌よりもむしろ落ち着きだったと伝えられる。

呉永輝が1年間集中して取り組んだのは大きく二つのことだった。一つは基盤モデルの実際の能力を引き上げ、研究効率を高めて定められた日程どおりに結果を出すことだった。もう一つは研究を優先する雰囲気を作り、一流研究チームという評判を築くことだった。彼の側近は、彼がGoogleから来た人間というよりByteDanceの中で育った人間のように、実利的でありながらもロマンチックな態度を持っていると語った。

組織を組み直す作業は2025年1月から始まった。Seedは「Seed Edge」という仮想チームを作り、3年間の評価基準を与えた。四半期ごとの業績評価に追われることなく、より根本的で長期的な汎用人工知能の課題に取り組めるよう保護膜を張った格好である。呉永輝はこのチームを編成する過程に直接関与した。

次に作られたのがフォーカス(Focus)というチームでした。既存の部門の境界や業務分担を取り払い、次世代モデルを世に出すために必ず解決しなければならない工学的なボトルネックに集中する精鋭組織でした。残りの人員はベース(Base)としてまとめられ、エンジニアリングとデータ、評価を担当し、現在使われている基盤モデルの開発を続けました。シードエッジで得られた成果はフォーカスを経てベースへと下り、フォーカスで見つけた長期的な課題は再びエッジへと上がっていくというように、三つのチームが噛み合って回る仕組みに設計されていました。

2025年3月には指揮系統も整理されました。当時LLMチームを率いていたチャオムーが社内通報を受けて業務を停止されると、大規模言語モデル傘下の事前学習、事後学習、ホライズンの各チームがすべて呉永輝に直接報告する体制に変わりました。前任のシード責任者だったジュウォンジャは数か月間呉永輝とともに部門を率いた後、巨大モデルの実際の応用を担当し、彼に報告する立場へと異動しました。

呉永輝は効率を高めるため、社内のデータとコードのリポジトリを透明に共有するよう推し進めました。それまでは、あるグループ

の研究者が別のグループの文書を見るには、文書の所有者とその上司それぞれの同意を得なければなりませんでした。時には部門責任者が動いても、すべての情報を把握しきれないことがありました。情報の壁を低くしたことで意思疎通は確かに速くなりましたが、2025年下半期にはインターンが関わる情報漏洩事件が少なくとも2件発生しました。これを受けてシードは内部文書の公開を義務づける方針を打ち切り、権限体系を見直しました。

組織図そのものは大枠を維持しました。マルチモーダル of 相互作用とワールドモデルのチームは2024年にアリババから移ってきたジョウチャンが率い、Doubao(豆包)のスマートフォン向け秘書モデルがこのチームの代表作でした。ビジュアル・マルチモーダル生成責任者のヤン・ジェンチャオの休暇と、ビジュアル基盤モデル研究責任者のポン・ジャスの退社が重なったことで、ジョウチャンが担う範囲は広がり、画像生成モデルのSeeDreamと動画生成モデルのSeeDanceも彼の管轄に入りました。インフラチームはシャンリャンが率い、バイトダンスのAIラボが科学向け人工知能、ロボティクス、責任あるAIという三つの方向性をそのまま抱えたままシードに統合されたことで、責任者のリー・ハンも呉永輝に報告するようになりました。

人員規模は1500人前後からは大きく増やしませんでした。外部から中堅の技術管理職を採用するケースはほぼなくなり、その代わりに新卒や若手人材を抜擢する方向へと重心が移りました。2024年に清華大学を卒業した博士課程の学生がジョウチャンと呉永輝の双方に報告する例が、こうした変化を象徴しています。

この1年で最も明確な成果が、まもなく登場するDoubao 2.0でした。Geminiに似た性格を持つマルチモーダルモデルで、パラメータ数は1兆に達し、シードが作ったモデルの中で最大の規模でした。この訓練の過程でシードのチームは、インフラ段階での容易ならざる問題に直面しました。過去2年間、追い上げに没頭するあまり基礎的な力を固める作業を相対的に後回しにしていたため、パラメータ規模を拡大する段階で不安定な現象が現れ、しばらく開発が滞りました。ある研究者はこれを、基礎が固まっていない家にレンガを積み増せば問題が起きるようなものだと表現しました。

