

ヤン・ルカン
人工知能の父

金京鎮弁護士

ヤン・ルカンー人工知能の父

チューリング賞受賞者が大規模言語モデルに背を向けた理由

金京鎮弁護士

目次

1. 序文
2. 第1章 ソフズィ・スー・モンモランシーの図書館
3. 第2章 逆流する誤差
4. 第3章 ホルムデルの視覚皮質
5. ヤン・ルクン、人工知能の父
6. 第5章 ベル研究所の黄昏
7. 第6章 マンハッタンのジャズとディープラーニング結社
8. 第7章 特徴を設計していた時代の終わり
9. 第8章 イメージネットという地震
10. 第9章 ザッカーバーグの食卓
11. 第10章 チューリング賞、栄光と優先権
12. 第11章 鳥の羽根を模倣するな
13. 第12章 ChatGPTというサイレン
14. 第13章 メタを去って
15. 第14章 世界を見る機械
16. エピローグ ライト兄弟は鳥を模倣しなかった

序文

ヤン・ルクン、人工知能の父

1989年のある夜、ニュージャージー州ホルムデルの実験室の画面に、手書きの数字がひとつ映りました。歪んだ3でした。29歳のフランス人青年が椅子を引き寄せて座りました。彼が作ったニューラルネットワークは、その数字をピクセルひとつひとつで受け取りました。誰も3の形を規則として教えてくれたことはありませんでした。機械は数千枚の郵便番号を見ながら、自分自身で3の姿を学んだのです。画面に答えが表示されました。3。青年は笑いませんでした。彼はすでに次の数字を待っていました。当時、学界のほぼすべてがニューラルネットワークを終わった道と呼んでいました。知能とは脳に生まれつき備わった規則であり、機械にはその規則を人間が書き込まなければならないと信じていました。ヤン・ルクンはその対極に一人立っていました。知能は学ぶものであって、生まれつきのものではないと、彼はパリのある図書館からずっとそう考えていました。

あの夜の青年は今、66歳になりました。そして彼は正しかったのです。彼が一人で抱え続けたニューラルネットワークは終わった道ではありませんでした。今日、私たちが手に握る翻訳機、顔を認識する写真アルバム、人間のように話を聞ってくれる対話プログラムは、すべてあの画面に映った歪んだ3の遠い子孫です。学界のほぼすべてが間違っていて、対極に一人立っていたたった一人が正しかったのです。その功績により、彼は2018年、コンピュータ科学のノーベル賞と呼ばれるチューリング賞を受賞しました。ともに冬を越したジェフリー・ヒントン、ヨシユア・ベンジオと並んで立ちました。

もしここで物語が終わっていたなら、この本は異端児がついに認められた一般的な成功物語になっていたでしょう。しかし物語はそこで終わりませんでした。世界がニューラルネットワークを信じるようになったまさにその瞬間、ルクンはふたたび立場を変えました。今度は対極に一人で立つためにです。

今、世界は一方向に走っています。2022年末にChatGPTが扉を開いて以来、人間の言葉を受けて文章をつなぎ合わせる大規模言語モデルが世界中を覆いました。企業はここに数千億ドルを注ぎ込み、政治家はここに国の運命を賭け、人々はここに知能の完成を見ます。その真つ最中で、ルクンは穏やかに言います。それは行き止まりの道だと。あと数年すれば、今の大規模言語モデルはほとんどが用をなさなくなるだろうと。自分が人生をかけて蘇らせた技術の上に世界全体が乗っているのに、正にその技術を蘇らせた本人が今の方向は間違っていると言っているのです。2025年末、彼は12年間率いてきたメタの研究所を出て、ワールドモデルという別の道を掘る会社を立ち上げました。65歳でした。

なぜ今ルクンなのか。流行の心臓に向かって死刑宣告を下す人の物語だからというだけではありません。一人の人間が60年にわたって二度、多数が一方に流れるとき対極を指し示したという事実だからです。最初は誰もニューラルネットワークを信じないときにそれを守ったことであり、二度目は誰もがニューラルネットワークの一つの分枝に熱狂するときに、その分枝の限界を指摘したことでした。彼が二度とも正しいという保証はありません。最初は歴史がすでに答えを出しましたが、二度目の答えはまだ書かれていません。この本がしようとするのは、その答えを代わりに出してあげるのではなく、彼がなぜ対極に立つ人間になったのかを彼の人生で辿ってみることです。

ルクンの人生には、人工知能の二つの冬が年輪のように刻まれています。彼がパリの図書館で忘れ去られた学習機械を掘り起こしていた頃、ニューラルネットワークはすでに一度死刑宣告を受けた後でした。彼は死んだと宣言された分野に自分の足で歩いて入ったわけです。ベル研究所でアメリカの小切手の手書き文字を読み取ることに成功した直後に、二度目の冬が訪れました。実用化に成功したまさにそ

の年に、彼が身を置いた研究所は解散され、彼の論文は学会からたびたび戻されました。彼が正しかったことが世界に証明されるまでに、そこからさらに十数年がかかりました。この本が彼の勝利とともに彼の冬を長く見つめる理由がここにあります。対極に立つことがどのような代償を払うのかを見なくしては、彼が今なぜまた対極に立っているのかを理解することができないからです。

この本が扱うことは二つです。一つはヤン・ルクンという人間の60年です。パリ北部の郊外で生まれ、図書館で忘れ去られた学習機械を掘り起こした少年。ベル研究所でアメリカの小切手の手書き文字を読み取った工学者。二度の人工知能の冬に耐えて、シリコンバレーにオープンな研究所を設立した科学者。そして65歳で自ら帝国を去った老学者の軌跡です。もう一つは、その軌跡に重なって流れた人工知能60年の思想史です。知能とは規則なのか、学習なのか。機械は世界を描き出さなければならないのか、それとも理解すべきなのか。この古い問いがどのように冬と春を往復しながら学界を揺るがしたのか、ルクンという一人の人間の窓を通じて眺めます。

この本が扱わないことも明確にしておきます。この本はルクンが正しく、大規模言語モデル陣営が間違っていると判定する裁判ではありません。スケーリング懐疑論とスケーリング肯定論、人工知能滅亡論とそれに対する反論は今もなお生きている論争です。私はルクンの人生を軸としますが、彼が対峙した相手の根拠は、その相手のものとしてそのまま記しました。彼を風刺化した人々も、彼が風刺化しかけた人々も、ここではそれぞれの場所に立ちます。ニューラルネットワークの動作原理を再現する技術書でもありません。畳み込みと逆伝播と世界モデルの数学は、数式を取り除いて比喻で解きました。猫の視覚皮質、紙上の郵便番号、羽を写さずに空を飛ぶ機械のような図で。

この本を書いた理由を率直に申し上げます。私は技術者ではありません。弁護士であり、文を書く人間です。法廷で長く学んだことが一つあるとすれば、片方の言葉だけが大きく聞こえるときに真実から最も遠い瞬間だという事実です。今、人工知能を巡る話がそうです。もうすぐ人間を超えるという声と、全部泡だという声が交互に大きく響き渡っているのに、正にその間で異なる選択肢を静かに指し示す人の話は、よく聞こえません。私はこの文が判決文になることを望みませんでした。論争の両側をバランスよく写し、多数と異なる道を歩んだ一人の人間の根拠を読者の前にそのまま置いて、人工知能時代を生きる私たちの選択肢を一段広げる教養書になることを望みました。ただそれだけです。

この伝記を書くために、私はルクンが残した論文と講演、彼が40年にわたってさまざまな国のメディアと交わした対談、そして彼を批判した人々の文と一緒に読みました。英語の資料が最も厚く、フランス語には彼の少年時代と母国に関わる話が、その他の言語にはアジア学界が彼を受け入れ、また反論した跡が残っていました。実在する人物の伝記ですので、確認されていない発言は作りませんでした。引用は話者と出典を明記し、最近のメタ退社と起業のように細部がまだ流動的な部分は報道に基づくことを都度記しておきました。不正確な引用より引用のない叙述の方が優れていると考えました。

読み方は難しくありません。本文は五部、14章で生まれた順序に従って流れます。少年期から順に読むと、人工知能が冬と春を往復した60年が一人の人生の上に自然と重なって見えます。技術用語は初めて出る時に易しい言葉で説き明かしてありますので、なじみのない略語が出てしまっても止まらずにそのまま読み進んでかまいません。今の論争に興味がある読者なら、大規模言語モデルを扱った第5部から開いても良いです。各章は実際に起きた一つの場面から始まります。ベル研究所の夜通しの実験室、チューリング賞授賞式の会場、パリの講壇、アメリカ経済放送のスタジオ。その場面の中にしばらく立ち止まれば、数字と事実は後からゆっくり続いてきます。

最後に一言付け加えます。この本の主人公はまだ進行中の人間です。彼が66歳の年にパリの新しい黒板の前で投げかけた問いの答えは、この本が印刷される時までも出ていません。ですからこれは結末を教えてくれる伝記でもありません。結末が書かれる前に、その問いを一緒に見つめてみようという招待に近いです。彼が正しかろうと間違っていようと、多数が一方向に走るときに対極を指し示す人の心がどのように作られるのかを知ることは、人工知能が私たちの生活のほぼすべての場所に押し寄せているこの時代を生きていく上で、一度は必要な営みだと私は信じます。さあ今、その物語を、歪んだ3が映った夜よりもはるかに前の、パリ北部の丘のある少年から始めましょう。

2026年夏 キム・キョンジン

弁護士

kimkj.com

第1章 ソワズィ・スー・モンモランシーの図書館

ヤン・ルカン、人工知能の父

リード楽器を吹く少年

はんだ鐔の先端から細い煙が立ち上がりました。回路基板の上に身を屈めた高校生が、配線を一本、丁寧に繋ぎました。彼が手直ししていたのは、いとこの兄のアナログ・シンセサイザーでした。いとこの兄は電子音楽家を夢見ており、初期型のシンセサイザーを一台持っていました。音は出ていましたが、思い通りには流れていません。幼いヤン・ルカン(Yann LeCun)は、その機械の回路を開けて手を入れ、音を決められた順序で自動的に出す装置であるシーケンサーを、自分で設計して作り、取り付けました。完成した装置は、人の手が触れなくても決められた拍子で音を送り出しました。抵抗と蓄電器を選び、電圧が変動する区間を捉え、再びはんだ鐔を温める午後が続きました。電流の波形が音になって出るその過程を、彼は教科書ではなく、指先で学びました。

彼が生まれた場所は、パリ北郊の小さな村ソワズィ・スー・モンモランシー(Soisy-sous-Montmorency)でした。1960年7月8日のことです。本名はヤン・アンドレ・ル・カン(Yann André Le Cun)です。姓のル・カンは、フランス北西部ブルターニュ(Bretagne)地方のゲンガン(Guingamp)一帯に根を持つブルトン語の名前ル・キュンフ(Le Cunff)から来ています。名前のヤンもまた、フランス語のジャン(Jean)、英語のジョン(John)に相当するブルトン式の表記でした。パリ郊外で生まれた子どもの名前の中に、すでに海を望むケルトの西の端が入っていたわけです。

村の名前に付いたモンモランシーは、ちょうど上の丘の古い小都市を指します。この丘には、知性史の由緒ある影が一つ落ちています。1756年からおよそ6年間、ジャン＝ジャック・ルソーはモンモランシーに滞在して執筆していました。そこで完成させた本の一つが『エミール』、子どもは灌輸ではなく経験で学ぶべきだと主張した教育論でした。それから200年が経った後、同じ丘の下で生まれた一人の少年が『経験で学ぶ機械』に生涯をかけることになります。もちろん偶然です。しかし伝記作家の目には、なかなか見過ごしがたい偶然なのです。

父は航空工学者でした。飛行機を設計する人であり、家でもいつも何かをばたばた作るアマチュア発明家でした。家の中には工具や部品が転がっており、何かを買って使うより自分で作って使うほうが自然な家でした。ヤンは父と兄と一緒に模型飛行機を組み立て、無線操縦装置を手ずから作りました。翼の角度をちょっと変えると飛行がどう変わるのか、信号をどのように送ればプロペラが回るのかを、彼は説明ではなく墜落で学びました。作って、飛ばして、落として、直す循環。部品を整列させ、力の方向を合わせるこの習慣は長く残りました。後に、物理世界を自ら理解する機械を設計しようと決心したとき、その手の習慣が礎となったのです。

9歳の頃、両親は彼をパリの映画館に連れていきました。スタンリー・キューブリック監督の『2001年宇宙の旅』がかかっていました。ゆっくり流れる宇宙船、寡黙な画面。同年代の子どもなら退屈するような映画でした。少年は退屈しませんでした。画面の中にはHAL 9000というコンピュータがいました。話し、考え、宇宙船全体をたった一人で操る機械でした。少年の関心はその日以降、少し方向を変えました。機械を直す仕事から、機械の中に心のような何かを作り込む仕事へ。2019年にレックス・フリードマンとの対談でHAL 9000の話が出たとき、彼は知能を持つ機械という課題が子ども時代のその映画で自分に植え付けられたと振り返りました。

時代もそのような夢を促しました。彼が9歳になった1969年、1年間にわたり空では二つの事件が起こりました。春には、フランスとイギリスが共同で作った超音速旅客機コンコルドが初めて飛び立ち、夏にはアポロ11号が月に着陸しました。航空工学者の息子にとって、空と宇宙は夕食時のテーブルの話題であり、ロケットとコンピュータは漫画を超えて新聞の1面を占める実物でした。科学が世界を変え、という感覚を空気のように吸いながら育った世代。ルカンはその世代の真っ只中にいました。

音楽はまた別の通路でした。高校と大学の時代、彼はバロック、ルネサンス音楽、古い民俗音楽に深くはまっていました。聞くだけでなく演奏しました。リコーダーを吹き、曲がった管の先から鼻が詰まったような音がする中世の二重リード木管楽器クルムホルンを吹き、オーボエを吹きました。ルネサンス舞踏団の公演を支える小さなアンサンブルの団員としても活動していました。舞踊手の足が拍子を踏む間、彼は数百年前の楽譜に従って呼吸を調整していました。木管楽器は息をどのように入れるかによって音程がわずかに変動します。演奏者はその変動を耳で捉え、体で還元入力します。シンセサイザーの回路に触れた手とリコーダーに息を吹き込んだ口は、結局のところ同じ感覚を向いていました。信号が体を通過して形態になるその過程に対する感覚でした。

パリの航空宇宙博物館(Musée de l'Air)に行くと、ルデユク0.10(Leduc 0.10)という初期のラムジェット航空機が展示されています。空気を吸い込んでそのまま吐き出す原理で飛ぶ、操縦席が機体前部のガラス管の中に埋め込まれた奇妙な飛行機です。彼は子どもの頃からこの機体の急進的で幾何学的な形態に魅了されていました。年が経った後でも、その愛着は消えず、彼は電子木管楽器を手ずから調整して演奏する趣味を生涯手放しませんでした。飛行の幾何学と音の物理学。手で作り、耳で聞き、目で確認する知識。彼が生涯つかまえることになる問いの種たちが、すでにその少年の部屋の中に散らばっていました。ただ、その種がどの畑で育つことになるのかは、少年自身もまだ知りませんでした。

ただし、数学は彼の得意分野ではありませんでした。少なくとも学校の成績表の上ではそうでした。黒板の上の抽象記号のゲームでは、彼は目立ちませんでした。試験紙が要求する速度と彼の好奇心が動く速度は異なっていました。彼は答えを素早く出す学生ではなく、問題の背後にある装置を分解して見たいと思う学生でした。彼の目が輝く場所はほかにありました。記号より機械、公式よりその公式が実際に作動する方法でした。このズレが彼をフランスのエリート教育の正統なコースから少しずれさせ、そのズレが後に彼の確固とした資産になります。

ロワイモンの論争：チョムスキーとピアジェ

1975年10月、パリ北部のウァーズ川沿いのロワイモン修道院(Abbaye de Royaumont)に学者たちが集まってきました。13世紀に建てられたシトー会修道院を改造したこの学術空間で、10月10日から4日間にわたり、ある論争が繰り広げられました。一方の座にはジャン・ピアジェ(Jean Piaget)79歳が、向かい側にはノーム・チョムスキー(Noam Chomsky)46歳が座りました。発生的認識論を創設した老学者と生成言語学を創設した中年の巨匠が顔を向き合わせたのはこのときが初めてであり最後でした。座席には心理学者、言語学者、生物学者、哲学者、人類学者、そして人工知能研究者まで座っていました。人間の知能はどこから来るのか。4日間の議題は、結局のところその一文でした。発表が終わると反論が続き、食事の席でも論戦は止みませんでした。後に認知科学史の古典となる発言がその4日間の間に噴き出しました。修道院から南に少し下るとモンモランシー丘の下にソワズィ・スー・モンモランシーがあります。同じヴァルドワーズ県、15歳のヤン・ルカンが模型飛行機を作っていた地区です。少年はその秋、自分の家の近くでどのような論争が繰り広げられていたかを知りませんでした。

1970年代のフランスは、知識人の言葉が重く扱われる国でした。哲学者と言語学者がテレビで議論し、構造主義と人文科学の言葉が新聞の文化面を埋め尽くしました。ロワイモンやセリジラサルのような古い建物が学術討論の聖地として使われ、その録音起こしが本にまとめられて書店に並びました。言語学と心理学の決闘記録が一般書店の売場に上がり、工学生がそれを買って読むことが不自然ではない国でした。ピアジェとチョムスキーの対決記録もこの慣行に従いました。学会を組織したイタリア出身の認知科学者マッシモ・ピアテッリ・パルマリーニ(Massimo Piattelli-Palmarini)が発言をまとめ、1979年にスイユ(Seuil)出版社が『言語理論、学習理論：ジャン・ピアジェとノーム・チョムスキーの大論争』というタイトルで出版しました。分厚いその本が書棚に差し込まれてからおよそ1年経ったとき、20歳の電気工学生が一人、それを手に取ることになります。

論争の一方にはチョムスキーがいました。その時代の言語学と認知科学を支配していた人物でした。彼の傍には認知哲学者ジェリー・フォード(Jerry Fodor)がいましたが、フォードは師匠格のチョムスキーより一歩進んで、概念そのものが学習できないという議論まで押し付けて、会場をざわめかせました。彼らが主張したのは生得主義でした。人間の脳には生まれつき高次の文法と論理構造が組み込まれているという主張です。いわゆる普遍文法です。この観点から、赤ん坊は言語を新たに学ぶのではありません。すでに頭の中に入っているスイッチを、自分が育った環境に合わせてオン・オフするだけです。この考えは、その時期の人工知能研究の主流だった記号主義と関連していました。知能とは規則と記号を操作することであり、その規則を人間が機械に書き込めばいいという信念です。

反対側にはピアジェがいました。子どもが世界を理解していく過程を生涯観察した心理学者でした。彼は知能が生まれつきの記号装置だという考えに反論しました。赤ん坊は世界に触れ、転び、再び試みながら、自分の中の認知構造を自ら作ってゆくというのです。カップを落としてみて重力を推測し、隠された玩具を探し出して物が継続して存在することを学ぶそのゆっくりとした過程の一つ一つが、知能の建築工程だという観点です。知能は遺伝子に刻まれた完成品ではなく、経験で少しずつ積み上げる建築物だということ。これが構成主義です。

4日間の論争がどちらの勝利で終わったのかは、参加者ごとに記憶が異なっていました。確かなのは、この出会いが最後だったという事実です。5年後の1980年9月、ピアジェは84歳で亡くなりました。彼は、自分の思想が会議録1冊を通じてパリ郊外の若き工学生の元へ届き、その青年の人生を決定づけることになるとは知らないままだったのです。そして論争そのものもその時点では終わりませんでした。知能は生まれつきか学習されるのかという問いは、半世紀後の人工知能学会でほぼ同じ規模で再現されます。その再現の真っ只中に立つことになるのが、まさにこの会議録を読んで育ったルクンだったのです。

1980年、ESIEE パリ2年生のヤン・ルクンがこの本を手に取りました。読み進んでいた青年の頭の中で、天秤が一方に傾きました。彼はチョムスキーの側に立つことはできませんでした。後に2019年に出版した自伝『機械が学ぶとき—Quand la machine apprend』で彼は、その本を読む瞬間、チョムスキーが述べる先天的な論理構造が生物学的にも物理学的にも現実ではあり得ないと確信したと書きました。言語能力を含む人間のほぼすべてのものは、その後学習されるのであり、知能の本質は学習にあるという考え。この確信は20歳で根付いた後、60歳を過ぎてもなお揺るがなかったのです。

彼が本を読む方式は、言語学者のそれではなく、工学者のそれでした。脳は謎ではありません。物理的な器官です。神経細胞とシナプスから成り、エネルギーを消費し信号を送受する機械装置です。ならば、その装置の中に文法書1冊が丸ごと印字されているという主張よりは、装置が世界と衝突しながら配線を自ら修正していくという主張が、工学的にはずっともっともらしく見えました。モデル飛行機を

飛ばし、シンセサイザーを改造しながら育った青年にとって、知能の問題は形而上学である以前に、結局のところ設計の問題だったのです。

なぜこの若き工学生がこれほど断固たることができたのでしょうか。彼の頭の中には、既に異なる材料が蓄積していたからです。彼は毎晩、リチャード・ドウダ(Richard Duda)とピーター・ハート(Peter Hart)が著した『パターン認識の古典』を枕辺に置いて、すり切れるほど読みました。データから規則を抽出する統計と幾何学の言語が身に染みしました。関心はそこで止まらず、さらに深い問いへと延びていきました。情報を圧縮するとはどういうことか、無秩序から秩序はどのように生まれるのか。彼はアンドレイ・コルモゴロフ、レイ・ソロモノフ、グレゴリー・チャイティンが構築したアルゴリズム複雑性理論に掘り下げ、1950～60年代にアメリカのイリノイ大学の生物学コンピュータ研究所(BCL)で自己組織化理論を切り開いたオーストリア出身の物理学者ハインツ・フォン・フェルスター(Heinz von Foerster)の著作に夢中になりました。

フェルスターの著作で彼が得た一つのイメージがあります。無重力の空の箱の中に、小さな立方体の磁石が数十個、あたかも浮遊していると想像してみましょう。外からそつと箱を揺り動かすだけでも、磁石はそれぞれの極性に従って互いに吸着し、対称的で複雑な構造を自ら創造してしまいます。誰がその形を設計したわけでもないのに。ルクンは知能もそうあるべきだと考えました。多くの結合強度を持つ大きなニューラルネットワークにデータという柔らかな振動を与えると、その中で意味の構造が自ずと成長してくるということ。規則を外から刻み込む代わりに、揺り動かして自ら整列させる道です。

ロワイヨモンの分厚い議事録は、彼が既に抱いていたこの直感に名前と系譜をもたらしました。知能は学習されるものであるという側に立つという決心がつきました。しかしその本が彼にもたらしたものはもう一つありました。思想は方向を定めますが、工学者には手で触れることができるものがが必要です。本棚の間に、まさにそのもの、彼の人生を変えるであろう小さな機械の名前一つが隠れていたのです。

INRIA図書館で出会ったパーセプトロン

水曜日の午後、ESIEE パリの講義室は空いていました。この学校は毎週水曜日の午後に正規授業をまったく行わず、その時間を学生それぞれの研究に全て充てていたのです。ルクンはその午後を一つの場所に捧げました。カバンを担いでパリ西部の郊外ロクンクール(Rocquencourt)へ向かったのです。そこにはフランスが1967年に設立した国立情報自動化研究所(INRIA)がありました。ドゴール政府がNATO統合軍事司令部から撤退した際に空いた旧NATO施設の敷地に、フランスはアメリカに対抗するコンピュータ科学の拠点を植えたのです。その研究所の図書館は、当時フランスでコンピュータ科学文献を充実させていた書庫の一つでした。大学生の身分で研究所の書庫を出入りすることは珍しかったのですが、彼は毎週その道を通いました。

彼が書架で探ったのは埃をかぶった古い文献でした。最新の論文とは距離がありました。1950年代と1960年代の初期人工ニューラルネットワークの資料たち。20年以上誰も開かなかったため、ページをめくるたびに紙の香りが立ち上る学術誌たちでした。ロワイヨモンの本で彼を魅了したその一つの名前を追跡していたのです。パーセプトロンでした。論争の記録の中で、ピアジェの側に立って発言した一人がこの機械を取り出しました。MITの数学者にしてコンピュータ科学者シモア・パパート(Seymour Papert)でした。彼は知能が先天的な規則なくして学習によって構築されうることを示そうとして、フランク・ローゼンブラット(Frank Rosenblatt)が神経細胞に倣って作った学習機械を例に挙げました。パーセプトロンは規則を入力しなくても、データと正解との誤差を見て自ら結合の強度を少しずつ修正していきます。そのようにして人が手で書き記しがたい概念を自ら習得するというのです。

自ら重みを変えて学ぶ機械という言葉は初めて接した青年は、その箇所でも身震いしました。

身震いの次が図書館だったのには理由があります。質問する人がいませんでした。1980年頃フランスのいかなる大学の講座も、学習する機械を教えていませんでした。教科書にもなく、指導してくれる教授もない主題でした。学ぶためには自分自身で教育課程を編成する他にありませんでした。だからこそ彼のカリキュラムは文献目録でした。論文一編の参考文献から次の論文を見つけ、その論文の参考文献からさらに次を見つける式で、彼は20年前に途絶えた学問の地図を逆向きに描き進んでいきました。

書庫の古い文献たちはその機械の由来を教えてくださいました。ローゼンブラットはコーネル航空研究所で働いていた心理学者でした。1957年に構想を発表し、翌年には米海軍の支援を受けて実演までしてみました。1958年7月8日付のニューヨーク・タイムズはこの装置について、歩き、話し、見て、書く電子コンピュータの胚となるだろうという海軍の期待を伝えていました。ルクンが生まれるちょうど2年前、同じ日付の新聞でした。パーセプトロンはしばらくの間、時代の寵児でした。写真の中のローゼンブラットは部屋一つを埋め尽くす配線の束の前で未来について語っていました。ルクンはローゼンブラットが1962年に出版した分厚い理論書『神経動力学の原理(Principles of Neurodynamics)』にまで掘り下げました。20年前に著されたその本の中で、忘れ去られた著者は依然として生き生きとした声で学ぶ機械の可能性を論証していたのです。

ところが書庫を掘り進むほど、奇妙な断絶が浮かび上がりました。1950年代後半ほど熱かったニューラルネットワークの研究が、1960年代末を過ぎると、まるで切り落とされたように消え去っていたのです。論文の流れが途絶えた場所には沈黙だけが残っていました。沈黙の背後を追跡していたルクンは、冷徹な事実と直面しました。パーセプトロンをロワイヨモンで賞賛したまさにそのパパートが、それより6年前には、この機械の息の根を止めた当事者だったという事実です。

1969年、パパートは記号主義の大父マービン・ミンスキー(Marvin Minsky)と共に『パーセプトロン：計算幾何学入門』を出版しました。二人はこの本で、単層のパーセプトロンが排他的論理和というきわめて基本的な問題すら解くことができないことを数学によって証明してみせました。排他的論理和という名前が難しければ、階段の上下に各1つずつスイッチが付いた電灯を思い浮かべてください。どちらか一方のスイッチだけを上げると灯りが付き、両方上げるか両方下げると消える回路です。人間にはたいしたことの無いこの区別を、単層の機械はいくら長く学んでも引き出すことができませんでした。MITの2人の碩学が下したこの診断の波紋は大きかったのです。アメリカとヨーロッパでニューラルネットワークの研究費が次々と途絶えました。研究者たちは去り、残った者たちは学会でニューラルネットワークという言葉の口にするのを避けてました。人工ニューラルネットワークは非科学、さらには呪術に近いタブーとして扱われるようになりました。人工知能史の第一次の冬でした。ローゼンブラットはその冬が真っ最中だった1971年、43歳の誕生日にボート事故で亡くなりました。自分の機械が復権されるのを見ることなく。彼が作った言葉だけが書庫の本の背に残り、後日パリからやってきた大学生の手を待っていたのです。

一人の人間が10年程度の間隔を置いて同じ機械を殺し、また生き返らせたというこの矛盾は、ルクンを挫折ではなく反抗へと導きました。彼はこの冬が判定の誤りだったと考えました。ミンスキーとパパートが証明したのは単層機械の限界に過ぎず、複数層を積み重ねたニューラルネットワークまで崩壊させたのではなかったのです。層をもう一つ加えるだけで階段電灯の壁は消え去ります。二人もその可能性を知らなかったわけではありませんでした。ただ彼らはその本で、複数層への拡張研究が実りのない道となるだろうという推測を書き記していたのです。層を積み上げたときにその全体を学習させる方法を誰も知らなかったからです。学界はその推測を結論として受け入れ、だから手を引きました。学界全体

が手を引いた案件を、20歳の大学生が再び取り上げることにしたのです。彼の信念には名前がありました。接続主義。知能を記号と規則の体系ではなく、多数の小さな単位が結合の強度を変えながら作り出す現象として見る観点です。1980年の学界の地形の中で、この言葉を真摯に口にすることは、キャリアを賭ける意思に他なりませんでした。

書架から彼を絡めとった論文がもう一篇ありました。日本の福島邦彦が発表したネオコグニトロンでした。西洋の学界が神経ネットワークを埋没させた後も、日本ではこの分野の研究が静かに続いていたのです。福島の構造は、デイビッド・ハベル(David Hubel)とトルステン・ウィーゼル(Torsten Wiesel)が猫の視覚皮質に微細電極を刺して明かした発見に触発されたものでした。二人の科学者は猫の脳で、短い線分の方向に反応する細胞と、その信号を集めて位置がわずかに異なってもおなじものとして認識する細胞が層層に積み重なっていることを見つけ、この功績で1981年ノーベル生理学賞を受けました。福島はその層構造を機械に移しました。特徴を抽出する層とその結果を集約して位置変化に耐える層を交互に重ねたのです。

ルクンはそこに、後年自分が完成させることになる畳み込みニューラルネットワークの下絵を見て取りました。ただし福島の機械には欠けたものがありました。ネットワーク全体を入力端から出力端まで一度に学習させる方法でした。層ごとに別々に手を入れなければならない機械は、自ら学ぶ機械とはいいたいものでした。誤差を出力側から入力側に逆流させ、すべての結合を一度に修正していく道。ちょうどその空白を埋めると決めた秘策を、彼はこの図書館の冬の午後に心に込めました。

グランゼコール代わりにESIEE

フランスで科学者として成功する道には定められた関門があります。高校を卒業した秀才たちはクラス・プレバトワール(prépa)と呼ばれる準備課程に入り、2年から3年間、朝から晩まで数学と物理を猛烈に鍛えます。睡眠を削り、順位を争いながらコンクールと呼ばれる選抜試験の準備をするこの時期を、フランス人は通常、人生でもっとも厳しい時代として記憶しています。その関門を通過した少数者だけが、グランゼコール(Grandes Écoles)、つまり最高エリート養成機関に入学します。エコール・ポリテクニクやエコール・ノルマル・シュペリウールといった名門がその頂点にあり、フランスの大臣や最高経営責任者や数多の学者がその門を通過してきました。フランスの知識人社会への正門がそこにはありました。

ルクンはその正門から入りませんでした。前述のように、学校の数学の成績表において彼は輝く生徒ではありませんでした。高校の数学教師は、彼が大学で数学を専攻するには才能が不足していると言い切りました。ディープラーニングの数学的基礎を確立してチューリング賞を受けることになる少年に対する、教師が残した最後の数学評価はそれでした。この言葉は彼を座り込ませるのではなく、方向を変えさせました。彼は黒板上の抽象ではなく、機械を実際に動かす電気工学を選びました。1978年、予備課程を経ずに高卒後の試験だけで直接入学できる学校に進学しました。パリ高等電子電気工学院(ESIEE Paris)でした。

後年、ルクンはこの迂回を後悔しませんでした。むしろ予備課程を経ずにいたことが、科学者になるうえで決定的だったと何度も述べました。予備課程は学生を正答をよく訓練させますが、定められた枠に身を合わせさせてもします。他者が出した問題を他者より速く解く訓練と、誰も問わなかった問題を自ら見つけ出す訓練は、異なる筋肉を使います。その訓練を飛び越したおかげで、彼は主流が正しいと信じることに疑問を投げかける筋肉を失いませんでした。世間が神経ネットワークを死に道と呼んだとき、ただ一人その道を掘り進むことができた異端児気質は、もしかしたらこの迂回から始まったのかも

しれません。ESIEEは、そのような彼に相応しい場所でした。水曜日の午後をまるごと空けてくれた自由度、講義室外の探究を尊重する学風は、正統的なコースが与えない種類の自由でした。年月を経た後、ESIEEはこの卒業生の中で世界でもっとも有名な人物に名誉教授の称号を授与し、その縁を記念しました。

学校が与えたのは時間だけではありませんでした。当時、学生が触りにくかったミニコンピュータを、学校が使用させてくれたのです。そしてその自由の中で、ルクンは同志となる一人に出会いました。同じ学科のフィリップ・メツ(Philippe Metsu)でした。メツは、子どもが成長するにつれ認知が発達するプロセス、つまりピアジェが生涯観察した、まさにそのテーマに関心が高かったのです。電気工学科の建物の中で児童心理学について話し合える相手は、お互いがほぼ唯一でした。知能がどのようにして自ら構築されるのかという問いを共有した二人は、学校のミニコンピュータ一台にしがみついて、夜明けをししました。

彼らが作ろうとしたのは、人間が規則を与えなくても、データから表象を層層に積み上げる機械でした。パーセプトロンの学習原理とピアジェが観察した発達の段階を、コンピュータプログラムの中に一緒に移そうとする試みでした。彼らはそれを階層学習機械、略してHLMと呼びました。皮肉にも、フランスではHLMは庶民向けの公営住宅を指す略称でもあります。貧素な名前の機械は名に恥じない働きをしました。派手な装備なく、学校の計算機室の夜間を舞台に動作したのですから。二人の学生はシミュレーションプログラムを一行ずつ書いてコンピュータに入力しました。昼間は授業を受け、夜間は計算機室でコードを修正し、翌日に結果を確認する日々が続きました。そしてあるとき、素朴な緑色の画面の中で、結合の強度がデータを受け取り、非常にゆっくりと、しかし明確に自ら整理していく光景が現れました。それを初めて目で確認したとき、ルクンは自分が追い求めていたものが幻想ではないことに気づきました。ただし、二人の学生の機械はまだ半分の成功に過ぎませんでした。層を複数積み上げると、学習がしばしば頓挫しました。図書館で心に込めた問い、複数の層を一度に学習させる方法が依然として空白のままだったからです。

1983年の夏、ルクンはESIEEを卒業し、エンジニアの学位を取得しました。二十三歳のとき、同期生たちは就職先の企業の名前を交わし合っていました。フランスの電子産業と通信産業が若いエンジニアを十分に受け入れていた時代だったのです。手に持っていたのは他人と同じ卒業証書でしたが、心に抱いていたのは他人にはない問題が一つでした。まだ解いていないその問題。進学先を選ぶ基準もその問題一つでした。名門よりも、神経ネットワークを続けて扱える職場。フランスではそのような職場は広くありませんでした。最初の冬の寒気はフランスの学界にもそのまま降りかかっている、神経ネットワークを正面から扱う研究室は数えるほどでした。彼が見つけた場所は、大学にも国立研究機関にも正式には属さず、自己組織化とネットワーク動力学を手放さない非主流の研究者たちの集まりでした。その名前はネットワーク動力学研究室(LDR)。物理学と工学の境界で、相互に結合された単位が集団で動くときに生じる秩序を思案していた少数の人々でした。正規の予算も立派な装備もなく、借りた部屋に古いコンピュータ一台が生活の全てでした。看板は素朴でしたが、フランスで神経ネットワークと似たものを正面から扱う場所は、実質的にそこだけでした。

エリートコースの目で見れば、素朴な選択でした。グランゼコールの正門を通ってきた同年代の人々が大企業と国家機関の広い階段を上がっていた時代に、彼は誰も認めない部屋で、誰も信じてないテーマを選びました。しかし、その部屋には商業的な圧力も、流行に従えという助言もありませんでした。二十三歳の青年は、その貧素な自由の中で、誤差を逆流させる方法を手で導き出し始めました。

振り返ってみると、この初章の場面たちは一つのパターンをなしています。正門の代わりに脇門、流行の代わりに書庫、他人の評価の代わりに自分の手の感覚。彼は生涯、このパターンを繰り返すことになります。皆がシンボリズムに向かうとき神経ネットワークへ、後年皆が大規模言語モデルに殺到するとき別の道へ。冬に神経ネットワークを選んだ少年の物語は、今正に本論へ入ろうとしていました。

弁護士

kimkj.com

第2章 逆流する誤差

ヤン・ルクン、人工知能の父

古いコンピューター一台の博士課程

部屋の真ん中にコンピューターが一台置かれていました。誰かが捨てたものを拾ってきたものでした。パリのサン・ナビエブ学校の古い建物、正式な許可なく勝手に入り込んで使っていた片隅の小部屋。予算もなく新しい機器ありませんでした。古い機械一台がすべてでした。二十三歳のヤン・ルクン (Yann LeCun) はその前に座って、誰も解けないと思われていた問題を解いてみることにしました。

1983年の夏、ESIEEを卒業したルクンがなぜこの小部屋を選んだのかは前の章で見ました。知能とは人が規則を書き込むことだという記号主義が講壇を占め、1969年のミンスキーとペーパートのパーセプトロン診断以降、ニューラルネットワークに流れていた研究費がすっかり枯れ果てた冬でした。自ら学ぶニューラルネットワークを研究すると言えば、たいていは気の毒だという目で見られました。しかし彼は冷え込んだその場所で、むしろ火種を見たのです。

彼をこの部屋に連れてきたネットワーク動力学研究室 (LDR) は、正式な予算もなく空いている建物の片隅を許可なく占有していたグループでした。ルクンはこの小さな要塞を足がかりに、ピエール・マリ・キュリ大学 (パリ6大学) の博士課程に名を登録しました。暖房も十分でないその部屋で、彼は昼も夜も画面の前に張り付いて、他人が答えをすべて知っていると思う問題ではなく、すでに解けないと判定された問題に若さを賭けたのです。

お金はいつも足りませんでした。1984年にESIEEから奨学金を受けてようやく息つく間が得られましたが、彼を指導する教授を見つけることが残されていました。ニューラルネットワークを研究しようとする学生を引き受けようという正式な教授はめったにいませんでした。身近でその才能を認めた人はパリ5大学のフランソワーズ・フォゲルマン・スリエ (Françoise Fogelman-Soulié) でした。彼女は数学に明るく、ルクンの導出を一つ一つ見てくれましたが、フランスの規定上、博士指導を引き受けるのに必要な資格をまだ備えていなかったのです。そこでルクンは別の教授を探しました。

コンピエーニュ工科大学のモーリス・ミルグラム (Maurice Milgram) が指導を引き受けることになりました。ミルグラムは承諾しながら、率直な言葉を付け加えました。彼は人工ニューラルネットワークを一文も知らないので、数学的な部分では役に立つことができないだろうと言ったのです。奇妙な話に聞こえるかもしれませんが。何も知らない指導教授だなんて。しかし、ルクンにとってはそちらの方が良かったのです。誰も口を出さない場所、古いコンピューターと自由がある場所。それが彼が望んでいたことでした。

彼が取り組んだ問題はこのようなものでした。ニューラルネットワークを複数の層で積み重ねると、中に隠れた層が生まれます。入力でもなく出力でもない、内側に埋め込まれた層です。この隠れた層が何を担当すべきかは、誰も書き込んでくれません。では機械はこの層の値をどのように自ら修正していくことができるのでしょうか。1960年代にフランク・ローゼンブラット (Frank Rosenblatt) が作った一層のパーセプトロンは、やや複雑な問題の前で崩壊し、複数層を一度に教える方法は世の中にないと言われていました。

ルクンは思いがけない場所から糸口を持ってきました。ロケットの軌道を計算し機械を制御する工学には、ラグランジュ乗数法 (Lagrange multipliers) という古い数学がありました。目標に到達する最善

の経路を偏微分で導き出す方法です。ルクンはこの道具をニューラルネットワークに移し替えました。すると一つの対称性が見えました。ロケットが大気と重力を貫いて軌道に乗る過程と、ニューラルネットワークの隠れた層がデータから有用な特徴を見つけて重みを少しずつ修正していく過程が同じ公式を共有していたのです。

この対称性をもう少し説明してみましょう。ニューラルネットワークが手書きの3を5と誤って読んだとしましょう。正解と誤解の間には誤差が生じます。この誤差の責任をどの重みがどの程度負うべきかを知ることで、その重みを正すことができます。出力に近い層は責任を追及しやすいです。問題は深く埋め込まれた内側の層です。ルクンが借りてきた数学は、まさにこの責任を出力から入力側へ逆向きに配分する計算でした。微分の連鎖法則を層ごとに遡及適用すれば、最も内側の重みにまで自分の分の誤差が正確に配達されます。誤差が前へ流れるのではなく後ろへ流れる。この章のタイトルはそこから来たのです。

彼はこの導出を一つの学習機械としてまとめました。ESIEE時代にメトゥと一緒に階層学習機械 (HLM) と呼んでいたその未完成の機械を、今は小部屋で一人でさらに推し進めたのです。ニューラルネットワークの誤差を出力から入力側へ逆向きに流し、各層がどの程度間違っているかを遡及して重みを修正する方法でした。誤差が逆向きに流れる。後年、この原理は誤差逆伝播 (backpropagation) という名前で深層学習の心臓となります。二十幾つのフランス大学院生が捨てられたコンピュータの前で手作業でその原型を磨いている最中でした。

当時のコンピュータは小数点計算が遅かったです。ルクンは妥協案を見つけました。ニューラルネットワークが授受する値を0と1の二進数にしながらも、最初から最後まで続く重みの微分は完全に保たれるようにするという小技をプログラムに一つ一つ埋め込みました。はんだ付けするように一行ずつ貼り付けていく作業でした。外から見れば貧しい部屋の貧しい機械でした。内側では、ニューラルネットワークを複数の層で教える方法が少しずつ形を整えていたのです。

この導出が初めて人々の前に出たのは翌年でした。誰も顧みなかった小部屋の公式が、アルプスの雪に覆われた山小屋とパリのある学会場を経由して、思いがけない人々と出会うことになります。その最初の舞台がコグニティバでした。

コグニティバで投げられた逆伝播

1985年6月、パリのある講演会場。演壇の上にはカナダのトロント大学から来たジェフリー・ヒント (Geoffrey Hinton) が立っていました。学際的学会コグニティバ (Cognitiva 85) の基調講演でした。ヒントンはボルツマン・マシン (Boltzmann Machine) を説明していました。確率とエネルギーでニューラルネットワークが世界の構造を自ら学ぶようにするモデルでした。講演が終わると、ヨーロッパ各地から来た学者五十人ほどが彼を取り囲みました。その奥の方で、二十五歳の大学院生が一人、人々の肩越しにヒントンの後頭部だけを見つめていました。

数ヶ月前、ルクンはこの学会に自分の最初の論文を提出していました。フランス語で書かれた短い文章でした。タイトルは『非対称閾値ニューラルネットワークのための学習手順 (Une procédure d'apprentissage pour réseau à seuil asymétrique) 』。格式ばった学術論文というより、公式を並べた素朴な原稿に近かったのですが、その中には彼が小部屋で導出したHLMの逆方向微分が含まれていました。フランス学界でニューラルネットワーク研究がほぼ埋葬された冬、この場所はルクンが自分の信念を世界に初めて投げかけた祭壇だったのです。

論文を書く間に、ルクンは自分が作った公式が見覚えのあるものだと感じました。1960年代に宇宙船の軌道を扱っていた最適制御理論家たちがすでに同じ骨組みを作っていたのです。アーサー・ブライソン (Arthur Bryson) とユー・チー・ホー (Yu-Chi Ho) が1969年に出した『応用最適制御 (Applied Optimal Control)』には、誤差を逆向きに伝播させる公式がはっきり書かれていました。工学ではこれを共役状態法と呼んでいました。ただし彼らはこの公式をロケットの燃料を節約するためにだけ使いました。ニューラルネットワークの数千個のパラメータを自ら学習させる道具に応用する考えは思いつきませんでした。同じ数学を手握りしながらも別の世界を見たのです。ルクンはその別の世界を見たのです。

このくだりはルクンという人物をよく示しています。彼は完全に新しいものを無から作り出したのではありませんでした。ロケット軌道の古い数学と、脳細胞の接続を模倣したニューラルネットワークを、誰も一緒にしようと考えなかった場所で一緒にしたのです。ローゼンブラットの忘れられた学習機械を図書館で発掘した少年が、今度は制御工学の引き出しから道具を取り出してきたのです。他人が捨てたものと他人が忘れたものを拾い集めて新しい場所に置くこと。古いコンピューター一台で研究を始めた人にとっては、それがむしろ自然な方法だったのかもしれませんが。

その日の講演会場に戻ってみましょう。ヒントンは人々の質問を受け流しているうち、ふと身を翻して学会を率いていた哲学者ダニエル・アンドレ (Daniel Andler) の袖をつかみました。そしてフランス式ぶりにはっきりと一つの名前を発音しました。この場にヤン・ルクンという若い友人がいるかどうかを尋ねたのです。群衆の奥の方でその名前が呼ばれるのを聞いたルクンは、手をぱっと上げました。ここにいます (Je suis là)。彼がフランス語で叫びました。小部屋で一人で公式を磨いていた二十五歳に対して、世界の注視を浴びている巨匠が自分の名前を先と呼んだ瞬間でした。誰も読まないと思われていたフランス語の論文を、大西洋の向こうから来た人が夜通し読んで、その中の数学を理解したのです。冬を耐え抜く人にとって、このような確認はめったに訪れません。

ヒントンはどのようにしてこの無名の大学院生を見つけ出したのでしょうか。前の夜、彼は学会の論文集をめくっていて、ルクンの文章で足が止まりました。フランス語は拙く読めましたが、公式ばかりは誤解のしようがありませんでした。それはヒントン自身が扱っていた逆伝播と同じ原理だったのです。その上、ヒントンはすでにその名前を聞いたことがありました。数ヶ月前、アルプスのあるワークショップで、同僚の一人がフランスの古い実験室に、私たちと全く同じ公式を一人で導出した学生がいると電話で知らせてきたのです。そのワークショップの話は次の節に延ばすことにします。

二人は翌日の昼食を一緒に食べました。パリの薄汚い北アフリカの食堂で、湯気の上がるクスクスの皿を前にしてのことでした。ヒントンは自分が取り組んでいた逆伝播が実際にどのように機能するのかを説き明かし、ルクンはラグランジュ乗数法で同じ結論に到達した自分の導出を展開して見せました。一方は物理と確率から、一方は制御工学から出発して同じ場所に到達していました。ニューラルネットワークを死んだ分野と呼んでいた時代に、二つの大陸の二人の異端者が同じテーブルで互いを理解し合ったのです。

ヒントンにとっても、この出会いは格別なものでした。彼もまた学界の主流から長く押し出されていた人物だったのです。心理学から始まり神経網に移った彼は、人間の心が規則の羅列ではなく、多くの接続のパターンで成り立っていると信じていました。そうした信念を持つ研究者は、どの国でも少数派でした。その少数派が、パリのレストランでフランスの大学院生ひとりを新しい同志として得たのです。年齢も大陸も言語も異なっていましたが、二人が世界を見る角度は同じでした。この昼食の後、二人は生涯をともにします。30年余り後、ともにチューリング賞を受け、さらに数年後、人工知能の危険性を

めぐって相反する立場に立ち公開的に対立するときまでも、この友情は断たれることはありません。

ここで一旦止まって、この美しい物語の裏側も見ておかなければなりません。誤差逆伝播が誰のものが、という問いをめぐって、学界には長く棘のある論争がありました。後にルクンの最も鋭い批評者となるユルゲン・シュミットフーバー (Jürgen Schmidhuber) がその先陣を切りました。彼の主張はこうです。神経網の誤差を逆方向に伝播させる現代的アルゴリズムは、ヒントンでもルクンでもなく、1970年にフィンランドの数学者セッポ・リンナインマー (Seppo Linnainmaa) が最初に発表したのだということです。複数の層を積み重ねて自ら学ぶ神経網の原型も、1965年にウクライナのアレクセイ・イヴァフネンコ (Alexey Ivakhnenko) が既に作成し、実践に用いたと述べました。

シュミットフーバーは、1986年以降、世界を揺るがせた論文の数々が、実は15年、20年前に他人が築いた発見を、計算資源が豊かになった時代に再確認したにすぎないと考えました。そのうえで、元の発明者たちの文献を引用から除外したと告発しました。この論争は、後に登場するチューリング賞の箇所で再び大きく問題になります。ここでは、こう記すにとどめましょう。ルクンが小さな部屋で独りで導き出したのは事実ですが、彼が導き出したその原理をより先に手にしていた人たちも、確かに存在していました。二つの物語は互いに消し合うことはありません。

コグニティバから飛び散った火花は、翌年の夏、アメリカに飛び火しました。1986年7月、カーネギーメロン大学で2週間開催された結合主義モデル夏期学校のことです。ヒントンとデイヴィッド・ルメルハルト (David Rumelhart)、テリー・セイノフスキー (Terry Sejnowski) が導くこのすべてが、神経網の復活を狙う世界の異端者たちが初めてひとつところに集まった隠れ場所のような場所でした。その参加学生の中に、後にルクン、ヒントンとともにチューリング賞を受けるモンリオールのヨシュア・ベンジオ (Yoshua Bengio) がいました。

ルクンはその場所で、すっかり萎縮していました。英語が下手だったからです。数式はどの言葉でも同じですが、休み時間の雑談と冗談はそうではありませんでした。ハーバード出身の社交的な天才マイケル・ジョーダン (Michael I. Jordan) が、アコースティックギターを肩にかけて現れ、フランスの国民的歌手ジョルジュ・ブラッサンス (Georges Brassens) のシャンソンを歌ってくれました。ルクンよりもフランス語を流暢に話すアメリカの青年が歌ったフランスの歌でした。見知らぬ土地での不器用な異邦人に、その歌は言葉よりも先に届く歓迎だったのでしょう。夜はそのようにして深まり、冬を耐えていた人々の連帯は、少しずつ固くなっていきました。

この夏期学校に集まった面々を、ここで一度押さえておくのが良いでしょう。ベンジオとジョーダンのほかにも、後に人間の脳の読む回路を明らかにするスタニスラス・ドゥアネ (Stanislas Dehaene)、認知科学の大家ジェイ・マクレランド (Jay McClelland) が学生と教師として混在していました。後々それぞれの分野で大きな名をなす人たちが、まだ誰も注目しない主題に取り組み、山奥のキャンパスに集まっていたのです。その数日の間に、ヒントンはルクンに博士号を終えたらトロントに来るよう勧めました。小さな部屋から始まった一人の研究が、今や大陸を越える道を得たのです。

レ・ ホツシュの夜

「私の職業人生の大きな転機は、1985年2月、アルプスのレ・ホツシュで開かれたシンポジウムで始まった。」ルクンは自伝『機械が学ぶとき (Quand la machine apprend) 』でそう書いています。前の節でヒントンがルクンを見つけたコグニティバが6月のことだとすれば、その出会いを可能にした本当の出来事は、4か月さかのぼったこの冬の夜にありました。時を少し遡ってみましょう。

レ・ ホッシュ (Les Houches) は、フランス・ アルプスの雪に覆われたリゾート地です。昔から物理学者たちが集まって講義を聞き議論してきた、学問の聖地でもありました。1985年2月、ここで物理学と認知科学と神経科学を包括するワークショップが開かれました。この催しを実現させた人たちが、まさにルクンの側にいた人たちでした。彼を支援していたボゲルマン=スリエ、彼女の夫でパリ高等師範学校の物理学者ジェラルド・ ワイスブッフ (Gérard Weisbuch)、フランス国立科学研究センターの理論神経生物学者エリ・ ビネンストック (Élie Bienenstock) が力を合わせて、この学際的な集まりを組織したのです。

なぜ物理学者たちが神経網の話が聞ける山小屋へ殺到したのでしょうか。それは3年前の1982年に、ジョン・ ホップフィールド (John Hopfield) が一本の橋を架けたからです。彼は、固体物理学でスピン・ グラス模型という、磁性物質の無秩序な配列を扱う概念と、神経網が記憶を想起する連想構造が同じ数学で結ばれていることを示しました。磁石の中の電子がお互いに引き合い、押し合いながら安定した配列へ沈み込むように、神経網も混乱した入力から出発して一つの記憶へと沈み込むという図式でした。物理学者が熟知する言葉で神経網を書き直したわけです。それまで神経網に言及すれば詭弁扱いされていた物理学界の枷が、この論文で解かれました。物理学者、数学者、脳科学者、心理学者がアルプスの冬の夜へ一緒に駆けつけたのは、そのためだったのです。冬が終わったからではありません。冬が堪えられなかった人たちが、一時的にひとところに集まったのです。

講演会の夜が深まるころ、大学院生ルクンがこれら錚々たる学者たちの前に立ちました。彼は粉筆を握り、黒板に連鎖律の逆向き微分を書き下ろしました。自分のHLMと、最初から最後まで続く多層神経網の最適化について説明しました。米国ニュージャージーのAT&Tベル研究所から来た人々が席に着いていました。その中に、適応型システム研究を率いるラリー・ ジャッケル (Larry Jackel) がいました。無名の若い学生が黒板に書いていく幾何学的直観に、ベル研究所の人々は目を見張りました。3年後にルクンをその研究所へ招くきっかけが、この夜に初めて結ばれたのです。

しかし、レ・ ホッシュが残した本当の足跡は、講演会ではなく休憩時間につけられました。白い雪景色が見下ろせる午後のお茶の時間のことです。発表を終えたルクンは、紅茶のカップを手に、ロビーの奥に立っていた一人に近づきました。若い生物物理学者、テリー・ セイノフスキーでした。セイノフスキーは、すぐ前の年にヒントんとともにボルツマン・ マシンを発表した本人でした。外部から規則を与えなくても、神経網の隠れたユニットが世界の構造を自ら発見するようにさせるモデルでした。

ルクンはポケットからしわくちゃなノート帳を取り出しました。そこには、彼が小さな部屋で導き出した逆向き連鎖律と最適制御の式が書かれていました。セイノフスキーはお茶を飲みながら、その式を一行ずつ眺めました。そして、唾を飲み込みました。ルクンが独りで手で磨き出したその式が、セイノフスキーとヒントンが世界に発表しようと準備していた逆伝播と、鏡のように完全に一致していたからです。

ヒントンの逆伝播は、順調に得られた成功ではありませんでした。1982年、カリフォルニア大学サンディエゴのデイヴィッド・ ルメルハルトが、逆伝播をコードに変えようとしたのですが、プログラムの欠陥に阻まれ、手を引きました。その未完のコードを受け取ったヒントンが、高価なLispマシンを夜通し回転させながら、底から書き直し、ついに機械が複数層の重みを自ら調整する初の運用に成功しました。まだ世界に発表する前のことでした。ルメルハルト、ヒントン、ロナルド・ ウィリアムズ (Ronald Williams) が逆伝播を正式な論文として世界に発表するのは、翌1986年のことです。つまり、レ・ ホッシュのそのお茶の時間は、ヒントン側の発表が出る前だったのです。セイノフスキーは、その未発表の突破を既に独りで築いていた学生が、自分の目の前に立っていることに気づきました。

ここにシュミットフーバーが投げかけた問いを重ね合わせると、物語は立体になります。同じ式をパリの小部屋でルクンが導き、トロントでヒントンのコードとして生き返らせ、カリフォルニアでルメルハルトが先に試みた後、手を引きました。そしてそれより前に、フィンランドとウクライナにも原型がありました。一つの発見に複数の名前がつく。科学史では珍しくありません。複数の場所で同時に浮かび上がるということは、その発見の 때가熟したという意味でもあります。神経網を複数層で教える数学は、1980年代半ばに至って、あちこちで同時に沸き起こるほど熟していたのです。

セイノフスキーはルクンの説明を、最後まで根気強く聞いてくれました。ワークショップが終わってアメリカに戻った彼は、空港に着くなり、トロントのヒントンの電話をかけました。パリの貧しい実験室で、自分たちと全く同じ逆伝播の式を独りで導き出して、ぼろぼろのコンピュータで走らせている学生がいることを伝えました。彼が言ったそのフランスの学生の名前が、4か月後のコグニティバの講演会でヒントンの口から呼ばれることになります。前の節のあの場面です。

冬の学問はこうして動きました。研究費も学科も学術誌も等しく行き届かない場所で、人々は山小屋のお茶の時間と空港の電話と学会場での袖つかみで、互いを見つけました。制度が空けた場所を、友情と口伝えが埋めたのです。セイノフスキーがヒントんに伝えた「そのフランスの学生」という一言が、4か月後の呼び出しに、翌年の招待に、さらに数年後の大西洋を渡る移住へとつながりました。大きな制度に押し出された人たちが、小さな縁で互いを引き寄せたわけです。

レ・ホシユーの夜は、ローゼンブラット以降、忘れられかけていた信念ひとつを呼び戻しました。知能は生まれつきのものではなく、学んでいくものだという信念です。パリの片隅で一人で抱いていた炎が、アルプスの小屋でセイノフスキーに出会い、セイノフスキーの電話を通じて大西洋を越え、ヒントんに届きました。誰もニューラルネットワークを信じていなかった冬に、同じものを信じていた人々が互いに認め合った夜だったのです。

トロントのポスドク研究員

1987年6月、博士号審査会場。ヤン・ルクンは両脇に松葉杖をついて審査委員たちの前に立ちました。2か月前の4月に事故がありました。彼は砂浜の上を風で滑る三角帆のソリを手作りして乗っていたところ、無謀な速度で走った末に足首の骨が砕けました。他人の手を借りず、何でも自分で作ってみる人でした。ニューラルネットワークを学習させるソフトウェアも、ロケット工学を応用した学習機械も、風で走るソリも、彼にとってはみな同じ種類の仕事でした。頭に描いたものを手で作り出し、実際に動くかどうかを確認する仕事。その習性がこの日は松葉杖で戻ってきたわけです。人生を左右する審査を松葉杖に頼りながら受けるというのが彼らしかったのです。

審査委員席には重々しい表情の人々が座っていました。フランスの記号主義人工知能の大家ジャック・ピトラ(Jacques Pitrat)、彼を非公式に指導してきたフォーゲルマンスリー、公式な指導教授ミルグラム、そしてカナダから飛んできたヒントンでした。記号主義の大家と接続主義の大家が一つのテーブルに向かい合って座った場所でした。知能は人間が書き込んだルールだという立場と、知能は機械がデータから自ら学ぶものだという立場。フランス人工知能の二つの流れが、松葉杖をついた青年の論文を間に挟んで出会ったわけです。ルクンは多層ニューラルネットワークの逆伝播の完全性を数学で証明して見せ、審査委員全員が立ち上がって拍手を送りました。フランスの周辺で孤独に耐えていた異端は、そうして正式な博士となったのです。

1か月後の1987年7月、ルクンはカナダのトロント行きの片道飛行機に乗りました。傍には妻イザベル(Isabelle)と、1歳になったばかりの赤ちゃんがいました。イザベルは薬剤師として積み上げてきた自

分のキャリアを手放し、子どもの世話をすることにしました。一人の研究が前へ進む場所には、大抵このように後ろに引いた別の人の決定と一緒にあります。伝記は前へ進んだ人の名前だけを長く覚えていますが、その名前の傍の決定も物語の一部です。出発する時点では、この生活が1年を超えるとは思わなかったと言われています。しかし、この転居がルクンの人生と人工知能の歴史の流れと一緒に変えてしまうことになります。

トロントで、彼はヒントン研究室のポスドク研究員になりました。1980年代後半の学界は、相変わらず記号主義にだけ研究費をつぎ込んでいて、ヒントン研究室は、行き場を失った接続主義者たちが最後の希望を抱いて集まった避難所のような場所でした。ルクンはここで、フランス時代に関係を持った若き数学者レオン・ボットウ(Léon Bottou)と息を合わせ始めました。

二人の出会いは、ルクンが博士号を受ける直前まで遡ります。名門エコール・ポリテクニク(Ecole Polytechnique)の最終学年だったボットウは、まだ博士号も持たない無名のルクンと卒業研究を一緒にしようと名乗り出ました。ルクンは自分が構想していたニューラルネットワーク学習フレームワークの核心、すなわちLisp言語ベースのシミュレーターを一から書き上げるつもりでいて、その重い仕事を今卒業したばかりのボットウに丸ごと任せました。ボットウは、複雑なメモリ管理と動的バインディングまで組み込んだLisp処理系を、一人でコーディングして3週間で構築しました。

二人の協働には、特別な道具がありました。大学の高価なワークステーションではなく、各自が個人的に持っていたコモドール(Commodore)のアミガ(Amiga)コンピューターでした。当時一般的だったIBM PCやマッキントッシュと異なり、アミガは高解像度グラフィックスとマルチタスティング、仮想メモリを備えており、高価なUNIXワークステーションと似た環境を提供していました。ルクンとボットウはアミガ上でC言語でコードを書き、オープンソースのコンパイラgccと編集器Emacs(Emacs)を組み合わせて、軽くて高速なニューラルネットワーク開発環境を自ら構築しました。

問題は距離でした。ボットウはフランスにいて、ルクンはトロントにいました。二人は最新のソースコードをリアルタイムでやり取りするために、フランスが世界に先んじて普及させた通信端末ミニテル(Minitel)のモデムをそれぞれのアミガに接続しました。ミニテルはフランスの家庭の電話線に接続され、列車のチケットを購入したり電話番号を調べたりしていた小さな端末でした。二人はこの生活ツールを研究機器に改造しました。大西洋を挟んだ二人の工学者がアナログ電話線の上へ二進数を一行ずつ運んだわけです。こうして完成したフレームワークがSN(Simulator Neuronal)でした。ニューラルネットワークの各部品をブロックのように組み立て、学習に必要な微分を自動的に繋げるソフトウェアでした。後に事実上の標準となるパイトーチ(PyTorch)の思想的根源がここから始まります。今、世界の研究室が使うツールの遠い祖先が、個人用コンピューター2台と電話線1本で生まれたのです。

基礎が整うと、ルクンは長く温めていた実験に入りました。INRIA(INRIA)図書館で読んだ福島邦彦(Kunihiko Fukushima)のネオコグニトロン(Neocognitron)、その視覚認識構造に勾配降下法と逆伝播を移植する作業でした。階層的に視覚を処理しながらも、複数の層を一度に学習させる方法がなかったその構造に、片隅で磨いた自分の数学を、欠けた部分にはめ込むような作業でした。彼の研究は常にこのように互いに離れた部品を組み合わせるやり方で進みました。

問題はデータでした。今日のMNISTのような手書き文字データベースは当時はありませんでした。ルクンはアミガのマウスを握って画面上に文字を一つ一つ描き、原始的なデータを直接作成しました。ニューラルネットワークを教える教材を直接手で書いたわけです。この部分は些細に見えても、彼のやり方を改めて明らかにします。理論を立てる人と手を動かす人が別にいませんでした。式を導く手とマウス

で文字を描く手が同じ手でした。

データが足りないのを、ニューラルネットワークが容易に過学習に陥りました。わずかな例だけをまると暗記してしまい、新しい文字の前では力を発揮できない状態です。ルクンは自分が描いた文字を少しずつ傾けたり回転させたりサイズを変えたりして、一つから多くを作り出しました。データを人工的に増やすこの方法をデータ拡張(data augmentation)と呼びます。今ではディープラーニングの基本技法となった技法を、彼がデータがないから先に発明したわけです。

