

AI創薬の パラダイム転換と 未来の融合エコシステム

データ主導の標的探索から自律駆動ラボまで

標的探索、創薬設計、臨床検証、そして規制まで

キム・ギョンジン 弁護士

目次

☒1☒ AI創薬の幕開けとパラダイム転換

- 1.1 R&D パラダイムシフト：死の谷（Valley of Death）の克服
- 1.2 生命科学のための人工知能: ファウンデーションモデル(Foundation Models)
- 1.3 多重オミクスの統合と新規バイオマーカーの発見

☒2☒ AI主導の創薬設計と物質探索

- 2.1 3Dタンパク質構造予測と構造ベース薬物設計 (SBDD)
- 2.2 ディープラーニングによる候補物質の生成 (De Novo Design)
- 2.3 インシリコ(In silico) ADMET及び毒性予測システム

☒3☒ 次世代モダリティとスマート薬物送達システム

- 3.1 バイオ医薬品および多重標的治療薬の設計 (ADC)
- 3.2 天然物新薬データベースのAI解釈とハイブリッド発掘
- 3.3 ナノ医学およびスマート標的送達システム (Nanomedicine)

☒4☒ 非臨床・臨床検証と精密医療の革新

- 4.1 前臨床研究のイノベーション：動物実験の代替および生物学的モデリング
- 4.2 デジタルツイン(Digital Twins)と次世代臨床試験の現代化
- 4.3 精密腫瘍学(Precision Oncology)と患者に合わせた治療

☒5☒ 規制科学、データ倫理、そして未来の融合エコシステム

- 5.1 AI新薬開発の規制フレームワークと品質管理ガバナンス
- 5.2 データ倫理、アルゴリズムの公平性、および連合学習エコシステム
- 5.3 医薬品業界の未来：自律化研究所(Self-driving Labs)

第1部

AI創薬の幕開けとパラダイム転換

1.1

R&D パラダイムシフト：死の谷（Valley of Death）の克服

2022年2月、特発性肺線維症の治療薬候補物質ISM001-055が初めて人間の体に投与されました。肺が徐々に固くなるこの病気は適切な治療法がなく、診断後3年から5年以内に患者の大部分が亡くなります。肺の組織が硬い傷跡に変わりながら呼吸がだんだん苦しくなりますが、正確な原因さえ明らかになっておらず、名前の通り特発性という言葉がついています。それまでこの病気の進行を遅らせることができる薬は世界で2種類しかなく、どちらも進行をある程度遅らせるだけで止めることはできませんでした。この物質を作ったのは香港に本社を置く Insilico Medicineです。研究陣は標的発掘プラットフォームPandaOmicsで線維化を引き起こすシグナル伝達過程を調節する酵素であるTNIKを新しい標的として見つけ出し、化学分子生成エンジンChemistry42でその酵素にぴったり合う分子構造を設計しました。

プロジェクトを始めてから30ヶ月も経たないうちに起こったことです。伝統的な方式であれば、標的を探して候補物質を一つ絞り出すためだけに4年から6年かかったはずの区間です。後にLentosertibという正式な名前を受け取るこの小さな分子は、このようにコンピューターの画面の中で生まれ、人間の血管の中へと歩いて入っていきました。それから3年後、この物質は業界全体の注目を再び集めるニュースの主人公になります。

新薬を作ることは、長い間、渓谷を挟んだ力比べのようでした。渓谷のこちら側には実験室で有望視される候補物質が数万個も立ち並び、渓谷のあちら側には患者に実際に処方される数粒の薬が置かれています。その間を渡ろうとして落ちる候補物質が圧倒的に多いため、業界はこの区間を死の谷と呼んできました。伝統的な新薬開発は、標的を見つけた瞬間から規制当局の承認を受けるまで平均10年から15年かかり、失敗した候補物質まで全て合わせた費用は26億ドルを超えました。

この金額には、成功した新薬一つの開発費だけでなく、一緒に試みられて失敗した数多くの候補物質の埋没費用まで含まれています。それなのに、渓谷を渡って市場に出る物質は最初の候補群の10パーセントにも満ちません。残りは効能が不足していたり、予想外の毒性が現れたり、体内で予想と異なって分解されて抜け出したりして、途中で座り込んでしまいます。

渓谷がこれほど深くなったのには理由があります。新しい薬はまずネズミやサルのような動物で試します。ところが、動物の体と人間の体は同じ物質を互いに異なって受け入れます。ネズミの肝臓では無難に分解されていた成分が人間の肝臓では毒性物質に変わることもあり、サルで確認された効果が人間には半分も現れないこともあります。さらに、同じ人間でも年齢や遺伝子、すでに患っている他の病気によって同じ薬に異なる反応を示します。動物実験の結果を人間に移して適用しようとするこの橋渡し研究の過程自体が、生まれながらの限界を抱えているのです。

小学生でも分かるような例えで説明すると、子犬に通用した訓練法を猫にそのまま使ってみる状況と似ています。見た目は似ているように見えても体内の反応は種ごとに異なるため、動物実験を通過した候補物質の大部分が、人間を対象とした最初の臨床試験の段階でそのまま崩れ去ります。業界ではこのような失敗を指して、ネズミには通用したが人間には通用しなかったという意味で種間障壁という言葉を使います。

この失敗の繰り返しは統計にもはっきりと表れています。新薬一つを承認してもらうためにかかる費用はおよそ9年ごとに2倍に跳ね上がってきた反面、同じ投資額に対して承認される新薬の数はむしろ減りました。研究者たちはこの流れに Eroom's Law という名前を付けました。半導体の性能が着実に良くなるという Moore's Law の綴りを逆さまにひっくり返して作った言葉です。Moore's Law が技術発展の象徴であるなら、Eroom's Law は製薬産業が半世紀以上抜け出せずにいる生産性退歩の象徴です。

専門家たちはその原因として、臨床試験の規模がますます大きくなり規制のハードルが高くなった点、すでによく効く薬が開発しやすい病気の大部分を占めてしまい、残った病気ほど攻略が難しくなった点を挙げます。新しく登場する病気の

標的ほど構造が複雑で扱いにくいものばかりが残っているという指摘もあります。お金はますます多くかかるのに成果はますます減る状況が長く続き、製薬会社はリスクの大きい希少疾患や難治性腫瘍の研究に気軽には飛び込めませんでした。

この流れをひっくり返す鍵として浮かび上がったのがデータです。ゲノムやタンパク質の情報、細胞一つ一つの反応を記録した単一細胞データ、患者の電子健康記録や画像資料、毎日あふれ出る数千編の医学論文まで、人間の頭では読み切れない情報が蓄積されてきました。この情報を人間が一つ一つ読んで互いに結びつけてみることは、最初から可能な規模ではありません。以前の新薬開発は、研究者の直感と大量スクリーニング、つまり数万個の化合物を無作為に試験してみて偶然引っかかる物質を探す方式に大きく依存していました。

人工知能はこの膨大なデータを学習し、病気の原因や薬物の反応をあらかじめ見通す方向へと研究方式自体を移しています。標的を無作為に探っていた探索がデータを根拠とする予測に変わったことで、新薬開発の最初の段階を支配していた試行錯誤の割合が目立って減りました。

二つのパイプラインを並べてみると、違いがさらに鮮明になります。伝統的なパイプラインでは10年から15年という時間と20億ドルを超える資本が入り、化学者と生物学者が一段階ずつ手で確認する順次的な検証を経て、その終わりには10のうち9は失敗に終わります。人工知能が導くパイプラインでは、この期間が半分にも満たないほどに減り、初期探索に入る費用が大幅に下がり、様々な段階が同時に回る構造の上で生成型人工知能が候補自体を絶えず整えていきます。

時間とお金を減らすことにとどまらず、最初から失敗する可能性が高い候補を序盤からふるい落とす方向へと、アプローチ自体が変わったのです。失敗を後から高い対価を払って確認する代わりに、安く素早くあらかじめ確認する方向へと重心が移りました。

伝統的な新薬開発パイプラインは、標的発見と化合物探索、最適化、前臨床試験、臨床第1相・第2相・第3相が一行に並んだ順序をそのままに従っています。前の段階が完全に終わらないと次の段階に進めない構造であるため、どの一つの段階でも時間が遅れると、後に続くすべての段階と一緒に遅れます。人工知能が入っ

たパイプラインはこの順序を大きく乱します。標的を見つける作業と候補物質を設計する作業、その物質が体内でどのように動くかを予測する作業が、複数のコンピュータで同時に実行されています。

生成型人工知能が数千種類の分子構造の代替案を一度に作り出す間に、別の研究チームは成功の可能性が高い候補の毒性を事前にシミュレーションします。一列に並んで順番を待っていた作業が複数の場所で同時に進行する作業に変わりました。この並列構造が開発期間を半分以上に引き下げた実質的な推進力です。

効果は数字に表れます。新しい疾患標的を見つけ、そこに適合する最初の化合物を選ぶのに従来は数年かかりましたが、人工知能プラットフォームを導入した生命工学企業はこの期間を12ヶ月から18ヶ月の間に短縮しました。確認のために実際に合成する必要がある化合物の数も数万個から数百個レベルに大幅に減らしました。世界的な経営コンサルティング機関は、人工知能を新薬開発全過程に導入すれば、医薬品業界全体に毎年600億ドルから1100億ドルに達する経済的価値が新たに生まれるだろうと予想しています。この価値の相当部分は、候補物質が遅い段階で脱落する際に一緒に消える埋没費用を減らすところから生じ、残りは開発期間が短縮されて特許で保護される期間内に医薬品をより長く販売できるようになったところから生じています。

インシリコメディシンナー社だけを見ても、標的発見から前臨床候補確定まで13ヶ月から18ヶ月かかり、ここにかかった費用は260万ドル程度でした。従来の方式であれば同じ区間に2年半から4年、数千万ドルがかかったはずで、2025年一年だけを見ても、人工知能新薬開発企業にかかったベンチャーキャピタル資金が3四半期まで27億ドルに達しました。資本がこの流れに先に乗ったということです。

このような削減効果は、業界全体の計算方法を変えてしまいました。以前は患者数が数千人に留まる稀少疾患に10年以上の歳月と20億ドルを超える資金を投資する名分を見つけるのが困難でした。失敗確率がとても高いため、ごく少数の患者のためにそのようなリスクを冒す医薬品会社は稀でした。しかし、初期発見費用が数百万ドルレベルに低下し、失敗をはるかに早い段階で除外できるようになったため、それまで顧みられなかった稀少疾患と小児がん、抗生物質耐性菌の治療

薬研究に再び資金が流入しています。

このような流れが重要な理由は、統計や株価のためだけではありません。特発性肺線維症のように患者数が少なく新薬開発の優先順位から外れていた疾患が、初期開発費が低下したおかげで、研究者の机の上に再び上がってきています。谷を渡る橋が広くなるほど、その橋を渡って患者に届く治療薬の種類も増える可能性が大きいです。

生成型人工知能という言葉は、正解一つを見つけ出すだけでなく、条件に合う新しい成果物を自ら作り出す人工知能を指します。絵を描いてほしいと言えば新しい絵を、文章を書いてほしいと言えば新しい文章を作り出すのと同じ原理で、新薬開発では標的に合う新しい分子構造を作り出します。ファンデーションモデルは特定の疾患一つだけを学んだモデルではなく、遺伝子とタンパク質と化合物と論文まで広範なデータを学習して複数の問題に一緒に使われる基礎モデルを意味します。

一つのモデルを作っておけば、標的探しと分子設計、毒性予測まで複数の作業に分けて使うことができるため、最初からモデルを新たに作るよりも効率ははるかに高いです。

この二つの技術が実際に新薬を設計する方式は、鍵の話で説明すれば理解しやすいです。以前の仮想スクリーニングは、既に存在する数百万個の化合物ライブラリの中から標的に合う鍵を一つずつ当ててみて探す作業でした。一方、最近の生成型人工知能は錠前、つまり標的タンパク質の立体構造を見た後、その穴にぴったり入る鍵を全く新しく削り出して作ります。研究者たちは生成的敵対的ニューラルネットワークと拡散モデル、グラフニューラルネットワークのような技術で、医薬品の結合力と選択性、体内での吸収と分解と排出のプロセスまでコンピュータで事前に描写します。

このように化合物を最初から新たに設計する方式をデノボ設計と呼びます。新薬候補発見が既存物質を探す探索の領域を脱出し、新しい物質を建てる創造の領域に移行したということです。世界のどこにも存在しなかった分子でも標的にぴったり合いさえすれば候補になりうるという点で、そもそも見つけることができる

候補の幅自体が広がりました。

このような設計方式がどのくらい速いかを示した初期の事例があります。イギリスの人工知能新薬開発企業Exscientiaは、日本の住友大日本製薬と協力して強迫性障害治療候補物質DSP-1181を設計しました。セロトニン受容体の一種を狙って設計されたこの物質は、探索研究を開始してからわずか1年未満で完成しましたが、従来の方式であれば平均4年半かかったはずです。

2020年1月、DSP-1181は日本で人間を対象にした最初の臨床試験に入りました。コンピュータ設計が始まってから12ヶ月で起こったことであり、当時の医薬品業界はこのニュースに大きく動揺しました。しかし、この物質はその後開発が中止されましたが、谷の入口を速く通り過ぎたからといって反対側の岸まで無事に到達するという保証はないという事実を後になって示した事例でもあります。

構造を読み解く人工知能も共に発展しました。タンパク質の3次元構造を原子単位で予測するアルファフォールド（AlphaFold）シリーズは、標的とするのが難しいと考えられていたタンパク質の隠れた結合部位まで見つけ出し、分子ドッキングシミュレーションの速度を大きく引き上げました。アルファフォールドを開発したデミス・ハサビスとジョン・ジャンパーは、この功勞により2024年のノーベル化学賞を受賞しました。半世紀以上もの間、生物学の難問として残っていたタンパク質の構造予測問題が人工知能によって解かれたことで、新薬設計の最初のステップである標的構造の把握にかかっていた時間が数か月から数日に短縮されました。Googleが発表した治療薬特化型言語モデルTx-LLMは、分子構造と標的タンパク質、臨床文献を1つのモデルの中で共に学習し、66の新薬開発評価タスクのうち22のタスクで最高水準の性能を記録しました。

2025年4月には、Googleがこのモデルを発展させたオープンソースモデルTxGemmaを新たに発表しましたが、同じ66のタスクのうち50のタスクで、従来の最高水準に匹敵するかそれを上回る成績を出しました。ゲノム全体を扱うファウンデーションモデルEvoもまた、単一細胞の反応とタンパク質間の相互作用を共に読み解き、標的発掘の過程を自動化しています。

このようにコンピューターの中で設計された分子は、今では人の手をほとんど經由せずに実際の物質として作り出されます。ロボットアームが試薬を運び、反応条件を調節する自動化実験室では、人工知能が提案した数百種類の合成経路を昼夜を問わず試します。合成された物質の効果を測定した結果は直ちに次の設計に反映され、設計と合成と測定が人の介入なしに繰り返される閉ループシステムを構成します。標準化されたロボットが実験を代行することで、人によって異なっていた手技のばらつきが減り、同じ条件で実験を繰り返したときに結果が食い違う問題も共に減少します。

研究者が直接試薬瓶を持ち、徹夜で実験をしていた風景は、人が目標だけを定め、機械が一晩中数百回の実験を代わりに回す風景へと変わっています。このような自律化研究所は、並列処理の範囲をコンピューターの画面の外、実際の実験台の上まで広げています。人が眠っている間にも実験が止まることなく行われるという点で、1日8時間余りだった実験室の実質的な稼働時間自体が大きく伸びました。

しかし、コンピューターの画面上で完璧に見える設計図が実際の建物になるには、依然として長い工事期間が必要です。インシリコシミュレーションで洗練された仮想の化合物であっても、生きている体の中に入れば予想外のオフターゲット反応や肝毒性、心臓毒性が現れる可能性があり、体外細胞実験や動物実験、オルガノイドを利用した前臨床検証を経なければなりません。

人工知能が候補物質を探す期間を数年から数か月に短縮してくれたとしても、人の体で安全性と効能を確認する第1相、第2相、第3相臨床試験は、患者の命に関わることであるため、飛ばすことはできません。この検証には短くても数か月、長ければ数年がかかります。谷の幅を狭めることはできても、谷自体をなくすことはできないということです。

前臨床検証は、大きく分けて5つのことを確認する過程です。薬物が体の中でどのように吸収され、どの臓器に広がり、肝臓でどのように分解され、体外へどのように抜け出るのか、そして予想外の毒性を引き起こさないかということをまとめて、吸収・分布・代謝・排泄・毒性と呼びます。過去にはこの5つを一つひとつ動物実験と細胞実験で確認しなければなりませんでした。今では人工知能モデルが分子構造を見ただけでも近い予測値を先に出してくれます。

この予測が実際の実験を完全に代用することはできませんが、そもそも成功の可能性が低い候補を実験台に上げる前にふるい落とせるというだけでも、研究者が費やす時間と動員される実験動物の数を共に減らしてくれます。

谷の中で候補物質が崩れる形も様々です。ある物質は培養皿の中の細胞では標的によくくっつきましたが、いざ人の体の中では望む組織まで到達できずに散ってしまいます。ある物質は肝臓でとても早く分解され、薬効が現れる前に体外へと抜け出ます。また別の物質は、標的を一つだけ触ろうとしたのに構造が似ている他のタンパク質まで一緒に触れてしまい、予想外の副作用を引き起こします。

人工知能は、このような失敗のパターンそれぞれに対応する予測モデルを別々に置くことで、候補物質がどの段階で崩れる可能性が高いのかを、初期の設計段階から推し量れるようにしてくれます。以前は実験室で何週間もかかってようやく明らかになっていたこのような弱点を、今では分子構造を描く段階で先に気付くケースが増えています。

この事実を痛烈に確認させる事例が2023年に出ました。住友ファーマが大塚製薬と共に開発した統合失調症治療の候補物質ウロタロントは、PsychoGenicsの人工知能基盤の行動分析プラットフォームによって発掘された物質です。従来の統合失調症薬物とは異なり、ドーパミン受容体ではない別の神経伝達物質受容体を狙うことで副作用を減らすという期待を集めました。ところが2023年7月、ダイヤモンド1とダイヤモンド2という名前が付けられた2件の第3相臨床試験において、ウロタロントはプラセボよりも症状をはっきりと改善することはできませんでした。

安全性自体には問題がありませんでしたが、いざ薬効を証明しなければならない最後の関門で崩れてしまったのです。新型コロナウイルス感染症の流行期と重なり、プラセボを投与された患者群の症状までもが例年より大きく改善した点が、結果の解釈をさらに難しくしたという分析も出ました。人工知能がどれほど精巧に予測したとしても、個別の患者の反応やプラセボ効果のような臨床現場の変数までを完璧に統制することはできないという事実を、この失敗は示しています。

臨床段階で起こるこのような不確実性を減らすために、業界が注目する解決策がデジタルツインと適応型臨床試験デザインです。デジタルツインは、ある患者の遺伝的、病理学的特性をコンピューターの中にそのまま移した仮想の双子を作る技術です。実際に薬を投与する前に、この仮想患者に様々な用量や治療シナリオをあらかじめ適用してしながら、腫瘍の異質性や薬物耐性がどのように変わるのかをシミュレーションします。

このようなアプローチは、患者を細かく分け、良い反応が予測される人に先に薬を試す患者層別化を精巧にします。実使用データを基に作った仮想の対照群を使って、臨床試験の規模と期間、倫理的負担を同時に減らす試みも増えています。病院に通わなくても参加できる分散型臨床試験も、こうした流れの上で増えています。

臨床試験が遅れるよくある原因は、参加者を時間通りに集められないことにあります。全体の臨床試験の80パーセント以上が、患者募集の遅れによって日程が遅れるという集計があるほどです。人工知能は、電子健康記録とゲノムデータを調べて条件に合う潜在的な参加者を事前に見つけ出し、試験の設計段階から募集しやすい基準を提案する方法で、このボトルネックを減らしています。

このような需要に後押しされ、臨床試験に特化した人工知能の市場規模は、2024年の28億ドルから2032年には540億ドル以上に大きくなると見込まれています。谷を渡る橋の幅を広げることと同じくらい、橋の上に人を時間通りに集めることもまた、新薬開発の速度を左右します。

そして2025年、実際に谷を渡った初めての事例が出ました。2022年に臨床第1相に入ったISM001-055、現在はLentosertibという名前を受けたこの物質の臨床第2a相の結果が、その年の6月3日に国際学術誌のNature Medicineに発表されました。企業の報道資料ではなく、査読を経た学術誌に結果が載ったということは、独立した専門家たちの検証を通過したという意味です。特発性肺線維症の患者71人を対象に12週間にわたって進められたこの試験は、医師も患者も誰が本物の薬を受け取ったのか分からないようにする二重盲検方式に、偽の薬を比較群として置くプラセボ対照方式を加えたものでしたが、Lentosertibはすべての投与量で安全であり、容量が増えるほど肺の機能がはっきりと良くなりました。

1日60ミリグラムを投与された患者群は、肺活量の指標が平均98.4ミリリットル増えた一方で、プラセボを投与された患者群はむしろ20.3ミリリットル減りました。この結果は同じ年のアメリカ胸部学会の学術大会でも一緒に発表されました。コンピューターの中で生まれた標的と分子が、実際の患者の体内で目で確認できる効果につながった、業界初の公開検証の事例です。

Lentosertib一つだけの話ではありません。2026年初めを基準に臨床段階に入った人工知能基盤の新薬候補は173件に達しますが、2023年末の24件と比較すると3年足らずの間に7倍以上増えました。様々な産業分析によると、人工知能がを見つけ出した物質は臨床第1相の通過率が80パーセントから90パーセントに達し、伝統的な方式の40パーセントから65パーセントを大きく上回ります。臨床第1相の目的は効能ではなく安全性の確認にあります。人工知能が設計段階から毒性が予想される構造を取り除き、体内の吸収と排出の過程をあらかじめシミュレーションしたおかげで、最初から安全性のハードルを越えやすい候補だけが臨床試験の舞台に上がるものと見られます。

ただし、臨床第2相に移ると格差が縮まります。これまでに第2相を終えた人工知能が発掘した物質10件のうち4件が成功して40パーセント前後を記録しましたが、これは業界平均である30パーセントから40パーセントとほぼ同じ水準です。臨床第2相からは、薬の設計の精巧さよりも、人によって異なる遺伝的背景や生活習慣、併存疾患のように、設計段階でコントロールすることが難しい変数の比重が大きくなるためだと見られます。規制当局の最終承認まで受け取った人工知能設計の新薬は、2026年の夏現在までまだありません。

谷の入り口は目立って広くなりましたが、谷の終わりの崖は依然として険しく残っているという意味に見えます。

制度もまた、この流れを後ろから後押ししています。アメリカの食品医薬品局は、2022年12月に食品医薬品局近代化法2.0を通過させ、新薬の承認申請に動物実験を必ず含めなければならないという数十年来の義務条項をなくしました。2025年4月にはモノクローナル抗体を皮切りに、動物実験をオルガノイドと人工知能基盤のシミュレーションに置き換えていく具体的なロードマップを出しました。後に続く食品医薬品局近代化法3.0は、2025年12月に上院を通過し、2026年7月20日に

は下院まで通過しました。

下院で修正された案を上院がもう一度議決しなければ大統領に渡らない仕組みであるため、2026年8月の現在、まだ署名は行われていません。この法律が発効すれば、新薬の承認規定のあちこちに埋め込まれていた動物実験という文言が消え、人の細胞で作った爪ほどの大きさの装置の上に肺や肝臓のような臓器の機能を再現する臓器チップや、コンピューターシミュレーションのような非動物試験がその文言に代わることになります。

大手製薬会社も先を争ってこの流れに合流しました。AlphaFoldを作ったGoogle DeepMindの子会社であるIsomorphic Labsは、Eli Lilly、Novartisとそれぞれ数十億ドル規模の新薬共同開発契約を結び、様々な疾患の標的に対する分子設計を一緒に進めています。Sanofiは研究開発の意思決定の過程全般に人工知能を結合させるという目標を掲げ、組織の構造自体を組み直しました。

同じ時期、Novo NordiskやAstraZenecaをはじめとする他の大手製薬会社も、自社の人工知能組織を新設したり専門企業と手を組んだりして、似たような再編に乗り出しました。人工知能の新薬開発の専門企業が独自に新薬を設計する方式とは別に、既存の大手製薬会社が自社の研究開発組織の中に同じツールを移植する流れも一緒に大きくなっています。

資本市場の反応もはっきりしています。Insilico Medicineは2025年末に香港証券取引所に上場し、22億7700万香港ドル、米ドルで約2億9200万ドルを調達し、その年の香港のバイオ企業の中で最大規模の新規株式公開を記録しました。公募の申し込み競争率は1427対1を超え、Eli LillyとTencent、Temasekを含む15の機関がコーナーストーン投資家として名前を連ねました。韓国の中でも似たような動きが速くなっています。人工知能基盤の新薬開発企業であるGaluxは、2026年2月にYuanta Investment、Hanguk Saneop Eunhaeng、Hanguk Tuja Jeunggwon、Mirae Asset Jeunggwonなどが参加した420億ウォン規模のシリーズB投資を誘致しました。

ブリーズバイオとファーストバイオセラピューティクス、エリシジェンのような企業も300億ウォンを超える投資を相次いで獲得し、企業公開に向けた基盤を整

えています。オンコビックスは自社の生成型人工知能プラットフォームで低分子標的治療薬を設計しており、アロンティアはセルトリオンと手を組み、新たなタンパク質構造基盤の候補物質を発掘しています。これらの企業のほとんどはまだ臨床初期段階に留まっていますが、投資規模と参加機関の顔ぶれを見ると、この分野に向けた国内資本の視線が以前と異なったものに変わったという信号として読まれています。

業界は2026年と2027年を本当の試験台と見ています。これまでは速度と初期段階の通過率を確認する時間だったとすれば、今後やってくる数年は、先んじたパイプラインたちが臨床3相という最後の関門の前で成績表を受け取る時間です。この成績表に従って、今後10年間、製薬会社たちが研究開発費をどこに配分するかが大きく変わることが予想されます。これは韓国の新興企業たちにも同様にやってくる試験台です。

レントセルティブが次の段階である臨床3相でも同じ効果を再現できるかどうか、そしてその結果が規制当局の最終承認につながるかが、人工知能新薬開発全体の信頼を測る物差しとなると見られています。死の谷は消えていませんが、その上を通る橋は今ようやく最初の通行を終えました。

韓国政府もこの流れに合わせて支援体系を整えていっています。保健福祉部と科学技術情報通信部は韓国製薬バイオ協会を事業主管機関とし、2024年4月から2028年12月まで348億ウォンを投じてK-Melody事業を進めています。ヨーロッパのMelodyプロジェクトを模範とするこの事業は、企業ごとに別々に蓄積した新薬候補物質データを一箇所に集めない代わりに、データは各機関にそのままとおき、人工知能モデルだけを移し回りながら学習させる連合学習方式を使っています。

製薬会社の立場からすると、失敗した実験記録が競争相手に漏れていくのを嫌う敏感な情報ですが、データを外に持ち出さなくても一緒に学ぶ道をこの方式が開いてくれたわけです。事業団を率いるKim Hwa-jong団長はあるインタビューで、成功した事例より失敗した事例がより多く蓄積されてこそ、人工知能が次の失敗を事前に選別できると述べています。失敗もデータということです。

食品医薬品安全処も2024年6月に医薬品開発時の人工知能活用ガイドを発表し、新薬審査過程で人工知能が作成した資料をどのように扱うべきかについての方向性を提示しました。Daewoong Pharmaceuticalは保健福祉部のK-AI新薬開発前臨床・臨床モデル開発事業に共同研究機関として参加し、抗がん薬と代謝性疾患治療薬を研究しながら蓄積してきた非臨床データを4年にわたって人工知能モデル学習に提供することにしました。

この企業は化合物分子モデル8億種を含む自社データベースを基盤に、新薬候補物質を探すプラットフォームDaisy-doを別途作って運営しています。新興ベンチャー企業だけでなく、長い歴史を持つ製薬会社も同じ道具を手にし始めた、ということです。

それでも、人工知能が設計した新薬のうち規制承認にまで至った事例がなぜまだないのかと疑問に思うかもしれません。最も先に臨床に入ったDSP-1181さえ2020年によく人間に初めて投与され、レントセルティブも2022年によく最初の臨床試験を始めました。臨床1相から3相、規制審査まで経るのに要する期間をそのまま当てはめると、今最先にある候補物質であっても、早くて2027年、遅ければ2028年にならないと結果が出ることになります。

人工知能が設計した新薬が特に多く落とされているのではなく、承認にまでつながるだけ十分な時間がまだ経っていない可能性が大きいです。計算機をいくら速く回しても、患者の体が薬に反応するのに要する数年の時間を早めることはできません。今後数年のうちに出来る臨床3相結果が、この判断が正しいのかを実際に確認してくれると見られています。

1.2

生命科学のための人工知能: ファウンデーションモデル(Foundation Models)

2025年2月、米国カリフォルニアの非営利研究所アーク・インスティテュート（Arc Institute）のコンピュータの前で、研究者たちは人間と米、大腸菌と古細菌を区別しない膨大な遺伝子配列が毎日休みなくモデルへと流れ込んでいく光景を見守っていました。画面に映る遺伝子の名前は馴染み深いものではなく、アデニン、グアニン、シトシン、チミンという4文字だけでできた配列でしたが、その中には人間の心臓病から微生物の抗生物質耐性まで、あらゆる生命現象の法則が混在していました。

彼らが作っていたプログラムは、特定の病気ひとつ、特定の生物ひとつを対象とした計算機ではありませんでした。わずか数年前まで生物学研究室で使われていた人工知能は、遺伝子ひとつ、標的タンパク質ひとつに執着する狭い道具に留まっていたため、種を区別しないで一つのモデルで答えを出そうというこの試み自体が異例でした。人間の遺伝病であれ、植物の干ばつ耐性であれ、産業用微生物の酵素設計であれ、質問が異なっても同じモデルひとつが答えを出すようにすることが目標でした。

医薬品を製造する人々にとって、ゲノムはこれまで参考資料程度にしか使われておらず、次の文を予測するように全体として計算して扱える対象だと誰も簡単には考えていませんでした。数週間後、世界に公開されたこのモデルの名前はEvo 2でしたが、その登場は医薬品と生命科学の分野で人工知能が歩んできた道の方角を明らかに変えてしまいました。

キムチチゲだけ上手に煮込める料理人と、包丁遣い、火加減、味加減という基本技術をまんべんなく習得した料理人は、新しい注文の前で結果が分かります。前者はキムチチゲは素晴らしく煮込みますが、パスタを注文されると手を放してしまいます。基本技術を習得した後者は、冷蔵庫の中の材料だけを見ても、スープであれ炒め物であれ、すぐに新しい料理を作り出します。初めて見る食材が入っ

てきても、後者は困惑せず、基本技術を応用して尤もらしい味を出します。ファウンデーション・モデルという言葉が指す人工知能は、後者の方に近いです。以前の人工知能プログラムは前者に似ていて、特定のタンパク質と特定の化合物の結合力を点数で評価するように訓練されると、その計算ひとつだけをこなしました。