こうしたインフラの問題はシードだけの話ではありませんでした。OpenAIのRLインフラ責任者ウォン・ザイは、あるポッドキャストでモデル企業は結局のところインフラのバグを直す速さで競っていると語ったことがあります。OpenAI自身も3年間使ってきたイン

フラを組み直す作業に着手し、アリババのQwenチームやテンセントもそれぞれインフラ組織を新たに作り直しました。シードのインフラチームは数百人規模で数十種類のモデルを同時に支えており、これを丸ごと再整備する作業は人員と信頼の両方を大きく消耗させるものでした。ある関係者はこれを、走行中の車のタイヤを交換するようなものだと例えました。

Doubao 2.0が抱えた問題は、結局複数のチームが3か月間力を合わせてモデル構造と訓練データを直し続けたことで解決し、春節前に予定どおり発表するという目標を守ることができました。

呉永輝が着任してから、会議のやり方も変わりました。以前は3か月に一度集まっていた研究チームが、いまでは月に一度、中核チームは隔週で意思疎通するようになりました。彼はチームの報告を聞いた後、たいてい3、4個の簡単な質問を投げかけましたが、答えを直接教えるのではなく、研究者自身がより本質的な問題を考えるよう仕向けるやり方だったといいます。食堂でトレーを持ち、研究者の向かいに座って進捗を尋ねる彼の姿もよく見かけられました。

報酬と勤務環境も同時に変わりました。2025年半ば、バイトダンスは既存のストックオプションとは別に「Doubao仮想株」という

新しい報酬資産を作り、研究者たちに支給し、給与も何度か引き上げました。シードの中では、インターンでさえ経営トップと直接顔を合わせられるほど自由な雰囲気が生まれました。あるインターンは、2か月近く会社に泊まり込み、明け方に浮かんだアイデアをコードに落としてはまた眠りにつくことを繰り返したとSNSに書き残しています。

呉永輝は研究成果を外部に知らせることに力を入れました。2025年3月の部門全体会議で彼は、シードは良い成果を数多く出しているのに外部ではあまり知られていないと述べ、論文とブログで成果を発表するよう促しました。研究者たちには個人のホームページを整えておくよう勧めました。彼が加わってから3か月間に発表された論文数は、2024年の1年間全体の発表数を上回りました。シードを去ったある研究者は、この時期を経て勤勉で研究に没頭したい人たちが残る方向へとチームが自然に選別されていったと振り返っています。

自律性が高まった分だけ、資源配分をめぐる緊張も生まれました。シードの中には、一部の研究者が直属の上司の同意を得ないまま上位の管理職に直接報告して承認を取り付け、新しい方向性を推し

進めるケースがありました。ある関係者はこうして使われる資源が全体の2割ほどになるだろうと推測しつつも、バイトダンスは資源が不足している会社ではないので、こうした無駄も答えになり得ると語りました。しかし、シードは競争から完全に自由な組織ではありませんでした。テンセントやアリババといった競合他社に対抗して会社全体に弾薬を供給しなければならない以上、資源は結局のところ短期的な成果を出すチームへと偏らざるを得ませんでした。

2025年第3四半期からは、バイトダンスの経営陣が論文発表について品質を保証し、現在開発中の中核技術とは無関係であることという条件を新たに課したため、月ごとの発表数はやや減りました。ある研究者は、呉永輝は研究的な雰囲気を作りたがっているが、締め切りに追われる習性を完全に避けるのは難しいと語りました。本来はアルゴリズムを自前で開発して問題を解こうとしていたチームが、サービスを急いで立ち上げるためオープンソースを手直しして使うよう指示された例もありました。

2026年1月、リャン・ルーボーは全社会議でその年のの中核目標を「頂点に上り詰めること」に定め、短期的にはDoubaoとその海外版である秘書アプリDolaを最高水準へと引き上げることに集中する

と述べました。1年前の同じ場で彼は、研究チームに知能の限界を
探るよう求めています。革新にはある程度のグレーゾーンと混乱
が必要だが、競争に勝つには秩序と規律が必要だという言葉のとおり、
この二つの要求のあいだで均衡を取ることが、呉永輝がシード
を率いながら解き続けなければならない課題として残っています。