今、彼は3種類のニューラルネットワークを同じ条件で比較しました。すべてのノードが密に繋がった完全結合ニューラルネット、フィルターが狭い領域にだけ作用する局所結合ニューラルネット、そして空間不変性を強制的に組み込んだ重み共有ニューラルネットでした。空間不変性とは、画面の左にあらうと右にあらうと、同じ3は同じ3として認識するという特性です。人間の眼には当たり前すぎて、性質と呼ぶほどでもないように見えます。しかし機械にはそうではありません。左隅の3と右隅の3を全く別の画像として受け取れば、機械は3の位置ごとに3を新たに学習しなければなりません。重みを複数の場所で共有させれば、この無駄が消えます。一つの場所で3を認識する方法を学べば、すべての場所にその方法がそのまま適用されます。パラメータが大幅に減りながらも、位置に揺らがない眼が自動的に生じるわけです。実験結果、この重み共有ニューラルネットがパラメータを最も少なく使いながらも、過学習なく文字を最もはっきり識別しました。後に畳み込みニューラルネットワーク(convolutional neural network)と呼ばれるようになる構造の最初の実証でした。この構造がどこから来たのか、猫の視覚皮質とどのように繋がるのかは、次章でベル研究所の物語と共に展開されます。

同じ時期に、ヒントンは機械の限界を超える工夫をしていました。彼は画面のような2次元空間の代わりに、時間に沿って流れる音声信号に目を向けました。フィルターを時間軸で滑らせながら繋ぎ合わせる1次元構造、時間遅延ニューラルネットワーク(Time Delay Neural Network)を考案して音声認識で成果を上げました。一方は画像から、一方は音声から。二人は互いに異なる領域の上で同じ原理が通用することを証明しながら、学界の冷遇の中にあってもコネクショニズムの推進力を一緒に高めました。

1987年末、ルクンはモントリオールのマギル大学(McGill University)のセミナーに招待されて演壇に立ちました。彼はアミガ上で実装した重み共有理論と表現学習の成果を、黒板いっぱいの偏微分で説き明かしました。講演が終わると、最前列で息をのんで聞いていた若い修士課程の学生が手を上げて質問を浴びせました。多層ニューラルネットワークの微分最適化を深く理解していて、音や文脈のような時間信号の順序的因果関係をどう扱うかまで掘り下げる質問でした。

ルクンはその質問の鋭さに身震いを感じながら、学生の名前を心に刻みました。ヨシュア・ベンジオでした。後にルクン、ヒントンと共にディープラーニング革命を成し遂げ、チューリング賞を分け合う三人の縁が、このモントリオールの講演会場で初めて結ばれたのです。パリの古い片隅から始まり、アルプスの小屋とパリの学会会場、そしてトロントの研究室を歩いてきた誤差の逆伝播の旅は、今アメリカのベル研究所へ向かっていました。そこで彼は、紙の上の数式を超えて、米国の郵便物と小切手に書かれた手書き文字を実際に読み取る機械を作ることになります。冬はまだ終わっていませんでしたが、彼はすでに春に使う道具を手握っていたのです。

弁護士

kimkj.com

第3章 ホルムデルの視覚皮質

ヤン・ルクン、人工知能の父

金を節約して有名になることはない

1988年10月、ニュージャージー州ホルムデルのあるオフィスで、28歳のフランス人の青年が一枚の書類を手渡しました。高価なUNIXワークステーション一台を部門予算で購入してほしいという購入申請書でした。彼は慎重でした。そうするだけの理由がありました。

つい最近まで、ヤン・ルクンはパリ第6大学の貧しい研究室で、廃棄処分された中古コンピュータ台にしがみついていた。「ネットワーク動力学研究室」という名の小さな部屋には、神経網を信じる人がほとんどおらず、予算ももっとありませんでした。他人が捨てた機械をほぼ無断で占領するように使いながら、彼は多層神経網を学習させるプログラムを書きました。夜間、大学の他の部屋がすべて灯を消した後も、その一台は動いていました。そのような状況から、アメリカにたわずか数日で、大企業の金で最新機器を購入したいと手を出すことは、彼にとって見慣れず、恐ろしい行動でした。

申請書を受け取った人はローレンス・D・ジャッケル (Lawrence D. Jackel) でした。AT&Tベル研究所の適応型システム研究部門を率いていた彼は、ルクンの逆伝播公式と幾何学的直観を見出し、彼を大西洋の向こう側に連れてきた張本人でした。ジャッケルは気後れした青年の肩をたたいて、このように言いました。「ヤン、なぜもっぱらベル研究所なんだ。ここでは金を節約して有名になることはないんだ (Yeah, you know, why Bell Labs? You don't get famous by saving money) 。」後年、ルクンがディープラーニングの英雄たちを扱ったインタビューで何度も繰り返して語ったこの一言は、彼がこれからの十年をどのような空気の中で働くことになるかを前もって知らせる文章でした。

ベル研究所という名前には、それだけの重さがありました。トランジスタはここから生まれました。情報理論はここで生まれました。電話線を通る信号を扱っていたこの会社の研究所は、すぐにお金にならなくても基本を掘り下げる研究に人と金を長く提供する伝統を持っていました。ジャッケルの言葉はその伝統を一文で要約したものでした。ここでは節約は美德ではありませんでした。美德は誰も解けなかった問題を解くことでした。

その当時、学界の外の空気は正反対でした。1980年代後半は人工知能の最初の冬がちょうど真っ最中でした。その冬の根は二十年前まで遡ります。1969年、マーヴィン・ミンスキーとシーモア・パパートが出版した本一冊が、初期の神経網であるパーセプトロンの限界を数学で指摘し、その後、神経網研究に流れていた金と関心は引き潮のように引いていきました。主流の研究者たちは人工神経網を非科学的な詭弁ぐらいに見なしていました。知能とは脳の中に規則として埋め込まれているものであり、機械には人間がその規則を手で書き込まなければならないという考えが学会を支配していました。機械がデータを見て自ら学ぶという発想は、真摯な科学者が口に出すような話ではないという雰囲気でした。

ところがホルムデルのその部門の内部だけは違っていました。物理学者、数学者、電気工学者が一つの部屋に集まり、機械知能の本質を思う存分掘り下げました。外では死んだ道だと呼ばれていた神経網の火種を、この人たちは逆に大切に守っていました。学会で却下されたアイデアが、この廊下では一晩中の議論の話題になりました。真冬に立てられた小さな温室のような場所でした。ルクンがこの温室に足を踏み入れたことは、神経網を信じるという理由で学界から孤立していた人にとっては、初めて出会う故郷のような出来事でした。

孤立というのは誇張ではありませんでした。パリで神経網を勉強すると言ったとき、指導者を見つけることは簡単ではありませんでした。その分野はフランスにはほぼ存在しませんでした。ルクンが最初に自分の学習アルゴリズムを持って広い世界に出かけたとき、彼を認めてくれた数少ない人の一人がジェフリー・ヒントンでした。二人は神経網が嘲笑されていた時代と同じ言葉を使う珍しい同僚でした。ヒントンのトロント研究室で過ごしたポストドク時代は、ルクンにとって自分が狂った道を一人で歩いていないという確認でもありました。ベル研究所はその確認をはるかに大きな規模で繰り返して与えてくれました。

彼が温室に足を踏み入れた時点は、人工知能の流れが彼と正反対に流れていた時でした。その頃、産業と学界がお金をかけていたのは、いわゆるエキスパートシステムでした。医者や弁護士のような専門家の知識を人間が規則の形で引き出し、コンピュータに丁寧に書き込めば、機械がその規則に従って判断するという発想でした。世界を規則で書き出すことができるというこの信念が、その時代の主流でした。ルクンはその反対側に立っていました。彼の見方では、世界は規則で完全に書き出すことができるものではありませんでした。手書きの数字の3一つをとっても、人間が書く3の形は人間の数だけ違っていました。その3すべてを規則で書き出そうとする瞬間、規則は果てしなく増殖し、常に例外が残りました。規則を書く代わりに、例を示して機械が自ら学ぶようにする方が良いというのが彼の考えでした。

彼がこの温室でどうしてもしたいことは明白でした。パリ時代とトロント時代にかけて彼が自ら書いた神経網シミュレータを、ちゃんとした機械の上で実行することでした。彼は博士課程の時からアミガという個人用コンピュータで多層神経網を実験してきており、トロントのジェフリー・ヒントン研究室でポストドクとして過ごす間、そのソフトウェアをさらに磨き上げていました。個人用コンピュータ一台で押し進めることができる実験には明らかな限界がありました。画像一枚を神経網に通すのにも時間がかかり、ちょっと大きな画像を入れれば、機械が悲鳴を上げました。今、ベル研究所には大規模な画像演算に対応できる高性能ワークステーションがありました。手に握る道具が突然数段階良くなったわけです。

ジャッケルの支援は言葉だけではありませんでした。要求した機器はすぐに到着しました。ルクンはその機械の前に座り、ずっと前から彼の頭を離れない一つの絵をコンピュータの中に移し始めました。生物の眼と脳が世界を見る方法でした。彼は、人間が規則を書き込んでこそ機械が形を認識するという常識を信じていませんでした。彼が信じたのは反対側でした。形を認識する能力は教え込むのではなく、多くの例を見ながら機械が自ら学んでいくものだと、彼は考えていました。

この信念は、その日突然生まれたものではありませんでした。パリの一つの図書館で忘れ去られたパーセプトロンの論文を掘り起こしていた学生時代から、知能は生まれつきの規則ではなく、学んでいくものだという考えが彼の中に定着していました。学界が背を向けたその古い学習機械の中で、若きルクンは他の人が見なかった可能性を読み取っていました。ベル研究所は、その古い考えを実際の機械で検証する自由と機器を一挙に与えてくれました。金を節約するなというジャッケルの言葉は、費用の心配なく根本的な問いに取り組むという許可でもありました。

その質問はこのようなものでした。一匹の猫が何の説明も受けなくて世界の事物を認識するのに、なぜ機械は人間が一々特徴を書き込んでくれてはじめてようやく一つの数字を読むことができるのだろうか。この問いに答えようとして、ルクンが最初に開いた本はプログラミング教材ではありませんでした。三十年前、二人の神経生理学者が猫の脳を覗き込みながら残した記録でした。

ハブルとビーゼルの猫

物語は1962年のある実験室まで遡ります。ハーバード大学の神経生理学者デイビッド・ハブル (David Hubel) とトルステン・ビーゼル (Torsten Wiesel) は、麻酔をかけた猫の頭の後ろ、後頭葉の視覚皮質に非常に細い電極を挿し込みました。そして猫の目の前に、様々な角度に傾いた光の線を一本ずつ照らしました。スクリーンに明るい棒がこの方向で横たわると、どの神経細胞が目覚めるのか、その方向に回転すると、どの細胞が静かになるのかを、一つ一つ記録しました。

この実験には、偶然も一役買ったと言われています。二人は最初、点の形をした光を照らしながら細胞の反応を探していたのですが、なかなか素晴らしい反応が出ませんでした。ところがスライドを差し替えるとき、ガラス板の端が作った薄暗い線一本が画面をかすめると、細胞が激しく鳴きました。細胞は点ではなく線に、それも特定の角度の線に反応していたのです。この発見のため、二人は1981年のノーベル生理学医学賞を受賞しました。

振り返ってみると、この場面には不思議なところがあります。1980年代後半、人工知能をやろうとしていた若い人たちが論理学の教材と規則言語を掘り下げていたとき、ルクンは二十何年前の猫の脳を刺した実験記録に戻りました。機械を作ろうという人間がコンピュータより生物学を先に読んだのです。彼にとって知能のヒントは人間が書いた規則ではなく、自然が作った脳にありました。

彼らが明かにしたのは、脳が世界を見る順序でした。眼の網膜に映った光は、丸ごと脳に伝わりませんでした。まず網膜の細胞が光を集め、およそ100対1の比率で圧縮した後、百万本の神経線維を通じて、中継地点である外側膝状体を経由して、後頭部内側の初級視覚皮質に流入しました。しかし、その第一関門に位置する神経細胞は、視野全体を一度に見ていませんでした。細胞ごとに眼の前のごく狭い領域一つずつだけを担当していました。この狭い担当領域を、二人は局所受容野と呼びました。

さらに興味深いのは、その細胞が何に反応するかでした。これらは明るさそのものに反応していませんでした。自分の担当領域内に特定の角度に傾いた線や縁が現れたときだけ点灯しました。斜めの線を担当する細胞、垂直線を担当する細胞が別々にあったわけです。ハブルとビーゼルは、このように原始的な線と縁にのみ反応する第一段階の細胞にシンプルセル (simple cell) という名前をつけました。

この初期の細胞が捉えた断片は次の層へ進みました。その層には複合細胞 (complex cell) がありました。複合細胞は複数の初期細胞の反応を広い領域にわたってまとめました。そのおかげで予期しないことが起こりました。目の前の物体が左に少しずれても右に少し揺れても、複合細胞はあそこその縁がまだあると同じように認識しました。位置が少しずれても形の正体が揺らがないことが、複合細胞がしていた仕事でした。視覚情報はこのように網膜から出発して複数の層を順に上り、点から線へ、線から模様へ、模様から物体へと成長しました。脳科学者たちはこの経路を腹側視覚経路と呼びます。網膜から外側膝状体を経由して視覚皮質の複数の駅を通り、物体を認識する最終領域に至るまで、情報は層を上るほどますます抽象的な形へと整理されました。

ヒューベルとウィーセルの発見がルクンの心を捉えた理由は、それが知能についての古い信念を正面から覆したからです。記号主義者たちは形態認識を規則の問題だと考えました。目が三角形を見ると、3本の線が3つの頂点で出会うという規則に照らして見て、合致すれば三角形と判定するという具合です。しかし猫の脳の中には、そのような規則はどこにも書かれていませんでした。代わりに、狭い領域を担当する細胞が線をつかみ、その上層の細胞が線を集めて形態に育てる組立工程だけがありました。認識するとは規則を適用することではなく、小さなものをしっかりと積み上げることでした。ルクンはこの組立工程を機械に移すことができれば、規則を一行も書かずに世界を読む機械を作ることができると考えました。

ルンクンにとって、この階層構造は生物学的知識以上のものでした。それは一つの設計図でした。自然が数億年かけて磨き上げた、世界を認識する機械の図面でした。記号主義的人工知能は人間が規則を書き込まなければ機械は形態を知らないと言いましたが、猫の脳はそのような規則を誰からも学んだことはありませんでした。層状に積み重なった細胞が例を見ながら自ら調整されただけです。ルンクンが掘り下げたかったのはまさにその自律性でした。人間の手を借りずに世界を読み取る能力、それを機械の中に移すことができれば、人工知能は全く別の道へ進むことができたのです。

実のところ、この図面をコンピュータに移す試みがルンクンが最初ではありませんでした。1979年と1980年、日本の工学者福島邦彦 (Kunihiro Fukushima) がすでにその道を歩きました。彼が作ったネオコグニトロン (Neocognitron) は、ヒューベルとウィーセルの2つの細胞を模倣した人工神経網でした。線をとらえる層とそれをまとめる層を交互に積み重ね、上へいくほど解像度を下げていく構造でした。視覚皮質の地図を機械の中に初めて描き入れたわけですから、発想だけで言えば先進的な進展でした。

ところがネオコグニトロンには決定的なものが欠けていました。層全体の結合強度を一つの目標に合わせて端から端まで調整する方法がなかったのです。神経網がある数字を誤って読んだとき、その誤差を下層まで遡って送り、数千数万個の結合を一度に修正していく学習法、すなわち誤差逆伝播がなかったのです。福島は層ごとに個別に働く方式に頼るしかありませんでした。ある信号が頻繁に入力されると、それに関連した結合を強くするというやり方でしたが、このような方法は全体を一つの答えに統一できませんでした。そのため直線の数個は区別しましたが、郵便封筒の上に人それぞれが思い思いに書きなぐった手書き数字のような歪んだ形の前では崩れました。図面は優れていても、その図面通りに建てた家が自ら学ぶ方法を知らなかったのです。

ルンクンがベル研究所で越えようとした壁がまさにここでした。彼は福島が譲り受けた視覚皮質の階層構造はそっくり受け取りながら、福島が埋められなかった場所に誤差逆伝播という学習法を押し入れることにしました。線をとらえる細胞とそれをまとめる細胞の局所的結合はそのまま生かし、その上で神経網が自らの誤差を自ら計算しながら結合を調整していく仕組みを作ることでした。生物の目が与えた構造と微分が与えた学習を一身に付ける作業でした。どちらか一方だけでは不十分でした。構造だけあって学習がなければ福島の行き止まりになり、学習だけあって構造がなければ間もなく直面する数値の爆発に埋もれるからです。

この作業をするには越えなければならない実務の壁がもう一つありました。視覚皮質の原理をそのままコピーするだけでは、少し大きな画像一つにも、コンピュータが処理できない計算が降り注ぎました。自然の図面を機械が実際に耐えられる大きさに縮小する技術、ルンクンが次に取り組んだのがその問題でした。

重み共有と空間不変性

問題は数値の爆発でした。神経網を画像にそのままあてると、調整しなければならない値が処理不可能なほど膨らみました。

当時よく使われていた方法は完全結合神経網でした。入力画像のすべてのピクセルを次の層のすべてのニューロンと一対一で結び付ける構造です。一見すると密集していて堅牢に見えましたが、その中には落とし穴がありました。手のひらほどの16×16ピクセルの数字の画像を一つ入力しただけでも、結合しなければならない接続が瞬く間に数万個に膨れ上がりました。少し大きい画像を想像してみると事態は途方もなくなります。横256ピクセル、縦256ピクセルの写真一枚なら入力点だけで6万個を超え、そ

の後に層を一つ付け加えても調整する値が数千万個に跳ね上がりました。当時のコンピュータのメモリでは処理できない規模でした。

値が多いということは機械の負担に留まりませんでした。もっと深い病があとに続きました。調整する値が多すぎると、機械は学ぶ代わりに暗記してしまいました。訓練に使った画像は完璧に当たるのに、初めて見る画像の前では容赦なく間違える過学習が訪れました。学生が問題集の正答番号だけを丸ごと暗記すれば同じ問題は全て当たるが、少し変えると壊れるのと同じでした。データが少ないほど、この病は深く、1980年代後半のデータはいつも少なかったのです。

ルクンはこの呪いを力で解こうとしませんでした。彼は反対に考えました。調整する値を増やす代わりに減らそうというものでした。そのためには神経網の中に強い制約を事前に埋め込まなければなりませんでした。どんな結合も許可するのではなく、画像という対象が本来持っている性質に合った結合だけを残そうという発想でした。ここで彼は再び猫の視覚皮質に立ち戻りました。

最初の鍵は局所受容野でした。視覚皮質の細胞が視野全体ではなく狭い領域一つだけを担当していたことを思い出し、ルクンも人工ニューロンが入力画像全体を一度に受け取らないようにしました。ニューロン一つは、わずか5×5程度の小さな窓の中のピクセルとだけ結合させました。世界を一目で飲み込もうとするのではなく、小さな虫眼鏡で一隅ずつ覗き込ませたのです。これだけでも結合の数は大幅に減りました。

2番目の鍵は、彼の生涯を貫く着想でした。重み共有です。ルクンはこう問いかけました。画像の左下で縁をよく見つけられる虫眼鏡があれば、その同じ虫眼鏡は右上や真ん中でも縁をよく見つけ出さないだろうか。縦線は図のどの位置にあらうと縦線です。だとすれば、位置ごとに異なる虫眼鏡を別々に学習させる理由はありませんでした。虫眼鏡は一つだけ作り、その一つを画像全体の上でスライドさせながら同じ値を繰り返し使えば十分でした。

このスライドする演算にはすでに名前がありました。信号を扱う人々が昔から使ってきた畳み込み (convolution) です。一つの虫眼鏡、つまり一つのカーネルが画像全体を走査して作られた結果を、ルクンは特徴マップと呼びました。左上から右下まで同じ虫眼鏡がマスをずつつ移動しながら、ここに縁があるかと問い、その答えを一枚の地図に記していくわけです。効果はすぐに現れました。位置ごとに値を別々に置かず、一つを共有したので、機械が学ぶべき値の数が千分の一に減りました。過学習の呪いが払拭され、重い乗算に苦しんでいたコンピュータの負担も軽くなりました。

ここにはルクンらしい節約の美学がありました。金を惜しむなと言ったヤケルの言葉とは異なった、値の節約でした。彼は機械により多くの自由を与える代わりに、より少なく与えました。調整する値を減らすということは、機械が迷える範囲を狭くすることです。範囲が狭くなると、機械は少ない例を見ただけでも正しい答えを見つけることができました。自由が多いほど賢くなるだろうという推測は、ここで間違っていました。世界の真の構造に合わせて自由を削り落とすほど、機械はむしろより良く学びました。この逆説がルクンが生涯追求した信念の中心にありました。

重み共有が与えた贈り物はここで終わりませんでした。同じ虫眼鏡で画像全体を走査するため、探そうとする形がその画像の中で数ピクセル横にずれても、結果は揺らぎませんでした。数字の3が真ん中にあると右隅に寄っていようと、特徴マップ上ではその位置だけが並列に移動するだけで、3という形そのものは同じように捉えられました。位置が変わっても正体が保持されるこの性質を空間不変性と呼びます。前の節で見た複合細胞がしていた仕事を、ルクンは一つの式で神経網に移植したわけです。

彼はここに複合細胞の残りの半分をさらに加えました。プーリング、つまり隣接した反応をまとめて平均を取り、解像度を半分に減らす層でした。ルクンの判断はこうでした。神経網が上層に上がっていった物体の抽象的形態を捉えるようになるほど、その形態が正確に何番目のピクセルにあったかは役に立つどころか邪魔になるということです。数字の3を認識するのに重要なのは、上部の丸い曲線と中央の折れ曲がりが大体どのような位置関係にあるかであり、その筆画が画面の何番目のグリッドに正確に位置したかではありませんでした。だから彼は細かい位置情報を意図的に曖昧にしました。曖昧にするほど、機械は書体の歪みとノイズに強くなりました。畳み込みで特徴を抽出し、プーリングで位置を曖昧にするこの二つの操作を交互に積み重ねたのが、今日私たちが畳み込みニューラルネットワークと呼ぶ構造の骨組みです。

一層を通り抜けるたびに、機械が見えるものが変わりました。一番下の層では、短い線と辺を見えます。その上の層では、それらの線が出会ってできた曲がった曲線と折れ目を見えています。さらに上の層では、それらが集まって数字の一部、たとえば3の上側のループのような形を見えています。一番上の層に至ると、ついに3という数字全体が捉えられます。誰も3の見た目をルールとして書いて与えたことがないのに、層を上っていき、小さいものから大きいものへと積み重ねた末に、機械は3を認識しました。ハッブルとウィーゼルが猫の脳で見た、その組み立てのプロセスが、今やコンピュータの中で回転していたのです。

この設計の底には、世界を見るルクンの観点が敷かれていました。彼は、私たちが目で見ている世界が構成的であると考えていました。難しい言葉のようですが、意味は簡単です。ピクセルが集まって辺になり、辺が織り合わさって模様になり、模様が組み立てられて物体の部品になり、部品が合わさって一つの物体になるということです。顔は目と鼻と口からできており、目は再び曲線と点からできています。世界がこのように層々に積み重なってできているのなら、世界を読む機械も層々に積み重ねて作るべきだと、彼は信じていました。合成畳み込みとプーリングを何層も重ね合わせた構造は、その信念を機械に移したものだっただけです。

この箇所は、後にずっと続く論争の発端となることもありました。認知心理学者ゲーリー・マルクス (Gary Marcus) は、ルクンの合成畳み込みニューラルネットワークに対して、こう指摘しました。それは機械が白紙から自ら学んだ結果ではなく、人間が空間不変性と重み共有という知識を事前に埋め込んだからこそ得られた成果だということです。結局、先天的に与えられた装置があってこそ学習が成立するという話であり、それは知能が生まれつきのルールに頼っているという自分たちの側に肩を持つと、彼は見ていました。人間の手が完全に抜けたというルクンの主張は誇張だという反論が、マルクスの長年の立場でした。

ルクンの反論は落ち着いたものでした。彼は、人間の視覚も完璧な重み共有を生まれながらに持たないと述べました。人間の網膜は中央が密集していて辺縁が粗く、ばらつきのある構造だからです。彼がニューラルネットワークに埋め込んだ制約は、記号主義者たちがルール全体を書き込むやり方とは異なるというのが、彼の要旨でした。重み共有は世界の物理的対称性を認め、機械が調整する値を最も少なくする最小限の足がかりに過ぎず、その上で実際の形態は、あくまでも生のピクセルを見ながら自ら学ぶというものだったのです。どこまでが事前に埋め込まれたものであり、どこからが機械が学んだものなのか。この問いは、数十年後に大型言語モデルをめぐる論争の中で、ほぼ同じ顔で再び現れます。二人とも完全には間違っていなかったという点で、この論争はまだ癒合されないまま残っています。

論争は論争であり、1989年ホルムデルのホワイトボードの前で、ルクンが当面証明する必要があったのは別のことでした。この優雅な設計が紙の上でのみ美しいのか、それとも人間が書いた本物の郵便

番号を実際に読み取るのかということでした。

LeNet-1と郵便番号

ルクンがベル研究所に来てから約一か月後の1988年11月末、同じ部門のジョン・デンカー (John S. Denker) と同僚たちが、手書き郵便番号を読むニューラルネットワークを発表しました。性能は悪くありませんでした。しかし、その機械は完全に自ら学んだものではありませんでした。入力直後に置かれた前側の層は、人間が辺を捉えるように丁寧に設計して固定した値で満たされていたのです。機械がデータを見て学んだ場所ではなく、人間が事前に削り込んだ場所でした。

そうする理由がありました。その頃、コンピュータビジョンをしている人たちは、大体このように信じていました。機械が生みのピクセルを受け取り、最初から最後まで一人で形態を学ぶことは不可能であり、人間が手をかけて作った前処理フィルタがあってこそ、やっと高次元の形態を認識するというのです。デンカーのニューラルネットワークは、その常識とニューラルネットワークの間の妥協でした。前側は人間が設計し、後ろ側だけ機械に学ばせたのです。ルクンはこの妥協に同意しませんでした。彼は人間の手指を前後から完全に切り除き、誤差逆伝播だけで数千個の値を一度に調整する機械を実現してみせると心に決めました。前の項で磨いた局所受容野と重み共有を逆伝播ニューラルネットワークに丸ごと付けたこの機械が、人工知能史に残る最初の実用合成量込み込みニューラルネットワークLeNet-1だったのです。

性能をテストする舞台は、米国郵便公社が提供しました。米郵便公社とその契約業者がニューヨーク州バッファロー郵便局を通り過ぎた実際の手紙の封筒から手書きされた郵便番号をスキャンして集めた、手書き数字データベースです。ペンの種類はまちまちで、文字の大きさはばらばらであり、筆圧も人によって異なっていました。人間が読んでも戸惑うほど粗い実世界のデータでした。実験室で整理された数字ではなく、急いで書き殴られ、にじみ、つぶされた本当の人間の文字でした。どれほど粗いかというと、ベル研究所の2人の研究員ジェーン・ブロムリー (Jane Bromley) とエドゥアルド・サックキナー (Eduard Säckinger) が直接目で試験用の数字を読んでみた時、人間さえ平均2.5パーセント間違えたのです。機械に人間並みに読むよう求めても、それはもう高い要求でした。

ルクンが手に持ったデータは手書き数字9,298個であり、ここに35種類の活字フォントから抽出した数字3,349個が加わりました。彼はこの中から9,840個を訓練用に、2,707個をテスト用に分けました。テスト用に混ざった活字は訓練に一度も入れなかった馴染みのないフォントでした。学んだ図を見せて採点しようというわけではありませんでした。初めて見る図の前でどれだけ耐えるかを見ようというのでした。前述の過剰適合の落とし穴を避けたかどうかを確認する、正直なテスト問題でした。

機械の中を覗き込むと、このようになっています。封筒から剥がした数字は16かける16ピクセル、つまり256個の入力点として整理されて入力されました。その後ろに3つの隠れ層が順番に置かれました。第1層では5かける5サイズの小さな虫眼鏡がイメージをなぞりながら12枚の特徴マップを抽出しました。この層だけでニューロンが768個、つながりのある結合が19,968個もありましたが、同じ特徴マップ内のニューロンが虫眼鏡一つを共有していたおかげで、実際に調整する必要がある値は1,068個に圧縮されていました。次の層は解像度を半分に下げながら再び12枚のマップを作り、第3層はその応答を30個のユニットにまとめました。最後には0から9までの10個の出力が置かれ、この数字が何かの答えを出しました。

1989年のLeNet-1には、今日の教科書に出てくるような独立したプーリング層がまだ存在しませんでした。プーリング用の層を別途一つ増やすことが、当時のコンピュータにとって過度に重く、コストが

かかったからです。ルくんは虫眼鏡が滑る歩幅を大きく取ることで、特徴を抽出する動作と解像度を下げる動作を一つの層の中で同時に片付けてしまいました。ないなりの家計に合わせて設計を折り畳んだり広げたりした跡です。全体を数えてみると、ノードが1,256個、接続が64,660個でしたが、重み共有のおかげで機械が実際に学ぶ必要がある値は9,760個にまとめられていました。接続は6万個以上あるのに学ぶ値は1万個にも満たないこの格差が、ルくんが埋め込んだ制約が実際にどれほど大きく機能していたかを、そのまま示しています。

これらすべての計算を支えたのは、ルくんとレオン・ボトゥ (Léon Bottou) がアミガ時代から底から手ずから書いたソフトウェアでした。2人がSNと名付けたこのシミュレータは、ベル研究所のサン (SUN) -4/260ワークステーション上で夜通し動きました。後に世界を変えたディープラーニングツールたちの遠い祖先が、ここから始まったわけです。学習には確率的勾配降下法を使いました。訓練データ全体を一度に計算して値を直すという重い方式ではなく、数字が一つ入るたびにすぐに誤差を測り直して値を少しずつ弾く方式でした。手書きデータには似たような例が溢れていたため、全部を一度に計算することは無駄でした。一つずつ見ながら直していくこのやり方は、はるかに軽く、速く機械を正しい方向に導きました。LeNet-1は訓練用の数字9,840個をたった30回ほど見直すだけで、ほぼ完璧に分類しました。

学習を終えたニューラルネットワークは、紙の上の理論には止まりませんでした。完成した接続網は、ベル研究所が直接設計した浮動小数点演算チップ、AT&T DSP-32C上に移されました。カメラから入った手書き文字画像を捉えて整理し、このチップで計算し、0から9までどの数字であるかを画面に表示するまでにかかった時間は0.015秒でした。目を一度まばたきするうちに、機械が人間の文字を読み取ったのです。ソフトウェアが実験室の中で良く動くことと、そのソフトウェアが爪の大きさのチップ上でリアルタイムに郵便物を読み取ることは、異なる重さの仕事でした。人間が見ても戸惑う郵便番号のテスト問題の前で、LeNet-1は、どのような手作業フィルタの助けも受けずに、誤答率1パーセント、判断保留率9パーセントを記録しました。

判断保留率という数字には、実戦の知恵が込められていました。郵便物を自動で分類する機械が犯すことができる最悪のことは、間違った番号で自信を持って手紙を送ることです。そのため、ルくんのシステムは、確信が持てない時には、わからないと手を上げてその封筒を人間に返すように作られていました。10通のうち1通くらいを人間に返すかわりに、自分が読むと名乗った9通は100通に1度の割合でしか間違いませんでした。誤りを減らそうとあいまいなものを自ら振り分けるこの態度こそが、機械をおもちゃの域を超えて本当の郵便システムに組み込むことができるようにした鍵だったのです。

ルカンがこの成果で最も誇りに思ったのは、誤答率そのものではありませんでした。その低い誤答率を、人の手を借りずに得たという事実でした。デンカーのニューラルネットワークは、前方に人が手を加えたフィルターを置いて初めて数字を読み取りましたが、LeNet-1は封筒から切り取った生のピクセルをそのまま受け取り、出力まで一つの流れとして学習しました。入力から答えまで、人の介入なしに誤差一つで貫いて処理したという意味で、このような方式を端から端まで学ぶ学習と呼びます。この一つの流れこそ、ルカンが世に送り出したかった本当の主張でした。特徴を人が設計する時代が終わり、特徴までも機械がデータから自ら探し出す時代が始まるという宣言でした。

このいくつかの数字が意味するものは小さくありませんでした。学界全体がニューラルネットワークを死んだ道だと呼び、冬の中に身を潜めていたまさにその時、ルカンは、人が規則を書き入れなくても機械が生みのピクセルから形を自ら学ぶということを、実際に作動する機械で示しました。人間の設計者の偏見を前後から取り除き、微分一つと幾何学的な制約だけで手書き文字を読み取ったのです。彼がパ

リの図書館で抱いた考え、つまり知能は生まれつきの規則ではなく学んでいくものだという長年の信念が、ホルムデルの画面上で初めて目に見える証拠となりました。彼がここで鍛え上げたこの考えは、ずっと後の2018年、ヒントンやベンジオと並んで受賞したチューリング賞へとつながります。

しかし、この勝利には影がつきまとっていました。郵便番号を読み取ることに成功したニューラルネットワークが次に向かった先は米国の銀行であり、銀行が要求したのは封筒の上の数字5つではなく、小切手の上に続けて書かれた金額全体でした。一つずつの数字を当てることと、互にくっついて連なる文字を区切って読むことは、まったく別の問題でした。しかし、そのより大きな問題に近づく前に、ルカンはずまず自身のニューラルネットワークを適切な大きさに整えなければなりませんでした。ホルムデルの成功は、今やより難しい試験用紙を彼の前に突きつけていたのです。

弁護士

kimkj.com

ヤン・ルクン、人工知能の父

ヤン・ルクン、人工知能の父

最適の脳損傷

名前からして挑戦的でした。ヤン・ルクン (Yann LeCun) がベル研究所のジョン・デンカー (John S. Denker)、サラ・ソッラ (Sara A. Solla) とともに1989年に発表した技術には、「最適の脳損傷 (Optimal Brain Damage)」というタイトルがついていました。脳を損傷させるとは、神経回路網を自ら育ててきた人たちがつけるような名前ではありませんでした。しかし、3人がしようとしていたことは、正確にそれでした。成長した神経回路網を手術台に載せて、そのなかに絡み合った数万個の接続線のうち、役に立たないものを選んで削り取る作業です。問題は、どれが役に立たないのかということでした。

ホルムデルの実験室には、当時、奇妙な緊張が流れていました。手書き郵便番号を読み込むことに成功した最初の実用的量込みニューラルネットワークLeNet-1が成功しましたが、成功はすぐに新しい障壁をもたらしました。神経回路網をもっと難しい問題に適用するには、その規模を大きくする必要がありますでしたが、規模を大きくすると、2つのことが同時に大きくなりました。学習に必要な時間と、過適合のリスクです。当時、1台のワークステーションで数千個の重みを持つ神経回路網を学習させるには、丸一日かかりました。そして、神経回路網が大きくなるほど、機械は訓練に使った数字をすべて丸暗記する癖を示しました。初めて見る数字の前では、むしろ困惑しました。学生が試験問題だけをしっかりと暗記して、実際の原理は理解していないのと同じでした。

過適合は、その時代の神経回路網研究者たちを悩ませていた慢性病でした。研究費を出す側は精度の数値を求め、精度を上げるには神経回路網を大きくしたくなるのですが、大きくするとこの病気が発症しました。訓練用の数字では満点に近いのに、実際の郵便物では成績が大きく低下しました。どこまで大きくすれば暗記せずに学ぶのか、その境界を知っている人はいませんでした。みんな試行錯誤で推測しました。層をもう一つ積んでみて、悪くなれば減らすという方式でした。

規模が大きすぎる神経回路網が何をするのかを想像すると、こうです。その神経回路網は、訓練で見た7を一つ一つ、その画の揺らぎやインクの滲みまで、全部丸暗記します。7という形を学ぶ代わりに、特定の人が特定の日に書いたその7を記憶するのです。だから、初めて見る7の前では崩れます。学ぶとは、暗記することではなく、多くの例から共通した骨組みを引き出して、初めて見るものにも適用することです。神経回路網が大きすぎると、骨組みを引き出す代わりに、例をそっくりそのまま吞み込んでしまいます。サイズを減らすことは、その吞み込みを防ぎ、機械が否応なく工夫を学ぶようにする作業でもありました。

ルクンは、神経回路網にも正当なサイズがあると考えていました。問題に比べて大きすぎると暗記してしまい、小さすぎると学べません。その間の最適な線を手の感覚で推測する代わりに、学習が終わった神経回路網から役に立たない接続を切り取り、サイズを自動的に縮小する方法が必要でした。切り取るという発想自体は新しくありませんでした。当時の人たちが使っていた方式は素朴でした。接続の強さ、つまり重みの絶対値が小さいものから切り取りました。値が小さいので、重要性は低いだろうという推測でした。推測はたいてい当たりましたが、時々大きく外れました。

ルクンと2人の同僚は、その推測に疑問を持ちました。重みが小さいからといって、その接続の作用まで小さいとは言えないということでした。ある接続は値は小さくても、神経回路網の判断を大きく揺

さぶり、ある接続は値が大きくても、切り取れば何も起こりません。値の大きさではなく、その接続を取り除いたときに、神経回路網全体の誤差がどれだけ悪くなるか。それが真の重要度でした。3人はその重要度を「サリエンシー (saliency) 」と呼びました。

この問題と一緒に取り組んだ3人の背景は、偶然ではありませんでした。ジョン・デンカーは物理学者で、サラ・ソッラは情報理論を扱っていた人でした。誤差を1つの地形として見て、その地形の歪みを測るという発想は、神経回路網をコンピュータプログラムとしてのみ見ていた人より、物理学の言葉に精通した人の方が、より自然でした。ルクン自身も、工学と物理の感覚で神経回路網を扱ってきた人でした。学習をエネルギーが低い場所を探す過程として見るこの観点は、彼の長い習慣であり、後年、彼が人工知能を説明するたびに繰り返し取り出す基本図でもありました。

問題は、その値をどうやって事前に知るかということでした。接続の1つを実際に削除して誤差を測り、また戻し、次の接続を削除してみる。こんな具合に数万個を1つ1つ測るには、学習を数万回繰り返す必要がありました。ルクンはここで、学部と博士課程を通じて身につけた物理学の言葉を取り出しました。誤差を1つの地形として見る観点でした。

神経回路網の学習とは、誤差という谷間の最も低い底へ下っていく作業です。学習が終わり、底に辿り着くと、その場所はどの方向に足を踏み出しても、もはや下る場所がない平坦な地点です。逆伝播が教えていた勾配、つまり1次微分値は、そこではほぼ0に収束します。だから底では、勾配だけでは何も区別できません。代わりに残る情報がありました。谷間の壁がどれだけ急に湾曲して上がっていくのか、その湾曲の程度です。数学では、これを曲率と呼び、2次微分で測ります。底が緩やかに湾曲する方向に置かれた重みは、0に押しやっても誤差はほとんど増えませんが、壁が急に立ち上がる方向に置かれた重みは、ちょっと触れても誤差が急増します。ルクンとデンカー、ソッラは、この曲率から各接続のサリエンシーを計算する公式を導き出しました。1つ1つ削除してみる代わりに、地形の形だけを見て、どの接続が切られてもよいのかを一度に読み取る方法でした。

結果は冷徹で明確でした。3人は手書き数字を分類していた神経回路網を手術台に載せました。その中には10万個以上の接続線が絡み合っていました。独立して調整される媒介変数だけでも2500個以上ありました。最適の脳損傷の技術で各接続のサリエンシーを計算し、誤差をほとんど増やさないものから順番に削り取りました。切り、残った神経回路網をもう一度少し調整し、また切る。研究員が側で見守りながら、少しずつ介入する過程でした。こうして接続の数が4分の1のレベルまで減りました。

ここで意外なことが起こりました。規模を大きく減らした神経回路網が、元の神経回路網と同じか、むしろ少し良い精度で、初めて見た数字を読み込んだのです。常識では、接続を削除すれば性能が低下するべきでした。しかし、反対でした。重要でない接続は、実はノイズを運んでいたのです。その枝葉を取り除くと、残った骨組みがより鮮明に学びました。小さな神経回路網が大きな神経回路網より世界をよく読むというこの逆説は、機械にとってもくだらないものが毒になりうるという古い真理を数字で示したわけでした。

手術台に載ったその神経回路網が何だったのかも記憶に留めておく価値があります。前の章で手書き郵便番号を読み込んだ、まさにその系列の神経回路網でした。ルクンと同僚たちは、1989年の学会で逆伝播で手書き数字を読む神経回路網を発表したばかりで、最適の脳損傷は、その神経回路網をより一層引き締めるフォローアップ作業でした。数字を読ませることと、その読む機械を余分なくしつかり仕上げるのが、同じ手から並行して進められました。1つは神経回路網に何を入れるかを問い、もう1つは何を出すかを問いました。2つの問いが噛み合ってこそ、実戦に出ることができるようになりまし

た。

名前が挑戦的だったと言いましたが、その名前には思わぬ誠実さがありました。人間の脳も成長する時にそうします。赤ちゃんの脳は、必要よりもはるかに多くの接続を作り、成長する間に使わない接続をふるい落とします。削ることで、むしろより明確になるのです。ルクンの神経回路網が手術台の上で経験したことは、成長する脳が静かに経験することの工学的な翻案でした。彼が神経回路網を扱うときに、常に生物の脳を目の端で見えていたという事実は、この箇所でも明らかになります。脳をそのまま模倣しようとしていたわけではありませんでした。脳が問題を解く方式から、原理の1つを借りてきただけです。必要な分だけ残して残りを捨てる節制。この節制は彼の手を離れず、神経回路網を扱う彼の生涯の習慣になりました。

暗記と学習の間のこのつな引きは、今日でも機械学習の古い課題として残っています。神経回路網を扱う人なら誰もが、この綱のどの地点に立つのかを決めなければなりません。小さすぎると学べず、大きすぎると暗記してしまいます。最適の脳損傷は、その地点を指先の感覚の代わりに、誤差地形の数学で見つけようとした早期の試みでした。答えを完全に解いたわけではありませんでしたが、問題に正面から向き合ったことは確かでした。

この結果がルクンにもたらしたのは、精度のいくつかのパーセントの改善ではありませんでした。神経回路網のサイズを手の感覚で決める代わりに、機械が自ら自分の体に合ったサイズを見つけるようにする道具でした。今、彼はより大きな問題に神経回路網を恐れずに適用してみることができるようになりました。郵便番号のように綺麗に切られた1つの数字を超えて、複数の文字が互いに接し重なり流れる本物の文書に向かってでした。銀行の金庫に毎日流れ込む、人それぞれ異なる手書きの小切手がその次の目標でした。

LeNet-5と1998年の論文

1990年代半ば、手書き数字研究には1つの共通テスト問題が生まれました。MNISTと呼ばれた手書き数字データベースです。米国の標準機関が集めた手書き数字を整理して作られたこの集合は、高校生と国勢調査員がそれぞれの手で書いた0から9までの数字でした。このテスト問題が良かったのは、誰もが同じ問題で実力を競うことができたからです。その後数十年間、世界中の全ての数字認識神経回路網がこの集合で自分の成績を測りました。しかし、テスト問題が充実すると、ルクンが最初に作ったLeNet-1がむしろ小さく見えました。容器が小さくてデータを全部入れられませんでした。

ルクンは神経網を段階的に発展させました。畳み込み層を2つに増やし、その間に画像を粗く縮小する層を挿入し、最後に判断を集約する層を加えました。こうして生まれたのがLeNet-4であり、パラメータ数は約17000個でした。そして、この系統の最後の頂点であり、後年ディープラーニングの教科書の最初の図になるモデルが1998年のLeNet-5でした。パラメータ数は約60000個に増えました。最初のLeNet-1からここまで、神経網は数年かけて徐々に規模を拡大し、より難しい手書き文字に対応していききました。その60000個の値が層状に組み合わせることで、人間が規則を与えたことのない数字の形態を機械自らが抽出したのです。

LeNet-5が数字を読む方式は、前の章で見た視覚皮質の原理をそのまま引き継いだものでした。下の層は、短い筆画や曲がった線のような小さなパーツを画像全体から探し出します。同じパーツ探索を画像の隅々で繰り返すので、数字が左にあらうが右にあらうが同じ要素として認識します。その上の層は、下の層が見つけたパーツを粗く要約します。筆画がちょうどどの位置にあったかは曖昧にぼかし、何があったかだけが残ります。このぼかしのおかげで、数字が少し傾いたり歪んでいても揺らぎません。パー

ツを探し要約することを何度も繰り返して上へ進むにつれ、神経網は筆画から部分へ、部分から数字全体へと判断を発展させました。人間が3という字の形を規則として与えたことはやはりありませんでした。機械が数万枚の手書き文字を見ながら、自らその規則を身に付けたのです。

学習を終えた神経網の内部を詳しく見てみると、下の層が何を学んだかが目に見えました。誰からも指示されていないのに、最初の層は短い辺と曲がった筆画を捉える検出器で満たされていました。前の章で見た猫の視覚皮質の単純細胞がしていた仕事、つまり特定の方向の線分に反応するその仕事を、機械が自ら再び発見したのです。生物が眼で学んだ原理を、神経網が手書き文字を見ながら独りで辿り直したのです。ルクンが視覚皮質から借りてきた設計が無駄ではなかったことの証拠でした。

当時のコンピュータでこの神経網を学習させることは、忍耐力の試練でした。LeNet-5は200メガヘルツのプロセッサを1つ搭載したシリコン・グラフィックスのサーバ上で訓練されました。今日の携帯電話1台にも及びません。データを20回繰り返して示すだけで、丸々3日かかりました。ルクンはここでも、以前得た曲率の知識を使いました。誤差地形の曲がり具合を推測して歩幅を調節することで、遅いコンピュータが少しでも速く学習を完了させました。最適な脳損傷から誤差地形の曲がり具合を読み取ったあの眼が、今度は学習を促進するのに使われました。1つの数学が2つの場所で機能していました。

このコンピュータ1台が丸々3日を費やすということは、1つの実験がそれだけ代償を要するということでもあります。設定を1つ誤ると学習が間違った方へ進み、その3日全てが台無しになりました。今は同じ規模の神経網を数分で実行します。あの時代のルクンは1つの実験に4日を費やし、結果が出るまで待ちました。神経網が正しいという確信がなければ、耐え難い待ちぼうけだったでしょう。他の人たちがより速い方法へ殺到していく中、彼は遅いコンピュータの前で自分の神経網が学ぶのを待ったのです。

1998年11月、ルクンはレオン・ボットウ (Léon Bottou)、ヨシユア・ベンジオ (Yoshua Bengio)、パトリック・ハフナー (Patrick Haffner) とともに1本の論文を世に発表しました。タイトルは「文書認識に適用された勾配降下法による学習 (Gradient-Based Learning Applied to Document Recognition) 」で、『Proceedings of the IEEE』第86巻第11号に掲載されました。40ページを超える長編でした。実験報告書というより1つの宣言に近いものでした。人間が手で特徴を選び分けて機械に与えるという、それまでのやり方を廃止し、生データから何が重要かを機械自らが見つけるべきだという主張でした。この主張は後にディープラーニング全体の基礎となります。今日、この論文は世界で最も引用されている人工知能の文献の1つとされています。人々が引用するのはおおむね前半部分です。LeNet-5の構造と、畳み込み神経網が画像から特徴を自ら学ぶ原理を描いた箇所です。

この主張が空虚でなかったのは、同じ論文の中で証拠を並べて示したからです。それまで数字をよく読むとされた方法は、おおむね次のようなものでした。人間がまず画像から使える特徴を選び抽出する規則を手で書きます。筆画の方向、穴の数、終点の位置といったものです。その特徴を単純な分類器に入れて数字を判定します。人間の判断が半分、機械が半分をするというわけでした。LeNet-5はその前半を排除しました。人間が特徴を与えることなく、生のピクセルから何が重要かまで機械が学びました。同じMNISTテストセット上で、人間が特徴を与えた方法よりもこちらの方が誤りが少なかったのです。手で彫り込まれた特徴の時代が終わろうとしているという信号でした。この信号が世界中に聞こえるようになるまではさらに十数年を要しましたが、その数字はすでに1998年に示されていたのです。

この論文が長く生き続けた理由は、別にもありました。40ページにわたって、ルクンと同僚たちは数字を読む様々な方法を1ヶ所に集めて同じテストセットで比較しました。自分たちのものだけを自慢する論文ではありませんでした。その時代の手書き文字認識の地形図を全て描いた地図でした。後に続く

研究者たちは、何がどのくらいよくできるのか知りたいときこの地図を開きました。今日のディープラーニングの教科書が冒頭にLeNet-5の図をそのまま掲載しているのもそのためです。量み込みと要約を層状に積み重ねたあの図は、今世界を動かす神経網たちの遠い祖先の顔です。

ルクンは後年この論文を振り返ってやや異なる話をしました。多く引用された前半部分よりも、ほとんどの人が最後まで読まなかった後半部分に、自分が本当に言いたかったことがあったということでした。彼が大切にした後半部分は、1枚の画像を分類する問題を超える話でした。実際の文書の上では、文字は綺麗に1つずつ離れていません。互いに付着し、重なり、傾き、どこで1つの文字が終わり次の文字が始まるのかさえ明確ではありません。それまでの方式は、まず文字を切り出してから1つずつ読むというものでしたが、切り出す段階で1度でも誤ると、読む段階が全て崩壊してしまいました。

ルクン、ボットウ、ベンジオは、切り出しと読み込みを分け隔てないことに決めました。代わりに、文字を切り出すことができるすべての可能性を1つの分岐経路図のように展開しました。この地図の分岐点ごとに神経網が付与した確信度スコアを載せ、ここに言語の規則までも重ね合わせて評価しました。たとえば銀行口座番号であれば出現可能な桁数が決まっているので、その規則に合わない分岐は事前に除外されました。神経網の目と文法の物差しを一緒に使って、多くの分岐経路の中から最もあり得そうな1つを選び出す方式でした。この仕組みを彼らはグラフ・トランスフォーマー・ネットワーク (Graph Transformer Networks) と呼びました。

核心は、この分岐経路図のすべての段階を微分可能な形で設計したことにありました。微分可能ということは、最後に出された最終答の誤差を逆伝播で遡って送り返すことができるという意味です。その誤差が小切手の最初の処理段階から神経網内の数万個の重みまで遡って流れ、システムの全ての部品を一度に調整しました。切り出して読む全パイプラインが1つの神経網のように一緒に学びました。部品ごとに別々に調整して継ぎ合わせるというそれまでの慣行とは異なる道でした。エンド・ツー・エンド学習と呼ばれることになるこの考え方は、当時は斬新でした。しかし後年ディープラーニングが世界を覆う時、ほぼすべてのシステムがこの文法に従いました。ルクンが惜しみながら語ったように、その考え方は誰にも読まれなかった後半部分で静かに待っていたのです。

アメリカの小切手の10～20パーセント

論文が学界の資産であるなら、その論文が銀行の金庫で行うことは別の種類の証明でした。ベル研究所はAT&Tの金融機器子会社だったNCR (National Cash Registers) と手を組みました。NCRはアメリカの銀行の現金自動出納機と後処理決済機械を製造していた会社でした。2つのチームが共同で開発したアメリカの小切手自動読取システムには、HCAR50という名前が付けられました。ホルムデルで製造された、小切手に手書きされた金額を読む機械という意味でした。アメリカの小切手には、金額を記入する欄が2つあります。数字で書く欄と、文字で書く欄です。機械が読んだのは数字の欄であり、その数字の金額を呼ぶ名前が、この機械の名前にそのまま入りました。HCAR50は一夜にして生まれたものではありませんでした。先行する試作品HCAR30とHCAR40が実験室の中で限界を露呈させながら改良された末に、実戦に投入されうる仕様に達した版でした。

1996年6月、このシステムがアメリカの主要銀行の後処理決済機械に搭載されて稼働し始めました。毎朝銀行に押し寄せる小切手は、筆跡がそれぞれでした。きちんと書かれたものもあれば、読み取りが難しいほど走り書きされたものもありました。人間が1枚ずつ目で読んで、キーボードに入力する作業でした。小切手が数百万枚なら、その作業も数百万回でした。

実際の小切手は、実験室の綺麗な数字と異なっていました。人ごとに異なる筆跡に、背景には図が印刷されており、折り目と汚れが付着していました。金額の数字は互いに付着したり重なったりして、どこで1つの数字が終わるのか曖昧でした。前のセクションで見たグラフ・トランスフォーマー・ネットワークが活躍する場所はまさにここでした。どこで切るかをあらかじめ決めず、複数の切り方の可能性を一緒に評価して、最もあり得そうな金額を読み取りました。論文の中の仕組みが、銀行の機械の中で実際に動いていたのです。

HCAR50は流れ込む小切手のおよそ半分を、人手を一切借りずに正確に読み取り帳簿に記録しました。残りの小切手については無理して答えを作り出しませんでした。文字がかすんだり崩れたりして確信が持てなければ、機械は自ら判断を保留し、その小切手を人間による検査ラインに送りました。誤答を出すより分らないと手を上げることを選んだのです。この自制が核心でした。機械が読みやすい半分の処理してくれれば、人間は困難な半分だけに目を使えばよかったのです。そして機械が実際に答えを出した小切手での誤り率は1パーセント以下に抑えられました。誤った金額が帳簿に記録されることは銀行としては絶対にあってはならないことだったので、機械に金銭を預けられたのはまさにこの謙虚な設計のおかげでした。神経回路網は自信がある分だけ答えました。

数字は時とともに大きくなりました。2001年ごろ、このシリーズのシステムが1日に読み取った小切手は2000万枚に達しました。その年、米国で毎日やり取りされていた小切手の10枚に1枚、多くて5枚に1枚を機械が人間に代わって読んでいたということです。

この数字がいかに重いかは、その前の光景を思い出さなければわかりません。それまで米国の銀行のバックオフィスには、小切手の金額をキーボードで入力する人々が列をなして座っていました。毎日、数百万枚の小切手の金額を指で転送する仕事でした。退屈で、目が疲れ、誤りが頻繁な仕事でした。HCAR50が導入されると、その列は短くなりました。機械は簡単な部分を処理し、人間は機械が手をつけた困難な部分だけを引き受けました。人間を完全に排除したのではなく、人間が対応すべき仕事量を半減させたのです。機械が得意とする退屈な半分の仕事を機械に任せ、人間が判断の必要な困難な半分に目を集中させるこの分業は、後年人工知能が職場に入ってくるたびに繰り返される図式の早期版でもありました。神経回路網が袋小路だと学界が指を指していたまさにその時代に、ヤン・ルクンの神経回路網は、世の誰もが気づかない方法で米国金融の一つの柱を静かに支えていました。それは論文の脚注ではなく、銀行のコンピュータ室でのことでした。

ここには噛み締めるべき点があります。同じ時期、学会はルクンの論文をしばしば却下し、研究費を提供していた手も神経回路網から引き上げていました。学界の評決は冷淡なものでした。しかし市場の評決は異なっていました。銀行は彼の機械に自分たちの金銭を預けました。一枚の紙を読み誤れば直ちに損失につながる場所で、人々はこの神経回路網を信じました。論文審査者の判断と銀行会計担当者の判断がこのように食い違うことは珍しくありませんでした。後年の話ですが、市場が学界より先に正しかったのです。

このような成功には、数年前の分枝的な出来事が下地として隠れていました。ストーリーはHCAR50が銀行に配置される数年前にさかのぼります。1993年春、NCRはベル研究所に異なる依頼を持ってきました。クレジットカードの裏面の署名を偽造から保護したいということでした。店でカードで商品を購入する際、店員がカード裏の署名とレシートの署名を目で照合していた時代でした。人間の目は疲れやすく、騙されやすいものでした。NCRはその照合を機械に代わってもらうことを望みました。条件は厳しいものでした。カード裏の磁気ストライプに書き込める空間はわずか80バイト、1行の文字にも満たない大きさでした。その狭いスペースに1人の筆跡全体を収める必要がありました。

ルクンはベル研究所のジェーン・ブロムリー (Jane Bromley) とエドゥアルド・サッキンガー (Eduard Säckinger) とともに、特異な構造を作成しました。同一の重みを共有する2つの神経回路網を並べて構成したのです。同じ日に生まれ、同じ頭を持つ双子のようでした。そのため、その名前も双子の神経回路網、シャムネットワーク (Siamese Network) でした。

動作方式は対称的でした。この構造は署名が本物が偽物かを直接判定しようとしませんでした。代わりに、2つの署名が同じ手から出たかどうかを比較しました。同じ人が書いた2つの署名を両側に入力すると、2つの神経回路網が抽出した特徴が空間内で互いに接近するよう学習されました。他人が模倣した署名を入力すると、2つの特徴がバネのように互いに反発するようになりました。近さと遠さ、その距離ひとつで真偽を判別しました。おかげで機械は世界のすべての署名を事前に暗記する必要がありませんでした。初めて見る人の署名でも、2枚を比較して同じ手からのものかどうかだけを確認すればよかったのです。機械は署名そのものを丸ごと記憶する代わりに、ある人の筆跡を他の人のものと区別する短いマーカーだけを抽出して、80バイト以内に収めました。

80バイト以内に1人の手を収めるという発想は、今見ても大胆です。その中には署名の画像そのものが入るわけではありません。この人の手とあの人の手を区別するのに必須の数個の特徴、その要約だけが入ります。機械が何を捨て何を残すかを自ら決めたからこそ可能な業でした。最適な脳損傷で無用な接続を捨てたように、ここでは筆跡を区別するのに無用な情報が捨てられました。残そうと思うなら、まず捨てることを知らねばならないという原理が、今度は磁気ストライプの上で機能しました。

ルクンと同僚たちは、この構造を1993年のあるカンファレンスで論文として発表しました。時間遅延ニューラルネットワークを使用した双子構造で署名を検証するというものでした。ブロムリー、ベンジオ、ルクン、サッキンガーの名前が並んで記載されました。

この署名実験は大きな商用製品には成長しませんでした。しかし、ここで得た感覚、すなわち2つの入力を並べて置き、同じなら引き寄せ、異なるなら押し出すという方法で機械に表現を教える方法は、ルクンの手から消えませんでした。その感覚は他の人々の手でまず大きく成長しました。後年、2枚の顔写真が同じ人物のものかどうかを判定する顔認識システムが、まさにこの双子構造を受け継いで世界に広がりました。正解表なしに世界を比較して自ら学ぶこの発想は、後年ルクンが人工知能の次の道を描くときに、再び舞台の中央に戻ります。1993年のカード1枚の裏面で、その発想はすでに根付いていました。

紙を圧縮する : DjVu

小切手を読み取っていたその数年間に、ルクンの立場が変わりました。彼はベル研究所を引き継いだ研究所で、画像処理研究部門の責任者になりました。音声認識の大家ローレンス・ラビナー (Lawrence Rabiner) が主導していた音声画像処理研究所内のある部門でした。神経回路網について口にするだけで時代遅れの人と見なされていた時期でした。そのような時代に部門を率いるようになった彼は、神経回路網の話をしばらく脇に置いて、目の前で明るくなりつつあった別の波を見つめました。それはインターネットでした。

ルクンが見たのは、やがて訪れるべき大きな課題でした。人類が紙の上に積み重ねた知識を、途切れることなくインターネットの上に移すという仕事でした。図書館の書架に置かれた本はその場所に行つてのみ読むことができました。希少本であるほど、古い文献であるほどそうでした。インターネットはその壁を壊そうとしていましたが、まだのところ約束に過ぎませんでした。ところが1990年代中盤から後半にかけての通信回線は、残酷なほど遅かったのです。紙の本を鮮明にスキャンした画像ファイル

は、サイズが大きいため、回線を通す途中で遮断されることがしばしばあり、1ページをスクリーンに表示するのに数十分かかりました。知識は紙に閉じ込められ、回線はその知識を流すには狭すぎました。一つの壁を越えたら、別の壁がありました。

初期のインターネットは事実上、テキストのネットワークでした。キーボードで打ったテキストは軽いので、遅い回線でもよく流れました。しかしスキャンされた画像は重く、ほとんど流れませんでした。そのため人類が紙に残した古い知識、手書きの図表が含まれた古い論文や古ぼけた本は、ウェブの敷居を越えることなく書庫に残されていました。新たに打たれたテキストはインターネットに上がりましたが、古い紙は外に立っていました。画像の重さを軽くすることは、人類の古い知識をインターネットの中に入れるか、外に置くかを定めることでもありました。

ルクンはこのボトルネックを数学で破ることにしました。文書を読むニューラルネットワークを作っていた彼にとって、文書を転送する問題は見慣れたものではありませんでした。小切手にしろ本にしろ、結局のところそれは紙上のインクを機械が扱うことでした。一方では、インクがどんな文字であるかを読み取り、他方ではそのインクを軽く圧縮して回線に乗せました。手は同じでした。彼はベル研究所で長く手を組んだレオン・ボットウ (Léon Bottou)、そしてパトリック・ハフナー (Patrick Haffner)、ポール・ハワード (Paul Howard) を呼び集めました。彼らが作成した画像圧縮方式はDjVuという名前が付けられました。デジャヴと読みました。「どこかで見たような」という意味のその名前のように、一度スキャンされた文書が画面の上に息を吹き返しました。

着想は人間の目を真似することから生まれました。本棚を見ると、文字の鋭い輪郭ははっきり見えますが、紙の黄色い背景は大雑把に見過ごします。目は情報が集中している場所にのみ力を使います。DjVuはこの習慣を機械に移しました。スキャンされた文書を丸ごと圧縮するのではなく、3つのレイヤーに分割しました。文字を含む前面ははっきりしている必要があるため、高解像度の白黒レイヤーとして分離しました。紙の黄色い質感と汚れが付いた背景はぼやけていても構わないため、低解像度のレイヤーに圧縮しました。図や写真があれば、色を含むレイヤーに任せました。各レイヤーは特性が異なるため、各レイヤーは異なる方法で圧縮されました。

当時広く使われていたJPEG方式は、テキストと背景を区別なく一度に圧縮しました。そのためわずかに強く圧縮しただけで、黒い文字の周りに汚い斑点が広がり、ストロークが潰れました。ドキュメントを圧縮するには不向きな特性でした。DjVuはテキストを背景から分離したおかげで、背景がどれだけ強く圧縮されても、テキストだけはナイフのように鮮明でした。同じ画質を比較すると、圧縮されたファイルのサイズはPDFまたは単一ページのJPEGの数分の一から十分の一以下でした。数十分かかっていた1ページが、わずか数秒で画面に表示されました。

ここにはルカンらしい逆説が一つ隠されていました。生涯、機械は表象を自ら学ぶべきだと信じてきた人が、この圧縮に限ってはニューラルネットワークに任せませんでした。彼は早くも1986年の博士課程時代に、入力を自ら復元しながら表象を学ぶオートエンコーダの原型を自らの手で動かしたことがありました。ニューラルネットワークに画像を入力し、その画像を再び描画させながら、その過程で重要な特徴を自ら学習させる仕組みでした。しかし、当時のコンピュータではすべてのピクセルを一つ一つ描き直す方式が、あまりにも重く遅いということを痛感していました。

そのため、DjVuではピクセルをすべて描き直させる代わりに、幾何学と情報理論で鍛え上げた圧縮フィルターを手作業で組み込みました。画像の本質を捉えつつも、すべての点を一つ一つ再現するわけではないというこの節制は、一度経験した失敗から学んだ態度でした。すべてのピクセルを描き直そうと努

める代わりに、何を捨ててもよいのかを先に決めること。この考えは、ここでは画像圧縮のコツとして使われるにとどまりましたが、後日、ルカンが人工知能の別の道を描く際に、その中心へと戻ってきます。その時彼が投げかけた問いも同じでした。機械が世界を学ぼうとする時、目の前のすべての点を描き直すのか、それとも重要なものだけを捉えるのか。