同じプログラムにこの化合物が肝臓に毒性を引き起こすかどうかを尋ねると、答えを出すことができず、最初から別のプログラムを新たに学習させなければなりませんでした。ファウンデーション・モデルは、膨大なデータから生命現象の基本文法をまず習得した後、標的（医薬品が付着すべき体内タンパク質）の予測であれ、新薬候補物質の設計であれ、臨床試験結果の推定であれ、新たに投げかけられた質問に合わせて実力をその都度取り出して使います。このような差異のため、医薬品企業は今では部門ごと、課題ごとに新しいプログラムを別々に書くよりも、基本技術を備えたファウンデーション・モデルひとつを複数の部門が共有する方向で研究組織自体を変えていっています。

新薬開発の現場でこの原理が力を発揮するには、乗り越えるべき壁がひとつありました。病院の電子健康記録とゲノムデータベース、医薬品企業の化合物ライブラリ、学術誌に散らばった論文は、長年にわたって異なるサーバー、異なるフォーマット、異なる所有権の下に閉じ込められていました。ある大学病院が集めた肺がん患者のゲノムデータは、その病院の中だけで回されていたし、ある医薬品企業が実験で確認した化合物の毒性データは、その企業の外に出ていませんでした。

規制機関が新薬承認審査過程で積み上げた有害事象報告書も、公開されたいくつかのデータベースを除けば、ほとんどが機関内部の書類箱に留まっていました。データがこのように分かれていると、どんなに優れたアルゴリズムでも、自分の前に置かれた狭い倉庫の中の情報だけで学習するしかありません。2023年ころから、医薬業界と学界、ビッグテック企業が同時に乗り出した作業は、このように分かっていた貯蔵庫をひとつに集める作業でした。疾病情報と患者のゲノム、臨床データを一つのモデルの中に一緒に入れて学習させると、モデルは肺がん患者の遺伝子変異ひとつを見ても、それと似たようなパターンを示した別の悪性腫瘍、別の医薬品、別の臨床試験結果を一度に思い出す汎用生物学推論能力を備えるよ

うになります。

この汎用推論能力こそが、ひとつの疾病、ひとつの標的だけを計算していた前の世代のプログラムとファウンデーション・モデルを分ける決定的な差異です。

生命現象を見つめるデータにも、複数の層があります。ゲノムは体を設計する全体的な設計図のようなもので、生まれてからほとんど変わらず、トランスクリプトームはその設計図の中で今この瞬間にどの遺伝子が実際に展開されて機能しているのかを示します。プロテオームはそのように展開された設計図に従って実際に作られた部品、つまりタンパク質リストを示し、メタボロームはその部品が機能しながら生み出した最終産物を示します。同じ建物でも設計図だけを見ては、今何階に照明が付いているのかわからないように、ゲノム情報だけでは、特定の細胞や特定の患者の中で今何が起きているのかをすべて知ることはできません。

ファウンデーション・モデルがゲノム、トランスクリプトーム、プロテオーム、メタボロームをいっしょに学習する理由は、設計図と照明の状態、部品リスト、排出物まで一度に見なければ、建物全体の状態を適切に診断することができないからです。ここに患者の診療記録と画像資料まで加わると、モデルは分子レベルの変化が実際の症状とどのように結びついているのかも一緒に推論できるようになります。

この統合に決定的な一撃を与えた洞察は、タンパク質配列と分子構造を人間の文と同じ方法で読むことができるという発見でした。タンパク質は20種類のアミノ酸が鎖のように繋がった構造ですが、これは26個のアルファベットが並んで単語と文を形成する方式と基本的に変わりません。化学者たちが長年にわたって図で描いていた分子構造も、原子と結合をひとつの記号配列に置き換えて記すことができます。もともと人間の文を翻訳し、次の単語を予測することを目的に設計されたトランスフォーマーという神経回路網構造は、文ではなくアミノ酸配列や分子記号配列を入力しても同じように機能しました。次に来るアミノ酸や原子を当てる訓練を繰り返す間に、モデルは自動的に生命現象の文法を習得し始め、研究者たちが特別に規則を教えなくても、折り畳み方や反応性のような化学的性質を自ら気付きました。

この文法を習得したモデルは、これまでになかったタンパク質の配列を新しく作り出したり、標的タンパク質のくぼんだ結合部位にぴったり合う分子を最初から設計したりすることができました。テキストを扱っていた生成AIの力が生物学と化学の言語にそのまま移り、新薬開発はすでに存在する候補の中から選ぶことから、この世になかった候補を作り出すことへと移行しました。

ファウンデーションモデルが作られる順序にもそれなりの論理があります。まず、膨大なデータを置いて特定の目的を持たずに生命現象の規則そのものを習得する事前学習の段階を経ます。この段階でのモデルは、小学校や中学校で国語や数学、科学の基礎を幅広く学ぶ生徒のようなもので、まだ特定の進路を決めてはいませんが、どんな問題が来てもひとまず手をつけることができる基礎を築きます。このように基礎を固めたモデルを特定の課題に合わせて再び訓練する微調整の段階を経ると、初めて新薬の候補物質の設計や臨床試験の結果予測のような具体的な任務に特化した実力がつきます。

TxGemmaはこのような順序をそのまま踏んだ事例であり、Googleの汎用言語モデルであるGemma 2が事前学習で固めた基礎の上に、治療薬関連のデータで微調整を加えて作られました。基礎を広くしっかりと固めるほど、その上に乗せる微調整の効果も大きくなるという点は、データ統合がなぜファウンデーションモデルの成否を分ける鍵であるのかを改めて示しています。

このように積み上げられたファウンデーションモデルの構造は、3階建ての建物に似ています。一番下の階には、化合物ライブラリやゲノム、トランスクリプトーム（細胞内で遺伝子が実際にどれだけ読み取られているかを示すデータ）、そして数百万件の臨床文献が基礎資材のように積まれています。この資材を整えて骨組みを立てる中間の階には、文章を理解する大規模言語モデルや、タンパク質の配列、分子構造を専門に読み取るファウンデーションモデルが置かれています。中間の階で整えられた骨組みは一番上の階に上がり、新薬の候補物質の生成、標的の優先順位付け、臨床試験の設計最適化という実際の成果として完成します。

例えば、下の階に積まれた肺がん患者のゲノムと関連論文を中間の階のモデルが読み、肺がんでよく目立つタンパク質を一つ選び出すと、上の階ではそのタンパク質にくっつく新薬の候補物質をいくつかすぐに描き出します。下の階の資材が

不十分であれば、中間の階がいくら精巧でも建てるものがなく、中間の階が崩れれば、下の階にいくら良い資材が積まれていても上の階まで力が伝わりません。3つの階が同時に丈夫であってこそ、初めて新薬開発の現場で役立つ結果が出るという事実を、2025年と2026年の間に次々と登場したファウンデーションモデルが一つずつ証明して見せています。

中間の階と上の階を繋ぐ力をはっきりと示した事例は、Googleが発表した治療薬特化の大規模言語モデルTxLLMです。TxLLMは、分子構造、標的タンパク質、細胞株、臨床試験のテキストという互いに異なる形式のデータを、一つのモデルの中に一緒に押し込んで学習させた結果物です。分子は原子記号で、タンパク質はアミノ酸の配列で、臨床試験の記録は人が書く文章でと、それぞれ表現方式が異なるにもかかわらず、TxLLMはこれらすべてを一つの共通した表現空間の中に溶け込ませ、互いに比較して連結できるようにしました。研究チームは、標的の発掘から臨床試験の結果予測まで、新薬開発の全過程にわたる66個の評価作業を置いて、TxLLMの実力を試験しました。

評価項目には、特定の化合物が臨床試験の3相まで生き残る確率を測る問題から、2つの薬物を一緒に投与した時に現れる可能性のある相互作用を見つけ出す問題まで混ざっていました。TxLLMはこのうちの22個の作業で、その時点までに出ていたどの専門モデルよりも優れた成績を記録しました。一つのモデルが、分子設計と臨床予測という互いに遠く離れた2つの仕事を同時によくやり遂げるという事実は、ファウンデーションモデルが論文の外へと歩み出て、実際の開発現場で使われることができるという証拠として受け止められました。

新薬開発チームが実際に直面する困難は、標的の候補が足りないことではなく、多すぎることにあります。人のゲノムにはタンパク質を作り出す遺伝子だけで2万個を超え、その中で新薬の標的になりそうな候補は、論文や特許、臨床記録を調べると数千個単位で溢れ出てきます。研究費と人材は限られているため、製薬会社はこの中から数十個、多くても数百個を選んで先に深く掘り下げるしかありません。かつてはこの選別作業を、経験豊富な数人の科学者の直感と論文検索に依存して進めており、その結果、すでに流行している標的ばかりに研究が集中する偏り現象が繰り返されていました。

TxLLMやTxGemmaのようなモデルは、標的と病気、既存の薬物の成否の記録を一度にざっと確認し、特定の標的が臨床まで成功する確率を点数でつけてくれるため、人が見逃しやすい候補も根拠と一緒に前へと引き上げることができます。標的の優先順位を決める作業に数週間から数ヶ月ずつかかっていた過去とは異なり、ファウンデーションモデルを使う今は、数日以内に順位表の草案が出ます。

TxLLMが打ち上げた可能性は、2025年3月にGoogleがTxGemmaを公開したことでさらに一歩進みました。TxGemmaは、Googleの軽量言語モデルであるGemma 2を骨組みとし、700万件を超える治療薬関連の学習事例で再び訓練したモデルであり、20億、90億、270億個のパラメータ規模の3つのサイズに分かれて登場しました。パラメータは、モデルが学習過程で自ら値を調整していく内部の取っ手のようなものであり、その数が多いほど概して精巧な判断が可能になります。予測専用のモデルとは別に、対話型で動作するTxGemma-Chatバージョンも一緒に登場し、研究者がチャットをするように質問を投げかけると、モデルが根拠を提示しながら答えるという方式も可能になりました。

最も大きいサイズの270億パラメータ版は、以前のTxLLMと同じ66個のタスク中64個で同等かそれ以上の性能を達成し、その中45個では既存の記録を塗り替えました。TxGemmaがTxLLMと異なる決定的な違いは公開性です。TxLLMがGoogle内部の研究結果のままだったのに対し、TxGemmaはモデルの重みをHugging FaceとGoogle CloudのVertex AIを通じて全面的に公開することで、大学の研究室や資金が十分でないバイオベンチャーも同じレベルの基礎モデルを使用できる道を開きました。

ゲノム全体を1つの言語で読み解く試みはEvo 2で頂点に達しました。Arc InstitutがNVIDIA、Stanford大学、California大学 San Francisco Campus、California大学 Berkeley Campusと共に開発したこのモデルは400億個のパラメータを有しており、1度に最大100万塩基（遺伝子を構成する化学的文字単位）の長さの配列を読み取ることができます。学習に使用されたデータは、人間と植物、細菌と古細菌を問わない10万種を超える生物から得られた9兆個を超える塩基でした。Evo 2は2025年2月にプレプリント論文の形で最初に公開され、その後検証を経て国際学術誌Natureに正式に掲載されました。

このモデルは、人間に病気を起こす遺伝子変異を事前に検出することから、新しい微生物酵素を設計することまで、さらには存在しなかったゲノム配列を最初から新たに創り出すことまで実行できることを示しました。Arc Instituteはこのモデルが人間の遺伝病研究だけでなく、農業育種や産業用微生物開発、新材料探索にも使用できると説明しました。1つのモデルが診断と設計という異なる2つの任務を同時に処理するという点において、Evo 2はゲノムデータ統合が新薬開発にどのような力を与えることができるかを示す象徴的な事例として位置付けられています。

Google DeepMindも2025年6月25日にAlphaGenomeという名称で類似方向のモデルをリリースしました。AlphaGenomeは100万塩基長のDNA配列を入力として受け取り、遺伝子発現、クロマチンアクセス性（DNAがどの程度簡単に開かれて読まれるかを示す値）、ヒストン修飾、ライゲーシオン位置など数千種類の調節シグナルを一度に予測します。人間のゲノム中でタンパク質を直接作らない領域、いわゆる非コード領域で起こる突然変異が実際に病気を引き起こすかどうかを判断することは、長い間遺伝学者を困らせてきた難問でしたが、AlphaGenomeは24個の評価項目中22個、26個の変異予測項目中24個で既存モデルを上回る成績を達成しました。

ただし、このモデルは個人のゲノム全体を入力して病気を診断するように設計されておらず、調節配列が10万塩基以上離れた遺伝子への影響を読み取る際にはまだ精度が低いという制限も明らかになりました。Google DeepMindは2026年1月、このモデルのソースコードと重みを非営利研究目的に限定して無料で公開し、その結果、世界中の遺伝学研究室が同じツールを使って異なる疾病の調節変異を比較できる共通基盤が整備されました。

病院ごとに散在していた患者ゲノムの読み取り結果がこの共通モデル上で再び出会うことにより、希少疾患診断のように単一病院のデータだけでは答えが得られなかった問題にも糸口が付き始めました。

タンパク質の分野では、米国のスタートアップEvolutionaryScaleが開発したESM3が同様の飛躍を遂行しました。従来広く使用されていたAlphaFoldシリーズモデルが与えられたアミノ酸配列を見て3次元構造を予測することに重点を置いて

いたのに対し、ESM3は構造や機能の側から条件を逆にして与えても、それに合致する新しい配列を作り出すことができるという点で、予測を超えて設計へとさらに一步進んだモデルです。ESM3はアミノ酸配列と3次元構造、そしてそのタンパク質が実際にどのような機能を果たすかを一緒に学習することで、3つのうちのどれか1つだけが条件として与えられても、残りの2つを自ら補完して完成させる能力を備えています。

このモデルの能力は2025年1月に国際学術誌Scienceに正式論文として掲載されることで検証されました。研究者たちはESM3に蛍光タンパク質が備えるべき条件だけを伝え、残りの配列を新たに作成させましたが、このようにして作成されたタンパク質は、自然界のどの蛍光タンパク質とも配列が明確に異なりながらも、明るく光を放ちました。研究者たちはこの結果について、自然界が5億年にわたる進化で到達したであろう境地にモデルが計算だけで到達したと評価しました。ESM3の重みとコードは学術研究目的に限定してGitHubを通じて公開されており、一般ユーザー向けのアプリケーションプログラミングインターフェイスにも無料で使用できるステージが用意されているため、論文を読むだけだった大学院生も直接新しいタンパク質を設計してみることができる環境が整っています。

ただし、商業的に大規模に使用するには別途契約が必要であるため、アクセス障壁をどこまで低くするかという問題はまだ検討中のままです。

細胞1つ1つの反応まで全て模倣しようとする試みは、2025年夏にArc Instituteが開催したVirtual Cell Challengeで明確なマイルストーンを残しました。細胞に遺伝子改変または薬物刺激を加えたときに出現する反応を最も正確に予測するモデルを選別するこの大会には、114カ国から5,000人を超える研究者が登録し、1,200個を超えるチームが結果を提出しました。優勝賞金10万ドルが掛かっているこの大会はArc Instituteが共に公開したVirtual Cell Atlasというデータリポジトリを基盤として実施され、このリポジトリには複数の実験室から得られた6億個を超える細胞の観察および刺激応答データが整理されて含まれています。

このデータには、バイオテック企業のTahoe Therapeuticsが提供した大規模な薬物反応データセットと、Arc Instituteが人工知能エージェントに公開された論文を調べさせて自動で集めた単一細胞データまで一緒に含まれています。大会の後、

Arc Instituteが発表したSTATEというモデルは、特定の細胞株に特定の刺激を与えた時、遺伝子の発現がどのように変わるかを実験なしで事前に計算して出します。1つの病気、1人の患者のデータに留まらず、様々な実験室、様々な細胞株、様々な刺激条件のデータを1か所に集めて初めて、細胞レベルの汎用的な予測力が生まれるという事実を、この大会は成績表として確認してくれました。

Arc Instituteはこの大会を2026年にも再び開くと予告しました。散らばっていた細胞実験データが1か所に集まることで、個別の実験室1つでは答えられなかった質問に糸口が見え始めたわけです。

このような流れは2026年に入り、個別のモデルの競争を超えて、産業全体の基盤施設を組み直すレベルへと規模を拡大しました。OpenAIは2026年4月16日、生命科学研究専用モデルGPT-Rosalindを発表しました。DNA二重らせん構造を明らかにするのに決定的な糸口を残した科学者Rosalind Franklinの名を冠したこのモデルは、その年の6月に性能を再び調整し、前のバージョンと比べてトークン（文章を処理する最小単位）の消費量を31パーセント減らしながらも、医薬品化学とゲノム分析の実力を引き上げました。ただし、このモデルは誰もがすぐに使える形では公開されず、AmgenやModerna、Thermo Fisher Scientific、Allen Instituteのように検証を経た機関にのみ、先にアクセス権限が与えられました。

このリストに名を連ねたAmgenとModernaは、すでに独自に膨大な臨床データを蓄積してきた大手製薬会社であるという点で、OpenAIのモデルが既存のデータを押し出すのではなく、その上に乗る形で定着する可能性を示しています。OpenAIはこれと共に、生命科学技術が誤った目的で使われる状況を防ぐためのRosalind Biodefenseという別の事業も一緒に始めました。生命科学専用モデルにまで会社の名前を掲げて参入したという事実自体が、このようなファウンデーションモデルがもはや学界と少数のバイオテックだけの領域ではないという点を示しています。

Amazon Web Services (AWS)も同じ年の4月15日、生命科学シンポジウムでAmazon Bio Discoveryを披露しました。このサービスは、40種類ほどの生物学ファウンデーションモデルを1つの画面に集めておき、研究者が自然言語でリクエストを投げれば、適切なモデルを選んで実験の設計まで立ててくれる人工知能秘書の形を備えました。合成と検証を担当する協力実験室のリストにはTwist

BioscienceとGinkgo Bioworksがすでに名を連ねており、A-Alpha Bioもすぐに合流する予定だとAmazon Web Services (AWS)は明らかにしました。このサービスは、研究者が自然言語で実験の手順を書き込めば、これを複数のモデルと分析段階からなる1つの作業順序に変えてくれる機能も備えており、コードを書くことができない生物学者でも複雑なパイプラインを自ら組み立てることができます。

Memorial Sloan Kettering Cancer Centerはこのサービスで抗体候補分子30万種を設計した後、そのうち10万種をTwist Bioscienceに送って実際の合成と検証まで終えましたが、数か月ずつかかっていた作業が数週間に縮まりました。Nvidiaはすでに2025年4月、小さな分子を扱う生成型モデルMolMIMをBioNeMoに追加したことがあり、2026年1月には同じプラットフォームにRNA構造予測モデルと合成可能性検証モデルを新たに乗せました。その月の12日には、Eli Lillyと新薬発掘の難問を一緒に解く共同革新研究所を設立する協力を、Thermo Fisher Scientificとは実験設備に人工知能を組み合わせる協力を並んで発表し、ファウンデーションモデルをコンピューターの外の実験室現場までつなぐ作業に飛び込みました。

予測と検証を直ちにつなぐこのような循環構造は、ファウンデーションモデルが論文の中の性能指標を超えて、実際の新薬パイプラインの中に定着したことを示しています。

このようなモデルを最初から訓練させるには、グラフィック処理装置数千台を数週間ずつ動かす計算量が入り、その電気料金と設備費用だけでも、一般的なバイオベンチャーの1年の予算を軽く超えます。その結果、新しいモデルを基礎から立ち上げる仕事は、Nvidia、Google、Amazon、OpenAIのように独自のデータセンターを備えた少数の企業と、政府の支援を受けるいくつかの研究所に集中する流れがはっきりしています。

規模の小さなバイオテックや大学の研究室は、これらの企業が公開したモデルの重みをダウンロードし、自分たちの特殊なデータで微調整する方を選ぶケースが増えています。TxGemmaやESM3、AlphaGenomeが重みを公開したのも、このような格差を少しでも縮めようとする試みと見ることができます。

これまでに出了ファウンデーションモデルは、それぞれ得意な領域が分かれていきます。Evo 2とAlphaGenomeはゲノム配列をよく読み、ESM3はタンパク質構造をうまく扱い、TxGemmaはテキストと臨床データをうまく結びつけます。1人の患者のゲノムと細胞反応、診療記録、映像資料をすべて読み取る1つのモデルに統合する作業は、まだ進行形です。Amazon Bio Discoveryのように複数のモデルを1つの画面に集めておき、必要な時に選んで使わせる方式は、完全な統合に至る中間段階と見ることができます。

異なる会社、異なる研究所が出したモデルを1つにつなぐ標準がまだ設けられていないという点は、この分野が抱えている目立った課題です。それにもかかわらず、わずか2、3年の間にゲノム1つだけを読んでいたモデルが細胞反応と臨床記録まで網羅する方向に急速に広がってきた流れを見ると、次の統合段階がそれほど遠くないところに来ていると推測することができます。

ファウンデーションモデルの価値は、アルゴリズムのコードそのものよりも、そのモデルが読み込んだデータの幅と、そのデータを実験で確認するスピードで決まります。同じトランスフォーマー構造を使っても、10万種の生物のゲノムを読み込ませたモデルと、ある一つの会社の化合物データだけを読み込ませたモデルとでは、まったく違う実力を見せます。ですから今、問いは「どのアルゴリズムがより賢いか」ではなく、「どの組織がより広くて深いデータ基盤の上に立っているか」へと移り変わっています。

データを幅広く握り、そのデータを実験としてフィードバックする循環構造まで備えた組織が、次の段階の新薬開発競争で優位に立つ可能性が高いと思われます。一つの標的を当てる計算機ではなく、病気やゲノム、そして臨床データを行き来する推論基盤を手にした組織は、新薬の候補物質の生成、標的の優先順位の判断、臨床試験の設計という3つの実務を一つに握ることになります。2022年にInsilico Medicineが1つの候補物質を30ヶ月で臨床1相まで進めたことが異例の成果として受け止められたとすれば、今ではそれだけのスピードを複数の病気で同時に再現することが、ファウンデーションモデルを備えた組織の前の次の課題として置かれています。

キムチチゲだけを作っていた料理人が、基本を広く備えた料理人になるために長い修業が必要だったように、生命科学のファウンデーションモデルが病床のそばの実際の処方へとつながるまでにも、データを集めて実験で検証する修業の時間は依然として必要です。

Gwagi-jeongtong-buは2025年10月、医科学分野を扱うファウンデーションモデル開発の主管機関として、医療人工知能企業のLunitが率いるコンソーシアムを選びました。ここにはKakao HealthcareやSK Biopharmのような企業、KAISTやSeoul Daehakgyoの研究チーム、病院9か所を合わせて23の機関が参加し、分子やタンパク質から臨床論文まで、様々な層の医科学データを一つにまとめて2026年9月までにモデルを完成させる計画です。Bogeon-bokji-buもまた、Hanguk Jeyak Bio Hyeophoeを総括機関として立て、31の機関と共に、前臨床段階から臨床段階までをつなぐK-AI新薬開発モデルを作っています。

Seoul Daehakgyo ByeongwonとSamsung Seoul Byeongwon、Hanguk Saengmyeong Gonghak Yeonguwonが分けて率いるこの作業は、病院ごとに違う場所に散らばった患者データを一つのサーバーに移すことなく、各病院の中で学習した結果だけを集めて一つのモデルに合わせる連合学習方式を使います。データが互いに異なるサーバーの中に閉じ込められていた古い問題を、データを移さずとも知識だけを集める方法で解いてみようとする韓国式の実験というわけです。

ファウンデーションモデルが出した予測が、実験台の上ですぐに覆された事例も出ました。スイスのバーゼル大学薬学部のマルクス・リル教授の研究チームは、AlphaFold3をはじめとする複数の構造予測モデルに意地悪なテストを仕掛けました。薬物がくっつく結合部位のアミノ酸を丸ごと変えたり、化合物の電荷を裏返したりして、実際には薬物がタンパク質にくっつくことができないようにした後、同じ問題を再び解かせたのです。

2025年10月、国際学術誌Nature Communicationsに載った結果を見ると、285個の複合体のうち40パーセントを超える場合で、モデルは結合部位が変わっていないかのように以前とまったく同じ答えを出しました。研究チームは、この結果について、モデルがなぜ薬物がタンパク質にくっつくのかを理解したのではなく、学習データで見た絵をそのまま繰り返したただけだと説明しました。

学習データと計算資源が少数の企業に偏る流れは、重みの公開だけではすべて解けない問題のように見えます。重みをダウンロードして微調整する側と、そもそも膨大な量の塩基配列や臨床記録を集めて最初からモデルを訓練させる側とでは、出発線そのものが違うからです。モデルを訓練させるときに使った元のデータの構成や出所まで一緒に公開する企業はまれで、外の研究者がそのモデルが実際に何を学んだのかを検証することも簡単ではありません。政府が予算を集めて病院や大学のデータをまとめてくれる韓国の連合学習方式のように、公共部門が立ち上がらなければ、この偏りは今後さらに広がるものと思われます。

1.3

多重オミクスの統合と新規バイオマーカーの発見

日本の複数の病院で進行された臨床第2相トライアンフ（TRIUMPH）試験は、ヒト上皮成長因子受容体2型、すなわちHER2遺伝子が大きく増加した転移性大腸がん患者を対象にしました。この患者たちはゴジラ（GOZILA）という循環腫瘍DNA選別検査研究を経て先に選ばれましたが、組織を再び切り取る代わりに血液一本で標的遺伝子の増幅の有無を見つけ出す液体生検方式でした。腫瘍組織を再び切り取ることが難しい患者にも標的治療の門を開いてくれたわけです。

研究陣は事前選別検査を受けた患者143名と実際の試験に参加した30名の組織スライドを人工知能病理判読プログラムに入力しました。このプログラムは韓国企業Lunitが作ったSCOPE HER2とSCOPE IOというソフトウェアで、スライド写真一枚の中でがん細胞がHER2染色薬にどれほど濃く染まったか、そして腫瘍の周りに免疫細胞や線維細胞がどれほどぎっしりと集まっているかを一度に数えました。病理医の判定と比べたとき、このプログラムの正確度は86.7パーセントであり、染色強度が一番濃い第3段階の患者を選び出すことでは100パーセント的中しました。

がん細胞の半分以上がその第3段階に染まった患者群は、そうでない患者群より無増悪生存期間が4.4ヶ月対1.4ヶ月と長く、全体生存期間は16.5ヶ月対4.1ヶ月に広がりました。この探索的分析結果は2025年国際学術誌である臨床腫瘍学精密医学（JCO Precision Oncology）に載りました。

この結果を正しく読み解くには、標的とバイオマーカーという二つの言葉から解き明かさなければなりません。標的は病気を引き起こす原因となるタンパク質や遺伝子を指します。錠前だと考えればよいです。薬はその錠前にぴったり合うように削って作った鍵であり、鍵が錠前の穴に入って回ってこそ病気を引き起こしていた扉が閉まります。バイオマーカーはこれとは違い、体の中で今何が起きているのかを教えてくれる信号です。

踏切の信号機や自動車の計器盤の警告灯のように、血の一滴や組織試料一つに込められた特定物質の量を測れば、この体がどんな状態でどんな薬が効くか、病気が今後どのように流れていくかをあらかじめ推測することができます。トライアンプ試験で人工知能が読み取ったHER2染色強度と免疫細胞密度が、まさにその信号機の役割をしたバイオマーカーでした。

人間の病理医が同じスライドを見ても互いに違う判定を下すことは珍しくありません。HER2染色強度を目分量で分ける2+と3+の間の境界は曖昧であり、このような曖昧な事例は別に蛍光in situハイブリダイゼーション検査を再び受けなければならない場合が多かったです。人工知能判読プログラムは人間の目が見逃しやすいこの境界を、染色強度と細胞比率という連続した数字に変えておき、判定者によって結果が揺れる偏差を大きく減らします。もちろんこの点数が病理医の代わりになるという意味ではありません。病理医が最終判断を下す前に、見逃しやすい微細な信号を先に選び出してくれる補助道具に近いです。

がんや認知症のような退行性脳疾患は、遺伝子一つが壊れて生じることは稀です。数百、数千個の遺伝子とタンパク質が絡み合っていて、特定の組み合わせでのみ病気として表面に現れます。組織検査一回、遺伝子検査一回ではこの絡み合った糸玉を解くことはできません。そこで研究者たちは、体を互いに違う層から覗き込む様々な技術の一つに集め始めました。ゲノム学、トランスクリプトーム学、プロテオーム学がそれです。ゲノム学は次世代塩基配列分析装備で細胞の中のDNA配列を丸ごと読み取ります。建物を建てるときに使う設計図だと見ればよいです。設計図は一度描かれればめったに変わりません。

トランスクリプトーム学はその設計図を置いて細胞がその日その日下ろす作業指示書に該当します。RNAシーケンシングという技術で、どの遺伝子が今どれほど多くオンになっているかを測ります。同じ設計図を持った細胞でも、肝細胞と神経細胞が互いに違って働く理由がここにあります。プロテオーム学は一步さらに進んで、実際に現場でレンガを運び釘を打つ働き手、すなわちタンパク質自体を質量分析器で一つ一つ確認します。同じタンパク質でもリン酸基のような小さなタグが付いたかどうかによって全く違って動きますが、このような翻訳後修飾まで調べる精密プロテオーム分析は、信号伝達経路で薬が実際に割り込むことがで

きる位置をはるかに細かく教えてくれます。

薬はたいてい遺伝子ではなくタンパク質に直接くっついて作動するので、この働き手たちが正確に何人いて、互いにどんな形で手を握り合っているかを知ることが新薬開発で要となります。問題はこれら三つのデータが測定方式も単位も規模もそれぞれ違うため、人間が並べて比較するのが非常に難しいという点です。

ここにメタボローム学とエピゲノム学まで加われば、絵がさらに細かくなります。メタボローム学は細胞が実際に何を食べて何を出すのか、ブドウ糖やアミノ酸のような小さな分子の濃度を測る作業であり、エピゲノム学はDNA配列自体はそのまま置いたまま、特定遺伝子に化学的タグを付けてその遺伝子をオンにしてオフにするスイッチの役割をする変化を調べます。

人間の遺伝子は2万個余りとして知られていますが、この遺伝子たちがトランスクリプトームやプロテオーム、メタボローム段階を経て作り出す組み合わせの数は、それより何桁も多くなります。この膨大な組み合わせを人間が手で対照することは最初から不可能に近く、まさにここで多重オミックスデータを丸呑みしてパターンを見つけ出す人工知能の役割が大きくなります。

ここで人工知能が割り込みます。変分オートエンコーダやトランスフォーマーのようなディープラーニングモデルは、ゲノム、トランスクリプトーム、プロテオームデータを一度に入れ込み、人間の目にはつかない潜在された規則を見つけ出します。三つのデータを無理やり平均を出したり並べて置いたりする代わりに、モデル内部に新しい座標空間を作り、患者や細胞一つ一つをその空間の一点に移しておく方式です。移された点たちの距離と配置を調べれば、伝統的な統計技法では捉えきれなかった絡み合った関係、例えば遺伝子一つとタンパク質一つが特定条件でのみ一緒に問題を引き起こす関係が現れます。

マーケットリサーチ企業は、このマルチオミックス分析市場が2025年の28億1000万ドルから2026年の32億ドル前後に拡大し、2035年までは毎年15パーセント以上の成長が見込まれていると予測しています。