中核モデルと開発の過程

呉永輝(吴永辉)という名前が中国のAI業界で重みを持つのは、彼
が作った会社ではなく、彼が17年間Googleで磨き上げたモデルに
あります。2008年に博士号を取得してGoogleに入った彼は、最
初の7年間でソフトウェアエンジニアとして過ごし、検索ランキン
グエンジニアリングの中核作業に携わりました。コードを扱う立場
から出発し、研究者へと移っていった経歴が、その後彼が率いる組
織の性格を決定づけることになります。

Google在籍時代に彼が残した大きな里程標は、Google・ニュー
ーラル機械翻訳、通称GNMTと呼ばれる2016年の研究です。既存の
統計的翻訳方式を一つのニューラルネットワークに置き換え、翻訳
品質を引き上げたこの成果は、機械翻訳の歴史における屈指の転換
点として残っています。検索ランキングエンジニアだった人物が、
言語をまるごと理解するモデルを設計する立場へと移った出来事で

あり、この移行がその後の彼のキャリア全体を貫く流れとなりました。

グーグル・ブレインに移った後も、彼の手を経た研究は積み重なっていきました。2020年には音響信号処理に使われるConformer構造を確立する作業に参加し、2022年には視覚と言語を一つの表現空間に整列させる初期モデルであるCoCaを発表しました。翻訳一つから始まった研究が音声や画像にまで広がった流れであり、この広がりが後にバイトダンスで彼がマルチモーダルモデルを統括することになる背景となります。

2023年、彼は巨大モデルの研究開発部門であるグーグル・ディープマインドの研究担当副社長の座に就きました。この立場で彼はGemini巨大モデルの開発と、グーグルがOpenAIを追いかけていた初期の追い上げ局面に直接携わりました。検索エンジニアから翻訳研究者へ、そしてマルチモーダル研究者を経て巨大モデル開発の統括へと上り詰めたこの軌跡は、17年という歳月のあいだ一つの中ですべて積み上げられたものでした。

2025年2月、彼はこの座を離れ、バイトダンスの巨大モデル部門であるシード(Seed)に基礎研究責任者として加わると公式に発表し

ました。LatePostの長文報道によれば、彼がシードを任された当時、部門は厳しい状況に置かれていました。数千人規模の研究チームと数百億元の資源を投じて2年間追い上げた末に中国トップクラスの基盤モデルを作り上げたものの、はるかに少ない資源で作られた数百人規模のチームのモデルにあっという間に追い抜かれた直後だったのです。

バイトダンスCEOのリャン・ルーボーは全社会議で、部門はもっとうまくやれたはずだと直接指摘し、部門責任者もこの指摘を受け入れました。呉永輝がシードに加わる際に背負った期待は、目先の成果改善にとどまらず、モデルの能力を中国国内首位へと引き上げ、世界の主要モデル企業と正面から渡り合える水準まで部門を押し上げることでした。シードの関係者たちが彼から受けた第一印象は、落ち着きだったと伝えられています。

加わった直後に彼がまず行ったのは対話でした。彼は100人あまりの中核研究者と密度の濃い1対1の面談を行い、この過程を通じてモデルアーキテクチャの方向性を導く研究者を新たに何人も抜擢しました。組織を再編する前にまず人を把握するやり方であり、このやり方がその後彼がシードで行った組織改編全体の出発点となりま

した。

彼が立てた目標は二つでした。一つは基盤モデルの能力を引き上げ、研究の効率を高めて製品発表を安定的に支えることでした。もう一つは研究志向の雰囲気を作り、一流の研究を行う一流のAI研究チームを作り上げることでした。彼の側近たちは彼を、グーグル出身というよりむしろバイトダンスの体系の中で育った人物のようだと評しました。これは彼が短期間のうちに組織文化に深く入り込んだということの表れと読めます。

組織改編の核心は、2025年1月に作られたシードエッジ(Seed Edge)という仮想チームでした。このチームは3年を期限とする評価の仕組みを設け、中核人材が目先の成果ではなく、より基礎的で長期的な汎用人工知能、すなわちAGIの課題を研究するよう促しました。呉永輝はこのチームの設計段階から直接関わり、バイトダンスはこのプロジェクトに限って四半期ごとのOKRと半期評価を全廃しました。