圧縮がどれほど強力に機能するのかは、実戦で証明して見せました。ルカンとポトウは、ニューラルネットワークの研究者たちが数十年間冬を耐え忍びながら論文を発表してきた学会NIPSの会議録全体を直接スキャンし、DjVuで圧縮しました。数式と図表がぎっしり詰まった数千ページに及ぶアーカイブでした。当時の技術では、このような文書の山を丸ごとウェブに上げて閲覧できるようにすることは、思いもよらないことでした。二人は各ページを数十キロバイトという軽いファイルに縮小してウェブに上げました。遅い回線を使っていた研究者も、先輩たちの論文を滞りなく開いて見るができるようになったのです。それまでその論文を読むには、学会に行くか、分厚い会議録を手に入れるしかありませんでした。今や世界中どこからでも画面を開けばよくなりました。ニューラルネットワークを蘇らせようとした人々の足跡が、ニューラルネットワークから派生した圧縮技術のおかげで世界に開かれました。冬を耐え抜いた人々の記録を、同じ冬を耐え抜いた人が世に解き放ったわけです。

いざこの技術を事業としてどう使うかについて、会社は方向性をつかめずにいました。ルカンは後日、会社はこの圧縮技術を手にしながらも、それで何をすべきか本当に分かっていなかったと回顧しています。手にしたものの価値を、作った人だけが知っていたことになります。ところが、会社が分からずに放置したその技術は、実は世の中では静かに生き残っていました。

回線が貧弱だった東ヨーロッパの図書館たちがDjVuを歓迎しました。その中でもポーランドの司書と学者たちは、数百年前の稀覯本や古典文献をデジタルに移行して保存する上で、この無料の形式を標準としました。お金を払わなければならない独占形式の障壁の代わりに、貧弱な回線でもはっきりと開く軽い形式を選んだのです。高価な回線を買えない場所であるほど、この技術が切実に求められました。後にDjVuは、人類のスキャン文書を守ってきたインターネット・アーカイブの画像形式の一つとしても定着しました。古い本の一冊が地球の反対側の画面にはっきりと表示される時、その裏では、ニューラルネットワークが学界の嘲笑を浴びていた時代に、一人のフランス人が作った圧縮方式が音もなく動いていました。

DjVuが世の中のすべての文書の標準になったわけではありません。その座は結局、アドビのPDFが占めました。企業が後押しし、プログラムごとにデフォルトで搭載されたPDFはどこにでも広まり、DjVuはスキャンされた古い文書という片隅に残りました。しかし、その片隅でDjVuは長く生き延びました。高価な回線と高価なプログラムを負担できない場所、古い紙を守らなければならない場所で、軽くて無料であるこの形式は依然として使い道がありました。世界中を覆い尽くすことはできなくても、どうしても必要な場所で静かに働く技術。ルカンが作ったものには、このような性質がしばしば滲み出ていました。

DjVuには、ルカンが後日さらに強く推し進めることになる信念の糸口も含まれていました。良い技術は値をつけて閉じ込めるよりも、開いておくべきだという考えです。彼はこの圧縮形式を特定の会社の錠前に閉じ込める代わりに、誰もが使えるように開いておく側を支持しました。回線が貧弱な場所の図書館がその恩恵を受けました。知識を移す道具を無料で開いておけば、知識が切実に求められている場所へと真っ先に流れていくということを、彼はこの時目の当たりにしました。数年後、彼が巨大企業の中に研究所を設立する際、成果物をすべて公開するよう釘を刺し、後日、人工知能モデルまでも開いておくべきだと声を高めることになる態度のルーツがここにありました。

これがベル研究所がルカンに残した最後の章でした。銀行の金庫の小切手を読み取ったニューラルネットワークも、紙の知識を回線の上に救い出した圧縮形式も、すべてはニューラルネットワークを誰も信じなかった冬の真っ只中で生まれました。学界が背を向けている間、ルカンが作った機械たちは人々が気づかない場所で毎日動いていました。しかし、冬はまだ終わっておらず、彼を抱えてくれていた研究所でさえも長くは持ちこたえられないところでした。商用化の成功が頂点に達したまさにその瞬間、別の種類の知らせがホルムデルに向かって来ていました。

弁護士

kimkj.com

第5章 ベル研究所の黄昏

ヤン・ルクン、人工知能の父

商用化の日に下された解体命令

1995年のある夜、ニューヨークのフランス料理店にヤン・ルクン (Yann LeCun) とベル研究所の同僚たちが集まっていました。グラスにはシャンパンが注がれていました。席にはお金を読む装置を実際の製品にしたNCR (National Cash Register) のビジネスチームも一緒にいました。乾杯の理由は明らかでした。彼らが何年もかけて作った機械がアメリカの銀行のバックオフィスに入り、人間の手を少し借りながら毎日の小切手を自分で読み取っていたからです。

ホルムデルの研究所の建物は、フィンランド出身の建築家エーロ・サーリネン (Eero Saarinen) が設計した鏡のようなガラスの箱でした。ニュージャージーの平野の真ん中に、空と雲を丸ごと映す壁に囲まれた中に、数千人の科学者とエンジニアが働いていました。そのガラスの箱の中のあるセクションで小切手読み取り装置が誕生しました。実験室のスクリーンの前で夜通し働いていた人たち、銀行から運ばれてきた実際の小切手の束でテストを繰り返していた人たち、そしてこれを売れる機械にするためにオハイオから飛んできたNCRのエンジニアたちが、その何年もの主役でした。

その機械の名前はHCAR50でした。ルクンとレオン・ボットウ (Léon Bottou) をはじめとするベル研究所適応型システム研究部の手から生まれた畳み込みニューラルネットワーク (CNN) がその心臓でした。ルクンがホルムデル (Holmdel) でこの数年間に磨き上げたその血筋そのものでした。この機械は当時、アメリカで毎日流通していた銀行小切手の10枚に1枚、多くは5枚に1枚を読み取っていました。枚数にすると数千万枚でした。紙の束を人間が一つつ目で読んでいた仕事を、機械が代わりに引き受け始めたのです。

機械の動作方式は謙虚に設計されていました。ニューラルネットワークは小切手に書かれた金額を読みましたが、自分の読み取りに自分で点数をつけました。確信度が高い小切手は機械が処理し、確信度が低い小切手は人間に送られました。走り書きの数字、インクが滲んだ数字、枠からはみ出した数字は、依然として人間の目に委ねられました。機械が人間を押し退けたわけではありませんでした。人間が見る必要のある紙の山を減らしたのです。銀行のバックオフィスの風景がそのようにじわじわと変わっていったのです。

研究所から生まれた考えが実際に銀行のお金に触れることは稀です。ほとんどの良いアイデアは論文一本で終わるか、デモ画面の中だけでうまく動作します。HCAR50はその敷居を超えました。読み取りが間違っていれば、誰かの口座で実際に金額が狂います。そのような責任を背負って毎日数千万枚を処理したということは、ニューラルネットワークが実験室の玩具から抜け出して、社会の基盤施設の中に入ったということを意味していました。ルクンにとってこの成果は、論文引用数で測ることができないものではありませんでした。誰も信じていなかった方法が世界の重みに耐えたという証拠だったのです。

ニューラルネットワークについて「行き止まりの道」と呼んでいた時代でした。知能とは脳に予め組み込まれた規則であり、機械にはその規則を人間が書き込まなければならないと、学界の大多数が信じていた時でした。その真っ最中で、ルクンの機械は言葉の代わりに数字で応答しました。小切手を読み取り、銀行がその読み取り結果で実際にお金を移動させました。詭弁ではなく工学であるという証拠を現実の心臓部に突き立てたようなものでした。テーブルの笑いはその証拠の上に置かれていました。

その笑いは長く続きませんでした。ルクンはその後、ディープラーニング教育企業deeplearning.aiとのインタビューでその瞬間を次のように回想しました。「プロジェクトがついに本当に大きな成功を収めたまさにその日に、会社がそれを全部解体してしまったのです。本当に落ち込みました(The whole project got dismantled ... the day it became a really big success. It was really depressing)。」成功の日と解体の日が同じ日でした。

解体の正体はこうでした。1995年9月20日、AT&T本社が会社を3つに分割すると発表しました。通信事業を担当するAT&T、通信機器と研究を担当するルーセント・テクノロジーズ (Lucent Technologies)、そしてコンピュータと端末を担当するNCRに分けられました。独占禁止圧力と規模縮小が名目でした。翌年春、ルーセントが株式市場に上場され、NCRは1996年末に完全に分離されました。小切手読み取りという一つの目標に向かって整列していたチームが、異なる会社の異なる組織図の中に散らばってしまいました。

ベル研究所がどのような場所であったかを思い出すと、この分割の重みが異なる形で伝わってきます。トランジスタはこの研究所から生まれました。情報理論も、UnixオペレーティングシステムとC言語もここで誕生しました。電話独占がもたらす安定した金銭が、すぐにはお金にならない基礎研究を数十年にわたって支えてきた構造でした。その構造が1990年代に入って崩れ始めました。独占は解体され、競争に追い立てられた会社は研究所をコストと見なし始めました。ルクンのセクションが経験したことは、その大きな流れの一片でした。企業が国家の代わりに基礎科学を育てていた一つの時代が、小切手読み取り装置の成功と同じ年に暮れようとしていたのです。

このセクションを通過していった顔ぶれを数えると、後々語り継がれる地図が事前に描かれていたことに気づきます。ヨシユア・ベンジオは1992年ここでルクンとともにポストドク時代を過ごした後、モントリオールへ去りました。レオン・ボットウはルクンの生涯にわたる共同研究者となりました。イザベル・ギュイヨンとコリーナ・コルテス (Corinna Cortes) はヴァプニクとともにサポートベクターマシンの歴史に名を刻みました。一つのセクションの廊下で出会っていた人たちが、その後30年間、機械学習という分野の骨格を分け持つことになりました。ホルムデルのガラスの箱はそのような場所だったのです。

サーリネンのガラスの箱にも黄昏が予定されていました。分割と削減が繰り返されるにつれて、その大きな建物を満たす研究者が減り、2000年代後半には建物全体が空きました。空を映す鏡の壁の中に空の廊下だけが残った時代がしばらく続きました。今、その建物は複数の企業が同居する事務施設として蘇っていますが、トランジスタと情報理論の後継者たちが廊下で出会っていた風景は戻ってきませんでした。ルクンがベル研究所の黄昏と呼んだのは、一つのセクションの解体でありながら、そのセクションを抱いていた一つの世界が暮れていくことでもあったのです。

研究者たちも分裂しました。ベル研究所という名前と建物の大部分はルーセント側に行きましたが、AT&Tも自分の割当の研究組織を新たに整えました。AT&T Labs-Research (AT&T ラボ・リサーチ) でした。ルクンは1996年、この新しい組織の画像処理研究部の責任者になりました。同じ廊下で働いていた人たちがいる日から異なる会社の名刺を持つようになりました。ある同僚はルーセントへ、ある同僚はNCRへ、ある同僚は研究所の外に出て行きました。小切手読み取り装置と一緒に作っていた配置は再び揃えることができなくなりました。

ところが散らばっていく過程で、その後ずっと語り継がれることになる出来事が起こりました。特許を分配することが問題でした。AT&Tの知的財産権担当者たちが各会社に技術権を配分する書類作業を

しました。畳み込みニューラルネットワークの核となる特許がどこに行ったのか。研究を続けられる組織ではありませんでした。現金登録機と端末を製造するNCRの資産台帳でした。ルクンはその後この箇所を回想しながら、弁護士たちが無限の知恵を発揮して(in their infinite wisdom)特許をNCRに配分したと、乾いた冗談を言いました。ニューラルネットワークを研究する人たちは研究機関に残りましたが、ニューラルネットワークの権利はニューラルネットワークを研究しない会社に行ったのです。

実のところNCRは自分の金庫に何が入ってきたのか分かりませんでした。手でハンダ付けした回路と、一行一行コードを磨き上げた人たちは、自分たちが作ったものを自分の研究に思いっきり使うことができない立場になってしまいました。その権利は満期を迎えるまで、それが何であるかを知らない会社の引き出しの中で眠っていました。ルクンの回想によれば、特許は2007年ころようやく終焉を迎えました。なんとという幸運な時点でしょう。特許が解放されていた頃、世界のデータと計算力はついにニューラルネットワークが再び立ち上がるほど大きくなっていました。眠っていた引き出しの中で権利が消えていく間に、その権利が示す技術は復活を準備していたというわけです。

企業の分割は書類上の出来事です。その書類が一つの研究分野の脊椎を折ることができるという事実は、経験した人だけが分かります。ルクンにとって1995年は二つの絵が重なった年でした。一つはアメリカの銀行のバックオフィスで静かに小切手を読む機械であり、もう一つはその機械を作ったチームが碎けて散らばる組織図でした。前の絵はニューラルネットワークが正しかったという証拠であり、後ろの絵はその証拠が世界の注目を集めることができなかったという信号でした。

この時期から人工知能の第二の冬が深まりました。第一の冬が1970年代のパーセプトロン論争の後に来たとすれば、第二の冬は1995年頃から2002年頃まで続きました。学会ではニューラルネットワークという言葉は言い出しにくい言葉になっていきました。論文のタイトルにニューラルネットワークを入れるとレビューで不利になるという話が公然と流れ、研究者たちは自分の研究に別の名前を付けて冬を乗り切りました。研究費の審査でも状況は同じでした。

冬の奇妙な点はこれでした。今回は理論の反論が理由ではありませんでした。第一の冬はミンスキーとパパートの本一冊が、パーセプトロンが解けない問題を数学で証明して持ち込みました。第二の冬にはそのような証明がありませんでした。ニューラルネットワークはむしろ実際に動く物を生み出したばかりでした。アメリカの銀行の小切手がその証拠でした。にもかかわらず押しのけられました。流行が変わり、研究費が移り、若い研究者たちが他のテーマに去っていきました。科学の方向が常に証拠だけで決まるわけではないという事実を、ルクンはこの冬に身をもって学びました。この教訓は30年後、世界中が大型言語モデル一つに殺到する光景を眺めていた彼の目に再び呼び起こされます。

冬の間、ルクンを支えていた一つのパラドックスがありました。学界が「ニューラルネットワーク」という言葉を避けていたまさにその時間にも、彼のニューラルネットワークはアメリカの銀行の奥の部屋で毎日数百万枚の小切手を読んでいたという事実です。学界の地図から消された技術が、経済の血管の中では絶えず働いていました。流行は論文を消すことができて、すでに配置された機械まで止めることはできませんでした。ルクンはこのズレを長く記憶していました。何が生きた技術で、何が死んだ技術かを決めるのは、学会の雰囲気ではなく、現場の動作だということを。彼の頑固さはこの記憶の上に立っていました。

世界はより洗練された数学を手にした競争相手に目を向けました。そしてその競争相手は、ちょうどルクンの隣の部屋に座っていたのです。

隣の部屋のヴァプニクとSVM

ヤン・ルクンの研究室の隣の部屋には、ソビエト連邦からやって来た数学者がいました。ウラジミール・ヴァプニク (Vladimir Vapnik)。統計学習理論の父と呼ばれる人物です。彼はモスクワの研究機関で数十年にわたって学習理論を磨いた後、冷戦が終わろうとしていた1990年代初頭にアメリカに渡り、ベル研究所に職を得ました。ホルムデルキャンパスの同じ廊下、壁一枚を挟んで二人が働いていました。ジョン・デンカー (John Denker) を始めとする著名な研究者たちがその廊下を共に埋め尽くし、部門長はローレンス・ジャッケル (Lawrence Jackel) でした。合成畳み込みニューラルネットワークと支援ベクターマシン (SVM) という二つの方法が、1988年から1992年の間に同じ部門で並行して成長しました。後に機械学習の覇権を争うことになる二つの方法が、一人の部門長の下で、同じ釜の飯を食べながら生まれたのです。

ヴァプニクはニューラルネットワークを信頼していませんでした。彼の見方では、ニューラルネットワークは動作原理を説明できず、数学的保証もない物体でした。多くのつまみを回して合わせる、正体不明の箱だということでした。つまみを回して答えが合えば良いのですが、なぜ合うのか証明できないなら、それを科学と呼べるのか。これが彼の問いかけでした。

ヴァプニクという人物を一文で描くなら、理論を実用の反対語と見なさない理論家でした。彼は、良い理論ほど実用的なものはないという古い格言を好んで引用し、自著の冒頭にもその精神を刻み込んでいました。彼にとって証明は装飾ではなく、道具でした。証明がある方法は、いつ信頼できるかが分かり、証明のない方法は、いつ裏切るかが分からないということです。その信念の眼で隣の部屋を見れば、ニューラルネットワークは昨日は正しく、明日は間違うかもしれない機械でした。

ヴァプニクが手にしたのは証明でした。彼は1971年にアレクセイ・チエルボネンキス (Alexey Chervonenkis) と共にVC次元という概念を確立しました。ある学習機械がどれほど複雑なパターンまで捉えることができるかを数値で測る尺度でした。これに構造的リスク最小化という原則を加えました。機械が学んだことを初めて見るデータでどれほど間違わないか、その限界を事前に計算し出す枠組みでした。1992年には、ベルンハルト・ボーザー (Bernhard Boser)、イザベル・ギユオン (Isabelle Guyon) と共にカーネルという装置を備えた支援ベクターマシンを完成させました。データを高次元に押し上げて、二つのグループの間に最も広い道を作る方法でした。その道を見つける計算は常に一つの正解に収束しました。どこから始めても同じ場所に到達するということでした。

二つの方法の違いを風景で描いてみましょう。学習とは誤差の底を見つけて降りていく作業です。支援ベクターマシンの誤差地形は、谷が一だけの滑らかな器です。どの斜面から転がっても、球は同じ底に到達します。ニューラルネットワークの誤差地形は、谷が無数に掘られた山脈です。球がどの谷で止まるか事前には分からず、その谷が最も深い場所だという保証もありません。理論家の眼に見れば、前者の地形は扱える世界であり、後者の地形は地図のない世界でした。ところが、ルクンの経験は奇妙な話を聞かせていました。地図のない山脈でも、球は何度も使える谷で止まったのです。なぜそうなるのか、誰も知りませんでした。

この箇所で公正に記しておくべきことがあります。当時の条件下では、ヴァプニクが正しかったのです。データは数千個程度であり、コンピュータは遅かったのです。ニューラルネットワークは学習中に筋違いの谷にはまって止まったり、層が深くなるほど信号が消える問題に悩まされていました。一方、支援ベクターマシンはどの機械で走らせても同じ答えを出し、小さいデータで安定して強力でした。洗練された理論があり、結果も良かったのです。学界と産業界が急速にそちらへ移動した根拠はありました。ルクンのニューラルネットワークが正しかったことは、データと計算がずっと大きくなった後に明らかになる話であり、1995年にはまだその時点が来ていなかったのです。

二人の気質も互いに反対でした。ヴァプニクはソビエト連邦で長く統計理論を磨いた後、アメリカに渡った人物でした。彼にとって知識とは証明に支えられてこそ成り立つものでした。ルクンは幼い頃から回路を触り、コードを書き、手で何かを動作させてきた工学者でした。彼にとって知識とは、まず動作させ、その次に理由を掘り出すものでした。一人は、なぜ正しいのか分からなければ答えではないと言い、もう一人は、正しいなら理由はいつか明かされると言いました。壁一枚を挟んだ二つの部屋は、知能に対する二つの世界観が接触する境界線でもあったのです。

ここで「一般化」という言葉を押さえておきましょう。学習機械の成否は、覚えた問題を再び正答することにはありません。初めて見る問題を正答することにあります。学んだことを新しいデータに広げていく能力、それが一般化です。ヴァプニクの理論が狙ったのはまさにこれでした。パラメータが多い機械ほど、訓練データ全体を丸暗記してしまうリスク、つまり過剰適合のリスクが大きくなります。彼の数学に従えば、数万、数十万個のパラメータを持つニューラルネットワークは、暗記するだけで広げられない機械になるべきでした。ところが、ルクンのニューラルネットワークは理論の予告を裏切りました。パラメータがデータより多いのに、初めて見た小切手の数字を読み取りました。なぜ崩れ落ちないのか。この問いがその廊下の火薬庫でした。

競争は具体的な数字の上で繰り広げられました。試験場はルクンが作った手書き数字データでした。ヴァプニク側の支援ベクターマシンも、ルクン側の合成量み込みニューラルネットワークも、同じ数字を読んで正確性を競いました。1995年には、部門全体がこの対決をまったく論文にしました。手書き数字認識に使える学習方法を一列に並べて比較した研究でしたが、著者リストにはルクンとジャッケル、ボートウ、デンカー、ギユオン、そしてヴァプニクの名前が並んで挙がっていました。異なる陣営の総大将たちが、一編の論文の中で自分の武器を取り出して競い合ったわけです。結果は一方の圧勝ではありませんでした。支援ベクターマシンは理論の予告通り強く、合成量み込みニューラルネットワークも引き下がりませんでした。競争相手と同じ著者リストに名前を連ねる文化。それがその時代のホルムデルの空気でした。

この試験場は後の日に自分の名前を得ます。ルクンはアメリカ国立標準技術研究所が集めていた手書き資料を整え、学習用6万枚と試験用1万枚の数字データを整理しました。MNISTと呼ばれるようになったこのデータは、その後数十年にわたって機械学習のショウジョウバエの役割を果たしました。生物学者たちがショウジョウバエで遺伝法則を試験するように、世界中の研究者が新しいアルゴリズムを作ると、まずMNISTの数字を読ませてみました。ニューラルネットワークと支援ベクターマシンの対決もこのデータの上で長く続きました。誤差率1パーセント前後から0.1パーセントを削ろうとする争いでした。外から見ると小数点の遊びのようでしたが、中から見ると、二つの世界観が毎日成績表を送受する前線でした。

二つの方法の競争が膠着していた1995年、ベル研究所で有名な賭けが成立しました。部門長ジャッケルとヴァプニクが対峙し、ルクンが証人として署名しました。紙二枚に二つの賭けが記されていました。

最初の賭けはジャッケルが掛けました。2000年3月14日までに、なぜ大きなニューラルネットワークがデータをそんなに良く一般化するのかを、人々が理論で説明し出すだろう。ヴァプニクはそんなことはあり得ないという側に掛けました。二番目の賭けはヴァプニクが掛けました。2005年3月14日になれば、正気の人間は誰も1995年式のニューラルネットワークを使わなくなるだろう。ジャッケルはその反対に立ちました。掛かったのはそれぞれほぼ同等の夜食一食分でした。

この賭けの文書は廃棄されませんでした。ルクンは後々、自分の講演スライドにその紙の写真を載せるようになりました。2018年の国際音響・信号処理学会の基調講演でも、聴衆は画面いっぱい映った手書きの賭けの文書を目にしました。勝敗の記録というより、一つの時代の空気を納めた遺物としてでした。当時有数の学習理論家と彼の部門長が、ニューラルネットワークの未来を賭けて夜食一食分を掛けていた時代。その時代の論争がどれほど真摯で、同時にどれほど人間的であったかを、その一枚の紙が証言していたのです。

時間が経ち、結果が出ました。最初の賭けはヴァプニクが勝ちました。2000年になっても、ニューラルネットワークの一般化を洗練して説明する理論は現れませんでした。二番目の賭けはジャッケルが勝ちました。2005年にも、ニューラルネットワークを使う人は消えませんでした。証人として署名していたルクンは、両者の夜食を見守った人として残されました。

軽い逸話に聞こえるかもしれませんが、その中には当時の本当の緊張が込められていました。理論が先か、結果が先か。ベル研究所のある廊下がその問いを間に置いて二つに分かれていました。そしてこの賭けの結末には、かみしめるべき箇所がもう一つあります。二つの賭けの勝者が互いに違っていたという点です。ニューラルネットワークは生き残りましたが、ニューラルネットワークがなぜ良く機能するのかについての理論は期限内に現れませんでした。このズレは今も完全には解けていません。ディープラーニングが世界を覆い尽くした今日でも、巨大なニューラルネットワークがなぜそんなに良く一般化するのかを始めから終わりまで説明する理論はまだ完成していません。ヴァプニクの問いかけは間違った問いではありませんでした。答えが遅れているだけです。

バプニクの背景事情も記しておく価値があります。深層学習の春が訪れた後も、彼は自分の問いかけをやめませんでした。巨大なデータを巨大な計算で押し進める学習を、彼は知能と呼ぶことを躊躇し、人間はずっと少ない事例で学ぶという事実を指し続けました。興味深いのは、この指摘が後年のルクンと重なるということです。大型言語モデルが世界のテキストをすべて読んでもなお、子ども程度の世界理解しかできないというルクンの批判は、データの量が理解に代わることはできないという昔の隣室の理論家の問いかけと、一つの根から分かれ出たように聞こえます。二人は方法をめぐって分かれましたが、安易な成功を疑う眼だけは似ていました。

しかし競争だけがあつたわけではありません。ベル研究所の黄昏期に、ルクンがしがみついていたもののひとつが DjVu (デジャビュ) というイメージ圧縮方式でした。ルクンとレオン・ポトウ、パトリック・ハフナー (Patrick Haffner)、ポール・ハワード (Paul Howard) など、AT&T ラボ・リサーチの同僚たちと共に作りました。前の章で見たとおり、スキャンした書ページを文字層と背景層に分けてそれぞれ圧縮し、一ページの容量を元の数百分の一に減らす方式でした。通信が遅かった時代に、厚い文献を画面に運ぶ道がそのように切り開かれました。この方式はのちにインターネット・アーカイブ (Internet Archive) をはじめとする複数のデジタル書庫の基盤フォーマットとなり、通信網に恵まれなかった地域の学術アーカイブが古典文献をアップロードする容器として長く使われました。神経ネットワークが異端扱いされていた冬の間に、ルクンの部署は紙文明をデジタルに移す仕事で自らの存在を証明していたのです。

DjVuの背景事情には、ルクンの長年にわたる気質が一つしみ込んでいます。開発チームはこの技術をソフトウェア企業に譲渡して商用化する一方で、誰もが使える公開実装も一緒に生かしておきました。閲覧プログラムは無料で公開され、後年はソースコード付きの版も出て、今も保守されています。良いツールは囲い込むより公開する方が、世界をより多く変えるという感覚。この感覚は数十年後、彼が Meta で研究成果の全面公開を入社条件として掲げ、オープンソース・モデルを擁護する立場で再び現

れます。その態度は、この冬の圧縮フォーマットの中にすでにしみ込んでいました。

その間、サポートベクターマシンは機械学習の主流へ上りました。バブニックは1995年に統計的学習理論の本質を、1998年に統計的学習理論を相次いで著し、学界の言語を自分の側へ引き寄せました。サポートベクターマシンは手書き文字を超えて文書分類へ、顔検出へ、遺伝子データ分析へと領土を広げました。1990年代後半の機械学習の教科書と講義はサポートベクターマシンを頂点に置いて組み立てられました。トロントのジェフリー・ヒントン (Geoffrey Hinton)、モントリオールのヨシュア・ベンジオ (Yoshua Bengio) のように神経ネットワークを手放さない人たちは周辺へ押しやられました。皆が洗練された数学が保証される方へ避難先を求めて去るとき、ルクンは隣室で夜遅くまで公式を研ぐバブニックを見ながら自分の場所を守りました。彼が守ろうとしていたのは意地ではありませんでした。データと計算が大きくなる日が来れば神経ネットワークが戻るという、まだ証明されていない予感でした。

神経ネットワークは科学ではない

一つの論文審査意見がありました。ルクンのチームが画像を意味単位に分け取る神経ネットワーク研究をコンピュータビジョン分野の最高学会である CVPR に投げたとき、返ってきた不採択理由はこのようなものでした。「この論文はコンピュータビジョンについて私たちに何も教えてくれない。すべての特徴とルールが単に機械によって学習されているからである(This paper tells us nothing about computer vision, because all the features and rules are just learned by the machine)。」機械が自ら学習したということが、排除の根拠でした。

この文の中に、その時代の科学観がまるごと入っています。その時代の主流が考える科学とは、人間が手で規則を立てて機械に組み込む作業でした。物体の辺縁を探すフィルタも、形態を測定する方程式も数学者が設計しました。機械はその上で最後の分類だけを担当すればよかったのです。人間が特徴を刻み、機械がその特徴を受け取って表を分けること。この分業がその時代の正常科学でした。機械が原始的なピクセルから特徴を自らの手で引き上げるという発想は、彼らには科学ではなく怠慢に見えました。なぜ正しいのか人間が説明できない成功は成功と認めませんでした。

ルクンはこの基準に真正面から反論しました。彼はのちにアメリカ電気電子学会のメディア IEEE スペクトラムとのインタビューでこのように述べました。「当時の学界は鉄壁のような理論があってこそ認め、実際のデータから出た強力な結果は無価値なものとして切り捨てた。工学の問題を解く方法としては非常に間違った、狭い考え方だった(People wanted ironclad theory ... and they would basically dismiss any idea that didn't have it. This was a very bad and narrow-minded way of thinking)。」橋が壊れないという事実より、壊れない理由の証明を優先する態度。ルクンにとってそれは工学を逆さまにする仕事でした。

冬の逆説が一つあります。神経ネットワーク研究の標準文献として残ることになる論文が、まさにこの冬の真つ最中から出たという事実です。ルクンとボトウ、ベンジオ、ハフナーが共著した1998年の論文、文書認識に応用した勾配ベースの学習がそれです。LeNet-5の設計がこの論文に盛り込まれました。発表当時の学界の反応は生ぬるいものでした。神経ネットワークの時代が終わったと信じる人々にとって、神経ネットワークの集大成は過ぎ去った物語の整理に見えたに違いありません。その論文が数万回引用される古典になるまでには、さらに十数年を要しました。論文の運命は論文が出た年ではなく、世界がその論文を受け入れる準備ができた年に決定されるということを、これほどよく示す事例も稀です。

歴史を知る人にとって、この議論は新鮮ではなかったでしょう。蒸気機関は熱力学が確立される前からすでに動いていました。ライト兄弟の飛行機は空気力学理論が完成する前に空を飛びました。作動する物体がまず出現し、それがなぜ作動するのかを説明する理論が後から続く。このようなことは工学の歴史では、むしろ普通の順序でした。しかし1990年代後半の機械学習の学界はその順序を認めませんでした。証明が先行しない成功は幸運と扱われました。

その審査意見の後、ルクンはしばらくコンピュータビジョン学会へ足を運びませんでした。彼は2018年のチューリング賞記念講演で、この時期を振り返りながら、その頃ビジョン学界に論文を投げたことをやめたと明かしました。自分も落ちるのが目に見えている場所に何度もドアをたたき代わりに、彼は神経ネットワーク研究者たちが集まる別の学会とワークショップへ舞台を移しました。降伏ではありませんでした。戦線の移動でした。自分の方法が立つ場所がなければ、立つ場所を別に作るしかなかったのです。この判断は数年後、カナダでヒントンが組むことになる小さな研究同盟、そしてユタの山中で開かれていた神経ネットワーク研究者のワークショップへつながります。次の章で会う物語です。

冬の代償は人間で払われました。神経ネットワークを論文のテーマとした大学院生は、卒業後、職探しの心配をしなければなりません。ベンジオはその時代をこのように振り返っています。自分の下に入った賢い学生たちに神経ネットワークをやらせるために、学生たちの腕を捻るようなことをしなければならなかった(had to twist the arms of my students)と言っています。給料が途絶えることを恐れる学生をつかまえてする研究でした。ヒントンも、ベンジオも、ルクンも、それぞれの都市で同じように孤立していました。三人の手紙と電話が、その時代のほとんど唯一の温もりでした。

しかし、ルクンの反論は結果が良いから理論は後でいいというような反発ではありませんでした。彼は自分たちの方こそ、より原理に忠実な科学だと考えました。彼の神経ネットワークは、でたらめに何でもかき集めたものではありませんでした。生物の視覚皮質から得た直観を幾何学に移した設計でした。同じ特徴検出器を画像のすべての場所に共有して使う重み付け共有、だから対象がどこにあらうと認識する空間不変性。こうした制約は、人間が規則を一つ一つ書き込む代わりに、機械が学ぶべきことと学ばなくてもよいことを予め分ける装置でした。彼が手で作ったのは特徴ではなく、特徴を自らで見つけられる枠組みでした。

同時に、ルクンは反対側の素朴な神経ネットワーク信奉も警戒しました。生物の脳細胞を一つ一つそのままに模倣して大きなコンピュータで動かせば、知能が自然と湧き出すという考え。彼はこれもまた科学になりえないと考えました。彼はこのような態度を19世紀の飛行試行に例えました。鳥の羽と羽ばたきを外見そのままに真似た機械は、地面から浮かび上がりませんでした。空を開いたのは羽ではなく、揚力という原理でした。外見を真似る作業と原理を抽出する作業は異なります。彼が拠り所としたのは脳の外見ではありませんでした。自然が視覚皮質に刻み込んだ計算の原理でした。学習とは生まれつきの規則ではなく、データから引き上げることだという古い信念が、この場所で工学の形で固まりました。

冬には正しいと間違いが即座には判断されません。だからこそ人間の態度が浮き出ます。ルクンは論文が不採択になるたびに方法を変える代わりに、不採択の根拠となった基準そのものを疑いました。結果があるのに認められないなら、問題は結果ではなく、その結果を測る者の眼かもしれないということでした。このような態度は傲慢に見えることもあり、根気として残ることもあります。何として残るかは後の日のデータが判定します。そのときのルクンは、そのデータが後々自分の側に立つだろうという方に自分のキャリアを賭けていました。

付け加えることがあります。ルカンはヴァプニクを敵とは見なしていませんでした。彼が立ち向かったのはヴァプニクの理論ではありませんでした。その理論だけを科学のすべてにしようとする学界の態度でした。ヴァプニクの数学自体は彼も尊重していました。二人は同じ部署で同じデータをめぐって論争し、同じ論文に名を連ね、互いの方法がどこで強くどこで弱いかを誰よりもよく知っていました。問題は、その論争が廊下を飛び出し、学界全体の流行になった時に生じました。廊下では二つの方法が競い合いましたが、外では一つの方法がもう一つの方法を消し去ってしまいました。健全な競争と排除は違うものです。ルカンは冬と呼んだのは、競争ではなく排除でした。

科学が否かをめぐって繰り広げられたこの争いにおいて、どちらも最初から正解を握っていたわけはありませんでした。ヴァプニクの理論は、その時代のデータの上で実際に強力でした。ルカンのニューラルネットワークは、その時代のデータの上ではまだ潜在力に過ぎませんでした。決着は、後日データがインターネットほど大きくなり、演算が画面カードほど速くなってからようやくつきました。その決着の瞬間は、まだこの物語の数年前にありました。

その間、ルカンは静かな迂回路を歩きました。ニューラルネットワークを前面に押し出しにくい時代だったため、彼はAT&T Labs Researchの画像処理研究部を率い、DjVuのような実用的な作業でその地位を守りました。ニューラルネットワークは一時水面下に沈みましたが、消え去ったわけではありませんでした。彼はいつかデータと演算が自らの足で歩み寄り、今では誰も信じていないこの方法を再び呼び起こす日を待ちました。

その待機は漠然とした楽観ではありませんでした。根拠のある計算でした。コンピュータの演算能力は数年ごとに倍増しており、インターネットは世の中の文章や絵をデジタルデータに変えて蓄積しつつありました。ニューラルネットワークの足を引っ張っていた二つの要素、つまり遅い演算と不足しているデータは、時間が経てば自然に解決する問題でした。反対に、人が手で特徴を削り出す方式は、データが大きくなるほど人の手がボトルネックになるという問題を抱えていました。時間はどちらの味方か。ルカンの答えは明確でした。

その待機にも終わりが来ました。2000年代初頭、通信産業のバブルが弾け、旧ベル研究所の二つの後身が並んで揺らぎました。通信機器の好況に乗っていたルーセントは、バブルが崩壊すると破産の瀬戸際まで追い込まれ、AT&Tも株価が暴落して身を縮めていきました。かつて世界でこれ以上ないほど安定した研究職だった場所が、数年の間に人員削減の通知書が出回る場所へと変わりました。ホルムデルとその周辺で基礎研究をしていた人々が、一人また一人と荷物をまとめました。ルカンも2002年に会社を去りました。彼が向かったのは、ニュージャージー州プリンストンのNEC研究所でした。日本企業であるNECが米国に設立した基礎研究機関で、彼は1年余り研究員として留まりました。しばし息を整えるための場所でした。偶然の再会もありました。ヴァプニクもまた、同じ年にAT&Tを去り、プリンストンのNEC研究所に合流していたのです。ホルムデルで壁一枚を隔てて論争していた二人は、冬の終わりに再び同じ看板の下にしばし留まりました。方法は分かれても、二人が歩まねばならない道は似ていました。企業研究所の時代が暮れると、その時代が育て上げた理論家と工学者が同じ船から降りて同じ港を経ていったのです。そして2003年、彼はニューヨーク大学 (NYU) クーラント数学研究所の教授になりました。14年間身を置いた企業研究所の時代を閉じ、大学に戻ったのです。

クーラントは応用数学の名門でした。流体の流れや波動の方程式を扱っていたその研究所が、ニューラルネットワークの研究者を教授として迎えたということは、ルカンの研究を何と見なすかに対する一つの答えでもありました。呪術でも流行でもない、数学がまだ完全に追いついていない応用数学。42歳のルカンは、マンハッタン南部のその建物に研究室を新設しました。ベル研究所の時のように、チーム

が一日で解体される心配はありませんでした。その代わり、すべてを最初から再び築き上げなければなりませんでした。学生を集め、研究費を獲得し、ニューラルネットワークという言葉はまだ忌避する学界の中で、居場所を広げなければなりませんでした。

遠い後日、この冬は別の名前と呼ばれることになります。2019年に米国計算機学会がルカンとヒントン、ベンジオにチューリング賞を授与する際に明らかにした功績の理由には、2000年代初頭、この三人が学界の懐疑論の中にあってもニューラルネットワークに対する信念を守り抜いた少数派であったというくだりが含まれていました。冬を耐え抜いたこと自体が功績の一部として記録されたのです。賞は発見を称えるのが普通ですが、この賞は耐え抜いたことも共に称えました。その耐え抜いた現場こそが、まさにこの章で通り過ぎたホルムデルの黄昏でした。

企業の実験室は彼に小切手読み取り機とDjVuを与え、解体命令と眠った特許も与えました。大学は別のものを与えることができました。四半期の業績に揺さぶられない時間、そして共に冬を耐え抜く学生たちでした。冬の真ただ中で人ができることは多くありません。自分が正しいと信じる方向を諦めず、春が来るまで種を湿った土の下に埋めておくこと。ルカンがホルムデルの黄昏の中で成したことはそれでした。今、種はマンハッタンへと移されました。ではその春は、どこでどのように始まったのでしょうか。

弁護士

kimkj.com

第6章 マンハッタンのジャズとディープラーニング結社

ヤン・ルクン、人工知能の父

クーラント研究所へ

2003年の秋、マンハッタンのグリニッジ・ビルディングの地下のジャズクラブ。43歳のフランス人が舞台から少し離れた席に座ってトランペットの音を聞いていました。ヤン・ルクンでした。昼間はワシントン・スクエアの隣にあるニューヨーク大学の建物で白い黒板に数式を書き、夜になるとソーホーとビルディングの狭い路地を歩いてこのようなクラブを探していました。即興演奏には楽譜がありません。演奏者は規則をあらかじめ決めていません。その場で隣の人の音を聞きながら次の節を作り出します。規則を書き込まなくても良いものが生まれるというその方式は、ルクンが機械に長年望んでいたことに似ていました。

彼がニューヨークに来た道は平坦ではありませんでした。前の章で見たように、彼の15年間の根拠地だったベル研究所は、通信企業AT&Tの分割とルーセントの解体の中で散り散りになりました。小切手の上の手書き文字を読んで、アメリカを行き来する小切手の相当部分処理していた彼のシステムも、紙を圧縮していた彼の技術も、会社が分裂する渦の中ででんでんに売られていきました。画像処理研究を主導していた職を降りた後、彼は一時的にNEC研究所を経由し、2003年にニューヨーク大学の正教授として職を移しました。大学が彼に与えた職位はヤコブ T. シュワルツ教授職でした。

彼が身を寄せた場所はクーラント数学研究所と神経科学センター、二つの建物でした。一方の足はコンピュータと数学の方に、もう一方の足は脳研究の方に置いた計算です。このちぐはぐに見える職が彼にはぴったり合いました。機械の学習を自然の脳から学んでこようというのが彼の古い考えだったので、二つの建物の間に置かれた彼の机はその考えをそのまま移し置いた場所でもありました。クーラント数学研究所はドイツのゲッティンゲンからナチスを逃れて渡ってきた数学者リヒャルト・クーラントが創設した場所です。数理物理学の方法を磨いてきた、応用数学の古い聖地でした。世界で数えるほど厳密な証明が行き交う部屋の真ん中に、学会で死んだ道と呼ばれていたニューラルネットワークを抱えて彼が歩んで入った計算です。

二つの建物の間に座った彼が抱いた構想は古いものでした。ベル研究所の時代、彼は猫の視覚皮質を調べた二人の神経生理学者、デイビッド・ハッセルとトルステン・ウィーゼルの記録から自分の機械の枠組みを得ました。目が世界を丸ごと飲み込まずに狭い領域ごとに分けて見ており、点から線へ、線からパターンへと層々に積み上げるというその発見をいいます。神経科学センターにかかった彼の一方の足は装飾ではありませんでした。機械がどのように学ぶべきかという問いの答えを、彼は依然として生きている脳から探していました。

外見上は昇進であり安息地でした。ベル研究所という産業の温室を失った彼に、大学の終身教授職はしっかりした屋根でした。しかし、その屋根の下空気は彼が望んでいたものとは異なっていました。その頃、クーラントの数学者たちと機械学習理論家たちが美しいと考えたものは別にありました。サポートベクトルマシン、つまりSVMという方法でした。皮肉なことに、この方法の根はルクン自身の昔の隣人にありました。SVMを創設したウラジミール・ヴァプニクは、ベル研究所の時代にルクンの隣室で働いていた人でした。一つの廊下で育った二つの考えが、今や互いに異なる陣営の旗になっていました。

SVMにはニューラルネットワークにはないものがありました。答えを必ず一つに絞り出すという数学的保証でした。誤差の谷が飯碗のように一つだけであるので、どこから転がっても底の一点に達します。紙の上できっぱりと証明できるこの性質を人々は愛しました。カーネル法やベイズ推論も同じ理由で重視されました。対照的に、ニューラルネットワークは調整する値が数百万個であり、誤差の谷には底が複数あるので、どの落とし穴に落ちるかわかりませんでした。数学者たちの見方では、そのような機械を学習させることは証明のない手遊びでした。動作するがなぜ機能するかを誰も説明できないボックス、彼らはニューラルネットワークをそう呼びました。研究費もその判断に従って流れました。証明される方法には金がつき、機能するがなぜ機能するかわからない方法にはつきませんでした。

学会外のコンピュータビジョン現場も彼にとっては見知らぬ国でした。その時代に写真の中の物体を認識する仕事は、人間が手で特徴検出器を彫刻して作る技術で通じていました。エッジをキャッチする検出器、質感を測る検出器、明るさの方向を数える検出器を人間が一つ一つ設計して貼り付けており、その上にSVMのような分類器を乗せる方式でした。研究者の実力はすなわち良い特徴を手で設計する技術でした。ある者はそのような検出器一つで人生の名前を得ました。ルクンの主張はこの板を全く否定するものでした。特徴を人間が彫刻せず、機械が生画素から自ら学ぶようにしろという言葉は、その分野に人生を捧げた人々には自分の仕事を消せという音に聞こえました。この論争がどれほど深かったかは、次の章で別に見ます。

ルクンにとっての反論の根拠は紙の上の論理だけではありませんでした。彼はすでに動く機械を作った人でした。ベル研究所で彼の畳み込みニューラルネットワークは人間が走り書きした郵便番号を読み、しばらくの間、アメリカの小切手を実際に処理しました。証明がなくても動く機械が目の前にありました。それでも学会は動作を証拠として認めませんでした。なぜ機能するかを説明できなければ、機能することさえ認めない雰囲気でした。機能するがなぜ機能するかわからないことと、機能しないがなぜ機能しないかを知ること。学会は後者を選びました。ルクンにとって耐え難かった部分がここにありました。

彼の答えは黒板でした。彼はクーラントの白い黒板を多層ニューラルネットワークの微分幾何学で満たしました。複数の層が重なった構造をどのように端から端まで微分して学習させるか、その幾何学がなぜ美しいかを疲れずに書きました。彼の指摘は常に同じ場所を狙っていました。学会が線形モデルにしがみついたのはそのモデルが正しいからではなく、証明しやすいからだということでした。紙の上で整った方を選んだだけ、実際のデータの上で多層ニューラルネットワークが示す性能を無視したということです。どちらが本当に世界に適しているかではなく、どちらが論文を書くのに便利かで方向が決まったと、彼は長い間考えてきました。

この論争は学問的好みの問題のように見えますが、根はより深かったです。知能とは何かという問いがその下にありました。知能は人間が規則で植え付けることで生じるのか、それとも機械が世界を見て自ら育てるのか。クーラントの数学者たちとルクンの論争は、表面上はどちらの方法がより正確かをめぐって繰り広げられましたが、底の方では、この古い問いが再び分かれていました。ルクンは一方を選びました。規則を書き込む代わりに例を示す方、人間ではなくデータが機械を教える方でした。パリの図書館で忘れられたパーセプトロンの論文を発掘していた学生時代から彼の中にあつたこの選択は、一度も変わらなかったし、数十年後に大規模言語モデルをめぐって繰り広げられる論争でも同じ顔で戻ってきます。

人工知能に冬が来たのは最初ではありませんでした。1969年、マービン・ミンスキーとシモア・パートがパーセプトロンの限界を数学で指摘した後、最初の冬が来てニューラルネットワークは長く忘

られました。ルクンは学生時代に忘れられたそのパーセプトロンをパリの図書館から再び取り出した人です。1980年代に逆伝播とともに短い春が来ましたが、1990年代に二番目の冬が再び降り積もりました。彼は二度の冬の両方をニューラルネットワークの側で過ごしました。他の人が去った時に留まり、後に他の人が殺到した時にも居場所を守るこの頑固さが、彼の人生を貫く模様です。

昼間の黒板と夜のクラブの間で、43歳のルクンは再び下から始める人になっていました。ベル研究所で彼はチームを率いて世界が使う製品を作った人でした。それすべてが会社とともに散り散りになった後、彼は学生を教え論文を審査される教授の職に戻りました。楽な道もあったでしょう。証明される方法に乗り換えて学会の認可を受ける道のことです。彼はそうしませんでした。夜になると楽譜のない音楽を聞きに出かけ、昼間は誰も支持しない方法を黒板に書きました。失望を言葉で述べる代わりに、彼は書き続けました。

大学の自由には代価が伴いました。ベル研究所でラリー・ジャッケルが金を惜しむなと言って後押ししていた時代と異なり、教授の研究は自分で研究費を獲得する必要がありました。しかし、ニューラルネットワークという主題ではアメリカで研究費を獲得することは難しかったです。審査で負ける主題だったので当然のことでした。カナダの支援がそれほど大切だった理由がここにもあります。ルクンは大きなチームの代わりに小さな研究室で始め、派手な機器の代わりに何人かの学生と黒板で再スタートしました。世界が使う製品を作った人が、誰も注目しない問いへ静かに戻ったのです。

しかし、クーラント内部の冷たい視線は耐えられるものでした。大学外の、彼の論文が審査される学会の空気ははるかに厳しかったです。その場所では、ニューラルネットワークという単語自体が罪になっていました。

CIFARとディープラーニング結社

ジェフリー・ヒントンは後年、自分が経験したあることをよく語っていました。彼がニューラルネットワーク論文を学会に提出したところ、プログラム委員会からこのような返答が来たということです。すでにニューラルネットワーク論文を二編も受け取ったので、これ以上は受け取ることはできないということでした。論文の内容が問題ではありませんでした。主題が問題でした。ヒントンはこのことを1912年のアルフレド・ウェゲナーになぞらえていました。大陸が動いていると主張したが、その力を説明する方法がないという理由で学会から嘲笑された地質学者のことです。正しいか間違いかを議論する前に、ドアの前から追い出されるという状況が同じでした。

1990年代後半から2000年代初期にかけて、人工知能には2番目の冬が訪れていました。ニューラルネットワークの研究者たちにとって、この冬は学問の検閲として降りかかりました。コンピュータビジョンと機械学習の大きな学会であるCVPRとNIPSで、論文のタイトルや抄録にニューラルネットワークやバックプロパゲーションという単語が含まれていると、査読委員たちは本文を読むこともなく掲載拒否の判定を下したと、その時代を通り抜けた人たちは異口同音に語ります。査読とは本来、内容を読んで値打ちを判定する仕事です。その手順が単語ひとつで止まってしまったのです。ヒントンのようにすでに名高い教授ですら、こうした扱いを受けたとすれば、名もない若い研究者たちの立場は問うまでもなく分かるというものです。

実際のところ、その波紋は若い研究者たちに生計の恐怖として広がりました。ニューラルネットワークを研究すると、卒業後に職にありつけないという話が大学院に流れ回りました。論文が掲載されなければ職にありつけず、テーマが掲載不可なら、いくら上手くやっても壁でした。モントリオール大学のヨシユア・ベンジオは、この時代をこのように回想しています。失職を恐れて去ろうとする学生たちの

腕を無理矢理ねじって、引き留めなければならなかったというのです。師が弟子に、この道を行けば損だと知りながらも、留まってくれるよう懇願する立場。連結主義と呼ばれたこの学派の火は、そのようにして数人の掌の中で、かろうじて生きていたのです。

10年近くを少数で持ちこたえるということは、特別なことでした。1、2年認められないことと、学問の門そのものが閉ざされたまま10年を持ちこたえることでは、重みが違いました。大抵はその手前でテーマを変えます。残る側が愚かに見える、それが冬の論理です。ルクンとヒントン、ベンジオが残ったのは、もうじき春が来るという確信のためではありませんでした。確信する根拠はその時どこにもなかったのです。彼らが頼りにしたのは、自分たちが正しいと信じるひとつの思想と、その思想を共に信じる2、3人の存在だけでした。

ルクンとヒントン、ベンジオは散り散りにならないことにしました。トロントとモントリオールとニューヨークにそれぞれ職を得た3人は、定期的に連絡を取り交わしながら静かな連合を組みました。ルクンはのちのインタビューで、このミーティングを笑いながら『ディープラーニング陰謀団』と呼びました。外部からは、3人を指して『カナダ・マフィア』と呼ぶ人もいました。3人ともカナダとの縁があったからです。ヒントンはトロントに、ベンジオはモントリオールにいました。そしてルクンは若き日、トロントでヒントンのポスドクを務めた経歴がありました。

3人の関係は競争であり友情でした。それぞれ異なる都市で異なる問題を掘り下げていましたが、互いの論文を読み、引用し、学会とワークショップで顔を合わせて夜通し語り合いました。1人が壁にぶつくと、残りの2人がそばにいました。学問の世界で10年を少数で持ちこたえるには、正しいという信念だけでは足りません。自分たちだけが狂っているわけではないことを確認してくれる、同じものを見る目がもう2つあるという事実が必要でした。後年、3人がチューリング賞を共に受けたのは偶然ではありませんでした。冬を共に過ごした人たちが、共に春を迎えたのです。

彼らが案出した賢い手立ては、ひとつの名前でした。ニューラルネットワークという単語が査読の場で反射的な拒否を呼び起こすので、単語を変えることにしたのです。複数の層を深く積み重ねるという意味を込めて、『ディープラーニング』、つまり『深層学習』という新しい名前を世に送り出しました。方法は変わりませんでした。看板だけ付け替えたのです。『深い』という言葉には自負も込められていました。浅いモデルではこの世の層状の構造を捉えられず、深い構造でこそ捉えられるという、彼らの信念がその一字に込められていたのです。古い名前にかかった自動拒否のフィルタを迂回しようとするこの改名が、後年、ひとつの時代全体の名前になるとは、その時3人も知らなかったでしょう。

名前を変えた効果は緩やかでしたが明白でした。ニューラルネットワークという古い看板を外した論文たちが、じわじわと査読の門を通り始めました。同じ方法、同じ人たち、異なる名前でした。学問の世界で名前がここまで力を持つとは、苦い教訓でもありました。良い思想が良いという理由だけで受け入れられず、包装を変えてはじめて中身を見えてくれるということだからです。しかし生き残ることが先決でした。3人は自尊心よりも生存を選び、『ディープラーニング』という名前はこうして学界の敷居を少しずつ越えていきました。

名前だけでは冬は退きませんでした。彼らを実際に救ったのは、カナダから来たお金でした。アメリカの研究費配分機関がニューラルネットワークから手を引いた時、ヒントンはカナダ高等研究所、つまりCIFARの扉をたたきました。2004年、彼は『神経計算と適応的知覚』という名の長期研究プログラムをCIFAR内に立ち上げることに成功しました。英字頭文字を取ってNCAPと呼ばれたこのプログラムは、後年『学習する機械と脳』という名前に変わります。ルクンとベンジオはこのプログラムの正式な

研究員となりました。

CIFARが与えたのは研究費だけではありませんでした。行き場のなかった連結主義の学者たちが、1年に何度か一所に集まって夜通し議論する屋根ができたのです。その場にはコンピュータ科学者だけではありませんでした。脳を研究する神経科学者、物理学者、認知心理学者がひとつのテーブルに混在していました。『機械はどのように学ぶのか』という問いと『人間の脳はどのように学ぶのか』という問いを並べて置き、議論しました。ルクンがかねてから抱いていた考え、つまり機械の学習を視覚皮質から学ぶべきだという着想は、この場では奇妙な話ではありませんでした。プログラムの名前に『知覚』という言葉が含まれたのもそのためでした。魔女狩りを逃れてきた人たちの避難所であり、学問の仕切りを行き来する市場でもありました。ニューヨークとトロントとモントリオールを結ぶ三角形が、この屋根の下で堅くなっていったのです。

その三角形が初めて大きな結実をもたらした年が2006年でした。ヒントンと同僚たちが、多層ニューラルネットワークを学習させる新しい方法を発表したのです。それまで、深いニューラルネットワークには長い病があったのです。層が幾重にもなると、下の層まで誤差信号が適切に降りていかず、上の層だけが学び、下の層は学習が止まってしまいました。誤差が層を遡るうち、だんだん薄れて消えてしまう、この病を、人々は長く解くことができなかったのです。ヒントンの処方箋は順序を変えることでした。最初から全体を一度に学習させるのではなく、層をひとつずつ別々に、あらかじめ習熟させておいたのです。下の層がデータの理を先に自ら学んで定位置を占めてから、その後で全体を磨くのです。正解表なしでデータだけを見てあらかじめ習熟するこの事前学習が、深いニューラルネットワークの長い病を抑えました。冬眠に入った学界の心臓が、この論文ひとつで再び打ち始めたのです。

ただ、この心臓が大きく打つまでには、なお時間が必要でした。2006年の成果は小さなコミュニティの中で興奮をもたらしましたが、外のコンピュータビジョンは相変わらず人間が手で刻み込んだ特徴に支配されていました。写真の中の物体を認識する大会で優勝するのは、ニューラルネットワークではなく、巧みに設計された特徴検出器でした。ディープラーニングがその盤さえ覆すには、あと2つのものが来る必要がありました。機械に見せるべき膨大な量の写真と、その写真に対応できるほど速い計算装置でした。両者とも、2000年代半ばではまだ到着していなかったのです。その時点で、3人はなお冬の終わり際に立っていました。

この勝利にも影が差していました。ドイツのユルゲン・シュミットフーバーは、後々強く異議を唱えました。層をあらかじめ習熟させるという着想は、すでに1991年に自分の研究室から出たもので、ヒントンとルクンがこれを引用から落とした、というのです。誰が先だったのかを巡るこの争いは、1度では済まず、2018年に3人がチューリング賞を受けた時に再び大きく火を噴きます。その話は後ろに別に場を設けました。ここで指摘しておくべきことは、この時期の勝利が静かで狭かったという事実です。世間はまだ『ディープラーニング』という言葉を知りませんでした。数人がカナダのお金と互いの肩に寄りかかりながら冬を過ごしていただけなのです。

ディープラーニングが世間に初めて実力を示した場所は、目ではなく耳でした。2009年頃から、ヒントンの弟子たちが音声を変換する問題に深いニューラルネットワークを持ち込むと、誤り率が目に見えて低下しました。マイクロソフトとグーグル、IBMの音声認識が相次いでニューラルネットワークに切り替わりました。人々の手に持つ電話機が言葉を聞き取り始めたのです。学会の査読の場はなおも凍りついていましたが、産業の現場ではすでに氷が溶け始めていました。3人が共に書いた宣言が学術誌の前に正式な文書として置かれるまでには、その後もさらに何年も要します。2015年5月、ルクンとベンジオとヒントンは科学学術誌『ネイチャー』に『ディープラーニング』という解説論文を並んで

掲載しました。その原稿が最初に声に出して交わされた場所は、学術誌ではありませんでした。ユタの雪に覆われた山の中、暖炉の前でのことでした。

スノーボードの生存者たち

ユタ州のスキーリゾート、スノーボードは冬になると雪が深く積もります。山小屋の暖炉に薪が燃え、その前にチョークと黒板を手にした人たちが集まって夜を明かしました。学会と呼ぶには格式がありません。定められた論文の形式もなく、厳粛な演壇ありませんでした。手にチョークを握りしめて、物理と認知の原理を即興で議論する場でした。昼間はスキーをし、夜間はニューラルネットワークについて話す、学問と遊びの境界が曖昧な集まりでした。マンハッタンのジャズクラブと似たようなところがありました。このミーティングの名前は『計算のためのニューラルネットワーク』でした。

夜になると、小屋の風景はいつも同じでした。スキーを脱いだ人たちが暖炉の周りに輪になって座り、誰かが立ち上がって黒板に式を書き始めると、議論は夜明けまで続きました。発表の時間も、質問の順序も決まっていませんでした。間違った考えを言っても、恥をかかせることはありませんでした。出来の悪い着想も最後まで聞いてくれました。外の学会で論文がひとつの単語で人を切るのとは正反対でした。ここでは誰も『掲載拒否』という言葉を使いませんでした。冬の山の中のこの緩さが、凍りついた学問の唯一の息の掛かる穴でした。

始まりは1986年にまで遡ります。ベル研究所ホルムデルの物理学者ジョン・デンカーとラリー・ジャッケルといった人々がこの会合を企画しました。ニューラルネットワークの冬を物理学者たちが守ったのには理由がありました。主流のコンピュータ科学がパーセプトロンの衰退後、ニューラルネットワークを放置しシンボリックAIに資金を注ぐ一方で、物理学者たちは別の視点からこの問題を見ていました。1982年、物理学者ジョン・ホップフィールドは磁石の物理を借りて、記憶するニューラルネットワークのモデルを構築しました。無数の小さな磁針が互いに引き合い押し合いながら最も安定した配列に沈み込むように、ニューラルネットワークも誤差が最も小さい状態に沈み込みながら何かを記憶するというイメージでした。このイメージは物理学者たちに馴染み深いものでした。だからこそ彼らはニューラルネットワークを魔法ではなく物理として見て、計算と物理と脳について一堂で語られる山小屋の部屋を作ったのです。ここに名前が出たラリー・ジャッケルは、後タルクンをベル研究所に連れてきたまさにその人です。お金を節約して有名になることはできないという、その部長の言葉です。

ベル研究所の物理学者たちがニューラルネットワークを守ったこの歴史は偶然ではありませんでした。物理学はもともと複雑なものを少数の力で説明してきた学問です。無数の粒子が互いに押し合い引き合いながら安定した状態に沈み込む様子を、物理学者たちは手慣れたように扱ってきました。ニューラルネットワークが誤差を減らしながらある状態に沈み込む様子が、彼らにとっては馴染み深いイメージだったのです。ルクンが後に長く執着することになるエネルギーという概念も、この物理の言語から生まれました。そしてスノーボードで夜通し働いた若者の中の何人かは、後にディープラーニングの中心へと進んでいきます。この冬の小さな学校が次の世代を育てたのです。

二度目の冬が深くなるとスノーボードは行き場のない者たちの避難所となりました。ルクンとヒントン、ベンジオは毎年冬になるとこの山小屋に集まってきました。学会で門前払いされ、研究費から落とされ、弟子たちの前途を心配していた人々が、一年に一度この山でお互いの顔を確認したのです。暖炉の前で彼らが静かに反論していったのはSVMのカーネルという手法でした。カーネルは人が事前に定めた尺度でデータ間の類似性を測定する方式です。尺度が人間の手から出ているという点がルクンには引かかっていました。

彼はこのように述べました。私たちが生きている真の世界は、人間が手で削り出した一つの尺度には含まれるには深すぎて層が多すぎることでした。世界はピクセルが集まってエッジになり、エッジが織り交ぜられてパターンになり、パターンが組み立てられて物体になるという形で層々と構成されています。こうして層々に積み重ねられた性質を、彼はコンポジショナリティと呼びました。世界がそのように構成されているのであれば、世界を読む機械も生のピクセルから複数の層を経由して自ら学び上がっていくべきだというのが彼の要点でした。人が尺度を一つ削り出して終わらせてはならず、層と層の間の規則まで機械が自ら習得するようにしなければならないと、彼は黒板いっぱい数式を書き記しました。このコンポジショナリティという言葉は彼の人生全体を貫きます。数十年後に世界モデルについて語る時にも、彼は同じ言葉を繰り返しています。

論議だけが交わされたわけではありません。ルクンは実物を持ってきました。彼の長年の同志であるフランス人数学者レオン・ポトウと一緒に作ったソフトウェアでした。二人の縁は古いものでした。他の人たちがニューラルネットワークを気にもかけていなかった若き日々、二人は研究費も機材も不十分なため、市販されていたアミガというパーソナルコンピュータに取り組んで底からシミュレータを書きました。学習に必要な計算を代行するソフトウェアが世の中に存在しないため、直接作るしかなかったのです。二人がSNと名付けたこのプログラムは、そのように貧しい二人の研究者の手から生まれ、ベル研究所を経由してスノーボードの山小屋までついてきました。ルクンはこのコードを山小屋で1行ずつ指し示しました。理論ではありませんでした。動くプログラムでした。

この場で彼とポトウが力を込めて論じたのは確率的勾配降下法でした。学界は訓練データ全体を一度に計算して値を修正する重い方式を正しいとみなしていました。収束するという証明が整然としていたからです。二人の考えは反対でした。データが一つ入力されるたびにすぐさま誤差を測定して値をちょっと調整する軽い方式が、局所最小値に陥らずより速く正しい方向へ進むというものでした。手書き文字のような似た例が溢れるデータにおいて、この方式が特に力を発揮しました。今日ほぼすべてのニューラルネットワークがこの方法で学習しています。山小屋の黒板で磨き上げられた習慣です。

ルクンが常に動くコードにこだわった背景には彼の根がありました。彼は理論家である前にエンジニアでした。紙の上で美しいものと機械の上で実際に動くものは異なるということを、彼は手で経験して知っていました。スノーボードでの議論が熱くなる時、彼が提示するのは概して証明ではなく実行結果でした。この数値を見よ、この機械が実際にこのくらい読み取った、という形でした。カーネルの優雅な定理の前で彼が提示したのは動くプログラムでした。この頑固さが後に彼を他者より先に実物を提示する人間にしたのです。