ターゲットを見つける作業がさらに一段階洗練されたきっかけは、単一細胞転写体分析（scRNA-seq）の登場です。以前は、組織の塊をまるごと砕いて、その中

に含まれるすべての細胞の遺伝子発現をひとまとめにして測定していました。複数の人が混ざった写真から平均的な顔を描くようなものですから、実際に隠れている少数のがん幹細胞や疲れ果てた免疫細胞は統計の中に埋もれて消えてしまいました。

単一細胞分析は、これとは異なり、細胞一つ一つにラベルを付けて、別々にインタビューするかのように遺伝子発現を測定します。腫瘍組織のように正常細胞と免疫細胞、突然変異細胞が混在した場所では、この方法でなければ細胞ごとの異なる事情を知ることはできません。

研究者は数万個の細胞をこのように一つ一つ読み取ることで、がんになる前の段階の細胞から始まり、細かく分かれた亜がん細胞群、そして力が抜けて本来の役割を果たせなくなった免疫細胞に至るまでの時間の流れをたどります。2025年から2026年の間に、イルミナは「ビリオンセルアトラス」というプロジェクトを通じて、わずか1年で20ペタバイトに達する単一細胞転写体データを新たに構築しており、究極的には50億個の細胞規模のデータベースを構築しようとしています。

このように蓄積されたデータで学習した単一細胞ファウンデーションモデルは、特定の細胞がどのような種類であるかに名前を付けることや、特定の刺激を与えたときに細胞がどのように反応するかを事前に予測することまでできます。

単一細胞データで学習する人工知能も、文献ベースの言語モデルと手を組み始めました。2026年に公開されたELISAというフレームワークは、単一細胞遺伝子発現を理解するscGPTのようなモデルとBioBERTベースの文献検索機能を1つに統合して、研究者が特定の細胞型で特に多く活動する遺伝子を見つけると、その直後に関連論文の根拠も一緒に表示します。実験データと文献知識がこのように1つの画面内で噛み合うことで、新しいバイオマーカー候補を目で確認し、すぐに妥当性まで検討することが可能になりました。

しかし、細胞一つ一つの正体を知ったからといって、話が終わるわけではありません。細胞は、どこに配置されているかによって、その役割が異なるからです。空間転写体学（Spatial transcriptomics）は、組織を碎かずに元の形のまま保ち、各位置でどのような遺伝子がどの程度活動しているかを地図に描きます。住民の

名前だけが書かれたリストではなく、何号室の誰が住んでいるかまで記載された地図のようなものです。10x GenomicsのVisiumや、複数の蛍光標識を同時に付ける多重エラー防止蛍光その場ハイブリダイゼーション法（MERFISH）といった技術が、この地図を単一細胞レベルの精度で描き出します。

この分野は2025年から2026年の間に急速に変わりました。10x Genomicsは2025年8月にXenium Proteinをリリースし、同じ細胞、同じ組織スライスからRNAとタンパク質を同時に測定する道を開きました。2026年4月には、Ateraという新しい装置を発表しました。既存のXeniumより4倍広い面積を2倍多くのスライドで、すべての遺伝子を対象に単一細胞単位まで測定しながら、グラフィック処理装置の力を借りて分析時間を時間単位から分単位に短縮しました。

その装置で得たデータから、低酸素腫瘍組織に集中したODAMや、侵襲性の強い分裂細胞に集中したC21orf58のような位置に束縛されたバイオマーカー候補が新たに確認されました。

人工知能は単一細胞データと空間データを重ね合わせることで、腫瘍細胞と周辺の結合組織細胞、免疫細胞の間の実際の物理的距離と信号のやり取りを計算します。2025年に発表された研究では、空間転写体と空間プロテオーム・データで一緒に学習したグラフニューラルネットワークが、免疫チェックポイント阻害剤への反応予測の精度指標であるROC-AUC 0.90を超える精度で予測し、細胞外マトリックスで厚く囲まれた腫瘍周辺領域が薬が効かない患者で特に顕著であるという事実を明らかにしました。

トライアンフ試験でのLunitのプログラムが読み取った免疫細胞密度も、このアプローチと変わりはありません。Lunitは2025年12月、胆道がん患者291人、スライド309枚を対象とした研究で、HER2判定結果が病理医の合意と83.5パーセント一致するという論文を国際学術誌に掲載し、2026年米国臨床腫瘍学会学術大会で5件、米国がん研究協会学術大会ではアジュ大学およびアジレントと共同した非小細胞肺癌研究を含む6件の研究を新たに発表しました。

腫瘍微小環境内のマクロファージも、人工知能が注視する対象です。マクロファージは性質が固定されていないため、腫瘍を助ける方向にも、腫瘍と戦う方向にも

方向を変えることができる細胞です。ある空間転写体研究では、CXCL16とSPP1という2つのシグナル分子がマクロファージのこのような方向転換を調整しているという事実を明らかにし、既に他の用途に使用されていたシロリムスとトリコスタチンA (trichostatin A) という2つの薬物がこのシグナル経路をターゲットにできる候補として浮上したと報告しました。

新しい化合物を最初から設計する代わりに、既に安全性が確認された薬物を別の病気に再び使用するこのような薬物再創出アプローチは、開発期間を大幅に短縮します。

これまで検討したデータはすべて、実験室装置が細胞と組織から直接抽出した数字です。しかし、人間が患う病気に関する知識の多くは、数字ではなく論文、特許、臨床試験報告書の中の文章に含まれています。毎年数百万の文献が流れ出てくるこの膨大な文献をすべて人間が読んで整理することは、そもそも不可能です。そこで登場したのが、BioBERTとPubMedBERTのように、膨大な生物医学論文で事前に学習させた言語モデルです。

このようなモデルは、一つの文章を読み、その中にどのような遺伝子、病気、薬の名前が出てくるかを自動的に見分け、それらが互いにどのような関係にあるのかまで抽出します。reguloGPTというモデルはさらに一歩進んで、一つの文章から、遺伝子が他の遺伝子の発現を抑えるのか高めるのかといった調節関係までも整理された表に変換します。

このように文献から抽出した関係と、実験室で得られたゲノム、トランスクリプトーム、プロテオームのデータを一つにまとめると、巨大な知識グラフになります。遺伝子とタンパク質、病気と薬がそれぞれ点になり、それらの間の関係が点を結ぶ線となる、クモの巣のように緻密な地図です。研究者たちはこの地図の上で、グラフニューラルネットワークや知識グラフ埋め込みといった手法を用い、網のように絡み合ったネットワークの中で特に多くの線が集まる遺伝子が何であるかを探します。まだ論文で報告されたことのない薬と標的の間の関係までも、このグラフの上で前もって推し量ることができます。

研究者たちはこのようなグラフ分析を疾患ネットワーク分析と呼びます。直接友達にならなくても友達の友達を通じて紹介されるように、グラフニューラルネットワークは、遺伝子Aと遺伝子Cが一度も同じ論文に並んで登場したことがなくても、どちらも遺伝子Bとよく関わっているという事実だけで、二つの間の隠れた関連性を見つけ出します。2025年に発表された知識誘導型グラフ学習手法であるKGDRPは、薬物反応予測と標的予測の二つの課題において、これまで知られているデータがほとんどない新しい病気についても予測性能を12パーセント引き上げました。

文献や臨床データがまだ乏しい希少疾患においても、人工知能が標的候補を最初から提示できるという意味であり、新薬開発の初期段階における試行錯誤を大幅に減らしてくれると見られます。

Insilico MedicineのPandaOmicsは、この方式をいち早く商業化したプラットフォームです。遺伝子発現とタンパク質のデータ、文献ベースの知識グラフを一つに結びつけ、肝細胞がんや特発性肺線維症をはじめとする様々な疾患において新しい標的の優先順位をつけ、実際の実験で検証してきました。2025年には、信頼度と商業的な可能性、薬にできる度合い、作用機序がどれほど明確であるかを一緒に評価する採点システムに磨きをかけ、TargetBench 1.0と、腫瘍学、代謝疾患、免疫疾患、線維化疾患、神経疾患の38の疾患を網羅するTargetProという枠組みを新たに発表しました。

同じ年、Insilico Medicineは胃がんと食道がんの標的を探すために世界的ながんセンターと手を結び、2026年にはPandaOmicsが選び出した脳疾患の標的を巡って、中国医療システムグループと1億6500万ドル規模の2度目の協力契約を結びました。

このプラットフォームが明確に成果を示した事例はrentosertib（開発コード名ISM001-055）です。PandaOmicsは線維化と老化に関する文献とオミックスデータに目を通す中でTNIKというタンパク質を見つけ出しましたが、このタンパク質はそれまで特発性肺線維症の治療薬の標的として注目されていた候補ではありませんでした。中国の22の病院で患者71人を集め、12週間にわたって偽薬と3つの用量を比較した臨床第2a相のGENESIS-IPF試験において、1日60ミリグラムずつ

投与された患者群は努力性肺活量（FVC）が平均98.4ミリリットル増えたのに対し、偽薬群は20.3ミリリットル減りました。

2025年6月3日にNature Medicineに掲載されたこの結果は、人工知能が最初から最後まで関与して見つけ出した標的と薬が、人間の体で実際に効果を出すという事実を業界で初めて証明した事例として挙げられ、その後第3相試験へとつながりました。

標的を探す方法には、合成致死（Synthetic lethality）分析もあります。飛行機にエンジンが2つ付いていると考えると理解しやすいです。エンジンが1つ止まっても残りの1つで飛行を続けることができますが、両方とも止まると墜落します。細胞の中にもこのようなペアを組んだ安全装置がいくつもあって、遺伝子が1つ壊れても他の遺伝子はその仕事を代わりに引き受けて細胞が生き残ります。

ところが、がん細胞はすでに遺伝子が1つ壊れたままで育った場合が多く、その残りのペアを組む遺伝子までも薬で防いでしまうと、正常な細胞は無事なのになんかがん細胞だけが死にます。人工知能は、汎がん多重オミックスデータと、遺伝子を1つずつ消してみるシミュレーションと一緒に学習し、このようなペアを大規模に見つけ出します。

このようなペアを探す人工知能モデルは、たいてい遺伝子依存性マップと呼ばれる大規模な公開データベースを学習材料として使います。アメリカのBroad Instituteが構築したCancer Dependency Mapが代表的ですが、数百種類のがん細胞株で遺伝子を1つずつ消していき、どの遺伝子がなくなればその細胞株が死ぬのかを実験で記録しておいた資料です。ここに各細胞株のゲノムとトランスクリプトーム、プロテオームの情報を重ね合わせると、特定の遺伝子変異を持った細胞株でのみ特によく死ぬペアを統計的に選び出すことができます。

BRCA遺伝子が壊れたがん細胞でPARPという酵素を防げば、がん細胞だけを選んで死ぬという事実が、初めてこのアプローチの可能性を証明しました。その後、臨床にまでつながったペアはそれほど多くありませんでしたが、2025年と2026年の間に新しい候補が相次いで出ました。MTAP遺伝子がなくなったがん細胞ではPRMT5という酵素を防ぐとがん細胞が死ぬという関係が確認され、IDEAYA

Biosciencesは2025年9月にIDE892というPRMT5阻害剤の臨床試験計画を提出し、同年第4四半期にMTAP欠損肺がん患者を対象に第1相試験を始めました。

ブリストルマイヤーズスクイブのBMS-986504は、MTAP欠損膵がん患者60名を登録する第2相試験につながりました。膵管腺がん研究では、単一細胞空間転写体学とプロテオミクス、競争的内因性RNAネットワーク分析、SELFormerというトランスフォーマーベースのディープラーニングを組み合わせたパイプラインが、TNFRSF10Aという遺伝子が産生するアポトーシス受容体TRAIL-R1を新しい標的として同定し、腎臓がん治療薬として既に承認されているテムシロリムスがこの受容体に結合することをコンピュータシミュレーションで確認しました。

発見されたバイオマーカーが実際の臨床現場で使用されている様子は、黒色腫治療薬ペムブロリズマブの大規模臨床試験であるKEYNOTE-002およびKEYNOTE-006で確認されます。研究者たちは575名の患者のCT画像から腫瘍の大きさと形の変化を示す特徴25個を抽出した後、ランダムフォレストアルゴリズムを使用してそのうち大きさ関連の2つと形の変化関連の2つ、合計4つの特徴を選別した画像バイオマーカーを開発しました。

これら4つの組み合わせは、治療開始6カ月時点で患者の全生存を相当に正確に予測しました。血液1滴も組織サンプルも必要とせず、既に撮影されたCT画像だけからバイオマーカーを抽出できるという事実は、この技術が病院の現場にいかに容易に組み込まれるかを示しています。

疾病ネットワークやマルチオミクスモデルがコンピュータ画面上で特定した標的候補は、それ自体は医薬品ではありません。細胞株やオルガノイド、実験動物を用いた実験でその標的を阻害した時に実際に病気が改善するかどうかを確認する段階を経なければならず、この段階を通過した候補のみが臨床試験設計へと進みます。トライアンフ試験やKEYNOTE試験のように既に他の目的で進行中の大規模臨床試験の組織スライドと画像データを人工知能で再び分析する探索的分析アプローチは、新しい臨床試験を最初から新たに設計することなくバイオマーカー候補を素早く検証でき、時間と費用を大幅に節約できます。

このように検証されたバイオマーカーが再び伴随診断の承認を受けて初めて、病院の診察室で実際の患者に使用できるようになります。

このように確認されたバイオマーカーは、臨床試験設計の方法そのものも変えています。以前は患者を無作為に分け、試験終了後にどの治療がより良かったかを判定していたのに対して、今は試験進行中もリアルタイムでバイオマーカー値を確認し、反応の悪い患者群は別の治療に切り替え、反応の良い患者群はそのまま続ける適応型臨床試験設計が増えています。

トライアンプ試験のように既に循環腫瘍DNAで患者を先に選別した後に本試験を開始する方法も、こうした流れの一つです。患者をより少なく迷わせながらも、より少ない人数でより迅速に結論に達することができるという点で、バイオマーカー発見は新薬開発のスピード自体を変える力を持っていると考えられます。

このように発見されたバイオマーカーが実際の処方につながるには、伴随診断（CDx）というゲートウェイを通過しなければなりません。特定のバイオマーカーを持つ患者にのみ特定の医薬品を使用するよう組み合わせて承認する検査ツールです。米国食品医薬品局が承認した伴随診断検査は2017年34件、2021年44件でしたが、2025年3月時点で188件となり、関連するバイオマーカーは45種類に増えました。2026年5月には、膀胱がん患者の血中循環腫瘍DNAで再発リスクを測定するナテラのシグナテラ伴随診断と、乳がん、大腸がん、非小細胞肺癌を同時にカバーするガルダントヘルスのガルダント360リキッド伴随診断が相次いで承認を受けました。

米国食品医薬品局はさらに、次世代シーケンシングベースの検査やポリメラーゼ連鎖反応ベースの検査など、既に安全性が確立されたがん伴随診断法を、現在の厳密な薬事申請前許可の代わりに、より軽い手続きで認可する方法について、2026年1月までの意見聴取を実施しました。しかし新薬と伴随診断が同じ時点で共に承認される例はまだ稀であり、医薬品は先に出たが指定する検査が後から続いたり、その逆の状況が実際の臨床現場で頻繁に起きています。

韓国でもこの流れに従う企業が増えています。アロンティアはマルチオミクスデータから新しい標的と候補物質を発見するREMEDY、低分子化合物を設計する

DRUG DESIGNER、病理画像から腫瘍微小環境を区別して読むPATHまで、新薬探索の全プロセスをつなぐプラットフォームを備えています。2026年米国臨床腫瘍学会ではルニットの他にもエスティキューブ、ティウムバイオ、イムノンシア、オンコニックセラピューティクスといった国内企業が人工知能バイオマーカーと合成致死率ベースの抗がん剤の臨床データを発表する予定です。

これほどの実績が積み重なったにもかかわらず、人工知能ベースの標的発見が病院の日常的な道具として定着するには、なお乗り越えるべき課題が残っています。最初に問題になるのは、学習に使用されるデータ自体の偏りです。現在蓄積されているゲノムおよび臨床データベースの大多数は西洋諸国または特定の先進国の集団に偏っており、その外の人種と地域の患者に同じモデルをそのまま適用すると、予測が外れることが頻繁です。データが一方に偏ると、モデルが出す答えも一方に偏らざるを得ません。健康格差をむしろ拡大する結果につながるリスクがあるわけです。

ディープラーニングモデルがなぜそのような答えを出すのかを人間が個別に追跡することが難しいという点も障害となります。複数の層に積み重ねられたニューラルネットワークと複数のモデルを混合する方法は予測性能に優れていますが、その中身を見るのが難しいため、一般的にブラックボックスと呼ばれています。患者の命がかかった処方の前で、そして新薬承認の可否を決める規制機関の前で、このような不透明さはそのまま受け入れられません。

そのため研究者たちはシャプレイ加法説明（SHAP）やライム（LIME）といった説明可能な人工知能技術と一緒に使って、モデルの判定がどの遺伝子変異やどの代謝経路から生じたのかを人間が理解できる言葉に変えようと努力しています。

法と制度の空白も依然として残っています。複数の病院や複数の国に散らばった患者データを集めて学習させる過程において、個人情報を守ることは最も敏感な事案です。欧州連合の一般データ保護規則(GDPR)や米国の医療保険の相互運用性と説明責任に関する法(HIPAA)を守るため、元のデータを外部に持ち出さず各病院内でのみモデルを学習させ、その結果値だけをやり取りする連合学習技術が徐々に広まっていますが、国ごとに異なるデータ基準と暗号化費用は依然として障害として残っています。

人工知能が自ら見つけ出した標的や設計したバイオマーカーをめぐって、誰を発明者として認めるか、その権利がアルゴリズムを作った会社とデータを提供した病院のどちらに属するのかをめぐる争いも、世界の至る所で起きています。米国食品医薬品局と欧州医薬品庁は、このような問題に対応し、モデルの透明性と品質管理、市販後の安全性の責任の所在を明確にする規制文書を継続的に見直しています。

業界は、2025年と2026年の間に、このように人工知能が見つけ出した標的から出発した薬物の臨床結果が10件以上出ると予測しており、その成績が蓄積されるほど、オミックスと文献データを結合して標的を探すアプローチ自体の信頼度も共に高まるものと見られます。

2025年と2026年を経て明らかになった流れを見ると、マルチオミックスと人工知能の結合は、計算能力を誇示する実験室のデモンストレーション段階を過ぎ、実際の患者の生存期間を分ける臨床現場へと移りつつあるように見えます。トライアンプ試験のHER2患者群、レントセルチブを服用した肺線維症患者群、MTAP欠損がんを狙った新しい阻害剤はすべて、異質なオミックスデータを一つにまとめ、文献の中の知識をナレッジグラフとして編み出した後に初めて得られた結果です。

異なる病院、異なる国から出たデータの標準を合わせてつなぎ合わせ、ブラックボックスモデルの判断の根拠を人間が理解できる言葉に置き換える作業が、今後数年の成否を分ける可能性が高いです。患者一人の組織スライド一枚、血の一滴から始まったシグナルが、地球の裏側にいる研究者のナレッジグラフとつながる現在の流れは、新薬開発という古い作業の裾野を再構築しているように見えます。

韓国政府も、マルチオミックスデータを国レベルで蓄積する事業を展開しています。疾病管理庁が率いる国家統合バイオビッグデータ構築事業は、国民100万人のゲノムと臨床情報を一つに集めることを目標に掲げました。2024年から2028年まで続く第1段階では、希少疾患の患者と重症疾患の患者、一般の参加者を合わせて77万2000人を集めることにしました。参加を希望する国民は、全国38の募集機関のうち1か所を訪れて同意書とアンケートを記入し、血液や尿などの検体を残した後、自身の臨床情報まで一緒に提供します。

事業団は、このように集めた資料が、韓国人に合った標的とバイオマーカーを探す基盤になるだろうと説明しています。このように集まったゲノムデータは、2026年から研究者たちに順次公開される予定であり、2026年4月末時点での参加者は16万5722人に増えました。

しかし、いざ新薬開発に役立つ重症疾患患者のデータは目標に遠く及びません。2026年4月末時点で、重症疾患の患者は目標の14万人のうち1万1274人しか集まらず、8パーセントをようやく超えました。希少疾患の患者もやはり、目標の4万7000人のうち4476人しか集まらず、9.5パーセントにとどまりました。事業団は、検体を集めて患者を説得する研究看護師が不足しており、医療陣が率先して参加する動機も弱いと、その理由を明らかにしました。

全世界のゲノム研究に使われたデータのうち、欧州系人口の割合が90パーセントを超えるという調査結果を思い浮かべると、韓国人と他のアジア人の遺伝子情報を埋め合わせるこの事業の重みは決して軽くありません。それにもかかわらず、実際に病気を患っている人の資料がこのように遅々として蓄積されないのであれば、データの偏りを正す機会もまた、長く手に入らないものと見られます。

発掘されたバイオマーカーが病院の処方までつながるためには、3つのハードルを順に越えなければなりません。第一は分析的妥当性であり、検査ツールが狙う標的を正確かつ一貫して捉えているかを見ます。第二は臨床的妥当性であり、その標的が患者の治療反応をどれほど上手く事前に教えてくれるかを確認します。第三は臨床的有用性であり、この検査を使えば患者の治療と結果が実際に良くなるのかを問います。米国食品医薬品局は、新薬候補の第3相臨床試験を設計する時から、コンパニオン診断検査と一緒に組み込むよう勧めています。

そうしてこそ、薬と検査が同じ日に並んで承認される可能性が高まるからです。検査の準備が遅れると、薬は先に出ても、その薬を使うべき患者を選び出すツールがなく、処方が遅れるという事態が起きます。韓国食品医薬品安全処も、コンパニオン診断医療機器として臨床的性能試験を行うには別途計画の承認を受けるようにし、似たような原則を守っています。

第2部

AI主導の創薬設計と物質探索

2.1

3Dタンパク質構造予測と構造ベース薬物設計 (SBDD)

2020年11月30日、コンピュータの画面の前には、世界中の構造生物学者数百人が集まっていました。第14回国際タンパク質構造予測学術大会、通常CASPと呼ばれるこの大会で、GoogleのDeepMindが開発した人工知能AlphaFold 2の採点結果が公開される瞬間でした。画面に表示された点数は90点をはるかに超えていました。エックス線結晶学や極低温電子顕微鏡で数ヶ月、場合によっては数年間かかって確認できたタンパク質の立体構造を、コンピュータが実験結果とほぼ同じように言い当てたのです。

半世紀近く生物学者を悩ませてきた難題の一つが、その日事実上解決されました。1970年代初頭、アミノ酸配列がタンパク質の立体構造を決定するという原理自体はすでに明らかにされていましたが、その配列だけを見て実際の形状を事前に計算することは、それ以来半世紀がたつまで未解決のままでした。

タンパク質は、アミノ酸という小さな部品が鎖のように長く繋がった物質です。細胞内で新しく合成された鎖は、糸玉のように無秩序に伸びていますが、水中で瞬時に精密に折り畳まれながら、固有の立体構造を持ちます。折り紙を思い浮かべると理解しやすいです。同じ正方形の紙でも、折る順番と角度が異なれば、鶴になることもあれば船になることもあります。

タンパク質も同様に、アミノ酸の並ぶ順序、つまり配列に従って折り畳まれ方が決まり、その形状がわずかにずれると、本来の機能を果たせなくなったり、見当違いの場所に付着して問題を引き起こしたりします。アルツハイマー病やパーキンソン病のように、タンパク質が誤って折り畳まれて凝集することで生じる病気が多く知られているのはこのためです。

製薬会社が新薬を設計する際に最初に知る必要のある情報は、まさにこの立体構造です。病気の原因となるタンパク質がどのような形をしているかを正確に知る

ことで初めて、そのタンパク質に結合して機能を阻害または促進する小分子を描くことができるからです。このように標的タンパク質の3次元構造に基づいて薬を設計する方法を構造ベース医薬品設計、略してSBDDと呼びます。過去には、この構造を得るために、タンパク質を結晶化してエックス線を照射し、回折パターンを読み取るエックス線結晶学、または試料を極低温で冷凍してから電子ビームを照射し、画像を重ねて見るCryo-EMと呼ばれる極低温電子顕微鏡実験を経なければなりませんでした。

どちらも熟練した労働力と高価な装置、そして数ヶ月から数年に及ぶ時間を必要とする作業でした。タンパク質を結晶化するプロセス自体がうまくいかず、アプローチを完全に変える必要がある場合も珍しくありませんでした。

AlphaFold 2がCASP舞台で証明した予測能力は、この古い作業方法を根底から揺さぶりました。2024年5月、GoogleのDeepMindとそのグループ企業であるIsomorphic Labsは、後継作であるAlphaFold 3を公開して、さらに一歩進みました。AlphaFold 3は単一のタンパク質の形状を予測するだけにとどまらず、遺伝情報を含むDNAとRNA、薬物候補となる小分子リガンド、金属イオン、さらには翻訳後の修飾までが混在する複合体全体の立体構造を、単一のニューラルネットワークで同時に予測します。

進化したEvoformerモジュールと拡散モデルを組み合わせた方法で、原子が散乱した雲のような状態から始まり、ノイズを徐々に除去しながら正解に近い配置を見つけていきます。公開時、学界では既存の方法と比べて相互作用予測の精度が50パーセント以上向上したという評価が出されました。

拡散モデルという言葉は耳慣れなくても、原理は昔のテレビのざあざあという画面を思い浮かべると理解しやすいです。画面いっぱいノイズだけが見えた状態からチャンネルが徐々に合わさり、明確な映像が現れるように、AlphaFold 3は原子がランダムに散乱したノイズに満ちた状態から始まり、段階的にノイズを取り除きながら実際の分子配置に近い映像を完成させていきます。この過程を数十回繰り返す間に、最初は曇っていた原子の雲がタンパク質とリガンドが互いに嵌合した明確な立体構造に絞られていきます。

こうして作られた予測構造は次々と積み重なり、AlphaFold Protein Structure Databaseという公共資産になりました。2025年にインターフェースを一新したこのデータベースには、2億1400万個を超えるタンパク質構造が登録されており、人間を含むほぼすべての生物種のタンパク質を1回の検索で見ることができます。実際に中国の研究チームはこのデータベースとAlphaFoldを用いて、ぶどう膜黒色腫患者に見られる異常な構造315個を発見し、新しい治療標的の手がかりを得ました。実験室でタンパク質構造1つを解明するのに数年かかった時代には想像しがたい規模です。

体感できる変化は速度から現れました。新薬候補発掘の第一関門である標的タンパク質構造確認作業は、AlphaFold 2が登場する前までは数ヶ月かかる作業でした。AlphaFold 2が登場してからこの作業は日単位に短縮され、前倒しされた時間は次のステップであるドッキングとシミュレーションにそのまま再投資されています。

最初に公開されたときAlphaFold 3は、研究者が直接コードをダウンロードして実行することができず、Googleが運営するAlphaFoldサーバーに配列を入力して結果のみを受け取る方式であったため、残念な点がありました。2025年に入ると、GoogleのDeepMindは非営利学術研究に限って推論コードをダウンロードしてモデルの重みをリクエストして得ることができるように門戸を広げ、AlphaFoldサーバーの1日の処理件数も既存の2倍に増やしました。その間に代替案を作る動きも活発でした。

中国のバイトダンスはAlphaFold 3を完全に再現したオープンソースモデルProteinixをリリースし、コロンビア大学のAlquraishi研究室はApache 2.0ライセンスの下でOpenFold3-previewを公開しました。学術目的だけでなく商用利用も無料で許可されています。一つの企業が独占していた技術が、この数年間で複数の研究チームが共有するオープンソースエコシステムに広がったわけです。

AlphaFold 3が精巧さで勝負するなら、Metaが開発したESMFoldはスピードで勝負します。大半の構造予測人工知能は、似たようなタンパク質を様々な生物種から見つけ出して並べる多重配列アライメントの作業を経てから、ようやく構造を計算します。この作業はデータベースを検索するだけでもかなりの時間がかかります。一方、ESMFoldは、数億個のタンパク質配列をあらかじめ学習しておいた

巨大言語モデルESM2をもとに、多重配列アライメントなしで配列を一つ入れるだけですぐに構造を提示します。

進化的に似た親戚がほとんどいないオーファンタンパク質や、新しく発見されたウイルスタンパク質のように、比較する対象が適当でない場合にESMFoldの強みが際立ち、大規模なスクリーニングや新規酵素の発掘現場で広く使われています。Metaは、この技術をもとに、自然界に存在するもののまだ知られていないタンパク質の構造を大規模に予測して公開したりもしました。

タンパク質構造予測のエコシステムをAlphaFoldとともに支えるもう一つの軸は、ワシントン大学のデイビッド・ベイカー教授の研究チームが作ったRoseTTAFoldです。その中でも最近のバージョンであるRoseTTAFold All-Atomは、AlphaFold 3と同様に、タンパク質と核酸、低分子リガンドが入り混じった複合体の構造を一緒に予測します。ベイカー教授の研究チームの本当の特技は、構造を当てることにとどまらず、希望する分子と結合する全く新しいタンパク質を白紙の状態から描き出すことにあります。拡散モデルを応用したRFdiffusion系列のツールがその役割を担います。

RoseTTAFoldは2021年、アミノ酸配列の情報と原子間の距離情報、3次元の座標情報をそれぞれ異なるニューラルネットワークの経路で同時に処理し、互いに情報をやり取りさせる3重ニューラルネットワーク構造として初めて発表されました。AlphaFold 2とほぼ同じ時期に出たこのモデルは、予測精度ではAlphaFoldに少し及ばないという評価を受けましたが、計算に必要なコンピューター資源が相対的に少なく、個別の研究室でも動かしてみることができるという実用的な強みを持っていました。

このツールは2025年12月3日、RFdiffusion3という名前でさらに一段階飛躍しました。原子単位まで扱うこのモデルは、タンパク質だけでなく、DNAや様々なリガンドと噛み合うタンパク質を一緒に設計することができ、以前のバージョンよりも推論速度が10倍近く速くなりました。自然界で長い時間をかけて洗練されてきた酵素に近づく性能の新しい酵素をコンピューターの中で描き出せるという発表は、学界に小さくない反響を呼び起こしました。産業用酵素を新しく開発するのに何年もかかっていた周期が、コンピューターの中で候補を先に描いてみて、

実験室では検証だけを行うという方式に変わる可能性があるという意味だからです。

同じ研究室はそれよりも先に、固定された形がなく扱うのが難しい本質的無秩序タンパク質にくっつく結合剤の設計に成功した論文を国際学術誌Natureに発表し、抗体の重い鎖と軽い鎖の全体をコンピューターだけで設計するRFantibodyの研究結果も一緒に提示しました。抗体はすでに様々な難病の治療に主力薬として使われているだけに、抗体の骨組みをコンピューターの画面上から設計できるという事実は、バイオ医薬品の開発戦略自体を変えてしまう潜在力を持っています。構造を読み取る人工知能から、希望通りに構造を作り出す人工知能へと重心が移りつつあるという意味です。