シードエッジと並んで、彼はフォーカス(Focus)という新しいチームを作りました。このチームは既存の部門の境界と分業構造を打ち破り、基盤モデルの難題に正面から挑み、次のバージョンのモデル

で向上させるべき点を集中的に研究する役割を担いました。残りの基盤モデル人員はエンジニアリング、データ、評価を含むベース(Base)チームに分かれ、現行世代のモデルの開発を続けました。三つのチームのあいだでは人材と課題が互いに入れ替わり、循環するように設計されていました。エッジで生まれた成果はすぐに下位のチームへと下り、フォーカスで見つかった長期的な課題は再びエッジへと上がっていく仕組みでした。

この組織再編には人事ラインの変化も伴った。2025年3月、当時LLMチームの責任者だったチャオムーが通報を受けて業務を一時停止された後、LLM傘下のプリトレインとポストトレイン、ホライズンチームはすべて呉永輝に直接報告する体制に変わった。同時期、前任のSeed責任者だったジュ・ウェンジャは呉永輝とともに数か月間部門を共同で率いた後、呉永輝に報告する構造へと移行し、大規模モデル応用部門を専任することになった。

科創板日報が伝えた全体会議での発言によれば、ジュ・ウェンジャと呉永輝はこの場で初めてともに登壇し、Seed部門の中核目標が知能の限界を探ることにあると明らかにした。二人はモデル研究とモデル応用が長期的には対立しないと説明した。知能水準が上が

らなければユーザーを満足させることはできず、応用サービスが活性化されてこそ研究への投資とフィードバックが続くという論理だった。ジュ・ウェンジャは、モデル応用はユーザーの視点から質疑応答や創作、問題解決、コーディングといった作業を完成度高く提供し、事前学習モデルの潜在力を引き出す方向へ進むべきだと付け加えた。

呉永輝が強調したもう一つの軸はインフラの透明性だった。彼は内部データとコードライブラリを部門内では公開しつつ、対外的には秘密を保持する方向で整理を進めた。以前は、あるグループの研究者が別のグループの文書を見るには、文書の所有者とその上司の両方の同意を得る必要があった。部門責任者が乗り出しても、すべての情報を確認できない場合があった。この壁を取り払ったことで意思疎通の効率は高まったが、2025年下半期には少なくとも2件のインターン情報漏洩事件が発生し、その後、文書公開の義務化措置の一部が撤回されることになった。

この1年間で呉永輝が出した目立った成果は、まもなく発表されるDouBao 2.0モデルだった。このモデルはGeminiに似た性格のマルチモーダルモデルで、パラメータ規模が1兆個に達し、Seedがこ

れまで訓練したモデルの中で最大の規模だった。訓練過程では会議の頻度が目立って増え、以前は3か月に一度集まっていた研究チームが毎月意思疎通するようになり、中核チームは隔週で会議を持った。彼が食堂でトレイを手に研究員と向かい合って座り、進捗状況を話す姿もたびたび目撃されたという。

訓練過程で研究チームが報告を終えるたびに、彼は普段、三つほどの簡単な質問を投げかけた。具体的な解決策をすぐに提示するよりも、研究員たちがより本質的な問いを自ら考えるよう導くやり方だった。この訓練過程でSeedが直面した本当の難関はインフラ層だった。過去2年間、追撃にばかり集中し、基礎能力の構築を相対的に疎かにしてきたため、パラメータ規模を拡張する過程で不安定さが生じ、進行がしばらく停滞することもあった。ある研究員はこれを、基礎が固まっていなければ上にレンガを積み増すたびに問題が起きる、と表現した。

この問題はSeedだけの悩みではなかった。OpenAIの強化学習インフラ責任者ウォン・ザイは、あるポッドキャストで、すべてのモデルチームのインフラにはバグがある、モデル企業は事実上そのバグを直すスピードで競争している、と語ったことがある。このスピ