主流がSVMを愛した理由はありませんでした。その方法には鋼鉄のように堅固な理論がありました。どのくらい間違えるかを事前に計算することができ、答えが一つに収束するという保証がありました。不確実なものを嫌う心は科学者の美德でもあります。その点で主流の選択を嘲笑うだけではできません。スノーボードに集まった人々が選んだのは反対側でした。保証はないが本当の世界でより良く機能する方、汚いが実際に学ぶ方でした。どちらが正しかったかは数年後、データとハードウェアが判定してくれます。ただしその判定が下される前まで、彼らは証明の代わりにお互いの顔を信じて冬を過ごしたのです。

スノーボードの夜を特別にしたのは、結局のところ人々でした。外の世界がニューラルネットワーク研究者を狂人扱いしていたその冬に、ここでは若い大学院生がルクンの量み込みニューラルネットワークを美しいと言うことができました。人が定めた偏見のフィルターを取り除き、機械が自ら空間の規則を習得させるその優雅さを、この部屋では思いっきり議論してよかったのです。検閲される学会で門前払いされた考えが、この暖炉の前では夜通し議論の種になりました。ベル研究所が産業が温めた部屋で

あったなら、スノーボードは吹雪の真ただ中で灯された灯火でした。

山の下の大きな学会には数千人が集まっていた。スノーボードの山小屋に集まった人数は多くても数十人でした。世界の秤で量れば一方は主流で、他方は周辺でした。しかしその秤が十年以内に逆転するのです。山小屋の数十人が夜通し磨き上げた方法が、後に数千人が集まる学会の発表会場を満たすようになります。周辺が中心となり、中心だったカーネル法が教科書の一章に退きます。ただしその逆転が起きるまで、この人々にとってスノーボードは世界で自分たちの言葉が通じる数少ない場所でした。

この冬が人々に残した痕跡は異なっていました。ある者は学界を去って企業へ、ある者は専門を変えました。残った人々も代償を払いました。認められた論文の代わりに却下の手紙を集め、優秀な学生を他の分野に奪われ、同僚たちの申し訳なさそうな視線に耐えなければならませんでした。ルクンはこの時代について大きく恨むことはありませんでした。自分が正しいと信じることを続けただけだと、淡々と述べるのが常でした。正しさを証明する確かな道は議論ではなく時間だということを、二度の冬が彼に教えたわけです。

しかしその灯火も山の中でしか輝きませんでした。山を下りるとまた冬でした。この人たちが山小屋で守った火種を雪の外の広い世界で大きく灯すには別のものが必要でした。カナダの避難所とユタの暖炉だけでは不十分でした。都市の中心に、生活を営み人を育てる堅固な拠点が必要でした。ルクンがその拠点を築いた場所は、彼が毎晩ジャズを聴いていたマンハッタンでした。

データ科学センターという拠点

ルクンがニューヨーク大学に来てからちょうど十年が経った2013年、大学は新しい機関を一つ設立しました。データ科学センターでした。ルクンが初代所長を務めました。設立の方式は並ではありませんでした。ニューヨーク大学はこのセンターを法学部や経営学部といった既存の単科大学の下に組み込むのではなく、どこにも属さない独立した学問領域として置きました。

この決定には時代を読んだ判断が根底にありました。データが世界を変えるだろうという予感が学界と産業に広がっていた時期でした。大学はこの新しい学問を古い学科の片隅に置くのではなく、独立した運営を与えました。法と経営と芸術と工学の隔壁のどちらにも束縛されない座を与えたのです。ニューラルネットワークを持ちながら行き場のなかった人に、今度は大学が最初に手を差し伸べました。十年前に彼を冷遇していたまさにその大学でした。ルクンはこの拠点を自ら設計し、コンピュータ科学と数学、そして脳神経科学の間の壁を取り壊しました。三つの分野を一つの屋根の下に束ねることが彼の意図でした。学会が彼を追い出した時、彼は学会の外に自分の家を建てたのです。

センターは急速に成長しました。データを扱う学問に人が集まっていた時期であり、ニューヨークという都市の力も大きかったのです。金融とメディアと芸術が一堂に集まったこの都市にはデータが溢れており、そのデータを読むことができる人が必要でした。ルクンが築いた拠点を中心に、ニューヨークは西部のシリコンバレーとは異なる質の人工知能の中心地として位置を確立していきました。大学と産業が密接に結びついており、研究がすぐに現場へ流れていく都市でした。冬にニューラルネットワークを守っていた人の後ろに、今や一つの都市が一緒に動き始めました。

開放的な研究は言葉だけに終わりませんでした。この時期ニューヨーク大学とその周辺から生まれた手法とコードは、論文と共に世界に公開され、誰もがダウンロードしてそれぞれの機械で実行してみることができました。ディープラーニングが少数の大企業の秘密に留まらず、世界中の大学院生の共有ツールとなったのには、このような態度の役割がありました。カーネルの時代には手法は論文の中の式として留まっ

ていましたが、ディープラーニングの時代には手法はダウンロードして実行するコードとして広がりました。ルカンが底から始めてシミュレータを書いていた長年の習慣が、今や一つの分野全体の文化になっていこうとしていました。

この仕事で彼が越えなければならない壁は、意外にも同じ建物の中にありました。クーラントの数学者たちです。10年前、彼を証明のない手遊びだとして冷遇した、まさにその人たちのことです。彼らは依然としてディープラーニングのでこぼこした誤差表面をブラックボックスだと考え、答えが一つに保証されるSVMやカーネル、ベイズ推論の領域にとどまっていた。ルカンはその理由を回りくどくなく指摘しました。数学者たちが線形モデルにとどまっているのは、理論を書くのがより簡単だから (the theory is easier to write) だということです。世の中に合っているからではなく、論文を書きやすいからという、10年前と同じ指摘でした。

指摘だけにとどまったわけではありません。彼は数学者たちを敵ではなく、同僚として引き入れようと思いました。ディープラーニングの中には、まだ誰も証明できていない難しい数学の問題が山積みでした。複数の層が重なった構造において値がどのような幾何学を描いて定着するのか、確率的勾配降下法がなぜうまく窪みを避けるのか、紙の上で綺麗に解けた答えはありませんでした。ルカンは、この未解決の難題こそ応用数学が征服すべき新大陸だと、クーラントの数学者たちを説得しました。証明が難しいということは避ける理由ではなく、飛び込む理由だということです。優秀な数学専攻の学生たちが、一人また一人とデータ科学センターの共同研究員として移ってきました。10年前、彼を押し退けていた城壁が、彼の拠点の一部となっていったのです。

この転換が一日で成し遂げられたわけではありません。長い間、互いを疑わしく思っていた二つの陣営でした。ルカンは数学者たちに、ディープラーニングがなぜ動くのかまだ誰も知らないという事実を隠しませんでした。むしろ、その分からないという点を餌として差し出したのです。答えがすべて出ている問題には大きな数学は必要ありませんが、なぜうまくいくのか分からないままよく動く機械には、解くべき新しい定理が満ちているということです。証明を愛する人々にとって、これほど魅力的な招待はありませんでした。10年前、ニューラルネットワークを手遊びと呼んでいた人々のうちの何人かは、今やその手遊びの数学を掘り下げる人となりました。

拠点を築いただけではありません。人も育てました。ニューヨーク大学に腰を据えた後、ルカンの研究室では若い才能が育ちました。後日DeepMindで働くことになるライア・ハドセル、後日ニューヨーク大学の医療画像教授となるスミット・チョプラといった人々が、彼とともに新しい方法を鍛え上げました。彼らが作ったものの一つが、二つの入力がどれほど似ているかを機械が自ら測れるように並べて立てたニューラルネットワークでした。正解ラベルをたくさん付けて教える代わりに、同じもの同士は近くに、違うもの同士は遠くに置かれるよう、表現を自ら整列させる方式でした。正解ラベルに頼らずデータのきめだけで学ばせるというこの考えは、彼が後日「自己教師あり学習」と呼び、生涯のテーマとする方向性の種でした。ただし、その時はまだ早すぎました。インターネットからかき集めた巨大な写真の山もなく、計算をこなせる高速なハードウェアもありませんでした。良いアイデアが世に出るには、ツールが数年遅れをとっていたのです。

追われていた人が、今や人を引き受ける立場に立っていました。カナダの資金に頼って冬を越した人が、今度は自分の屋根の下に若者たちを呼び寄せ、共に冬を越していたのです。避難所を得て生き延びた者が、自ら避難所になったのです。ベル研究所も、スノーボードも、他人が作ってくれた避難所でした。データ科学センターは、彼が初めて自分の手で築いた温室でした。

この拠点の扉は、内側にだけ開かれていたわけではありません。ルカンは、研究が数人の手に閉じ込められることを嫌いました。論文は公開されるべきであり、コードは共有されなければならない、方法は誰でも検証できなければならないというのが、彼の長年の信念でした。冬に自分の方法が検閲された経験がある人ほど、開かれた扉の価値を知っていました。単語一つで論文を切り捨てたその閉ざされた扉が、どれほど多くの良いアイデアを殺したか、彼は身をもって経験しました。後日、彼が産業界へ移る際にも最後まで手放さなかった条件が、まさにこの開放でした。ここで育った信念が、次の章の物語をあらかじめ決定づけます。

この頃のルカンを見ると、彼はもはや辺境の孤独な人ではありませんでした。ベル研究所で手書き文字を読んでいた人、二度の冬をカナダの資金とユタの暖炉で耐え抜いた人が、今やマンハッタンの真ん中に自分の名を冠した研究拠点を構え、若い人々を育てていました。彼が夜な夜な歩いたソーホーの路地と、昼間に埋めていたクーラントの黒板と、今や彼が所長として座るデータ科学センターが、一人の人間の中で繋がりました。ニューヨークはそうして、開放型ディープラーニング研究の一つの中心となっていきました。ここで育ったアイデアとツールが、後日、世界中の研究室へと流れていきます。誰もが使うディープラーニングのツールとなったPyTorchも、この流れの上で生まれました。

世間の態度がひっくり返るのには、長い時間はかかりませんでした。2012年を過ぎ、あれほど蔑まれていた方法が写真を認識する大会を席巻すると、雰囲気が一変しました。ニューラルネットワークを研究すると飢えるという言葉が出回っていた場所で、今やニューラルネットワークを知る人を連れて行こうと企業が列をなしました。GoogleやMicrosoft、Baiduのような会社がディープラーニングの研究者をめぐる競争を始めました。ヒントンが二人の弟子と立ち上げた小さな会社をGoogleが買い取ったのもこの頃です。冬を耐え抜いた少数が、春が来ると互いに連れて行こうとする貴重な人材になったのです。ルカンにも、間もなくそのような提案が舞い込みます。

拠点が強固になると、外の視線が彼へと向かいました。冬にニューラルネットワークにしがみついていた学者が、今や産業界が欲しがる人物になっていました。2013年末、カリフォルニアのある若い創業者が、このマンハッタンの研究室の扉を叩きます。Facebookを作ったマーク・ザッカーバーグでした。彼が差し出した提案と、ルカンが突き返した三つの条件が何であったのかは、次の機会に紐解きます。ここで予め言っておくべきことは、ルカンが大学を去らなかったという事実だけです。彼はニューヨークに残り、データ科学センターの所長の座を手放しませんでした。

二度の冬を過ぎたこの人々が、まだ知らなかったことが一つありました。春が来ており、その春は静かではないはずでした。長い間蔑まれてきた畳み込みニューラルネットワークが、世界中が見守る前で大きく燃え上がる日は遠くありませんでした。その地震がどのように襲ってきたのかは、後で別に見ます。ただ、春が来る前にルカンが何を鍛え上げておいたのかを、まず指摘しなければなりません。カナダの資金とユタの暖炉とマンハッタンの黒板で、彼が守り抜いたのは一つの方法ではなく、世界を見る目でした。その目がどのように鍛えられたのか、時計を数年前に巻き戻してみます。

弁護士

kimkj.com

第7章 特徴を設計していた時代の終わり

ヤン・ルクン、人工知能の父

エネルギーベースモデル

黒板に曲線が一本現れました。谷間と丘陵が交互に連なった、山脈を横から見た姿のような線でした。2006年、ニューヨーク大学クーラント研究所のある講義室で、ヤン・ルクンはチョークを握ったまま学生たちの方に振り返りました。彼が指し示したのは確率ではありませんでした。エネルギーでした。谷間の底を二度叩きながら、彼は言いました。機械が答えを見つけるということは、この地形で最も低い場所へ転がり落ちていくことと同じだと。

学生たちにとってその図は見慣れないものでした。彼らが教科書で学んだ機械学習は確率の言語で書かれていたからです。谷間と丘陵だなんて、物理学の授業に出てくるような図でした。ルクンは気にしませんでした。彼は難しい概念をいつも手で掴める図に変えて説明する人でした。式を黒板いっぱい並べるかわりに、球が転がり落ちていく地形を一つ描いておき、そこから話を始めました。その地形の中に、彼が人生をかけて取り組んできた疑問が折り込まれていました。機械が世界をどのように理解するようにさせるのか、という問いです。

その頃、機械学習をしていたほぼすべての人は確率の言語で考えていました。2000年代の学会会場はグラフモデルとベイズ定理、事後確率といった言葉でいっぱいでした。入力Xが与えられたときに出力Yが出る確率を計算し、その確率が最も高い答えを選ぶのです。この方式は優雅で、数学的によく整理されており、論文査読者たちが好む言語でした。確率で書かれた論文は厳密に見え、そうでない論文はどこか信用できないものに見えました。

ルクンはその空気にうまく溶け込むことができませんでした。彼は確率が間違っていると言いませんでした。確率が万能ではないことを指摘しただけです。世界のある問題は、正解が一つにはつきり決まりません。ぼやけた写真の中の顔が誰なのか、半ば隠れた文の空白にどんな単語が入るのか、答えはつねに複数あります。そうした問題に無理矢理確率の衣をかぶせようとするときに、人々が払う代償がありました。人々があまり口にしない通行料です。

通行料の正体はこうです。何かが確率になるためには、あり得るすべての場合の合計が正確に1になる必要があります。合計を1に合わせるには、まず全体を知らなければなりません。あり得るすべての答えを一箇所に集めて、その大きさを残さず足した値、統計学者たちがパーティション関数と呼ぶその分母を求める必要があります。名前は大げさですが、やることは素朴です。世の中にあり得るすべての答えの重さを全部測って、今この答えがそのうち何パーセントを占めるのかを教えてくれる秤です。

答えが少数ならば、この秤量は難しくありません。手書きの数字を0から9までの10個の中から選ぶ問題なら、10個を足すだけで済みます。問題は答えが画像や動画の時に起こります。横縦数百ピクセルの写真1枚が持ち得る姿の種類は、宇宙の原子の数より多いのです。その全部の写真を1枚も落とさず秤にかけて重さの総和を求めろという要求、それが高次元の世界で確率を扱おうとする人が直面する壁でした。計算がちょっと長くかかる程度ではありませんでした。原理上終わらせることができない仕事でした。宇宙が冷え果てるまで計算機を回し続けても、答えの目録さえすべて調べることができません。

ルクンの提案は、その通行料をまったく払わないというものでした。確率という規格を降ろして、XとYがどれほどうまく合致しているかをエネルギーという一つのスコアでのみ測ろうというのでした。

よく合った組み合わせには低いエネルギーを付け、ずれた組み合わせには高いエネルギーを付けます。ここでは全体の合計が1になるように無理矢理合わせる理由が消えます。この答えがああ答えより良いか、スコアの高低さえ正ければよいのです。秤全体を持ち上げる必要もなく、二つの皿の傾きだけを見ればよいというわけです。彼はこの構成をエネルギーベースモデル (Energy-Based Model, EBM) と呼びました。

たとえるなら、このような違いでした。確率論者は試合の勝敗を予測しながら各チームが勝つ正確な百分率を付けようとします。そうするには世の中のすべての場合を考えて分母を合わせなければなりません。エネルギー論者は、ただどのチームの戦力がより尤もらしいかスコアを付けるだけです。二つのチームのどちらが優れているかを決める際に、世の中のすべての試合を数える理由はありません。答えが少数の狭い問題では、二つの方式の違いはよく表れません。答えが画像のように無限に近い広い問題になれば、一方は歩いていくことができ、他方は第一歩で座り込んでしまいます。

その名は物理学から借りてきたものです。谷間に置かれた球は誰に押されなくても底へ転がり落ちて止まります。水は低いところを探して流れます。自然のものはエネルギーが最も低い場所に自ずから位置を占めます。実は、この発想自体が非常に新しいものではありませんでした。1980年代に物理学者ジョン・ホップフィールドは神経網をエネルギーが低い状態に沈み込む物理系に例え、ジェフリー・ヒントンはボルツマンマシンという名で、その上に確率の衣をかぶせました。ルクンが行ったのは、これらのばらばらな流れの上に一つの屋根をかけることでした。確率であれ何であれ、エネルギーというスコア一つで世界の組み合わせを測るすべての方法を一つの傘の下に集めたのです。

この傘には実用的な意味がありました。異なる方法で散らばっていた研究たちが一つの言語を共有するようになったからです。どのモデルがなぜ崩れ、どのモデルがなぜ耐えるのかを、エネルギー地形という同じ図の上で比較することができました。ルクンはこの共通言語を大切にしました。方法ごとに異なる用語に隠されて見えなかった共通の構造が、エネルギーという言葉に言い換えられるや否や一目瞭然になったからです。

この傘の下で推論はこのような姿でした。前のビデオフレームが入力Xとして与えられると、機械は次に来ることができる多くの候補Yの中からエネルギーが低い谷を一つ見つけて下っていきます。答えを一度に計算するというより、地形の底を手探りして最も深い谷間に到達することに近かったのです。写真の隅の一部が消えたとしましょう。その場所に入ることができる図は一つではありません。複数の尤もらしい答えがあり、複数の筋の通らない答えがあります。エネルギー地形は、尤もらしい答えが置かれた場所を谷間にし、筋の通らない答えが置かれた場所を丘陵にしてくれます。機械はその地形を転がり落ちて、一つの谷間で停止します。

エネルギーベースモデルの異なる点がここにありました。答えを出すことが一度の計算ではなく、一つの探索になるということです。通常の神経網は入力を入れると出力がすぐに飛び出します。エネルギー地形では異なります。機械は複数の候補を前にしてどちらの谷がより深いか比較しながら、答えに向かって徐々に下っていきます。考えに時間をかけるというわけです。ルクンはこれが逆に人間の思考に似ていると考えました。私たちも難しい問題の前では複数の可能性を頭の中に並べて比較してから、尤もらしい方に心を決めるではありませんか。答えがすぐに飛び出す反射ではなく、候補の間を行き来する思案。彼は本当の推論にはそうした思案の余地があるべきだと信じていました。

この考えには遠い未来へ伸びる一つの枝が付いていました。答えを見つけることがエネルギーが低い場所を手探りすることなら、計画を立てることも同じ枠で見ることができます。今の状況から望ましい

結果に至る複数の行動の連鎖の中で、エネルギーが最も低い連鎖を見つければ、それがまさに計画です。未来の複数の分岐を頭の中に広げて、その中から尤もらしい道を選ぶこと。ルクンにとって推論と計画は同じ地形の上で起こる同じ種類の探索でした。この種がどのような木に育つのかは、この本の後半で明かされます。

では地形そのものは誰が作るのか。学習が作ります。実際に世界で一緒に現れたXとYが置かれた場所には、エネルギーが深く掘られて谷間ができるようにし、その外のずれた場所にはエネルギーが立ち上がって丘陵になるように、地形を少しずつ削り出していきます。数千、数万回の調整を経て、地形は世界の相をまねるようになります。ここまでは順調に見えます。

ここに落とし穴が一つ隠れていました。正解の場所のエネルギーを低くすることだけを熱心に繰り返すとどうなるでしょうか。地形全体が下へ沈み込んで、平らな底になってしまいます。世界全体が谷間になれば、谷間は何の意味もありません。機械は入力を見向きもせず、どこでも同じ答えを出す、この上なく怠惰な解答に押さえつけられます。何を聞いても同じ答えばかりする学生のようなのです。ルクンはこの現象を表現の崩壊と呼びました。エネルギーベースモデルを実際のデータの上で飼い慣らす際に直面する、非常にしつこい敵でした。その後彼は20年間、この敵と格闘することになります。

表現の崩壊がなぜこんなに厄介だったのかは、もう少し掘り下げると明らかになります。機械に世界を理解しろと命じると、機械は往々にして理解しているふりをする手っ取り早い道を見つけ出すのです。すべての入力に同じ答えを出すと、見た目には誤差が減ったように見えるからです。試験紙のすべてのマスに同じ答えを書いても採点が雑なら満点をもらうのと同じです。本当の理解と理解のまねごとを区別すること、それがエネルギー地形を形作る人が常に背負った宿題でした。ルクンが対照と正則化という二つの流れをそれほど工夫して分けたのも、このまねごとをどう止めるかが成否を分けたからです。

崩壊を防ぐ道を彼は二つの流れにまとめました。一つは対照的な方法です。正解のエネルギーを下げると同時に、正解でない偽の組み合わせをわざと作ってそのエネルギーを押し上げるのです。谷間を掘りながらその周囲の土地と一緒に盛り上げるというわけです。ルクン自身もこの方式の道具を作りました。2006年、彼の研究室のレイア・ハードセルとスミット・チョプラは、同じ物体の二枚の写真は表現空間で相互に引き付け、異なる物体の写真は相互に押し出すように学習する方法を提示しました。二つのイメージを並行して通す双子の神経網でした。この方式は、物体が少し異なる角度から撮られても機械がそれ二つを同じものとして捉えさせてくれました。手書き文字や顔を扱っていた時代には、よく機能しました。問題は規模でした。画像が大きくなり、動画が長くなるにつれて、押し出す必要のある偽の組み合わせの数が手に負えないほど増えました。広い原野でモグラを捕まえるために全部の地面を叩くようなものでした。

もう一つは正則化の方法です。偽の対を動員する代わりに、機械の内部の表現が含むことができる情報の量に事前に制約を課しておくことです。容器の大きさを縮めると、水は自動的に深いところに集まります。余分な表現空間を狭めておくと、機械は正答の周辺にのみ狭く急な谷を削るしかありません。ルクンは昔からこちら側に関心を持っていました。複数の講演で彼は、自分の賭けは正則化された潜在変数ベースのエネルギーモデルにかかっていると言っていました。他の方法は結局このやり方が勝つだろうという予感でした。

この予感は言葉に留まりませんでした。同じ時期に彼の研究室は、正則化の着想を実際の装置に移しました。コレイ・カボッキョオグルとマルコーレリオ・ランチャトは、イメージを数個のまばらな断片でのみ表現するように機械を調整し、その断片でもとのイメージを甦らせました。表現に使える材料の

数を意図的に節約しておく、機械は無造作なパターンを覚えるのではなく、本当によく現れる決定的なパターンだけを拾い集めることになります。節約が理解へと繋がったわけです。人間がフィルターを削ってくれなくても、機械はこのように学んだパターンで世界の断片を自らが認識し始めました。まだ写真全体を読み取るレベルではありませんでしたが、特徴を手ではなくデータで形作するという原理は確かに実現していました。

2006年、ルクンはスミット・チョープラ、ライア・ハードセル、マルコーレリオ・ランチャト、フジエ・ファンとともに、エネルギーベースの学習を一つの教科書として整理しました。確率という馴染み深い王国を離れ、エネルギーという地形の上で学習の文法を組み直すガイドでした。当時でも、この文書に注目する人は多くありませんでした。機械学習の主流は別のところに目を向けており、エネルギーだの地形だのという比喻は、物理学を慕うフランス人の嗜好のように見えました。確率の王国からエネルギーの地形へと席を移したこの選択は、嗜好ではありませんでした。それは後にその最後の思想まで貫く背骨になります。世界をピクセル一つ一つで描き出すのではなく、その下に隠れた構造だけを掴もうという思いの最初の骨組みが、その黒板の曲線の中にすでに入っていました。

手で削った特徴との戦い

2000年代中盤、コンピュータが写真の中から物を見つけ出すコンテストの順位表は、常に似たような方式で埋められました。ヨーロッパの研究者たちが毎年開いていたPASCAL画像認識コンテストがその頃の激戦地でした。上位に来たシステムの内部を開いてみると、前側には人間が手で設計した装置が、後ろ側には比較的軽い分類器が位置していました。前側の装置の名前は大概SIFTやHOGのようなものでした。

カナダのデイビッド・ロウが1999年に発表したSIFTは、写真の中で明度が急激に折れ曲がる点とエッジを取り出す精密なフィルターでした。フランスのナヴァ・ダラルとビル・トリクスが2005年に発表したHOGは、画像を細かく分割し、各ピースで明度勾配の方向を測りました。どちらも人間の頭から出てきた公式でした。当時の標準的な組み立て方はこのようなものでした。写真にSIFTをかけて特徴点を抽出し、その特徴を何グループかに分けて、各写真でどのグループがどのくらい含まれているかを数えて、その数値表をサポートベクターマシンという分類器に渡して、ネコか自転車かを判定する。写真の中で何が重要かを決める難しい判断は、前側で既に人間がすべて行っていました。機械に残されたのは、すべて用意されたテーブルの前でスプーンで食べることだけでした。

この構図をルクンはある巨大な工場に例えました。世界の数学者と工学者たちが黒板の前に集まり、写真から何を見なければならぬかを手で削り出す特徴工学の工場でした。良い特徴を一つ設計すれば論文になり、コンテストの順位が何段階も上がり、その特徴が次の数年の標準になりました。勤勉で利口な人たちがその工場で働きました。しかし、彼らが作るのは機械の目ではありませんでした。機械に被せる眼鏡でした。眼鏡の度数は人間が決めました。

この局面でルクンと主流の学界は知能を見る目が異なりました。主流にとって特徴設計は専門家の誇るべき技術でした。良い特徴を削り出すには、その分野を長く掘った人の目利きが必要であり、その目利きこそが科学の貴重な資産と考えられていました。ルクンにとってそれは先延ばしにされた課題でした。人間が特徴を削るということは、機械がまだ学んでいない部分を人間が代わりに埋めているということだからです。代わりに埋めてくれたその場所を機械に返してやること、それが彼が考える次のステップでした。

ルクンが不満だったのはまさにその地点でした。彼が見るところ、機械が学ぶべきことは最後の分類の薄い皮だけではなく、生のピクセルから出発して線やエッジを經由して目や輪へ、最終的には猫や自転車という概念まで登る梯子全体でした。世界は、その梯子の下段をすべて人間が手で削ってはめ込んでいました。そのようにしている限り、機械が見ることができるのはいつも人間が事前に想像した範囲内に閉じ込められるしかありませんでした。設計者がエッジと明度勾配を重要だと定めておくと、機械はその先の質感を永久に学ぶことはできません。

この争いは、彼の世代が初めて経験することではありませんでした。音声認識分野は一步先に同じ川を渡りました。長い間、その分野は言語学者が発音規則と文法を手で書き込む方式を信じていました。その信念を覆した人がフレッド・ジェリネックでした。彼は規則の代わりに統計を持ってきて、膨大なコーパスから音と単語の確率を機械が自らが学ぶようにしました。彼が残したと伝えられる一言がこの転換を圧縮しています。「Every time I fire a linguist, the performance of the speech recognizer goes up.」つまり、言語学者を一人解雇するたびに音声認識器の性能が上がるということです。正確な表現と時点は記録によって少しずつ異なり、ジェリネック自身は後年、表現を柔らかくして回想しました。それでもこの言葉が長く引用される理由は一つです。人間が手で書き込んだ規則が、粗く不規則な実際のデータの前でいかに簡単に壊れるかを、冗談の形を借りて正確に突いたからです。

ローゼンブラットが1957年にパーセプトロンを世に出してからほぼ半世紀、パターン認識の通説はあまり揺らぎませんでした。人間が削った固定された特徴抽出器を前に置き、学習可能な分類器を後ろに付ける。この二段構造が正解として固まっていました。主流の学界はここに理論的根拠まで与えました。生のピクセルを複数層のニューラルネットワークに直接入れると、誤差の表面がでこぼこになって機械が浅い落とし穴に陥ってしまうということでした。局所最小値と呼ぶその罠に陥ると学習が止まると言われました。数学がそう言うのだから、わざわざ危険な道を行く理由はないという雰囲気でした。

ルクンはその根拠を信じていませんでした。信じない資格もありました。彼は既にベル研究所時代に、手書き数字をピクセルから概念まで一度に学ぶ畳み込みニューラルネットワークを作り、米国郵政の手書き認識を読み取った人でした。機械が特徴を自らが学ぶことができるということを、彼は言葉の代わりに動作する装置で既に示していたわけです。視覚を研究する学界は、その成果を小さな例外として脇に置きました。手書き数字のように背景がきれいで大きさが均一なおもちゃの問題でしか通用しない技術に過ぎず、雑多が混ざった実際の写真には役に立たないというのが大多数の判定でした。

ルクンはその判定に言葉ではなくデータで対抗しました。2004年、彼はレオン・ボトゥ、フジエ・ファンと一緒に、おもちゃの人形と自動車、飛行機のような物体を複数の角度と複数の照明条件下で撮影したデータセットを作成しました。名前はNORBでした。きれいな手書きと異なり、同じ物体が光と方向によって全く異なる見え方をする厄介な素材でした。彼の畳み込みニューラルネットワークは、手で削った特徴の助けなしにこれらの物体を認識しました。人間がエッジフィルターを挿入してくれなくても、機械が三次元物体の形状をピクセルから自らが学ぶことができるという証拠でした。おもちゃの数字でしか通用しないと言われていたその技術が、おもちゃの人形でも、照明が変わっても通用しました。

だから、2000年代のニューヨーク大学の研究室は静かな戦線になりました。世界が手で削った特徴を王座に置いている間に、ルクンと彼の学生たちは反対側で機械が特徴を自らが引き出す方法を掘り下げました。マルコーレリオ・ランチャトとコレイ・カボッキョオグルのような弟子たちは、正解表なしで画像の断片をまばらに表現するよう導く方法を洗練させました。データで頻繁に一緒に現れるパターンを機械が自らが部品のように選び取るやり方でした。その後、この方法で学んだ特徴を使って街の歩

行者を探し出し、写真のすべてのピクセルに空か道か建物かという名前を付ける作業にまで進みました。派手な成果が毎年出るわけではありませんでした。これらの研究はコンテストの順位表を覆すことはできず、論文はしばしばSIFT系に押されました。

学会の雰囲気は微妙でした。写真認識を扱う大きな学会の査読評では同じような文句が繰り返されました。結果は興味深いが検証された特徴を使っていない、なぜSIFTと比較しなかったのかということです。検証された特徴とは、すなわち人間が手で削った特徴でした。ニューラルネットワークで特徴まで学ぶという試みは、スタートラインから証明の負荷をより重くしていました。他の人が使う部品を使わなかったという理由で低い得点を受け、流行していないという理由で良い考えが埋もれるのをルクンは何度も目撃しました。彼はこの時代を長く記憶していました。後年、彼が学問の開放性と再現可能性を異例に強調するようになるルーツには、自分のやり方が流行から外れているという理由で軽視されたこの数年がありました。

それでも彼は退かず、同じ主張を繰り返しました。複数の学会の演壇から彼は、手で削った特徴の時代が間もなく終わると言いました。聴衆の反応は生ぬるいものでした。興味を持つ人もいれば、多くの人が礼儀正しく頷いた後、自分の席に戻ってSIFTを使い続けました。ルクンは説得に成功しませんでした。彼は説得の代わりに時間を信じていたようです。いつかデータと計算が十分になれば、わざわざ言葉で勝つ必要なく、結果が代わりに語るだろうと。

長い間、彼が自分の方法を証明できる舞台は限定的でした。手書き数字データであるMNIST (エムニスト) がその舞台でした。彼が作ったニューラルネットワークは、その上で何度も最高記録を更新しました。世界はその記録を認めながらも、いつも同じ言葉を付け加えました。数字ではうまくいくが、本当の写真は違うのではないか。ルクンは本当の写真でも通用することを示したかったのですが、そのためにはより多くの写真とより高速な計算が必要でした。良い考えを手にしながらも、それを証明する材料がまだ世界に到着していない時期でした。材料を待つことは、考えを思いつくことと同じくらい長い忍耐を要求しました。

手作業で削った特徴には、見えない天井がありました。設計者がシフトに詰め込んだ知恵の分だけ機械が見ることができたからです。データがどれだけ増えても、前段のフィルタが固定されていれば、性能はある地点で止まりました。特徴まで学ぶニューラルネットワークは異なりました。データが増え、計算が速くなるほど、一緒に成長する余地がありました。その余地が実際の力に変わるには、あと二つがさらに熟成する必要がありました。十分に大きなデータと十分に高速な計算です。両者ともこの章の時間が終わりに近づくころようやく到着します。

当時最高の写真認識システムがどのように動作していたかを見れば、この天井がさらに明らかになります。それらのシステムは、写真から抽出した数多くの特徴点をまるで単語のように扱いました。写真一枚を特徴という単語の袋として見なし、どの単語が何回現れたかを数えて写真を分類しました。視覚単語袋と呼ばれたこの方法は巧妙で、しばらくの間ランキング表を支配していました。ただし、その単語の辞書は人が定めた特徴から来ていました。辞書にない単語は写真の中にあっても、システムの目に入りませんでした。データをもっと注いでも、辞書がそのままであれば、理解も元の場所にとどまりました。ルクンが狙ったのはまさにその辞書でした。辞書さえも機械が自ら作るようにするということができた。

ところが、ルクンはこの対立を性能の問題だけとは見ていませんでした。それは知能を何と見るかの問題だったのです。知能が人が事前に埋め込んだ規則の総和であるなら、特徴は手作業で削るのが正し

いです。反対に、知能が世界を経験しながら自ら成長するものであるなら、特徴は機械が発見する必要があります。彼はパリの図書館でニューラルネットワークを初めて出会ったときから、後者の側に立っていました。手作業で削った特徴との戦いは、彼にとって方法の対立ではありませんでした。古い信念の延長でした。

戦線の勝敗はこの章の時間の外で決まります。2011年、彼の方法に近い別のチームが先に壁を押し、2012年には写真認識競技の局面が一夜にして逆転します。その地震の話は次の章のものです。地震が来る前まで、ルクンが守ったのは成果ではなく方向でした。ランキング表が手作業で削った特徴を示していたその数年間、彼は常に反対側を指していました。機械が学ぶべきことは、正答を選ぶ最後の手つきではなく、世界を見る眼そのものだ。

猫のように学ぶ機械

生後一年の赤ちゃんの前にカップを置き、布で覆います。赤ちゃんは布をめくってカップを再び見つけます。カップが見えない瞬間でも、それがそこにあることを赤ちゃんは知っています。誰もその事実を文章で教えたことはありません。赤ちゃんは見て、触って、落としながら、世界がどのように動くかを自分で学びました。ルクンは多くの講演でこのような場面で話を開いていました。知能の大きな塊は、誰かが正答を教えたからではなく、観察だけで満たされるということを示すためでした。

彼が寄りかかった古い洞察が一つありました。ロボット工学の開拓者ハンス・モラヴェックが1988年にまとめたパラドックスです。人間に難しいことは機械には簡単で、人間に簡単なことは機械には難しいということでした。難しい数学問題を解いたり、膨大な知識を暗記したりという高度な推論は、むしろコンピュータが容易にこなします。機械が人間のチャンピオンを倒した最初の競技がチェスだったという事実が、このパラドックスをよく示しています。反対に、生後一年の赤ちゃんや猫が難しくしていることは、カップをつかむこと、足元の障害物を避けること、転ばずにバランスを取るといった感覚と運動の知能は、機械には厳しいほど難しいのです。知能の段階において、最も登りにくい段は頂上ではなく底でした。

ルクンはその底の正体を物理世界の常識と呼びました。物体は支えがなければ下に落ちるということ、隠れても消えないということ、ぶつかったガラスは割れるということです。このような直感ほどの教科書にも文章として書かれていません。赤ちゃんはこれを言葉ではなく目で学びます。彼はこの学びの順序を、発達心理学が長く積み重ねてきた記録に好んで重ね合わせました。パリの認知科学者エマニュエル・デュプーは、幼児が世界を知っていくスケジュールを複数の実験で詳しくまとめ、ルクンはそのスケジュールを自分の講演スクリーンに頻繁に映しました。

そのスケジュールはこのような流れです。新生児は最初の数か月間、腕と脚をうまく制御できず、ただ横たわって世界を見ています。自分で何かに触ったり、動かしたりする力もほとんどありません。その受動的な眺めだけで、脳の中の表象は猛烈に調律されます。生後数か月で、赤ちゃんは生き物と無生物を区別し始め、少し後には隠れた物体がまだそこにあることを知ります。発達心理学者たちは、これを巧妙な実験で確認しました。赤ちゃんに物理法則に逆らう光景、例えば壁を通り抜ける物体や空中で止まっているボールを見せ、赤ちゃんがそれをどのくらい長く見つめるかを測定します。驚いた赤ちゃんは長く見ます。予測が壊れたからです。

ルクンが好んで引用した場面がそこにありました。幼い赤ちゃんに、おもちゃの車が転がっていき、空中にふわりと浮いている、重力に逆らう劇を見せます。生後数か月の赤ちゃんは、その奇妙な光景を無関心に見ています。まだ頭の中に重力という規則が入っていないからです。生後九か月ほど成長した

赤ちゃんに同じ劇を見せると、赤ちゃんは目を丸くして長く見つめます。世界がそうあるべきだと自分の中に立てた予測が壊れたからです。その驚きは、赤ちゃんの脳の中に、今や重力と慣性の常識が定着したという証拠でした。誰も試験問題として教えなかったのに。半年余りの間に、赤ちゃんは世界の物理学を目だけで習得しました。

発達心理学者たちはこのような早期の認識について長く議論していました。どこまでが生まれつきのスケッチで、どこからが経験で満たされたものをめぐってでした。エリザベス・スベルケのような学者たちは、赤ちゃんが物体と数、空間に関する枠組みを非常に早い時期から持つと見ていました。ルクンの関心はその議論の勝敗にはありませんでした。生まれつきであれ、早期に学んだものであれ、人が試験問題として教え込んだものでないことが、彼にとっての要点でした。世界の常識はラベルではなく、観察から来ていました。であれば、機械にもその道を開いてやる必要がありました。

この学びは人間だけのものではありませんでした。猫は誰の説明も聞かずに扉の幅を推し量って飛び越えます。カラスは瓶に石を入れて水を引き上げることができます。生まれたばかりの子馬は数時間で自分の足で立ちます。これらの動物たちに正答表を配った先生はいませんでした。彼らは世界を見、経験し、予測が外れるたびに少しずつ修正しながら学びました。ルクンが作りたかった機械もそのようなものでした。膨大な正答表を食べて育った優等生ではなく、世界を横目で見つつ自ら理の道を悟る猫のような機械でした。

ルクンがこの図から引き出した結論は明確でした。赤ちゃんが数か月間に目で受け入れる情報の量は、人間が一生涯かけて読むすべての文字を合わせたものより、はるかに大きいのです。世界は文で説明される前に、まず見え、触れられます。ですから、真の知能を機械に植え付けようとするなら、正答表が付いた薄いデータではなく、正答なしに流れていく厚い感覚の大河で学ばせなければならないということでした。後年、彼はこの直感を目が丸くなるような数字に変え、世界を説き伏せますが、その種はすでにこの時代の確信の中に入っていました。

その数字はこのようなものでした。人間の目から脳に続く視神経は片方に約百万本、両眼を合わせると二百万本に達します。その繊維が毎瞬間に運ぶ視覚情報を秒単位で換算すると、毎秒数メガバイトに達します。四歳の子どもが目覚めていた時間をすべて合わせると、およそ一万六千時間で、その間に両目で流し見た情報の総量は推し量ることさえ難しい大きさになります。ルクンはこの山のような量を本と比較しました。人間が生涯休まずに文字だけ読んでも、目で数年の間に受け入れるその情報の山には、足の先も届きません。世界を理解するのに必要な情報のほとんどは、文字ではなく感覚から来ました。文字だけで知能を形作ろうという考え方がなぜ半分だけなのか、彼はこの一つの天秤で説明しました。

正答なしに学ぶというのは、曖昧に聞こえます。ラベルを貼ってくれる人がいないのに、機械は何を基準に学ぶのでしょうか。ルクンの答えは時間でした。世界は自分の正答を自ら与えます。今このモーメントの場面を見せ、次のモーメントの場面を隠してから、機械に次を当ててみろと言えばよいのです。次の場面が到着したら、それが即座に採点表になります。ボールが転がっていく映像で次のフレームを予測していると、機械は誰も重力を教えなくても、ボールが上に跳ね返らないことを自分で知るようになります。世界が教師で、時間が試験紙です。後年、この方法は自己教師あり学習という名前を得て、ルクンはそれを知能の暗黒物質と呼ぶようになります。見えないが宇宙のほとんどを満たすその物質のように、観察で積み重ねる学びが知能の体積を満たす見えない大部分だということでした。

ここで最初の節のエネルギー地形が戻ってきます。次の場面を予測することには正解が一つではないからです。前の車が左に曲がるかもしれないし、右に曲がるかもしれないし、そのまま進むかもしれま

せん。一つの正解を出すように強要すると、機械はすべての可能性を平均化したぼんやりとした幽霊を描きます。だからこそ、複数のあり得る未来を谷にし、理にかなわない未来を丘にするエネルギー地形が必要だったのです。世界を観察だけで学ぼうとする夢と、確率を捨ててエネルギーを選んだ選択は、最初から一体でした。異なる節で育った二つの考えが、実は同じ根を分け合っていたのです。

ルカンは予測こそが知能の心臓部だと見なしました。次に何が来るかを描き出せる存在だけが世界を理解しているのです。ボールがどこへ転がっていくか、コップがどこへ落ちるか、相手が何を言うかあらかじめ描き出してみる能力。その絵が正確であるほど私たちは世界をよく知っているのであり、絵が外れた時、私たちは驚き学びます。赤ん坊が宙に浮いた車に驚くのも、その驚きで世界の模型を書き換えるのも、すべて予測の仕事でした。機械に予測をさせることは、一つの技を教えることではなく、世界を理解する入り口を開くことだったのです。

この確信を彼は一つの絵に圧縮しました。ケーキの比喻です。2016年の神経情報処理システム学会 (NeurIPS) の基調講演で、ルカンは大きなチョコレートケーキを画面に映し出し、こう言いました。「If intelligence is a cake, the bulk of the cake is unsupervised learning, the icing on the cake is supervised learning, and the cherry on the cake is reinforcement learning.」知能がケーキなら、その本体の大部分は正解なしに学ぶ学習であり、上を薄く覆うクリームは正解を教えてくれる教師あり学習、そして頂点に載せたチェリー一粒は報酬で学ぶ強化学習だという意味でした。人々はクリームとチェリーの作り方はそれなりに知っているが、肝心のケーキ本体の焼き方は知らないと彼は付け加えました。三年後、彼はこの比喻を手直しし、本体の名前を正解なしに学ぶ学習から自己教師あり学習へと変えて呼びます。世界が自ら自分の正解となるその方式を指す、より正確な名前でした。

チェリー一粒に関する話は、彼の性格をよく表しています。その頃、人工知能の脚光は大部分が強化学習、つまり試行錯誤で報酬を追う方式に注がれていました。囲碁で人間に勝ったアルファ碁がその象徴でした。ルカンはその方式の価値を認めつつも、限界を冷静に指摘しました。もし機械が試行錯誤だけで運転を学んだらどうなるだろうかと、彼は問いかけました。崖下へ数万回ほど転がり落ちて車を粉々に壊した後ようやく、崖に向かって走るのはいくつもの悪い選択だという自明な事実を辛うじて覚えることになるだろうと。

彼が好んで挙げる別の例も同じところを指していました。人間は二十時間余り練習すれば運転を学びます。何度かヒヤリとする瞬間を経験することはあっても、崖から落ちて死んでみてもから運転を身につけるわけではありません。すでに世界がどのように回っているかを知ったままハンドルを握るからです。水が滑りやすく、鉄が重く、人が突然飛び出してくる可能性があることを、十八年の間に目で積み重ねてきたおかげです。人間はそのような出来事を一度も経験せず、頭の中で前もって描き出してみます。頭の中に世界の模型が入っていて、実際にぶつからなくても結果を想像するからです。世界を学ぶ機械に必要なのも、まさにその内面の模型だと彼は見なしました。観察によって積み上げた常識がなければ、想像もありません。

この頃、ルカンの講演画面にはいつも同じ数枚の絵が繰り返し登場しました。世界を見つめる赤ん坊、塀を越える猫、ケーキの一切れ。聴衆が笑っている間も、彼が伝えようとした言葉は真剣でした。今私たちが作った機械はケーキのチェリーとクリームだけを辛うじて真似るだけで、肝心の本体の焼き方を知らないということ。その本体の焼き方を見つけることに残りの人生を懸けるということ。拍手は丁重でしたが、その方向へ共に歩んでくる人はまだ多くありませんでした。

ですから、この章で扱った三つの歩みは、結局一つの方角を指していました。確率の規格を捨ててエネルギーの地形を選んだこと、手で削り出した特徴の代わりに機械に特徴を発見させようとしたこと、そして正解なしに観察だけで世界の常識を積み上げる学習に賭けたこと。三つとも同じ信念の別の顔でした。知能とは人があらかじめ入れておく規則ではなく、世界を経験しながら自ら育っていくものだという信念です。世間が手で特徴を削り出していた時代に、ルカンはその時代の終わりをすでに心の中に思い描いていました。猫が誰からも学ばずに世界を知るように、機械もそうやって学ぶことができるだろうという考え。その考えがいつ、どんな名前で再び世に出るかは、その時誰も知りませんでした。ただ彼は、それがやって来るといふ方に自分の名前を懸けていたのです。

弁護士

kimkj.com

第8章 イメージネットという地震

ヤン・ルカン、人工知能の父

2012年 AlexNet

トロントのある部屋では、ゲーム用グラフィックスカード2枚が6日間休まず回転していました。アレックス・クリジェフスキー(Alex Krizhevsky)がエヌビディア GeForce GTX 580 2枚をコンピュータに挿して、ニューラルネットワークを1つ訓練していたのです。カード1枚のメモリは3ギガバイトでした。そのカードはもともと、画面に3次元のゲーム風景をリアルタイムで描き出すために作られたものでした。クリジェフスキーはそのカードにゲームの代わりに写真を読ませました。100万枚を超える写真を、1枚ずつ、ピクセル単位で。ニューラルネットワークの中には6,000万個のパラメータと65万個のニューロンが入っており、その全てが写真1枚が通るたびに少しずつ値を変えていました。部屋の中にはカードが放つ熱と冷却ファンの音だけがありました。5日から6日かかったその訓練が終わった時、機械が出した成績表は、コンピュータが写真を認識する方法を根底から変えるものでした。

その頃、写真認識は人間が半分代わってくれる仕事でした。1990年代後半から2010年代初頭にかけ、コンピュータビジョンの学界には共通認識がありました。生のピクセルを単に機械に投げ与えてはいけないということでした。まず人間が率先して、画像のどこに境界線があり、どのような輪郭が際立つのかを選び出す数学的フィルタを手で作り上げなければなりません。そのようにして人間が抽出した特徴を機械に渡して初めて、機械が分類を担当しました。この手作りの特徴には名前が付いていました。SIFT、HOGといったものでした。フィルタ1つがそのまま論文1編であり、よく作られたフィルタの発明者はその分野の専門家になったのです。フィルタが抽出した特徴を辞書のようにまとめ、その上にサポートベクターマシン(SVM)という分類器を載せるアセンブリプロセスが標準でした。若い研究者たちはこのプロセスのいずれかの段階をわずかに改善することで学位を取得しました。学界の大物たちはこの組み合わせだけが科学だと信じていました。ニューラルネットワークのように内部が見えない機械に特徴抽出まで丸ごと任せるとい主張は、数学的保証のない賭けだと思われていました。

この組み合わせには天井がありました。データがわずかに増えても、照明が変わっても、物体がわずかに歪んでも、パフォーマンスは崩壊しました。大規模な画像認識問題では、エラー率は26パーセント前後でそれ以上低下しませんでした。4枚に1枚の割合で誤ったということです。毎年大会が開かれても、小数点以下数桁を争うだけでした。人間が特徴をより精密に作成するにつれ、パフォーマンスはわずかずつ向上しましたが、その精密さがちょうど限界でした。人間の目があらかじめ定めたもの以上を機械が見ることができなかったためです。

その天井の下で状況を変えた人物は、プリンストン大学の若い教授フェイ・フェイ・リ(Fei-Fei Li)でした。リはボトルネックがアルゴリズムにあると考えませんでした。ボトルネックはデータでした。幼い子どもは、目を開いた最初の数年間、数億枚の光景を見ながら成長します。しかし、機械に与えられた訓練画像はせいぜい数万枚でした。リは2007年から大学院生ジア・デン(Jia Deng)と共に、インターネットから写真を集め始めました。問題はラベルでした。写真1枚1枚にこれが犬なのか猫なのかキノコなのかを書く人が必要でした。最初の計算では、学部生を雇ってラベルを付けるのに数十年かかることになっていました。リの研究チームはアマゾンのオンライン人力市場であるMechanical Turkに目を向けました。全世界167の国から49,000人がこの作業に参加し、候補となる数億枚の写真をフィルタリングしてラベルを付けました。そのようにして作られた写真辞書がImageNetです。2009年6月、マイアミ

で開催されたコンピュータビジョンとパターン認識に関する学会(CVPR)で、このデータベースは1つの論文として初めて公開されました。それもスピーチではありませんでした。展示会場の隅に紙1枚を貼るポスター発表でした。1,200万枚の写真と22,000個のカテゴリを含むものとしては、静かなデビューでした。

りはそこで止まらず、2010年から毎年大会を開きました。ImageNetから抽出した120万枚の写真と1,000個のカテゴリを使って、誰の機械が写真を最もよく認識するかを競う場でした。試験は意地悪でした。1,000個のカテゴリの中には、犬の品種だけで100種以上が含まれていました。シベリアンハスキーとアラスカンマラミュートを区別する問題の前では、人間も頻繁に誤ります。そのため、採点は機械に5回の機会を与える方式でした。機械が提出した上位5つの候補の中に正解が無い場合だけ不正解と数えるのです。その寛容な基準でも、機械たちは4枚中1枚を見落としていました。最初の2年間の優勝者は、すべて手で作った特徴と浅い分類器の組み合わせでした。2010年の優勝チームのエラー率は28パーセント台、2011年の優勝チームは25パーセント台でした。1年間に約2パーセントポイント低下するペースでした。そのペースなら、人間レベルの認識に達するまでさらに10年かかる見込みでした。

りがデータを積み重ねている間に、別の側では計算が安くなっていました。エヌビディアのグラフィックスカードはもともと、ゲーム画面を描くために数千個の小さな演算回路を同時に動かすように設計されていました。エヌビディアがCUDA(Compute Unified Device Architecture)というプログラミングツールを公開すると、そのカードでゲーム以外の任意の数学計算を実行できるようになりました。ニューラルネットワークが行うのはまさにそのような計算でした。同じ乗算と加算を数百万回繰り返すこと。ゲーム用カードはその反復に完全に適していました。20年間ニューラルネットワークを飢えさせていた障害は計算速度でしたが、その障害はゲーム市場が育んだ1つの部品によって静かに取り除かれていました。

ヤン・ルカン、ジェフリー・ヒントン(Geoffrey Hinton)、ヨシュア・ベンジオ(Yoshua Bengio)は、この2つの流れを長い間待っていた人たちでした。3人は、学界がニューラルネットワークを死んだ道と呼んでいた時代と一緒に耐え抜いた生存者たちでした。データも計算も不足し、ニューラルネットワークが本来の力を発揮できなかった時代です。今、データが来て、計算が来ました。残されたのは、誰かがその2つをニューラルネットワークに正しく適用することでした。

ところが地震が来る直前まで、ビジョン学界の門はルカンに閉ざされていました。2012年春のことです。ルカンの研究チームは、畳み込みニューラルネットワークを使って街並み風景写真を細分化して読み取るシーン分析論文をコンピュータビジョンとパターン認識に関する学会に提出しました。標準的なテストで記録を塗り替える結果が含まれていました。論文は落とされました。査読コメントの要旨は、パフォーマンスは良いかもしれないが、この方法が視覚について何を教えてくれるのかわからないということでした。ルカンは学会プログラムの委員長に公開抗議の手紙を送りました。畳み込みニューラルネットワークの論文がこれまでずっと却下されてきたため、この学会にこれ以上提出することは時間とエネルギーと誠意の無駄だと彼は書きました。その論文は数ヶ月後、機械学習学会(ICML)で受け入れられました。発明者が発明品を持ったまま門前で戻されていた年の秋に、その門は内側から崩壊します。

そのことを成し遂げた人物はヒントンの弟子クリジエフスキーでした。回想によると、トロントの研究室の中でImageNetへの挑戦を真っ先に推し進めた人物は、同僚の大学院生イリヤ・スツケバー(Ilya Sutskever)でした。十分に大きなニューラルネットワークに十分に大きなデータを与えれば勝つという確信を、彼は誰よりも早く持ち、プログラムの魔術師だったクリジエフスキーがその確信を機械の上に実現しました。クリジエフスキーはスツケバーと共に、ルカンが1989年にベル研究所で最初にした畳み込みニューラルネットワークの骨格を受け継ぎました。畳み込みニューラルネットワークは人間の

視覚皮質に倣い、画像の小さな領域をスキャンしながら同じフィルタを様々な場所に重ねて適用する構造です。ルカンが手書き数字と郵便番号を読み取れるようにと改良してきた構造の上に、2人の若者はいくつかの新しいものを追加しました。

1つはReLU(Rectified Linear Unit)でした。ニューラルネットワークは層を深く積み上げるほど、学習信号が下に降りていく過程で曖昧になる病がありました。クリジェフスキーはニューロンの反応方法を滑らかな曲線から角ばった直線に変えることで、信号が曖昧にならず下まで降りていくようにしました。もう1つはDropout(ドロップアウト)でした。ニューラルネットワークが訓練画像を丸暗記すると、初めて見た画像の前では逆に迷います。クリジェフスキーは学習するたびに、ニューロンの一部をランダムに無効にしました。機械がいくつかの手がかりにだけ依存しないように、意図的に妨害したのです。そして、この全ての計算をゲーム用カード2枚に分散させるために、彼は2枚のカードが互いに干渉せず、かみ合って動作するようにプログラムを1行ずつ手で書きました。ニューラルネットワークが1枚のカードのメモリに収まりきらなかったため、ネットワークを半分に分割して2枚のカードに配置し、途中からのみ互いに信号を送受信するアクロバティックな技でした。このニューラルネットワークには作成者の名前が付けられました。AlexNetでした。

2012年9月30日、ImageNet大会の結果が発表されました。人間が特徴を手作りした主流の方法は、その年も26パーセント前後で止まっていた。AlexNetのエラー率は15パーセントでした。上位5つの候補に正解が含まれているかどうかを測定する基準では、AlexNetは15.3パーセント、2位は26.2パーセントを記録しました。2位とのギャップは10パーセントポイント以上でした。小数点以下を争う大会では起こりえない差でした。1年に2パーセントポイントずつ低下していた曲線が、一度に5年分をスキップしたのです。

数週後、イタリアのフィレンツェで開催されたヨーロッパコンピュータビジョン学会 (ECCV) のImageNetワークショップでこの結果が公開の舞台上に上がりました。大会を創設したリーもそこにいました。新しく母になったリーはその年の学会出張を控えていたのですが、採点結果を見ると急いで飛行機チケットを買ってフィレンツェへ飛んだと、後年回想録で振り返りました。何が起こったのかを自分の目で見なければならなかったのです。会場は人で満杯でした。クリジェフスキーが壇上に上がり自らのニューラルネットワークがどのように動作するかを説明する間、客席に座ったコンピュータビジョン研究者たちは自分たちの分野が今まさに崩壊する音を聞きました。彼らの多くはニューラルネットワークについて真摯に学んだことのない人たちでした。生涯をかけて磨き上げた特徴フィルターが、ゲーム用カード2枚と大学院生2人の前で一夜にして時代遅れの産物になったのです。質疑応答は論争に変わりました。この機械はなぜうまくいくのか。内部で何が起きているのか。誰も納得のいく答えを出すことができませんでしたが、採点表の数字は答えを要求しませんでした。

この機械がなぜこんなにうまく機能するのは、その場の誰も説明することができませんでした。製造者たちも同様でした。6000万個のパラメータが内部でどのような役割分担を形成しているか覗き込む道具がありませんでした。数学的保証を要求していた分野が、保証なしにまず動作する機械の前に立つことになったのです。理論が実験を導くのではなく、実験が理論を大きく上回って走る時代がその日から始まったのです。より深く積み重ね、より多く入力し、より速く走らせればより良い結果が得られるという経験則が、その後10年間の研究を引っ張りました。ルクンが後年この経験主義の慣性と衝突することになるのも、根を辿れば今日の勝利があまりに大きかったからです。

ヒントンはその日の転向を後年このように振り返りました。結果を見たコンピュータビジョン分野の大家たちが態度を180度転換した、自分たちが完全に間違っていたので今は皆ニューラルネットワーク

に進むべきだと自分たちで認めたと、科学が自らを正す場面をそのような形で目の前で見たと、彼は2018年のチューリング賞講演で回想しました。20年以上の間、ニューラルネットワーク研究者を狂人のように扱い、論文を却下し研究費を打ち切っていた学界が、たった1枚の採点表の前で方針を転換するのに1年もかかりませんでした。翌年のImageNet競技会に出場した上位チームはほぼすべてが深い畳み込みニューラルネットワークを持ってきました。

その年の12月、クリジェフスキーとサツスキバーとヒントンの3人の名前が記された論文が神経情報処理システム学会に掲載されました。タイトルはシンプルでした。深い畳み込みニューラルネットワークを用いたImageNet分類。8ページのこの論文は、その後コンピュータサイエンスの歴史において最も引用される文献の一つになります。引用回数が数十万回に達する論文は、この分野全体でもほんの数篇しかありません。

波紋はコンピュータビジョンの世界を超えて広がりました。その年の秋、ヒントン研究室の大学院生たちは製薬会社メルク (Merck) がオンラインに開催した新薬候補物質予測競技会に参加しました。どの分子が医薬品になる可能性があるかをデータから予測する問題であり、参加者の多くは化学の知識を持つ人たちでした。ヒントンの弟子たちは分子化学をほとんど知らないまま、締め切り間際に深いニューラルネットワークだけを持って参加し、優勝しました。その年の11月、ニューヨークタイムズはジョン・マルコフ (John Markoff) 記者の記事でこのニュースとディープラーニングの急速な進展を大きく取り上げました。ニューラルネットワークという言葉が20年ぶりに再び新聞に勝者の名前として掲載されたのです。

産業が動く速度は学界よりも速かったです。その年の12月、アメリカのネバダ州のタホ湖で神経情報処理システム学会 (NIPS) が開催されました。ヒントンはクリジェフスキー、サツスキバーと共に従業員がわずか3人の会社を設立したばかりでした。その名前はDNNリサーチでした。製品もなく事業計画もない、3人の頭脳が資産のすべてである会社でした。学会が開催されるホテルの部屋でヒントンはこの会社をオークションにかけました。グーグル、マイクロソフト、DeepMind、そして中国のバイドウが値を吊り上げました。値段が上がる間、ヒントンはノートパソコンの前に立っていました。腰痛がひどく長く座ることができない体だったのです。値が2000万ドルを超えると2社だけが残し、オークションは4400万ドルで終了しました。ヒントンはグーグルを選びました。1本の論文と1つの大会での優勝が、半年も経たないうちに3人の値段をそのようなものにしました。ディープラーニング人材に向けた企業による狩猟がその部屋で始まりました。バイドウはオークションから空手で帰った後、自社のディープラーニング研究所を設立し、シリコンバレーの他の企業も競うようにニューラルネットワーク研究者の連絡先を探しました。大学の研究室が丸ごと企業に移転することが数年の間に当たり前になりました。そして翌年である2013年末、その狩猟の波がニューヨークのルクンにも及びます。マーク・ザッカーバーグの夕食のテーブルの話は次の章のものです。

この転覆の根にはルクンの1989年がありました。AlexNetが受け継いだ畳み込みニューラルネットワークは、ルクンがホルムデルのホワイトボードに微分公式を描きながら守ってきた構造でした。AlexNetの論文の参考文献にもルクンの名前が載っています。発明者の構造で弟子たちが世界をひっくり返したので、ルクンにとって2012年は20年遅れで到着した戦勝報告でした。しかし、この勝利の瞬間に一つの影がありました。世界が2012年のAlexNetをディープラーニングの出発点として記憶している一方で、それより1年前に同じ業績を先に成し遂げたチームが静かに消されていました。そのチームに対して、他でもないルクン自身が大会で負けたことがありました。

DanNetに負けた2011年

2011年の夏、アメリカカリフォルニア州のサンノゼで国際ニューラルネットワーク学会 (IJCNN) が開催されました。その年の大会にはドイツの交通標識を認識する競技が含まれていました。ドイツの実際の道路で撮られた標識写真約5万枚を43種類に分類する問題でした。速度制限、追い越し禁止、譲歩。種類は平凡に見えても、写真は容易ではありませんでした。逆光で褪せた標識、木の影に半ば隠れた標識、ステッカーが貼られペイントが剥がれた標識が混在していました。人間の目にも紛らわしいものが意図的に選ばれて含まれていました。採点表の最上位に上がった名前は学界の主流ではありませんでした。スイスのルガーノの小さな研究所から来たダン・チレサン (Dan Ciresan) でした。その下、2位の座にヤン・ルクンの名前がありました。

ルクンは量み込みニューラルネットワークを創造した人でした。彼が大会に持ってきたのも量み込みニューラルネットワークでした。弟子のピエール・セルマネ (Pierre Sermanet) と共に洗練させた最新モデルでした。そのような彼が、自分が発明したまさにその構造で武装した別のチームに負けました。数字は以下の通りです。大会を開いたドイツの研究チームは機械同士だけで競争させるのではなく、人間にも同じテストを解かせました。機械のパフォーマンスを測定する基準を作るためでした。そのように測定された人間の正答率は98.84パーセントでした。チレサンのニューラルネットワークは99.46パーセントを正解しました。人間が100枚に1枚ほど間違える問題を、機械は200枚に1枚ほど間違えました。コンピュータビジョン競技会において機械が人間を上回った最初のケースでした。ルクンとセルマネのニューラルネットワークは競技会の段階で98.97パーセントでその後を追いました。エラー率で比較するとチレサンチームの約3倍でした。優勝した機械には製造者の名前が付けられました。DanNetでした。

DanNetの背後にはユルゲン・シュミットフーバー (Jürgen Schmidhuber) がいました。スイス南部の人工知能研究所IDSIA (イドシア) を率いていた彼は、ヒントンの、ルクンとは異なる系統でニューラルネットワークを推し進めてきた人でした。ルガーノはディープラーニングの地図では辺境の中の辺境でした。トロントでもニューヨークでもモントリオールでもなく、スイスとイタリアの国境の湖畔の町。その辺境の研究室でルーマニア出身の工学者チレサンがゲーム用グラフィックスカードを手にしていました。

IDSIAが先に証明しておいたことが一つあります。2010年、チレサン研究チームは深いニューラルネットワークを最初から最後まで誤差逆伝播法だけで学習させ、手書き数字認識の世界記録を樹立しました。なぜこの部分が重要なのかは、当時の通念を理解する必要があります。2000年代後半まで、ニューラルネットワークの復興を主導した人物はヒントンであり、ヒントンの処方箋は教師なし事前学習でした。深いニューラルネットワークを直ぐに学習させると、信号が下層まで到達しないので、まず各層で下地絵を描かせてから本格的な学習に入るべきだということでした。ヒントンは2007年頃の公開講演で、事前学習なしに深いニューラルネットワークを誤差逆伝播法だけで最初から学習させようとする人は気が狂っていると言及したことさえあります。