このような技術を長く待ち望んできた場所こそが、標的化不可タンパク質と呼ばれる領域です。はっきりと窪んだ結合ポケットがなく、薬物が掴むべき取っ手が見当たらないタンパク質、細胞の中で形を絶えず変える本質的無秩序タンパク質、表面が広くて平らであるため小さな分子一つでは到底防ぎきれないタンパク質とタンパク質の間の結合部位、このような対象が長い間、新薬開発の候補リストから除外されてきました。静的な結晶構造の写真一枚だけでは、このようなタンパク質を設計の対象にするのが難しかったためです。数百万個の配列と構造のデータを学習した生成型人工知能は、この壁を別の方式で乗り越えています。

本質的無秩序タンパク質という名前は馴染みがなくても、意味は難しくありません。大半のタンパク質が折り紙のように決まった形にピタッと折りたたまれるのとは違い、このタンパク質は固定された形態がなく、水の中でずっとゆらゆらと揺れ動いています。手に掴めるような形がないため、合わせてみる鍵穴も特になかったことになり、だからこそ長い間、新薬開発者たちの関心の外に留まっていました。2026年には、ポケットが浅い標的だけを別に集めて性能を競うShallowBenchというベンチマークも登場し、標的化不可タンパク質を扱う人工知能モデルの実力を綿密に見分け始めました。

タンパク質には正門だけでなく裏門もあります。薬物がよく狙う活性部位が正門だとするなら、アロステリック部位は活性部位から離れたタンパク質の表面にある裏門に該当します。この裏門に小さな分子が少し手を触れると、タンパク質全

体の立体の形が変わり、酵素の活性や信号伝達が調節されます。タンパク質言語モデルが抽出した特徴と結合ポケットの形を一緒に分析するDeepAlloのようなツールは、結晶学の写真一枚では現れないキナーゼやGタンパク質連結受容体のアロステリック部位を見つけ出します。正門が固く閉ざされていて手を出せなかったタンパク質も、裏門を利用すれば機能を選択的に調節できるようになるわけです。

実際の成果も続きました。AlphaFoldが明らかにした上皮成長因子受容体、すなわちEGFR突然変異の構造とKRASタンパク質の立体の形の変化情報は、肺がんと乳がんの患者に使われる精密標的治療薬であるソトラシブとエルロチニブの効能を整えるのに実際に使われました。長い間、KRASは表面が滑らかで取っ手がないタンパク質として悪名高かった標的です。KRASががんに関連する遺伝子だという事実自体はずっと前に明らかになりましたが、手に掴める結合ポケットを見つけられず、数十年間、標的化不可能な遺伝子というレッテルを剥がせずにいました。構造情報が綿密になるにつれ、このようなタンパク質も少しずつ攻略する方法を見つけ出しています。

最近では、AlphaFold系モデルからさらに一歩進んで、一つの正解構造ではなく、タンパク質が行き来できるいろいろな形を一度に予測するアンサンブル生成研究も活発です。タンパク質が結合前と後に異なる形を行き来するという点を考えると、写真一枚より複数のスナップショットを並べる方が、実際の動きにより近い試みであると言えます。

ただし、AlphaFoldが出す構造は、一瞬を撮った写真一枚に近いです。実際のタンパク質は、写真の中の静止した姿ではなく、体温と水分子に囲まれたまま、絶え間なく揺れ、ねじれている生きた構造物です。写真ではなく、動画で見てこそ本当の動きが見える、それと同じ理由です。この動きを計算で再現する技法が分子動力学、略してMDシミュレーションです。

MDは、ニュートンの運動法則に従って、原子一つ一つの位置と速度を非常に短い時間間隔で絶えず更新しながら、ピコ秒からミリ秒に至る時間の間に、タンパク質がどのように揺れ、リガンドが結合する瞬間に周囲の形がどのように変わるかを追跡します。

人工知能ベースの分子ドッキングが結合のスナップショットを速く複数枚撮り出す役割をするのに対して、MDシミュレーションはそのスナップショットが実際の体内環境でどの程度長く、どの程度安定的に維持されるかを検証する役割を担います。リガンドが結合することでタンパク質の形が一緒に変わる誘導適合現象や、水分子が結合部位を出入りして作る微細な効果まで、MDシミュレーションは捉えます。ドッキング結果だけを見て有望だと判断していた候補物質が、MDシミュレーション段階でわずか数時間で結合が解けてしまう場合も珍しくありません。このような検証を経た後によりよく、候補物質は次の段階へ進む資格を得ます。

化学結合が新しく作られたり、切れたりする瞬間、つまり酵素が反応を触媒したり、薬物が標的タンパク質に共有結合で付着する瞬間の電子の動きは、MDシミュレーションでも正確に計算することができません。このような微細な電子の世界を扱うには、量子力学方程式が必要です。しかし、タンパク質全体を量子力学で計算しようとする、スーパーコンピュータでも対応するのが難しい時間がかかります。そのため、結合が実際に起こる活性部位周辺だけを精密な量子力学で計算し、残りの巨大なタンパク質骨格と水分子は計算コストが安い古典力学で処理する、量子力学と分子力学の混合計算、略してQM/MM技法が使われます。

建物全体を同じ精密度で建てるのではなく、人間の手が直接触れるドアノブだけに精密部品を使い、残りの壁と柱は標準素材で建てるのと似た発想です。

最近では、量子機械学習とグラフニューラルネットワークを組み合わせた計算法が、QM/MMにかかる費用を大きく下げています。POSTECH研究チームは、電荷平衡原理とグラフニューラルネットワークを結びつけた機械学習ポテンシャルを開発して、量子力学レベルの精度に近づきながらも、計算速度は非常に速い方式を提示しました。2026年のある新薬開発カンファレンスで、Schrödinger が公開したワークフローは、物理ベースのシミュレーションと生成形人工知能を結びつけて、わずか4時間30分で1万2000個を超える化合物候補を作り出したと発表し、注目を集めました。

前臨床段階で、化合物の毒性と薬動学的特性までコンピューター内で事前に選別するインシリコ検証の信頼度がそれだけ高まっているという意味です。

このようにいろいろな技術が積み重なることで、人工知能が新薬候補を見つけていく手順は、およそ四段階を経ます。研究者はまず、病気を引き起こすタンパク質のアミノ酸配列を入力します。次にAlphaFoldのような人工知能がその配列だけでタンパク質の3次元構造を予測します。構造が手に入ったら、分子ドッキングプログラムが候補化合物一つ一つをその構造の結合部位に合わせて見ながら、よく合致する候補を選別します。

最後にQM/MM計算とMDシミュレーションが結合部位周辺の電子の動きと時間による安定性まで精密に再び確認しながら、最終候補を整えます。配列一つから出発して検証された候補物質に至るこの過程が、過去には何年もかかっていたのに対し、今は数週間か数か月単位に圧縮されています。

このフローの中で、第3段階である分子ドッキングは、鍵穴に鍵を一つ一つ合わせてみるのと似ています。標的タンパク質の表面にある溝を鍵穴だとすると、候補物質はその穴に合う鍵を見つけるタスクに当たります。ただし、新薬開発の現場では、一つの鍵穴に合う鍵を何百万、何億個の候補の中から同時に試さなければならないという点が異なります。過去には、AutoDock VinaやGlideのように、物理的エネルギー関数と経験的スコアリングに依存する古典的プログラムがこの仕事を担当していました。

最近では、深層学習とグラフニューラルネットワーク、拡散モデルを組み合わせたDiffDock、SurfDock、Chai-1、Boltz-1のような人工知能ドッキングモデルが次々と登場して、標的タンパク質の表面形状と静電特性を学習した後、柔軟な結合姿勢をはるかに速く見つけます。何億個に達する化合物ライブラリーを調べる仮想スクリーニング時間が、このおかげで大幅に短縮されました。

タンパク質とタンパク質がかみ合うドッキングでは事情がやや異なります。ZDOCKのように物理法則に基づく伝統的プログラムは、結合過程で起こる立体構造変化を細密に反映する一方で、複数の深層学習モデルは結合前後の形が固定されていると仮定する静的な形状合わせに止まり、実際の結合環境を歪める場合があります。ただし、タンパク質と小さい分子を結ぶドッキングでは、最近の雰囲気異なります。

2025年に発表されたPoseXという新しいベンチマークは、訓練データになかった見知らぬタンパク質とリガンド組み合わせを混ぜて試したのですが、ここで一部の人工知能モデルが伝統的な物理ベースの方法を上回る結果を示したこともあります。どちらか一方が完全に優位を占めたと言うのはまだ時期尚早です。

2025年6月、Massachusetts Institute of Technologyと新薬開発企業のRecursionが共同で発表したBoltz-2は、この流れに新たな動きを加えました。構造を当てることにとどまらず、結合の強さ、つまり結合親和性までも同じモデルの中で予測する、初めての共同折り畳みモデルだからです。自由エネルギー摂動と呼ばれる精密な物理ベースの計算に匹敵する正確さを、最大1000倍の速さで出すと評価され、学界と業界の関心を同時に集めました。

Boltz-2は最初オープンソースで公開されましたが、2026年上半期に出た後続バージョンのBoltz 2.1は、APIの形でのみ提供される閉鎖型へと転換されました。技術の流れが開放から再び商業化へと移っていく様子を示す事例です。

公開範囲をめぐる綱引きは、この分野全体を貫く流れのように見えます。

AlphaFold 3が学界に門戸を広げ、ProtenixやOpenFold3-previewのような代案が後に続く一方で、肝心の結合親和性まで予測していたBoltz-2は後続バージョンで門を閉ざしました。優れた予測モデルを一つ作るのと同じくらい、そのモデルを誰にどれくらい開いておくかを決めることも、この生態系の勝負を分ける要因になっているようです。

華やかな成果の裏には、まだ解決されていない欠陥もはっきりしています。

PoseBustersという厳しい検証基準でAlphaFold 3の結合姿勢の予測結果を調べてみると、基本データセットでは76パーセント余りであり、採点基準を緩く変えても84パーセント前後にとどまります。本当に深刻な問題は別にあります。ディープラーニングモデルが出した結合姿勢の中には、原子と原子が物理的に重なってしまったり、鏡に映った形のように方向が反対である光学異性体を区別できなかったりする、現実にはあり得ない構造が珍しくなく混ざっています。

学習データと似た結合のタイプに出会うと、人工知能は驚くほどの正確さを見せますが、一度も見たことのない見知らぬタンパク質や全く新しい化合物の前では、

予測力ががたと落ちる様子も、いくつかのベンチマークで繰り返し確認されました。そのため、人工知能が出した結合姿勢を物理ベースの計算でもう一度整え、不自然な構造を正す手順が、正確さを引き上げる実質的な解決策として定着しつつあります。

この問題を正面から狙ったのが、2021年にGoogle DeepMindから分社して発足したIsomorphic Labsです。2026年2月にこの会社が公開したIsoDDEは、訓練データとの配列類似度が20パーセントにも満たない、まさに未知のリガンド結合の事例で、AlphaFold 3よりも約2倍高い正確さを見せたと発表しました。Isomorphic LabsはNovartis、Eli Lilly、Johnson & Johnsonのような大手製薬会社と手を結んで新薬候補を共同で開発する一方で、腫瘍学や免疫学の分野で独自の候補物質も育ててきました。

Demis Hassabisは2026年1月のDavos Forumで、当初2025年に予定していた初の臨床試験の日程が2026年末に遅れたと明かしながらも、患者への投薬が遠くないという考えをはっきりさせました。2026年初めにはJohnson & Johnsonが、AlphaFold 3をもとにタンパク質とタンパク質の間の結合を防ぐ新しい阻害剤を設計する、より深い協力に乗り出したと発表したりもしました。

国内でも関連研究が続いています。KAISTの研究陣は、薬物が標的タンパク質にくっつくかどうかを超えて、くっついた後に実際に望む通りに働くかまで一緒に予測する人工知能を披露しました。結合するかどうかだけを見ていた従来のモデルの目線を、一段階引き上げた試みです。Seoul National University化学部のSeok Cha-ok教授が2020年に設立したスタートアップGaluxは、AlphaFoldが登場する前から国際タンパク質構造予測大会で最上位圏の成績を収めた経歴をもとに、物理の法則を骨組みにした人工知能でタンパク質を設計します。LG Chemとは抗がんタンパク質を共同で開発する協力を進め、これまでに集めた投資金は680億ウォンに達します。

同じ時期に国内で門を開いたSimplexも注目に値します。グローバル製薬会社のGSKやBMS、Amgenで20年以上も新薬研究プラットフォームを開発してきたJo Seong-jin代表が設立したこの会社は、予測結果が出た理由を一緒に見せてくれる、説明可能な人工知能を前面に押し出し、臨床段階に入る前に失敗する可能性が高

い候補をあらかじめふるい落とすことに強みを持っています。

構造を直接確認する実験技法も人工知能と手を結びました。極低温電子顕微鏡の映像を処理するソフトウェアのRELIONやCryoSPARCには、粒子を選び出して分類し、3次元で再び組み立てる作業を自動化する人工知能が導入され、何週間もかかっていた処理時間を数時間に減らしました。形がぼやけて解釈しにくい密度マップの前では、AlphaFoldが出した予測構造が参考マップの役割を果たし、モデルを組み立てる過程を助けます。実験機器が作り出したぼやけた生データと、人工知能があらかじめ描いておいた正解に近い地図が、互いに前後を合わせていく構造に変わったと言えます。

Pfizer、Novartis、AstraZenecaをはじめとする大手製薬会社が、AlphaFoldとRoseTTAFoldを社内の新薬開発パイプラインに取り入れたのも、こうした流れと通じています。構造生物学の部署が、予測モデルを動かす電算チームと同じオフィスを使う光景も、今では珍しくありません。実験と計算が交互に互いを検証する循環構造が、一つの標的タンパク質を把握するのにかかる時間を減らし続けています。

これらすべての予測の土台には、ディープラーニング特有の不透明さがまだ残っています。ニューラルネットワークがなぜ特定の原子配置を正解として選んだのか、その判断の内側の論理を、研究者が一行一行たどりながら説明することはまだ簡単ではありません。結合親和性の点数一つがどんな物理的な根拠から出たのかを振り返りにくいこのような特性のため、計算結果は最終的な答えではなく、実験で絞り込むべき候補のリストとして扱われています。

このような一連の流れ全体に公式な指標を立てた出来事が、2024年のノーベル化学賞です。スウェーデン王立科学アカデミーはその年のノーベル化学賞の半分をデイビッド・ベイカーに授与し、残りの半分をデミス・ハサビスとジョン・ジャンパーに共同で授与しました。ベイカーは自然に存在しない新しいタンパク質を設計した功績が、ハサビスとジャンパーはアミノ酸配列だけからタンパク質の構造を予測した功績が認められました。賞金1100万スウェーデン・クローナが実験室の研究者1人ではなく、人工知能企業の最高経営責任者と研究員、そして大学教授に並んで手渡されたという事実そのものが異例的でした。

化学と生物学の長年の課題をコンピュータの計算で解き明かした成果を、科学界そのものが最高の権威として認めたわけであり、その後、製薬業界全般が人工知能ベースの構造予測を周辺的な実験ではなく医薬品開発の標準的な手順として受け入れる流れに速度がつけました。かつてコンピュータ科学側の実験的試みとしてのみ見なされていた方法論が、わずか数年のうちにノーベル賞を受ける正統な科学の仲間入りをしたわけであり、このような認定は逆に、より多くの若い研究者をこの分野へ引き寄せる役割も果たしています。

構造を得る方法が実験1つから計算層へと増えていくにつれ、医薬品開発初期段階のボトルネックは、もはや構造を知らないことにあるではありません。溢れ出ている予測構造と結合候補の中から何を実験台の上に最初にのせるかを選ぶ判断へとボトルネックが移っていつているようです。

タンパク質1つの配列を入力して短い時間のうちに予測構造を受け取ること、その構造の上で数億個の化合物を試験してみること、有望な候補を物理計算で再度検証することが、今はいくつかの製薬会社とスタートアップの日常的な業務として確立されました。ただしコンピュータ画面の中で完成した結合のポーズが、実際の人間の体の中でもそのまま維持されるかどうかは、結局のところ表面プラズモン共鳴やX線結晶学のような伝統的な実験台の上でもう一度確認する必要があります。予測と実験がやり取りするこの往復作業が速くなるほど、長い間手をつけることができないと思われていたタンパク質のリストも少しずつ短くなってきています。

構造を正確に当てることと、その構造の上で結合の強さを正確に測ることは別の問題です。AlphaFoldやBoltz系列のモデルが原子1つ1つの位置を実験結果とほぼ同じように当てたとしても、その結合が体の中でどのくらい強くどのくらい長く維持されるかを示す結合親和性の予測は、依然として答えを出すのが難しいままです。結合親和性は静止した写真1枚では捉えられない値だからです。水分子が結合部位を出入りする程度、鎖が折りたたまれながら失う自由度、アミノ酸とリガンドがそれぞれどのような電荷状態にあるかまで、すべて足したり引いたりして初めて1つの値が出ます。

Boltz-2はある検証では物理ベースの計算と並ぶ成績を出しましたが、形態が絶えず変わるタンパク質を対象とした別の検証では、実際の活性がある化合物さえランダムに選んだものと同様の成績を出すこともありました。実験で結合の強さを測定する際にさえ1kcal/mol前後の誤差が含まれるという事実を考えると、計算がいかに精密になったとしても、この下限を下回るのは始めから置かれている壁であり、そう簡単には下がりにくい状況にあるわけです。

このような予測ツールが広く広がるにつれて、全く異なる懸念も一緒に育ってきました。自然に存在しないタンパク質を描き出すツールが誰もが手の届く場所に置かれると、危険なタンパク質を配列をわずかに変えて監視網を避ける道も一緒に開かれるからです。アメリカ政府は2025年5月に行政命令を通じて生物学研究の安全と保障を改善するよう指示し、政府支援研究に使う合成核酸の審査基準を改めて定めました。5カ月後の10月には、マイクロソフト最高科学責任者エリック・ホビッツが率いた研究チームがその審査網に開いた穴を国際学術誌『サイエンス』に発表しました。

彼らはリジンとボツリヌス神経毒素を含む危険なタンパク質72種の配列を公開されている人工知能設計ツールで変更してコンピュータの中で76,000個以上の設計図を作成してみました。そしてその中でも相当数が遺伝子合成会社が使っていた既存の監視ソフトウェアをそのまま通過するところだったという事実を確認しました。実際にはタンパク質を合成しなかったコンピュータ内での実験でしたが、研究チームは国際遺伝子合成コンソーシアムと共にソフトウェアを改善し、その後も変形された配列の約3パーセントは依然として検出されませんでした。

計算が出す答えと、その答えを実際に信頼してもよいかどうか確認する能力の間のズレは、この分野全体で繰り返される問題のようです。結合力を測定する際にも危険な配列を選別する際にも、作り出す能力が検証する能力より先に走って行く構図は、今後もしばらくの間続く可能性が大きいです。

2.2

ディープラーニングによる候補物質の生成 (De Novo Design)

2022年、中国の上海にある製薬会社の実験室で、1台のコンピューターが特発性肺線維症という病気を治す一つの物質を描き出しました。特発性肺線維症は、肺の組織が原因もなく徐々に硬くなり、息をするのが苦しくなる難病です。診断を受けた後の平均生存期間が3、4年程度にとどまるほど重い病気であるため、新しい治療薬に向けた患者たちの待ち望む思いは切実でした。この物質を設計した会社はインシリコ・メディシン(Insilico Medicine)であり、標的としたのは細胞の中で炎症と線維化の信号を運ぶ酵素であるTNIKで、このようにして生まれた物質にはISM001-055という名前が付けられました。

標的を決めた時から人を対象とした第1相臨床試験に入るまでにかかった時間は30ヶ月、3年にも満たないものでした。人工知能が標的の発掘から分子の設計までを引き受け、臨床試験の進捗にまで至った事例は、当時としては世界的にも数えるほどでした。2025年、国際学術誌のネイチャー・メディシンと米国胸部学会の学術大会で第2a相臨床試験の結果が公開されました。中国の21の病院で患者71人を集めて行われたこの試験は、安全性という目標を満たし、肺がどれくらい多くの空気を吸い込んで吐き出すかを測る指標である努力肺活量も、投与量が増えるにつれて一緒に良くなりました。

2026年7月7日には、第3相臨床試験に入るという発表が出ました。人が何十年も試行錯誤して探し求めていた新薬の候補を、コンピューターが白紙から描き出し、その結果を患者の体で一つずつ確認していっています。この物質はその後レントセルチブという国際一般名まで受けましたが、国際一般名は世界保健機関が付与する公式の薬物の名前であり、一つの会社の実験室のコードネームを超えて、全世界が一緒に呼ぶ正式な医薬品として扱われ始めたという意味でもあります。

新薬の候補の一つが標的の発見から動物実験を通過するまで、伝統的な方法では4年から6年を超える時間と莫大なお金がかかります。製薬業界でよく引用される推

算では、新薬の一つを市場に出すまでにかかる費用が20億ドルを超えます。この過程で大部分の候補物質は途中で脱落してしましますが、業界の人々はこの区間を死の谷と呼んでいます。このような流れを指して、業界ではエルームの法則 (Eroom's Law) という言葉も使います。半導体の性能が絶えず良くなるというムーアの法則を逆さまにひっくり返した名前で、新薬の一つを開発するのにかかる費用と時間は、技術が発展してもむしろますます増えるという意味です。

臨床試験の段階まで上がった候補物質でさえ、最終承認まで生き残る割合は10に1つにも満たないというのが、業界で長く通用している経験則です。人工知能が初めてこの産業に入ってきた時に任された役割は、ZINCやEnamine REALのように数十億個の化合物が積まれた倉庫を素早く探しまわることでした。コンピューターは数兆個にのぼる化合物を瞬きする間にざっと見て有望な候補を選び出すことができましたが、このようにして見つけ出した物質はどうせ誰かがすでに作って倉庫に入れておいたものにすぎませんでした。倉庫の中に正解がなければ、いくら早く探しても正解を見つけることはできません。

ディープラーニングが新薬開発に新しく持ち込んだのは、倉庫をより早く探す技術ではなく、倉庫にないものを最初から新しく作り出す能力です。科学者たちは、薬になり得る性質を持った分子が理論的に10の60乗から10の100乗個にのぼると推算しています。宇宙にある星の数が10の24乗ほどと見積もられるので、化学が抱いている可能性の大きさは、夜空の星よりも数え切れないほど大きいです。このような数字は組み合わせの力から出ます。レゴブロック10種類だけを持って、20個を繋ぎ合わせる方法の数は、すでに私たちの頭では数えられないほどに膨れ上がります。

実際にレゴの会社が計算した結果、完全に同じ形をしたブロックを6個だけ積んでも出てくることができる形は9億通りを超えます。ところが、実際の分子を構成する原子の種類は100種類を超え、これらの原子を鎖で、枝で、輪で繋ぎ合わせる方式まで加えると、その組み合わせはレゴよりもはるかに猛烈な速度で膨れ上がります。人が実験室で1世紀の間に作って試験してみた化合物の数は、この大きな空間に比べると数え切れないほど小さいです。

この広大な空間でコンピューターが新しい分子を描くには、まず分子を数字と文字に書き写す方法がなければなりません。広く使われた方法がSMILESです。SMILESは炭素や酸素、窒素のような原子と、その間の結合をアルファベットと記号のいくつかで書き取る表記法で、分子を短い文章のように1行で解きほぐして書きます。人が言ってあげる通りに書き取るように、コンピューターはSMILESという文字を見て、分子の姿を読み書きする方法を身につけます。例えば、コーヒーやお茶に入っているカフェインはCN1C=NC2=C1C(=O)N(C(=O)N2C)Cという文字列一つで丸ごと表現されます。

炭素と窒素、酸素がどんな順序で繋がり、どこで輪を作るのかが、この短い行の中に余すところなく込められているのです。問題は、文字の行がとても敏感だということにあります。一つの文章で綴りが一つだけ間違っても意味が通じないように、SMILESの文字列もたった1文字だけずれると、化学的に存在できない偽物の分子を作り出します。生成モデルが数万個の候補を吐き出すたびに、このような偽物が混ざって出てくるせいで、研究者たちは長い間この表記法の弱点と格闘しなければならませんでした。

SELFIESという表記法は、この弱点を正面から狙って作られました。SELFIESで書いた文字列は、どのように組み合わせても、必ず実際に存在できる分子に書き移されるという点を数学的に保障します。端の形が互いにずれないようにあらかじめ削っておいたパズルのように、SELFIESは間違っても繋ぎ合わせる余地そのものをなくしてしまったのです。

この表記法が広く使われ始めることで、人工知能モデルを訓練する過程が目に見えて安定し、文字列の一つ一つを信じて分子に変えて使えるようになりました。ただし、SMILESであれSELFIESであれ、1行の文字列だという点は同じで、原子たちが実際の空間でどんな形に置かれるのかは、別途教えてあげることはできません。

分子を1行の文字列ではなく、網のように絡み合った構造そのものとして扱う技術がグラフニューラルネットワーク(GNN)です。原子ひとつひとつを点として置き、原子間の化学結合をその点を結ぶ線として描くと、分子は地下鉄の路線図に似たような網目の形になります。グラフニューラルネットワークはこの路線図をその

まま読み込み、各原子周辺の電子環境と立体的な形状と一緒に学習します。グラフ畳み込みニューラルネットワークやグラフアテンションネットワークと呼ばれるモデルは、こうして習得した情報に基づいて、ある分子が標的タンパク質にどれほど良く付着するか、水にどれほど良く溶けるかを、かなり正確に事前に教えてくれます。

数万から数十万に及ぶ候補分子を人が一つ一つ目で見ることなく、コンピュータが先に選別する仮想スクリーニングの速度と精度は、この技術のおかげで大幅に高まりました。

グラフ畳み込みニューラルネットワークは、隣接原子の情報を複数の層に渡って集め、分子全体の特徴を要約する方式で動作し、グラフアテンションネットワークはここにさらに、どの隣接原子がより重要かを自ら重みとして付与します。標的タンパク質の結合ポケットの真ん中に置かれた原子と分子の端に置かれた原子は、結合力に与える影響が異なりますが、アテンション方式はこのような違いを自ら区別します。このような予測能力のおかげで、研究者たちは数百万の候補を一つ一つ合成して試験管に入れて試さなくても、コンピュータ内で大部分をまず選別し、本当に可能性のある数百個だけを実験台に乗せることができるようになりました。

分子を読み書きする言葉を持ったなら、次はその言葉で実際に新しい文を作る仕事です。この創作を担当する3つの技術が変分オートエンコーダ(VAE)、生成的敵対ネットワーク(GAN)、そして拡散モデルです。変分オートエンコーダは、多数の化合物を学習して、分子ひとつひとつの複雑な形状を1つの座標に圧縮するエンコーダと、その座標を再び完全な分子に蘇らせるデコーダで構成されています。このように圧縮された座標空間の中では、座標をほんのわずかに動かしても、元の分子と骨組みは似ていながら、溶けやすさや毒性といった性質がわずかに改善された分子を得ることができます。

この方式は、料理の材料の配合をちょっとずつ変えながら新しい料理を探す料理人のノートに似ています。塩を小さじ半杯だけ減らし、砂糖をちょっと足すと、全く違う料理ではなく、元の料理と味の骨組みは同じでありながら、ずっと改善された料理が出てくるようなものです。ジャンクションツリー変分オートエンコー

ダという発展したモデルは、原子を一つずつ付ける代わりに、すでに化学的に意味が通じる小さな塊を単位として分子を作るため、結果がずっと現実的です。

生成的敵対ネットワークは、これと異なる方式で競い合います。偽りの分子を絶えず生み出す生成器と、その分子が本物の医薬品データベースにあったものなのか、生成器が作り出したものなのかを見分ける判別器が互いに対抗しながら一緒に実力を高める構造です。生成器は判別器を騙そうと、ますます本物と区別しがたい分子を生み出すようになります。インシリコメディスンは、この生成的敵対ネットワークに強化学習を加えたGENTRLというモデルで、線維症と神経変性疾患で重要な標的であるDDR1キナーゼを阻害する候補物質をわずか21日で見つけた実績があります。

もっとも、この技術には明らかな弱点もあります。訓練過程が数学的に不安定なため、モデルが多様性を失い、いくつか似た形の分子ばかりを繰り返し出すモード崩壊という現象に陥りやすいのです。このような不安定さのため、ここ数年の新薬設計分野では、重心が生成的敵対ネットワークから拡散モデルへと移って行く流れが明らかです。

拡散モデルは、写真を復元する過程からヒントを得た技術です。きれいな写真にノイズを少しずつ加えて完全にぼやけさせ、そのノイズを逆に一層ずつ剥がしながら元の写真を蘇らせる訓練をします。この原理を分子にそのまま当てはめると、ランダムなノイズの塊から出発し、ノイズを段階的に取り除く間に、はっきりとした分子構造が浮かび上がるようになります。薬が体内で実際に働くには、平面ではなく立体で標的タンパク質にぴったり合致する必要がありますが、最近の3次元拡散モデルは、分子を回転させたり移動させても、その性質が乱れないように数学的に整えられた構造を備えており、非現実的にねじれた分子を最初から選別し出します。

GeoDiffやDiffDock、DiffSBDDといったモデルは、標的タンパク質の結合ポケットの形を読み取り、その凹凸にぴったり合う原子配置と3次元座標を一緒に描き出します。DiffLinkerというモデルは、別々に離れた化合物の断片を3次元空間で自然に繋ぎ合わせる構造さえも作り出します。2026年初めに発表されたGlintDMモデルは、結合力と薬物らしさを同時に最適化しながらも、計算コストを大幅に削

減し、同じ年に出たDiffGuiモデルは、コンピュータ内にとどまらず、実際の実験室で物質を合成して試験する湿式実験まで経て、設計した分子が本当に働くという事実を確認しました。

標的タンパク質の結合ポケット情報を直接利用して分子を作る方式を構造ベース生成と呼び、標的の立体構造なしに、すでに知られた活性物質の特徴だけで似た性質の新しい分子を作る方式をリガンドベース生成と呼びます。タンパク質構造が実験で解明されていれば、構造ベース方式がより正確な結果を出しますが、構造をまだ知らない標的も多いために、両方式は今日でも並んで一緒に用いられています。

拡散モデルや敵対的生成ネットワークがいくら多くの分子を生み出しても、その中には標的によく結合するものの体には有害な物質が混ざっているものです。複数の条件を同時に満たす方向へと導いてくれる技術が、強化学習です。強化学習とは、コンピュータ内のエージェントが分子構造を少しずつ変えてみて、その結果を採点する評価関数から報酬と罰を受け取り、点数を多く稼ぐ方法を自ら習得していく学習方式です。このとき、点数は薬効だけでつけられるわけではありません。体内でどれくらいよく吸収され、分布し、代謝されて排出されるか、毒性はないかを網羅する指標であるADMET、すなわち吸収・分布・代謝・排泄・毒性が、水に溶ける度合いや脳への関門である血液脳関門の透過性、心臓への負担とともに点数に反映されます。

AstraZenecaの研究陣が改良したREINVENT 4というプラットフォームは、言語モデルのように分子を扱う人工知能に強化学習を組み合わせ、標的への効果は高めながら、見当違いの場所に作用して生じる副作用のリスクは下げる方向へと分子構造を繰り返し整えていきます。2025年以降には、強化学習を拡散モデルと完全に一つにまとめ、複数の標的に同時に適合する分子を立体構造まで備えた状態で設計する研究も続いています。