ードが単位時間内に検証できるアイデアの数を決め、アイデア自体は人材密度で解決できるという説明だった。OpenAIもまた、3年間使ってきたインフラ体系を再構築する作業を進めていた。アリババのチェンウェンチームやテンセントも、それぞれインフラ組織を新たに編成し、同じ課題と格闘していた。

Seedのインフラチームは規模が数百人に達し、部門内部で数十種類のモデルの研究開発とテストを同時に支えており、経営陣はこの組織を国内最上位クラスと評価していた。ただし、あるSeed関係者は、インフラを再整備する作業には膨大な人員と物的資源、そしてかなりの信頼コストが要求されると述べ、これを走行中の車のタイヤを交換することにたとえた。結局、複数のチームが3か月間、モデルアーキテクチャと訓練データの両面から集中的に問題解決に取り組んだ。 DouBao 2.0は春節前に正常にサービスを開始することができた。

この訓練過程と並行して、Seedはマルチモーダル領域でも複数のモデルを相次いで発表した。テキストから画像を生成するSeedreamは4.0を経て5.0まで進み、テキストから動画を生成するSeedanceも1.0から2.0へと発展した。全二重方式で自然な対話を処理する

音声モデルSeeduplexは、DouBaoアプリに搭載された後、対話の滑らかさを高めたと伝えられている。

マルチモーダル研究を率いているのは、2024年にアリババからByteDanceに移ってきたジョウ・チャン(周暢)であり、彼が率いるチームの代表的な成果がDouBaoスマートフォンアシスタントモデルである。この1年の間に、ビジュアル・マルチモーダル生成を率いていたヤン・ジェンチャオが休職に入り、ビジュアル基盤モデル研究を率いていたポン・ジャスが退社したことで、ジョウ・チャンが管理する範囲はSeedreamとSeedanceにまで広がった。インフラ組織はシャン・リャンが率いている。ByteDanceのAIラボが担当していたサイエンスAIとロボティクス、責任あるAIの三方向はSeedに統合され、責任者のリー・ハンが呉永輝に直接報告する構造に組み込まれた。

マルチモーダル基盤モデルの中で技術的に最も具体的な事例はSeed1.5-VLである。このモデルは5億3200万パラメータ規模のビジョンエンコーダーと、200億パラメータの混合エキスパート(MoE)構造の言語モデルバックボーンを組み合わせた形で設計された。教師あり微調整段階では、5万件のSFTサンプルと長い推論過程を含むL

ongCoTデータを2エポックにわたり、ビジョンエンコーダーを固定した状態で学習させた。訓練には13万1,072トークンの長文脈長とAdamWオプティマイザーを使用し、ベータ値は0.9と0.95、重み減衰は0.1に設定された。

このモデルは公認ベンチマーク60種のうち38種で当時最高水準の性能を記録した。画面操作やゲームプレイといったエージェント領域では、OpenAIのGPT-4やClaude 3.7を数値で上回った。Seed1.5-VLはByteDanceのクラウドプラットフォームである火山引擎(Volcano Engine)に公式ID「doubao-1.5-thinking-vision-pro-250428」として配備され、実サービスで使われている。関連コードと文書はGitHubを通じて公開されている。

Seedが発表した学術研究の中にはモデル統合技術もある。事前学習の途中で一定の学習率で訓練された複数のチェックポイントを結合し、モデルの性能がその後どう変化するか、すなわちアニーリング挙動をあらかじめ予測する手法を扱った研究である。このほかにも、大規模分散強化学習システムであるDAPOや、高難度推論を強化学習で扱うVAPOなど、システムレベルで知能の限界を広げようとする研究がSeedの中で続いた。

呉永輝が部門全体会議で繰り返し強調したのは、研究成果を外部に発信せよという要請だった。彼はSeedが優れた成果を数多く出しているのに、外部ではそれがあまり知られていないと述べ、研究者たちに論文やブログで成果を発信し、個人のホームページも整えるよう勧めた。彼が加わってから3か月間で発表された論文数は2024年一年間の発表量を上回り、これは研究員たちの自発的な動機を大きく高めたと伝えられている。