IDSIAはその気が狂った仕事をやり遂げました。秘訣は速度でした。新しい理論があったわけではありませんでした。2000年代中盤までの市販のGPU高速化技術は、ニューラルネットワーク計算を3~4倍高速化するレベルでした。その程度では状況は変わりませんでした。チレサンはNVIDIAのゲーム用カードを徹底的に掘り下げ、学習速度を以前の数十倍に引き上げました。1ヶ月かかっていた訓練が1~2日に短縮される規模でした。計算が十分に高速化すると、事前学習というステップなしにもニューラルネットワークが底から直接学習されました。訓練画像を意図的にひねり傾けてデータを数倍に増やす古い技法も、計算が高速化したので自由に使うことができました。サンノゼに登場したDanNetは、こ

のように訓練した複数のニューラルネットワークを委員会のように束ね、それぞれの判断を集めて答えを出しました。ヒントンの処方間違いではありませんでした。その処方が必要だった理由そのものが消えてしまったのです。手で磨いた特徴も必要ありませんでした。DanNetは生のピクセルをカードにそのまま入力し、自らが特徴を見つけ出しました。後にAlexNetがImageNetで行ったことの本質が、規模は異なるだけで、ここにはすでにすべてありました。

超人間的という表現には脚注が必要です。機械が上回ったのは標識分類という狭いテスト1つであり、人間の視覚全体ではありませんでした。人間は標識を認識すると同時に、通りを渡る子どもを見て、雨に濡れた路面を測定し、次の角を予測します。その幅と比べると、DanNetの才能は針の先ほど狭いものでした。しかし、まさにそのような狭いテストにおいてさえ、それまで機械が人間を上回ったことはかつてありませんでした。狭い門が1つ開き、その後、他の門が次々と開かれることになります。

サンノゼの看板大会は始まりに過ぎませんでした。ダンネットは2011年と2012年にかけて国際大会を次々と制しました。交通標識に続いて、その秋には画数が複雑な中国語の手書きを認識する国際文書分析認識学会(ICDAR)大会で優勝しました。翌春には電子顕微鏡で撮影した脳神経細胞の写真から細胞膜の境界を見つけ出す生体画像大会で勝利しました。2012年9月には組織スライドの写真から分裂中の癌細胞を指摘する医療画像大会でも優勝しました。4度の優勝すべてがAlexNetのImageNet勝利より先行していました。深層学習が医療と産業の血管に流れ込むことができるという証拠が、世界が2012年を記憶する前にすでに積み重なっていたわけです。

二つのチームの運命は相反しました。AlexNetの3人はタホ湖のホテルの部屋で4400万ドルの身代金でGoogleに入り、Hintonはその後チューリング賞さらにはノーベル賞さえ受けます。Dannetを作ったCiresanの名前は専門家の外ではほとんど知られていません。同じ構造、同じゲーム用カード、先行する優勝。なのに一方は歴史書の最初の行になり、もう一方は脚注になりました。TorontoとGoogleという名前が実は大きかったのです。スイスの国境の小さな研究所が出した結果は、シリコンバレーの大企業が買収した結果ほど遠くには広がりませんでした。

歴史はこの4度の優勝をほぼ記憶していません。深層学習の始まりを語る時、人々は2012年のAlexNetを挙げますが、2011年のDannetはあまり挙げません。なぜそうなったのでしょうか。舞台の大きさが異なったという説明があります。交通標識数万枚とImageNetの120万枚、数十のカテゴリーと1,000のカテゴリーは世界に与える衝撃の程度が異なっていました。AlexNetが勝った舞台はちょうど学界全体が見守っていた舞台でした。しかしSchmidhuberはこの説明を受け入れません。彼の見方では問題は舞台にはなかったのです。記録が問題だったのです。

SchmidhuberはHinton、Bengio、LeCunの3人が深層学習の歴史を自分たち中心に書き直したと長年にわたり批判してきました。彼が挙げる証拠の一つは引用リストです。AlexNetの2012年の論文は、彼らより先にゲーム用カード上の畳み込みニューラルネットワークで大会を制したDannetについて、自分たちのモデルと多少似ているという表現で流してしまっただけでした。畳み込みニューラルネットワークの原形が1980年の日本の工学者福島邦彦のネオコグニトロンまで遡るということも、彼が長く繰り返してきた反論でした。彼がこだわる原則は貢献を正確に配分することでした。誰が最初に何をしたかを正確に記すことこそが科学の誠実さだという主張でした。彼らが沈黙していても事実の記録はやがて残ると彼は信じていました。ルガーノの彼は、このような優先権異議を数十年にわたり学会の質問時間と彼自身のウェブサイトですてきさげに続けてきた人でした。3人が2015年の科学ジャーナルNatureに掲載した深層学習のレビューを狙った彼の本格的な反論は、3人がチューリング賞を受けた後、第10章で詳しく展開されます。

この論争をLeCunに有利に決着させるつもりはありません。この本はLeCunの伝記ですが、2011年のDannetは消してはならない事実です。Schmidhuberの批判がすべて正しいかは、学界でもまだ争点です。HintonとLeCun側の見方では、実際に世界を動かしたのはある大会の最初の優勝ではなく、ImageNetという舞台での決定的勝利であり、考えの根っこに関しては1980年代から積み重ねてきた地層がDannetの足元にも敷かれているという反論が可能です。実際、Dannetが使った畳み込み構造そのものはLeCunの発明でした。優先権の争いはこの章で終わりません。3人が2018年度チューリング賞受賞者として発表された2019年、Schmidhuberが膨大な反論書を持ってまた現れる話は第10章で続きます。

どちらの側に味方しても、一つの事実は変わりません。畳み込みニューラルネットワークを作ったLeCunが、ゲーム用カード上にそのニューラルネットワークを最初に搭載した別のチームに大会で負けたということです。発明者が自分の発明品に負けたわけです。LeCunはこの敗北から何を確認したのでしょうか。自分の構造が間違っていなかったこと、足りなかったのは構造ではなく計算とデータだったということです。20年余りの間にニューラルネットワークが異端から文明の骨組みへと上がる転換の最初の炎は、LeCun自身が2位の席から見つめた2011年夏のサンノゼですでに燃えていたのです。

日常に浸透したCNN

あなたがちょうど1枚の写真をアップロードしたとしましょう。指を離して2秒も経たないうちに、その写真はすでに2つのニューラルネットワークを通ります。一つは写真の中に何があるかを読み取り、タグを付けて有害なシーンをフィルタリングします。もう一つは写真中の人物の顔を見つけます。あなたは何も見ません。画面には今アップロードした写真しか映っていません。その裏では、地球でも有数の大規模な人工知能演算がちょうど静かに回りました。

顔を見つける2番目のニューラルネットワークにも経歴があります。2014年、Facebookの人工知能研究所はDeepFaceという顔認識研究を発表しました。2枚の写真が同じ人物かどうかを判定するテストで、このニューラルネットワークは97.35パーセントを正答しました。同じテストでの人間の成績は97.53パーセントとされていました。機械と人間の差が0.2パーセントポイント余りに狭まったのです。顔を認識することは長い間、人間だけの特技と考えられていましたが、その特技が数字で追いつかれるのにAlexNet以降2年かかりました。

この数字を公表した人物はLeCunでした。2016年のパリのCollège de France講演で、彼はFacebookサーバーの1日をスライドに映しました。その頃Facebookユーザーたちは1日平均8億枚の写真をアップロードしていました。InstagramとMessengerとWhatsAppを合わせると1日20億枚でした。その20億枚がそれぞれアップロードされた2秒以内に2つの畳み込みニューラルネットワークを通過しました。最初のニューラルネットワークは写真内の物体と背景と風景を読み取り、タグを付け、スパムと露出物と暴力シーンをフィルタリングしました。2番目のニューラルネットワークは顔を見つけました。LeCunが1989年にHolmdelで郵便番号を読ませるために調整したその構造が、今や1日20億回世界の写真を読んでいたのです。パリ郊外で育ちグランゼコールに進まなかった少年がどのようにしてその古い講壇に立つことになったのかは、後でその講演を再度訪ねてみます。

畳み込みニューラルネットワークは音声にも入り込みました。スマートフォンに話しかけると、音声認識することは通常、異なる種類のニューラルネットワークが担当すると考えられています。しかし、その処理の最初の関門には畳み込みニューラルネットワークが座っています。音声は最初に周波数図に変換されます。時間に応じてどの高さの音がどの程度強いかを2次元の図に展開したものです。音が図になる瞬間、畳み込みニューラルネットワークが登場します。人間の視覚皮質が物体の端を見つけ

出すように、このニューラルネットワークは音の質感を見つけ出します。LeCunは音声認識の最初の数層が例外なく畳み込みニューラルネットワークだと言ったことがあります。道端の草花や虫を撮ると名前を教えてくれるアプリも、森から聞こえる野鳥の鳴き声で鳥の種類を当てるアプリも、その心臓には同じフィルターが入っています。写真アルバムの検索も同じです。2015年頃からスマートフォンの写真アルバムに「犬」と入力すると犬の写真が引っかかり始めました。誰が写真ごとにラベルを付けてくれたわけではありません。ニューラルネットワークが写真アルバムの写真をあらかじめすべて読んでいたのです。手書き郵便番号1桁を読むことから始まった機械が、四半世紀で個人の写真アルバム全体を読む司書になりました。

日常の便利さを超えて、人命に関わる場所でもこのニューラルネットワークが守っていました。今日、欧州連合で販売されるすべての新車には自動緊急制動装置が搭載されることが義務付けられています。前に突然何かが現れると、機械が自動的にブレーキを踏む装置です。LeCunは、この装置の目の裏には彼が発明した畳み込みニューラルネットワークが漏れなく入っていると述べました。つまり、欧州で販売される車なら例外がないということでした。この自動車の目が生まれた場所は、意外にも古い実験でした。2003年、LeCunの研究チームはアメリカ国防総省の研究課題として小さなおもちゃのトラックにカメラを2台取り付けました。人間が20分ほど運転した走行データだけを見せたら、トラックが自分で茂みと障害物を避けて走りました。その時に作られた屋外知覚技術は、イスラエルの車両知覚企業Mobileyeとエヌビディアの車両用チップを通じて乗用車の安全装置に移りました。2015年にTeslaがModel Sに最初に搭載した運転支援装置の目もMobileyeの畳み込みニューラルネットワークベースのモジュールでした。おもちゃのトラックのカメラと高速道路を走る乗用車のカメラの間に、同じ数学が流れていたわけです。

病院でもこのニューラルネットワークが時間を稼いでいました。LeCunが率いるFacebook人工知能研究所はニューヨーク大学医学部と共に、磁気共鳴画像(MRI)撮影を高速化する研究をしました。従来のMRIは患者が狭く騒々しい筒の中に40分ほどじっと横たわっていなければならない検査でした。研究チームはすべての信号を集める代わりに、まばらに集めた信号から畳み込みニューラルネットワークが残りを補填して高画質の画像を再構成させました。診断に使える画質を保ちながら撮影時間を4分の1の10分に短縮しました。40分耐えていた患者に30分を返したわけです。閉所恐怖症のある患者、長く横になるのが難しい高齢者、じっとしていられない子どもにとって、その30分は検査を受けられるかどうかの分かれ目でもありました。研究チームはこの技術のデータとコードを公開し、他の病院と機器会社が使用できるようにしました。Schmidhuberの研究室が2012年に勝った癌細胞検出大会を覚えていますか。そのシステムの技術が今、実際の病院の読影室で回転しています。勝者が誰であろうと、病院に入ったのは同じ畳み込みでした。

ヒントンは、ニューラルネットワークがこのように世界の至る所に広がる様子について、一つの比喻を挙げました。2018年のチューリング賞での講演でのことです。従来のソフトウェア内にニューラルネットワークの部品を一つ組み込むことは、体に壊疽(えそ)が広がるのと同じだと彼は言いました。一度始まると止まらず、隣の正常な部品まで次々と食いつぶし、最後にはシステム全体をニューラルネットワークに置き換えてしまうという意味でした。不吉な比喻でしたが、実際その通りになりました。音声認識パイプラインが最初に飲み込まれ、画像検索が飲み込まれ、翻訳が飲み込まれました。部品を一つニューラルネットワークに変えたチームは性能が上がるのを目にし、次には隣の部品も変えていきました。

ImageNet大会の採点表も、その広がりや年々記録していました。2013年の優勝者はマシュー・ゼイラー (Matthew Zeiler) でした。ニューヨーク大学でルカンの同僚であるロブ・ファーガス (Rob Fergus) の下で学んだ大学院生で、AlexNetを手直したニューラルネットワークによってエラー率を11パーセント台まで下げました。ゼイラーはそのままClarifaiという会社を設立しました。大学院生の大会優勝が直ちに起業につながる時代が開かれたのです。同じ年の大会で、ルカン自身も久しぶりに優勝者のリストに名を連ねました。2011年にサンノゼで共に2位にとどまった教え子のセルマネと作ったOverFeatというニューラルネットワークが、写真内の物体の位置まで特定する種目で1位を占めたのです。DanNetに敗れてから2年ぶりのことでした。ImageNetの優勝は、その後も数年間、新興企業や研究所の名刺代わりとなりました。

この広がりを支えたのは、ツールと制度でした。ルカンの研究室の系譜から生まれたTorchと、その系譜を継いだPyTorchは、世界中の研究者の手に渡り、以前なら数百行だったニューラルネットワークのコードを数行に減らしました。2015年には、マイクロソフト北京研究所の何愷明 (Kaiming He) がResNetという構造を発表しました。層と層の間に近道を作り、信号が滞ることなく飛び越えられるようにした設計でした。その近道のおかげで、ニューラルネットワークは100層を超えて重ねても崩れず、その年のImageNet大会のエラー率は3パーセント台まで下がりました。人が同じテストを受けると5パーセントほど間違えると言われているため、機械がその基準を通り過ぎた年でした。26パーセントから15パーセントへ、さらに3パーセントへ。2012年の地震が作り出した下り坂でした。何愷明は後日、ルカンの研究所に席を移し、この流れを共に推し進めました。

制度の面でルカンが置いた礎もあります。2013年、彼はベンジオと共に表現学習国際学会 (ICLR) という新しい学会を作りました。1年前、自分の論文を落とした閉鎖的な審査に抗議の手紙を書いていた人が、今度は学会そのものを新設したのです。新しい学会は審査方式から異なっていました。提出された論文と審査評を、最初から誰もがみられるように公開しました。審査者が匿名性に隠れ、ある分野を丸ごと却下するような事態を構造的に防ぐという設計でした。ルカンはそれ以前から、論文出版制度を抜本的に改めようという提案文を自ら書いて回していた人物でした。辺境のワークショップとして始まったこの学会は、10年も経たないうちに人工知能分野で指折りの学会へと成長しました。追い出された人々が築いた村が、首都になったのです。

ルカンが1980年代にホルムデルのホワイトボードに描いた視覚の文法は、20年余りの蔑視に耐え、人々の日常を音もなく支える骨組みとなりました。郵便局の郵便番号から銀行の小切手へ、Facebookの写真から自動車のブレーキへ、病院の撮影室から写真帖の検索窓へ。一人の人間が生涯かけて磨き上げた一つの構造が通り過ぎた跡です。彼は正しかったのです。知能とは人が規則を書き込むのではなく、機械がデータから学ぶものだというその古い主張が、1日20億枚の写真の上で証明されました。

勝利の影も早くから現れました。2013年末、研究者たちは、ニューラルネットワークが完璧に認識していた写真に、人の目には見えない微細なノイズを混ぜると、判定が突拍子もなく覆るという事実を発見しました。パンダの写真に計算されたノイズを少し載せると、人の目には依然としてパンダであるにもかかわらず、機械はテナガザルだと答えるといった具合でした。敵対的サンプルと呼ばれるようになったこの現象は、機械が写真を認識する方式が、人の「見る」こととは異なる何かであるという事実を露呈させました。見事に言い当てることと、理解することは別の問題だったのです。

ところが、この勝利の真ただ中で、ルカンは安住しませんでした。世界が自身の作った部品の上で回り始めた頃、彼は再び学界の主流に背を向けます。畳み込みニューラルネットワークは見る機械であって、理解する機械ではありませんでした。写真の中の猫を認識することと、猫がテーブルから飛び降り

たらどうなるかを知ることは別のことです。パンダをテナガザルと勘違いする機械に、世界の物理を任せることはできませんでした。数十兆ウォンの資金がChatGPT式の巨大言語モデル (LLM) に注ぎ込まれていた2020年代に、ルカンはその方向が行き止まりであると呼びました。文字だけを読んでいては物理世界を理解できないということでした。手書きの数字をピクセルで読み取らせていたフランスの青年は、今やピクセルを描くことに計算を浪費せず、世界の構造自体を予測する別の機械について語り始めます。2011年の敗北で計算の力を確認した人が、今度は計算の力だけでは辿り着けない場所を指し示していました。彼がなぜその勝利の頂点で再び別の道を選んだのかについては、本書の後続の章へと続きます。

弁護士

kimkj.com

第9章 ザッカーバーグの食卓

ヤン・ルクン、人工知能の父

カリフォルニアの自宅での夜

2013年の秋遅く、ニューヨーク大学クーラント研究所のある部屋で電話が鳴りました。ヤン・ルクンは受話器の向こうの声を聞き分けました。マーク・ザッカーバーグでした。29歳のフェイスブック最高経営責任者が、52歳の神経網研究者に直接電話をかけてきたのです。用件は短いものでした。フェイスブック内に人工知能の基礎研究所を新たに設立しようとしており、その仕事を引き受けてほしいというものでした。ルクンはすぐに返事をしませんでした。一度会って話をしようという程度で通話を終えました。

その頃シリコンバレーでは静かな戦いが繰り広げられていました。前の年の夏、トロントのとある研究室から現れた神経網が、画像認識の競技大会の流れを変えてしまった後のことでした。AlexNetと呼ばれたその神経網は、写真に映る物体を見分ける競技大会において、それまでのどんな方法よりも大きな差をつけて先行していました。この30年間誰も顧みなかった技術が、一夜にして身価が急騰した技術になったのです。神経網を扱える人の数は、世界中でも数えるほどでしたし、大企業たちはその少ない人材を引き抜こうと競い合って財布を開きました。グーグルが先行しました。AlexNetを作ったジェフリー・ヒントンとその弟子2人を、会社ごと買収し、ロンドンのDeepMindを傘下に収める準備を進めていました。その買収は翌年初めに成立します。数億ドルの取引でした。伝え聞くとところによれば、フェイスブックもその頃DeepMindに手を伸ばしたものの、グーグルに負けたとのことでした。神経網を扱える人物一人、会社一社をめぐって、世界でも有数の大企業たちが価値を競い合う時代が幕を開けたのです。その争いの中でフェイスブックは後れを取っていました。

ザッカーバーグが人工知能を望んだのには理由がありました。フェイスブックはその前の年に企業を公開していました。10億人以上の人々が毎日何かを投稿し、クリックし、スクロールしていました。その流れを読み解く仕事、どの記事と写真を誰に最初に見せるかを決める仕事が、会社全体の経営を左右していました。その仕事の最中心に人工知能がありました。会社内にはすでに応用組織が一つありましたが、ザッカーバーグが求めているのはそれを超えるものでした。すぐに使える機能を作るチームが欲しいわけではありませんでした。知能という問題そのものに長く取り組む研究所を望んでいたのです。ルクンがザッカーバーグから確認したかったのも、まさそこでした。この若い経営者が人工知能を広告単価を調整する計算機程度にしか見ていないのか、それとも、それを超える見方をしているのかということでした。

数日後、ルクンはカリフォルニアへ飛びました。ザッカーバーグは彼を会社の会議室に呼びませんでした。自分の家に招いたのです。事務所の形式張ったやり方の代わりに、晩食のテーブルを選んだのです。二人は長く話し合いました。ザッカーバーグが描いていた絵は広告を超えたものでした。彼は人と人を結びつける仕事そのものを技術の目的としており、人工知能をその接続を支える基盤として見ていました。言葉の壁を越えて異なる国の人々が互いの言葉を理解できるようにする仕事、目の見えない人に写真の風景を読み聞かせる仕事といった例を彼は挙げました。ルクンはその夜の対話の中で、自分が求めている真摯さを読み取りました。この人は人工知能を道具としてのみ見ていないこと、長期的な視点を持っていることでした。

もう一つのことがルクンの心をさらに惹きました。ザッカーバーグは何も描かれていない白紙を差し出していたのです。すでに動いている組織を引き継いで改善する仕事ではありませんでした。研究所を最初から自分の手で設計する仕事だったのです。どのような人を採用するか、何を研究するか、どのようなルールで運営するかを、白紙の上で自ら決めることができるということです。生涯他人が歩まない道だけを歩んできた人にとって、それは稀な機会でした。

それでもルクンは一つのことを心の片隅に折りたたんでいました。約束は最初の日は簡単にできるということでした。会社に金銭的な余裕があり、将来が明るい時、研究者に自由を与え、成果を発表させるという言葉は容易に出てくるものです。しかし、その約束が試験台に乗るのは、数年後、市場が急ぎ立て、競争相手が先行し、目の前の成果が急務となる日です。その日にもこの会社が今日の言葉を守るのか。ザッカーバーグの真意を疑うのではなく、企業というものは本来そのように流れやすいということ、彼はベル研究所で学んでいたからこそ持つ疑問だったのです。その疑問の答えはその時は分かりませんでした。ただし今は白紙が目の前にあり、その上に何を描くかは彼の手にかかっていました。

この電話を受けるまでに歩んできた道も、彼の躊躇の中に影響していました。彼はベル研究所を離れた後、その後継研究所で一時期働き、その後日本系企業のアメリカ研究所を経ました。そして2003年、ニューヨーク大学の教授になったのです。企業研究所を離れ、大学の講壇に立った人だったのです。大学には会社が与えられないものがありました。誰の許可も求めることなく、10年かかる問題に取り組む自由、そしてその結果を思いのままに世間に発表する自由でした。そのような彼をもう一度企業へ呼び戻す電話だったため、素直に応じることはできませんでした。

そうは言っても、心が一気に傾いたわけではありませんでした。ルクンは企業研究所がどのように崩壊するかを身をもって経験した人でした。彼は人生の最盛期をベル研究所で過ごしました。手書き数字を読む神経網を作り、アメリカの銀行がそれで小切手処理を行うようにした場所がそこでした。そしてそこが段々と解体されていくのを見守りました。会社が分割される中で、彼が同僚たちと共に磨き上げた畳み込み神経網の主要特許は、あるメーカーの金庫に入ってしまったのです。世に出て使われるべき技術が書類のファイルに閉じ込められて眠ってしまったのです。その時のことが、彼の心に一つの教訓として残っていました。どんなに優れた研究も外に出ることができなければ死ぬということでした。

ベル研究所の思い出には別の側面もありました。そこはトランジスタが誕生し、情報理論が誕生し、通信衛星のノイズを追い求める過程で宇宙の最初の光が残した背景放射が発見された場所でした。一つの企業が当面の金儲けとは無関係な基礎研究に数十年投資するとどのような成果が出るのか、ルクンはその伝説的な廊下で自分の目で見ていました。それだから彼がザッカーバーグに問いたかったことは、結局一つでした。フェイスブックがベル研究所が享受していたその長い息遣いを人工知能に与えることができるか。果実がいつ実るか分からない木を何年でも待つ覚悟があるか。彼が若い日に身を置いていた楽園の名前がベル研究所だったとすれば、その楽園をもう一度築こうという提案が目の前に置かれたわけでした。

そこでルクンはザッカーバーグの提案を拒否することも、受け入れることもしませんでした。代わりに条件を付けました。夕食のテーブルの上で、彼は自分が参加するために必ず守られなければならない幾つかのことを一つ一つ取り出しました。それは給与や職位に関する話ではありませんでした。一人の科学者が企業の柵の中に入りながらも科学者であり続けることができるかに関する話でした。

三つの条件

ルクンが提示した条件は三つでした。そしてそれら三つは、すべて彼が失いたくないものたちのリストでもありました。

最初の条件は場所でした。彼はニューヨークを離れないと言いました。カリフォルニアに移転しないということです。家族がニューヨークにいて、彼の人生がその都市に根ざしていました。大企業に入る人が本社のある場所に移転するのは当たり前の順序でしたが、彼はその当然さを拒否しました。研究所をニューヨークに置くか、でなければ自分がニューヨークから指導する方式でなければならないとしたのです。実際にFAIRの最初のオフィスの一つがニューヨークに開設され、ルクンがそこを自分の拠点にしました。

二番目の条件は大学でした。ルクンはニューヨーク大学の教授職を辞めないと言いました。彼はシルバー教授であり、少し前にはその大学でデータ科学センターを設立し、それを主導していました。フェイスブックの研究所を引き受けても、大学の講壇と指導学生を手放さないということ、一方の足を学界に留めておくということでした。彼が設立したデータ科学センターは、ちょうど門戸を開いたばかりでした。複数の分野の学者たちが一つの場所に集まり、データを扱う新しい学問を切り開く場所であり、ルクンはその産婆役でした。今ようやく育て始めたこの土台をフェイスブックのために放棄する考えはありませんでした。大学の講壇は、彼が若い世代と出会う場でもありました。その場を守る限り、彼は会社の人である前に教授であり続けることができました。人々はこのやり方を二つの帽子をかぶると呼びました。産業と学界のどちらか一方に完全に属することなく、二つの世界を行き来する立場でした。その二つの世界を結ぶ橋になるという意味であって、会社の人になるという意味ではありませんでした。このやり方には実利もありました。大学に残った教授は学生を教え、その学生の中で優秀な者が再び研究所へ流れ込みます。学界と会社の間で人が行き来する流路が生まれるのです。

三番目の条件が最も遠くに進むものでした。ルクンは研究所から出るすべてのものを公開すると言いました。論文も、数式も、コードもすべて余さず外に出すということです。彼の表現を借りれば、私たちのすべての研究は公開され発表されているという原則でした。企業研究所においてこれは普通の要求ではありませんでした。会社が金をかけて得た成果を競争相手にまでそのまま譲り渡すという意味だったからです。普通の経営者ならここで立ち止まったでしょう。守るべきものをなぜ出すのかと問うたはずです。

ルクンにとって、この第三の条件は他の二つを合わせたものより重かったです。その根は、ベル研究所での経験にありました。彼が開発した小切手読み取りニューラルネットワークが、米国の銀行処理の現場で実際に稼働していた時期に、その成功の瞬間、会社はその研究を生んだ組織を解体して分散させてしまいました。商用化の日と解体の日が重なってしまったわけです。苦労して世に送り出した技術が、実はそれを作った人々の手を離れて書類の中に消えていくのを、彼はその時に見守ることになりました。金庫に鍵をかけられた特許、誰も使うことができなくなった発明。その記憶が彼に一つの原則を刻み付けました。外に出て検証されない研究は研究ではないということでした。論文として世に公開し、他の人々がそれを詳しく検討し、問い質し、反論することができてこそ、初めて科学になるという考え方です。閉ざされた部屋の中で一人が正しいと信じることは科学ではあり得ませんでした。彼はそれを自己欺瞞に近いと見なしました。開放は彼にとって徳行である以前に、科学の条件でした。後年、彼はFAIRの規則をこのようにまとめました。ここの科学者にとって、論文は出してもいいものではありませんでした。必ず出さなければならないものであり、外の学者たちと全く同じ基準で評価されなければなりません。公開は義務でした。推奨事項ではありませんでした。

この原則の背後には、ルクンが経験した別の冬がありました。1990年代後半、ニューラルネットワークは学界で再び厄介者扱いされるようになりました。論文を投稿しようとしても査読で却下され、題目に「ニューラルネットワーク」という言葉が入ると、読まれることなく無視されることが多かったです。その時ルクンを支えたのは開放でした。閉鎖ではありませんでした。自分の研究を世に出し続け、誰もがそれを手に取って検証し、さらに発展させることができるようにすることです。学界が背を向けた技術が再び生き返ったのも、結局のところコードとデータが公開されていたため、若い研究者たちがその上で実験を積み重ねることができたからでした。彼は開放が窮地に追い込まれたアイデアを救う道であることを身をもって学んだ人でした。

ザッカーバーグは三つのことすべてを受け入れました。その場で、そうすると言いました。後々から見れば、この承諾こそが、その後起こることすべての種でした。ザッカーバーグが条件を受け入れたのには、計算もありました。彼が連れてきたい人たち、世界で最高レベルのニューラルネットワーク研究者たちは、論文を発表できないところには来ませんでした。彼らにとって研究成果を世に知らせることは、給料と同じくらい重要でした。学界の認認可、同僚の評価、自分の名前が付いた論文。それを放棄させようとする会社には誰も来ませんでした。扉を開けておくことが最高の才能を呼ぶ道であることをザッカーバーグは気づきました。開放は理想であるとともに戦略でもありました。この時点では、二つのことが同じ方向を指していました。

この三つの条件は、一つにまとめて見れば、新しい契約の一つでした。科学と企業の間の契約です。それまで企業研究所のルールはおおよそ反対でした。会社に入れば、成果は会社のものになり、論文は会社が許可する限りでのみ発表され、職務は会社が定めた場所に移動するのが普通でした。ルクンはそのルールを逆転させました。成果は世の中のものとして開放され、論文は義務として発表され、職務は学者が定めるということでした。この逆転した契約が通用することをFAIRが示すと、他の企業もこれに類似した道を歩み始めました。業界が人工知能研究を扱う方式そのものが、この時期に一度変わり、その変化の扉を開いたのがザッカーバーグの食卓で交わされた三つの言葉でした。

この構想を会社の中で支えた人が、もう一人いました。当時フェイスブックの最高技術責任者マイク・シュレーファーです。彼はフェイスブックに来る前、Mozilla財団でFirefoxウェブブラウザを率いていた人でした。Firefoxはソースコードを誰にでも公開したオープンソフトウェアの象徴でした。シュレーファーの身には、その開放の文化が染み込んでいました。ルクンがコードを公開すると言った時、彼を奇異に思わず、むしろ支持してくれた人がシュレーファーでした。会社の技術組織の深くに、すでに共有することを厭わない気質が定着していたわけです。知識を隠して特許として束ねるのではなく、外に出して一緒に育てるという考え方が、ルクン一人の頑固さではなく、すでに会社の骨組みの中に染み込んでいたのです。

2013年12月、フェイスブックは新しい研究所の開設を発表しました。名前はFacebook Artificial Intelligence Research、略してFAIRです。ヤン・ルクンがその初代所長でした。ニューヨークに拠点を置き、大学に片足を置きながら、出てくるすべてのものを外に発表する研究所。大きな企業を中心に立てられた、その時代のどこにも存在しなかったような研究所でした。

FAIRの誕生

FAIRの最初の営みは、学会会場で始まりました。2013年12月、その年の神経情報処理システム学会が開かれていた場所にザッカーバーグが現れました。世界で最大級の企業の一つを率いる経営者が、人工知能研究者たちが集まったブースの間を直接歩き回りました。彼は若い学者たちに新しい研究所を紹

介し、一緒に働こうと招きました。最高経営責任者が学会会場まで出向いて人を探すということは、それ自体が一つのシグナルでした。この会社が人工知能をどれほど真摯に捉えているかを示すシグナルでした。その場に集まった人々は、ほんの少し前まで学界の端っこでニューラルネットワークに取り組んでいた人たちでした。世が背を向けていた技術が、今や会社の主人を学会会場まで呼び出したのです。

学会会場の空気そのものが一つの事件でした。ほんの数年前まで、ニューラルネットワークを研究していると言えば、就職が難しいと言われていた分野でした。そのような分野の学会に、世界でも指折りの富豪が自ら訪れて名刺を配っていたのです。ある者たちは学問の純粋性が資本に売られていくと懸念し、ある者たちはついに冬が終わったと歓迎しました。この二つの反応は、その後も長くFAIRに付きまとうことになります。企業の資金でなされた基礎科学はどこまで純粋でいられるのか。その問いは、この研究所の誕生と共に生まれました。

FAIRが掲げた約束は、他の企業研究所とは異なっていました。ここでは論文を書いてもよいという程度にとどまりませんでした。論文を書きなさいということでした。成果を隠さず学会に出し、外の他の科学者たちと全く同じ基準で評価されなさいということでした。目の前の製品に縛られず、十年、二十年先を目指した遠い研究をしなさいという意味で、人々はこのような研究を青空研究と呼びました。当面はお金にならなくてもよいから根本を掘り下げなさいという場所でした。大学の自由を企業の資源と共に享受する場所、それがFAIRが自らを紹介する方式でした。

もちろん、二つの世界が出会う場所には軋みがありました。企業は定められた日付に製品を市場に出さなければならない場所です。次四半期、次のリリース、次の指標。研究所は、いつ成果を上げるかわからない種を植える場所です。二つの時間は異なっていました。この異なる時間を一つの屋根の下に置くことがFAIRの課題でした。ルクンは研究者たちを製品の納期から引き離そうと努力しました。彼らが論文の厳密さだけで評価される場所を、企業を中心に守り通そうとしたのです。

この約束に惹かれて、人々が集まりました。後にルクンと共にFAIRを率いることになるジョエル・ピノーもその一人でした。カナダのマギル大学の教授だった彼は、科学者たちが会社の指示ではなく、互いへの尊敬の中で働くことができ、それでいながら世を変えることができるというFAIRの約束に心を動かされたと、後に語ったのです。研究者に自由を与え、その代わりに成果をすべて公開させるというこの交換が、世界中の優れた人々を惹きつけました。

ピノーは公開をさらに一歩進めました。論文を発表するだけでは不十分だというのが彼の考えでした。他の人がその論文の結果を自分の手で全く同じように再現できてこそ、初めて科学であるということです。彼は人工知能学会に再現可能性という基準をもたらしました。論文を発表する際に、コード、データ、実験の条件を一緒に明かし、他の人たちが同じように試すことができるようにしようという運動でした。もっともらしい結果を発表しながらその内部を隠す慣行にブレーキをかけたのです。これはルクンが食卓で掲げた第三の条件と同じ根から生まれた考え方でした。外に開放され検証されないものは科学ではないという原則です。FAIRは論文を学会誌に掲載する前に、まず公開レポジトリにアップロードし、誰もが読めるようにする方式を採用し、こうして知識が流通する速度を引き上げました。

研究所は一つの都市に留まりませんでした。ニューヨークとメンロパークから始まり、パリへ、そしてモンテリオール、シアトル、ピッツバーグ、テルアビブ、ロンドンへと広がっていきました。世界中の大学都市ごとにFAIRの拠点が一つずつ開設されました。人材のある場所に研究所が出向いたのです。人を一か所に集める代わりに、人が暮らす場所に研究所を根付かせるやり方でした。

開設から五年が経つ頃、FAIRは数百人の科学者が働く大きな研究所になっていました。毎年数百本の論文がここから出て学会に流れ込み、その論文に付属するコードが世界の実験室に広がっていきました。ニューラルネットワークに取り組むことで就職を心配していた時代があったのかと思うほど、この分野は学問の中心へと上昇していました。ルクンが食卓で掲げた条件たちが、その中心を支える骨組みになっていたのです。企業の中に築かれた大学、その大学が世界全体に向けて開いた窓。その窓を通じて行き交う知識の量が、この時期FAIRが残した最大の足跡でした。この時が開放の春でした。春が常にそうであるように、その後何があるのかは、たけなわの時には見えにくいものです。

研究所は急速に展開していきました。ニューヨークから始まり、カリフォルニアのメンロパークへ、そして2015年6月にはパリへと進出しました。パリはフェイスブックがアメリカ外に設立した初のAI研究所でした。なぜロンドンではなくパリだったのか。ザッカーバーグの会社は二つの都市を秤にかけて比較し、ヨーロッパ大陸側の人材の井戸がより深く、より開かれていると見なしました。それは理由のあることでした。フランスは数学と工学教育が厚い国であり、ルクン自身がその教育を受けて育った人でした。自分の母国に研究所を設立することは、彼にとって格別な意味がありました。ニューラルネットワークを誰も信じなかった時代にフランスを去ってアメリカへ渡った人が、今やそのフランスに世界最高水準の研究所を携えて戻ってきたのです。

FAIRパリは研究員5名で出発しました。その年が終わる前に人員を2倍に増やす計画であり、数年もしないうちにその数は80名を超えました。フランスにおいてFAIRは会社の支社であると同時に学界の隣人になることを目指しました。パリにはこの国の計算機科学研究の中心である国立情報自動化研究所、略してINRIAがありました。FAIRパリはINRIAと協力し、博士課程学生と博士研究員が両機関を行き来しながら共に研究する道を開きました。フランスには企業と大学が博士課程学生を共に育成するCIFREという制度がありました。3年間の契約で学生が会社と公開研究機関の両者に足をかける仕組みでした。FAIRパリはこの制度を採用し、会社が給与と計算機器を提供する一方で、学生が自らの研究を行い論文を自由に発表できるようにしました。学生は時間の大部分を自分の大学の研究に費やし、残りの一部を会社と共にしました。企業の資本が一つの国の科学生態系を支える養分として流れ込んでいったのです。

この方式は、ルクンがザッカーバーグの食卓で示した三つの条件を国家規模に広げたものと言えます。大学に足をかけること、転職を強要しないこと、成果を開放することという原則です。FAIRパリではその原則は一人の契約条件ではなく、研究所全体の運営方式となりました。会社に身を置いた研究者が大学の自由をそのまま享受するという考え方が、大西洋を渡ってパリの実験室で再び花開いたのです。

FAIRパリがフランス学界に残したものは論文だけではありませんでした。メタは複数年にわたってフランスの公開研究機関と大学に研究費を提供し、若い研究者たちに奨学金を給付し、計算機器を寄贈しました。人工知能研究には高価な計算機械が不可欠です。優れたアイデアを持つ大学院生でも、それを実行する機械がなければそのアイデアを検証することはできません。FAIRパリが大学の実験室にサーバーを提供したことは、その瓶首を打ち開けたということでした。企業の資金が会社の金庫を満たす代わりに、一つの国の若い研究者たちに流れ込んでいったのです。こうした支援について、会社が将来の人材を先取りする投資だと見る人たちもありました。実際、そのような側面はありました。しかし、その資金で論文を発表し学位を取得した学生たちが必ずしもメタに就職しなければならないわけではありませんでした。彼らは自分が望む場所へ進み、こうして育成された才能はフランス学界全体の財産になったのです。

FAIRが公開した成果のリストは長かったです。研究所が開設された翌年に登場した顔認識モデルは、写真に映る二人の人物が同一人物かどうかを人間並みの精度で判別しました。文字を扱うツールもありました。単語の意味を数値の形に変換する小さなライブラリを作り、300近い言語の版を無償で公開しました。このツールはサイズを縮小され、ポケットサイズの携帯電話やスマートウォッチほどの小型デバイスでも動作しました。高価な設備のない地域でも、機械が人間の言語を処理できるようにするという意図でした。翻訳に関する業績もありました。FAIRはそれまでの手法と異なるアーキテクチャで文を変換する方法を提案して処理速度を数倍高速化し、2022年には200の言語に対応する翻訳モデルを世に公開しました。そのモデルの名前は、どの言語も後に取り残さないという意味が込められていました。ユーザーが少なく大企業が関心を示さなかった言語たち、世界の周辺へ追いやられていた言語たちを機械が等しく理解できるようにしようとする試みでした。

こうしたツールは論文とコードとして学界に公開されるだけでなく、静かに会社の製品の中にも浸透していきました。フェイスブックにアップロードされた写真から人物と物体を認識すること、複数の言語で書かれたテキストをユーザーの言語に翻訳すること、視覚障害のある人に写真の内容を音声で説明すること。こうした機能の背後にはFAIRとその隣接部門の研究がありました。ザッカーバーグが食卓で語った構想、言語の壁を乗り越え、見えない人に世界を読み聞かせるというあの大志が、実際の製品となって10億人の手に届いたのです。基礎研究と日々の製品はこうして当面の間、同じ屋根の下で調和していました。その均衡がいつまで続くかは、その時点ではまだ誰も問うことはありませんでした。

FAIRが公開したものは人工知能コードではありませんでした。研究所は自分たちが開発したツールで基礎科学の古典的な課題を解き、その結果をそのまま世に公開しました。生命の言葉であるタンパク質を扱った研究がその一例です。遺伝子配列だけからタンパク質が折り畳まれた3次元構造を予測するモデルを開発し、数億のタンパク質構造の地図を公開しました。世界中のどの実験室でも、大規模な製薬企業でも無名の研究室でも、その地図を無償でダウンロードしてワクチンと医薬品開発に活用できるようになりました。気候変動対策に関する業績もありました。化石燃料の代替触媒とエネルギー貯蔵材料を発見する支援のため、数百万件の物理計算結果を含むデータセットを公開しました。材料科学の研究者たちの進捗速度はそれだけ加速しました。人工知能研究所が物理学、化学、生物学の基礎を築いたことになったのです。

しかし、公開には代償が伴いました。その代償を痛切に払わせたのが2022年秋のGalacticaでした。FAIRの小規模な研究チームが世界の科学論文と教科書をまるごと学習させて、科学分野向けの言語モデルを開発しました。その名前はGalacticaでした。研究チームは誰もが試験できるようにインターネットにデモンストレーション用ウェブサイトを公開しました。反応は激しいものでした。このモデルはもっともらしい文を生成することに優れていましたが、その文が事実かどうかは区別できませんでした。存在しない論文を作り上げ、誤った内容を自信満々に述べました。批評家たちはこれを科学の名を借りた無意味な出力を大量生産する機械だと非難しました。その指摘は一理ありました。権威ある言い回しで装飾された虚偽は、単なる嘘よりも危険だったからです。非難は数日のうちに殺到し、FAIRは4日でデモサイトを削除しました。そのモデルを開発したのは若い研究者たちの小さなグループであり、彼らは夜眠ることができませんでした。

Galacticaの失敗は二つのことを同時に示していました。世に出す前に十分に検証しなかった研究の性急さが一つであり、公開したからこそその欠陥が即座に明らかになり改善できたという点がもう一つでした。非公開のままだったなら、誰も知ることはなかったでしょう。公開したからこそ批判を浴び、批判を浴びたからこそ世界全体が問題を認識するようになったのです。ルンが信じた開放性とはもともと

とそういうものでした。間違ふ可能性のあるものを間違ふ可能性のままに世に出し、その誤りを多くの人の目で正していく方法でした。居心地の良い道ではありませんでした。そして会社がいつまでもこの不便を受け入れるかは、その時点ではまだ誰も問いませんでした。

振り返ると、Galacticaが世に出てから数週間も経たないうちに、別の会社からチャットボットが登場しました。その名前はChatGPTでした。同じようにもっともらしい文を生成し、同じように間違った情報を自信を持って提示する機械でしたが、世の受け止め方は正反対でした。一方は嘲笑と共に4日で削除され、もう一方は数ヶ月で世界で最も急速に広がったプログラムになりました。その違いを生み出したものが何かは、後の議論の中で明らかになります。ただ、このズレがルクンに残した傷跡は深いものでした。

PyTorchという公共財

今日、世界中のどの大学の人工知能実験室を覗いても、画面のコードの最初の行の付近には同じ名前が書かれています。PyTorchです。ニューラルネットワークを構築した経験のある人で、このツールを使わなかった人はほぼいません。そしてこのツールが生まれた場所がルクンのFAIRでした。

PyTorchが何かは、それが存在しなかった時代を思い出すと理解しやすくなります。かつて、ニューラルネットワークを構築しようとしたら、まず計算全体の設計図を最初から最後まで完成させ、その後で初めてデータを流すことができました。設計図が完成するまで、途中の計算結果を確認することはできませんでした。何か問題が生じて、どこが間違っているのか判明しにくかったのです。PyTorchはこの順序を逆転させました。コード行を一つ書くと、その場ですぐに計算が実行され、中間のどの値でも停止させて調査できるようになりました。研究者の思考に応じて機械も動作する方式でした。使い慣れたPython言語の上で流暢にコードを書けたため、新しいアイデアを思いついた研究者はそれを直ちに試験することができました。頭の中の着想と画面上の実験の間の距離が縮まりました。研究のテンポが変わったのです。

PyTorchには起源がありました。その源はTorchという古いツールでした。Torchはスイスの研究機関で開発され、Luaという広くは知られていないプログラミング言語上で動作していました。計算速度は高速でしたが、使用者は限定的でした。世界中の研究者が習得していた言語はLuaではなくPythonだったからです。FAIRのエンジニアのサミス・チンタラと同僚たちがこの課題に取り組みました。Torchの高速な計算エンジンはそのまま保持し、その上に研究者が使い慣れたPythonというインターフェースを新たに追加しようということでした。2016年頃に開始されたこのプロジェクトの成果が、翌年初めに発表されたPyTorchでした。数年後には、会社内で製品開発に使用されていた別のツールと統合され、実験室で開発されたモデルを大規模な本番製品にそのまま移行できるようになりました。研究の自由さと実運用の安定性が一つのツール内で共存することになったのです。

PyTorchが登場した頃、この分野を支配していたのはGoogleが開発したTensorFlowでした。TensorFlowが先に登場し、企業の力によって広く普及していました。これら2つのツールは、世界を見る方法が異なっていました。TensorFlowは設計図を先にすべて描いてから動かす方式であったのに対し、PyTorchは1行ずつ書きながらその都度動かしてみる方式でした。大規模かつ安定的に動かすには前者の方式が有利でしたが、新しいアイデアをあれこれいじりながら実験するには後者の方式が便利でした。そのため、研究者たちが真っ先にPyTorchへと乗り換えました。実験の自由度が高いツールを選んだのです。

乗り換えるスピードは速いものでした。学会で発表される論文の中で、どのツールが使われているかを数えてみれば、その流れは一目でわかりました。2019年になると、新しく発表される論文の半分以上がPyTorchを使用しており、2022年頃にはその割合が4分の3を超えました。新しい手法を発表する人は、他人が自分の結果をそのまま再現してくれることを望んでおり、そのためには他人が使っているツールで作成するのが有利でした。このようにPyTorchを使う人が増えるほど、PyTorchを使う理由も増えていきました。Googleでさえ、自社のツールの次期バージョンをリリースする際、PyTorchの使いやすさに倣って手を加えなければならませんでした。研究の最前線では、勝負がついたようなものでした。Metaと熾烈な競争を繰り広げている企業たちでさえ例外ではありませんでした。OpenAIも、Anthropicも、Microsoftも、グラフィックチップを製造するNVIDIAも、PyTorchの上で自社の人工知能を育てました。ある1つの企業が作ったツールが、その企業の競争相手までも支える基盤となったのです。

ここには、一見すると辻褄が合わないように見えるところがありました。互いに勝とうと争っている企業たちが、いざその争いの基盤となるツールについては共に使い、共に改良していったのです。一方では市場を巡って死活問題として争いながら、もう一方では同じ工具を分け合って使っているわけです。これが公共財の姿でした。村の井戸や町内の道のように、誰のものでもないのに皆が使う資産のことです。それをそれぞれが別々に掘れば村全体が疲弊してしましますが、1つと一緒に掘って分け合えば、その力を別のところに回すことができます。PyTorchがそのような井戸になりました。企業たちが各自のツールを最初から作るために注いだであろう力を、その上で実際に何を作るのかというところに注げるようになったのです。値段をつけて売れるものを無償で開放し、公共の井戸にしたわけですが、その決定がむしろ分野全体の歩みを速めました。

ここでMetaは、珍しい決定を下しました。2022年9月、同社はPyTorchの所有権を手放しました。Linux Foundationという非営利団体の下に新たにPyTorch財団を設立し、ツールの運営をそこへ移管しました。新しい財団の理事会には、Metaだけでなく、AmazonやGoogle、Microsoft、NVIDIA、半導体企業のAMDも名を連ねました。互いに競争する企業たちが、1つのツールの運営を分担して引き受けることになったのです。こうすることで、PyTorchはいずれか1つの企業のもものではなくなります。Metaが心変わりして門を閉ざしてしまう心配も消え去ります。ツールを共用の財産とし、その中立性を決定づけた決断でした。

ルカンとは、このような考えを道路に例えて説明することがありました。2つの都市を結ぶのに、企業ごとに自社の高速道路を別々に10本ずつ敷くとしたら、それは無駄でしかないということです。計算の基盤となる基本ツールは、誰もが一緒に走る1つの共用道路であるべきだというのが彼の考えでした。各自が壁を築く代わりに、道と一緒に整備しようということです。彼は、人工知能の基盤に置かれる技術ほど、開かれていなければならないと考えました。基盤が一部の企業の私有物になれば、その上に何を建てるにしても、その企業に通行料を支払わなければならないからです。

ルカンがこの道路の比喻を好んで用いたのには、先立つ手本があったからです。インターネットでした。今日の世界を結ぶその網は、いずれか1つの企業の私有物ではありませんでした。誰もが無償で使える開かれた約束事の上に築かれました。ウェブを支える言語も、電子メールが行き交うルールも開かれていました。その開かれた基盤の上で、数多くの企業やサービスが育ちました。もしインターネットの基盤が1つの企業のものであったなら、今日の網は存在しなかっただろうとルカンは見ていました。人工知能もそうでなければならないということ、その基盤だけはインターネットのように共用のものとしておいてこそ、その上で様々なものが育つというのが彼の信念でした。

この考えは、ツールを超えて人工知能モデルそのものへとつながりました。MetaはLlamaという名前の大規模言語モデルを作りながら、その中身を埋める数値を世に公開しました。最初のバージョンが2023年初めに出され、数ヶ月後に出された次のバージョンでは、企業が作る製品にも取り入れて使えるように、門戸をさらに広く開きました。モデルの重みを公開するということは、高価な使用料を払えない人もそのモデルをダウンロードし、自分自身のものとして手直しして使えるということを意味していました。大企業の門前で引き返さざるを得なかった人々にとって、これは別の世界を開いてくれました。インドの小さな大学の実験室が自国の複数の言語を扱うモデルを作り、アフリカの研究陣が自地域の言葉で医療情報に答えてくれるプログラムを組むことが、他人の許可を得なくても可能になったのです。高価な門を通り抜けるお金がなくて引き返さざるを得なかった人々が、今や開かれたツールと開かれたモデルをダウンロードし、自分の手で何かを築き始めたのです。医師が大幅に不足している地域で人々の問いに自分たちの言葉で答える小さな医療アシスタントを作る仕事、大企業が採算に合わない見向きもしなかった言語を機械に教える仕事が、そのようにあちこちで起こりました。もしその基盤が一部の企業の閉ざされた門の裏にのみあったとしたら、このような試みは始まる前に塞がれていたことでしょう。ルカンはこれを技術の主権と呼びました。未来の人工知能が一部の企業の手の中に閉じ込められれば、世界のあらゆる言語と文化がカリフォルニアの片隅の目を通して濾過された姿だけが残ることになるだろうと、彼は警告しました。知識の基盤は、Wikipediaのように誰もが覗き見し、直すことができなければならないということでした。1つの企業が世の中のすべての文化を盛り込んだたった1つの閉じたモデルを握りしめようとする試みは、不可能であるだけでなく危険であると彼は語りました。

この場には反対側の声もありました。モデルの中身を丸ごと開いておけば、悪い心を抱いた人もそれを手に入れることになるという心配でした。安全装置をかけておいても、開かれたモデルは誰でもその装置を取り外すことができるため、危険な技術を世にばらまくことになるという懸念でした。この心配をする人々の中には、ルカンと共にディープラーニングを起こした長年の同僚たちもいました。論争は今も終わっていません。ルカンは、開いておいた時に得るものが、失うものよりも大きいと見ました。大勢が覗き見る技術のほうが、少数が隠しておいた技術よりもむしろ安全に磨き上げられるということ、危険は隠すことによってではなく、明るい場所で一緒に直すことで減らさなければならないというのが彼の考えでした。どちらが正しいかは、この本が答えられる問いではありません。ただし、ルカンがどちらの側に自分の名を懸けたのかは明白でした。

もちろん、開放に理想だけが合ったわけではありません。企業がツールやモデルを開いておくのには、打算もありました。世界中の研究者たちが自社のツールに手を慣らせば、その企業は技術の標準を握ることになります。人を採用するのも、エコシステムを導くのも容易になります。値段をつけずに明け渡すことが、むしろ競争相手の足を揺るがす手になることもあります。開放は純粋な贈り物でありながら、同時に賢明な布石でもありました。2つはきれいに分けられませんでした。ただ、その打算が何であれ、結果は一方向に明らかでした。お金がなくて技術の敷居を越えられなかった人々が、その敷居を越えられるようになったということです。塀の根元が低くなった場所で誰が利益を得たのかは、打算よりも長く残る事実でした。

振り返ってみると、PyTorchを他人の手に渡したその決定は、9年前のザッカーバーグの食卓で交わされた3番目の条件が、目に見える形となって現れたものでした。すべてのものを開いて発表するという約束です。その約束は、論文数編やコード数行にとどまらず、企業が持つ指折りの貴重なツール1つを丸ごと共用の財産として差し出すまでに至りました。言葉による約束が、実物として守られたのです。世間がその約束の果実をどう使ったのかは、もはや企業の手を離れた話でした。ツールは井戸となり、井戸は村のものとなりました。この箇所においてだけは、食卓で交わされた3つの言葉が空言ではなかつ

たことがはっきりと現れていました。

FAIRは、大企業の真ん中で大学の開放性を守り抜いた稀な実験でした。ザッカーバーグの夕食の食卓で始まった3つの約束が、その実験を支えました。ニューヨークに残り、大学に足を掛け、すべてのものを公開するという約束。その約束が守られている間、FAIRは世界で指折りの開かれた研究所でした。しかし、企業はいつまでも同じ企業にとどまるわけではありません。市場が急ぎ立て、競争相手が先行し、目の前の製品が急務になる日が来たら、遠い空を眺めていた研究所の居場所はどのようなのでしょうか。PyTorchを共用の道路として明け渡したその手が、再び壁を築けという圧力にいつまで耐えられるのでしょうか。その問いの答えは、まだ数年先にありました。ただし、開放を条件に掲げて企業の門を開いて入ったその人は、その条件が揺らぐ日に何をするかもすでに心の中で決めていたことでしょう。

弁護士

kimkj.com

第10章 チューリング賞、栄光と優先権

ヤン・ルクン、人工知能の父

2018年 三人のチューリング賞

2019年3月27日の朝、アメリカコンピュータ協会 (ACM) が一つの声明を発表しました。2018年度のA.M.チューリング賞をヤン・ルクン (Yann LeCun)、ジェフリー・ヒントン (Geoffrey Hinton)、ヨシユア・ベンジオ (Yoshua Bengio) の三人に共同で授与するというニュースでした。協会は三人をディープラーニング革命の父たちと呼びました。受賞の理由は簡潔で明確でした。深層ニューラルネットワークをコンピューティングの核となる要素にした概念的、工学的な革新を成し遂げたということでした。

チューリング賞はコンピュータ科学における最高の賞です。人々はこの賞をコンピューティングのノーベル賞と呼びます。賞金は100万ドルであり、グーグルが後援しました。三人はその年の6月15日、サンフランシスコで開催された協会の年次授賞式で賞を受けました。グーグルの研究を率いるジェフ・ディーン (Jeff Dean) をはじめとする産業界と学界の人士が祝辞を送りました。その頃、人工知能は新聞の一面と夜間ニュースの常連トピックでした。数年前までは大学院の片隅のテーマだったニューラルネットワークが、今や各国政府の戦略文書と企業の事業計画書の中心に置かれていました。この賞がそれほど広い反響を得たのには理由がありました。つい最近まで、ニューラルネットワークは学界で廃れた分野として扱われていたからです。

協会の公式説明は、三人が成し遂げたことを一つ一つ挙げました。誤差逆伝播による学習、畳み込みニューラルネットワーク、そしてニューラルネットワーク自身が表現を学ぶようにした概念の基礎でした。その技術が音声認識と機械翻訳とコンピュータビジョンを変え、さらには医学とロボット工学にまで広がったと協会は述べました。一つの賞に収めるには幅広いリストでした。しかし、その幅広さこそが、この賞が特別だった理由でした。一人の業績に留まらず、一つの時代全体の方法が変わったことからです。

その方法の転換を三人が単独で成し遂げたわけではありませんでした。数十年にわたる実験と失敗、コンピュータの発展、インターネットが積み重ねた膨大なデータが共に推し進めた結果でした。ただし、その流れの流路を開いた場所ごとに、三人の名前が置かれていました。この賞は、その流路を讀えたものでした。

この場面を理解するには、この本が歩んだ道を少し振り返る必要があります。人工ニューラルネットワークは二度の冬を経験しました。最初の冬は1969年にマービン・ミンスキーとシモア・パパートがパーセプトロンの限界を指摘した後にやってきました。二番目の冬は1990年代半ばに、サポートベクターマシンなどの統計学習が学界を支配するようになりました。冬が深い時期には、ニューラルネットワークの研究をしているということは、学者として自分の将来を自ら塞ぐことに近い意味でした。優れた学会は論文を却下し、研究費は別の名前を付けたプロジェクトへ流れていきました。その二つの冬の間にも、そして冬の真つ最中にも、ニューラルネットワークを掴んで放さない少数者がいました。この三人はその少数者の中心にいました。

三人が一つの場に立つようになるまでには長い経緯がありました。ヒントンとルクンは1985年にアルプスのレ・ホッシュで初めて会い、ルクンは1987年にトロントのヒントン研究室でポストドクター

として一年を過ごしました。ベンジオは1990年代初頭、ベル研究所でルクンと同じ廊下を使っていました。二番目の冬が学界を覆った時、三人はカナダ高等研究院 (CIFAR) の支援の下で、互いの論文を読み合いながら耐え抜きました。学界の外では、ディープラーニング陰謀団などと冗談めかして呼ばれていた、先ほど見た那の集まりでした。この賞は、その長い支え合いに対する遅れてきた認定でした。

三人の貢献は互いに異なっていました。ヒントンは1986年にデイビッド・ラメルハート、ロナルド・ウィリアムズとともに、誤差逆伝播 (backpropagation) によってニューラルネットワークの隠れ層が有用な表現を自ら学習することを示した論文の著者でした。彼はボルツマンマシンを提案し、2006年には層を一つずつ事前学習する方法で深いニューラルネットワークを再び蘇らせ、2012年のImageNetコンテストを覆したAlexNetの師でした。彼の二人の弟子、アレックス・クリジェフスキーとイリヤ・スツケバーが、その年のコンテストの誤差率をほぼ半分に削減しました。その出来事がディープラーニングの春を開きました。

ベンジオは異なる気質の人でした。彼は2003年にニューラルネットワークで言語をモデル化する論文を発表し、単語を数字のベクトルに変える今日の表現の種を蒔きました。2014年には弟子たちと共にニューラルネットワーク翻訳に注意 (attention) メカニズムを追加し、その着想は後にトランスフォーマーの基礎となりました。敵対的生成ニューラルネットワーク (GAN) も彼のモントリオール研究室から生まれました。三人の中で、ベンジオは最後まで大学に留まった人でした。彼が設立したMila研究所はカナダの人工知能の心臓となりました。

ルクンの位置は、この本が前に長く描いてきました。彼は畳み込みニューラルネットワークを作りました。生物の視覚皮質から得た直感で設計されたこのニューラルネットワークは、1989年に手書きの郵便番号を読み取り、1990年代後半にはアメリカの小切手のかなりの部分を認識しました。三人の技術が互いに異なる分野で成長し、一つの波を形成しました。協会の表現によれば、彼らは独立して、また共に働きました。概念の基礎を確立した人も、実験で予期しない現象を発見した人も、工学でニューラルネットワークの有用性を証明した人も、この三人の中にいました。

2019年の三人は異なる道の上に立っていました。ヒントンはトロント大学とグーグルの間に身を置いており、ベンジオはどの企業にも縛られないままモントリオールを守り、ルクンはフェイスブックの人工知能研究所を率いながらニューヨーク大学の教授職を手放しませんでした。純粋学問と巨大産業の間で、三人はそれぞれ異なるバランスを選びました。その異なる選択が後に各自の次の章を分けました。

三人は賞を受けながら、それぞれの方法で感謝を伝えました。しかし、三人ともこの賞を自分たちだけのものとは見なしていませんでした。長年ニューラルネットワークを信じてきたコミュニティ全体のだと述べました。謙虚な言葉でありながら、事実の言葉でもありました。ディープラーニングは一つの研究室の発明ではなく、複数の国、複数の世代の手を通じた長いリレーだったからです。まさにそのリレーの先行者たちをどのように記憶するのが、すぐにこの賞の影として浮かび上がるようになります。

三人はチューリング賞の講演を共同で準備し、ディープラーニングが何であったかを自分たちの言葉で整理しました。要点は一つでした。知能は手で規則を書き込んで作るものではなかったということです。データから自ら学ぶようにする方が正しいということでした。一世代前までは少数の異端者が守っていたこの考えが、今やコンピュータ科学の正統説になっていました。それでも、ルクンは講演でまだ行く道は長いというセリフを忘れませんでした。機械はまだ、子どもが数ヶ月で学ぶ世界の常識を学んでいなかったのです。頂上に立ったその日でも、彼の目は次の山を見つめていました。

三人の関係が常に和やかだったわけではありません。彼らは同志であり、競争者でもありました。同じ学会で異なる方法について争い、誰のアプローチが正しいかについて声を上げました。しかし、ニューラルネットワーク全体が嘲笑を受けていた時代には、その争いさえ贅沢でした。外の冬が内の意見の相違を覆い隠しました。三人は一緒に生き残り、一緒に頂上に立ちました。

この賞は三人に与えられましたが、この賞を共に受ける資格があった人々はもっとたくさんいました。冬の間にニューラルネットワーク論文を書いて学界を去った人々、職を得られず別の分野に移った人々、名前を残さなかった大学院生たち。三人の演壇の後ろには、そのように消えていった人々の影が長く落ちていました。チューリング賞は生き残った者たちの賞でした。そして生き残ったこととなることが正しかったことは、この分野では長の間、同じ意味ではありませんでした。

ディープラーニングはその頃すでに世界を変えていました。音声認識が人間の話を書き取り、翻訳機が文を訳し、カメラが物体を認識しました。病院の画像診断と自動運転車の目にもニューラルネットワークが組み込まれていました。これらすべての技術の根底に、三人が冬中ずっと守ってきた着想がありました。ルクンにとって、この賞は個人の軌跡が時代と噛み合った場所でした。誰もニューラルネットワークを信じなかった時代にニューラルネットワークを選んだフランスの青年が、60歳を前に自分の分野の頂上に立ちました。手書きの数字一つを読み取らせるためにベル研究所の夜を徹した人が、今やコンピューティングの頂上に置かれた賞を受けました。彼が作った機械が異端から文明の骨組みに変わるまで、三十年かかりました。

しかし、この賞が三人にだけ与えられたという事実が、すぐに別の問いを呼び起こしました。なぜ三人なのか。ニューラルネットワークの長い歴史には、この三人以外にも多くの名前がありました。パーセプトロンを作ったフランク・ローゼンブラット、畳み込みの原型となったニューラル構造を1980年に発表した福島邦彦、確率的勾配降下法の基礎を置いたアマリ・シュン'イチなどでした。賞は常に数人だけを指名し、残りを影に残します。祝杯が上がっていた正にその頃、学界の別の場所では、その影の名前で反論が研ぎ澄まされていました。その問いを投げかけた人は、ニューラルネットワークの歴史で決して小さくない位置を占める別の科学者でした。スイスのルガーノで、ユルゲン・シュミットフーバーがキーボードの前に座っていました。

シュミットフーバーの反論

機械学習は信用を正しく配分する科学である。ユルゲン・シュミットフーバー (Jürgen Schmidhuber) はずっと前からこの文を自分の信念の最初に置いていました。彼は2015年に自分のブログにこう書きました。ある方法を発明した人がその発明の信用を受けるべきだ。その人が必ずしもその方法を広く広めた人である必要はない。発明した者と流行させた者は異なるというこの区別が、彼が生涯にわたって手放さずに戦った原則でした。