2025年から2026年の間には、強化学習を他の技術と完全に一つにまとめる研究がぐっと増えました。MDRLという方法は、拡散モデルと強化学習を丸ごと結びつけ、互いに異なる複数の標的に同時に適合する分子を、立体構造まで備えた状態で設計します。POLOという方法は、人間の好みを何度も対話するように反映さ

せながら候補物質を整えていき、MolMemという方法は、過去に試みた結果を記憶しておき、少ない試行回数でも目標に近づけるよう助けます。試行錯誤をコンピューターの中で前もって経験させ、その記憶を次の設計に使う方向へと、強化学習自体も急速に進化しています。

このように分子を造る技術は、小さな化合物にとどまりません。2025年、Chai Discoveryという会社が発表したChai-2モデルは、抗体を完全に白紙の状態から設計し出しました。すでに知られている抗体や結合物質を模倣せず、標的としたタンパク質の情報だけで、新しい抗体とナノ抗体を描き出したのです。ナノ抗体はラクダやアルパカのような動物の体にだけ自然に作られる、一般的な抗体よりもはるかに小さくて扱いやすい結合タンパク質です。研究陣が一度も扱ったことのない標的52個を選んで試験した結果、標的ごとに20個前後の候補しか作らなかったにもかかわらず、その半数で実際にくっつく抗体を見つけ出し、設計を始めてから実験室で結果を確認するまで2週間もかかりませんでした。

前の世代のコンピューター設計技術と比べると、成功率は100倍以上も跳ね上がりました。小さな分子を新しく造っていた生成技術が、今では体を守る複雑なタンパク質の武器である抗体設計にまで手を伸ばしました。このように完成した抗体候補のうち86パーセント以上は、実際の治療用抗体に匹敵する安定性と生産可能性を備えているという評価まで受け、同じ年の11月には、それに続く研究で、より扱いにくい標的を相手にした抗体設計の結果まで追加で公開されました。

表現が土台なら、生成はその土台の上で実際に手を動かす作業です。人間の化学者が数ヶ月かけて手作業で分子を描き、整えていた仕事を、今ではコンピューターが1、2日の間に数万通りの代案として代わりに描いて見せてくれます。拡散モデルと強化学習が手を取り合って目標に合った新しい分子を造り出した後、残る仕事は、その結果を世の中へと押し出すことです。

新しく造った分子をまだ誰も行ったことのない化学空間のほうへと押し進めていくことと、すでに世に出ている薬をまったく別の病気に再び使う道を探ること、この二つが共に成し遂げられる段階がまさに拡張です。この段階に至ってこそ、実験室の中に閉じ込められていた人工知能の想像力が、ようやく患者の前に置かれる薬へとつながるのです。

薬を再び使う人工知能は、病気と標的、薬物の関係を一つの大きな知識の網に編み上げるところから出発します。グラフニューラルネットワークと自然言語処理技術が数百万件に及ぶ医学論文や遺伝子発現のデータを調べ、人間の研究者がまだ目を向けていなかった薬と病気との隠れたつながりを見つけ出します。このような知識の網の一つには、数千種の疾病と数万種の薬物、そしてその間をつなぐ数百万個の関係が盛り込まれることもあります。イギリスのBenevolentAIは、この知識の網を分析し、関節リウマチ治療薬であるバリシチニブがコロナ19の引き起こす過度な炎症反応を鎮めると同時に、ウイルスが細胞の中へ入る道筋まで塞げるという事実を見つけ出しました。

この予測は実際の大規模な臨床試験につながり、結局、コロナ19の治療薬として緊急使用承認を受けるまでに進みました。すでに安全性が確認された薬に新しい使い道を付け加えるこのアプローチは、新しい分子を最初から作る道よりもはるかに短い時間で患者の前に届くことができるという点で魅力的です。

しかし、拡散モデルや強化学習がいくらもっともらしい分子を描き出しても、その絵を実験室で作り出せなければ何の役にも立ちません。コンピューターの画面上で原子と結合が論理的に辻褄が合うからといって、その分子を実際に合成できるという意味ではありません。建築設計図の上ではどんな形の建物でも描けますが、いざその土地に建物を建てるとなると、資材を調達できるか、その重さに耐える柱を立てられるかを再び確かめなければならないのと同じ理屈です。研究者たちの間では、この問題について「実験室の中の象」という言葉まで出ています。誰もが知っていながら、なかなか取り扱えない大きな問題という意味です。この隙間を埋める技術が逆合成計画です。

目指した複雑な分子を市場で安く入手できる原料物質まで逆に1段階ずつ細かく分解していく方式で、合成経路を見つけることも化学空間を探索することと変わらないため、候補となる経路の数は多くは数百万種に至ることもあります。

AiZynthFinderやAskCosのようなツールは、この数百万件の化学反応データを深層学習とツリー探索で分析して、安価で収率の高い経路を見つけってくれます。このツリー探索は複数の経路を先に試してみて、成功の可能性が高い道を選ぶ方式で機能しているのですが、強化学習が良い選択に報酬を与えながら学ぶ方式に似

ています。

合成可能性を高めることは言うほど簡単ではありません。既に良く知られた反応と材料だけで分子を作らせると作成は簡単になりますが、その結果は既存の化合物に似て新しさを失いやすいです。反対に化学空間の見慣れない隅々まで欲張って突き進めば、独創的な分子は得られても実際に作れるかどうかは保証しがたくなります。2025年のある研究は、この問題を多様性と合成可能性、生物学的効能という3つのウサギを同時に追わなければならない課題と整理したのですが、1つを追うともう1つを見落としやすいという点が、この分野の古い課題です。

2026年7月、インシリコ・メディシンはソウルで開催された国際機械学習学会（ICML）でChemCensorという新しい評価基準を発表しました。正解をどの程度当てただけを機械的に数える既存の方式の代わりに、実際の化学者が反応中心と官能基を見て判断する方式に近い形で逆合成経路の良し悪しを区別しようとする試みです。最近ではPrexSynやSynLlamaのように、合成可能性を結果物を選別する最後の段階ではなく、分子を作る過程そのものに最初から組み込もうとする研究も続いています。

一部の実験室はこのプロセスを完全にロボットに任せて、設計と合成と試験を人間の手を介さずに次々と進める自動化実験室を試験しています。人間が徹夜で実験を監視しなくても、ロボットが定められた順序を自ら繰り返す程度に、候補の1つを検証するのに要する時間も著しく短くなっています。

2026年現在、人工知能を通じて生まれた新薬候補パイプラインは、全世界で170以上に達するという集計もあります。しかし、薔薇色の見通しだけがあるわけではありません。2020年、人工知能が糖尿病治療候補物質の山から選び出したハリシンという物質は、強い抗生効果を持つ新薬候補として、メディアと学界の関心を大きく集めました。2025年に発表された後続研究は、ハリシンが複数の抗生物質に耐性を持つ細菌18種のうち17種を抑制するという事実を再度確認してくれました。

それでもハリシンは依然として人間を対象とした臨床試験に進むことができていません。体内への吸収度が低く、体外への排出が非常に速く、感染状況に応じて

効果が不安定だという弱点が後になって明らかになったからです。コンピュータの画面上では完璧に見えていた結合力が、生きている体の前では機能しない場合がこのように実際に起こっています。

人工知能の新薬開発をリードしていると自負していた企業も、この壁の前では例外ではありませんでした。米国のリカージョン・ファーマシューティカルズ（Recursion Pharmaceuticals）は2025年の1年間で、同社が発掘した複数の候補物質の開発を中止しました。脳血管奇形である海綿状血管腫を狙ったREC-994は、中間段階の臨床試験で期待に及ばない有効性が現れたことで開発が中止され、神経線維腫症2型を狙ったREC-2282も同じ道を辿りました。

計算能力がどんなに優れていても、その結果が直ちに良い薬に結びつくわけではないという事実を、この会社の軌跡が示しています。新薬開発において人工知能が減らすことができるのは、候補物質を見つけて磨くのに要する時間であって、人間の体が要求する安全性と有効性の閾値そのものではありません。

こうした支障が繰り返されるのには共通した理由があります。深層学習モデルは既に合成に成功して臨床試験も終えた医薬品データを学習資料とするのですが、このデータは化学空間全体で見ると非常に狭く偏った領域に集中しています。モデルがこの慣れた領域を離れて、一度も行ったことのない見慣れない分子を描き出すほど、予測が実際と食い違う危険も一緒に大きくなります。

研究者たちはこれをドメイン転移と呼び、コンピュータが出した結合力と安全性の予測を実験と臨床で再度確認する手順を省くことができない理由もここにあります。ですから新薬開発の現場では、コンピュータが選んだ候補1つだけに頼らず、性質が異なる複数の候補を同時に実験台に乗せてリスクを分散させる戦略も一緒に使っています。

この競争は韓国企業にも他人事ではありません。ハンミ薬品は独自の人工知能プラットフォームであるHARPで肥満治療候補物質HM17321を設計しました。食欲を減らす既存系の医薬品とは異なる経路で作用して、脂肪は減らしながら筋肉は保つ効果を狙い、米国食品医薬品局から臨床1相進入の承認を取得しました。シンテカバイオは既に広く使用されている抗体医薬品であるキイトルーダよりも優れ

た性質を持つように再設計したバイオベター候補物質を独自プラットフォームで複数確保し、スタンダ임はハンミ薬品と手を組んで候補物質の発掘に乗り出しています。

Google DeepMindのスピンオフであるアイソモーフィック・ラブス（Isomorphic Labs）も2026年2月、AlphaFold 3を足がかりとしたIsoDDEという設計エンジンを披露し、抗体を白紙から描き出す機能を備えていると紹介し、人間を対象にした最初の臨床試験は2026年中に開始される予定だと明かしました。この会社は2025年の1年間だけで6億ドルを超える投資を新たに獲得もしました。このエンジンは新たに作られた結合予測試験台であるRuns N' PosesでAlphaFold 3より2倍を超える正確度を示し、セレブロンという蛋白質で、それまで知られていなかった結合部位を配列情報だけから事前に見つけ出し、その予測が後に続く実際の研究で確認されてもいました。

新薬候補を最初から描き出す競争は、今や数社の先導企業だけの話ではなく、世界各地で同時に起こる流れになりました。

組み立てる手はますます速くなっています。2019年にGENTRLがDDR1阻害剤の候補を21日で探し出したことが驚きをもって受け止められたとすれば、2026年のISM001-055は、その指先から生まれた物質が臨床第3相試験まで進めることを示しました。しかし、ボトルネックは消えたのではなく、別の場所に移っただけです。以前は使える分子を想像すること自体が難しかったとすれば、今は想像した分子を実際に作れるかどうか、そしてその分子が生きている体内で本当に安全で効果があるのかを確認することがボトルネックとして残っています。表現と生成、拡張という3つの足が揃ったことで、新薬の候補を描き出す速度はすでに人の手を追い越しており、残る戦いは、その絵を現実の化学と生物学の前に立たせ、最後まで持ちこたえさせることです。

2022年に上海の実験室で始まった1枚の絵が2026年に臨床第3相試験という敷居まで上がったように、今後数年の間には、最初からコンピューターが描き出した薬が薬局の棚に並ぶ瞬間に直面する可能性が高いです。化学者の直感と指先の感覚に頼っていた新薬の設計は、このようにデータとアルゴリズムが共に作り出す作業へと重心を移しており、この変化は製薬産業全体の研究方式を再び組み立てる

流れと繋がっています。

候補物質を描き出す手は速くなりましたが、その絵を確認する手はまだそれほど速くありません。2026年に出たある産業分析は、新薬開発のボトルネックが候補物質を探す段階から実験室での検証段階へと移ったと明らかにしました。アメリカのEli Lillyはすでに2020年から設計と合成、精製と分析を一つにまとめ、人の手をほとんど経ずに回るロボット実験室をStrateosという会社と共に運営しています。Hanguk Bogeon Saneop Jinheungwonも、ロボットと人工知能が設計と合成と実験を人の手なしに引き継いで回る自律実験室を備えることを、これから解決すべき課題として挙げました。

この流れの後ろには、あまり知られていない事情が一つあります。人工知能が試行錯誤を減らすには、何が通じたのかだけでなく、何がなぜ通じなかったのかも共に学ばなければなりません。失敗した化合物の記録は論文にも公開データベースにもなかなか残りません。企業は失敗した候補を静かに引き出しにしまっしまい、その失敗の理由を世の中に知らせません。成功した事例だけを集めて学んだ人工知能は、失敗を目の前にしてもその気配をあらかじめ読み取れないことが多いです。この資料不足は、自律実験室がいくら増えても簡単には埋まらないと見られます。

生成モデルは、標的によくくっつく分子だけを組み立てるのではありません。すでに登録された特許の請求項を避けて通る分子を組み立てることにどんどん上手になっています。最近では、強化学習の点数計算自体にこの目標を埋め込む方法まで出ました。既存の特許化合物と構造が近い分子には罰点を与え、まだ誰も権利を主張していない空いた場所を探し出した分子には賞を与える方式です。この道具を使うのは、新薬を初めて作った会社だけではありません。複製薬を作る会社も同じ技術で特許を避けて通る似た物質を設計し、オリジナル会社の特許を相手にした争いを早めています。アメリカの裁判所は2022年、人工知能は特許の発明者になれないと判決を下しました。

ところで、分子を描いた側が人工知能で人は合成経路だけを整えた時、その発明者の欄に誰の名前を書くべきかについては、まだはっきりとした答えがありません。特許制度は、新しい発明を世の中に公開した代償として一定期間の独占権を

与える約束ですが、機械が請求項の隙間を探して回る速度が人が法を直す速度よりもはるかに速ければ、この約束の重さはますます軽くなる可能性が高いです。安い複製薬を早める効果があるとしても、発明を保護して研究費を取り戻せるようにしてくれる特許本来の目的とこの流れがぶつかる問題は、これからもずっと騒がしくなるものと見られます。

2.3

インシリコ(In silico) ADMET及び毒性予測システム

1990年代の米国の薬局では、セルデイン(Seldane)という商品名で売られていた抗ヒスタミン剤のテルフェナジンが、季節性アレルギーを患う人々に広く処方されました。くしゃみや鼻水を抑える効果は優れていましたが、この薬を飲んだ一部の患者が、特定の抗生物質や抗真菌剤、さらにはグレープフルーツジュースと一緒に飲んだ時、心臓が不規則に脈打ち、倒れる事故が続きました。原因は、心臓の細胞のカリウムチャネル、よくhERGチャネルと呼ばれる扉が薬で塞がれ、心臓の鼓動の電気信号が乱れる不整脈でした。

開発当時は、この薬が肝臓で分解されず血液中に溜まり、心臓を脅かすとは誰も予想できませんでした。結局、この薬は市場から静かに姿を消し、製薬業界は一つの新薬が人の体内でどのように動き、どこで危険になるのかを、臨床試験の前にあらかじめ見つけ出す方法を長い間探し求めました。2020年代後半の人工知能は、まさにこの問いに対する答えとして、候補物質が人の体に入る前に、コンピュータの中でその運命をあらかじめ見せてくれるインシリコ予測システムへと育ちました。

一つの新薬が実験室の化合物から薬局の棚に並ぶまでには、平均して10年を超える時間と数千億ウォンのお金がかかります。ところが、臨床試験に入った候補物質のうち、最終承認まで生き残る割合は10に1つにもなりません。2010年から2017年までの臨床試験の失敗事例を分析した研究によると、失敗の原因の40から50パーセントは薬効が期待したほど現れなかったためであり、30パーセント近くは体が耐えられない毒性のためであり、10から15パーセントは薬物が体内に吸収されて消える仕組み自体が間違っていて設計されていたためでした。毒性と薬物動態の問題を合わせると、臨床失敗の半分近くが吸収、分布、代謝、排泄、毒性という五つの文字、すなわちADMETの領域で起きていることになります。

米国国立衛生研究所は、新薬の研究開発費用の22パーセントが、まさにこのADMETの問題を確認し解決するために使われていると明らかにしました。製薬業

界では、新薬の開発費用が時間が経つにつれてかえって増えるという逆説的な流れを、半導体の性能が絶えず良くなるというムーアの法則をひっくり返して、イルームの法則(Eroom's Law)と呼んでいます。候補物質の毒性が、動物実験や人を対象とした臨床試験にまで進んでようやく明らかになるなら、それまでにかかった時間とお金はそのまま消えてしまいます。

長い間、製薬会社はこの危険をあらかじめ取り除くために、ネズミやサルのような動物に候補物質を飲ませて反応を見守る方法に頼ってきました。問題は、動物の体と人の体が、遺伝子や生理構造において少なからず違っているという点です。動物実験で安全だと確認された薬物が、いざ人を対象とした臨床試験では予想外の副作用を起こすという事例が繰り返され、学界ではこの隙間を橋渡し研究の隙間と呼び、長い間頭を悩ませてきました。米国食品医薬品局は、2022年の食品医薬品局近代化法2.0を通じて、新薬の承認申請に動物実験を義務的に含めなければならないという84年続いた規定を法的に取り除きました。

2025年4月にはさらに一歩進み、抗体治療剤を皮切りに、今後3年から5年以内に動物実験を例外的な手続きにするというロードマップを示し、その空席を人工知能に基づく毒性予測モデルやオルガノイド、臓器チップのような新しい評価方法で埋めると明らかにしました。この法律を実際の審査規定に完全に反映するように釘を刺す、食品医薬品局近代化法3.0の法案も、2025年12月に上院を、2026年7月に下院を順に通過し、2026年8月現在、上院の再議決を経て大統領の署名に引き継がれる手続きだけが残っています。

臨床試験という高価な関門に入る前に、問題のある候補物質をあらかじめ取り除く、いわゆる早期フィルタリング(Early Filtering)が、実験室のガラス瓶の中の動物がしていた仕事を引き継ぎ、新薬開発の最初の関門になったわけです。

ADMETという五つの文字は、それ自体が人の体内で薬が経験する旅の物語です。最初の文字Aは吸収で、錠剤が胃や腸を通り過ぎて血液の中に染み込む過程を言います。Dは分布で、血液に乗って全身を回り、肝臓や脳、筋肉のような場所に薬物がとどまる過程です。Mは代謝で、肝臓がまるで工場の分解ラインのように薬物を細かく砕いて、別の物質に変える過程を意味します。

Eは排泄で、使い終わって残った薬物のカスが、腎臓を経ておしっこやうんちとして体外に出る過程です。最後のTは毒性で、このすべての旅の途中で、薬物が肝臓や心臓、腎臓のような臓器を傷つけないかを指します。五つの段階のうち、どの一箇所で問題が起きても、どんなに効果の高い候補物質であっても、結局は薬として生まれることはできません。

分子一つの構造だけを見て、体内でどのように動くかをあらかじめ推測する学問を、定量的構造活性相関、略してQSARと呼びます。分子の形と性質の間に、一定の規則があるという発想から出発した理論です。初期のQSARは、リピンスキーの5法則のように、分子量や脂溶性といったいくつかの数字だけを計算して、薬らしさを測る統計の公式レベルにとどまっていた。しかし今日、新薬の候補として検討すべき化合物の化学空間は数百万、数千万個に達するほど広くなり、この広大な空間を古い線形の公式だけでは扱えなくなりました。そこで研究者たちは、分子を原子という点と化学結合という線でできた網の目として認識するグラフニューラルネットワーク、略してGNNと、分子をスマイルス(SMILES)という文字列で書き、その文法を読み取る大規模言語モデルやトランスフォーマー構造を、QSARに結びつけ始めました。

その結果として出たADMET-AIは、数百万個の化合物の吸収と毒性をたった数時間で見渡し、ディープ・ピーケー(Deep-PK)は、グラフに基づく分子署名を利用して、73種類に及ぶ薬物動態や毒性、物理化学的な指標を一度に導き出します。このようなモデルは、ケンブル(ChEMBL)やドラッグバンク(DrugBank)、トックス21(Tox21)のように、誰でもダウンロードできる広大な公共のデータベースを学習の材料として育ちました。

医薬品に混入する可能性のある不純物が、遺伝子を損傷してがんを引き起こす危険性はないかを判断する国際ガイドラインICH M7は、QSARの予測値を実際の規制審査資料として認めた代表的な事例です。このガイドラインでは、原理の異なる2つの予測モデル、すなわち専門家の化学的知識を規則に移したモデルと統計的に学習したモデルを同時に動かし、異なる結論が出た場合は人間が直接判断するように規定しています。2025年には、見慣れない化学構造に対しても従来の機械学習手法より感度が30パーセント近く高くなったAmesNetが登場し、長年の障害

であった新規骨格化合物の変異原性予測を一段階引き上げました。

このようなモデルが学習する基盤には、米国環境保護庁が運営するToxCastとTox21データベースがありますが、2025年8月に公開された最新版は、1万個に近い化学物質に対する試験管内反応の情報を収めており、毒性予測人工知能を訓練する代表的な公共資産として数えられています。製薬会社は、このような公開データを自社の化合物ライブラリと一緒に学習に使用し、変異原性や発がん性が疑われる不純物を、合成前の設計段階から取り除こうとしています。

候補物質の性質を予測するだけでなく、最初から良い性質を備えた分子をコンピュータの中で新しく描き出す試みも活発です。Retro Drug Designフレームワークは、薬効や生体利用率、代謝安定性、低い毒性のように、互いに衝突しやすい6つのADMET条件を同時に満たす方向へ自ら分子構造を整えていき、実際にこのように設計したキナーゼ阻害剤の候補が実験による検証でも高い成功率を示しました。生成型人工知能のREINVENTとADMET-AIを一つにまとめたADMETriXは、分子を新しく描くとすぐに複数のADMET指標と一緒に計算し、設計と評価をリアルタイムで行き来できるようにしてくれます。

DrugEx v2は、複数の目標が互いにぶつかる時、どれか一つを諦めることなくバランスの取れた点を探すパレート方式の強化学習を薬物設計に取り入れ、標的タンパク質にはよくくっつきながら、見当違いの標的は避ける分子を選び出します。料理人がレシピを考える時に味や栄養、調理時間を一度に考慮するように、これらのモデルは効能と安全性という複数の料理を同時に並べて出すようなものです。

分子一つ一つの性質を調べることに、その分子が実際の人の体内で時間とともにどのように上下するのかを描き出すことは、別の次元の問題です。この隙間を埋める技術が、生理学的薬物動態モデリング、略してPBPKモデリングです。人の体を心臓や肝臓、腎臓、筋肉などの複数の都市からなる地図に見立て、血流をその都市を結ぶ道路として、薬物という自動車が時間の経過とともにどの都市にどれくらいとどまるかを計算する方式だと考えれば理解しやすいです。人工知能は、分子のSMILES構造や指紋情報を分析し、分配係数や水溶性、肝ミクロソームでの分解速度、血漿タンパク質と結合する割合などの重要な数値を瞬時に抽出し、この数値をPBPKシミュレーションエンジンに入れると、静脈注射や経口服用の後に

血中濃度が時間とともにどのように上下するかを示す曲線が描かれます。

化学言語モデルやグラフ畳み込みニューラルネットワークをPBPKと結合した検証研究では、最高血中濃度やその到達時間、曲線下面積、半減期といった指標の予測値が実際の値と決定係数0.61から0.93に至る相関関係を示し、誤差の大きさを表す絶対平均倍数誤差も1.84から2.8の間に留まり、従来の動物実験の外挿法に劣らない精度を示しました。米国では、2020年から2024年の間に承認された新薬245件のうち65件、割合としては26.5パーセントがPBPKモデルを審査の重要な根拠として提出しましたが、その中で薬物間相互作用の評価に使われたケースが81.9パーセントと目立ち、抗がん剤の分野では42パーセントと特に活発に使われました。

このような精度は、新薬を人に初めて投与する臨床試験において安全な開始用量を決めるにとどまらず、肝機能が低下した高齢者や体重、臓器の大きさが大人と異なる小児のように、臨床試験に直接参加することが難しい集団の薬物曝露量まであらかじめ推測する、患者に合わせた精密医療の基礎として使われています。米国食品医薬品局は、このような流れに合わせて、薬物相互作用の評価におけるPBPKモデルの役割を公式に認めるガイドラインを別途設けています。

心臓毒性の予測は、新薬開発の現場で急いで扱うべき課題として挙げられます。テルフェナジンの事例で見たように、薬物が心臓細胞のカリウムチャネルであるhERGチャネルを塞いでしまうと、心拍の電気信号が乱れ、命を脅かす不整脈につながる可能性があります。胃腸の運動を助けていたシサプリドも似たような理由で市場から撤退しましたが、hERG関連の不整脈が抗ヒスタミン剤の一つの系統にとどまらず、複数の治療分野にわたって繰り返し現れる危険であることを示した事例でした。Mamoshinaの研究チームが作った人工知能ベースの定量的連続採点システムは、薬物銀行や医学辞典データベースに収められた膨大な化合物情報を学習し、検証データにおいて曲線下面積79パーセントのレベルで潜在的な心臓毒性をあらかじめ見分けました。

2025年には、複数の特徴を一度に結合するグラフニューラルネットワークモデルであるhERG-MFFGNNが、判断の根拠まで説明してくれる形で登場し、信頼できる心臓毒性データを標準化して集めたDICTrankというデータセットも一緒に発表

されました。同年、バランスの取れた学習データと様々な変数を組み合わせたXGBoost合意モデルは感度0.83、特異度0.90という成績表を受け取り、大規模言語モデルの埋め込み技術を取り入れたhERG-LTNまで登場するなど、予測方式は年々多様化しています。

関連研究25件を集めて検討した論文では、外部検証でも曲線下面積が0.70を超えたり、精度が0.75を上回るモデルが複数確認され、実際の現場で使えるレベルに大きく近づいたという評価を受けました。

肝臓への毒性、特に薬物誘発性肝損傷、つまりDILIは発病の原因が一つに限定されないため、予測が特に困難です。人によって肝臓の酵素の量と種類が異なり、同じ薬物でも、ある人には何も影響せず、別の人には肝臓を損傷させるというように反応が分かれるからです。シトクロム系酵素遺伝子における一般的な多形性、例えばCYP2D6やCYP2C19の活性が人によって異なって生まれつき決まる違いは、同じ用量の薬であっても、ある人には過剰に、別の人には不足するようにしてしまいます。

最近では、こうした遺伝型情報を化合物構造情報とともに学習し、特定の遺伝型集団でのみ顕著に現れる肝毒性リスクまで識別しようとする試みが続いています。最近の人工知能システムは、薬物の物理化学的特性とともに遺伝子、タンパク質、代謝産物を含む多重オミクスデータ、そして標的を外れた相互作用ネットワークまで一度に観察することで、薬物誘発性肝損傷を最大89パーセントの精度で識別する成果を上げました。

2026年初頭には、ChemBERTaやRoBERTa、SMILES2Vecのような異なる言語モデルが抽出した情報をまとめるマルチモーダルディープラーニング方法が学術誌に発表され、分子構造だけを調べる方法よりもはるかに詳細な予測が可能になりました。象徴的な出来事もありました。2025年、アメリカ食品医薬品局はAbsentia Labsが提出した人工知能ベースのデジタル肝臓モデルを、革新的な新薬開発アプローチを審査するISTANDプログラムの認証対象として受け入れましたが、これは人工知能が作成した予測ツールがこの認証手続きに上った最初の事例でした。

まだ三段階の認証手続きの最初の関門を通過したばかりですが、コンピュータ内の仮想肝臓が実際の人間の肝臓に代わって新薬審査書類の一部を埋め始めたという点で、意義深い出来事として記憶されました。

薬物は、そもそも標的とした標的タンパク質一つにだけおとなしく結合しません。一つの鍵が複数のロックに中途半端に合うように、新薬候補物質が本来の目標とした他の受容体や酵素にも引っかかりながら予期しない副作用を起こすことを、標的外相互作用と呼びます。Liu研究チームは、一つのモデルが複数の予測課題を同時に処理するマルチタスクグラフニューラルネットワークを使って、化合物の標的外相互作用を精密に捉える方法を提案し、この方法で薬物誘発性腎損傷のような深刻なリスクを候補物質の最適化段階で事前に選別できることを示しました。

構造と毒性の関係を追求するProTox系のウェブサーバーも急速に規模を拡大しました。最新版のProTox 3.0は、半数致死量と発がん性、突然変異誘発性、免疫毒性を含む61種類に及ぶ毒性指標を分子類似性と機械学習を組み合わせで計算します。膨大な費用がかかる動物毒性試験に入る前に、こうした無料のウェブツール一つだけでも深刻な欠陥を持つ化合物を候補から選別できるようになったわけです。

高齢人口が増え、慢性疾患を患う患者が増えるにつれて、一人が複数の薬を一度に服用することが一般的になり、それだけ薬と薬が体内で相互作用する、つまりDDIを事前に知ることが患者の生命に直結する問題として浮上しました。テルフェナジンが抗生物質またはグレープフルーツジュースと出会ったときのように、一つの薬物が別の薬物を分解する肝酵素を阻害または促進すると、二番目の薬物の濃度が予想より大きく上昇または低下することが起こります。過去には、こうした相互作用を既存の文献を一つ一つ検索して確認する段階にとどまっていたのですが、今は膨大な医学論文と電子健康記録、化合物の3次元構造を一度に読み取る人工知能がこれに代わっています。

BERT-DDIのような大型言語モデルは、数百万件に及ぶ臨床記録と有害事象報告書の中に埋もれていた薬物間のリスク信号を文単位で見つけ出します。薬物と標的タンパク質の関係をグラフで描き、関係ネットワーク上で新しい相互作用を推論していたDTiGEMS+の発想はSumGNNのような後続の知識グラフニューラルネッ

トワークへと繋がり、まだ臨床で報告されたことのない新しい薬物組み合わせのリスクまで事前に知ることができるレベルに至りました。こうして得られた予測結果は、医者が処方箋を書く瞬間に画面に警告を表示する臨床決定支援システムにもすでに接続されており、コンピュータの中で終わる計算ではなく、診療室で実際に機能する安全装置になっていっています。

薬が市場に出た後も、電子健康記録や保険請求資料のような実際の臨床データを人工知能が継続的に眺めながら、臨床試験段階では明らかにならなかった薬物組み合わせのリスク信号を遅れてでも見つけ出す市販後薬物監視作業も同時に行われています。

分子のSMILES文字列や3次元構造、そして融点や脂溶性のような物理化学的性質を入力すると、反対側ではADMET特性予測値とともに、この予測をどこまで信頼してよいかを示す適用範囲、つまりApplicability Domain(適用範囲)と一緒に出力されます。適用範囲は、今入力した分子がモデルが学習した化合物とどの程度似ているかを示す尺度であり、経済協力開発機構もQSARモデルを検証する際にこの範囲を必ず明確にするよう勧告しています。地図にない見知らぬ道に入ると、ナビゲーションのガイダンスが急に狂い始めるように、学習データから遠く離れた新しい化学構造に出会うと、いかに精巧なモデルであっても予測が大きく揺らぐことがあるということです。

現在広く使われているハイブリッドプラットフォームの一つであるADMETlab 3.0は、ChEMBLやPubChem、OCHEMなどの公開データベースから40万件を超える化合物情報を集め、スマイルスを標準入力値として119種類の性質を一度に計算します。複数の予測課題を同時に処理するマルチタスク指向性メッセージパッシングニューラルネットワークに、伝統的な分子記述子を並行して組み合わせた構造のおかげで、速度と正確さという二兎を同時に得たと評価されています。ただし、2026年に発表されたある批判的検討論文は、代表的なベンチマークリポジトリであるセラピューティクス・データ・コモンズ(Therapeutics Data Commons, TDC)のADMETリーダーボードに載った22の標準評価項目すべてにおいて、再現に成功した上位モデルは数えるほどしかなく、いくつかのモデルはコードが公開されていなかったり、学習データとテストデータが混ざっていたりする欠陥まで