ただし2025年下半期からは、経営陣が論文発表に高い品質を保証し、現在開発中の中核技術とは無関係であるべきという条件を新たに課した。このため月間論文発表数はやや減った。Seedの内部では、呉永輝が研究の雰囲気を作りたいと考えているものの、締め切り志向を完全には避けられていない、という評価が出た。2026年1月、リャン・ルボは全社会議で、その年の中核課題は頂点に立つことだ、短期的にはDouBaoとその海外版であるドラ(Dola)アシスタントアプリを業界最高水準に引き上げることだ、と述べた。1年前の同じ場で、彼はAI研究チームが知能の上限を探るべきだと語っていた。この二つの要求の間でバランスを取ることが、呉永輝がSeedで解き続けなければならない課題として残っている。

現在の立ち位置と中国AIにおける意味

Seedの基礎研究総括責任者の座に就いた呉永輝は、2026年現在もByteDanceの中で最も目立つ研究指揮官であり続けている。彼が担っているのは、画面に見えるチャットボット一つをきれいに磨き上げる仕事ではなく、次世代の基盤モデルがどこまで到達できるかを見極める基礎研究全体を率いる仕事である。

この座の重みは、彼が歩んできた道にそのまま表れている。彼はGoogleで17年働き、その最後の肩書きはGoogleのエンジニアリング組織が与える最高の栄誉であるGoogle Fellowだった。ByteDanceに移る直前にはGoogle DeepMindの研究担当バイスプレジデントとして、Geminiの大規模モデルの研究開発を率いていた。

2025年2月、彼がByteDance入りを正式発表した際に就いた肩書きは、DouBao大規模モデルチーム、すなわちSeed部門の基礎研究責任者だった。入社の背景を彼は、世界トップクラスの研究をし、最高のAI研究チームを作りたいという言葉で説明した。

Seedの中での呉永輝の立ち位置を理解するには、ジュ・ウェンジャというもう一人の人物も併せて見る必要がある。二人はそろってByteDanceのCEO梁汝波(リャン・ルボ)に直接報告する構造をなし

ている。ジュ・ウェンジャはモデルを実際の製品へと移す応用研究を統括し、呉永輝は知能の限界を押し広げる基礎研究を指揮する。

2025年3月に開かれたSeed部門の全体会議は、この二重構造が初めて公にされた場だった。ジュ・ウェンジャと呉永輝が初めてともに登壇し、部門の目標を語り、この会議を機に二人の役割分担が対外的に確認された。

その場で二人が明らかにしたSeed部門の最も重要な目標は、知能の限界を探ることだった。同時に、組織文化を固め、技術をより開かれたものにし、オープンソースを検討するという方針も併せて示された。

知能の限界を探るという目標は、言葉だけに終わらなかった。この目標は、先に発表されていたAGI長期研究計画であるSeed Edgeプロジェクトとして具体化された。能力とアイデアを備えた研究者たちが、不確実性を抱えながらも大胆な課題に長く取り組めるようにすることが、この計画の核心だった。

呉永輝は、この長期研究の重要性をとりわけ強調した人物だった。彼はSeed Edgeに投入する計算資源をさらに増やし、社内外から

潜在力と好奇心を備えた研究者を引き続き招き入れると述べた。

ByteDanceがSeed Edgeという名のAGI長期研究計画を正式に発表した時点は2025年1月だった。この発表は、ByteDanceが基礎研究から技術開発、製品応用に至る全区間にわたる布陣を終えたことを意味し、同社がAIにどれほど確固たる決断を下したかを示す場面でもあった。

この決断が下された背景には、その年の旧正月連休に登場したDeepSeekがあった。少ない人員と比較的小さな資源で強力な性能を発揮したDeepSeekは、業界全体が自らの組織の研究理念と人材育成の方式を見直すきっかけとなった出来事だった。

呉永輝が着任した後、ByteDanceの大規模モデル組織では実際に組織改編も行われた。市場では一時、これまで朱文佳に報告していたアルゴリズムや技術の責任者の一部が、呉永輝に報告する体制へと変わったという話が広まった。

Seed組織が再編される過程で手が加えられたのは、人員と目標だけではなくだった。ByteDanceは安定した研究環境を守るとして、Edgeプロジェクトに課されていた四半期ごとのOKRと半期評価を廃