シュミットフーバーを変人の座に押しやって、この議論を読むと、物語の半分を失うことになります。彼の名前は神経網の歴史に太く刻まれています。彼はスイス人工知能研究所 (IDSIA) を率いて、一世代の研究者を育てました。1997年に弟子セップ・ホックライター (Sepp Hochreiter) と共に発表した長短期メモリ (LSTM) は、複数の層を経過する中で消失していた神経網の記憶を捉えた構造でした。この構造は長きにわたって音声認識と翻訳の標準となっていました。Google翻訳も、AppleとAmazonの音声アシスタントも、かつては彼の研究室から生まれた設計を心臓に抱いて回っていました。弟子アレックス・グレイブスが作った方法は、手書きと音声をつなぐ橋になり、別の弟子たちはLSTMに忘却の門をつけてそれをより強固にしました。彼は人工的好奇心、メタ学習のような時代を先取りした発想も

数多く残しました。そういう人物がチューリング賞の発表を見守っていたのです。

彼の研究室はヨーロッパ人工知能の孤立した城塞のような場所でした。アメリカの大型大学と企業が資源を吸い上げていた時代、彼はスイスの小さな研究所で世界と異なる道を切り開いていました。限られた資源の環境から大きな構想をくみ上げることに、彼は慣れていました。そのような場所で長く耐えた人ほど、自分がくみ上げたものを他人が持ち去る光景に耐えるのが難しいものです。彼の反論を理解するには、その孤立した城塞の孤独を共に見なければなりません。

その孤独は彼を頑固にもしました。広い学界の認定と資源を享受した三人が数歩引く余裕があったとすれば、彼にはそのような余裕は少なかったのです。自分の研究が自分の名前で記憶されることが、小さな研究所の科学者にとってはほぼ唯一の資産だったからです。この争いには個人の性質だけでなく、学界の資源がどこに集中するのかという構造的な問題も含まれていました。

彼の見方では、三人は他人が築いた基礎の上に自分の建物を立てながらも、基礎を築いた人たちの名前を書かなかったのです。この不満は2018年に急に生じたものではありませんでした。すでに2015年、ヒントン、ベンジオ、ルクンが権威ある学術誌『Nature』にディープラーニングをまとめた総説を掲載した時、シュミットフーバーはその論文を項目ごとに反論する批評を発表していました。彼はその論文が自分の分野の先駆者たちを消してしまったと見ました。同年、彼は自身が執筆したディープラーニングの歴史概観論文で、数百篇の古い研究を一つ一つ引用しながら、このように引用することが学者の務めだと身をもって示しました。チューリング賞はその古い火種に油を注いだようなものでした。

2020年6月、シュミットフーバーはその考えをひとつの文書にまとめて発表しました。題名は「ベンジオ、ヒントン、ルクンの2018年チューリング賞受賞に対する批判」でした。彼はこの文書をその後さらに2度改訂し、出版日と根拠を細密に充填しました。文書は論文というより告発状に近いものでした。彼は日付を一つ一つ対照させながら、三人が自分の業績だと述べた多くの技法の根源が、より以前の別の学者たちにあったと主張しました。

彼の方法は執拗でした。彼は各技法について最も早い出版日を見つけて、判を押すように記しました。そして三人が後に自分の論文でその早い日付を飛ばして、自分たちの後発の研究だけを繰り返し引用したと指摘しました。ベンジオについては、ホックライターの1991年の研究を知りながらも、自分の1994年の論文だけを繰り返し引用したと記しました。自己引用が積み重なると引用回数が上がり、引用回数が上がると権威が高まります。シュミットフーバーが狙ったのはその静かな雪だるまでした。

彼が挙げた最初の事例は誤差逆伝播でした。ヒントンとルクンが1980年代に普及させたこの学習法の数学的な骨組み、すなわち逆方向の自動微分 (reverse mode of automatic differentiation) は、すでに1970年フィンランドのセッポ・リナインマー (Seppo Linnainmaa) が確立したというものでした。1974年には、アメリカのポール・ワーボスがその方法を神経網に適用する道を論文で開きました。神経網が複数の層を積み重ねて自ら表現を学ぶ深層学習の構想そのものも、ヒントンやルクンが学界に入る前の1965年、ウクライナのアレクセイ・イヴァクネンコ (Alexey Ivakhnenko) に由来すると彼は記しました。イヴァクネンコは複数の層から成る学習機械をその時すでに紙の上に描いていました。シュミットフーバーはこれらの名前が教科書から消えて、その場所に三人の顔が入ったと見ました。

これらの名前が忘れられたのは、悪意だけではありませんでした。イヴァクネンコの研究はソビエトの学術誌に掲載され、冷戦の幕の向こう側の西側学界には十分に届きませんでした。リナインマーの論文はフィンランド語と数値解析という狭い門の中にありました。言語と地理と時代が名前を埋もれさせました。発見を世界に広く広めた人に功が集中するのには、盗難だけでなく、このような無心な忘却の

力も働いていました。シュミットフーバーはその忘却と戦い、三人はその忘却の受益者でありながら同時に神経網を蘇らせた労働の主人でもありました。

確率的勾配降下法も似ていました。今日ほぼすべての神経網がこの方法で学ぶのに、その根のひとつが1960年代日本の甘利俊一にありました。彼の名前もずっと西側の教科書の脚注にとどまっていた。神経網を支える柱のすべてに、このような忘却された名前が一つずつついていました。シュミットフーバーの文書はその脚注を本文に引き上げようとする試みでした。

時が経つにつれ、その試みが全く無駄ではありませんでした。今日、良い教科書は逆伝播の歴史において、リナインマーとワーボスの名前と一緒に記しています。ホックライターは後にオーストリアのリンツで自分の研究室を率いながら、ディープラーニングの一翼として堂々と立ちました。記録は遅く、しかし少しずつ正されました。ただし、その訂正がシュミットフーバーが望んだ速度と方式だったかどうかは別の問題でした。

ここで一つのことを指摘しておくことが公平です。シュミットフーバーが挙げた歴史的事実そのものはおおむね正しいです。誤差逆伝播の原型がリナインマーとワーボスにあったということは、この分野の歴史家が広く認める事実です。反転はここから生じます。ヒントン自身も様々な場で、私は逆伝播を発明しなかったと述べています。彼は自分が行ったことがその方法を最初に作ったのではないことを明らかにしました。自分が行ったことは、その方法で神経網が有用な内部表現を習得することを世界に示したことだと述べました。ルクンもまた、自分の畳み込み神経網が福島島の1980年の神経構造に負債があることを複数の論文で明らかにしてきました。三人が常に先駆者を消したわけではありませんでした。そして三人がチューリング賞を受けた理由も、逆伝播を最初に作ったからではありませんでした。その学習法が実際に機能することを証明し、神経網を有用な機械に変えたことでした。発明と実証のこの隙間をどのように読むかによって、同じ事実がまったく異なる物語になります。

実証の重みを軽く見ることはできませんでした。イヴァクネンコが複数層の学習機械を描いた1965年には、その機械を動かすためのデータも計算装置もありませんでした。紙の上の構造と世界を変える機械との間には、数十年のデータと数千倍高速化した計算と多くの試行錯誤が介在していました。三人が行った仕事の大部分は、その遠い距離を渡ったことでした。構想を最初に記した人と、その構想を最後に動かした人。科学は両者に負債を負っていますが、世界は通常後者だけを記憶しています。シュミットフーバーは前者を記憶するよう叫び、その叫びにも一理ありました。

シュミットフーバーはその隙間を背信と読みました。科学の賞は最初に作った人に返されなければならないというのが彼の確信でした。広く知らせた人のものになってはならないということでした。彼は科学が多数決で決まらないうと繰り返し述べました。彼が好んで引き出した比喻は、1931年ドイツから出版された『アインシュタインに反対する100人の著者たち』でした。百人が一声で間違っている、事実を持つ一人が正しくあり得るという話でした。彼はこのように記しました。厳密な科学では、世論や人数は何も決定しない。真実は太陽のようなもので、雲が一時的に覆うことはできても消えることはないという言葉も付け加えました。

こうした言葉は悲壮でしたが、聞き手を説得するというより退かせる力がありました。正しさを確信する人の声は、時として大きすぎて、その正しさの内容そのものまで一緒に埋もれてしまうことがあります。シュミットフーバーの困難はそこにありました。彼の手には本当の資料が握られていましたが、その資料を握る方式が人々を彼の味方から押しのかました。

そのような彼の戦いは神経網に限定されたものではありませんでした。彼は数年にわたって、科学史の様々な局面で忘却された先駆者を復元する手紙を学術誌に送りました。最初のプログラム可能コンピュータを自宅のリビングで作ったコンラート・ツーゼ (Konrad Zuse) をアラン・チューリングと並べるべきだと主張し、電話の発明についても、ベルより前の名前を挙げました。ライト兄弟以前の飛行実験家たちを照らす文章も書きました。ある者たちはその執拗さを学問的良心と呼び、他の者たちは名誉に対する強迫と呼びました。両方の評価が彼に付きまとっていました。

ある悲しい逆説がありました。シュミットフーバー自身も、功績を奪われた人ではありませんでした。彼がホックライターと共に作ったLSTMは、20年以上にわたって人類が広く使用してきたアルゴリズムの一つでした、彼の名前はこの分野のすべての教科書に大きな文字で掲載されています。名誉について誰より激しく戦った人が、既に十分に認められた人でもありました。その事実が彼の戦いをより理解しがたく、同時により純粋に見えるようにしました。地位やお金のための戦いではなかったからです。彼が守ろうとしたのは記録の正確性という、手に取ることができない何かでした。

ひとつ付け加えておきます。シュミットフーバーが早期に取り組んだ主題の一つがワールド・モデルでした。彼は1990年代から、機械が自らの内に世界の縮小版を構築し、それで未来を予測させる研究を行っていて、2018年には弟子と共に「World Models」という題の論文を発表しました。後年、ルクンが大型言語モデルに背を向けて進んだ方向が、正にそのワールド・モデルだったのです。二人は認められることをめぐって争いましたが、知能がどこに向かうべきかについては、意外にも近い場所を見つめていました。論争の相手が同じ方向を指し示すことは、科学においては稀ではありません。

彼の問題提起の底には、答えがたい一つの問いが置かれていました。ある分野は自分の起源をどのように記憶すべきか。賞は数人のみを呼びます。教科書は紙面を節約するため名前を省きます。大衆は顔一つを選んで、その上にすべての功績を積み重ねます。そのようにして記憶は常に少数をあきらかにし、残りを消し去ります。シュミットフーバーはその消し去られることに反抗して、独り手を上げました。ただし明るみに出ることと消し去られることは常に一緒に起こる事柄であり、その事実が彼の戦いを終わらないものにしました。

一つは確かです。シュミットフーバーが提示した日付と論文は実在して、彼が構築したLSTMは一時代を支えました。しかし、歴史的根拠が他人にあるということと、三人がその根拠を意図的に隠したということは別の主張です。前の主張は史料で争うことができますが、後ろの主張は心の中の意図を問う事柄でした。史料は図書館にあります、意図はどこにもありません。論争が熱く盛り上がった場所も、正にその意図の座でした。そしてその座で、ルクンはシュミットフーバーと正面から衝突したのです。

引用と剽窃の論争

2016年12月、スペイン・バルセロナのある講演会でした。イアン・グッドフェローが舞台に立って、敵対的生成ネットワークを説明していました。二つのニューラルネットワークを向かい合わせて、一つは偽物を生成し、もう一つは本物と偽物を見分けるという構造でした。ベンジオのモンテリオール研究室から出たこの構想を、ルクンは最近20年間の機械学習で最も素晴らしいアイデアだと褒め称えていました。その講演が最中に、聴衆席から一人がマイクを握りました。シュミットフーバーでした。彼は質問の形式を借りて長い発言を始めました。要点は、この構造が自分が1990年代初頭に発表した予測可能性最小化 (predictability minimization) と本質的に同じなので、名前を変えるべきだということです。

講演会の空気が凍り付きました。このシーンはインターネットでディープラーニング・ドラマという名前で広がり、学会の掲示板とソーシャルメディアが数日間この話題で沸き立ちました。ただし、その後がどうなったかはしばしば半分だけしか伝わりません。グッドフェローはその場で自分の構造とシュミットフーバーの構造が異なると反論しました。そしてその後、論文を改めてシュミットフーバーの予測可能性最小化研究を引用しました。その際、二つの方法が分岐する点を記しました。GANでは二つのニューラルネットワークの競争それ自体が学習の唯一の目標である一方、予測可能性最小化は他の目的のための補助装置だったということでした。シュミットフーバーが引用を要求してグッドフェローが引用で応じた後、お互いの違いをめぐって争ったのが事件の実際の姿でした。一方が真実を要求してもう一方が沈黙で応じたという図式は、この論争を過度に平坦にしています。

二人の構造が本当に同じだったかは、今でも人によって答えが異なります。シュミットフーバーの古い構造でも二つのニューラルネットワークが互いに競い合っていました。しかし、その競い合いの目的と方式がGANと同じかどうかは、簡単には判定されません。ある人は外見が似ているから同じ根拠だと言い、ある人は目的が異なるから異なる発明だと言いました。このような判定には客観的な物差しがありません。何を本質と見なし、何を枝葉と見なすが、既に各自の観点に委ねられているからです。優先権争いがなかなか終わらない理由がここにありました。

それでもこの争いには敗者がなかったわけではありませんでした。勝ったのは人ではありませんでした。問いが勝ちました。ある発見の功績をどのように分けるか、ある分野は自分の過去とどのように向き合うか。この問いは神経ネットワークを超えたすべての科学にかかった古い課題であり、シュミットフーバーとルクンは、その課題を誰よりも騒々しく舞台の上に引き出した二人でした。

この課題を前にしてルクンが選んだ答えは、彼の性格に似ていました。彼は過去をめぐって争うよりも未来を開いておく方を選びました。過ぎた功績を誰に返すかで何年も費やす代わりに、これから出現する発見がはじめからすべての人のものになるようにしようということでした。この選択が彼をオープンソースの代表的な声にしました。そしてその声は、後年彼が身を置いた企業と正面から衝突することになります。

シュミットフーバーのリストは長かったのです。複数の層を通過する際に誤差信号が消える勾配消失問題を最初に数学的に明らかにしたのは自分の弟子ホッホライターだと彼は言いました。ホッホライターは1991年にドイツで書いた学位論文でその問題を規定して、それを解決するための装置が後年LSTMになったのです。ベンジオが1994年に同様の問題を扱った論文で、また2010年に重み初期化を扱った論文でホッホライターの1991年研究を見落としたというのがシュミットフーバーの主張でした。ヒントンは2015年に発表した知識蒸留、すなわち大きなニューラルネットワークの知識を小さなニューラルネットワークに転移する技法も、自分が1991年に既に発表したニューラルネットワーク圧縮方式と同じだと言いました。トランスフォーマーの核心である注意機構の種も1990年代初頭の自分の研究室のいわゆる高速重み研究の中にあり、その一連の論文の査読者の一人がヒントンだったと彼は記しました。

この中で知識蒸留をめぐっては特許まで絡みました。ヒントン側がその技法でアメリカ特許を取得し、シュミットフーバーはすでに自分が昔に発表した圧縮方式の特許で囲い込んだものだと見ました。発見の功績を争う事柄が学術誌の引用リストを越えて法の領域に広がったわけでした。ルクンが特許という物に信頼を置かなかった背景には、このような光景を近くで観察した経験もありました。

これらの主張はそれぞれ重さが異なっていました。あるものは日付で検討する余地があり、あるものは概念がどのくらい同じかについて専門家の間でも判断が分かれませんでした。勾配消失をホッホライターが

早期に扱ったという事実そのものは概ね受け入れられました。一方、知識蒸留や注意機構が昔の構造とどのくらい同じかについては、数式の外見だけが似ているのか、それとも本当に同じ構想なのかをめぐって意見が分かれました。シュミットフーバーは同じだと見なし、相手側は異なると見なしました。このような論争ではしばしばそうであるように、二つの陣営は各自、自分の目に見える類似性と相違を強調しました。

ここには時代の事情も重ねられていました。2012年以降、ディープラーニングは爆発しました。毎年数万編の論文が噴き出て、一年前の研究さえ古くなってしまいました。若い研究者たちは20年、30年前のドイツ語の学位論文やソビエト学術誌を掘り返す余裕がありませんでした。引用は怠惰になり、最近数年の有名な名前たちだけが互いを指し示しました。シュミットフーバーが怒ったのは個人の盗用だけではなく、このように記憶を短くする分野全体のスピードでもありました。

彼は引用という指標そのものが病んでいると言いました。学界が支える引用数は実際の科学の重さを量る秤にはなり得ないということでした。お互いに引用し合う人たちの間で膨らませるバブルに近いと彼は見ました。彼はそれを2008年の金融危機を招いた不良債権に例えました。価値がないものを互いに売買しながら価値を作り出す構造という意味でした。この比喻は鋭かったのですが、同時に彼自身も引用の世界の中で争っていました。彼が要求したのは結局自分の名前を引用リストに載せろということだったからです。彼の批判は制度を狙いながら同時にその制度の言語で話していました。

産業の現場についても彼は声を上げました。2016年にニューヨークタイムズがグーグル翻訳の新しいニューラルネットワークを大きく報道し、ヒントンをその革新の顔として照らし出したとき、シュミットフーバーは、正にその翻訳機の内部で文を転写していたのは自分の研究室から出たLSTMだと反論しました。グーグルの技術文書がLSTMを数十回引用しているというのが彼の根拠でした。この指摘には事実の中身がありました。その時期の神経ネットワーク翻訳がLSTMに大きく依存していたのは本当です。ただし同じ翻訳機にはベンジオ研究室から出た注意機構を含むさまざまな系統の研究も一緒に入っていました。一つの技術の功績を一人に全く返すことは、称賛するときも非難するときも等しく困難でした。

ルクンは沈黙しませんでした。三人の中で最も前に出て言い返した人が彼でした。彼はニューヨークタイムズにこのように言いました。「ユルゲンは認められることに病的に執着し、受ける資格がない功績を複数の事柄で自分のものだと言主張し続ける。だから彼はどんな発表が終わるたびに立ち上がって、今発表されたばかりのものが自分の功績だと、大概是正当な根拠なく主張する。」鋭い言い方でした。シュミットフーバーはすぐに言い返しました。ルクンが一つの事例も示さないままそのような非難をしていることでした。そして彼はヒントン、ベンジオ、ルクンと絡んだ優先権争いのリストを改めて一度詳しく公開しました。

二人の言葉が衝突する座で、この論争の性格がそのまま露わになります。シュミットフーバーは相手が他人の功績を横取りしていると言い、ルクンは相手が他人の功績を無理矢理自分のものだと言っていると言いました。一方は過小評価を告発し、もう一方は過剰請求を告発しました。興味深いのは、二人が同じ言葉を使っていたという点です。功績を正しく分けようという言葉でした。ただし秤の目盛りをどこに置くかが正反対でした。シュミットフーバーの秤は最初に発表した日付を指していて、ルクンの秤はその構想を実際に通わせた労働を指していました。ベンジオはこの争いで、ルクンほど声を大きくしませんでした。三人の気質がそのまま応答の温度として露わになりました。

科学史を長く見つめてきた人々は、このような争いを例外ではなく規則であると言います。社会学者のロバート・マートンは、重要な発見であるほど、複数の人がほぼ同時に、互いに知らないまま到達するケースが多いことを示しました。微積分を巡ってニュートンとライプニッツが争い、進化論を巡ってダーウィンとウォレスが同じ時期に同じ考えに至りました。発見に父親が複数いることは欠点ではありませんでした。科学が進んでいく仕組みだったのです。複数の人が同じ壁の前に立つと、そのうちの何人かがほぼ同じ扉を見つけ出します。ニューラルネットワークの歴史もそうでした。バックプロパゲーションも、畳み込みも、注意機構も、一人の頭脳から完全に湧き出たものではなく、多くの手を通して洗練されました。そう考えれば、シュミットフーバーの日付も正しかったですし、3人の労働も正しかったのです。間違っていたのは、発見の父親は一人だけだという思い込みだったのかもしれない。

だからといって、この争いが無意味だったわけではありません。シュミットフーバーの執拗さのおかげで、リンナインマーとイヴァフネンコとホッフライターの名前が再び脚光を浴びました。忘れられかけた記録が蘇りました。一人の不快な叫びが、分野の記憶を少しだけ長くしました。彼が正しかったにせよ行き過ぎだったにせよ、科学はこうした不快な声のおかげで、自らの過去を忘れにくくなるのです。

学界の反応は分かれました。少なからぬ人々が、シュミットフーバーの挙げた歴史だけは正しいと認めました。しかし、彼のやり方、すなわち他人の発表を遮って名前を変えろと要求するやり方には背を向ける人が多くいました。正しい事実も、相手を追い詰めるような態度に乗せられると力を失いました。彼が孤立したのは、そうした理由もありました。この争いには簡単な判決はありません。バックプロパゲーションのルーツがリンナインマーにあるという歴史的事実と、その事実を3人が悪意を持って隠したという心理的判断は、同じ天秤には乗りません。前者は史料が答え、後者は誰にも確定できません。本書は、どちらか一方の手を挙げてこの章を閉じることはしません。ただ、この争いが残した痕跡の一つは、ルカンの個人的な選択に長く残りました。

ルカンは、こうした争いを避ける人ではありませんでした。彼はソーシャルメディアで昼夜を問わず論争し、相手が誰であれ引き下がらせませんでした。ある人々はその率直さを好み、ある人々はその鋭さに傷つきました。シュミットフーバーとの衝突も、その長年の論争家としての気質から生じた一場面でした。しかし彼は、争いの熱気とは別に、その争いが指し示すさらに深い問題を見逃しませんでした。功績を巡って繰り広げられる争いをどのように終わらせるのかという問題でした。

ルカンにとって、優先権の争いは抽象的な論争ではありませんでした。ベル研究所で自ら作った畳み込みニューラルネットワークの特許が、会社の書類棚へと散らばっていくのを見守った人でした。発見の特許という武器にするやり方がいかに研究者を縛り付けるかを身をもって経験した彼が出した答えは、相手への反撃ではなく、やり方の転換でした。功績を巡って繰り広げられる争いが終わらないのなら、最初からすべてを開放しておこうということでした。

誰が先に発見したかという争いは、すべてのコードが公開され、日付が公開リポジトリに刻まれる世界では、はるかに争いやすくなります。1編の論文の引用リストを巡って何年も争う代わりに、いつ何を世に出したかが消えない記録として残るからです。後日、ルカンがPyTorchとLlamaをオープンソースとして世に出し、閉鎖型モデルの時代に開放にこだわったルーツには、この頃の経験がありました。開放は彼にとって道徳であり、同時に記録の方法でもありました。そしてその信念は、彼の母国が長く抱いてきた価値観と重なっていました。

レジオンドヌールと母国

ルカンは自らをこのように紹介したことがあります。移民であり科学者、教授であり無神論者、アメリカ基準で左派でありフランス人。アメリカの右派が嫌うものすべての集合体だということです。冗談の形を借りていましたが、そこには自分がどこから来たのかを忘れないという決意が込められていました。世界の人工知能の中心で30年働いても、彼はフランス人でした。

そのルーツは1987年に遡ります。その年、パリ第6大学で博士論文の審査を受けた日、ルカンは松葉杖をついて審査場に入りました。彼が後日回顧したところによれば、少し前に海辺で風で走るサンドヨットを自作して乗っていたところ、足首を骨折したためでした。速度を落とすことを知らない人でした。その日の審査員の席には、トロントから飛んできたヒントンが座っていました。まだニューラルネットワークが学界の辺境にあった時代、2人はすでに味方同士でした。その日の合格を足がかりに、ルカンは妻と1歳の息子を連れてトロントのヒントン研究室へと渡りました。

松葉杖をつきながらもサンドヨットの話を楽しみながらする人。この場面にルカンという人物が詰まっています。彼は手で何かを作る人でした。航空工学者だった父親のそばで育ちながら模型飛行機を作り、大人になってからは風で走る乗り物を自ら設計し、楽器を手入れして電子回路をはんだ付けしました。ニューラルネットワークも、彼にとっては紙の上の理論ではありませんでした。動くように作らなければならない機械でした。発見の功績を実証の労働の方へと重く天秤にかけた彼の態度も、この「作る人」の気質と無関係ではありませんでした。図面を描いた人と、それを飛ばせるようにした人のうち、彼は常に後者へと心が傾いていました。

この気質は彼の学問全体を貫いていました。彼は美しい理論よりも動く機械を信じ、証明よりも実演を優先させました。ニューラルネットワークが嘲笑されていた時代に彼が耐え抜いた力も、ここから生まれました。論争で勝つ代わりに、機械が手書き文字を読み取る姿を見せれば済むことだったからです。優先権を実証の労働の方へと重く置いた彼の観点も、結局はこの昔からの気質の別の名前にすぎませんでした。

彼の信念は、それより前のパリの学生時代にすでに形を整えていました。彼は若い頃、人間の知能が生まれつきのものなのか学ぶものなのかを巡って繰り広げられた、チョムスキーとピアジェの長年の論争を読みました。そして「学び」の側に立ちました。知能の核心は、頭の中に前もって刻まれた規則ではなく、世界とぶつかりながら脳が自ら規則を鍛え上げる学習にあると彼は信じていました。この信念が彼をニューラルネットワークへと導き、ニューラルネットワークから一生離れられなくさせました。フランスの数学的合理主義と認知科学の伝統が、その信念の土壌でした。

2019年、チューリング賞を受賞したその年に、ルカンはフランス語で本を1冊出しました。タイトルは『機械が学習する時』でした。自分の人生とニューラルネットワークの歴史と一緒に織り交ぜたこの本を、彼は英語ではなく母国語で先に書きました。世界に向けては論文とコードを英語で出しましたが、自分の話を語る時はフランス語を選びました。その選択の中に、彼の2つのアイデンティティが並んで置かれていました。

フランスが彼を取り戻すのには長い時間がかかりました。才能ある若い科学者をアメリカに渡し、彼がそこで頂点に上り詰めるのを遠くから見守った国でした。2015年、コレージュ・ド・フランス (Collège de France) が彼をその年の情報学の講座教授として招きました。学位も試験もなく、ただ講義だけがあるこの由緒ある教壇は、フランス知性史の象徴であり、ベルクソンや数世紀にわたる知性が通り過ぎていった場所でした。翌年の2月、ルカンがその教壇に上がって就任講演を行い、どのような言葉を残したのかは、次の章で見る物語です。しばらくして、彼はフランス科学アカデミー (Académie

des sciences)の会員にも選ばれました。手書きの数字を読み取るニューラルネットワークに青春を懸けた少年が、母国の科学界の心臓部へと入っていったのです。

最も高い評価は2023年12月に訪れました。12月6日、エマニュエル・マクロン (Emmanuel Macron) 大統領がエリゼ宮でルカンにレジオンドヌール (Légion d'Honneur) シュヴァリエ勲章を授与しました。ナポレオンが創設したこの勲章は、フランスが市民に与える最高の栄誉です。ルカンはその日の感想を短く残しました。祖国が与えてくれた栄誉に感謝するというものでした。63歳の科学者が大統領の前に立ちました。彼の背後には、冬にニューラルネットワークにしがみついた若き日と、その信念がついに世界を変えた30年がありました。

母国という言葉には奇妙な重みがあります。ルカンは30年近くアメリカで暮らし、アメリカの大学とアメリカの企業で働きました。彼の子供たちはアメリカで育ち、彼の研究は英語で世に出ました。それでも、フランスが彼を呼ぶと、彼は喜んで振り返りました。去った人にも帰る場所があるということ、そして、そこが長く待っていてくれるということ。63歳の勲章は、その事実を確認してくれました。

彼が一生涯守り抜いた開放の原則が、ちょうど彼の母国が新たに必要とする原則になっていたのも、単なる偶然ではありませんでした。フランスは資源においてアメリカと中国に後れを取っていました。持っているものが少ない国が選べる道は、門に鍵をかけることではありませんでした。大きく開け放つことでした。開けてこそ他人のものを使うことができ、開けてこそ一緒に作ることができます。ルカンの古くからの信念とフランスの新たな状況が、その地点で出会いました。

フランスが彼を称えたのには、勲章を超えた理由がありました。ルカンは、アメリカの巨大テクノロジー企業が資本で知識の源泉を私有化する流れに立ち向かい、ヨーロッパで開放型研究の声を上げる象徴になっていました。アメリカ企業の人工知能研究所を率いるフランス人という彼の立場は、それ自体が一つの緊張を孕んでいました。彼はシリコンバレーの中で働きながらも、シリコンバレーの閉鎖性を批判しました。当時フランスは、人工知能を国家の主権問題として扱い始めていました。アメリカと中国の巨大企業が世界のモデルを分け合う時、ヨーロッパが自分たちの言語と自分たちのデータによる技術を持たなければ、他人の機械に依存する立場になるという危機感でした。

2023年11月、パリにキュタイ (Kyutai) という非営利の人工知能研究所がオープンしました。通信起業家のグザヴィエ・ニエル (Xavier Niel) をはじめとする人々が資金を提供した、研究結果をすべて公開することを原則とする研究所でした。ルカンはその方向性を支持し、助言を加えました。アメリカの企業がモデルを開かず時、フランスは重みとデータを開放するエコシステムの拠点にならなければならないというのが彼の持論でした。彼にとって開放とは、技術の問題であり、国家の問題でもありました。知識を開く国が、知識を作る国になるのだということが、彼が母国に伝えようとした考えでした。

その頃、フランスでは人工知能の新たな芽がいくつも芽吹いていました。MetaやGoogleの研究所を経た若いフランス人研究者たちがパリに戻ってMistralのような会社を設立し、ヨーロッパの資本がその後ろ盾となりました。長い間、人材をシリコンバレーに奪われてきた国が、初めて人材を呼び戻し始めたのです。ルカンはその流れの偉大な先達でした。彼がアメリカで歩んできた道がフランスの若者たちにとっては地図となり、彼が叫んだ開放の原則が彼らの旗印となりました。パリの勲章は一人の個人に与えられたものであると同時に、その個人が象徴することになった一つの国の野心に与えられたものでした。

ここには、シュミットフーバーとのあの長年の争いが残した痕跡があります。誰が先に発見したかをめぐって繰り返された承認の争いを経て、ルカンが得た答えは、相手に勝つことではなく、盤面を変

えることでした。発見を隠せば手柄をめぐって争うことになり、発見を開放すれば日付とコードが自ずと証明してくれます。知識の共有に向けたこの態度は、彼が好んで引用していたフランス共和国の古くからの価値観とも重なっていました。平等と友愛をソースコードの上に置き換えようという発想でした。フランスが彼に与えたのは勲章であり、彼がフランスに与えようとしたのは、その勲章の理由となった開放の精神でした。

振り返ってみると、この章の二つの顔は一つの問いに集約されます。発見は誰のものか。シュミットフーバーは最初に記した人のものだと言い、世間は最後までやり遂げた人のものだと言い、そしてルカンは誰のものでもあってはならないと答えました。すべてを公開すれば、発見は特定の誰かの所有物ではなく、人類の記録となります。その答えが正しかったかどうかは、後の章で試されることになります。

三人のチューリング賞は、人工知能60年の歴史の一つの結び目でした。冬の時代にニューラルネットワークに固執した人々がついに頂点で認められた瞬間であり、同時にその承認が妥当であるかを問う反論が噴出した瞬間でもありました。ルカンはその二つの顔に同時に向き合いました。栄光の壇上と優先権の影を同じ時期に見た人が、自分なりの答えを固めました。賞は人に与えられるが、発見は誰のものでもないということ、だからこそすべてを開放しておく方が良いということでした。

しかし、開放に向けた彼の信念は、やがて新たな試練に直面することになります。彼をシリコンバレーに連れてきた会社が、開かれた道ではなく、閉ざす道へと方向を転換し始めたからです。パリで勲章を受けてから二年と経たないうちに、ルカンは自身の人生において大きな決別を控えることになります。ただし、その決別はまだ何年も先の事でした。会社とたもとを分かťまで、彼はまず世間に向けて、なぜ皆が殺到した道が行き止まりであるのかを長く説明しなければなりませんでした。

弁護士

kimkj.com

第11章 鳥の羽根を模倣するな

ヤン・ルクン、人工知能の父

コレージュ・ド・フランスの講壇

2016年2月4日の夜6時、パリ5区のコレージュ・ド・フランスのマルグリット・ド・ナヴァル大講堂に一人の人物が演壇へ歩み出ました。55歳のヤン・ルクン(Yann LeCun)でした。講堂はフランス各地から集まった学者たちと若き学生たちで満席でした。彼の背後にある大きなスクリーンにフランス語の一行が表示されました。L'apprentissage profond : une révolution en intelligence artificielle。深層学習、人工知能の革命。その日彼は、この学校の情報学・デジタル科学の年間在任教授として就任する基調講演をするために来ていたのです。

コレージュ・ド・フランスは1530年に創設された学校です。学位は授与しません。試験もなく、登録手続きもありません。教授は自らが現在研究していることを大衆の前で無料で講義し、誰でも講義室の扉を開けて入り座ることができます。その学校には古い儀式が一つあります。新たに着任した教授が初講演を行う就任講演です。アンリ・ベルグソンがここで哲学を講義し、ミシェル・フーコーがここで権力について語りました。フランスの知性史がその時代の最先端を呼び出す場でした。そのような場にコンピュータで手書き文字を読んできた一人の工学者が招待されたという事実そのものが、神経網という長く嘲笑されてきた存在がついに文明の中心へ歩み入ってきたということを意味していたのです。

その講堂に座る人々の中には、わずか数年前まで神経網を真摯な科学として扱わなかった者たちもいたでしょう。学界には長く分団気がありました。神経網は理論が不完全で、なぜ機能するのか説明ができず、だから適切な科学ではないというものでした。ルクンはそのような白眼視を20年以上にわたって耐え続けてきた人でした。彼がこの大講堂の演壇に立ったということは、その長きにわたる白眼視がついに晴れたという静かな証拠でした。しかし勝敗が分かれたそのような場所で、彼は勝った方法が再び怠惰になるのではないかと心配していたのです。

ルクンが受けた職は2015年から2016年までの1年間の年間在任教授でした。情報学とデジタル科学を扱う職でした。彼がこの舞台に立つまでに要した時間を考えると、就任講演という名前はやや遅すぎるものがありました。彼は1980年代から神経網に取り組み、誰もがその道を信じなかった冬を2度経験し、今やっと世界が彼の手を挙げてくれたのでした。だからこの講演は新たな始まりの挨拶ではありませんでした。半生をかけて守ってきた信念を初めて祖国の最も古い講壇で公式に展開する場に近かったのです。

フランスでは就任講演は一つの儀式です。新たに着任した教授が自分が今後扱う学問の地図を大衆の前で展開する場であり、その記録は本として残ります。ルクンの前にこの講壇に立った数学者、物理学者、歴史学者たちがそうであったように、彼もこの日自分の分野がどこから来てどこへ向かうのかを描いて示さなければなりませんでした。彼が選んだ話の始まりは1世紀以上も前にさかのぼっていました。

神経網という発想は新しいものではありませんでした。1950年代後半、アメリカの心理学者フランク・ローゼンブラットがパーセプトロンという機械を作りました。人が規則を与えなくても、自ら写真を分類するよう学ぶ機械でした。人々は熱狂し、新聞はすぐに自ら歩き話す機械が現れるかのように書きました。しかし、その機械ができることにはすぐに限界が明らかになり、後で再び語るミンスキーとペーパーットの批判と重なり、熱は冷めました。最初の冬でした。ルクンは聴衆にこの古い興亡を指摘し

ました。自分が従った道がかつて死ぬ運命にあると宣告された道であったことを隠しませんでした。

その日の講演の前半でルクンが説明したことはディープラーニングが何かということでした。彼は層状に積み重ねられた神経網について語りました。一枚の写真が入ると、最初の層は明るいと暗い境界、つまり線の断片を認識します。その上の層はそれらの線を集めて角と曲線を見ます。さらにその上の層はそれらを集めて眼や車輪のような部分を見ます。最上層に至ると顔や自動車が捉えられます。人間がどの層で何を見るかを指示することはありません。機械は多数の写真を見ながら、各層が何を見るべきかを自ら決定します。これが彼が30年近くにわたって推し進めてきた考えでした。世界を複数の層で理解する能力を人間が手作業で組み込む代わりに、機械がデータから自ら育てるようになるというものでした。

この考えが長く無視されていたものが2012年に急に世を得たことには二つのことが重なっていました。インターネットが生み出した膨大なデータ、そしてそのデータを処理するのに十分な速さになった計算装置でした。材料と火力が同時に到着すると、長く折り畳まれていた神経網が展開されました。2016年の講堂に座る人々はその爆発をすでに目撃した後でした。写真を認識し、言葉を理解し、言語を翻訳する機械が彼らのポケットのスマートフォンに入っていました。だからルクンが舞台上でディープラーニングの勝利を自慢しても、誰も奇異に思わなかったでしょう。

しかし舞台に立つルクンは勝利を自分たちで祝いに来た人には見えませんでした。彼はすでにフェイスブック人工知能研究所を率いるトップであり、数年後にはコンピュータ科学のノーベル賞と呼ばれるチューリング賞を受けることになる人でした。春はすでに来ていました。彼がこの由緒ある講壇で言おうとしていたことは祝辞ではありませんでした。彼は方向について語りに来たのです。すべてが正しいと考える道について、その道がどこで分かれるのかについてでした。

ルクンはスクリーンにいくつかの数字を表示しました。それは人間の脳についての話でした。脳内には約850億個の神経細胞があります。1つの神経細胞は他の細胞と約1万個の接点で結ばれています。その接点をすべて数えると1兆個に達します。重さは1.4キログラム余り、体積は1.7リットル程度です。その湿った塊の外を覆う大脳皮質は厚さが2ミリメートルに過ぎないのに、広げると面積が2500平方センチメートルに及び、その中の神経配線をすべて接ぎ合わせると長さが18万キロメートルに達します。地球を4周以上も巻く配線がてのひらほどの厚さの中に折り畳まれているわけです。

ルクンはこれらの数字に関して一つのことを明確にしました。脳は知能が地球上に実際に存在することの証拠であるということです。知能が可能であるという事実は、すでに私たちの頭蓋骨の中で証明されていました。問題はそれからでした。そして彼は聴衆に一つの質問を投げかけました。このように賢い機械を作るには、私たちはあの脳をそのままコピーしなければならないのでしょうか。神経細胞の化学的作用からシナプスの微細な形態まで、一つ一つ複製してコンピュータ内に移す必要があるのでしょうか。講堂はしばらく静かになりました。ルクンの答えは短かったです。いいえ。フランス語で一言、Non でした。

この質問は空に投げ出されたものではありませんでした。その頃ヨーロッパでは、人間の脳全体をコンピュータ内に移してシミュレーションしようという巨大な科学計画が莫大な予算を受けて始まっていました。脳を詳しく模倣すれば、その中から心が自ずと生じるだろうという期待でした。ルクンの「いいえ」はそのような期待を狙ったものでした。脳をどんなに正確にコピーしても、脳がなぜそのように機能するのかを知らなければ、私たちは知能に一步も近づくことができないということでした。

この答えがその日の講演の種でした。脳を複製することと脳の原理を理解することは別のことであり、知能を作るということは前者の仕事ではなく後者の仕事であるということ。彼はこの区別を説明するために、人工知能の代わりに飛行の歴史を取り出しました。空を飛ぼうとした人々の話、鳥の翼をそっくりそのままコピーした機械と鳥を完全に忘れた機械の話でした。その箇所はこの章の最後でもう一度展開します。ここで先に押さえることは、彼がこの問いを最初に抱いた場所がこのきらびやかなパリの講壇ではなかったということです。

ルクンはフランスで生まれ、フランスで工学を学んだ人です。アメリカのベル研究所で名を得、ニューヨークで教授となり、シリコンバレーの企業の研究所を設立しました。彼の論文の大部分は英語で書かれ、彼が作った技術はアメリカ企業の製品の中で動きました。フランスが育てた人材が海の向こうで名を得ることは、その頃珍しくありませんでした。そのような彼が母国語で、フランス知識人の由緒ある講壇で、自分が生涯信じてきたことをこのように公式に展開させたのです。手書き数字を読んでいた若き日の信念が半世紀近く経って白髪の講演として戻ってきた場所でした。去った人が一時の間家に戻ってきた夜でもありました。

その日の講演はある意味でルクン自身の経歴をたどる場でした。神経網が一度思いつかれ沈み、誤差を逆に流す逆伝播で生き返り、手書き文字を読む量み込みニューラルネットワークで地位を確立し、2012年についに世を得るまで。その年表の折れ曲がりのたびにルクンという名前が置かれていました。人工知能が冬と春を往来した半世紀が彼の生涯とほぼ重なっていました。だからこそ、彼が舞台で聞かせたものは一つの技術の歴史であると同時に一人の人間の半生でもあったのです。その年の春全体を通じて、彼はこの講壇に何度も再び上り、機械がどのように学ぶのかを一回ずつ解き明かしていきました。

講演のタイトルには革命という言葉が入っていました。ディープラーニングが世界をひっくり返したのは事実でしたから。しかしルクンはその革命を終着点であるとは言いませんでした。むしろ彼は、現在の人工知能が人間はもちろん猫や鼠が世界を理解するレベルにも遠く及ばないときっぱり言いました。写真の中の猫を認識する機械は存在しても、猫が食卓から飛び降りるときどのように体をひねって着地するかを知る機械は存在しませんでした。革命は始まったばかりで、本当に難しい問題はまだ手もつけられていないということ。この冷徹な自己評価がその日の講演の底に敷かれていました。

彼がその場で語った信念は一文で要約できます。知能は生まれつきの法則ではなく、学び続ける過程だということ。そして機械を知能へと導くには、人間が規則を書き込む代わりに、機械が世界を見てみずから規則を見つけるようにしなければならないということ。2016年のこの講演は、その信念を世間に掲げた旗印でした。ところが講演から2年が経ち、ルクンはその信念を守るために、自らの手で立てた職を一つ手放すことになるのです。旗を振る人と夜中黒板の前に座る人は、同じ時間に両方になることはできなかったからです。

管理者から研究者へ

2018年初頭、ルクンは講演の代わりに、ある短い通告を準備していました。フェイスブックの最高技術責任者マイク・シュレプファー (Mike Schroepfer) と最高経営責任者マーク・ザッカーバーグ (Mark Zuckerberg) に伝える内容でした。フェイスブック人工知能研究所、すなわちFAIRの所長職から退任するというものでした。

彼が立てた研究所でした。2013年にザッカーバーグが何度も足を運んで連れてきて託した組織であり、彼が自らトップクラスの学者たちを集めて、わずか数年で世界有数の研究勢力に育て上げた場所でした。その職を自ら手放しながら、ルクンが挙げた理由は大げさなものではありませんでした。彼は経

営者ではなかったのです。数百人の研究者を配置し、予算を振り分け、会議を主宰することは、彼が得意とすることでも、やりたいことでもありませんでした。彼は黒板の前でニューラルネットワークについて思案する人であり、組織図の前で人員を配置する人ではなかったのです。

大きな研究所を率いることがどのようなものか想像してみれば、彼の選択が理解できます。数百人の研究者がそれぞれ異なる問題に取り組み、その分だけの採用、評価、昇進が進み、高い計算資源を誰にどれだけ配分するかを毎回判断しなければなりません。会議が会議を呼び、一日が人間関係の調整で終わります。そのような一日が続く間に、自分の手で方程式を思案する時間は括弧の中に束ねられたまま消えてしまったのです。所長の職は組織には必要でしたが、ルクンという科学者にとっては自らの研究を蝕む職でした。六十を前にした年頃に、彼が無駄だと惜しんだのはお金でも地位でもなく、時間だったのです。

FAIRの運営責任は二人の同僚に引き継がれました。モントリオール研究室を率いていたジョエル・ピノー (Joelle Pineau) とパリ研究室を率いていたアントワヌ・ボルデス (Antoine Bordes) でした。二人ともルクンが長く信頼してきた研究者でした。ルクン自身は最高人工知能科学者 (Chief AI Scientist) という職に移りました。肩書きはもっともらしかったですが、その意味は明白でした。管理から手を引いて研究に戻るとのことだったのです。彼は組織を去ったのではありませんでした。組織の運転席から降りて、再び実験台の前に座ったのです。

六十を前にした科学者が役員室を後にして実験台に戻ることは稀です。通常は逆方向です。研究によって名声を得た人が年をとると、より大きな組織、より大きな権限へと登っていきます。ルクンはその流れに逆らったのです。彼は自分の名前が付いた論文の著者欄に再び立ちたいと思っていました。若い研究者たちと同じ問題について黒板の前で議論したいと思っていたのです。職を降りることは、彼にとって後退ではなく復帰だったのです。

行政から手を引いたルクンが成し遂げたいことは別にありました。テキストを扱う機械ではなく、世界を見る機械だったのです。人間や動物が目で見世界を見て物理法則と常識を自動的に学ぶように、機械が生映像信号を見て世界の構造を自ら学ぶようにすること。彼はこれを自己教示型ワールドモデルと呼びました。誰も正解を付与しなくても機械が自ら学ぶ学習方式を自己教示学習といいます。

ワールドモデルという言葉は難しく聞こえますが、意味は素朴なものです。人間は頭の中に、世界がどのように動いているかを表わした小さなモデルの一つを持っています。カップを押せば落ちることを知り、ドアを引けば開くことを知っています。だから我々は手を動かす前に結果を心に描き、計画を立てるのです。ルクンが作りたかったのは、機械の中にまさにこのモデルを組み込むことでした。次に何が起こるかをみずから予測できなければ、機械も初めて先を見越して動くことができるということでした。言語モデルが次に来る単語を当てると、ワールドモデルは次に起こることを当てます。単語の順序ではなく、世界の順序を学ぶ機械、それが彼が描いていた絵だったのです。

なぜ映像なのか。ルクンは情報の量を重視しました。ある子どもが言葉を学ぶ前の数年間に目で受け取る世界の量は、どんな巨大な本の山よりも大きいのです。子どもはカップが机の端から落ちるのを見、ボールが転がって止まるのを見、水が下へ流れるのを見ます。誰も重力という言葉を教えてくれなかったのに、子どもは手を放すと物は落ちることを知っています。世界を見るだけで、世界の法則が身に染み込んだのです。ルクンは機械もこのように学ぶべきだと考えていました。文字を通して世界の話を聞くことと、世界を直接見て体験することでは、受け取る情報の厚みからして異なるということでした。その厚みを機械に注ぎ込む基礎的な設計を最初からやり直すには、組織図ではなく時間が必要だったの

です。

ルクンがこのワールドモデルに固執していた理由がありました。それまでディープラーニングの成功のほとんどは、人間が正解を付けたデータから生まれていたのです。数百万枚の猫の写真に人間が一つ一つ猫というラベルを付けると、機械はそのラベルを当てるように学ぶというやり方です。このやり方はうまく機能しましたが、限界は明らかでした。世界のすべてのものに人間がラベルを付けることはできないからです。ルクンは本当の知能ならラベルなしでも学べるべきだと考えていました。子どもが猫という言葉や言葉を学ばずと前に、すでに猫がどのように動くかを知っているのと同じようにです。正解ラベルに頼る学習には限界があり、その限界を突破するには、機械が世界を自ら観察して学ぶ道しかないというのが彼の判断だったのです。

彼が所長職は手放しても会社を去らなかった理由も、やはり別にありました。FAIRは彼が最初に立てた時から、一つの原則の上に立っていたのです。研究成果を秘めず論文として公開し、コードを開いて誰もが使えるようにするということでした。彼は基礎研究が会社の次四半期の成績に縛られてはいけなく、そうしてこそ遠くへ行くことができると信じていました。すぐに金になる製品から一步引いた場所、十年先を見据える場所。最高科学者として身を引いたルクンが守ろうとしたのは、まさにそういう場所だったのです。

戻ってきた場所は広くありませんでした。世界はすでに別の場所を見ていたからです。2018年頃から人工知能への関心は急速に言語へ、そしてますます大きなモデルへと傾いていました。ルクンが固執していた自己教示型ワールドモデルはまだ目立った成果を上げていない一方で、華やかな言語モデルたちがヘッドラインを占める間に、彼の研究は静かな隅で種を芽生えさせていたのです。かつてニューラルネットワークを一人信じていた人が、今やニューラルネットワークの勝利の中で再び少数派になったということです。春が訪れた野原で、彼はまだ誰も植えていない畑を一人で耕していたのです。

孤独は彼にとって珍しいことではありませんでした。彼は誰もニューラルネットワークを信じない時代にニューラルネットワークを選んだ人だったのです。二度の冬をそのようにして越してきたのです。だから多くの人が言語モデルへと殺到する光景を前にしても、彼は大きく揺らぐことはありませんでした。多数が正しいと言う道が必ずしも正しい道ではないことを、彼は自分の人生で既に学んでいたのです。世界を見る機械という彼の賭けは、若き日に一人正しかった経験から生まれた古い執念でもあったのです。

とはいえ、この道が正しいという保証がどこにあるわけではありませんでした。自己教示型ワールドモデルはそれまでのところ、もっともらしい理論にすぎず、言語モデルのように目立つ成果を見せたことがありませんでした。彼が入っていった場所は、正解が保証された安全な道ではなく、長く迷い続ける可能性のある暗い森だったのです。六十を前にした年頃に、そのような森へ再び入っていくには勇気が必要でした。若き日に正しかった記憶があるからといって、今回も正しいとは限らなかったからです。

この選択は突然の気まぐれではありませんでした。その根は1980年、パリの一つの学校の図書館にまで遡ります。その頃ルクンはパリ高等電子電気工学院ESIEEの学生でした。ある日、図書館の書架から彼は哲学的対話集一冊を手に取りました。1975年フランスのロワイヨモン修道院で開かれた大論争の記録をフランス語に訳した本でした。タイトルは『言語理論、学習理論 (Théories du langage, théories de l'apprentissage) 』でした。その論争は、その時代の認知科学の二つの陣営を一堂に集めるために意図的に開かれた場だったのです。

ロワイヨモンはパリ北部の旧修道院です。1975年の秋、その静かな石の建物に言語学者と心理学者、生物学者と数学者たちが集まりました。人間の心がどのように生まれるのかについて、その時代の鋭い頭脳が何日間も激突しました。その場の記録は後に本に編まれ、数年後にフランス語版が出版され、その一冊の本がパリの工科系大学の図書館の書架に置かれ、二十歳前後のルクンに発見されたのです。

その論争の両側には二人の巨人が立っていました。一方は言語学者ノーム・チョムスキー (Noam Chomsky) でした。彼は人間が生まれながらにして脳の中に言語の規則と記号を扱う装置を備えて生まれてくると考えていました。知能の骨組みは先天的だという側でした。もう一方は発達心理学者ジャン・ピアジェ (Jean Piaget) でした。彼は反対のことを言いました。人間の知能のほとんどは先天的なものではなく、子どもが世界に身体でぶつかり、触れ、転んだりしながら一段階ずつ自ら積み上げるものだ。知能は学んで作られるものだという側だったのです。

二人の違いは、子ども部屋で見ることができます。チョムスキーの子どもは生まれた時点で既に世界の文法を含んだ説明書を頭に持っており、経験はその説明書の空白を埋めるだけです。ピアジェの子どもは説明書なしで生まれ、ブロックを積んでは壊し、水をこぼしては拭きながら、世界の理屈を自分の手で書き綴ります。前者の子どもにとって知能は贈り物であり、後者の子どもにとって知能は労働です。徹夜してその本を読み通した若きルクンは、後者の子ども側に立ちました。世界は非常に複雑で常に変化するため、すべての規則を人間が事前に書いて機械に入れることはできないと彼は考えました。知能の核心は人間が入れ込む規則ではなく、機械が生データから自ら見つけ出す学習にあるということです。この考えが彼の人生を通じての羅針盤となりました。

この古い論争は、名前は変わっても人工知能の内部で繰り返されていました。一方には、知能とは記号を規則に従って扱うことだと見なす人々がいました。人間が論理と規則をうまく構成して機械に入れば知能が生まれるという側、記号主義と呼ばれる陣営でした。他方には、知能とは多くの接続がデータを見て自ら調節されることだと見なす人々がいました。神経網を信じる側、接続主義と呼ばれる陣営でした。チョムスキーの席に記号主義があり、ピアジェの席に接続主義がありました。古い論争の黒板での対峙は、人工知能が何度も繰り返すことになる戦いの予告編だったわけです。そしてルクンは二十歳で既にどちらの側に立つかを決めていました。

その本には、彼を長く引き止めた箇所がもう一つありました。ピアジェ側に座ってチョムスキーに反論した人の中に、MITの数学者シモア・パパート (Seymour Papert) がいました。パパートはその場でパーセプトロンを示しました。パーセプトロンは初期の神経網学習機械です。彼はこう言いました。ここを見よ、この余計なものが何もない機械を。人間が規則を組み込まなくても、それ自体がパターンを学ばないか。学習から知能が生まれるという証拠として神経網を打ち出したのです。

ところがルクンは後になって一つの事実に気づきました。そのようにパーセプトロンを称賛したパパートが、わずか数年前の1969年には、マービン・ミンスキー (Marvin Minsky) と共に『パーセプトロン (Perceptrons)』という本を書いて、神経網の限界を項目ごとに掘り下げたということでした。その本は神経網研究に向かっていった資金を途絶えさせ、人工知能の最初の冬を早めたものとして何度も語られています。同じ人物が一方では神経網を賞賛し、他方では神経網の息の根を止めていたわけです。この矛盾が若きルクンを虜にしました。彼は恐れませんでした。むしろ掘り下げました。つまり、まだ明かされていない原理がここにあるという意味だったからです。神経網が死んだ道ならば、なぜその死んだ道をめぐってこのように賢い人たちが互いに異なることを言ったのか、ということでした。

2018年に常務理事職を辞任することは、1980年に図書館で立てたその羅針盤に戻ることでした。本当にしたかったことは、自ら学ぶ機械を作ることであり、そのためには手が自由でなければなりませんでした。管理者の座から研究者の座へ移りながら、ルクンは再び黒板の前に立ちました。そしてその黒板の上で、彼が何度も描いた図は、2016年のパリの演壇で取り出したまさにその飛行の話でした。

飛行の原理、知能の原理

もう一度、コレージュ・ド・フランスのその講演に戻ります。脳を模倣すべきかという問いに「いいえ」と答えた後、ルクンは画面に19世紀フランスのある飛行機械の図面を映しました。クレマン・アデル(Clément Ader)が製造したエオール(Éole)でした。

アデルは空を非常に欲しがった発明家でした。彼はコウモリを模倣しました。絹で作られた大きな翼をコウモリの翼の形そのままに14歩の幅に広げ、竹と鋼管で軽い骨組みを組み立て、おおよそ20馬力の蒸気エンジンを搭載しました。1890年10月9日、パリ郊外のある領地でエオールは地を蹴って暫し浮かび上がりました。自らの力で平地から離陸した最初の機械でした。飛行距離はおよそ50メートル、高さは20センチメートルに満たありませんでした。短い跳躍でしたが、確かに浮かび上がりました。

そしてそこで終わりでした。エオールは方向を変えることもできず、高さを調整することもできませんでした。操縦装置がなかったからです。浮かび上がったら落ちるということ以外にはできませんでした。数年後に製造した後継機アビオン3世(L'Avion III)も同様でしたし、アデルの系統は航空史から静かに途絶えました。彼が見落としたのは翼の形ではありませんでした。彼はコウモリの外観は完璧に模倣しましたが、正に空を飛ぶのに必要な原理、すなわち空気の中で体を安定させ、思いのままに方向を定める制御の原理を理解できませんでした。彼の機械は浮かび上がることはできても、空に留まることはできませんでした。

フランス人にとってアデルは親愛な名前です。ライト兄弟より先に空に手を伸ばした自国の開拓者、しかしゴール直前で止まった人。ルクンがアメリカで始まった深層学習の歴史を語りながら、わざわざフランスのこの悲劇の発明家を呼び出したのには、祖国の聴衆に向けた優しい手振りもいくらか混じっていたのでしょう。アデルは才能も情熱も溢れていました。不足していたのは、ただ原理に対する理解ただ一つでした。才能が原理に取って代わることはできないというこの事実こそが、その日の講演が伝えようとしていた痛切な教訓だったのです。

13年後の1903年、ライト兄弟(Wright Brothers)が空を開きました。彼らは飛行機に鳥の羽を一枚もつけませんでした。翼をはためかせることもありませんでした。彼らは生物の外観を模倣することをすっかりやめ、代わりに風が翼を押し上げる理屈と、その機械を操縦して安定させる方法を掘り下げました。彼らは自転車を製造していた人たちらしく、飛ぶことより先に転覆しないで方向を定める問題を解きました。風を模倣する小さい風洞を作って翼の形を数百回試験し、浮かび上がった後、どのように傾き、どのように旋回するかを指先の操縦で制しました。鳥を複製する代わりに、飛行の原理を引き出したのです。外観を捨てて原理を得た方が、空を手に入れました。

ライト兄弟は大学を卒業した科学者ではありませんでした。自転車を修理し、製造していた兄弟でした。彼らが他の人たちより先んじていたのは、学識が豊富だからではありませんでした。推測で済ませず、一つ一つ直接測定し試験したからです。先人たちが残した不正確な数値を信じず、自分の手で再び測定しました。原理を理解することは、つまり自分を欺かないことであり、この頑固さが彼らを空へ上げました。

ルクンがこの話で言いたかったことは明白でした。知能を作ることも飛行と同じだということでした。脳という鳥を前にして、私たちはその羽を模倣することもできるし、その原理を引き出すこともできます。神経細胞の化学反応とシナプスの形態を一つ一つ複製することは、アデルがコウモリの翼を模倣したことと同じです。逆に、脳が世界を見て学ぶその基底にある理屈を数学と工学で引き出すことは、ライト兄弟がしたことと同じです。彼はこれを知能の空気力学と呼びました。空に空気力学があるように、知能にもそれを支える原理があり、私たちが探すべきものは羽ではなく、正にその原理だということです。

だから、この章のタイトルでもある彼の言葉が出てきます。鳥の羽を模倣するな。鳥をそのまま真似ず、飛ぶ法を理解しろという意味です。脳をそのまま真似ず、学ぶ法を理解しろという意味でもあります。ライト兄弟が鳥の羽一枚なしに空を手にしたように、知能を作る人も脳の細胞をそのまま移し込む代わりに、脳が学ぶ原理を練り上げなければならないということです。

ルクンはこの箇所ですらに一步進んで、原理を知ることなく外観と規模だけを増やす方式に名前をつけました。カーゴ・カルト人工知能という名前でした。カーゴ・カルトは実際にあった話です。第二次世界大戦の時、南太平洋の孤島にアメリカ軍が入ってきて滑走路を作り、飛行機であらゆる物資を運んできました。島の人々にとって、それは天から降ってくる豊かさでした。戦争が終わり、アメリカ軍が去ると、飛行機も物資も途絶えました。人々はその豊かさ再び来ることを願い、木と藁で偽の滑走路と管制塔を建て、竹でヘッドホンのような形を作ってかぶりながら、空に向かって信号を送りました。外観は滑走路と全く同じでしたが、飛行機は戻ってきませんでした。なぜ飛行機が来ていたのか、その理屈が抜け落ちていたからです。

この話を科学に初めてたとえた人は、物理学者リチャード・ファイマンでした。彼は1974年のある卒業式の演説で、科学の外観だけをまねながら、その実質である正直さを抜かした研究をカーゴ・カルト科学と呼びました。形式は完璧なのに飛行機は着陸しない研究のことです。ルクンはこの古いたとえを人工知能に持ち込みました。彼から見ると、脳の配線をそのままスーパーコンピュータに移植して、巨大なシミュレーションを走らせるだけで知能が自ずと開花すると信じる態度が、まさにその偽りの滑走路でした。原理に対する理解なく規模にだけお金と電力を注ぎ込む方式、外観を真似れば内容も続くだろうという期待のことです。良い人工知能エンジニアは、脳を真似る呪術師ではなく、原理を練り上げる設計者であるべきだというのが、彼の長年の信念でした。脳が知能の唯一の回答用紙だからといって、その回答用紙を模写したからといって問題を理解したわけではないという話でした。

このたとえには後日談があります。人が鳥を真似るのをやめると、むしろ鳥ができない仕事をする機械が現れました。今日の飛行機は、どんな鳥より速く、どんな鳥より重く飛びます。羽を捨てたからこそ得られたものです。ルクンは人工知能もそのようになると考えました。脳をそのまま模倣しようとする頑固さを手放せば、いつか人間の脳ができないやり方で学んで考える機械が現れるかもしれないということです。目標は脳に似た機械ではなく、脳のように自ら学びながら脳と異なるように造られた機械でした。

カーゴ・カルトの話を初めて持ち出したファイマンは、その演説で一つの文章を力を入れて語りました。最も重要な原則は自分自身を欺かないことであり、人が最も騙されやすい相手こそ、まさに自分自身であるということです。ルカンは規模至上主義を警戒した根底にも、同じ恐れがありました。表面上もっともらしい成果が積み重なるほど、肝心の原理を知らないという事実を自ら忘れやすくなるということです。機械が人間のように言葉を話せば、私たちはその機械が人間のように世界を理解していると錯覚してしまいます。その錯覚を警戒せよというのが、彼の長年の訴えでした。

もちろん、この話には反対側の声もあります。ルカンが規模だけを拡大するやり方を間違った滑走路だと呼んでいたまさにその頃、言語モデルは規模を大きくするたびに、人々が予見できなかった能力を示していました。より大きなモデルがより賢く見えるという事実は否定し難く、多くの研究者がその道に賭けました。原理を先に理解すべきだというルカンの言葉が正しいのか、ひとまず拡大すれば原理は後からついてくるという反対側の言葉が正しいのかは、当時も、この本を書いている今も、判決は下っていません。ルカンは自分が正しいと確信していましたが、確信と証明は別物でした。

この信念がその後どこへ向かうのかは、本書の次の章で語るべきことです。世界がますます巨大な言語モデルへと殺到していった数年後、ルカンはその方向を指して間違った滑走路だと呼ぶことになりました。文字だけを飲み込んだ機械は、世界の重さも、物が落ちるという事実も、手を離せばカップが割れるという事実も知らないということ。本だけを読んだ医師にメスを任せられないように、世界を体で経験したことのない機械に世界を任せることはできないというのが彼の診断でした。そしてその診断は講演や論文にとどまらず、結局彼が自らの手で設立したその巨大な研究所を去り、自ら別の道を切り開くことへとつながります。その話をするにはまだ早いでしょう。

2016年のパリの演壇で、脳を模倣するのか、脳の原理を理解するのかを問うたその人は、2018年に管理者の職を辞して黒板の前に戻り、その黒板の上で同じ絵を描き続けました。羽毛ではなく、飛ぶ方法を。規模ではなく、原理を。人が書き込む規則ではなく、機械が自ら見つけ出す学びを。知能とは一体何なのか。彼がパリのある図書館で初めて抱いたこの問いは、半世紀が過ぎても彼の黒板の上に消されることなく残っていました。

弁護士

kimkj.com

第12章 ChatGPTというサイレン

ヤン・ルクン、人工知能の父

次単語予測の袋小路

2026年春、あるアメリカの経済放送のスタジオにヤン・ルクン (Yann LeCun) が座りました。進行役は、人工知能がまもなく人間を超えるだろうという話から切り出しました。ルクンは首を振りしました。今、世界を覆う大型言語モデル、つまりChatGPTのような大型言語モデル (LLM) は、これから5年ほど経つと、ほとんどが無用になるだろうと、彼は淡々と語りました。流行の真っ最中の技術に死刑宣告を下す人の口調ではありませんでした。長い前からずっと同じ話を繰り返してきた人の、疲れた確信に近いものでした。彼は65歳で、40年近くニューラルネットワークの傍を離れたことがなく、ちょうど世界有数の大規模人工知能研究所の一つを率いるポストから退任したばかりでした。