あったと指摘し、華やかな性能数値をそのまま信じる前にもう一度吟味すべき理由を残しました。

予測と検証をやり取りするループも次第に緻密になっています。人工知能が有望だと選び出した候補物質をロボット実験室がすぐに合成し、試験管試験まで自動で進めた後、その結果を再び予測モデルにフィードバックする、いわゆる設計、製作、試験、分析の循環構造が多くの研究機関に広がっています。人が一つ一つ化合物を注文して実験の予定を立てていた昔の方式とは異なり、この循環は1日にも何度も回りながら予測モデル自体を継続的に改善していきます。

ADMET予測が精巧になるほど、この循環の速度も一緒に速くなり、このように蓄積された新しい実験データは、逆に次世代の予測モデルをさらに緻密に鍛え上げます。

これほど速く精巧に見える人工知能の予測も、臨床や規制の現場に入るためには、越えなければならないハードルがまだ高いです。内部試験では決定係数が0.8前後だった優秀なモデルが、既存のデータベースになかった新しい骨格の化合物に出会うと、外部検証の性能が0.3台まで急落することは珍しくありません。これまで蓄積された毒性と薬物動態データの大部分が、すでに合成しやすく商用化に成功した特定の薬物群や、欧米圏の人口集団に偏っているという点も障害です。このデータで学習したモデルは、少数人種特有の代謝特性や、細胞・遺伝子治療薬のような新しい治療方式を評価する際に見当違いの答えを出す危険を抱えています。特定の人口集団でのみ目立つ代謝酵素の変異が学習データに十分に反映されていなければ、その集団でのみ特に大きく現れる副作用をモデルが見逃すこともあります。

数千万個のパラメータで絡み合ったディープラーニングモデルが、なぜ特定の化合物を肝毒性物質として指摘したのか、どのような構造的特徴が不整脈を引き起こすのかを、人が理解できる言葉で説明できないブラックボックス状態にとどまっている場合も多いです。患者の命を扱う審査の過程で、このような不透明さはそのまま受け入れられにくいため、人工知能が出した予測値は、オルガノイドや臓器チップのような物理的な生物学的システムを経て、もう一度確認されて初めてその役割を果たしたと認められます。

米国食品医薬品局は2025年1月、医薬品と生物学的製剤の規制判断に人工知能をどのように使うべきかを扱うガイドラインの草案を出し、モデルが使われる状況に応じて危険度を分け、信頼性を検証する7段階の手順を提示しました。このガイドラインは、非臨床と臨床、市販後監視、製造段階まで人工知能が介入する段階を幅広く網羅していますが、安全性や品質に影響を与えない業務効率化までは規制しないと一線を画しました。欧州医薬品庁は2024年9月、医薬品のライフサイクル全般にわたる人工知能の適用に関する考察文書を出し、既存の優秀臨床試験管理基準や品質管理の原則の中に人工知能を取り込みながら、追跡可能性と透明性を強く勧告しました。

大韓民国食品医薬品安全処も2026年から、許可資料の要約や専門翻訳、検討書の草案作成に人工知能を段階的に導入し始め、これを足がかりに新薬の許可審査を240日以内に終える目標を立てる一方で、今後5年間で人工知能新薬・バイオ分野に集中した規制科学の専門人材1,100人を育成すると明らかにしました。各国の規制機関は、それぞれ異なる人工知能の審査基準をそのままにしておけば、複数の国で一緒に進行する臨床試験の資料が国ごとに異なる評価を受ける可能性があるという問題意識のもと、国際調和会議(ICH)などの協議体を通じて共通の基準を合わせていく議論も同時に始めました。

コンピュータが出した予測値の一つが実際の処方箋につながるまで、今では技術に劣らず、このような規定や手続きが速度を左右しています。

どんなに優れたモデルでも、学習するデータが不足していれば力を発揮できません。しかし、病院や製薬会社は自分たちが苦労して集めた臨床記録や化合物データを重要な営業秘密と考え、なかなか外に出そうとせず、ここに米国の健康保険の携行性と責任に関する法律や欧州の一般個人情報保護法のように、患者情報を守る規制まで重なり、データが多く組織の中に閉じ込められたまま散らばっている、いわゆるデータサイロ現象が固まりました。

元のデータは病院や会社のセキュリティサーバーの中にそのまま置き、そのデータを学習した人工知能モデルの重みだけを暗号化して中央サーバーに送り、統合する方式です。大韓民国では2024年4月から2028年12月まで、科学技術情報通信部と保健福祉部が共同で348億ウォンを投じてK-MELLODDY事業を推進しており、

20以上の国内機関が持つ実験室データと臨床データを連合学習で結びつけ、ADMETと薬物動態を一緒に予測するモデルを鍛え上げる作業が2025年から本格的な新規課題として拡大されました。

ヨーロッパで先に複数の製薬会社が一緒に進めていたMELLODDYプロジェクトがその基礎となり、二つのプロジェクト両方とも患者情報流出や営業秘密侵害なしでも予測モデルを一緒に育てることができるという事実を証明してみました。国内製薬会社のうち一つはInsilico Medicineと協力してから4ヶ月で吸収と分布、代謝、排泄、毒性という5つの要件を均等に備えた新薬候補物質を見つけたと明かしたのですが、以前のような数年かかっていた仕事が数ヶ月に減ったわけです。

この流れはADMET予測という狭い領域に留まらず、遺伝体とタンパク質体情報を一緒に読む多重オミックス分析、オルガノイドと臓器チップのような物理的検証装置、そして実際の診療現場の電子記録までを一つに織り交ぜる更に大きな融合生態系に広がっていています。新薬候補物質1つについてコンピュータ内の予測と実験室の検証、そして病院の処方記録が相互にデータをやりとりする構造が整えられると、ADMET予測の精度はどれか1つの技術の発展ではなく、この生態系全体と一緒に厚くなるスピードに従って上がっていきます。

国内でも政府予算で動くK-MELLODDYと民間企業の協力事例が並んで積み重ねられながら、このような融合は少数の大型製薬社だけの武器ではなく、複数の機関と一緒にアクセスできる公用基盤施設に変わっていています。ADMET予測技術の次の勝負処は、どの会社がより精密なアルゴリズム1つを持っているかではなく、どれほど多様なデータと検証手段を密集させて織り出すかに掛かっているわけです。

今この瞬間にも世界中のサーバーでは、候補物質数百万個のADMET性質が夜通し計算され、その中で危険信号が明らかな化合物は、動物1匹、患者1人に会う前にも静かに候補目録から消されています。テルフェナジンが薬局の棚に置かれていた時代には、想像しがたかった風景です。もちろん今の人工知能がすべての危険を完璧に濾し出すと保証することはできません。学習できなかった新しい構造の前では依然として揺らぎ、なぜそのような判断を下したのか説明できないことも多く、結局はオルガノイドと臓器チップのような物理的検証を経て初めて信頼で

きる値として認められます。

それでも臨床試験の敷居を越える前の段階で、危険な候補物質を濾し出す早期フィルタリングの網目は、年が経つにつれて密集してきており、その網を織る糸1本1本がQSARとPBPK、グラフニューラルネットワークと大型言語モデル、そして国境を行き来する連合学習につながっています。新薬候補物質が人体に触れる前にコンピュータシミュレーションという関門を先に通過しなければならない手順は、いくつかの実験室の試験事業を越えて新薬開発の標準段階として固まっていています。

2025年9月、遺伝子編集治療薬企業Intellia Therapeuticsが開発していたnex-zを投与された患者1人で、4等級、つまり生命を脅かすレベルの肝数値上昇が現れました。患者は病院で治療を受けました。この薬は心臓と神経にアミロイドタンパク質が蓄積する遺伝性疾患トランスチレチンアミロイドーシスを目指して、体の中に直接入って遺伝子を書き直す3相段階治療薬MAGNITUDEとMAGNITUDE-2臨床試験の主人公でした。米国食品医薬品局は10月29日、2つの臨床試験全体に臨床保留命令を下し、11月には80代男性患者1人が薬物投与後に亡くなることまで重くなりました。

調査の結果、この死は肝損傷ではなく感染のためであることが明かされましたが、後期臨床試験全体が一瞬に止まったという事実だけはそのまま残されました。この治療薬は化合物の化学構造ではなく遺伝子ハサミそのものが体の中で作動する方式でした。SMILES文字列と分子指紋を学習した既存予測モデルの目には、そもそも引っかけられない対象だったわけです。2つの臨床試験は肝数値をより頻繁に検査し、危険群患者を事前に濾し出す安全装置を加えた後、2026年に入ってようやく1つずつ再び歩みを進めました。

細胞と遺伝子治療薬のような完全に新しい治療方式が増えるほど、分子構造を基盤とする予測モデルの適用範囲は逆に狭くなるという逆説が起きているように見えます。化合物の毒性を読み出す人工知能は数十年分の小分子データを踏んで立っています。しかし遺伝子編集や細胞治療が体の中で起こす免疫反応と標的外編集まで含むデータは、まだ世界のどこにも十分に蓄積されていません。当分の間は、いくら精密な予測モデルを前に出したとしても、患者を直接見守りながら異常信

号を即座に捉える臨床試験現場の監視体系に代わることができない可能性が大きいです。

韓国でも予測の幅を広げようとする動きが新たに始まりました。食品医薬品安全処は2025年12月3日にソウル大学校データサイエンス大学院と協約を結び、国内医療機関の診療データと多重オミックス情報を織り交ぜて、人によって異なって現れる薬物反応を地図として描く人工知能プラットフォームと一緒に構築することにしました。新薬候補物質を実験室段階で事前に濾し出すK-MELLODDYと異なり、今回の協約は薬が患者の体に入った後に起きることまでデータに移し入れようとする試みです。

2つの機関は共同研究に留まらず、データ品質を管理し分析人材を育てるワークショップまで一緒に開くことにしました。2つの事業が並んで積み重ねられながら、予測モデルが扱う時間の幅も一緒に広がっていています。

第3部

次世代モダリティとスマート薬物送達システム

3.1

バイオ医薬品および多重標的治療薬の設計 (ADC)

2026年6月22日に開催され、4日間にわたってアメリカのカリフォルニア州サンディエゴコンベンションセンターは、世界中の医薬品とバイオ企業関係者で足の踏み場もないほど混雑していました。世界最大のバイオ産業見本市であるバイオ・インターナショナル・コンベンション、通称バイオUSAと呼ばれるこのイベントで、韓国企業セルトリオンはブースを出展し、4日間で180件を超えるビジネスミーティングをこなしました。ブースを訪れた来場者は2,000人を超えました。セルトリオンが前面に出したのは、これまで主力としてきたバイオシミラー、つまり特許が切れたバイオ医薬品の後発医薬品ではなく、人工知能ベースの新薬開発と抗体-薬物接合体、そして複数の標的を同時に狙う多重抗体技術でした。

低分子化合物1錠をコンピュータで設計していた時代から、巨大な蛋白質分子を人工知能で丸ごと新しく描き出す時代へと医薬品産業の重心が移行しているという証拠でした。

ほんの数年前までは、抗体や細胞治療薬のように体積の大きい医薬品を設計することは、研究者の直感と反復実験に頼る苦しい作業でした。候補抗体1つを作り、実験皿に入れて反応を見守った後、失敗すれば最初からやり直すという試行錯誤を何百回、何千回も繰り返さなければなりませんでした。その過程で10年以上の年月と兆に近い費用がかかることも珍しくありませんでした。ところが蛋白質のアミノ酸配列と立体構造をコンピュータ画面の中で原子単位まであらかじめ描いて試験してみる技術が発展すると、話が変わりました。

巨大言語モデルと蛋白質言語モデル、複数層の生体データを一度に読み出す多重オミックス分析がその中心にあります。実験皿の上で数年を費やさなければならなかった設計作業が、今ではサーバーの中で数日で終わる場合が増えています。

抗体は私たちの体の免疫系が作り出すY字型の蛋白質です。細菌やウイルスのように体に侵入した異物を認識し、そこにぴったり付着して印を残す役割を果たしま

す。錠前に合う鍵のように、抗体1つ1つは決められた標的とのみ結合するようにできています。私たちの体は数百万種類を超える異なる抗体を作り出すことができ、その中から特定の標的1つにだけ付着する抗体を選び出して大量に複製して薬として使うのが単一クローン抗体（mAb）です。リウマチ関節炎治療薬から各種がん治療薬まで、今、病院で使用されている抗体医薬品の大部分はこの方法で作られました。

抗体1つが標的1つにだけ付着することから一歩進んで、異なる2つ以上の標的に同時に付着するように設計した抗体も次々と登場しています。このような抗体を多重抗体または二重特異性抗体と呼びます。一方の腕で癌細胞表面の特定のタンパク質をつかんで、もう一方の腕で体内を巡回するT細胞をつかみ、ちょうど2人の手を両手でそれぞれ握って互いに向き合わせるような仲人のように、T細胞を癌細胞のすぐ隣まで引き寄せるやり方です。

セルトリオンがバイオUSA 2026で新たに披露した技術も、このような多重抗体を設計し開発可能性をあらかじめ評価する人工知能ツールでした。標的を1つではなく2つ、3つに増やすほど、抗体が認識しなければならない組み合わせの場合の数が幾何級数的に増えるため、人間が1つ1つ手で候補を絞り出していた従来の方式では対応しきれなかった設計空間を人工知能が代わりに探索してくれるのです。

抗体-薬物接合体（ADC）は、この抗体に強力な毒性物質を取り付けた精密誘導医薬品です。宅配便の箱を思い浮かべると理解しやすいです。箱の中には癌細胞を殺すほど毒性の強い化学物質、つまりペイロードが入っています。この箱を無差別に投げ出したら大変なことになるため、正確な住所を知っている配達人が必要です。その配達人の役割を果たすのが正に抗体です。

癌細胞表面にだけある特定のタンパク質を認識し、その家にだけ箱を配達します。そして箱をしっかりと結んで配達中は決して開かないようにして、正確な家に到着したときだけ解けるようにする紐がリンカーです。抗体とリンカーとペイロードという3つの部品が息を合わせることで、初めて癌細胞の中だけで毒性物質が爆発する精密爆弾が完成するのです。

既存の抗がん化学療法はこれとは異なるやり方で機能してきました。毒性の強い薬物を全身の血管に流して、速く増殖する細胞を手当たり次第に攻撃するやり方でした。その過程で髪が抜け、腸粘膜が傷つき、白血球数が低下する副作用を避けられませんでした。いわばカーペット爆撃に近かったわけです。抗体-薬物接合体は正反対に座標を正確に打って、そこにだけ落ちる誘導ミサイルに近いです。

同じ毒性物質を使用しても、正常細胞には触れず癌細胞にだけ集中的に作用するため、はるかに少ない用量ではるかに強い効果を発揮できます。この精密さを実際に実現することが決して容易ではないという事実が、長らく抗体-薬物接合体開発の足かせとなってきました。

このような人工知能設計フローが本格化する前から、抗体-薬物接合体はすでに臨床で成果を証明してきたモダリティーでした。乳がんと胃がん治療に使用されるエンハーツ、乳がんと膀胱がん治療に使用されるトロデルビのような薬物は、既存の抗がん剤ではもう手の打ちようがなかった患者たちに新しい選択肢を開き、この分野全体の信頼を築き上げてきました。人工知能は、このようにすでに検証されたモダリティーを次世代へとより速くより精密に押し上げる役割を担っています。

正にこの点で人工知能が抗体-薬物接合体設計の様相を変えています。研究者たちは抗体の結合部位とリンカーの化学構造、ペイロードの毒性の強さを1つ1つ別々に調整する代わりに、3つの要素を同時に動かす強化学習方式を使い始めました。人工知能モデルが候補の組み合わせを1つ提示すると、コンピュータの中でその組み合わせが血液の中でどの程度安定して持ちこたえるのか、癌細胞に到着したときどの程度正確に切れるのか、正常組織にはどの程度漏れ出すのかを仮想的に計算します。

その結果に従って点数をつけ、点数が低ければ構造を変えて再度計算する過程を数万回繰り返します。このように、抗体の標的結合力を引き上げ、全身毒性を下げるという二つの目標を同時に追い求める循環の輪が作られます。抗体とリンカーとペイロードが互いの弱点を補い合いながら、共に進化する構造というわけです。

抗体とリンカーとペイロードが作り出すこの循環の輪は、事実上ひとつの多重標的最適化ループと呼ぶにふさわしいものです。この循環構造をもう少し具体的に覗いてみると、人工知能に与えられる報酬信号は大きく二つに分けられます。一つは、抗体ががん細胞表面の抗原にどれほど強く、どれほど長くくっついているかを表す結合指標であり、もう一つは、ペイロードが正常細胞ではどれほど静かにしており、がん細胞の中でどれほど激しく働くかを示す選択的毒性指標です。

二つの指標を同時に高める組み合わせを見つけ出す仕事は、人が手で計算するには変数が多すぎる多次元最適化問題です。人工知能は、数多くの抗体とリンカー、ペイロードの組み合わせを同時に試しながら、二つの指標を共に満たす組み合わせを探し出し、その結果として出た候補だけを実際の実験室での検証段階へ回します。化学者と免疫学者が各自別々に働いていた昔の方式ではなかなか出にくかった組み合わせが、このような過程からいくつも発見されています。

この仮想計算には、定量的システム薬理学という技法が共に使われます。薬物が体の中に入った後、吸収され、広がり、分解される全過程を、数学モデルで予め描いてみる方法です。ペイロードが細胞を殺す性能とリンカーの安定性をこのモデルで予めふるいにかければ、実際の動物実験や臨床試験に入る前に、失敗する組み合わせをかなりの数、選び出すことができます。このような設計能力を持った会社を買い取ろうとする動きが、2025年と2026年の製薬業界を熱くしました。Gilead Sciencesは、2026年4月、ドイツの抗体薬物複合体専門企業Tubulisを、契約金31億5000万ドルに最大18億5000万ドルの段階別マイルストーンを加えて、総額50億ドル規模で買収することにしました。

Eli Lillyも、抗体薬物複合体の新興企業Crossbridge Bioを買い取り、同じ流れに合流しました。二つの取引はどちらも、会社が保有する特定の抗がん剤候補一つではなく、リンカーとペイロード、接合化学を設計する人工知能プラットフォーム自体を狙った買収でした。

この市場における中国のバイオ企業たちの存在感は圧倒的です。全世界の抗体薬物複合体の技術輸出契約のうち、中国から出た候補物質が90件中9件の割合を占めているという評価が出るほどです。韓国企業の中では、LigaChem Bio、旧名で

はLegoChem Biosciencesが先頭を走っています。この会社は、独自開発したリンカーとペイロードの源泉技術を前面に押し出し、Amgenに約1兆6000億ウォン、Janssenに約2兆2000億ウォン、日本のOno Pharmaceuticalに9400億ウォンを超える規模で技術を輸出し、公開された契約だけを合わせても9兆3000億ウォンを超えます。

非公開の契約まで加えれば、累積の技術輸出規模が10兆ウォンを超えるという集計もあります。LigaChem Bioが抗体開発会社のNextCureと共に開発中であるB7-H4標的の抗体薬物複合体LNCB74は、グローバル臨床第1相で高用量投与群を追加すると米国食品医薬品局に申請し、次の段階へと進んでいます。市場調査機関のEvaluatePharmaは、抗体薬物複合体市場の全体規模が2025年の150億ドル前後から2032年には570億ドルまで大きくなると見込んでいます。

抗体自体を設計する段階においても、人工知能はすでに実戦に入っています。ディープラーニング基盤のタンパク質言語モデルは、膨大な抗体配列データを学習した後、特定の標的に強くくっつく新しい抗体配列を最初から新しく作り出します。米国の人工知能新薬開発企業Absciは、生成型人工知能プラットフォームで設計した抗体ABS-101を炎症性腸疾患の治療薬候補として臨床第1相に進入させました。

TL1Aという標的タンパク質にくっつくように設計されたこの抗体は、人工知能が丸ごと描き出した後、実際の人間に初めて投与された生物医薬品であるという点で業界の注目を集めました。4ヶ月に一度だけ皮下に注射すればよいように設計されており、免疫原性の危険を下げることに焦点を合わせました。2025年の下半期には、初期の安全性と薬物動態学、免疫原性の資料が出る予定でした。

抗体設計を専門とする人工知能スタートアップにも莫大な投資金が集まっています。OpenAIの投資を受けた米国スタートアップChai Discoveryは、2025年12月に1億3000万ドル規模のSeries B投資を誘致し、企業価値を13億ドルと認められ、その後、企業価値を34億ドル水準に上げて4億ドルを追加で調達する交渉を進めていると知られています。この会社が出したChai-3モデルはPfizerと技術契約を結び、2026年1月にはEli Lillyともパートナーシップを発表しました。

業界全体で見ると、2025年6月から2026年5月までの一年間で、抗体をはじめとする生物学的製剤を人工知能で設計する企業にだけ、9億8000万ドルを超える投資金が20件以上の取引を通じて流れ込みました。2026年の第1四半期にも、Arendilが7億8700万ドル、EmendoBioが8000万ドルを新たに誘致するなど、投資の熱気は冷めていません。

抗体の立体構造を描き出す人工知能モデルも、設計速度を目に見えて引き上げています。タンパク質のアミノ酸配列さえ入れれば立体構造を予測してくれた既存の技術から一歩さらに進み、望む標的の形にぴったり合う結合部位の立体構造を逆算して描き出し、そこに合うアミノ酸配列を埋め込む逆方向の設計が可能になりました。すでに体にある服の中から一つを選んで着るのではなく、お客様の寸法を先に測り、それに合わせて生地を裁断するオーダーメイドの洋服店と似た方式です。

このような構造基盤の設計技法は、抗体だけでなく、ナノボディ、ペプチド、PROTACの二重フック構造を作るのにも共に使われながら、様々な次世代モダリティを貫く共通の基盤技術になりつつあります。

抗体よりもはるかにサイズが小さいナノボディも、人工知能の恩恵を大きく受けている分野です。ラクダやアルパカのような動物の体からのみ発見されるこのミニ抗体は、サイズが小さいため硬い腫瘍組織の奥深くまで入り込みやすく、微生物を利用して大量に生産するのにも適しています。最近、ある研究チームは人間の代わりに複数の人工知能エージェントからなる仮想研究チームを編成し、新型コロナウイルスを無力化するナノボディの配列をコンピューター内でゼロから設計させることにしました。

人工知能エージェントたちが標的タンパク質の構造を分析し、結合候補を提案して互いの設計を検討する過程を経た後、実際の実験室でこの配列を合成して検証することにまで成功しました。人間の研究者が標的を定めて方向性だけを示せば、残りの設計作業は人工知能チームがほぼ自律的にこなせるという事実を示した事例として挙げられます。

ペプチド、つまり数十個のアミノ酸が短くつながった小さな分子を設計する分野でも似たような傾向が見られます。マスク化言語モデリングという自然言語処理の手法を借りてきたペプエムエルエム(PepMLM)という人工知能ツールは、標的タンパク質の3次元構造情報が全くなくても、アミノ酸の配列情報だけで標的に特異的に結合する線形ペプチドを設計し出します。これまでペプチドの設計は、タンパク質の折りたたまれた構造を精密に知っていなければ始められない作業でしたが、構造情報がなくても結合力の高い候補を選び出せることが確認され、設計可能な標的の範囲そのものが広がっています。

抗体や抗体薬物複合体のようにサイズの大きな分子を作るのが負担になる状況で、ペプチドはより軽く、より安く作れる代案として注目を集めています。

抗体やCAR-T細胞が標的を認識するには、ぴったりくっつくための取っ手のような配列区間が必要ですが、これをエピトープと呼びます。広い塀の中でもドアの取っ手を正確に見つけてこそドアを開けられるのと似ています。問題は、私たちの体の免疫システムが、新しく入ってきた治療用の抗体や細胞そのものを見知らぬ侵入者だと誤認して攻撃してしまうケースが少なくないという点です。

これを免疫原性と呼び、この副作用がひどいと薬効が消えたり危険なアレルギー反応につながったりします。人工知能は、抗体の配列を人間の目では見当もつかないレベルまであらかじめ調べて免疫系を刺激する配列区間を見つけ出し、その区間を人間の体にすでに馴染みのある配列に置き換えるように案内してくれます。

このような免疫原性の予測を実際に実行する人工知能ツールも急速に精巧になっています。抗体の配列全体と個別のアミノ酸部位別の免疫原性スコアを一緒に計算するシータ(SITA)という予測システムは、人間に由来する抗体とそうでない抗体をはっきりと区別し出します。2万6千個余りに及ぶペプチドと主要組織適合遺伝子複合体の結合データを学習したイムノストラクト(ImmunoStruct)というマルチモーダルディープラーニングモデルは、配列情報と構造情報、生化学的な特性を一度に結合して予測の正確度を引き上げました。

ネットエムエイチシーツーパン(NetMHCIIpan)、フーマブ(Hu-mAb)、バイオフィイ(BioPhi)のようなツールは、抗体を人間にとってより親しみやすい配列に整え

るヒト化作業に実際に使われています。新薬の候補が動物実験の段階に入るずっと前に、配列だけで危険信号を取り除けるようになったという点で、開発費用と期間を同時に減らす効果を生み出しています。

がんワクチンを設計する際にも、エピトープの予測は重要な関門です。同じがんでも、患者ごとにがん細胞の表面に現れる突然変異タンパク質、いわゆる新生抗原の形がすべて異なります。研究者たちは、タンパク質言語モデルと畳み込みニューラルネットワークを結合したディープラーニングモデルを使って、患者一人一人の腫瘍組織からどのような新生抗原が免疫細胞によく認識されるかを予測しています。黒色腫の患者の腫瘍検体から抽出した数千個の突然変異タンパク質の中から、実際に免疫反応を引き起こす候補を絞り込むのにこのようなモデルが使われており、その結果をもとに、患者一人のためだけのオーダーメイド型がんワクチンや免疫抗がん剤を設計する研究が複数の大学や企業で進められています。

細胞の表面に抗原を運ぶ主要組織適合遺伝子複合体1型に提示されるCD8陽性T細胞の標的エピトープの予測正確度も、ここ数年の間に目立って向上しました。

このように標的エピトープを見つけ出し、免疫原性をあらかじめ取り除く人工知能のパイプラインは、抗体だけにとどまりません。CAR-T細胞に新しく取り付ける人工受容体の外側の認識部位を設計する時にも同じ方式の予測モデルを回し、患者の免疫系がその受容体自体を見知らぬ物質だと誤認して拒絶反応を起こすリスクをあらかじめ計算します。プロタックのように低分子に近い分子であっても、体内で代謝されて新しく生まれる副産物が思いがけない免疫反応を引き起こす可能性を排除できないため、代謝物予測モデルと毒性予測モデルと一緒に回してこのようなリスクを事前に取り除く手順が、今や開発初期段階の基本的な点検項目になっています。

標的を見つけ出す段階から免疫原性を検証する段階まで、互いに異なるモダリティであっても人工知能が実行する検証手順の骨組みはかなり似ています。

薬の標的となるタンパク質を扱う方法において、まったく異なる発想をとる技術も急速に成長しています。プロタックとは、タンパク質分解誘導キメラという意味の英語の頭文字を取った名前です。既存の薬の大部分は、悪いタンパク質の動

きを抑え込むことにとどまっていた。薬の分子がタンパク質にくっついていて間だけ効果があり、薬が体から抜けたりタンパク質の構造が少しでも変わったりすると、再び活動を始めた。

プロタックはこの方法を完全に捨てました。悪いタンパク質の活動を抑え込む代わりに、細胞内のリサイクル処理班と言えるプロテアソームにそのタンパク質を丸ごと引き渡して分解させるようにします。片方の手ではなくすべきタンパク質を捕まえ、もう片方の手では分解の目印をつける酵素を呼び寄せる二重のフックのような分子です。問題は、この二重フックの構造がサイズも大きくて重いため、細胞膜をうまく通過できず、体内の吸収率も低いという点でした。

最近では表皮成長因子受容体をターゲットにしたEGFR-PROPKのような機械学習ベースの薬動学予測モデルが登場し、PROTAC分子が体内でどの程度よく吸収され拡散するかを事前に計算して分子構造を細かく調整できるようになっています。

PROTACが注目を集める理由は、既存薬が手をつけなかったターゲットまで触ることができるからでもあります。ある種のタンパク質は表面が滑らかなため、薬が強固に付着する場所がないことがあり、このようなタンパク質を新薬開発者たちは長い間触ることができないターゲットと呼んできました。PROTACは活動を阻止するほど強固に掴む必要がなく、ほんのわずかに引っかかるだけで、後は細胞自身の分解装置に任せるため、これまで新薬開発者たちが諦めていたターゲットの相当数が再び候補リストに上がっています。

人工知能は、このようにわずかに引っかかることができる結合部位を見つけ出す作業にも使われています。タンパク質表面の数多くの凹凸をスキャンしてどの部位にどのような形のフックを引っかけるのが良いかを提案し、その部位にどのようなリンカーをどの程度長く接続すれば2つのタンパク質が自然に近づくかを併せて計算してくれます。

このような人工知能設計ツールの力を得て、PROTACは実際の臨床現場でも成果を出しています。ファイザーと共に乳がん治療薬Befedegstantを開発してきたArvinasは、2025年米国臨床腫瘍学会で3相Veritac-2試験結果を発表しましたが、以前に治療を受けた経験のあるホルモン受容体陽性、HER2陰性乳がん患者に対し

て既存治療薬Fulvestrantより無進行生存期間を統計的に明確に延長させました。両社は2025年6月米国食品医薬品局に新薬承認を申請し優先審査対象に指定されました。Chimera Therapeuticsが開発した経口用STAT6分解薬KT-621は、アトピー性皮膚炎臨床1b相試験で血液と皮膚組織でそれぞれ96%、94%に達する標的タンパク質分解率を示しました。

Beigeneの BTK ターゲット PROTAC BGB-16673は、以前に治療を受けたが再発した慢性リンパ球性白血病患者を対象にした初期試験で78%の奏効率を記録した後、2025年に3相に進みました。Bristol-Myers Squibbのアンドロゲン受容体ターゲット PROTAC も、前立腺がん患者を対象に初の3相試験を開始しました。2025年中盤の時点で臨床段階に進入した PROTAC プログラムは全世界で28件に達し、そのうち10件はすでに2相ないし3相にまで進んでいます。

細胞自体を遺伝的に改造して治療薬として使う CAR-T 細胞治療薬はさらに進んだアプローチです。正式な名称はキメラ抗原受容体を付けた T 細胞治療薬という意味です。患者の血液から免疫を担当する T 細胞を取り出した後、遺伝子を操作して癌細胞表面タンパク質だけを認識する人工受容体をつけ加えます。あたかも武装警備員に特定の容貌が記載された手配書を渡して再び巡回に送り出すようなものです。