止した。

四半期ごとのOKRと半期評価は、大手IT企業の間でもByteDance特有の管理方式として知られていた制度だった。もともとは事業を頻繁に点検し効率を高めるために作られた仕組みだったが、頻繁な評価がもたらす週次・月次報告書の作成負担は、少なくない社員を疲弊させていた。呉永輝がこの制度を取り除いたのは、短期的な成果に追われることなく、長く取り組むべき研究にふさわしい環境をつくるための選択だった。

研究と製品の間での資源配分をめぐっても、二人の考え方の違いが表れた。呉永輝は、長期的な視点で見れば研究と応用は対立しないと考えていた。知能水準が上がらなければユーザーを満足させることはできず、応用が活発であってこそ投資も続き、モデルへのフィードバックも大きくなるという論理だった。

朱文佳もまた同じ文脈で、核心は結局のところ知能の限界であり、モデルの応用は長期的に見てモデル自体の能力に従わざるを得ないと語った。彼が思い描く応用の姿は、ユーザーの視点から見て質疑応答や創作、問題解決、コーディングといった作業を完成度高くこなし、検索やエージェントのようなツールを賢く使いこなして、

事前学習モデルの潜在力を引き出す方向だった。

呉永輝は、Seedが良い成果を上げながらも外部にはあまり知られていない点を残念に思っていた。今後は論文やブログで成果をより発信し、コミュニティが扱いやすい規模の中小型高密度モデルをオープンソースとして公開する案も検討すると明らかにした。

こうした方向性が実際のモデルとして形になった事例が、2025年5月に公開されたSeed1.5-VLである。視覚と言語を同時に扱うこのモデルは、5億3200万個のパラメータ規模のビジョンエンコーダーと、トークン処理時に実際に稼働するパラメータが200億個規模のMoE(専門家混合)構造の言語モデルを組み合わせた形で設計された。

規模を大きくすることなく、このモデルは複数のベンチマークで良好な成績を収めたと報告された。開発陣は、60項目に及ぶマルチモーダル評価のうち多くで最高水準の性能を記録したと明らかにした。仮想画面をマウスとキーボードで操作する課題のようなエージェント型の作業でも、競合モデルと比肩し得る結果を示したと説明した。このモデルはGitHubとSeedの公式ページを通じて配布され、開発者がすぐに利用できるようになった。

同年4月には、画像生成を扱うSeedream 3.0の技術報告書が発表された。この報告書もまた、Seed部門が言語と視覚にまたがる複数種のモデルを同時に積み上げていることを示す事例だった。

呉永輝が米国の中核的な研究現場を離れByteDanceへ移った出来事は、個人の転職を超えた意味を持つものとして読み取れる。先端半導体とコンピューティングハードウェアの導入が年々厳しさを増す状況の中で、長年Googleで実績を積んできた研究者が中国企業の基礎研究を指揮する立場に就いたという事実自体が、人材の流れがどちらの方向にも向かい得ることを示したと言えそうだ。

ByteDanceが彼の着任を機に手を加えたのは組織文化だった。短期評価制度を取り除き、長期研究を保証するSeed Edgeのような枠組みを打ち立てたことは、速い成果に偏りがちだった中国の技術業界の中で、別の方式の研究運営が可能であることを示した事例として評価に値する。

呉永輝と朱文佳がともに描いたオープンソース戦略も注目に値する。Seed1.5-VLを公開された形で世に出し、中小型高密度モデルの公開までも検討する姿勢は、重みを非公開にして商業利用のみに限定する方式とは異なる道をSeed部門が歩んでいることを意味して

いると読み取れる。

Seed部門が歩んできた道のりを振り返ると、呉永輝の役割は結局二つに要約される。一つは、知能の限界に向けて長い時間を要する問いを投げかけ続ける研究文化を守ることである。もう一つは、そうして得られた成果をSeed1.5-VLやSeedream 3.0のように実際に使えるモデルへとつなげていくことである。この二つの役割が噛み合って動いている姿こそが、2026年現在の中国AIの中で彼が占めている立ち位置を最もよく説明している。

Support This AI Library Book

If this book was useful, you can support future translations, public editions, and open AI knowledge projects through PayPal.



PayPal

[**https://paypal.me/kimkjj**](https://paypal.me/kimkjj)

Scan the QR code or open the link above.