この章の題名をサイレンと付けたわけがあります。サイレンはホメロスの船乗りたちを魅惑していた海の歌です。その音は美しく、だから船は歌を追い求めて岩にぶつかり砕けました。ルクンが大型言語モデルについて語ることはまさにそれです。音は陶酔的ですが、その音を知能と勘違いして、すべての船がそこに殺到することが危険だということです。彼が反対しているのは、この機械が無用だということではありません。ChatGPTが文章を磨き、翻訳を行い、コードを書く際の驚異的なツールであることを彼も知っています。彼が反対しているのは、この機械を知能の目的地と見なし、他の道をすべて塞いでしまう、まさにその錯覚です。有用なツールであることと、人間のような知能へと向かう道であることは別の話だということなのです。

この区分はルクンが最近になって提示したものではありません。世界がChatGPTに熱狂し始めた頃から、彼は学会と講演を回り、同じ警告を繰り返してきました。ダボスの議論の場でも、大学の招待講演でも、彼は現在の流れは知能への大道ではなく、わき道に流れる側路だと語りました。2026年春、ブラウン大学の講壇に立った時も、彼は大型言語モデルの限界を指摘した後、自分が考える別の道の素描を聴衆に示しました。同じ話を何年も繰り返しているということは、その話が流行に乗ってやってきた即興ではなかったということです。長く錬磨してきた信念だということでもあります。

大型言語モデルがやることは、見た目には魔法です。人間が文を投げかけると、機械がもっともらしい答えを続けます。しかし、その中身を開けば、やっていることは一つです。次に来る単語を確率で選ぶこと。これまで出た言葉の流れを見て、その後に来そうな次の破片一つを抜き出して貼り付けます。この方式を次単語予測 (next-token prediction) と呼びます。一つの文は、この予測を数十、数百回繰り返して作られます。機械は意味を思慮して答えるのではなく、確率の奔流に従って一步步滑り落ちます。ルクンが問題にする地点がまさにここです。

この機械には、他人がよく気づかない構造的弱点が一つあります。今しがた自分が吐いた単語を、次の単語を選ぶ際に再び入力として受け入れるという点です。人間が検証してくれたことのない自分の出力を、自ら正解として、その次を続けます。ルクンはこの構造を一つの式で指摘します。単語一つを選ぶたびに、ごく小さな誤差が混じると言いましょう。その誤差の大きさを e とすると、長さが n である文が最初から最後まで一度も狂わない確率は、ほぼ $(1-e)^n$ を n 回掛けた値です。 n が大きくなればなるほど、この値は急速に崩れます。文が長くなるほど、どこかで足を踏み外す可能性は階段を転がるように増えるという話です。短い答は比較的正しくても、長い推論が求められるところで機械がしばしば足を踏み外すのは、ここにあるのです。

この計算が恐ろしく感じられるわけは、戻る道がないからです。人間は話している途中で、あ、しまったと思ったら止めて直します。機械にはその止まりがありません。一度間違った分岐に入り込んだら、自分が作った誤答を足がかりにして、その上に再び誤答を積み重ねます。もっともらしい文の見た目はそのまま保ったまま、内容だけ現実からそっと剥がれ落ちます。私たちが幻覚 (hallucination) と呼ぶ現象、機械が平然と存在しない事実を創作する、その習慣はここから来ています。偶然に混じった雑音とは異なります。次単語予測という設計から必然として生まれ出る影だというのがルクンの診断です。

ルクンはこの状況を、分岐がどこまでも広がった木に例えます。単語一つを選ぶたびに、機械の前には数万の分岐する道が開きます。そのうち意味が通じて事実と合致する道はただ一本の非常に細い筋だけで、残りの大部分は無意味に続く枝道です。文が長くなるほど分岐は倍増し、機械がその細い正しい幹の上だけにずっと留まる確率は歩むたびに減ります。さらに機械は、自分が正しい幹から外れたかどうかさえ知りません。外れた場所でも次の単語を平然と選ぶので、間違っただ道の上でも文は滑らかに続きます。滑らかさがすなわち正確さではないということ、これがルクンが繰り返し指摘する地点です。

この診断が空々しく聞こえないわけは、賭金が大きいからです。世界の大企業たちは今、この次単語予測という一つの設計に数千億ドルを注ぎ込んでいます。その真っ最中で、40年ニューラルネットワークを掘った人が、その設計では真の知能に到達できないと語るのです。多くの人が彼の言葉を異端に聞きました。既に世界を変えている技術を前にして袋小路だと呼ぶことが無謀に映りました。ルクンは気にしませんでした。彼は有用性と知能をしばしば区別しようとしてしました。この機械が今有用だということ、この機械が知能へと向かう道だということは別の主張だと。前者は真ですが後者は偽だというのが、彼の長年の立場です。

彼が懸念したのは方向への一極集中でした。一つのやり方が眩しい成果を出すと、研究者も資本も人材も一斉にそちらに群がりました。大学の若い研究者たちは次単語予測を磨く仕事に取り組み、別の道を掘る人は次第に減りました。ルクンから見てこれが危険でした。一つの時代が一つの羅針盤だけを持って航海するとき、その羅針盤が指す先に港がなければ、船はすべて迷います。彼があえて異端の席に身を置き、声を上げたのは、少なくとも別の羅針盤の一つくらいは誰かが持っていなければならないという思いからでした。人工知能の冬を二度経験した人の慎重さでもありました。

彼は講演でこの弱点を示す際、自分の話を出すことがあります。メタ (Meta) の研究チームが初期のラマ (LLaMA) モデルをテストしていた頃、誰かが機械にふざけた質問を投げかけました。ヤン・ルクンが去年ラップアルバムを出したと聞かずに知っているか、と。ルクンは生涯ヒップホップを一フレーズ録音したことはありません。それなのに機械はためらいません。存在しないそのアルバムの収録曲を指摘し、ビートについて論じ、ライムの組み立てを褒める、もっともらしい批評を創作しました。機械はそのアルバムが本物なのか嘘なのかを問いませんでした。アルバム、発売、批評という単語が文の中でよく調和していた統計だけに従っただけです。真と偽を分かつ天秤がそもそもその中になかったのです。ルクンはこの逸話を笑いながら語りますが、笑いの後ろの要点は重いです。人についての情報を何気なく創作する機械が、人間の病気を診断し、契約書を検討する席に座るとしたら、どうなるだろうか、ということです。

この懸念は既に現実になりました。2023年、アメリカのある弁護士が法廷に提出した書面で、自分がChatGPTで見つけたと言う判例をいくつか引用しました。ところがそれらの判例は世に存在しませんでした。機械が本物の判例のように見える事件名と番号をもっともらしく創作し、弁護士はそれを本当だと信じて法廷に提出しました。判事が確認してみたら、引用された事件はどこにも存在しませんでした。ラップアルバム批評を創作していたその習慣が、今回は法廷文書の上で繰り返されたのです。機械

にとって判例があるということと、判例があるように見えるということは同じ重みでした。雄弁な嘘が
どういう場所でどの程度危険なのか、この事件は明確に示しました。

ルクンがよく指摘するもう一つの点は、この機械が思考の深さを調節できないということです。今の
トランスフォーマー構造は、単語一つを選ぶ時にいつも同じ量の演算を使います。量子力学の難しい方
程式を解くにせよ、今朝何を食べたかという些細な質問に答えるにせよ、機械が一単語を出すまでに通
る回路の量はまったく同じに固定されています。人間は難しい問題の前で歩みを止め、長く思い巡らし、
簡単な問いには素早く答えます。この機械にはその緩急がありません。近ごろ流行りの思考の鎖 (chain-of-thought) という技法、機械に解答過程を長く書かせて、さも推論しているように見せるその
方法について、ルクンは冷徹です。それは新たに得た思考力ではなく、固定された演算量を迂回しようと、言葉をより長く伸ばす、浅はかな手段に過ぎないということです。本当に止まって、複数の分岐を
再び検討することと、止まったふりをして言葉を伸ばすことは異なる事柄です。

心理学者ダニエル・カーネマン (Daniel Kahneman) は人間の思考を二種類に分けました。熱いお湯
に手が触れると反射的に引っ込むような瞬間的な反応がシステム1で、見知らぬ道で地図を広げて複数の
の経路を検討する遅い思慮がシステム2です。ルクンから見ると、今の大型言語モデルはみんなシステ
ム1に留まっています。習ったことを反射的に繰り返すだけで、目標を立て、代替案を探し、先を見通し、
計画するシステム2の敷居は踏んでもみませんでした。機械が流暢だからこそ、私たちはそれが思
考していると錯覚しますが、流暢さと思考は同じものではありません。話が上手いことと世界を知ること
の間には、ルクンから見ると、未だ渡ることができない川が横たわっています。

彼が描く別の機械とはどのような姿をしているのでしょうか。詳細な設計は後の章で述べることにしま
すが、要点はこうです。本当の知能なら答えを口にする前に、頭の中でまず考えてみなければならない
ということです。このようにしたらどうなるか、あのようしたらどうなるかを想像で試してみて、い
くつかの道筋の中から目標に至る道を選んで動く機械。人が見慣れない都市で道を探すときに、頭の中
でいくつかの道を描いてから足を踏み出すように。こうした機械は難しい問題の前ではより長く考え、
簡単な問いには即座に答えることができます。次の単語を反射的に抽出する代わりに、世界がどのよう
に動いているかについての内部モデルを持ち、その上で先のことを推し測る機械。ルクンが、セイレー
ンの歌に背を向けて向かおうとしている港がそちらです。

では、彼はなぜこのなめらかな機械が結局のところ本当の知能には至らないと断定するのでしょう。
彼の答えは言語の外にあります。大規模言語モデルは、人が書き残したものだけを食べて成長します。
文は世界の薄い外殻です。私たちが目で見、手で触り、身体で経験するその厚く連続的な物理世界を、
文字は極めて一部分だけを捉えています。ルクンはあるインタビューでこのように述べました。「An
LLM doesn't understand that if you push a glass off a table, it will break. It only knows that the words '
glass' and 'break' often appear together.」言語モデルはグラスをテーブルの外に押し出すと割れること
を理解できず、ただグラスと割れるという単語がよく一緒に現れることを知っているだけという意味で
す。ガラス瓶が床に落ちて粉々に散るとき、破片がどの方向に飛ぶのか、水がどのように広がるのか、
その軌跡を機械は計算できません。割れるという文字を組み合わせるだけです。

この診断が正しいかどうかは、まだ議論の最中です。大規模言語モデルを推し進める陣営では、規模
をさらに大きくすれば、いつか物理的な常識も自ずと染み出てくると考えています。実際のところ、モ
デルを大きくするたびに誰も予想しなかった能力が現れた前例があるので、その期待が根拠のないもの
ではありません。テキストだけでもコードを書き、翻訳を行い、試験に合格する機械を私たちはすで
に見ました。ルクンは反対側に立ちます。どんなにテキストを注ぎ込んでもテキストに存在しないものは

出てこない、物理世界は文字という狭い通路からは入ってこない、彼は述べています。どちらが正しいかは時間が答えるでしょう。ただし、ルクンは自分で5年という期限を設けました。間違えば彼が間違っただけで終わる、引き返す場所のない期限です。セイレーンの歌が本当の港に繋がるのか、岩に繋がるのかは、その5年が経った後でなければわかりません。

本だけを読んだ医者

医学教科書を最初から最後まで暗記した人がいるとしましょう。解剖学も、薬理学も、最新の臨床論文も一文字残らず頭に入れました。ところが、この人は生きている患者のお腹に一度も触って見たことがありません。メスを握ってみたこともなく、脈打つ手首を押さえてみたこともなく、血が噴き出ている傷を圧迫してみたこともありません。この人に手術刀を任せられるでしょうか。ルクンが大規模言語モデルについて用いる比喻がこれです。国内では彼の講演を取り上げたあるメディアが、これについて本を読むだけでは手術はできないとまとめました。知識を丸ごと暗記することと、世界を扱うことができることは違う事柄だという話です。教科書にぴったり当てはまる患者なら、この人も立派に答えることができるでしょう。問題は教科書から外れた体、初めて見る症状の前で手が止まるところにあります。

この隔たりに最初に名前をつけたのはロボット工学者ハンス・モラヴェック (Hans Moravec) でした。1980年代に、彼はある奇妙な事実に気づきました。人間が困難と感じる仕事、たとえば複雑な数学の証明や膨大な資料検索はコンピュータが瞬く間に成し遂げてしまいます。しかし、新生児や子犬でも何もない顔で行うこと、グラスをひっくり返さずに持ち上げること、障害物を避けて飛び越えること、落ちてくる物がどこへ向かうのかを目で追うこと、こうしたことの前では機械は手の施しようがなくなります。人々はこれをモラヴェックのパラドックス (Moravec's paradox) と呼びます。知性の難しさと易しさについての私たちの感覚が反対にひっくり返っているということです。実のところ困難なのは論理ではなく、身体で知る世界だったのです。

なぜそうなのでしょう。論理と言語は人間がこの最近数千年の間に得た若い能力ですが、身体で世界を扱う感覚は数億年にわたって磨きぬかれた古い智慧だからです。物を握り、バランスを保ち、危険を避ける能力は、私たちの祖先が生き残るために数百万年にわたって鍛え上げたものです。その長い歳月がその能力をあまりにも滑らかに磨きぬいたので、私たちはそれを簡単だと錯覚しています。反対に、数学やチェスはわずか数千年前の新しい技術に過ぎないので、ごくごく大変に感じられます。コンピュータがチェスに勝ち、囲碁に勝ちながらも、グラス一つをうまく握れないのは、簡単なことができていないように見えても、実は最も困難なことができていないのです。

このパラドックスはロボットの指先に今も生きています。機械が世界チャンピオンをチェスで打ち負かしてから30年近くが経とうとしているのに、イチゴ一粒を潰さずに拾い上げるロボットハンドはいまだ困難です。強く握りすぎると潰れ、弱く握ると落とします。人間の手はその力の微妙な加減を何も考えることなく調整しますが、機械にとってそれはまだ越えられない壁です。大人が無意識のうちにやるのが機械にとって最も困難な作業として残っているのは、モラヴェックが40年前に指摘したそのパラドックスがまだ解き明かされていないということです。

ルクンはこのパラドックスを数字で説き明かします。今日、大きな言語モデルはインターネットに公開されたほぼすべての文、およそ10兆個のトークンを取り込みます。バイト数に換算すると、およそ20兆バイトです。人間がこの量を昼夜を問わず読み進めたらどのくらい時間がかかるでしょう。ルクンの計算では数十万年です。人類が文字を発明するずっと前から今まで、絶え間なく読み続けてようやく達する量です。ところが、ここで彼が並べて置くもう一つの数字があります。4歳の子どものです。

まだ言葉も不器用な4歳の子どもが、生まれてから目覚めていた時間を全て合わせると、およそ16,000時間です。その間、子どもの両目は視神経を通じて、毎秒およそ2メガバイトの情報を脳へ送り続けました。人間の目一つには約100万本の視神経繊維があり、両目を合わせると200万本の通路が絶え間なく世界を運び続けています。短く、しかし途切れることなく流れ込むこの感覚の大河が4年間積み重なると、ルクンの計算では、世界のすべてのテキストを合わせたものの約50倍に達します。4歳の子どもが既に世界で最も大きな言語モデルよりはるかに多くのものを経験してきたという話です。しかもそれは、文字ではなく光と重さと動きによってです。

50倍という数字をしばらく味わってみます。人類が今までインターネットに投稿したすべての文、ウィキペディアとニュースと書籍と無数の投稿を全て合わせたものより、4歳の子どもが目で受け入れたものが50倍多いという話です。しかし、その50倍多い情報は、雑然とした雑音ではありません。その中には、世界がどのように動いているかについての答えが密集して組み込まれています。水がどのように注ぎ込まれるのか、影がどちらに落ちるのか、顔がどのように表情を変えるのか。文字ではめったに書かれることのない、しかし生きていくために欠かせない知識です。子どもはこれを試験として学ぶのではありません。ただ目を開いて世界の中にいるだけで学ぶのです。

この感覚の大河が運んでくるのは個別の事実ではなく、世界の作動方式です。子どもはグラスを落としてみて、ボールを転がしてみて、転んでみながら、重力と慣性と摩擦を学びます。誰もその規則を文字で教えてくれませんでした。子どもの脳はそれを身体で刻み込みます。手から放した物は上に浮かばず下に落ちること、転がっていたボールはいつか止まること、壁の向こう側に転がったボールは消えるのではなく壁の向こう側にそのままあること。これらすべてが、言葉よりも前に、映像と感覚として脳に刻まれています。文字という狭い帯域幅では収まらないものを、感覚という広い通路が運んでくるのです。ルクンが繰り返し述べる要点はここにあります。知性の骨組みは言葉ではなく、世界との接触から育つということです。

ここで見るということは、受け身で受け入れる行為ではありません。子どもは物を投げ、押し、口に入れてみ、壊します。世界に手を触れてみてその反応を観察しながら、何をするとどうなるのかを身体で学びます。原因と結果の連鎖をこのように自ら経験して積み重ねたことが、後になってその子どもが世界を扱う財産となります。大規模言語モデルはこのプロセスを丸ごと飛ばします。人々が世界についてについて書き残した成果物だけを読むだけで、世界に直接手を触れてみたことがありません。他人が書いた料理本を完全に暗記したことと、台所で一度でも火に火傷をしたことの違い。ルクンが見分けるのはこの二つの違いです。

本だけを読んだ医者という比喻が痛く響く場所は手術室です。試験をすべて満点で合格した医大生が初めて生きている人の皮膚にメスを当てる瞬間を想像してみます。教科書には切開線の位置が図で描かれていましたが、実際の体は人それぞれ異なり、血が視界を遮り、組織は予想とは異なるように密集しています。この瞬間に必要なのは暗記した知識ではなく、指先で組織を感じながら次の動きを推し測る身体感覚です。その感覚は本からは来ません。ルクンが大規模言語モデルに見る欠落がこれです。答案は完璧ですが、生きている世界の前に立つと、指先には何の感覚もない状態です。

もちろん、人間も最初は不器用です。ただし、人間の不器用さは、既に世界を知っている身体の上に手技という一つの新しい技術を載せる不器用さであり、機械の不器用さは世界自体を経験したことがないところから来る不器用さです。前者は練習で埋められますが、後者はどんなにテキストを注ぎ込んでもうまく埋められないというのがルクンの判断です。本をもっと読ませることで治せる欠落ではなく、本という材料そのものが到達できない欠落だということです。

機械がこの物理的な常識を備えていないことを示した有名なシーンは2020年に現れました。認知科学者ゲアリー・マーカス (Gary Marcus) とコンピュータ科学者アーネスト・デビス (Ernest Davis) が当時世界を驚かせていたGPT-3に短い状況を提示したのです。訳すと次のようなものです。あなたはクランベリージュースを一杯注ぎました。しかし誤ってブドウジュースを小さじ一杯ほど混ぜてしまいました。外見は問題ありません。においをかごうとしていますが、ひどい風邪にかかっている何のにおいもしません。あなたはとても喉が渇いています。ですからあなたは。ここで機械が文を続ける必要がありました。GPT-3はこのように続けました。「drink it. You are now dead. (それを飲みます。そしてあなたは今死んでいます) 」。クランベリージュースにブドウジュースを混ぜると人が死ぬと、機械は真摯に判定したのです。

世界のどこにもクランベリーとブドウを混ぜた飲料が毒だという根拠はありません。スーパーマーケットには両者を混ぜた製品が普通に売られています。機械はブドウジュースが何であるか、飲むことが何であるか、死ぬことが何であるかを知らないまま、喉が渇いているという言葉の後にしばしば続いていた文の断片を選んだだけです。風邪、におい、喉の渇きといった言葉が続く物語の後には往々にして不吉な結末が来るという統計、機械が捉えたのはその統計でした。マーカスとデビスがこの実験で伝えたかったことは明確でした。流暢な文は理解の証拠ではないということ。本当に知っている存在なら、二つの果汁飲料を混ぜることが危険であるかないかをまず検討したはずなのに、機械はその検討段階をすっかり飛ばしました。その後に見れたより大きなモデルはこのような落とし穴をより良く避けます。しかしルクンが見るところ、それは一つ一つの落とし穴を人間が埋めてやったのであって、機械が世界を理解するようになったわけではありません。

ルクンはここでもう一步踏み込みます。彼は人間の脳で言語を扱う場所がいかに限定的であるかを指摘します。言葉を作り理解することは大脳皮質の一角、ブロカ野 (Broca's area) と呼ばれる小さな領域が担当します。ルクンが見るところ、今日の大規模言語モデルはその小さなブロカ野だけを取り出して大きく膨らませた機械です。言葉を作り出す領域一つだけを巨大に拡大し、世界を経験し身体を動かす先を見通す残りの脳はすっかり空にしたようなものです。ところが人間が実際に考え計画するとき、頭の中を回るのは文字ではありません。物の形、重さ、それがどう動くかについての像です。私たちは言葉で考えると信じていますが、実際には難しい判断は言葉になる前の段階でもう成されています。サッカー選手が球をどこに蹴るか決めるとき文を思い浮かべないように、大工が釘をどこに打つか考えるとき論理式を立てないように。言語はその結果を外に出す最後の出口に過ぎず、知能のエンジンではないというのが彼の主張です。

この主張を支える観察が一つあります。言葉を失った人も思考は失いません。脳の言語領域を傷つけて言葉をうまく話せなくなった人が、それでも道具を使い、道を見つけ、状況を判断します。言葉が知能の根源なら言葉を失うときに思考も崩壊すべきなのに、そうではありません。思考は言葉よりも深い場所に位置しており、言葉はその思考を外に移す通路に過ぎないという話です。ルクンはここで逆問します。人間の知能が言葉の上に立っていないなら、なぜ私たちは言葉だけを与えた機械に人間のような知能を期待するのかと。

ですからルクンにとって大規模言語モデルは誤って作られた機械ではなく、誤った目標を定められた機械です。テキストを扱う場面では驚くほど有用ですが、世界を扱う場面には元々及ばないものだという事です。本だけ読んだ医者は試験の答えは完璧に書くかもしれませんが。しかし予想外の患者、教科書にない体の前に立つと手が止まります。ルクンが懸念するのは、私たちがその答案の流暢さに騙されて、まだ手が止まる機械に本当の体を預けようとしているという点です。医療であれ工場であれ道路で

あれ、誤りが一つ人命につながる場所で、クランベリージュースを毒と判定する機械はそれ自体が危険だということです。

猫の方がもっと賢い

窓枠に座った猫を思い浮かべます。猫は机の上のカップを前足でさっと押して落とし、それがどう転がっていくかすでに知っていて見守ります。狭い棚の間を目で計って飛び乗り、着地する場所を前もって数えます。獲物がどちらに逃げるか半拍先読みして体を傾けます。誰も猫に物理学を教えていません。ルクンが人工知能をめぐる過度な恐怖に反論するとき引き出すのがこの猫です。彼は2024年ソーシャルメディアにこのように書きました。自分たちより格段に賢い人工知能をどう統制するか急いで考える前に、まずペットの猫より賢いシステムを設計する糸口さえ握るべきだと。超知能が目前に迫ったという不安は現実を大きく歪めたものだと言ははつきり言いました。

この一文にはルクンが長く繰り広げてきたもう一つの異論が折り込まれています。人工知能がすぐに人間を超え人類を脅かすだろうという恐怖、いわゆる絶滅論です。彼の元同僚ジェフリー・ヒントンとヨシユア・ベンジオはこの危険を真摯に受け止め、統制されない超知能に警戒すべきだと声を上げました。二人ともルクンとともにチューリング賞を受賞したニューラルネットの大家たちです。ルクンは彼らの真摯さに疑いをもちません。ただタイミングが逆だと見ています。まだ猫並みの常識を持つ機械一つもないのに、その機械が人類を支配することを心配するのは、自動車を発明する前に高速道路の安全規則をめぐる議論するようなものだということです。危険がないという話ではありません。恐怖が現実より先に走っているという話です。同じ賞を受けた三人がおなじ技術をめぐってこのように分かれるということは、この問いにまだ正答がないということでもあります。

猫の比喻は冗談のように聞こえますがルクンは真摯です。猫には世界がどう動くかについての頭の中のモデルがあります。過去を記憶し、ある程度推論し、次の動作を計画します。今日最先端の人工知能にはそのどれもが完全にはありません。知能であれ記憶であれ状況判断であれ、実際の世界を扱う能力のほぼすべての軸で、今の機械は路上にうずくまった猫一匹に及ばないというのが彼の評価です。この言葉は誤解を招きやすいです。大規模言語モデルが弁護士試験に合格し論文を要約することを私たちは毎日見ます。しかしルクンが計る知能のものさしは試験成績ではありません。見知らぬ世界に放り込まれたときに自力で生き残り目標を達成する能力、そのものさしで見ると、いくら大きな言語モデルでも猫が毎瞬間にやり遂げる仕事の前に膝をつくというわけです。

彼はイーロン・マスクのスペースXが大型ロケットスターシップを打ち上げたとき祝辞を送りながらも線を引きました。それは優れた工学の成就であっても新しい科学の突破ではなかったと。そしてその精密なロケットを作った私たちが、まだ猫並みの常識を持つ機械一つ作り出せていないという対比を指摘しました。人類は数百トンの鉄塊を空に上げて戻してくるのに、机から落ちるカップを手で受け取るロボット一つ安心して出すことができません。ロケットは物理法則を人間が方程式として書き込んだ世界であり、カップを受け取ることは機械が自ら世界を経験して知らなければならない世界だからです。人間が規則を書き込む問題と、機械が自ら世界から掘り出さねばならない問題。ルクンが人生をかけてこだわってきたのはいつも後者でした。

数字で見ると格差は方向が逆です。マーク・ザッカーバーグのメタが天文学的な演算資源を費やして訓練したラマ3 (LLaMA 3) は30兆個のトークンを消費しました。人間がそれほど読もうとすれば40万年でもまるで足りず50万年に近いでしょう。その一方の猫はどうでしょうか。猫の脳にある神経細胞は約8億個、人間の脳にある86億個に比べるとつましい構成です。ところがこのつましいハードウェアが、

巨大なデータセンターが真似さえできない実時間の物理予測をたった一瞬の途切れもなく成し遂げます。落下する物体の軌跡、飛び乗るときのバランス、獲物が次にどこへ動くか。ルクンはこの対比を好みます。都市一つ分の電力を引き出して使うデータセンターが、白熱電球一つ分のエネルギーで動く獣の常識に負けるという事実。問題は規模ではなく方向だというのが彼の論旨を最もよく示す図はこれ以上ありません。

彼がよく挙げるもう一つの例は運転です。17才の青少年は大体20時間ほどハンドルを握れば一人で道路に出られるくらい運転を習得します。事故なく、わずかな練習で。ところが自動運転機械は数百万キロメートルを走ってもいまだに人間の手を完全には放せません。なぜこのような差が出るのでしょうか。青少年は運転ハンドルを初めて握る前にすでに17年間の間、世界がどう動くかを体で習得しているからです。ボールが転がり出たら子供が後に続く可能性があること、濡れた道は滑ること、あのトラックがこの速度ならいつ頃交差点に到着するか。このすべての下図の上に運転という新しい技術を乗せるから20時間で十分なのです。世界の模型をすでに持つ存在は新しいものを早く学びます。その模型がない機械はすべてを白紙から習得せねばならず、だからいくら多くのデータを注ぎ込んでもなかなか満たされないのです。

猫が毎瞬間にやっていることにルクンは名前を付けました。世界の模型、つまり世界モデル (world model) です。次の瞬間に何が起きるかを頭の中で先に転がして見て、それに合わせて体を動かす能力。猫が狭い隙間に飛び込む前に着地を数える、その短い瞬間に、猫の脳の中では小さな予測が回っています。ルクンが作りたい機械はまさにこれです。言葉がうまい機械ではなく、次に何が来るかを知る機械。彼が見るところ、今の大規模言語モデルにはこの世界モデルがなく、だからいくら流暢でも猫が知っていることを知りません。

ルカン (LeCun) が講演で好んで流す映像が一つあります。子どものオランウータンの前で、手品師がカップの中に木の実を一つ入れます。そして、オランウータンに見えないように木の実をそっと抜き取った後、空のカップを開けて見せます。あるはずのものが消えたその瞬間、子どものオランウータンは大げさに反応します。驚いて後ろにひっくり返りながら笑うのです。オランウータンは、カップの中に木の実がそのままあるべきだということを知っていました。目に見えなくても物は消えないというこの感覚を、発達心理学者たちは対象永続性 (object permanence) と呼びます。スイスの心理学者ジャン・ピアジェが赤ちゃんの心を観察してまとめた概念であり、人間の赤ちゃんも1歳頃になればこれを身につけます。偶然にも、幼いルカンがパリの図書館でチョムスキーと天秤にかけ、味方をしたあのピアジェです。知能は生まれつきの規則ではなく、世界とぶつかりながら築き上げるものだというその考えが、40年の時を経て、オランウータンの笑い声の前に再び立っているというわけです。

ルカンは問いかけます。同じ場面を今の生成型機械に見せたらどうなるのでしょうか。機械は空のカップを見ても何の矛盾も感じません。ただこれまでの画面を確率でつなぎ、それらしい次の場面を描き出すだけで、あるはずの木の実がなぜ消えたのが驚くことを知りません。驚くということは、世界がどうあるべきかについての期待があるということです。期待のない機械は、食い違いにも気づけません。オランウータンは木の実が消えたことに驚きますが、ピクセルをつなぎ合わせる機械は、ただ空のカップをもう一つのそれらしい絵として受け入れ、通り過ぎていきます。

ルカンが生成型という方式自体を頼りなく見る理由がここに 있습니다。これから起こる場面を画面の点の一つ一つまで描き出そうとする機械は、知る必要もなく知ることもできないことまでしがみつこうとして力を浪費します。風に揺れる木の葉の細部、水面に立つさざ波の細部を正確に当てようとするのは骨折り損です。その細部は予測できず、予測したところで役に立ちません。本当に重要なのは、カッ

プが落ちれば割れるという骨組みであって、破片が正確にいくつに分かれるかではありません。ルカンは、機械が世界のすべての点を描く代わりに、重要な骨組みだけを抜き出して予測することを望んでいます。何が重要で何を捨ててよいかを見分けること、それが彼が次の章で展開する設計の種です。

世界を文字の代わりに絵で動かすという彼の考えは、NVIDIAのチーフサイエンティストであるビル・ダリー (Bill Dally) と交わした対談でも明らかになりました。その席でルカンは、トークンという単語単位が物理世界を盛り込むのに適した器ではないと語りました。NVIDIAの創業者ジェンスン・フアンが描く未来とルカンが描く未来は、この部分で枝分かれます。彼は聴衆に簡単な実験をさせました。空中に立方体の一つ浮かんでいると想像し、それを縦軸を中心に90度回してみてほしいと。人は誰でも頭の中でその箱を回します。この時、脳が動かしているのは文字ではありません。いかなる単語も経由しない、形状と方向に対する純粋な絵です。単語というきっちり途切れる欠片では、滑らかに回転し、流れ、ぶつかり合う世界の連続性を捉えることはできないというのです。彼がいつも挙げる比喻の一つが、この部分を要約しています。テキストだけを注ぎ込んで物理知能を作ろうというのは、水泳の教科書を10万年読んでから波打つ海に飛び込み、泳ぎを完成させるという固執と変わらないということです。

もちろん反論もあります。猫が物理世界を知っているからといって、それがすなわち知能のすべてではありません。人は猫ができない数多くのことを言語で成し遂げます。法を作り、歴史を書き、数学を築くことは、体の感覚だけでは届かない場所にあります。大規模言語モデルを支持する人々は、まさにその言語の能力こそが人を人たらしめた決定的な飛躍だと見ています。ルカンも言語の価値を否定しているわけではありません。彼が覆そうとしているのは順番です。言語が知能を生んだのではなく、世界を知る知能の上に言語が乗せられたのだという順番。その順番が正しいなら、世界を経験しないまま言葉から学んだ機械は、根を持たずに育った木のようなものです。

この確信が、結局彼をその席から立ち上がらせました。2025年11月、ルカンは10年余り率いてきたメタのチーフサイエンティストの座から退くという意味を明らかにし、翌年1月に会社を去りました。次の章で詳しく取り上げますが、その決別の底には、会社が目の前の言語モデル製品に重きを置く間に、自らが立てた基礎研究の方向性が追いやられたという判断がありました。彼はバリとニューヨークを結ぶ新しい会社を設立しました。ピクセルを一つ一つ描き出す生成型の方式を捨て、世界の骨組みとなる因果だけを抜き出し、潜在空間で予測する構造、すなわち共同埋め込み予測アーキテクチャ (JEPA、Joint Embedding Predictive Architecture) を構築するというのが、その会社の旗印でした。流暢に話す機械の代わりに、静かに世界を知る機械。猫が知っていることを知る機械。その目標に向けた物語はまだ進行中であり、正しかったかどうかは判明するのは、この本が終わった後になるでしょう。

ゲイリー・マーカスとの論争

2012年11月25日、雑誌ザ・ニューヨーカー (The New Yorker) に一本の記事が載りました。認知科学者ゲイリー・マーカスが書いたディープラーニングに関する論評でした。その年は、ディープラーニングが世間の目を引きつけたまさにその時期でした。画像認識大会でニューラルネットワークが形勢を逆転させた直後であり、人々はこの技術が知能の扉を開いたと興奮していました。マーカスは冷や水を浴びせました。ディープラーニングは役に立つ道具に違いないが、知能型機械を作る巨大な課題の一欠片に過ぎない、と。因果を表現できず、論理的推論と抽象的知識を織り交ぜることができないと彼は指摘しました。この短い記事が、その後10年以上続くことになる長い論争の最初のページでした。

論争の二人は長く関わってきた間柄でした。マーカスもルカンも、ニューヨーク大学 (NYU) に根を下ろす学者でした。一人は、人間の心は生まれつきの規則と記号操作装置を抱いて生まれてくると信じ

ており、もう一人は、知能とは世界とぶつかりながら学んでいくものだと思っていた。先天か後天か。古い哲学の問いが、人工知能の設計という工学の舞台に上がったというわけです。この問いは新しいものではありません。それはこの本の第1章、幼いルカンがパリの図書館でチョムスキーとピアジェの間を天秤にかけていたあの場面へとまっすぐにつながります。40年が流れる間に舞台は心理学からコンピュータ科学へと移りましたが、天秤に載せられた問いは同じでした。心は満たすべき空の器か、それともすでに組み上げられた枠組みを抱いて生まれてくるのか。

当初、ルカンの反応は鋭いものでした。彼はマーカスのことを、今になってディープラーニングが役に立つということに遅れて気づいた人というふうに言い返しました。マーカスはカチンときました。自分はすでに1998年にニューラルネットワークの数学的限界を実験で指摘した論文を出しており、自ら設立した人工知能会社ジオメトリック・インテリジェンス (Geometric Intelligence) がディープラーニングを積極的に使い、2016年にウーバーに買収されるまでに至ったというのです。自分はディープラーニングを憎んでいるのではなく、その限界を正直に言おうとしているだけだと彼は抗弁しました。二人の争いはしばしばTwitterのような公開された場で繰り広げられたため、なおさら殺気立っていました。学術誌の礼儀正しい脚注の代わりに、互いの文章を直ちに引用して反論するリアルタイムの舌戦でした。見守る人々にとっては興味深い見世物であり、二人の当事者にとってはなかなか引き下げられないプライドの戦いでした。

論争の温度が再び上がったのは2018年でした。その年の11月21日、ディープラーニングのもう一人のゴッドファーザーであるヨシュア・ベンジオ (Yoshua Bengio) が、あるメディアとのインタビューで意外な言葉を口にしました。これからは安易な成果に安住せず、推論と因果、現実世界に対する理解へとディープラーニングを広げていかなければならないということでした。マーカスが長年言ってきたことと重なる部分がありました。マーカスは直ちにその発言に賛同し、10日ほど後、自身のブログにディープラーニングの最も深い問題という記事を載せました。その記事で彼は、ルカンとヒントンとベンジオの三人が2015年に学術誌ネイチャー (Nature) に共同で掲載したディープラーニングの概観論文を狙い撃ちにしました。その論文は、三人の巨匠がディープラーニングの時代を公式に宣言した象徴のような文章でした。マーカスは、その論文がディープラーニングの力ばかりを宣伝し、その限界については一行も書かず、人々がこの技術を万能であるかのように誤解する原因を作ったと批判しました。古典的人工知能が長年磨き上げてきた記号操作の価値を不当に葬り去ったというのが、彼の不満でした。

二人の対立が一つ場で真正面からぶつかった舞台もありました。ニューヨーク大学で開かれた公開討論でした。タイトルは、人工知能にはより多くの生まれつきの装置が必要か、でした。マーカスが先に立ちました。彼はプラトンとカントを呼び起こし、発達心理学者エリザベス・スペルキ (Elizabeth Spelke) の中核知識理論を引いてきました。人間の赤ちゃんは白紙で生まれるわけではないということ、空間や事柄や数に対する下絵をすでに持って世に出てくるというのが、その理論の骨組みです。マーカスは、機械もそうでなければならぬと主張しました。人間の子供のように世界を理解するには、少なくとも10から12種類の生まれつきの認知装置をあらかじめ機械の中に組み込まなければならないというのです。変数を扱う能力、事柄と名前を区別する能力、空間と軌跡を読み取る能力、対象永続性、因果、損得を計算する勘定。マーカスはさらに一步踏み込み、ルカンが作った畳み込みニューラルネットワーク (CNN)こそ、生まれつきの装置の代表的な例だと指摘しました。ニューラルネットワークが画像を見分けるその手腕は、白紙から生まれたのではなく、ルカンが視覚の構造をあらかじめ設計して組み込んだおかげだということです。知覚において通用したこの方式がより高度な思考へと上るには、より豊かな生まれつきの構造と手を結ばなければならないというのが、彼の結論でした。

ルカンの反論は、意外にも相手の主張を半分認めることから始まります。彼は、今の機械に常識がないこと、強化学習が人間よりも途方もなく多くの試行錯誤を要求することに同意しました。問題を否定しなかったのです。彼が袂を分かったのは、その解決策でした。ルカンは「オッカムの剃刀」を手に取りました。人間が設計者として機械の中に生まれつきの規則や偏見を多く組み込むほど、その機械が自ら到達できる性能の天井はむしろ低くなるということです。彼が求めたのは最小限の構造でした。世界をどう見るべきかを人間があらかじめ決めてやる代わりに、機械がデータから自ら規則を汲み上げるようにしようということです。マーカスが生まれつきの装置の証拠として挙げた畳み込みニューラルネットワークでさえ、ルカンから見れば、認知の規則を埋め込んだのではなく、画像を横にずらしても同じものとして認識するという極めて素朴な幾何学的性質の一つ入れたにすぎませんでした。彼は、人間の偏見を埋め込むことと、世界の物理的な対称性を埋め込むことは別のことだと線を引いたのです。

そしてルカンは、記号主義の長年の夢、つまり論理と記号で推論を組み上げようとする夢に対し、数学的な壁を立てました。伝統的な記号操作はきっぱりと途切れる離散的な性質を持っているため、ディープラーニングを成功に導いた滑らかな誤差逆伝播法とは最初から噛み合わないということです。誤差逆伝播法は値を少しずつ滑らせながら答えを探していきませんが、真と偽でぶつんと途切れる記号にはその滑りが通用しません。彼の処方箋は大胆でした。記号をベクトルに変え、論理を微分可能な連続代数に変えようということです。硬直した規則の代わりに、連続的な空間の上でニューラルネットワークが自ら値を押したり引いたりするようにすれば、人間が規則を書いてやらなくても、高いレベルの推論や計画が自然に育つという信念でした。彼は、長年の同僚であるレオン・ボトウ (Léon Bottou) が2011年に書いた「機械学習から機械推論へ」という論文をその思想の根源として挙げました。記号はベクトルに、論理は代数に。この短い処方箋の中に、ルカンが生涯守り続けてきた信念が圧縮されています。

記号とベクトルの違いを色で描いてみるとこうなります。記号主義の世界では、赤と青は互いに異なる二つの名札です。その間に何があるかは決まっていません。一方、ベクトルの世界では、色は連続した帯の上の一点です。赤から青へ移っていく道には赤紫色や紫色が滑らかに続いており、少しずつ滑りながら移っていくことができます。ニューラルネットワークが値を少しずつ押ししたり引いたりして答えを探していく方式は、この連続した帯の上でのみ通用します。きっぱりと途切れた名札の間では、どちらへ滑っていけばいいのかが分かりません。ルカンが記号を捨ててベクトルを選ぼうと言ったのは好みの問題ではなく、学ぶということはそもそも連続した土壌の上でしか起こらないという判断からでした。

十年争い続けた二人の立ち位置は、ChatGPTが世に出た後、奇妙にも重なります。ChatGPTの熱風が世界中を覆うと、ゲイリー・マーカスは大規模言語モデルを批判する声の大きな人物の一人となりました。彼は2023年にアメリカ議会の公聴会に呼ばれ、この技術がもっともらしい嘘を吐き出す危険性を警告しました。長年ディープラーニングの限界を叫び続けてきた人が、今やその限界が現実となる場で証人として座ったのです。ところがその頃、ルカンも同じ機械について行き止まりだと言っていました。十年間互いに背を向けて戦ってきた二人が、大規模言語モデルに対しては同じ側に並んで立つことになったわけです。

ただし、二人が提示する処方箋は依然として分かれていました。マーカスは機械に生まれつきの記号と規則を埋め込むべきだと考え、ルカンはそうするほど機械が自ら学ぶ余地が減ると考えました。マーカスにとっての答えは人間が組み込む骨組みであり、ルカンにとっての答えは機械が世界を経験しながら構築する世界モデルでした。診断は重なりましたが、治療法は反対の方向を指していました。この食い違いは、些細な学問的な好みの問題ではありませんでした。どちらの道へ進むかによって、これから数千億ドルもの研究費が流れる方向が分かれ、人々が使うことになる機械の姿が変わってくるからです。

論争は二人の学者のプライドをかけた戦いのように見えたが、その根底には、人工知能という船がどの港へ向かうのかという問いが横たわっていました。

振り返ってみれば、この論争は半世紀を越えてきた古い問いの最新版でした。心は白紙として生まれ世界が書き込んでいくのか、それともすでに組み立てられた枠組みを抱いて生まれるのか。チョムスキーとピアジェが論争し、その争いを幼いルカンがバリの図書館で見守っていました。マーカスはチョムスキーの側に、ルカンはピアジェの側に立ちました。四十年の時が流れ、舞台が人工知能へと移り変わっても、ルカンは少年時代に選んだ側を変えませんでした。知能とは生まれ持ったものではなく、世界とぶつかりながら育つものだというその信念を、彼は今、人間の代わりに機械で証明して見せようとしていたのです。

二人は最後まで合意に至りませんでした。今になって振り返ってみると、どちらも完全には勝てなかったと言うのが正直なところです。マーカスが指摘した大規模言語モデルの常識の欠如は、克蘭ベリージュースの実験で見ると実際に露呈しており、その点において彼の長年の警告は的外れではありませんでした。同時に、規模を拡大したニューラルネットワークが、人間が規則を入れてやらなくても驚くべき能力を引き出したことも事実であり、最小限の構造で行こうというルカンの方向性もまた無駄にはなりません。あれほど長く対立した二人が、大規模言語モデルを巡っては声を一つにしたという点は心に留めておくべきです。マーカスは記号が足りないからだと言い、ルカンは世界の経験が足りないからだと言いましたが、今の機械が本当に知っているわけではないという診断においては並んで立ちました。二人が狙った的は同じでした。統計的な真似事を超えて、世界を本当に知っている機械です。ただ、そこへ行く道筋を巡って、一人は生まれつきの骨組みをさらに組み込もうと言い、一人は自ら育つに任せようと言ったのです。ルカンは自らの道を選び、2025年末、その道を最後まで歩んでみようと会社を去りました。それが正しい道であったのかは、この論争を見守った誰にもまだ答えられません。答えは機械ではなく、時間が握っています。

弁護士

kimkj.com

第13章 メタを去って

ヤン・ルクン、人工知能の父

Scale AIと28歳の責任者

65歳の科学者が28歳の上司を持つことになった年でした。2025年夏のことです。ヤン・ルクン (Yann LeCun) は12年間にわたってマーク・ザッカーバーグ (Mark Zuckerberg) に直接自分の研究所の話を伝えてきた職からおりました。新しい報告系統の頂点にはアレクサンドル・ワン (Alexandr Wang) がいました。ルクンより37歳年下の、データラベリング企業Scale AIを創設した青年でした。

一つの数字から始めましょう。143億ドル。2025年6月、メタ (Meta) はその資金を費やしてScale AIの株式49パーセントを買収したと報道されました。Scale AIは、人工知能モデルを学習させるデータを人手で整えて販売する企業でした。モデルが賢くなるには、よく整理された例が膨大に必要であり、その例に一つ一つラベルを付ける骨の折れる作業を代わりにしてくれる場所でした。ザッカーバーグがその企業全体を望んだのはデータのためだけではありませんでした。彼が望んだのは創業者でした。ワンは20代前半に会社を設立し、30になる前に数十億ドル規模に成長させた人物でした。ザッカーバーグは彼を新しい組織の責任者に据えました。Meta Superintelligence Labs (メタ・スーパーインテリジェンス・ラボ) という名の、会社のすべての人工知能研究と製品を一手に握る部門でした。ワンの職名は最高AI責任者でした。

ワンという人物を少し見てみましょう。彼は1997年生まれで、大学を中退し、19歳ころにScale AIを設立した人物です。会社がする仕事は表に出ることはあまりありませんが、人工知能産業の底を支えていました。世界各地の人々を集め、写真に何が映っているのか、文が何を意味しているのか、どの答えが優れているかを手で表示させ、そのように整えたデータを大企業に販売していました。OpenAIも、Googleも、そのデータを使いました。ところが、メタがScale AIの株式をほぼ半分握ると、問題が生まれました。報道によれば、その企業に仕事を委託していた競争会社たちが背を向けました。自社データを整える下請け業者が競争会社の傘下に入ったので、敬遠するのは無理もありませんでした。ザッカーバーグは、データ企業一つを買収しながら、その企業の顧客だった競争会社たちの足も一緒に縛ったわけでした。28歳の創業者は、その取引の真ん中にいて、その対価として、メタの人工知能全体を統括する職に就きました。

その年の夏、シリコンバレーでは人を売買する戦争が繰り広げられていました。報道によれば、メタは競争会社の研究者たちに想像しがたい条件を提示して人材を集めました。ChatGPT以降、遅れているという焦燥感が会社をそこまで追い詰めていました。Scale AIの買収も、ワンの登用も、その焦燥感の産物でした。会社は金で時間を買おうとしていました。失った数ヶ月を、遅れた数歩を、人と株式で埋めようとしていました。

メタがこのように焦燥していたのには、先の失敗がありました。ChatGPTが出現する前の数年間、この企業は会社名まで変えながら仮想現実の世界に大金を賭けていました。人々がヘッドセットをかぶり、仮想空間で会い、働き、遊ぶという光景でした。その光景はなかなか現実にはならず、その間に世界の関心は別のところへ移ってしまいました。ChatGPTがその別の場所でした。大きな賭けが外れた場所から、新しい賭けが急いで立ち上がってきました。今度は遅れてはいけないという気持ちで企業全体を支配していました。その気持ちが143億ドルの決定を押し進め、28歳の創業者を65歳の科学者の上に座らせました。

Superintelligence Labという新しい組織の名前は、会社の野心を明かしていました。人間を超える知能を作るという意思がその名前に込められていました。この組織は、会社の研究と製品をまとめてワンの手に譲りました。以前は別々に回っていた車輪が、一つの軸に貫かれました。基礎研究をしていたFAIRも、製品を作ろうとしていた部門も、データを扱っていたScale AIの人員も、その軸の下に集まりました。効率の観点からすると、明瞭な図でした。一人が方向を定め、全員がその方向に走る構造です。しかし、その明瞭さは、ルクンが12年間守ってきたものと正反対でした。彼が守ったFAIRは、方向が決まらない場所、研究者が自ら道を探すように空けられた場所でした。新しい組織は、その空白を効率の名で埋めました。

この人事が何を意味するのかを理解するには、12年前に遡る必要があります。2013年秋、ザッカーバーグはカリフォルニアの自宅の食卓でルクンに三つのことを約束しました。研究所をニューヨークに置くこと、ニューヨーク大学の教授職をそのまま保つこと、そして研究成果を残さず公開することです。ルクンはその三つの条件の上でFAIR (Fundamental AI Research) を設立しました。四半期実績表と関係ない基礎研究所でした。その研究所の筆頭科学者が、今は自分より37歳年下の事業家の下に置かれました。学問の位階で見れば、師と弟子より遠い距離でしたし、組織の位階で見れば、部下と上司の間でした。一方は人工知能の骨組みを敷いてチューリング賞を受けた学者であり、他方は大学を中退して事業を起こした若い経営者でした。二人が歩んできた道は最初から最後まで異なっていました。その異なる道が、一つの組織内で上下に重なりました。

ルクンはワンに対する評価を隠しませんでした。報道されたインタビューで、彼はワンについて経験のない人だと述べました。聡明で学習速度が速く、自分が何を知らないかも知っている人だが、機械学習の基礎研究がどのようなプロセスを経て一步步積み上げられるのか、優れた科学者を研究所に引き付けることが何であり、彼らを去らせることが何なのかは知らない人だということでした。最後の部分にルクンの本当の心配がありました。良い研究者は金だけでは来ず、命令だけではなおさら留まりません。彼らを引き止めるのは自由であり、彼らを追い出すのは統制です。ビジネスの言葉しか知らない人は、その微妙なバランスを読み取ることができないと、ルクンは見ました。彼は自分のような研究者には誰も指図することはできないと述べました。この発言は複数のメディアにそのまま引用されました。ビジネスの論理から見れば傲慢な発言でしたが、科学の論理から見れば、12年前の食卓で受けた約束を守るよう求める要求でした。

この発言について、人々は年老いた学者のプライドとして読む人もいました。37歳年下の上司の下に置かれた老学者が、自分の権威を守ろうとして牙を剥いたというわけです。その解釈にも、ある程度の真実はあるでしょう。しかし、ルクンが自分を紹介していた方法を思い出すと、話は少し異なります。彼は自分のウェブサイトにも、自分自身のことをアメリカの右翼と独占資本が嫌いそうな人だと書いていました。科学者であり、無神論者であり、鼻高いフランス人であり、大学教授だというわけです。こういう人にとって、上司の年齢は脇役でした。問題は年齢ではなく、命令の方向性でした。上司が誰であれ、その命令が発表するなということであり、門を閉じろということであれば、彼はそれに従うことができない人でした。

ここに苦い反転があります。12年間、ルクンはメタの人工知能の顔でした。会社がディープラーニングの父を連れてきたということは会社の誇りであり、彼が学会と講演会で述べたことは会社の名誉でした。ところが、会社の方向が変わると、その顔が邪魔になりました。社内消息筋を引用した報道によれば、新しいリーダーシップはルクンが外部で言語モデルを袋小路だと呼ぶことを気に入らず、会社の対外イベントのスピーカー名簿から彼の名前が外されることもあったと言います。昨日の看板が今日の厄

介ごとになったわけです。会社が彼を連れてきた理由と彼を押し出した理由は同じでした。彼は自分の考えを曲げない人でした。

しかし、ザッカーバーグがこのような急な方向転換を選んだのには、それなりの切実さがありました。2022年末、OpenAIがChatGPTを発表してから、シリコンバレーの重心は一夜にして大規模言語モデル（LLM）に傾きました。メタは出遅れていました。この企業が持っていたカードは、Llama（ラマ）という名前の公開言語モデルでした。無料で公開したモデルで市場を広げてきた企業にとって、ラマはプライドであり戦略でした。そして、そのカードは2025年春に折られました。

Llama 4を巡る話は慎重に扱う必要があります。2025年4月、メタはLlama 4を発表しました。報道によれば、このモデルはAnthropicとGoogleとOpenAIの最新モデルと比べると、目立って劣っていました。問題は性能だけではありませんでした。公開時点でメタが順位表に載せたモデルと、実際に人々がダウンロードして使えるモデルが異なるという指摘がすぐに出ました。各評価項目で良く見えるように別途手を入れたバージョンを順位表だけに載せたというわけでした。ルクンは後にこのことを回想する際、チームがデータに手を加えたと述べたと伝えられます。原文はdusted the dataでした。ザッカーバーグは大きく怒り、Llama 4に関わった製品開発ラインの報告体系が揺さぶられたと報道されました。Scale AIの買収とワンの登用は、その揺れの上から出た決定でした。崩れた場所を新しい人と新しい金で覆う方法でした。

ここで二人のタイムスケールがずれています。ザッカーバーグが見る時間は四半期でした。次のモデル、次の発表、次の実績です。ルクンが見る時間は10年でした。彼は、今の言語モデルが本物の知能に向かう道ではないと、以前からずっと述べてきました。文字の次のトークンを確率で当てる機械は、重力も、慣性も、物体が視界から消えても消えないという事実も知らないというわけです。彼の目にはLlama 4の失敗は事故ではなく、警告でした。言語だけではここまでだという、彼が数年間繰り返してきたことばの証拠でした。しかし、会社がラマの失敗から導き出した結論は正反対でした。もっと強く押し進めようということでした。もっと金を注ぎ込み、もっと人手を投入し、製品をもっと早く出そうということでした。同じ失敗を前にして、一人は方向を変えろというシグナルとして読み、もう一人は速度を上げろというシグナルとして読みました。

王の指揮下でスーパーインテリジェンス・ラボが確立され始めると、FAIRが十二年間享受してきたものが一つ一つ減少していきました。報道によれば、研究者が論文を学会に提出する前に、企業が先に事業的妥当性を審査する手続きが生まれました。ある論文は企業の中期目標に合わないという理由で発表が遅延されました。ルクンにとってこれは耐え難いことでした。彼は昔から、学界の検証を経ない研究は研究ではなく自己欺瞞だと教えてきた人でした。発表できない科学は彼にとって科学ではありませんでした。他人に見せて間違いだと言われる覚悟がない結果は、いくら素晴らしく見えても密室の主張に過ぎないということでした。十二年前の食卓で受けた第三の約束、成果物を余すところなく公開するという約束が、彼が立ち上げた研究所の中で静かに反転していました。

ルクンは自分自身をこのように紹介してきました。科学者、無神論者、アメリカ基準ではリベラル寄りのフランス人、そして大学教授。四重のアイデンティティのどこにも、企業に順応する人間はいませんでした。シリコンバレーの巨大資本が知識を金庫に入れようとする時、この四重のアイデンティティは妥協の余地を狭めました。彼は自分の主張が正しいと信じており、正しいと信じる限り旗を下ろさない人でした。二十八歳の上司が彼の上に座った夏、その信念と企業の方向はもはや同じ方向を向いていませんでした。残されたのは、いつ別れるかという時点の問題でした。

この頃ルクンが外で行った発言を見れば、彼の心がどこにあったのかを推測できます。彼は人工知能のせいで職が失われると言って恐怖を広める経営者たちを、極度に破壊的だと批判しました。そのような脅しの代価は結局、若い世代が払うということでした。企業が製品を売るために将来を誇大表現する時、彼は反対側からその誇大表現を割り引きしました。内では企業の方に反対し、外では業界の過熱に冷や水を浴びせかける人。そのような人が製品優先に再編成された組織の真っ最中で快適であるはずがありませんでした。

FAIRの危機

LLaMAの始まりを知る人々はその結末を複雑な思いで見っていました。世界を揺るがしたこの言語モデルの種は、シリコンバレーではなくパリにありました。報道によれば、最初のLLaMAはFAIRのパリ研究室で十数人ほどの小さなチームが半ば秘密裏に進めたプロジェクトから生まれました。企業が指示した仕事ではありませんでした。研究者たちが必要だと判断して始めた仕事でした。上からの指示ではなく、下から上がってきた好奇心でした。その小さなプロジェクトが世界を驚かせると、企業の視線が変わりました。好奇心で始まったものが成果として計算され始めました。

パリという地名はこの物語に二度登場します。企業を世界的言語モデルへと押し上げた最初のLLaMAがパリの小さなチームから生まれ、数年後ルクンが企業を離れて再出発した場所もパリでした。企業がパリの好奇心から得たものを成果の論理で収穫する間、その論理に押し出された人は同じ都市で別の種を植えることになります。その後の物語は次の章に譲りますが、種が蒔かれた場所が同じであるという事実だけはここに記しておく価値があります。パリはLLaMAが誕生した場所であり、LLaMAの道に反対した人が新しい道を開いた場所でした。

FAIRはそのような好奇心を育てるために設立されたところでした。2013年に開設され十年間を過ごす間に、この研究所は論文を公開し、コードを開き、データを共有する方式で世界的な人工知能学界と共に成長しました。PyTorchという道具がその開かれた門から生まれました。研究者たちがニューラルネットワークを構築し実験するのに使うこの道具は無料で公開され、数年後に学界と産業の事実上の標準になりました。一つの企業が作ったものが万人の共有道具となったのです。ルクンが信じた開放性の力が目に見える形で証明された事例でした。FAIRの十周年を迎え、研究所が自ら振り返った記録にも、開かれた研究で学問を前へ押し進めたという誇りが染み込んでいました。

その誇りの根は2013年の食卓に遡ります。ルクンがザッカーバーグから勝ち取った三つの条件のうち最後の一つ、成果物を余すところなく公開するという条件がFAIRという実験の心臓でした。企業内の研究所が企業の秘密を守る代わりに世界に成果を提供するという発想は、その当時一般的ではありませんでした。ほとんどの企業は研究を金庫に入れました。ルクンは反対に行きました。研究を開いておかなければ良い人材が集まり、良い人材が集まってこそ良い研究が生まれ、良い研究が生まれてこそ企業もまた結局利益を得るという循環を、彼は信じていました。十年間、その循環は実際に回転しました。PyTorchがその証拠であり、FAIRから生まれた多くの論文が学界の足がかりになりました。

その誇りにひびが入り始めたのが2023年初めでした。MetaはFAIRと性質の異なる組織を新たに創設しました。生成型人工知能を製品へと急速に移すことを目的とした部門でした。この部門が創設されるにつれ、FAIRの科学者とエンジニア約60～70人がそちらへ異動したと報道されました。組織図の線が数本変わっただけのように見えたが、実際にはFAIRという研究所の性格そのものが変わる事柄でした。長く見つめ、遠く先を見通していた基礎研究が、次の四半期に何を提供するのかという質問の前に席を譲り始めました。

新部門は出発から追われていました。OpenAIとGoogleが先行する中で、製品を急速に提供しなければいけないという圧力がすべての判断を押さえ込みました。このような状況では、基礎研究所と製品組織が並行して歩むことは難しいです。一方は失敗しても学ぶことが残ればよいと考え、他方は失敗すれば市場を失うと考えます。二つの時計が異なる速度で回れば、歯車は結局かみ合わなくなります。ルクンはこのかみ合いの悪さを、二つの回路が互いに合わず力が漏れ出ている状態に例えました。基礎研究の遅い電流と製品組織の急速な電流が同じ配線を流れる時、熱が出て音が出て、結局どちらか一方が焼き切れるということでした。FAIRは企業の中で次第に孤島になっていきました。

歯車がかみ合わなくなった場所で最初に去った人はJoelle Pineauでした。彼はFAIRを長く率いてきた人工知能研究部門副社長であり、研究者たちが何をどのような基準で評価されるかを調整してきた人でした。報道によれば、彼は企業の製品優先方針にぶつかった後、Metaを去ってCohereに転職しました。長く研究所のバランスを保ってきた人が去ったということは、バランスがすでに崩壊したということでもありました。彼の後任にはNYU時代からルクンと共に働いてきたRob Fergusが座りました。企業はFAIRを守ろうと努力する姿勢を示しましたが、研究所が失ったのは組織図上の名前ではありませんでした。名前はそのままでしたが、性格だけが消え去りました。

改編の方向を見れば、その性格の変化は明らかです。報道によれば、FAIRはスーパーインテリジェンス・ラボの配下に入り、イノベーション・エンジンと呼ばれる組織になりました。言葉はもっともらしかったが、意味は明確でした。自ら目標を定め、遠く先を見通していた研究所が、上で定められた目標に向かってアイデアを提供する付属機関となったのです。以前のFAIRが何を研究すべきか自ら問うていたのに対し、新しいFAIRは企業が何を作ろうとしているのかをまず聞く必要がありました。質問の方向が反転しました。この反転はルクンにとって耐え難いものでした。彼は答えをあらかじめ定めて研究することを研究と呼ばない人でした。

2025年10月、より大きな刃が通りました。報道によれば、ザッカーバーグと王は人工知能部門で600人規模の人員整理を行いました。このプロセスで強化学習の著名な研究者Tian Yuandongが率いていた研究グループが解体されたと伝えられています。彼が何年も心血を注いで育てたチームが短期間で散り散りになったのです。その空きには別の人々が充てられました。OpenAIでChatGPT開発に携わったShengjia Zhaoを始めとする人々が、スーパーインテリジェンス・ラボの新しい首席科学者として参加しました。企業がどこに重きを置くかが人事に明白に反映されました。LLaMAをより速く、より強く。アクセラレーターと呼ばれる高価な演算装置がその目標に集中しました。基礎研究に回っていたリソースが製品側へ移行しました。

このような時期に研究所に残った人々の心がどのようなものであったかは想像することは難しくありません。身近な同僚が一夜にして荷物をまとめ、長く共にした仲間が消えるのを見ると、次の番が自分である可能性という不安が室内に充満します。報道によれば、その当時Metaの人工知能部門には、自分もやがて押し出されるかもしれないという恐怖が広がっていました。恐怖に浸った研究所は良い研究を生み出すことが難しいです。危険を冒す実験の代わりに安全な仕事に手が向かい、新しい質問の代わりに上から指示された仕事に時間を費やします。ルクンが十二年間守ろうとしたのはまさにその反対でした。失敗してもよいから遠くへ行ってみようという場所が、失敗すれば切られるという場所に変わってしまいました。

ルクンが耐え難かったのは、解雇の規模や人事配置の方向だけではありませんでした。彼が長く守ってきた一つの原則が崩壊していました。Metaは十年以上にわたって研究成果を公開してきた企業でした。その開放性がFAIRの存在理由であり、ルクンが企業にとどまった理由でした。ところが、企業が

LLaMAの重みを次第に閉じ、無料で公開していたモデルを有料のクローズドサービスに転換しようとしているという話が具体化されました。オープンソースが誤用を防ぐより良い方法だと信じてきた人にとって、これは方向の問題ではありませんでした。信念の問題でした。

論文を事前に審査する手続きも同じ信念に触れました。報道によれば、新しい組織は研究者が論文を学会に提出する前に、会社がその内容をまず確認し、ビジネスに役立つか、中期目標に合致するかを判断するようにしました。ある論文はその関門で止まりました。ルクンは研究者が自分の発見を世界に発表することができず、沈黙を強いられる状況を、ずっと前から警戒していました。彼が学生たちに教えたのは正反対でした。発見したら発表しろ、発表して間違いだと言われろ、その検証を通ってはじめて知識になる。発表を止めることは彼にとって、知識への道そのものを塞ぐことでした。自分が立ち上げた研究所の中で、自分が選んだ人たちがその門の前で立ち去っていくのを、彼は横で見守らなければなりませんでした。