このように改造された T 細胞を体に再び注入すると、血液がん細胞を正確に見つけ出して攻撃します。白血病やリンパ腫のような血液がんではすでに複数の CAR-T 治療薬が承認を受けて使用されていますが、胃、肺、大腸にできる固形がんは事情が異なります。固形がんは腫瘍周辺に免疫細胞の接近と活動を阻止する防御壁が厚くとり囲まれているため、CAR-T 細胞が到達しても効力を発揮できない場合が多くありました。

人工知能はこの固形がん障壁を超えるための CAR-T 細胞設計にも深く関与しています。CAR-T 細胞が体に入る前に受容体自身が信号を誤って活性化させて疲弊してしまう現象をトニック信号伝達と呼びますが、人工知能シミュレーションモデルはどの受容体構造がこのような早期枯渇を引き起こすかを事前に計算して、より長く耐える構造を選別するのに役立ちます。患者一人ひとりの身体状態を反映したデジタルツインとロボット装置を組み合わせることで CAR-T 細胞を製造する製造プ

ロセス自体を自動化する試みも続いています。

細胞を患者の体に注入する前に、腫瘍細胞と免疫細胞が実際にどのように戦うか、サイトカインという信号物質が一気に放出されて高熱と低血圧を引き起こすサイトカイン放出症候群がどの程度重篤に現れるかをコンピュータ内で事前にシミュレーションし、リスクを低減する方向に細胞設計を調整します。

実際の臨床現場でも固形がんを対象にした CAR-T 細胞治療薬が次々と成果を出しています。Allogene Therapeuticsが開発した off-the-shelf CAR-T 細胞 ALLO-316 は、腎臓がん患者を対象にした初期試験で長く持続する反応を示し、固形がんでも CAR-T 細胞治療薬が奏効する可能性があるという信号として受け入れられています。STAR-101 1相試験では、KIR という新たなターゲットを対象にした CAR-T 細胞が42人の患者を対象に低用量でも安全でありながら効果シグナルを示しましたが、これまで承認された細胞治療薬が存在しなかったがん種で得られた結果であり、その意義は大きかったのです。

中国の UTC Therapeutics は、環状メッセンジャーRNA を利用した CAR-T 細胞で中国臨床試験許可を取得し、2026年米国がん研究学会で固形がん1相結果を発表しました。米国と英国では、横紋筋肉腫とユーイング肉腫のような小児固形がん患者60人を対象に、2つのターゲットを同時に標的にする次世代型 CAR-T 臨床試験も新たに開始しました。

このような華やかな設計成果にも限界は確かに存在します。コンピュータ内で完璧に計算されたリンカーの安定性や結合親和性が、実際に生きている細胞と動物モデルでも同じように再現されるという保証はありません。予測できない標的外毒性や免疫反応が後になって現れて大規模臨床試験が中止される事例がまだ報告されており、人工知能モデルを学習させるデータ自体が特定の人種または特定の臨床環境に偏っているため、別の環境の患者に同じように適用することが困難な場合も存在します。

設計段階でいくら緻密な計算を行ったとしても、結局のところ実際に人体に直接投与して確認する手順をスキップすることはできないという事実が、この分野全体の上に重くのしかかっています。

抗体や抗体薬物複合体、PROTAC、CAR-T細胞までを見てみると、人工知能が担う役割には共通の流れが見えます。標的となるタンパク質やエピトープを見つけ出すこと、その標的にくっつく分子や細胞を最初から新しく設計すること、免疫系を刺激しないかあらかじめ分けること、そして体内で実際にどのように動くのかシミュレーションであらかじめ点検することまで、4つの作業が一つの循環構造の中で噛み合っていて回っています。

1種類の低分子化合物を精巧に整えていた過去の新薬開発とは異なり、今は抗体やリンカー、ペイロード、遺伝子を操作した細胞のように、複数の部品が混ざり合った複合治療薬を、多重標的の観点から一度に設計する作業へと重心が移っています。この作業を実際にこなす人工知能の設計力とデータを併せ持つ企業が、次世代の新薬開発の主導権を握る可能性が高いと思われます。

今、臨床試験の段階に上がった抗体薬物複合体やCAR-T細胞治療薬の候補が一つ出るまでに、コンピューターの中ではすでに数万個の候補の組み合わせが作られては捨てられる過程を経てきました。患者の血管に実際に注射針が入る前に、抗体の結合力やリンカーの安定性、ペイロードの毒性、細胞受容体のシグナル伝達まで、インシリコ環境で先に検証する手続きが、例外ではなく標準として固まりつつあります。低分子の飲み薬一つを最適化していた時代の新薬開発とは、まったく異なる産業が今作られています。

抗体薬物複合体が歩んできた道が順調だったわけではありません。2000年に白血病治療薬として承認されたマイロターゲットは、初の抗体薬物複合体という名声を得ましたが、薬物をぶら下げたリンカーが不安定で、正常な細胞でも毒性物質が漏れ出し、肝静脈閉塞症などの副作用に加えて、後続の臨床試験で生存の利益を再び証明できず、2010年に市場から自主撤退しました。今広く使われているエンハーツでも、患者10人に1人の割合で間質性肺疾患が現れます。毒性物質が標的細胞を越えて周辺組織に漏れ出すバイスタンダー効果が原因として挙げられます。2025年12月には、第一三共とMSDが共同で開発していたイフィナタマブ デルクステカンが、小細胞肺がんの第3相試験で間質性肺疾患による死者が多数出たことで、世界中の臨床試験が一斉に止まりました。

米国食品医薬品局は保留を解除しましたが、両社はリスクの大きい患者を試験対象から外し、臨床試験の担当者を再教育する安全対策を出さなければなりません。欧州の規制当局はまだ扉を開いていません。コンピューターの画面の中でいくら安定的に計算されたリンカーでも、人の肺組織の中でどのように反応するかまでは、人工知能があらかじめすべてふり分けることはできないようです。

いくら薬効が良くても、値段が負担しきれないものなら、結局患者の手には届きません。トロデルビやパドセブのような抗体薬物複合体は、1回打つごとに数百万ウォンがかかり、健康保険の給付が確定するまでその費用はそっくり患者が抱え込みます。Geongang Boheom Simsa Pyeonggawon Yakje Geubyeo Pyeonggawiweonhoeは、トロデルビを審査する際、費用対効果の指標がすでに審査を終えた他の薬よりも悪く出たにもかかわらず、代わりとなる治療法がなく、生存期間が明らかに延びるという点を根拠に給付を認めました。

2026年5月19日、Gukhoe Uiwon Hoegwanで開かれた女性がん政策討論会で、Yubangam Hwanu ChongyeonhaphoeのChoi Seung-ran会長は、遺伝子変異がある転移性乳がんの患者たちが、治療法を目の前にしながらも、費用と制度のために治療の崖っぷちに立たされていると述べました。設計段階で人工知能が開発期間と費用を大きく減らしてくれたとしても、その節減分が薬価の引き下げに直結しない限り、患者が実際に享受する恩恵はゆっくりとしか増えない可能性が高いです。

韓国の政府と企業もこの流れに歩調を合わせています。Sikpum Uiyakpum Anjeoncheoは、2025年11月28日、抗体薬物複合体の品質を審査する際に何を見るべきかをまとめた案内書を初めて出しました。抗体とリンカー、ペイロードそれぞれの品質資料をどのように分けて作成するか、製造方法が変わった際に何を再び検討すべきかが盛り込まれています。

Ligachem BioとCelltrionのほかにも、リンカー技術を前面に押し出したIntoCellは、Samsung Bioepisと共同開発を進める一方で、中国のFrontline Biopharmaと候補物質2種を共同で開発する契約を結びました。Pinotbioは、抗体薬物複合体の候補物質5つのうち3つをCelltrionに引き渡し、残りの2つは自ら開発を続けています。

3.2

天然物新薬データベースのAI解読とハイブリッド発掘

1971年の晩秋、北京郊外のある実験室で、トゥ・ヨウヨウは西暦340年頃に葛洪が書いた医学書『肘後備急方』のページをめくっていました。この実験は個人の研究ではなく、1967年に始まった中国の国家プロジェクト、いわゆる「523任務」に属するものでした。マラリアで倒れる兵士が続出したため、政府が全国の研究者を集めて抗マラリア物質を探すように指示した事業でした。当時、マラリアは毎年数百万人の命を奪う病気であり、従来の治療薬であったクロロキンに耐性を持つ菌株まで広く蔓延しているところでした。

トゥ・ヨウヨウのチームはすでに数百回の実験を繰り返していましたが、はっきりとした成果を得られていない状況でした。クソニンジンをお湯で煎じるのではなく、汁を絞り出すように低い温度で抽出せよという古い一節が、彼女の目にとまりました。熱を加えると薬効成分が壊れてしまうという事実を、彼女はどのように古い文献の中から読み取ったのです。この洞察から生まれた物質がマラリア治療薬のアルテミシニンであり、トゥ・ヨウヨウは2015年にノーベル生理学・医学賞を受賞しました。

古代エジプトやギリシャの人々も、ヤナギの樹皮を噛むと熱や痛みが和らぐという事実を経験として知っていました。19世紀末、ドイツのバイエル社の化学者フェリックス・ホフマンは、その樹皮の中の成分を化学的に調整し、アスピリンを作り出しました。実際にアスピリンが体内のどの酵素を阻害して痛みや熱を和らげるのか、いわゆる作用機序が科学的に解明されたのは、薬が市販されてから70年近く経った1970年代のことでした。1928年、ロンドンのある病院の実験室では、アレクサンダー・フレミングが休暇を終えて戻り、一つの培養皿がアオカビに汚染されているのを見つけました。

カビの周囲だけで細菌が繁殖できなかったその偶然の痕跡から、ペニシリンが生まれました。1960年代には、アメリカ国立癌研究所から依頼を受けた植物学者た

ちがタイヘイヨウイチイの樹皮を剥がして抗がん成分を見つけ出し、この物質はパクリタキセルという名の抗がん剤として世に出ました。現在、薬局で売られている低分子医薬品の多くが、このように自然から生まれた構造にルーツを持っていると知られています。

このリストは四つでは終わりません。ヨーロッパに自生するジギタリスという植物の葉から抽出されたジゴキシンは、古くから心不全患者の心拍を調節する薬として使われてきました。マダガスカルのキョウチクトウ科の植物からは白血病の治療に使われるビンクリスチンが生まれ、ケシの実の乳液からは現在でも強力な鎮痛剤とされるモルヒネが生まれました。コレステロールを下げるスタチン系薬物も、もとはといえばカビが自分を守るために作り出した物質から出発しました。

このいくつかの事例だけでも、自然が人類の薬箱の役割を果たしてきた歳月の厚みを推し量ることができます。アスピリンもペニシリンもパクリタキセルもアルテミシニンも、人が自然を長く観察し、偶然と根気を積み重ねて得た結果でした。

現在、この発見の過程を人工知能がコンピュータの画面の中で繰り返しています。自然界には植物や微生物、海洋生物が作り出した化学物質が少なくとも数十万種存在すると推定されていますが、人がその一つ一つを実験台に乗せて薬効を確認しようとするれば、何世代かかっても終わりません。理論的に組み合わせ可能な薬物候補分子の数は宇宙にある星の数よりも多いという比喻がケモインフォマティクスの分野でよく使われるほど、人が手作業で探すにはその空間は想像しがたいほど広いのです。

これまで人類が化学的に確認して名前を付けた天然物は数十万種にとどまりますが、まだ構造すら明らかになっていない化合物まで加えれば、その数は何倍にも増えると推定されています。この広大な空間に対処するために登場した技術が、グラフニューラルネットワーク（Graph Neural Network, GNN）です。

グラフニューラルネットワークという人工知能技術は、分子一つについて原子を点として、化学結合をその点を結ぶ線として表現した網の構造、すなわちグラフに変換し、この網の模様を丸ごと学習します。鍵のでこぼこした溝が錠前の中の構造とぴったり合って初めて扉が開くように、分子の立体構造が体内のタンパク

質の特定の溝とぴったり合って初めて薬効が現れますが、グラフニューラルネットワークはこの適合度を原子単位で計算し出します。最近の研究では、この方法で作られたバーチャルスクリーニングモデルが、実際に薬効のある化合物を見分ける精度の指標で0.90を超える成績を出したと報告されました。

自己注意機構を結合したDTIGNのようなモデルは、化合物と標的タンパク質の結合力を予測する課題において、前の世代のモデルより一段階進んだ性能を見せたと知られています。ただし、このようなモデルはたいてい、すでに知られている化合物と標的の結合情報だけを持って学習するため、訓練の過程で一度も見たことのない全く新しいタイプの標的に出会ったと、予測力ががくっと落ちる可能性があるという限界も同時に指摘されています。人が実験室で数百種類の候補をふるい分けるのに何年も費やしていた仕事を、このようなモデルはコンピュータの中で数日のうちに終わらせます。

植物化学物質、つまり植物が自らを守るために作り出す化学物質の構造を人工知能が一つずつ読み解くことで、これまで分からなかった薬理機序も明らかになってきています。潰瘍性大腸炎に古くから使われてきた天然化合物ベルベリンが、体内のどのようなシグナル伝達経路を調節するのかは、ネットワーク薬理学と人工知能の分析を通じて、ようやく具体的に描かれようとしています。ケルセチン誘導体であるフィセチンも、定量的構造活性相関モデル、すなわち分子の構造と薬効の間の数学的関係を計算するモデルを通じて、標的結合力と毒性が同時に評価されています。天然物由来の分子の大きさと枝分かれ、そして立体異性構造を扱うように別途訓練された変分オートエンコーダモデルは、複雑な3次元の骨格を維持しながらも薬物としての性質を改善した新しい天然物類似構造を生成し出すまでに進歩しました。

全世界の実験室が共に積み上げてきた質量分析データベースGNPS2には、2025年7月の時点で240万件を超えるスペクトルが集まっていますが、新しく得た試料のスペクトルをこのデータベースと照合すれば、すでに知られている化合物を再び発見するために時間を無駄にせずに済みます。人が数百種類の混合物から有効成分を一つ分離し出すために何年も費やしていた仕事を、人工知能は構造式を入れるだけで数時間のうちに候補群へと絞り込んでくれます。

これらすべての予測は、結局のところ、よく整理されたデータベースに支えられてこそ可能なことです。COCONUT（ココナッツ）という名前で知られた公開天然物データベースは、全世界の文献に散らばっていた天然物の構造情報を一箇所に集め、誰もがダウンロードできるように整理しました。天然物アトラス（Natural Products Atlas）は、微生物が作り出した天然物の出所と生産菌株の情報を標準化された形式で整理しておきました。データが散らばったままであれば、いくら優秀な人工知能モデルでも、学習させるための材料を得ることができないので、このような公開データベースを着実に更新し、誤りを取り除く作業は、華やかな新薬発表の後ろに隠れてよく現れないものですが、なくてはならない基礎工事に該当します。

GNPS2のような質量分析データベースとこのような構造データベースを相互に接続すると、初めて見る試料一つのスペクトラムだけでも、その中にどのような化合物が入っているか事前に絞ることができます。

同じような解読作業は、目に見える植物を超えて、土の中の微生物のゲノムにも広がっていきます。微生物は、自分が作る化合物の設計図を、生合成遺伝子クラスターという遺伝子の束の中に保存しておきます。この束をゲノム全体から探し出す作業をゲノムマイニングと呼びます。2025年に発表された、ある自動化ゲノムマイニング研究は、ストレプトマイセス菌株から、自分自身を保護する耐性遺伝子を手がかりにして、生理活性天然物を高速で探し出す方法を提示しました。新たに開発されたゲノムマイニング・ツール、NPtagMは、酵素予測の精度を、既に広く使われていたツール、antiSMASHの3倍に引き上げながらも、実行時間を8分の1に削減したと報告されています。

e-DeepBGCやBiGCARPのようなディープラーニング・モデルは、遺伝子配列だけを見て、その中にどのような種類の生合成経路が隠れているのか、抗生物質なのか抗がん物質なのかを分類して出します。深海熱水噴出口や極地方の氷のように、人が接近するのが難しい環境に住む微生物は、ほとんどが実験室でさえ培養されていないことが知られていますが、ゲノム配列さえ確保すれば、培養することなく、人工知能がその中に眠っている化合物の設計図を事前に読み出すことができます。これまでゲノム配列が確保された菌株の中でも、その中に含まれた生合成

遺伝子クラスターの実際の産物まで明らかにされた例は、一部に過ぎないことが知られており、ゲノムマイニングが今後探し出す化合物の余地は、依然として広く残っています。

新しい抗生物質がなければ、普通の細菌感染さえ治療できない時代がやってくるという警告がずっと前から出ていたので、この土の中の化学工場を読み出す能力は、それ自体で緊急の課題です。

ハリシンとアバウシンは、どちらも人工知能が先に探し出した抗生物質の候補物質です。マサチューセッツ工科大学のジェームス・コリンズ研究室は、その緊急性に早くから応答してきたところです。2020年、この研究チームは、もともと糖尿病治療薬の候補として開発されていた物質ハリシンが、意外にも多くの細菌を殺す抗生効果を持っていたという事実を機械学習で探し出しました。2023年には、抗生物質に効きにくいアシネトバクター・バウマニ菌に効く物質アバウシンを探し出し、この物質が細菌の脂質タンパク質輸送を担うLoIEタンパク質を抑制して細菌を殺すという作動原理まで明らかにしました。

この研究チームは2025年には、既存の化合物リストをさかのぼる方式から離れて、生成型人工知能で存在したことの無い分子を新たに描き出しました。3600万個を超える仮想化合物をコンピュータで作り出して、抗菌性をスクリーニングした末に、治療が難しい淋菌と抗生物質耐性黄色ブドウ球菌、すなわちMRSAを狙う候補物質を探し出しましたが、既存の抗生物質とは構造も作動方式も異なっていました。この成果を足がかりに、アメリカ保健高等研究計画局から新規抗生物質候補15個を設計する事業の支援を受け、非営利団体フェアバイオと一緒に抗生物質人工知能プロジェクトも続けています。

2025年10月には、同じ研究チームが生成型人工知能で新たに探し出した抗生物質候補が、腸内細菌を正確にどのような経路で攻撃するのかまで地図を描き出しましたが、通常は数年かかっていた作用機序の解明を数ヶ月に削減した事例として紹介されました。

マサチューセッツ工科大学の外でも、抗生物質を人工知能で探す研究は複数の方式で進行しています。2026年2月、ある科学専門誌はセサル・デラフエンテ研究

チームを照らし出し、このチームが化合物リストや微生物ゲノム一つに留まらず、様々な生物の遺伝情報にまで至るまで人工知能で広く調べて、抗菌性を持つ短いペプチドを探し出すと紹介しました。人間の腸の中や皮膚に住む微生物群集もまた、それ自体で膨大な化学情報の倉庫であり、人工知能はこの情報の中から病原菌だけを選んで殺す短いアミノ酸配列を探し出すのに、ますます上達しています。

抗生物質の材料を特定の植物や特定の微生物一つに限定せず、地球上のすべての生物の遺伝情報を候補の倉庫として見るこのような視点は、データを扱う人工知能があるからこそ、初めて現実になっています。

自然の化学的多様性の中でも、厚い地層一つは、数千年の間に積み重ねられた伝統漢方医学の処方記録です。一つの薬草ではなく、複数の薬材を定められた比率で混ぜて複数の標的に同時に作用させるこの処方体系は、長い間経験医学としてのみ扱われ、科学的根拠を認めてもらうことが難しかったです。完成した結果物を逆にばらばらにしながら設計原理を遡る逆設計技法をここに適用すると、処方の中の複数の薬材がそれぞれどの標的を押さえ、その力をどのように合わせてシナジーを出すのかを一つ一つ遡ることができます。老化を遅くする成分一つを実験で確認するには、数百個の処方と数千個の成分を一つ一つ細胞に振りかけて試してみなければならないのですが、人工知能はこのプロセスをデータの上でまず濾し出します。

人工知能を使った新しい薬の開発企業であるInsilico Medicineは、人とハツカネズミとマカクザルなど、いろいろな種の組織から取り出した120万件を超えるオミックスデータ、つまり遺伝子とタンパク質と代謝物質を大量に測定した実験データを学習させて、Precious3GPTという基盤モデルを作りました。いろいろな課題に広く使えるように、あらかじめ幅広く訓練しておいたモデルという意味で、このようなモデルをファウンデーションモデルと呼びますが、Insilico Medicineはこのモデルを、韓医学の成分と標的の情報を集めたBATMAN-TCM2データベースと結びつけました。この結びついたモデルは、人和其他の動物の種にまたがって共通して現れる老化に関連する遺伝子13個を見つけ出し、これらの遺伝子と重なる韓薬の成分34種類を選び出しました。

細胞が分裂するたびに、染色体の端の部分にあるテロメアが少しずつ短くなり、ある瞬間に細胞そのものが老けてしまいますが、その結果、Hwasanoojadan (Hua Shan Wu Zi Dan)という古い処方に入っているいろいろな成分が、肝臓と肺と筋肉の組織でこのテロメアが短くなるのを遅らせる方向に一緒に働くという事実が、分子のレベルで確認されました。

一つの処方を読み解くことからさらに進んで、韓医学の知識全体を検索できる地図に変えようとする試みも続いています。データベースであるTCMBankは、文献マイニングによって韓薬材9192種と成分6万1966種、標的タンパク質1万5179種、関連する病気3万2529種の間のつながりを整理しておきました。これにとどまらず、各成分の3次元構造をDrugBankやPubChemのような西洋圏の新しい薬のデータベースと互いに結びつけておき、韓薬の成分とすでに承認された合成医薬品の間の化学的な似ているところまで、一目で比較できるようにしました。このような比較のおかげで、研究者たちは新しい化合物を最初から合成しなくても、すでに他の病気に使われていた薬が韓薬の成分と構造が似ているという理由だけで、新しい適応症の候補として再検討する薬物再創出の作業を一緒に進めることができます。

TCM-GPTは、韓医学の文献だけで事前学習した最初の70億パラメータ級の言語モデルで、韓医師の資格試験の問題と診断の課題で、これまでの汎用の言語モデルを上回る成績を出しました。北京の人工知能企業であるShuimu Molecular (水木分子)は、清華大学のAIR研究所と手をつないで、1000億個のパラメータを持つ巨大言語モデルをもとにChatDDという対話型の新しい薬の開発秘書を世に出しましたが、標的を見つけ出すことと候補物質を探すことを助けるモジュールと、臨床試験の設計を助けるモジュールを分けて提供し、韓医学の近代化の研究にも応用を広げています。成分と標的と病気をそれぞれ別々に見ていた西洋医学の単一の標的へのアプローチでは解けなかった問題が、いろいろな成分が網のように絡み合って複数の標的を同時に押さえる韓医学の原理を人工知能が数値に移すことで、再び開かれています。

マカオ大学の研究チームは、アルツハイマー病のように原因となる標的がいくつも絡み合った神経退行性の疾患を狙って、韓医学の文献と化合物のデータを統合

した別の人工知能プラットフォームを構築し、複数の標的にアプローチする方法がこのような治りにくい病気にも通じるかどうかをテストしています。

このような成果があっても、韓医薬のデータを人工知能がそのまま飲み込むのは簡単ではありません。脈を測り、舌を見て、症状を尋ねる方法で積み重なってきた診断の情報は、最初から数値として残っていませんし、処方記録も病院や韓医院ごとに書き方がバラバラなので、一つの表に整理するだけでもかなりの手直しが必要です。データベースを構築する研究者たちが論文や古い本を一つひとつ読みながら、成分と標的の情報を人の手で2回ずつ検証する理由もここにあります。

標準化されていない資料を無理に学習させると、人工知能は存在しない相関関係を実際のように提示する危険があるため、資料を整える手直しの作業は華やかなモデルの開発に劣らず重要です。

Hanguk Gwahak Gisul Yeonguwonは、独自に持っている天然物の試料の情報をデジタル化してタンパク質の標的を予測する人工知能プラットフォームであるNPI-finderを開発し、柿の葉のエタノール抽出物がドライアイの症状を改善する効能を証明して、Sikpum Uiyakpum Anjeoncheoの承認と技術移転まで引き出しました。2020年にSon Mi-won代表が設立したEmterapamaは、いろいろな成分でできた天然物の働く仕組みを精密に明らかにするビッグデータプラットフォームであるSyMthomicsを構築し、このプラットフォームから出た糖尿病性の神経障害の治療薬の候補物質は、アメリカの臨床第2相試験で痛みを48パーセント減らすという結果を得て臨床第3相試験を準備しており、パーキンソン病の治療薬の候補物質は、2022年にアメリカの臨床第1相試験を終えてAbbVieやMSDのようなグローバルな製薬会社の関心を集めました。

Yeongnam Daehakgyoの医生命工学科のChoe In-ho教授の研究チームは、人工知能とインシリコ分析を結びつけてサルコペニア（筋肉減少症）を治療する天然物の候補物質を見つけ出したと2026年に発表しましたが、実験をせずにコンピューターの分析だけで候補を絞り込んだ後、細胞と動物の実験で確認するという順番にそのまま従いました。政府も2026年から2030年まで続くJe 5-cha Haniyak Yukseong Baljeon Jonghap Gyehoekに、韓医薬の人工知能・デジタル大転換を重要な目標として釘を刺し、この流れに力を添えています。しばらくの間、国内

の製薬業界で後回しにされているという評価を受けていた天然物の新しい薬は、見つけ出す費用と時間を大きく減らしてくれる人工智能のおかげで、再び注目される財産として浮び上がっています。

人工智能がいくら多くの候補を選び出しても、最後の確認はやはり培養皿と動物のモデルを経なければなりません。マサチューセッツ工科大学の研究チームが3600万個の仮想の化合物のなかから抗菌効果が期待される候補を絞り込んだ後も、実際に化合物を合成して細菌にさらして生き残るかを見守る実験はそのまま必要でした。Hanguk Gwahak Gisul YeonguwonのNPI-finderもやはりコンピューターの予測だけで終わらず、柿の葉の抽出物を実際に適用する試験の段階を経て初めて、Sikpum Uiyakpum Anjeoncheoの承認を受けることができました。

予測と実験を交互に繰り返し、実験結果が予測とずれるたびにその誤差を再びモデルに学習させるこの循環構造を、ハイブリッド探索と呼びます。人が到底処理しきれなかった探索範囲を人工智能が先に絞り込み、そのように絞り込まれたリストだけを人が手作業で確認するという役割分担が、現在この分野の標準です。人工智能が提示した予測の一つ一つがすぐに正解になるわけではなく、実験で確認すべき仮説であると認めてこそ、この循環は初めて空回りすることなく、次の候補物質へとつながります。

植物化学物質データと韓医学の処方データ、そしてディープラーニングと結合親和性予測アルゴリズム、すなわち薬物分子が標的タンパク質にどれほど強く結合するかを計算するアルゴリズムが共に出会う時、自然模倣構造という新しい設計方式が生まれます。最初から化合物を白紙から描く代わりに、自然が数億年にわたって形作ってきた骨組みを人工智能が学習し、その骨組みを変形させて新しい候補物質を提示する方式です。フェリックス・ホフマンがヤナギの樹皮一つからアスピリンを作るには、数十年の化学知識が蓄積されなければなりませんでした。今では同じ系列の骨組みを数千個の変形体として同時に描き出し、その中で有望度の高い候補を数日以内に選び出すことができます。

最近公開されたAlphaFold3系列の構造予測モデルを、化合物とタンパク質の結合の強さを評価する課題に合わせて調整したモデルは、これまで数日ずつかかっていた自由エネルギー摂動計算と同程度の正確度を、はるかに少ない演算で達成し

たと報告されました。抗生物質や抗がん剤、脂質降下剤をはじめ、すでに承認されているさまざまな治療薬が、突き詰めれば自然に由来する骨組みから直接または間接的に派生した物質であるという点を考えると、人工知能がその骨組みの文法を学習して新しい構造を提案する現在の流れは、自然な次の段階であるように見えます。

天然物が新薬の材料として格別な理由は、数億年にわたる進化がすでに一度ふり落とした構造であるためです。植物や微生物は捕食者を防いだり競争者を抑えついたりするために化学物質を作り出してきましたが、その過程で生きている細胞のタンパク質と噛み合うように洗練された立体構造が自然に蓄積されました。実験室でコンピューターを用いて最初から設計した化合物が、しばしば細胞膜を通過できなかったり体内で直ちに分解されたりするのは異なり、天然物の骨組みはすでに生物学的環境で生き残るという試験を通過した構造であると言えます。

抗生物質や抗がん剤、免疫抑制剤のように、現在病院で使われている複数の薬物系列がこのような自然由来の骨組みから派生したという事実は、この強みを裏付けています。化学情報学の研究者たちの間では、無作為に生成した仮想の化合物よりも、自然に由来する骨組みを出発点とした方が、新薬候補が臨床段階まで生き残る確率がさらに高いという見解が力を得ています。

しかし、同じ理由で天然物は扱うのが困難です。一つの植物の中にも数百種類の成分が混ざっており、どの成分が薬効を示すのかを見分けるのが難しく、有効成分が抽出物全体の100万分の1のレベルでしか存在しない場合もよくあります。自生地的气候や土壌によって成分の含有量が年ごとに変わるため再現性を確保するのが難しく、構造が複雑な分子は化学的に大量合成するのも容易ではありません。実際に、パクリタキセルを初めて生産していた時期には、タイヘイヨウイチイ数千本の樹皮が剥ぎ取られ、その木自体が資源枯渇の危機を迎えることもありましたが、有効成分を採取しようとして原料となる植物や菌株を絶滅の危機に追い込んでしまうという逆説は、現在も完全には解決されていません。

古典的な探索方式は同じ化合物を繰り返して再び見つけ出すことが多く、試料をいくらたくさん集めても、本当に新しい構造に出会う確率が低いという限界も長い間指摘されてきました。人工知能は、すでに知られている化合物の質量分析の

指紋をあらかじめ取り除くという方法でこのような重複発見の無駄を減らしてくれますが、そうして見つけ出した物質を実際の植物や微生物から十分な量で再び抽出する作業は、依然として人の手を必要とします。

天然物が体内に入った時にどのように吸収され、分解され、毒性を引き起こすのかを予測する作業、いわゆるADMET予測においても、人工知能の用途が広がっています。様々な成分が入り混じった抽出物は、成分同士が互いに吸収を助けたり妨げたりする相互作用が絡み合っており、単一の化合物を扱うように毒性を計算することが困難でした。最近では、個別の成分の構造だけでなく成分の組み合わせの比率まで共に入力し、肝毒性や心臓毒性が現れる確率を計算する多成分予測モデルが次々と発表されています。

精製された単一の成分ではなく、元々の複雑な抽出物をそのまま規制機関から承認してもらおうとする天然物新薬の開発会社にとって、このような予測能力は、臨床試験を設計する段階から失敗の確率を下げる実質的なツールとなっています。

自然から得た知識をデジタルデータに変えるこの作業は、その知識が本来誰のものであったかという問いをも同時に呼び起こします。2010年に採択された名古屋議定書は、ある国の生物資源とそれに付随する伝統的知識を利用して利益を得たならば、その利益を資源の原産国や地域共同体に分配するように定めた国際条約です。問題は、最近の人工知能が扱うものが、植物の標本や種子のような物理的資源ではなく、その遺伝子配列を書き写したデジタル配列情報であるという点にあります。種子は税関を通りますが、データはそうではありません。生物試料を一つも国境を越えさせることなく、公開されたデータベースから配列情報だけをダウンロードして人工知能に学習させることができるようになったことで、原産国にどのように利益を分配するのかを巡る論争が長く続いてきました。