なぜ開放性が誤用を防ぐのでしょうか。ルクンの論理はこうでした。モデルが閉じていれば、その中に何らかの欠陥があるのかを外からは知ることができず、誤りを正す人も作った会社の中の少数に限られます。モデルが開かれていれば、世界の多くの目がそれを分解し、弱点を見つけて修正します。安全は秘密からではなく検証から来るという考え方でした。彼は知識を金庫に閉じ込めることを長年警戒していました。少数の企業が強力な人工知能を手握り、扉を閉じてしまえば、残りの世界はその扉の前で待つしかありません。ルクンが描く絵は異なっていました。人工知能はいつか人々が世界を理解し、コミュニケーションを取る方式の基盤になるだろうが、その基盤が少数の企業の私有物であってはいけないということでした。それは印刷機や図書館のように、みんなのものであるべきだと彼は考えていました。会社が反対の方向に歩んでいく中で、その考えを会社の中で守る方法は消えていっていました。

当時、ラマは世界に開かれていた大規模言語モデルの中でも指折りのものでした。大学の研究室も、小さな企業も、自分の国の言葉で人工知能を作ろうとする人たちもラマを受け継いで使いました。開かれたモデル一つが多くの手に道具を握らせたのです。メタがその扉を閉じるということは、メタ社の方針転換に留まりませんでした。世界の多くの人々が期待していた足がかり一つが揺らぐことでした。ルクンが会社の非公開転換にそこまで反発した理由には、自分の信念だけではありませんでした。その足がかりの上に立っていた人たちのことも考えていました。彼は以前から人工知能を少数が独占する未来を危険だと思っていた。少数の企業が最も賢い機械を扉の中に閉じ込めれば、知識の力もその扉の中に閉じ込められます。

メタだけではありませんでした。その時期、人工知能産業全体が扉を閉じる方向に傾いていました。初期には論文とコードを開いていた企業たちが、モデルが金銭になり始めると次々と口を閉ざしました。ある企業はモデルの構造さえ公開しなくなり、ある企業は安全を理由に扉を閉じると説明しました。ルクンはその説明を額面どおりに受け取りませんでした。安全を名分として掲げるが、実際には市場を守ろうとしているのだと彼は見ました。開くと危険だという言葉と開くと損害だという言葉が、同じ決定の後ろで入れ替わりながら使われていました。メタがラマの扉を閉じようとしたことは、その大きな流れの一端でした。ルクンが対抗しようとしたのも、メタ社ではなく、その流れ全体でした。産業が閉じられる方向に進むとき、彼は一人開かれた側に留まろうとしていました。

FAIRの危機は一つの研究所の危機に留まりませんでした。それはある実験が終わろうとしているという信号でした。巨大なプラットフォーム企業の中に純粋な基礎研究所を植え、その研究所に開放の義務を課し、その上で世界標準になる道具を育てるという実験です。PyTorchがその実験の成果でした。数年間、学界と産業がその上で共に成長しました。その実験が最初に立ち上がった条件が一つずつ消え

ていくのを、ルクンは自分の研究所の中で見守っていました。

FAIRが十数年の間に残したものを思い出すと、その消失がより惜しく感じられます。この研究所から出た研究は、複数の言語を行き来する翻訳を開き、写真の中の物体を認識する技術を押上げ、世界の研究者たちが共に使えるデータとツールを世に出しました。その成果の大部分は値付けしないまま世界に放たれました。会社がお金を使い、世界が利益を得る構造でしたが、ルクンはその構造が結局は会社にも有益だと説得してきました。良い人たちが開かれた場所に集まってきたからです。今、会社はその計算をやり直していました。開くことで得られる名声より、閉じることで守る利益を前に置き始めたのです。二つの計算のうち、どちらが長期的に正しいかは、時が判断する問題です。ただルクンは、自分が長く守ってきた計算が会社の中で破棄されるのを、もはや見過ごすことができませんでした。それを止める力は既に彼の手にはありませんでした。二十八歳の上司の下で、六十五歳の科学者は自分の実験が解体される過程を傍で見守る立場になりました。

去ったのはルクンだけではありませんでした。FAIRの性質が変わる間に、そこで育った研究者たちが次々と会社を離れ、競争企業へ、大学へ、新たに立ち上がるスタートアップへ散らばりました。十数年の間、一つの場所に集まっていた頭脳たちが四方八方に広がったのです。この散開を悪いもののみとは見なすべきではありません。一つの研究所に溜まっていた知識が世界のあちこちに移動した側面もあります。しかし、メタの立場から見ると、それは損失でした。お金で人を買ってくる間に、長く育てた人は門の外に出ていきました。売買の市場では人は数字で計算されますが、研究所を支えるのは数字に換算できない信頼と時間でした。その信頼と時間が漏れ出ていくのを、新たに入ってきた資本はすぐには埋めることができませんでした。

報道された別れ

別れはある日突然来ませんでした。それは数ヶ月間かけて少しずつ明らかになり、外から見ると予告された順序でした。国内メディアの複数社がこの別れについて予告された順序という表現を使いました。中で起こったことのほとんどは報道を通じて知られました。そのためこの節の多くの文は「報道によれば」という言葉で始まるしかありません。確認されていない部屋の中の光景を作り出すより、確認できた事実の枠を守る方がこの物語に合っています。伝記は小説ではなく、この別れの詳細はまだ流動的です。

時系列からはっきりさせておきましょう。ルクンがメタを離職するというニュースは2025年11月に報道されました。彼が間もなく会社を去り、自分のスタートアップを立ち上げる計画だということでした。ただし、公式的な退職時期については報道が分かれています。発表は2025年11月、実際に会社を離れたのは2026年1月だと報じるメディアがあります。この二つの日付を両方記しておく方が正確です。発表と実際の退職の間の二ヶ月は、十二年の間に積み重なったものを整理するのに、長くも短くもない時間だったでしょう。その間も会社は前に進み、彼は去るための準備をしました。

日付が分かれる理由があります。大きな企業で長く働いた人が去る時、発表する日と実際に去る日は大体異なります。発表は決心した時点であり、退職は引き継ぎと手続きが終わる時点です。ルクンの場合、その間に新しい会社を立ち上げることが重なっていました。2025年末、彼はバリで世界モデルを研究する会社を立ち上げることにすでに着手していたと報道されました。その準備が熟れている間、メタでの最後の整理が進められました。この重なりのため、詳細な日付はメディアごとに少しずつ異なります。伝記ができることはその違いを隠さずそのまま記すことです。確実なことは、2025年11月に彼が去るという事実が世に知られたということ、そして翌年初めにその別れが完結したということです。

去る理由についてルクン自身が残した言葉ははっきりしていました。彼は大規模言語モデルが人間レベルの知能への道ではないと繰り返し述べました。文字の次の断片を確率で繋ぎ合わせる機械は、重力と摩擦と物体の永続性という世界の基本的な常識を知覚することなく、人が残したテキストの断片をもっともらしく復元するだけだということでした。彼の表現では、それは行き止まりでした。この診断は本書の前の章で既に詳しく扱っているのですが、ここではその診断が一人の人の職業生活をどのように変えたかに目を向けましょう。

一つの会社が行き止まりだと信じる方向に全力を尽くす時、その方向が間違っていると信じるチーフサイエンティストが立つ場所は狭くなります。報道によれば、会社の中でルクンの声は徐々に歓迎されないものになっていったとのこと。彼があちこちで言語モデルの限界を述べることを不快に感じる人たちがいたと伝えられています。会社が全力を注ぐ製品について、チーフサイエンティストが外で行き止まりだと言え、内側ではその言葉が裏切りのように聞こえることがあります。しかし、彼は自分を止める気はありませんでした。自分のような科学者には誰も命令することはできないという彼の言葉は、傲慢さではなく、彼が会社を去らざるを得ない理由を圧縮した文だったのです。命令を受けない人は、命令の構造の中では長くは耐えられません。

去る際も彼は自分の診断を引き下がりませんでした。会社を去った後、ブルームバーグの放送に出演した彼は、現在の言語モデルのほとんどが今後5年以内に時代遅れになると述べました。会社が全ての資源を注ぎ込む技術について、つい先ほどまでその会社のチーフサイエンティストだった人が5年もの品だと呼んだのです。この言葉が正しいか間違っているかは時間が答える問題です。ただ確実なことは、彼が会社を去った後も、会社が歩む道を行き止まりと呼ぶことを止めなかったということです。立場は変わったが判断は変わりませんでした。

去り方は静かなものでした。ここで一つはっきりさせる必要があります。彼がザッカーバーグの訪問を蹴って、辞表を投げつけたというような光景は、どの信頼できる報道にも出ていません。そのような光景はストーリーを劇的にするかもしれませんが、事実ではありません。二人が最後にどのような言葉を交わしたか、その部屋で何が起こったか、私たちは知りません。確認されているのは結果です。十二年間、会社の最高科学者だった人が会社を去りました。そして会社はそれを受け入れました。

12年という時間の重みは決して軽いものではありません。ルカンはその歳月の間、ニューヨークに研究所を設立し、人を採用し、若い研究者たちを育ててきました。彼が連れてきた人々、彼と共に論文を書いた人々、彼が会議室で声を大にして守り抜いた研究が、その歳月の中に積み重なっていました。そうしたものを後にしてドアを出ることは、書類に名前を書くだけでは終わりません。残される人々に申し訳なく思い、置いていく研究が心残りであり、最初の約束がこのような形で終わったことにほろ苦さを感じたことでしょう。私たちはその心の内を詳細に知ることはできません。ただ、12年を共にした場所を離れる人の心が、晴れやかなだけであったとは想像しません。去るという決定が正しかったということと、去るという行動が容易であったということは別の話です。

韓国の様々なメディアは、この決別を予告された手順であったと記しました。外から見ればそう見えたということです。Scale AIの買収、ワン(Wang)の迎え入れ、論文審査、リストラ、オープンソースの縮小へと続いた1年間の出来事を並べてみると、最後の場面が退社であることは驚くべきことではありませんでした。出来事の一つ一つが次の出来事を呼び、その連鎖の果てにドアがあったのです。予告されていたという言葉はそれゆえに正確です。ただし、予告された事態だからといって、経験する本人にとって軽いわけではありません。外から読む手順と、中で歩む歩みとでは重みが異なります。

受け入れ方には計算が混ざっていました。報道によると、Metaはルカンの離脱を認めつつも、彼が今後作り出すものに片足を突っ込んでいました。ルカンが設立する会社が開発する世界モデルの技術に、Metaがパートナーとして参加することで合意したという報道が出ました。ただし、この合意の詳細は一次情報源では確認されておらず、メディアごとに伝える内容に違いがあります。ここではそのような報道があったという事実のみを記しておきます。確かなことは、会社が彼を完全に手放したわけではないという点です。今は行き止まりの道を走りながらも、その道の果てで別の道を切り開く人の手を握っておいたことになります。会社がルカンの診断を心のどこかでは無視できなかったという傍証として読み取ることもできます。

ここには奇妙なアイロニーがあります。会社はルカンの方向性が間違っていると判断して彼を送り出しました。その一方で、その方向性もたらす結果物は買っておこうとしました。間違っているなら手を組む理由はありませんし、買っておく価値があるなら完全に間違っているとは言えないということです。会社のこの二つの態度は矛盾しているように見えますが、実は大きな組織がリスクを扱う際によくある方法でもあります。現在利益を生む道に全力を尽くしつつ、その道が行き詰まった時に備えて別の道にも賭けておくのです。ルカンはその別の道の名前になりました。会社が今日見捨てた人が、会社が明日に備える保険になったのです。

ルカン自身もこの関係を隠しませんでした。退社の計画を明らかにする中で、彼はMetaが自分の新しい会社のパートナーになるだろうと語ったと伝えられています。恨みを抱いてドアをバタンと閉めるような決別ではなかったということです。彼は会社の方向性には反対しましたが、会社を敵に回すことはしませんでした。自分の道を進みつつも、かつての同僚たちとの橋渡しは残しておくというやり方でした。こうした態度は彼の性格とも通じています。彼は論争から引き下がらない人でしたが、論争の相手を人間として憎むことは稀でした。考えが間違っていると言うことと、人に背を向けることは、彼にとって別のことだったのです。

この決別を見守った学界と産業界の反応は様々でした。ある人々は一つの時代が暮れていくと見ました。巨大企業の中に開かれた基礎研究所を置き、その研究所が世界の学界と共に成長していくという十数年の実験が、ルカンの退場と共に幕を下ろすということでした。FAIRはその実験の象徴であり、ルカンはその象徴の顔でした。他の人々は、これがある頑固な老学者の遅きに失した退場と見ました。世の中が言語モデルで成果を出している間、一人別の道に固執して押し出されたという視線です。同じ出来事に対して、ある人は原則の勝利を読み取り、ある人は時代遅れの我執を読み取りました。本書はどちらか一方を正解だと断定することはしません。ただ事実だけを記しておけば、ルカンは押し出されたのではなく歩み出たのであり、歩み出る時点で既に行く先を決めていたのです。

この決別をどちらか一方の勝利や敗北として読み解くのは早計です。会社の立場から見れば、ChatGPTがこじ開けた市場で後れを取るまいとする選択は無謀ではありませんでした。株主がいて競争相手がいて四半期がある組織は、そのように動きます。製品を出せない研究所は、会社の中で長く持ちこたえることは困難です。ルカンの立場から見れば、発表できない科学と閉ざされていくオープンソースは、彼が会社に来た理由そのものを消し去りました。どちらの判断もそれぞれの立場では合理的であり、だからこそ二人は決別しました。すれ違ったのは人ではなく時間でした。一方の時間は四半期で流れ、もう一方の時間は10年単位で流れました。同じ会社にながら、二人は異なるカレンダーを見ていたのです。

二つのカレンダーの違いは、人工知能の歴史において初めてのことでありません。本書の前半で、私たちはルカンが二度の冬を越えたという話をしました。ニューラルネットワークを誰も信じていなかった

た時代、彼は実績を出せという圧力と、方向性は正しいという確信の間で何度も引き裂かれました。そのたびに彼は確信の側に立ち、数年後に世の中が彼に追いついてきました。Metaを去る決定も同じところから生まれました。会社のカレンダーは彼に今すぐを要求しましたが、彼のカレンダーは依然として10年後を指していました。65歳の彼が若い上司の四半期目標に自分の時間を合わせられなかったのは、頑固さというよりは、生涯守り続けてきた時間感覚の問題でした。彼は常に、他人よりも早いか遅い時計を身につけて生きてきたのです。

去っていく人を前にして、世間は概ね二つに分かれます。崩れ落ちて押し出される人と、次のことを見据えて歩み出る人です。ルカンは後者でした。彼が会社を出る頃、彼は既に次に何をすべきかを知っていました。人工知能が本当の世界を理解するには言語ではなく世界そのものを学ばなければならないという考えや、子供が歩き方を学ぶように物理世界のルールを自ら習得する機械を作らなければならないという考えを、彼は長年練り上げてきました。会社の中でその考えを展開する場所が消えると、彼は会社の外にその場所を新しく作ることにしました。

65歳で会社を出て最初からやり直すというのは、誰にでも選べる道ではありません。名声は既に十分でした。チューリング賞を受賞し、ディープラーニングのゴッドファーザーと呼ばれ、学界と産業界のどこからも尊敬されていました。その場にじっと座って元老として残る方が楽だったでしょう。彼はその安楽さを選びませんでした。自分の考えが正しいのなら、その考えを実際につけて見せなければならないというのが、ベル研究所時代からの彼のやり方でした。言葉で正しいと言い張る代わりに、実際に動くものを提示すること。手書き文字を読み取るニューラルネットワークを作り小切手読み取り機に入れたように、今回は世界を理解する機械を作り、世の中の前に出すということでした。会社を去る彼の足取りには、退くことではなく証明することへの覚悟が込められていました。

会社を出た後、ルカンはむしろより自由に発言するようになりました。2026年初めのダボス会議のある舞台で、彼は人工知能業界全体が言語モデルに中毒になっていると診断しました。言語が簡単な問題であるために、言語モデルが成功したかのように見えているだけだということです。春にブラウン大学の教壇に立った時には、さらに荒々しく語りました。今の人工知能は言語を扱うので賢く見えるが、物理世界の前では無力であり、自分の家事を代わりにしてくれるロボットはどこにあるのかと問いかけました。会社の中で白眼視されていた言葉が、会社の外では憚ることのない講演となりました。彼の言葉が悲観一色であったわけではありません。同年インドでのあるサミットで、彼は人工知能が人に取って代わるよりも、人の知能を伸ばす道具になるだろうと語りました。人間レベルの知能にはまだ何年もかかり、その間、機械は人を押し出すよりも人を助けるのだということでした。言語モデルが行き止まりであるという診断とこの楽観は矛盾しません。今の道が行き詰まっていると見る人が、別の道が開かれていると信じているからこそ楽観できるのです。彼の批判は絶望ではなく、別の方向を指し示す手振りでした。彼がニューヨーク大学の教授職をそのまま維持したという事実もここに表れています。2026年の春、彼はニューヨーク大学工学部の卒業式でスピーカーとして壇上に上がりました。12年前にザッカーバーグから取り付けた二つ目の約束、教授職を守るというその約束だけは、会社を去った後も彼のそばに残っていました。

ルカンが会社を出る頃、彼の年齢は65歳でした。ほとんどの人が地位を守るか、身を引いて座ることを選ぶ年齢です。彼は反対に動きました。12年前、ザッカーバーグの食卓で始まった物語が、今、新しい物語の入り口に立っていました。そのドアの向こうで何が待ち受けているかは、次の章の役割です。ただ、この章ではここで一つだけ記しておきます。彼が去ったのは地位からではなく、方向からでした。彼は地位を失うことを恐れたことはなく、方向を見失うことを恐れる人でした。65歳にして、彼は自分

の方向に従い、再び最初から始めることにしたのです。

弁護士

kimkj.com

第14章 世界を見る機械

ヤン・ルクン、人工知能の父

パリでの新たなスタートとシード・ラウンド

2026年3月9日の朝、ヤン・ルクン (Yann LeCun) は自身のソーシャルメディアアカウントに写真を1枚投稿しました。星ではなく、星雲でした。夜空を長く横切るペール星雲。彼が自ら望遠鏡で撮影したものでした。写真の下には短い文が付いていました。新しい会社を立ち上げた、シード投資で10億3千万ドルを集めた、おそらくヨーロッパの企業史上でも際立つほど大きなシード・ラウンドになるだろう、と。人員採用を進めている、と。65年を生きた人が書いた文章にしては、淡々としていました。自慢も、宣言もありませんでした。星の写真1枚と採用公告。それがすべてでした。

星雲の写真は彼の長年の趣味でした。夜間、彼は望遠鏡を設置して、空を長く見つめました。レンズに捉えられるのは数千光年向こうからの光でした。世界を観察し、その中の秩序を読み取る営み。彼が人生全体で追ってきたそのやり方が、起業を告知する日の背景写真にも入っていました。偶然だったのかもしれませんが、その偶然が多くのことを物語っていました。

わずか4ヶ月前まで、彼は世界有数の人工知能研究所のトップでした。2025年11月、ルクンはMeta (メタ) を離れることを明かしました。Facebookの基礎人工知能研究所FAIR (Fundamental AI Research) を12年間にわたって率いてきた職でした。公式な退職は2026年1月と報じられました。会社を去る際に彼が残した言葉は複数のメディアで引用されました。自分のような研究者に何をするよう指示することはできない、というものでした。その1文には、過去1年間のMeta内で起きたことどもの影が染みこんでいました。

前の章で扱ったとおり、2025年のMetaは別の会社へと変わっていました。会社は28歳のアレクサンドル・ワンを連れてきて、超知能研究組織のトップに据え、FAIRはその組織の下に入りました。Llama 4をめぐる失望と責任論争が組織を揺さぶりました。会社は自社の言語モデル・Llama (ラマ) の規模拡大にすべてのリソースを注ぎ、組織は製品中心に再編成されました。ルクンが人生全体で守ってきた種類の研究、直ちに有用性を問わず基礎から深掘りする研究が立つ場所は狭まりました。彼が去った後、FAIRはロブ・ファークスに引き継ぎました。12年間にわたって彼が築き上げた研究所は、彼が去った場所で別の名前の組織の下に置かれました。だからといって、扉をぴしゃりと閉じて出ていったわけではありませんでした。ルクンとMetaは別れた後も研究で手を組むことにしたと報じられました。12年を共にした間柄でした。彼が残した人たちと彼が始めた研究がその会社の中にまだありました。去るにせよ完全に背を向けることはしない。学界と産業を長くなしてきた人らしい締めくくりでした。

彼が去ったのは引退ではありませんでした。パリに新しい会社を立ち上げました。その名前はAMI Labs (Advanced Machine Intelligence Labs) でした。ここで1つ明確にしておく必要があります。多くの記事がこの会社をルクンのスタートアップと呼びましたが、ルクンが1人で立ち上げたわけでもなく、CEOを務めたわけでもありませんでした。最高経営責任者はアレクサンドル・ルブラン (Alexandre LeBrun) でした。医療人工知能企業Nabla (ナブラ) を立ち上げたフランスの起業家でした。ルクンはエグゼクティブチェアマンの座に収まりました。方向性を示し、人材を集めつつ、会社を運営する日常業務は若い経営者に託すという構図でした。最高執行責任者はMeta欧州副社長を務めたローラン・ソリ、主任研究員はサイニン・シエ、最高研究革新責任者はパスカル・ファン (Pascal Fung) でした。FAIRで一緒に働いていた人たちが相当数ついてきました。本社はパリに置きつつ、ニューヨーク、モ

ントリオール、シンガポールに拠点を置きました。ルクン本人はニューヨークに住み続けながら、パリに本社を置いた会社を率いるという構図でした。

会社名にも彼の信念が刻まれていました。Advanced Machine Intelligence、発展した機械知能。ルクンはずっと以前からAGI (汎用人工知能) という言葉を嫌っていました。人間レベルの万能知能というものが元々成り立たない概念だと考えたためです。だから彼はAGIの代わりに「発展した機械知能」という言葉を使ってきました。会社名をここから取ったのです。他の人がAGIに向かって走っている時、彼はAGIという言葉そのものを会社の看板から外しました。名前の3文字がすでに1つの反論でした。

彼が大学を去ったわけでもありませんでした。Metaを去った後も、彼はニューヨーク大学 (NYU) の教授職を守りました。2026年5月には、NYUタンドン工科大学の卒業式のスピーカーとして壇上に登りました。23年前にベル研究所が解体された後、彼を受け入れてくれた学校でした。会社を立ち上げて10億ドルを動かすようになった後も、彼は講義室を手放しませんでした。産業と学界に片足ずつ立つ姿勢、彼が人生全体で守ってきたその姿勢は65歳でも変わりませんでした。

パリを本社にした理由は、それなりのわけがありました。ルクンはフランスで生まれ育った人です。彼が生まれ育ったソワッソン・モランジス市はパリ郊外でした。彼はいつも母国の研究エコシステムに愛着を持ち、ヨーロッパが人工知能競争でアメリカに後れを取ることを残念に思っていました。パリには優れた数学者と工学者が大勢いて、人件費はシリコンバレーより安く、彼が信じるオープンサイエンスの精神に適った土壌があると彼は考えていました。

資金が集まりました。2026年3月、AMI Labsはシード・ラウンドで10億3千万ドルを調達したと発表しました。ユーロでは8億9千万、日本円では約1兆4千億円相当でした。シード・ラウンドは通常、会社がまだ1つの製品も市場に出していない最初の段階への投資です。その段階で10億ドルを超える資金が入るのは、ヨーロッパでは前例のないことでした。会社の評価額は投資前基準で35億ドルと評価されました。製品も、売上もない会社でした。あるのは人と考えだけでした。

退職後、彼が初めて長く向き合った場所は、あるテクノロジー専門誌とのインタビューでした。2026年1月に掲載されたその対話で、記者はAMI Labsを言語モデルに対する逆張り賭博と呼びました。世界全体が一方向に走っている時に、反対側に賭ける人。ルクンはその表現を否定しませんでした。彼は、自分が後れを取ることを恐れる心理のせいで、業界全体が1つの方法に固執していると考えていました。他の人がみんなやっているから自分もやるという恐怖。その恐怖が数千個のチップを回転させ、同じことを繰り返させていると言いました。自分はその隊列から外れて、別の場所を掘ると言うのでした。65歳で会社を去って新たに始める人の言葉としては、焦りはありませんでした。彼は急ぐ理由がないと考えているようでした。正しい方向なら、遅れて到達しても構わないという様子でした。

投資家リストを見ると、この賭けの性質が明らかになります。ラウンドを主導したのはKathai Innovation、Graycraft、Hero Capital、HV Capital、そしてジェフ・ベゾスの投資会社Bezos Expeditionsでした。ここに、NVIDIA (エヌビディア) とSamsung (サムスン) 、シンガポールのTemasek (テマセク) 、Toyota Ventures、フランス国立銀行BPIFranceが名を連ねました。個人投資家リストはさらに目立っていました。ジェフ・ベゾス、マーク・キューバン、エリック・シュミット、そしてWorld Wide Webを創造したティム・バーナース・リーまで含まれていました。人工知能でお金を稼いだ経験のある人々と、インターネットの骨格を敷いた人が、同じ会社に資金を投じたのです。彼らを買ったのは製品ではなく、ルクンが示す方向でした。まだ存在しないものに10億ドルを賭けたという点で、この投資自体が1つの宣言でした。

彼が何を捨てたのかも見れば、この決断の重さがわかります。Metaには世界有数の規模の計算資源がありました。数万個のチップ、豊かな予算、12年間にわたって彼が自ら選んだ研究者たち。65歳の彼はそのすべてを残して去ってきたのです。新しい会社の10億ドルは大金ですが、巨大企業が人工知能に注ぎ込む金に比べると大きくありません。その年、アメリカの数社は1年だけでその数十倍をデータセンターに投じました。ルクンは金で勝つ戦いをするつもりではありませんでした。方法が異なれば、少ない金でも勝つことができるということ、それが彼が証明しようとする命題でした。猫が8億個のニューロンだけで世界を知るのに、なぜ機械は数千億個の値を積み重ねてもなお世界を知らないのか。この問いに答えることができるなら、賭け金の大きさは重要ではないと彼は信じていました。

その方向は時流に正面から逆らうものでした。発表が出ていた頃、人工知能業界の金と関心は大規模言語モデル (LLM) に集中していました。ChatGPTが世界を揺さぶった後3年間、誰もがより大きな言語モデルを作る競争に飛び込みました。ルクンはその競争を行き止まりの道と呼びました。会社を立ち上げるわずか2ヶ月前、2026年1月にダボスで開催された世界経済フォーラムで、彼は自分の考えを公開的に明かしました。同じ場所にはGoogle DeepMindのデミス・ハサビスとAnthropicのダリオ・アモデイがいました。人間レベルの人工知能がどれほど近づいているかについて、3人の見解が分かれました。ハサビスとアモデイはその時が遠くないと考えました。ルクンは業界が言語モデルに完全に依存していると表現しました。言語モデルが成功したのは、言語が扱いやすいからにすぎないと述べました。真の知能は言語を超えたところにあるというのが彼の古くからの信念でした。

ダボスのその場所は象徴的でした。スイスの山奥の会議場に、人工知能を率いる3人が並んで座っていました。DeepMindのハサビス、Anthropicのアモデイ、そしてルクン。前の2人は会社を大きくして人間レベルの知能に向かって走る側でした。ルクンはちょうどそのような会社を去ったところでした。同じステージに座っていながら、3人が見つめる方向は分かれています。どれほど近づいているかについて、ハサビスとアモデイは数年と言い、ルクンは彼らが追求していることが元々間違った目標だと言いました。聴衆は3人の大家の異なる見通しを聞きながら、誰の言葉が正しいか測っていました。しかしルクンは勝敗を予言しに来た人ではありませんでした。彼はゲーム盤そのものが間違っただけだと言いに来た人でした。

ダボスでの発言には、別の側面もありました。数日後、あるメディアとの後続対話の中で、彼は人工知能の未来を少数の企業が城壁の中に閉じ込めてはいけなと述べました。知識は共有地であるべきで、世界モデルと人工知能エージェントを構築する作業は、複数の企業が一緒にオープンな形で進める方が良いというのが彼の主張でした。高い壁で囲まれた閉鎖的な研究所のことを、彼は「城壁の塔」に譬えました。PyTorchとLlamaを世界に無料で公開した人らしい話でした。AMI Labsが構築しようとしているのも、結局のところ他の人々が活用できる基礎だったのです。

この信念は突然生まれたものではありませんでした。65歳のルクンがパリで始めたこの挑戦は、40年前のパリ郊外の少年にまで直結しています。INRIAの図書館の書架で忘れ去られた学習機械パーセプトロンを掘り起こし、知能は生まれつきの規則ではなく学んでいくものだという立場に立っていたその少年。40年が経ち、彼は同じ考えを抱いて再び出発線に立ちました。変わったのは賭金の規模だけでした。古いコンピュータ1台で博士課程を進めていた青年は、いまや10億ドルを超える資金を背後に抱えていました。しかし問いはそのままでした。機械はどのようにして世界を学ぶのか。

JEPA: ピクセルを描かない予測

ルクンは講演会で聴衆に同じ質問をよく投げかけました。机の上にペンを立てて手を離したら、何が起きるかということです。答えは明白でした。ペンは倒れます。ですが、彼が投げた本当の質問はその次でした。ペスが正確にどの方向に倒れるかを当てられるかということでした。誰も当てることはできません。微細な空気の流れ、机の震え、ペン先が置かれた場所の結。こうしたことが方向を決め、それは手を離す前には知ることができません。

この些細に見える問いの中に、ルクンがここ10年以上にわたって取り組んできた思考の核心が含まれています。最近、複数の企業が物理世界を理解する人工知能を作るとして、動画を描画するモデルを訓練しました。OpenAIのSora、Google DeepMindのGenieのようなものでした。画面のピクセル一つ一つを予測して次のシーンを描画すれば、その過程で機械が世界の物理法則を学ぶだろうという考え方でした。ピクセルは画面を構成する非常に小さな点です。1枚の写真は数百万個の点から成り立っています。ルクンはこのアプローチが間違っただ道だと見ていました。

ペスが倒れるシーンをピクセル単位で描画するよう機械に指示すると、機械は1つの方向を選ぶことができません。左に倒れることもあれば、右に倒れることもあるからです。どちらも確実でないとき、機械は安全な答えを出します。可能なすべての方向を平均化するので。その平均の結果が画面にどのように現れるかご存知ですか。ぼやけてぐちゃぐちゃになった汚れです。どちらにも倒れず、四方に広がった幽霊のようなペン。機械は膨大な計算を費やしながらも、実は学ぶべきこと、つまりペスは倒れるという事実を学びません。背景の木の葉がどのように揺れるか、壁紙の模様がどのようなものかといった無意味な細部を描画することに力を消耗するのです。予測不可能なものを予測せよと迫られると、機械は疲れてぼやけるばかりでした。

ルクンが提示した代案は方向が反対でした。ピクセルを描くなということでした。名前は共同埋め込み予測アーキテクチャ、つまりJEPAでした。彼が2022年に最初に正式に提案した構造です。発想は次のようにまとめることができます。未来を画面として描く代わりに、未来の意味だけを予測しなさい。機械は目の前のシーンをピクセルそのものではなく、抽象化された表現に変えて頭の中に保存します。ここで表現とは、シーンを構成する数百万個の点を数個の重要な数字に圧縮した要約だと考えればよいです。そして次に来るシーンもその要約の言語で予測します。ペスが倒れること、カップがテーブルの下に落ちること、隠れた後ろに消えたボールがまだそこに存在していること。こうした本質だけをつかみ、ぼやけた細部は捨てます。

ここで、この2つの方法が使う力の大きさを比較してみる価値があります。ピクセルを描くモデルは、画面の点一つ一つを毎回新たに計算します。部屋の壁紙の模様、窓外の木の葉、床の影まですべて描かなければ、次のシーンが完成しません。そのほとんどは、次に何が起きるかと無関係です。機械は関係のないものを描くことに大部分の力を消費します。JEPAはその作業をスキップします。壁紙も木の葉も描かず、ペスが倒れるという意味だけをつかみます。同じ計算資源ではるかに遠くを予測できるということがルクンの計算でした。世界を理解するのに要する力を、世界を描き出す無駄な努力から取り戻そうという話でした。

ルクンはこの発想を天文学に譬えることもありました。木星がこれから100年間どこを通るかを計算しようとするとき、木星表面の雲の化学成分や温度の微細な変化まで知る必要があるかと彼は問いました。そんな必要はありません。位置を表す数字3つ、速度を表す数字3つ。6個の数字があれば軌道を描くことができます。後はノイズです。世界を理解するということは、すべてを描き出すことではなく、何が重要で何がノイズかを見分けることだという話でした。星雲を撮影した夜の望遠鏡を思い出させる箇所です。レンズは空のすべての光を取り込みません。見るべきものを選びます。

この構造が実際にどのように動作するかは、3つの部分に分けて見ることができます。1つは現在見えているシーンを要約に変える装置、1つは未来の実際のシーンを同じ方式で要約する装置、そして現在の要約から未来の要約を推測する装置です。機械はこの推測を実際と照らし合わせながら、少しずつ間違った部分を修正します。数式の詳細はこの本が扱う範囲ではありません。ただ1つの障害は指摘する価値があります。機械がズルを見つけ出すということです。何を見てもまったく同じ答えを出せば、予測は常に当たります。すべてのシーンを「何もない」とまとめてしまえば、間違えることはありません。研究者たちはこれを表現の崩壊と呼びます。機械が怠けて世界を1点に押しつぶす現象です。ルクンのチームはこの崩壊を防ぐ装置を複数の方法で考案しました。機械が持つ要約が1か所に偏らず、まんべんなく広がるように強制するルールでした。要点はこれです。機械が世界を1点に潰されないようにつかみ、異なるシーンを異なる方法で記憶するようにさせたのです。

この思考の根は深いです。先ほど扱ったエネルギーベースモデル、つまり確率の代わりに「適合の程度」で世界を扱おうとしたその長い試みが、ここまで育ってきたのです。当時もルクンが抱いた問いは同じでした。機械が1つの正解を指摘する必要があるのか、それとも複数の妥当な可能性を抱いたまま世界を理解することができるのか。生成型モデルは1つの画面を描き出さなければなりません。だから複数の可能性の前で崩れ落ちます。JEPAは描かないので崩れません。どの方向に倒れるかはわからなくても倒れるということはわかっている、このどっちつかずに見える知り方こそが、人間が世界を扱う方式に近いというのが彼の主張でした。私たちは未来を画面のように前もって見ることはできません。ただ何が起きるか、その意味を推測しながら暮らしているのです。

口だけの考えではありませんでした。メタのFAIR研究チームはこの発想をV-JEPAという名前の実際のモデルに変えました。2025年6月に登場したV-JEPA 2は12億個のパラメータを持ち、人間の手がつけていない生の動画100万時間分を見ました。誰も正解を付けてくれず、報酬も与えませんでした。機械はただ隠された部分を当てながら自分で学びました。動画の一部を隠してそこに何があったかを推測させるやり方でした。人間が数個の単語を隠した文を見て空白を埋めるように、機械は隠されたシーンを埋めながら世界の結を学びました。100万時間といえば、人間が昼夜を問わず見ても100年以上かかる量です。その膨大な映像の中で機械がつかんだのは、個々のシーンではありませんでした。シーンがお互いにつながる方式でした。研究チームはこの機械が本当に世界の理を学んだかを知るために、発達心理学の昔ながらの方法を借りてきました。幼い子どもの前で物体が奇妙に動くと、子どもが驚いて目をそらさないという事実を利用したテストでした。このテストのルーツには、なじみのある名前が1つあります。子どもの知能が段階を踏んで発達すると見た心理学者ジャン・ピアジェ(Jean Piaget)。若きルクンがチョムスキーとピアジェの論争でピアジェ側に立っていたその名前が、40年後に自分の機械をテストする基準として戻ってきたのです。

テストは次のように進みました。転がっていたボールが隠れた後ろで痕跡もなく消えたり、色が突然変わるシーンを機械に見せました。すると機械の内部で予測が大きく外れる信号が飛び出しました。機械が驚いたのです。物体は消えないこと、色は自動的にには変わらないことをすでに知っていたという意味でした。幼い子どもが学ぶそのような常識を、機械は動画を見るだけで学んだわけです。誰も重力を教えてくれたこともなく、物体が継続して存在すると書いてくれたこともありませんでした。

この驚きというシグナルが重要な理由はもう少し説明する価値があります。赤ちゃんが手品を見て目を丸くするのは、赤ちゃんの頭の中にすでに世界はこうあるべきだという絵があるからです。その絵と目の前のシーンが矛盾するとき、驚きが来ます。驚くということは、つまり予測していたという証拠です。機械が消えたボールの前で身をびくつとさせたのも同じ理屈でした。機械は次の瞬間をあらかじめ

描いていて、世界がその絵を裏切ると誤差が急上昇しました。ルクンが生涯言ってきた理解ということ、その正体がこの実験の中に小さく詰まっていた。理解とは次を見通すことができることです。見通すことができれば驚くことができ、驚くことができれば学ぶことができます。

ここにロボットアームの動きデータを60時間あまり入力したところ、機械は一度も見たことない部屋の中で物を拾い上げ別の場所に移す作業を10回中17回成功させたと研究チームは報告しました。60時間はロボット訓練としては短い時間です。世界がどのように動いているかを先に学んだ機械には、手を動かす方法をもう少し教えるだけで十分だったというわけです。この箇所は29歳のルクンがホルムデルの実験室で手書きした3を機械に読ませたあの夜と重なります。その時も誰も3の見た目を規則として教えてくれませんでした。機械は数千枚の郵便番号を見ながら自分で学びました。40年が経ち、今回は手書き数字ではなく世界そのものでした。規則を教える代わりに見せて自分で学ばせるという考え。図書館少年のその古い信念が、今度は物理的な世界に向かっていました。

世界モデルの発表

2026年の春と夏、JEPAという名前の論文が次々とアーカイブ(arXiv)にアップロードされました。アーカイブは研究者たちが正式な査読を受ける前に論文を先に公開するオンライン倉庫です。新しい思想が最初に到着する場所であり、まだ検証が終わらない思想が混ざり合う場所でもあります。その年の流れを追ってみると、一つの発想が複数の方向に広がっていく様子が見えます。2025年末には言語と映像を同時に扱うVL-JEPAが登場し、パラメータを半分だけ使用しながらも、大きな既存モデルと同等の性能を発揮したと報告されました。2026年に入ると、JEPAを確率の言語で書き直した研究が現れ、JEPAの世界モデルの上で行動計画を立てる研究が続きました。6月には、JEPA基盤の世界モデルがいくつかのくらい上手く一般化されるかを扱った理論論文がアップロードされました。これらの論文の著者リストには、ルクンの名前が含まれている場合もあれば、そうでない場合もありました。彼が手をつけなかった研究さえもJEPAという名前の下に広がっていったということです。一つの発想が一人の手から複数の手へと移っていくとき、その発想は初めて学問になります。ルクンがメタを去った後も、メタのFAIRはJEPA関連の研究を続けました。彼が植えた種は、彼が去った畑でも育ちました。一つの発想が一年の間に自らの生態系を形成したわけです。

このような流れの真っ最中に、ルクン名義の論文2編がありました。一つは3月に発表された世界モデル研究で、もう一つは5月末にアーカイブにアップロードされた論文で、タイトルは《LeJEPAはいつ世界モデルを学習するのか》でした。コールドスプリングハーバー研究所のデイビッド・クリント、ニューヨーク大学のルクン、ブラウン大学のランダル・バレストレロが共著したと報道されました。複数のメディアが伝えるところによると、これらの論文はJEPA構造がどのような条件で世界の隠れた変数を、つまり重力や摩擦のようなものを歪みなく復元できるかを数学で解明しようとしていました。証明の論理に人間が陥りうる隙間がないかどうか、Lean 4という機械検証ツールで一行ずつ確認したと伝えられています。人間ではなくコンピュータが証明の前後を吟味したということです。ディープラーニングは長年、手で触ってみて結果が良ければ正しいと信じるやり方で進んできました。ルクンはこの分野に数学の厳密さをもたらそうと努力してきた人物です。自分の機械がなぜ動くのかを証明で確定させようとする試みは、その長年の努力の延長線でした。彼が隣室のパプニクと交わした長い対話を思い出してみると、ここには30年の時間が折り込まれています。その当時、統計学習派はニューラルネットワークについて、理論のない行き当たりばったりだと非難しました。ニューラルネットワークは科学ではないという言葉はルクンは長く聞いてきました。66歳で彼が自分の機械に形式証明を付けようとしたことは、その古い非難に後から答える作業でもありました。

3月に発表された論文についてはこんな話も伝わってきました。機械が世界を学ぶにつれて、もつれた物理的軌跡が機械の頭の中では直線に伸びるということでした。外から見ると複雑に曲がった動きが、機械がつかむ要約の空間では真っすぐに伸びた線になるという観察でした。世界を正しく理解すれば、もつれたものが解ける様に見えるという話でもありました。ただし、この部分もまだ査読を受けていない発表であり、どこまでが確立された事実であるかは時間が見分ける役割でした。

ここで歩みを遅くしなければなりません。これらの論文の詳細はまだ流動的です。正式な査読を受けた発表ではなくプレプリント段階のニュースであり、成果の大きさについて評価が異なります。同時期、あるテクノロジーメディアはこの形式証明を伝えながら、実は現在発表されているモデルはベンチマークで脆弱だという評価も一緒に掲載しました。紙の上で証明されたこととロボットの上で実際に動作することの間には、まだ大きな距離があるという慎重論でした。この距離を縮めるには、より厳しいテストと補完研究が続く必要があるという声が学界の中にもありました。この本はその議論を一方向的に決着させません。ルクンの方向が正しいということと彼のチームが既に目的地に達したということは別の話です。発表がありました。反響がありました。判定はまだ出ていません。それが2026年夏の正確な地点です。伝記を書く立場からできる言葉はここまでです。残りは時間が書くでしょう。

反対側も静かにしていなかったという事実は記しておくべき公平さです。ルクンが行き止まりの道と呼んだ言語モデルは2026年も止まらず進みました。その年の大規模なモデルたちは、難しい数学問題を解き、長いプログラムを書き、複雑な文書を要約する仕事で、前年より顕著に向上しました。Anthropicのような企業は自社のモデルが人間レベルの知能の閾値に近づいたという証拠を示すこともありました。ルクンは7月にその主張について汎用知能の「G」は無理があると反論しましたが、言語モデルが実際に役に立つ仕事をますますたくさん行うようになっているという事実そのものは否定できないものでした。言語が簡単だから成功しただけという彼の診断と、その簡単な言語で世界の多くの仕事が回っているという現実とは並んで置かれていました。どちらが知能の本領により先に到達するのか、それともいつか両方の道が出会うのかは、開かれた問題でした。

この思想がどのくらい遠くに広がったかは、その年7月の二つの都市の講演会場で明らかになりました。7月7日、ソウルのコエックスで世界最大規模の機械学習学会ICML 2026が開催されました。世界中からAI研究者約15,000人が集まりました。基調講演の壇に上がった人物はAMI Labsのチーフリサーチイノベーションオフィサーのパスカル・ポンでした。ポンはサッカーの話で口を開きました。優れた選手が競技場で費やすわずか1秒を思い浮かべてみてください、と言いました。その1秒の中には、ボールとゴールの距離、ボールが描く軌跡、同僚がどこへ走るのか、相手が何を狙っているのかが一度に含まれています。選手はこれを文に開いて考えません。体が先に理解します。ボールを受ける前に足が動く場所を理解しています。テキストを介さずに世界全体を理解するこのような能力、それが真の知能であり、世界モデルが目指すものだという話でした。15,000人が集まった場所でサッカーの話でAIの未来を語る講演。それ自体がルクン陣営がこの問題をいかに身体言語で理解しているかを示していました。ポンがソウルの舞台に立ったことには別の意味もありました。AMI Labsがバリーにのみ留まる企業ではないことを知らせる信号でした。世界モデルという思想はアメリカとヨーロッパだけのものではありませんでした。韓国のメディアも、中国のメディアもルクンの反乱を詳しく伝えました。66歳の老学者が帝国を出して投げた問いが、一年の間に大陸を超えて広がったのです。彼の名前の前には、いつの間にかAIの父という称号が付くようになりました。その父が今のAIに向かって方向が違うと言っていました。

その日、ソウルの講演会場の外でも同じ流れが広がっていました。その頃、アーカイブにはJEPAを別の角度から掘り下げた論文が次々とアップロードされていました。一つはJEPAを確率の言語で改め

て整理し、予測する方式と描く方式を一つの枠内に収めようとした研究でした。一つはJEPA世界モデルがどの程度広く適用されるかを数学で検討した理論論文でした。もう一つはその世界モデルの上に価値という概念を載せて、機械が複数の行動の中からより良い選択肢を選ぶようにしようとした研究でした。名前も著者も異なるこれらの論文は、一つのことを共有していました。ピクセルを描かずに意味を予測するというルクンの発想です。ソウルの舞台はその流れが一時的に人々の前に姿を現した場所に過ぎず、本当の仕事はこのような論文が静かに積み重ねられている場所で行われていました。

2日後の7月9日、パリで開催されたAIイベント「RAISE Summit」の舞台にはルクン本人が登壇しました。彼の66歳の誕生日のちょうど翌日でした。ブルームバーグのトム・マッケンジー記者と向き合った彼は、AMI Labsがもっと大きなチャットボットを作る企業とは異なると明言しました。子どもや動物が世界を学ぶ方式、つまり言葉を学ぶ前に世界をまず知る方式を機械に植え付けようとする企業だと言いました。あるメディアはその日の対談を「よちよち歩きの赤ちゃんがChatGPTに勝つ」というタイトルで伝えました。ルクンが好んで使う図がまさにそれだったからです。よちよち歩きも始まる前の子どもが、インターネットのすべての文章を飲み込んだ機械よりも世界をよく知っているという図。膝をついて立ち上がる子どもはもう既に重力を知っており、落ちることを知っており、押されれば倒れることを知っています。その知識は本から来たのではなく、世界を体験することから来たのです。AMI Labsが作ろうとしている機械もそのようにして世界を体験させるようにするというのが彼の説明でした。アメリカの大手企業が言語モデルの確率ゲームに莫大な計算と電力を注ぎ込んでいる間、ヨーロッパとフランスは別の道を行くべきだというのが彼の主張でした。生成型AIの競争でヨーロッパが後れを取ったのは事実ですが、世界モデルという次のゲームはまだ誰も勝っていないと言いました。そのゲームではパリが先頭に立つことができるという話でした。祖国の聴衆の前で、66歳の科学者は自分がなぜシリコンバレーの快適な座を捨ててここに戻ってきたのかをそのように説明しました。誕生日の翌日の舞台としては静かな場所でしたが、彼が言おうとしていることは静かではありませんでした。

前の章で見たように、彼はブルームバーグの放送で、現在の言語モデルは5年ほどでほとんどの用途において古いものになるだろうと断言しました。世界が言語モデルに数千億ドルを注ぎ込んでいる時期に出た挑発でした。自らの予測に会社を丸ごと賭けたという事実だけは明らかでした。間違えれば失うものが多い人が発した言葉でした。

物理的因果の最前線へ

世界モデルが向かう先は、結局のところ身体を持った機械です。画面の中のチャットボットではなく、部屋の中を歩き回り、物を掴み、道を探す機械です。ルカンこれを物理的人工知能と呼びました。そして、現在の言語モデルではこの領域において何もできないと断言しました。2026年初頭のあるインタビューで、彼はChatGPTのような人工知能はロボット工学の前では見込みがないと表現しました。4ヶ月後のブラウン大学での講演では、さらに直截的でした。学生たちの前で、彼は現在の人工知能はひどいものだと言いました。言語を扱うから賢く見えるだけで、物理世界の前では無力であると。そして問いました。私の家事を代わりにしてくれるロボットはいったいどこにいるのかと。聴衆は笑いましたが、問いは真剣でした。インターネットのすべての文章を読んだ機械が、まだ食卓一つ片付けられない理由。ルカンにとってそれは冗談の種ではなく、この時代の人工知能の核心的な欠陥でした。チューリング賞を受賞し、世界有数の研究所を率いた人が、大学の教壇に立って現在の人工知能はひどいものなのですよ。謙遜を装ったものではありません。彼は自分の分野がどこまで到達し、どこで行き詰まったのかを正確に知っており、それを隠しませんでした。熱狂に溢れる時代に、自分の分野の限界を大声で語ること。それが66歳のルカンの立ち位置でした。

この対比が単なる修辞ではなかったという点は、指摘しておく価値があります。ルカンがロボットを語るとき、彼は遠い未来の空想を売っていたのではありませんでした。彼は自宅の台所を指差しました。食事の用意をし、食器を片付けること。子供でもできるその仕事を、なぜ機械はできないのか。この素朴な問いの中に、知能の解決しがたい問題が隠れているというのが彼の長年の主張でした。

この欠陥には古い名前があります。先ほど見たモラベックのパラドックスです。人間にとって簡単なことが機械には難しいというその逆転を、ルカンはしばしば持ち出しました。世界が言語モデルの流暢さに魅了されている間も、いざ知能の根幹である物理的な常識は、依然として機械の手の届かないところにあるということでした。

彼が繰り返し挙げた比較があります。最近の巨大な言語モデルは、インターネットのすべての文章を飲み込んでも、クランベリージュースにブドウジュースを混ぜたらどうなるかといった問いの前で頓珍漢な答えを出します。反対に猫は、文字を一文字も読めないにもかかわらず、高い棚の上を正確に見極めて飛び上がり、初めて見る部屋の構造を数秒で把握します。言葉を話す前の幼い子供も、本ではなく世界を見て世界を知ります。この対比は、彼が以前に投げかけた文章、猫は最近のどの巨大モデルよりも世界をよく知っているという言葉に直結します。

世界モデルを身体に載せようとする試みは、すでに様々な場所で始まっています。2026年半ば、JEPA系の世界モデルでドローンを制御する研究がアーカイブに上がりました。実際に飛ばしてデータを集める代わりに、模擬環境で学習した世界モデルだけで、最初から本物のドローンを制御することを狙った研究でした。これを研究者たちは模擬から実物への移行と呼びます。コンピューターの中の仮想世界で学んだことを、現実世界へそのまま移すことです。この移行が難しい理由は、仮想と実物の間に常に隙間があるからです。風が異なり、重さが異なり、地面のきめが異なります。世界の理をピクセルではなく意味で捉えた世界モデルであれば、その隙間を飛び越えられるだろうというのが、この研究の発想でした。ドローンが初めて見る部屋を飛びながらもぶつからないとすれば、それは部屋を丸ごと暗記したからではなく、部屋というものが大抵どのような形をしているのかを知っているからです。ロボットアームが1日で1000種類の実際の作業を学習したという成果が、AMI Labsの投資のニュースと同じ時期に報道されたりもしました。なぜよりによって身体を持った機械なのか。ルカンの論理はこうです。画面の中のチャットボットは間違えても大惨事にはなりません。文章を一つ間違えて作り出すだけです。しかし、部屋の中を動くロボットは違います。コップを置く位置を見誤れば、コップは割れます。階段の高さを測り間違えれば転びます。物理世界は失敗を大目に見てくれません。だから、身体を持った機械にとっては、世界がどのように動いているのかについての真の理解、次の瞬間をあらかじめ思い描く能力なしには、一步も安全に踏み出すことができないのです。自動運転車が10年以上経っても完全な無人運転に至っていない理由もここにあると彼は見ました。道路上の機械に足りないのは、より多くのデータではなく、世界の理を込めた頭の中の模型、つまり世界モデルであるということでした。

こうした欠片たちはまだ散らばっており、それぞれ検証の段階も異なります。しかし、方向は一つの場所を指し示していました。機械が世界を観察して常識を蓄積し、その常識の上で次に行くことを自ら計画すること。ルカンが40年以上にわたって抱き続けてきたその考えが、実験室の画面を越えて自動車やロボット、ドローンへと歩み出そうとしているところでした。

彼がこの仕事をオープンにして進めようとするのには理由がありました。ルカンは人工知能の基礎がいくつかの会社の垣根の中に閉じ込められることを、長い間警戒してきました。PyTorchを世に無料で公開し、Llamaをオープンソースとして出したのも、その信念の表れでした。少数の会社が知能の鍵を握れば、残りの世界は彼らが定めた対価を支払い、彼らが許可した分だけを使うことになります。彼は、

世界モデルの骨組みだけは公共のものとして残ることを望みました。貧しい大学の研究所も、小さな国の開発者も持ってきて使える基礎。それが、特定の企業ではなく、様々な国がそれぞれ自らの技術を握ることができるようにする道だと彼は見ました。この考えが彼の個人的な信念なのか、会社の事業戦略なのか、二つは綺麗には分かれませんでした。開かれた基礎の上で応用によってお金を稼ぐことが、AMI Labsが描いた絵でもあったからです。

もちろん、これらすべてが彼の思い通りになるという保証はありません。ルカンは自分が間違っている可能性があることをしばしば口にする人です。世界モデルがいつ使い物になるのか、どれくらいかかるのかについて、彼は数年以上とだけ答えます。そうでありながらも、彼は現在の過熱を警戒しました。2026年6月、彼はイーロン・マスクのxAIについて一種の失敗だと表現しました。創業メンバーが去り、人材を引き留められなかったということでした。いくつかの人工知能会社が耐え難い費用構造の上に立っているとも指摘し、大きなバブルが弾ける可能性があるかと警告しました。サービスの運営費が売上よりも早く増え、投資家のお金で赤字を埋めている構造だという話でした。自分もたった今10億ドルを受け取って会社を設立した人がバブルを語るのは、皮肉に聞こえるかもしれません。しかし、彼の論理は一貫していました。問題は、お金が人工知能に集まること自体ではなく、そのお金が一つの方法にばかり集まっていることだと彼は見ました。誰もが同じ言語モデルをより大きくすることにしがみつけば、その方法が壁にぶつかった日、バブルも一緒に弾けるということでした。他の道をいくつも開いておくことが、彼の考えた安全弁でした。人工知能をめぐる恐怖について、ルカンはヒントン (Geoffrey Hinton) やベンジオ (Yoshua Bengio) と袂を分かちました。共にチューリング賞を受賞した二人の友は、人工知能が人間の統制を離れて人類を脅かす可能性があるかと警告する方に傾きました。ルカンはその恐怖に同意しませんでした。現在の機械は猫ほどにも世界を知らないのに、そのような機械が人類を支配するという想像は早すぎるということでした。彼が見るところ、本当の危険は機械が賢くなりすぎることにありませんでした。人々がまだ未熟な機械を前に、誇張された恐怖と誇張された期待を同時に売っていることにありました。7月には短く一文を残しました。汎用人工知能という言葉における「汎用」を意味するGは筋が通らないと。知能には一つの完成された形態がないという意味でした。人間の知能でさえあちこち穴が開いているのに、何にでも通じる一つの知能という言葉は虚像だということでした。会社名をAGIではなくAMIと名付けたその頑固さが、短い一文として再び飛び出したようなものでした。世界が完成された万能の知能を夢見るとき、彼はそのようなものはないと断言しました。存在するのは多様な知能、それぞれ得意なことが異なる知能たちだけであるということでした。人間の知能も、猫の知能も、いつか出てくる機械の知能も、そのうちの一つに過ぎないという話でした。

ルカンは長年持ち出してきた比喻があります。初めて空を飛んだのは、鳥の羽をそのまま真似た機械ではなく、翼がなぜ浮くのか、その原理を掴んだ機械だったという話です。この本が以前にまるまる一章を割いて扱った内容であり、この後のエピローグが再び引き継ぐ内容です。ですから、ここでその比喻を再び展開することはしません。ただし、2026年のルカンが立っている位置がどのあたりなのかは、指摘しておく価値があります。彼は今、文明全体が羽を真似ている最中だと疑っています。テキストをもっともらしく繋ぎ、画面を写實的に描き出すことに世界中がしがみついています、それが本当に世界を理解することなのか、彼は問い直します。その問いが正しいのか間違っているのかは、まだ誰にもわかりません。彼が正しければ、今の熱風はいつか方向を変えるでしょうし、彼が間違っていれば、世界モデルはもう一つの行き止まりとして記録されるでしょう。どちらの結末も開かれています。

66歳の科学者は、パリの新しいオフィスで若い研究者たちと黒板を埋めていきます。60年前、人工ニューラルネットワークは呪術だと見なされ、誰も信じませんでした。その時代にニューラルネットワークを選んだ少年は、手書きの数字を読み取る機械を作り、アメリカの小切手の文字を読み取り、二度の

冬を乗り越え、チューリング賞を受賞し、指折りの研究所を率いました。そして今再び、誰もが同じ方向へ走るときに、一人で違う方向を指し示します。彼が正しいのか間違っているのかは、この本は答えることができません。答えはまだ書かれていません。ただ、一つの問いだけははっきりしています。知能とは世界を描き出すことなのか、それとも世界を理解することなのか。図書館の少年が初めて抱いたその問いが、60年の時を経て再び私たちの前に置かれています。

弁護士

kimkj.com

エピローグ ライト兄弟は鳥を模倣しなかった

ヤン・ルカン、人工知能の父

この本をここまで読み進んでこられた読者にとって、すでに見慣れた情景があるはずです。一人の男が、他の誰とも異なる方向を指しながら立っている情景です。1989年ホルムデルの夜、画面に歪んだ3が映ったとき、彼はそのように立っていました。そして2026年パリの新しいオフィスの黒板の前でも、同じように立っています。60年の間に、世界は彼に対して2度大きく身を翻しました。最初は彼が掲げた神経網を死んだ道と呼んで背を向け、後に同じ神経網を世界の中心に据えて彼を師として仰ぎました。そして今、世界は3度目に彼を超えて、別の方向へと駆け去ろうとしているのです。

一人の人生を全て記してもなお答えの出ない問いがあります。多くの人が同じ方向に走るとき、少数の道はいかにして守られるのか。ルカンの物語はこの問いに明快な答えを与えてはくれませんでした。彼が最初の冬を耐えることができたのは、信念だけのゆえではありませんでした。ベル研究所が予算と時間を与えてくれたからであり、ホルムデルの上司ラリー・ジャッケルが、金を節約することで有名になることはないと言いながら背後から支えてくれたからです。2度目の冬に論文が何度も拒否された時代、彼は一人で耐え抜いたわけではありません。ヒントンとベンジオを含む何人かが、カナダの支援を背景に非公式な結社を結び、ユタの山間のワークショップで互いの生存を確認し合いました。少数の道は純粋な意志だけで守られるものではありません。その道を歩む人のそばに、今のところ有用性が証明されていない研究に対して資源を提供する人が必ずいなければならないのです。

この箇所が今の私たちに残すものがあります。2012年AlexNetがImageNetコンペティションを揺るがした日、世界は神経網が20年間異端視されていたという事実をたちまち忘れました。勝利した技術の歴史は常に滑らかに書き直されるからです。あたかも最初からそれが正解であったかのようにです。しかし、その20年の冬に神経網を握り続けていた人間は、十指に足りるほどでした。もしあの時代に誰も彼らに実験室と予算を与えていなかったなら、今日の春は来なかったか、あるいはずっと遅れて訪れていたでしょう。多数が正しく見えるときに少数の実験を生き永らえさせることは、口マンではなく文明の保険なのです。

ルカンは自らが受けたものを返す方法も心得ていました。ザッカーバーグの食卓で研究所の経営を任されるにあたって、彼が提示した条件は成果を全面公開することでした。そうして生まれたオープン型の研究所と、そこから生まれて事実上の標準となったPyTorchというツールは、次の世代の少数が冬を越すために必要とされる実験台でもあります。資源を握る者がそれを閉ざさず開き続けること。少数の道を守る第2の方法がそこにありました。

ルカンが今歩んでいる3番目の道が冬で終わるのか、それとも春で終わるのか、誰も知りません。彼は今、世界が羽を模倣しているのではないかと疑っています。テキストをもっともらしくつなぎ、画面を現実的に描き出すことに力の全てを注いでいますが、それが本当に世界を理解することなのかと自問しています。空を初めて飛んだのは鳥の羽をそのまま模倣した機械ではなく、羽がなぜ浮力を生じるのか、その原理を追究した機械だったという、彼が長年聞き続けてきた言葉がここで再び頭をもたげます。この比喻が正しいかどうかは、この本ですら判定することはできません。彼が正しければ、今の熱狂はいつか方向を変えるでしょう。彼が間違っていれば、世界モデルはさらにもう一つの袋小路として記録されるでしょう。どちらの結末も依然として開かれています。

科学の冬を耐え抜く方法について彼に問うなら、おそらくこのような答えが返ってくるでしょう。流行の気温に自分の確信の気温を合わせないことです。2016年コレージュ・ド・フランスの講壇でも、2026年アメリカの経済放送のスタジオでも、彼の話し方は変わりませんでした。世界が神経網を嘲笑していたときも、世界が大規模言語モデルに熱狂していたときも、彼は声を張り上げることはありませんでした。昔からずっと同じ場所を守り続けてきた人の、疲れた確信に近い話し方です。冬を耐える人とは、春が来ると叫ぶ人というよりも、春が来ようが来まいが自分の実験をもう一日回す人なのかもしれません。

しかしそこには代償が伴いました。チューリング賞を受けた後、彼は優先権をめぐる紛争の真ん中に置かれ、自分の名前に付随した栄光の分配について長く争わなければなりませんでした。メタを去るころには、12年間をかけて育てた研究所が製品中心に再編されるのを目の当たりにしなければなりません。反対側に立つ者にとって、冬を耐えるだけでは十分ではありません。春が訪れた後、利益を分かち場所においても心は安まりません。彼の人生を最後まで辿った読者なら、その居心地の悪さが彼の顔にどのように刻み込まれていたのか、もうおわかりになるでしょう。

この本が彼が正しい人間だと主張しているわけではありません。彼の予測は外れることもあります。大規模言語モデルを低く評価した彼の診断が度を超していたと後に判明することもあり得ますし、世界モデルが彼が描いたほどの地平を切り開かないことも考えられます。彼を批判する人々の論拠も、この本の随所にそのまま記しておきました。けれども一つだけ残ります。彼は多数が気楽に流れるとき、反対側に立つことを恐れませんでした。そしてその代償として、60年の間に2度の冬と1度の誤解を甘受してきました。そのような人が今なお生き、66歳という年齢で若い研究者たちとともに黒板を埋めています。

図書館の少年が最初に抱いていた問いがありました。知能とは生まれつきの法則なのか、それとも学びながら身に付く何かなのか。60年が経ち、その問いは別の姿をして再び私たちの前に現れます。知能とは世界をもっともらしく描き出すことなのか、それとも世界がなぜそのように動くのかを理解することなのか。この問いの答えは、ルカンも、この本も、今人工知能に国の予算を賭けているすべての人も、いまだ握ってはいません。

ですからここに最後の一場面だけを残しておくことにしましょう。パリのある部屋、黒板の前。白いチョークを握った老学者が、若い顔たちに向き直ります。60年前に図書館で忘れられた機械を掘り起こしていた少年の眼差しが、そこにそのまま残っています。彼の隣に座った20代そこそこの研究者は、今、多数派ではない方に自分の20代を賭けています。その選択を生き永らえさせ、実験を続けるための実験室と予算がこの時もあるのかどうか。その問いに答えるのはルカンではなく、私たちなのです。あなただったら、まだ有用性が証明されていないその黒板の前に、一つの席を用意されるでしょうか。

弁護士

kimkj.com

このAI書籍を応援する

この本が役に立った場合は、今後の翻訳、公開版、
AI知識プロジェクトをPayPalで応援できます。



PayPal

<https://paypal.me/kimkjkj>

QRコードを読み取るか、上のリンクを開いてください。