マダガスカルのコウチクトウ科の植物から作られた抗がん剤ビンクリスチンは、こうした懸念を示す事例としてよく取り上げられます。欧米の製薬会社がこの植物から莫大な売り上げを誇る抗がん剤を開発しましたが、肝心の原産国であるマダガスカルにはそれにふさわしい報酬がほとんど還元されなかったという指摘が長い間続いてきました。交渉の場で、遺伝資源が豊かな開発途上国はデジタル配列情報も物理的資源と全く同じように扱うべきだと主張した半面、この情報を主

に扱う先進国の製薬会社や学界は、過度に厳格な規制はかえって研究の速度を遅らせると対立しました。

2025年2月25日、イタリアのローマでは、生物多様性条約名古屋議定書の締約国がカリ基金（Cali Fund）という国際資金メカニズムを立ち上げました。その2日後には、デジタル配列情報に関する公式決定文も採択されました。遺伝子資源のデジタル配列情報を利用して利益を得る企業は、売上の0.1パーセント、または利益の1パーセントをこの基金に納めるよう勧奨されており、集まった資金は生物多様性に恵まれた開発途上国と先住民共同体の支援に充てられます。

まだこの勧奨に従う企業が多くなく、法的拘束力も弱いですが、デジタル情報だけで新薬を設計する時代に、原産地の国に利益をいかにして返すかという国際社会の第一歩という点で意味があります。漢方処方を逆算し、微生物ゲノムをマイニングするすべての作業の基底には、結局のところ特定の土地と特定の共同体が長年守ってきた生物資源があるという事実を、この基金は数字で改めて確認させてくれます。

漢方処方ひとつをグラフに変え、微生物ゲノムひとつを言語モデルに入力する作業は、今や技術的には難しくありません。実際に長くかかる作業は、そうして見つけた候補物質を実際の患者に安全に使える薬にする最終段階と、その発見の根となった知識と生物資源から得た利益を正直に分け合う手続きです。グラフニューラルネットワークも大規模言語モデルも、結局のところ人類が何千年もの間観察し、記録し、失敗しながら蓄積してきた資料の上で動作しています。ヨモギを低い温度で煮出せという古い句が、ノーベル賞につながったように、今データベースのどこかに眠っている処方と配列ひとつも、次の新薬へとつながる可能性があります。

数億年の進化と数千年の経験が残した記録を人工知能が読み取る今、その記録を最初に残した土地と人々に正当な利益が返されているかどうか、この技術の次の成否を左右するだろうと考えられます。画面内のグラフと実験台の培養皿、そしてその根が育った森と村までもすべてがつながってこそ初めてひとつの薬が完成されるという事実は、人工知能がいかに精巧になったとしても変わらないだろうと考えられます。

韓国漢方医学研究院は、この10年間、人工知能関連の課題と体質ゲノム疫学の課題を通じて、約5万人に近い臨床データを集めてきました。漢方医の経験と主観的判断に依存する診断方式は、漢方医学を科学として認められるようにしようとする努力において、長年にわたって課題として指摘されてきましたが、研究院はゲノムと顔写真とウェアラブルデバイス測定値と一緒に集めて、この判断を数値で支える方向を選択しました。

これに加えて、漢方界は国際医療情報標準であるFHIR形式に合わせて臨床記録を整理する作業も自ら始めました。この基準に従えば、漢方医院での少ない診療記録も、病院の電子カルテや健康保険資料といった他の医療情報とつなげて見ることができるようになります。研究院はこのようにして整備した資料を、漢方精密医療技術の開発に使う計画です。

成分を一定に合わせる作業は、規制の障壁を超えるのにも同様に課題となります。2015年、監査院は天然物新薬制度が国内企業に過度な優遇を与えていると指摘しました。その後、食品医薬品安全処は、天然物新薬という別の分類そのものを廃止し、審査基準を新薬と同じレベルに引き上げました。米国食品医薬局も2016年に天然物医薬品の開発ガイドラインを発表してこの問題を解決しようとしてみましたが、植物抽出物を加工せずそのまま薬として承認された事例は、2006年のベレゲンと2012年のクロペレメル程度で、今でも限定的です。

採取時期と産地により成分含有量が異なる物質について、審査官が毎回同じ薬だと判断するのは簡単なことではないからです。人工知能が成分偏差を事前に計算して補正値を提示したとしても、その値に基づいて承認の可否を決める作業は、今後も人間の審査官が担当する可能性が高いです。

伝統知識が新薬へつながった事例ひとつは、利益をどのように分け合うかを考えるための糸口を与えてくれます。ペルー領アマゾンの先住民たちは、クロトン・レクレリ樹から出る赤い樹液を、傷と腹痛を治す薬として昔からずっと使ってきました。この知識に目をつけた製薬会社は、19の先住民共同体、30余りの場所から樹液を持続可能な方法で集める関係を数十年間続けてきました。このようにして作られたクロペレメルは、2012年にエイズ患者の慢性下痢を治す薬として米国食品医薬局の承認を得ました。会社はその見返りとして、採取地域に教育と給水

事業を支援してきたと明らかにしています。

契約書1枚よりも数十年にわたって築かれた信頼がこのような分け合いを支えてきたわけですが、配列情報だけでも新薬を設計できるような現在のような時代には、会社の善意だけで原産地の共同体を長く守るのは難しいだろうと考えられます。

3.3

ナノ医学およびスマート標的送達システム (Nanomedicine)

2020年12月、アメリカとヨーロッパの空港には、マイナス70度に耐える特殊な冷凍コンテナを搭載した貨物機が次々と降り立ちました。箱の中には、ファイザーとバイオエンテック(BioNTech)が一緒に開発したコロナ19ワクチンBNT162b2が積まれていて、そのワクチン一滴の中には目に見えない油の粒が数十億個浮かんでいました。この油の粒一つ一つがまさに脂質ナノ粒子(LNP)です。遺伝子配列が公開されてから25日で初期設計を完了させたモデルナ(Moderna)も、8ヶ月という記録的な速度で緊急承認を獲得したファイザー・バイオエンテックも、勝負を決めたのはワクチンの核心材料であるメッセンジャーリボ核酸(mRNA)ではなく、その脆弱な遺伝物質を体内まで無事に運んでくれた油の粒、つまりナノ粒子でした。

mRNAは単独では体の中の酵素にすぐに分解されてしまう弱い分子でした。この油の粒が傘のように包み込んでくれなかったら、ワクチンは世に出ることはなかったでしょう。この日以降、世界中で数十億人が腕をまくり上げてこの油の粒を打たれ、脂質ナノ粒子は実験室の見慣れない素材から、人類全体が身をもって経験した最初の商用ナノ医薬品へと一気に躍り出ました。この油の粒技術の根底には、カナダで数十年間脂質ナノ粒子を研究してきたピーター・カリス(Pieter Cullis)と、彼が設立したアキュイタス・セラピューティクス(Acuity Therapeutics)があり、ファイザー・バイオエンテックが使用した脂質ナノ粒子の中核特許の多くがこのカナダ企業から生まれました。

その日滑走路に降り立ったのは単なるワクチンの箱ではなく、ナノ医学が実験室の塀を越えて全人類の腕に注射される瞬間でした。

1ナノメートルは1メートルを10億個に分割した大きさです。ガラス玉一つを1ナノメートルだと考えれば、1メートルは地球一個と匹敵する大きさになるほど、非常に小さい世界です。ワクチンに使用された脂質ナノ粒子は直径がおよそ100ナノメートル前後で、人間の髪の毛の太さの千分の1にも満たないのです。私たちの

体の赤血球一つがおよそ7000ナノメートル程度なので、このナノ粒子を一行に並べると赤血球一つを横切るのに70個程度が必要な計算になります。肉眼でも、通常の顕微鏡でも個別には見ることはできない大きさの世界で、ナノ医学のすべての物語が始まるのです。

ところがこのような小さい世界では、大きさが数ナノメートル異なるだけでも、粒子が肝臓へ行くのか肺へ行くのか、細胞の中に入るのか外側で濾過されるのかが全く変わってきます。表面にどのような分子を付けたのか、表面が正電荷を帯びるのか負電荷を帯びるのかも同じく目的地を変える決定的な要因です。郵便局の小包がサイズと外観表記によって異なる窓口に分類されるのと同じく、ナノ粒子もサイズと表面特性によって体内のそれぞれ異なる臓器へ分けられます。体内の血管と細胞膜に開いた隙間のサイズもそれぞれ異なるので、大きすぎると全く入ることができず、小さすぎると腎臓で濾過されてすぐに体外に排出されてしまいます。

ナノ粒子を利用して薬を運ぶという考え方自体は新しくありません。1995年、アメリカ食品医薬品局は抗がん剤ドキソルビシンを油膜で幾重にも包んだドキシル(Doxil)という薬を承認し、ナノテクノロジーが実際の患者に使用された最初の事例を残しました。その後、アルブミンという血液タンパク質に抗がん剤パクリタキセルを付着させたアブラキサン(Abraxane)といった後続薬が出てきてはいますが、その進歩は遅いままでした。2025年一年を基準に、世界全体で臨床に使用することが承認されたナノ医薬品を全て合わせても50個から80個程度に留まり、そのほとんどが数十年前の化学的基本原理から大きく逸脱していなかったのです。

新しい組み合わせ一つを試験管で作ри、細胞と動物に入れてみて結果を得るまでに数ヶ月がかかり、結果が良くなければ最初からやり直さなければならなかったのが理由です。30年近く人間の直感と繰り返し実験に頼ってきたこの遅い歩みが、人工知能に出会うことでようやく速まり始めたのです。

2025年、ナノサファリ(NanoSafari)という名前の人工知能エージェントが2万編を超えるナノ科学論文を昼夜を問わず読み続けました。このエージェントは論文ごとに散らばっていたナノ粒子合成条件、サイズ、表面特性データを一つにまとめ、大規模言語モデルと接続した知識データベースを構築しました。研究者一人

が生涯をかけて読むべき量を人工知能が一晩で読み切るということなので、数ヶ月かかっていた文献調査が一日で終わるのと変わりがありません。以前は研究者が実験台の前で数ヶ月、数年を試行錯誤に費やしてナノ粒子一つの最適な配合を見つけ出すことができたのです。

今は原子と分子を煉瓦のようにひとつひとつ積み上げて、望む機能を備えた構造物を構築するナノアーキテクトゥクス(Nanoarchitectonics)というアプローチが人工知能の助けを借りて急速に広がっています。建築家が設計図を描いてから初めて建物を建てるのと同じように、人工知能がまず分子の設計図を描き、その設計図の通りナノ粒子を組み立てる順序が標準として定着しつつあります。このような設計方式は新薬開発に留まらず、化粧品や精密化学素材を製造することにも広がっており、ナノ単位設計全般を包括するアプローチとして領域を広げていく様相を呈しています。

成均館大学の研究チームが2025年に国際学術誌アドバンスド・マテリアルズ(Advanced Materials)に発表した論文は、このような流れを明確に整理して示しています。病気を引き起こす体内の標的を生物情報学で見つけ出す段階から出発し、機械学習でナノ粒子の表面を精密に設計する段階を経て、インシリコ、つまりコンピュータの中でこの粒子が体内のどこに広がり、いつ薬物を放出するかを事前に計算してみる段階へと続く一つの連鎖です。この連鎖は一方向で終わらず、再び最初に戻ってそれが次の設計に反映される閉じた循環構造を成しています。

以前であれば数週間かかっていた1回の設計と合成、評価の周期が、コンピュータの中では1〜2日で何度も繰り返すことができます。この循環が密になるほど、見当違いの臓器に薬物が蓄積するオフターゲットの副作用は減り、患者一人ひとりの体の状態に合わせた設計も一層近づいています。研究室のベンチの上で繰り返していた試行錯誤を、コンピューター画面上のシミュレーションに移し替えたという点で、これはナノ医学研究の速度を根本的に変えた転換点と言えます。

このように設計されたナノ粒子を一つ紐解いてみると、まるで丹精込めて包装されたプレゼントの箱に似ています。一番外側の殻には、特定の組織や細胞にだけくっつく標的リガンド、つまりタンパク質や糖鎖のように体内の特定の受容体に正確にくっつく小さな目印が付いています。この目印が、腫瘍微小環境(TME)や

血液脳関門(BBB)のような難しい目的地を見分けて扉を開けてくれます。その内側には、機械学習によって厚さや配列、構成比率まで細かく計算された骨組みが置かれ、粒子全体の大きさと硬さ、分解速度を決定します。

そして一番内側の中心部には、実際に体に作用する薬物、つまりペイロードが入っており、目的地に到着した後になって初めてその役割を果たし始めます。3つの層を切り離して設計できるという点が重要で、中身の材料はそのままにして、外側の目印だけを付け替えれば、全く別の臓器を狙うナノ粒子を新たに作り出すことができます。この3つの層のうち、どの1つの層でも設計が狂うと、ナノ粒子は見当違いの臓器で薬物を漏らしたり、目的地に届く前に体外へ洗い流されたりしてしまいます。

6万個余りに及ぶ仮想脂質化合物のライブラリーをアジャイル(Agile)というディープラーニングのプラットフォームが自ら学習した後、実際の実験室で大量に合成して測定したデータで再び見直す過程を経て、イオン化脂質が細胞の中へ遺伝物質をどれほど上手く運ぶかをあらかじめ見極めます。2025年8月の国際学術誌ネイチャー・ナノテクノロジー(Nature Nanotechnology)には、コメット(COMET)というTransformer構造のニューラルネットワークモデルが紹介されましたが、このモデルはランス(LANCE)という過去最大規模の脂質ナノ粒子データセットで学習し、見慣れない組み合わせの脂質ナノ粒子に対しても効能を予測し出すことができました。

同じ年の11月には、ジョンズ・ホプキンス(Johns Hopkins)の研究陣が数百種に及ぶ脂質ナノ粒子の配合を高速で作って試験した後、機械学習によって遺伝子編集の効率を高める配合条件を見つけ出したと明らかにしました。2026年には、複数の大規模言語モデルのエージェントが役割を分けて協業しながら新しい脂質構造を提案し、互いの設計を交差検証して安全性まで一緒に確認するリポエージェント(LipoAgent)というシステムも学界に登場しました。実験室のベンチがデータ工場に、コンピューターが設計者に変わっていく流れをよく示している事例です。

脂質ナノ粒子をどの臓器に送るかを決めるためには、選択的臓器標的化(SORT)という技術が広く使われています。鍵の束からどの鍵を選んで握るかによって異なる扉が開くように、脂質ナノ粒子の表面に加える脂質分子の性質によって、粒子

が向かう臓器が変わります。正の電荷を帯びる脂質を加えると肺へ、負の電荷を帯びる脂質を加えると脾臓へ、イオン化するアミノ脂質を加えると肝臓へと粒子が集まっていくという事実が実験で確認されました。粒子の表面を覆っていたポリエチレングリコール成分が体内で剥がれ落ちると、血液中の特定のタンパク質がその位置に新しくくっつき、このタンパク質が各臓器の細胞の受容体と対になりながら目的地が定められるという順序です。

臓器を選んで狙えるようになったことで、肝臓でだけ働くべき遺伝子治療薬が肺や心臓まで広がって引き起こしていた副作用も減る傾向にあります。人工知能は、この微妙な化学的組み合わせを人が手で一つ一つ実験する代わりに、数多くの候補の中から素早く選び出す役割を担っています。

脂質ナノ粒子に劣らず、ポリマーやミセル系の運搬体も人工知能の手を経て精巧になっています。微細な流路で粒子を作り出すマイクロ流体力学の機器と高解像度のイメージングデータを一緒に学習した機械学習モデルは、粒子の大きさと表面電荷、形まで自ら調整し出します。ナノ粒子が血流に入った瞬間、血漿中のタンパク質が粒子の表面に群がってくっつき、まるでコートを羽織るように新しい表面を作る現象をタンパク質コロナと呼びますが、このコートが何で編まれるかによって、粒子がどれくらい長く血液中に留まり、どの細胞に吸収されるかが分かれます。

研究者たちは、817種のナノ粒子の配合と2497種のタンパク質の吸着データを集めたデータベースをもとに、複数の小さな判断規則を集めて結論を出すランダムフォレストやエックスジブーストのような解釈可能なアルゴリズムを訓練し、粒子の大きさと表面電荷、血液中に留まった時間がタンパク質コロナを左右する核心的な変数だという事実を明らかにしました。最近では、ディーエヌエー(DNA)の鎖を折りたたんで作ったはるかに複雑なナノ構造物でもタンパク質コロナを予測する研究が続いており、予測の対象が脂質や高分子を越えてだんだんと広がっている傾向にあります。

この予測が精巧になるほど、長期間血液中に留まらなければならない抗がん剤の運搬体と、標的細胞に素早く吸収されなければならないワクチンの運搬体を、異なる表面設計で区別して作ることができます。臨床試験に入る前に、コンピュー

ターの中でこのコートの形をあらかじめ描いてみるできるようになったわけです。

脳を治療の対象にする場合、はるかに高い壁が待ち受けています。脳血管の壁には、密集した網目のような血液脳関門（BBB）が取り囲んでいて、経口投与される治療候補物質100個のうち、この障壁を通過するのは多くても6個、少なくとも2、3個にすぎません。アルツハイマー病やパーキンソン病を狙って作られた新薬候補は、動物実験では効果を示しながらも、人間の体でこの障壁を越えられず、臨床試験の終盤で崩れ落ちた事例は数えきれないほど多くあります。トランスフォーマー系モデルであるMegaMolBARTとXGBoostを組み合わせた予測モデルは、約8,000個の化合物の通過可否を整理したデータベースで学習して、予測が実際のどの程度一致するかを示す曲線下面積指標で0.88というスコアを出しました。

2025年8月、アメリカのランターンファーマ（Lantern Pharma）はpredictBBB.aiというAIモジュールを公開しましたが、化合物の構造情報のみで血液脳関門を通過するかどうかを94パーセントに達する精度で判別できると発表しました。グラフニューラルネットワークで分子構造を読み出すこのような予測モデルは、新薬候補を動物実験に入れる前に、コンピュータ画面上でふるい分ける門番の役割を果たしています。このような予測ツール一つがふるい分ける失敗作のおかげで、動物実験に入っていた予算と時間が大幅に削減されているという評価が出ています。

化学的にこの障壁を通過することが難しい大きな分子は、まったく別の道を選びます。脳血管の細胞表面に多く現れるトランスフェリン受容体のような通路蛋白質に付着する抗体や小さな蛋白質をナノ粒子の表面に付けると、ちょうどトロイの木馬の中に兵士を隠すように、その受容体がもともと行っていた物質輸送の仕事にそっと乗っかって、脳の内側まで一緒に運ばれていくことができます。AIは、抗体のどの部位をどのように変更すれば受容体にぴったり付着しながらも、体の中で思わぬ免疫反応を引き起こさないかを、数多くのアミノ酸の組み合わせの中から探し出すために使われています。

このような受容体媒介通過技術は予測モデルと組み合わせさせて、どの抗体設計が通過確率を大きく高めるかをコンピュータ上で先にふるい分けてから実験で確認

する手順が新しい標準となっています。

合成ナノ粒子だけでは突き抜けるのが難しいこの障壁の前で、科学者たちが注目したのが、生きた細胞が自ら放出するナノサイズの小袋、エクソソームです。体が自ら作って使っていた小胞なので、免疫系が異物と誤認して攻撃するリスクが合成ナノ粒子より低いという点が、臨床の場でエクソソームを好む理由の一つです。脂肪から取り出した幹細胞や骨髄の間葉系幹細胞から得たエクソソームにAIが計算した最適な比率で薬物を詰め込んだ後、鼻にスプレーして投与する臨床試験が複数の場所で進行しています。鼻の中の粘膜と接触した嗅覚神経束を通して入ると、全身を循環する血液循環を経ることなく、脳の病変部位まで直接達することができるからです。

血液という大きな道を経ることなく、嗅覚神経という抜け道で直接脳に入っていく経路というわけです。光で動作をオンオフする蛋白質配送技術や、遺伝子発現を沈黙させるクリスパー干渉技術をエクソソームに乗せて、パーキンソン病にかかった実験用ネズミの脳神経回路でアルファシヌクレインという病を引き起こす蛋白質の遺伝子スイッチを切った実験も報告されています。体がもともと作って使っていた運搬体を借りて使うという点で、エクソソームは脳まで薬物を届ける次世代の運搬体として位置付けられています。

ただし、エクソソームが研究室を抜け出して薬局の棚に上るまでには、まだ距離が残っています。2026年の基準で、世界全体で開発中のエクソソーム治療薬候補は120種類前後で、このうち40パーセント程度が臨床試験の段階に入りましたが、581億ドル規模と推定される市場展望とは異なり、アメリカの食品医薬品局の承認を受けたエクソソーム治療薬はまだ一つもありません。脊髄損傷や視神経損傷を狙ったExoPTENという候補物質も、人間を対象とした臨床試験への進入を目標に前臨床段階を踏んでいて、後続く候補たちの歩みも忙しくしています。

エクソソームを分泌する心臓由来細胞治療薬デラミオセル（Deramioce）を開発してきたキャプリコ・セラピューティクス（Capricor Therapeutics）は、2026年7月29日に開かれた食品医薬品局の諮問委員会で、デュシェンヌ型筋ジストロフィー心筋症に対する効果を実証できなかったという9対3の否定的な投票を受けました。そして8月22日に予定された最終審査の結果を待っています。諮問委員会の

投票には強制力はありませんが、どんなに精巧なAI予測と動物実験の結果が積み重なっても、人間を対象にした最後の関門の前では、依然として慎重な尺度が適用されているという事実を、この事例は示しています。

業界では、デラミオセルの最終的な結論がどちらの方向になろうとも、その結果がエクソソームという治療方式全体を見る規制当局の目線に長期間にわたって影響を及ぼすだろうと見通しています。

脳ではなく固形がんを目を向けると、腫瘍微小環境（TME）という別の課題が待ち受けています。腫瘍周辺の組織は血管が歪に伸びていて、酸素が不足していて、酸性度まで正常な組織と異なるため、どれか一つの方式だけでは薬物をくまなく届けるのが難しいです。血管壁の隙間が特に粗い腫瘍組織にナノ粒子がよく浸透して長く留まる現象を、浸透・蓄積効果と呼びますが、研究チームは378個の腫瘍データセットと200余りの論文から抽出した生物学的応答情報を、ナノ粒子の化学構造と生物学的効果の間の関係を含む構造活性相関モデルと、体の中で薬物が吸収されて広がっていくプロセスを式で解き明かした生理学的薬物動態モデルと一緒に入れて、この現象を24時間後の基準で計算して出しました。

予測値と実際の値がどれほど一致するかを示す決定係数が0.83に達するほど、精度が高くなりました。ここからさらに一歩進んだ生体材料の逆設計は、順序を完全に逆転させます。腫瘍細胞を殺したい、あるいは免疫細胞を目覚めさせたいという最終目標を先に入力すると、アルゴリズムが膨大なデータベースを遡り、その目標に合った高分子の比率や表面リガンド、さらには腫瘍の細胞膜や赤血球膜を模倣した偽装コーティングまで逆算して見つけ出します。腫瘍が放出する特定の酵素や低い酸性度にのみ反応し、そこだけで薬物を放出する賢い運搬体が、このように作られています。

ただし、腫瘍ごとに、さらには同じ腫瘍でも患者ごとに微小環境の詳細な条件が異なるため、一人の患者で検証された設計を他の患者にそのまま適用するのは容易ではない場合があります。

治療薬を体内に送り込むことと同じくらい、病気がどれほど進行しているかを事前に察知する役割も、ナノテクノロジーと人工知能が共に担っています。表面増

強ラマン散乱や局在表面プラズモン共鳴のようなナノ光学技術は、血液や脳脊髄液の中にほんのわずかだけ漂っている病気の信号さえも捉え、光の形で増幅してくれます。簡単に言えば、懐中電灯の微かな光を凹面鏡で集めて明るく照らす理屈と似ています。人工知能はこのように増幅された複雑な光のスペクトルと蛍光信号をリアルタイムで読み取り、アルツハイマー病のマーカであるアミロイドベータやタウタンパク質が凝集し始める初期の瞬間を誰よりも早く発見します。

症状が表に現れる数年前、患者自身も気づかない時点で病気の芽を見つけ出すという点で、従来の診断法とは一線を画しています。高解像度のナノ画像装置とディープラーニングによる画像分析を併用すれば、体内に入れたナノ粒子が実際に病変部位にどれくらい集まったのか、薬物がどのような速度で放出されているのかを体外からでも観察することができます。

ここからさらに一歩進むと、体内に埋め込んだナノセンサーが患者の生体信号をリアルタイムで人工知能に送り返し、人工知能がその信号に合わせて薬物の放出速度を自ら調節する、自律駆動型の精密医療が次の目標として浮上しています。血糖値を測ってインスリンの量を自動で調節する機器がすでに糖尿病患者のそばにるように、がんや退行性脳疾患でも体内のナノセンサーとナノ運搬体がペアになり、人の手を介さずに一日中薬の量を合わせてくれる日が遠くないという見通しが出ています。

ただし、このような自律駆動型のシステムが実際の病室に導入されるためには、センサーが誤った信号を送った時に薬物が過剰に放出されないように防ぐ安全装置も一緒に検証されなければなりません。診断と治療、そしてその間をつなぐ自律調節機能を一つに収めたナノシステムは、ナノ医学が次に進むべき方向を示す例として挙げられます。

人工知能モデルを訓練するのに十分なほど質の高いナノ毒性のデータは圧倒的に不足しており、研究室ごとに測定方法や基準がバラバラであるため、データを一箇所に集めることからして容易ではありません。同じナノ粒子であっても、どの研究室で測定したかによって毒性の数値が変わるケースも珍しくなく、異なる結果のうちどちらを支持すべきか判断が難しい場合も生じます。複数の大学や企業がそれぞれ異なるプロトコルで蓄積してきたデータを一つの標準にまとめる作業

でさえ、今はまだ歩き始めたばかりの段階に留まっています。予測の精度がいくら高くても、ディープラーニングモデルは、なぜ特定のナノ粒子の構造が毒性を引き起こさないと判断したのかを、臨床医や規制担当者に筋道立てて説明できないというブラックボックスの問題を抱えています。

物質の動きを微分方程式で解き明かす伝統的な生物物理学モデルと、データから直接答えを見つけ出す人工知能モデルが互いに歩調を合わせられなければ、コンピューターの画面上では完璧だった予測が、実際の体内では免疫拒絶反応や予想外の副作用へと広がる危険が常に残っています。さらに、人工知能が新たに提案したナノ粒子の構造であっても、人の体に入れるまでには数年にわたる前臨床試験や臨床試験の費用をそのまま負担しなければならないため、設計段階が早くなったからといって、新薬が市場に出るまでの全体の日程までがその分縮まるわけではありません。

制度の時計は、技術の速度にそのまま追いつくことはできていません。2025年に米国食品医薬品局が承認した新薬46件のうち、半分以上が世界初のメカニズムを持つ治療薬だったという事実は、規制当局もそれなりに審査方式を変えつつあることを示していますが、ナノ医薬品のように物質自体が何層にも重なってできている場合には、依然として別の基準が必要です。米国食品医薬品局は近代化法2.0を通過させ、動物実験を必ず経なければならないという長年の規定を見直し、臓器チップやコンピューターシミュレーションで前臨床の毒性試験を代替する道を開いてくれました。

しかし、欧州医薬品庁をはじめとする多くの国の規制機関は、エクソソームのように生きた細胞から出た運搬体や、血液脳関門を行き来する先端ナノ医薬品に対して、それぞれ異なる審査基準を求めており、一つの国で通過した製品が別の国では最初から資料を準備し直さなければならないことがよくあります。

人工知能が患者のゲノム情報を大量に学習する過程で、個人情報についてどこまで同意を得るべきかという境界線もまだはっきりしていません。人工知能が自ら設計し、用量まで決めたナノ粒子が患者に深刻な副作用を引き起こした時、その責任をソフトウェアの開発者と製薬会社、処方した医療スタッフの間で、誰がどれだけ分担すべきかについても決まった答えはありません。従来の製造物責任の

法理は、設計の欠陥や製造の欠陥を、特定の時点、特定の主体の判断の誤りに絞って問う方式に慣れています。人工知能が膨大なデータを学習しながら自ら設計変数を変えていくナノ医薬品では、その誤りが正確にいつ、どの段階で生じたのかを遡ることは容易ではありません。

設計アルゴリズムを作った会社と、それを持ってきて特定の患者に合わせて微調整した病院、そして最終的な処方を出した医療スタッフの間で、責任を分ける基準から新しく作らなければならないという声が、法曹界の内外で大きくなっています。アルゴリズムの判断の根拠を人が見てわかるように解き明かす、説明可能な人工知能の技術を固め、細胞実験と動物実験を細かく経て、国ごとに異なる規制を一つにまとめる作業が一緒に行われてこそ、ナノ医学が臨床現場の標準として定着できると思われます。

この流れは韓国の製薬・バイオ産業にもそのまま続いています。2026年に開かれたNano Koreaは、ナノバイオと3次元印刷、精密計測機器と一緒に集めた世界的な規模の行事として開かれ、人工知能で設計した標的伝達体と診断機器が主な話題として扱われました。国内でも、抗がん新薬や遺伝子治療薬の分野を中心に、ナノ粒子の設計段階から人工知能の道具を引き入れる企業が一つ二つと現れています。

国内の規制機関もやはり、アメリカとヨーロッパの審査基準の変化を注視しながら手続きを整えており、国内の製薬会社と研究機関、そして彼らに助言する法律家たちがこの流れに遅れを取らないためには、海外の臨床データと規制の動向をこつこつと追う努力が続かなければならないと思われます。

2020年の冬、マイナス70度の冷凍箱の中で世界に初めて存在感を知らせた脂質ナノ粒子は、今では脳の奥深くの病変と腫瘍の真ん中を正確に探し出す人工知能設計の運び手に育ちました。その間にナノ粒子を作っていた手は、試験管を握った研究者の手から、データを読み取る人工知能の計算へとかなりの部分が移ってしまいました。だからといって、人がすべき仕事が無くなったわけではありません。どのデータを信じるか、どの危険を引き受けるか、患者に何を先に尋ねるかを決める仕事は、依然として研究者と医療スタッフ、そして規制の担当者が引き受けています。

残された課題は、この精巧な設計図が実験室の画面を抜け出し、病院の診察室と薬局の棚まで無事につながる橋を架けることです。その橋が完成する時期は、結局のところデータの品質と規制の速度、そして患者のそばで安全性を最後まで確認する人の判断が一緒に決めると思われます。その判断が一つ一つ積み重なってはじめて、ナノ医学はコロナ19ワクチンのように危機の瞬間にだけきらっと光る技術ではなく、普通の診察室で毎日使われる道具として残ることができると思われます。

国内に目を向けると、ナノ運搬体を直接作って売った企業の歴史も短くありません。Samyang Biopharmが2007年に国内で許可を受けた抗がん剤Genexol PMは、パクリタキセルを高分子ミセルの中に閉じ込めた薬であり、世界で初めて承認された高分子ミセル抗がん剤として挙げられます。国内の製薬会社がすでに20年近くナノ運搬体を扱ってきたことになります。最近ではST Pharmが独自に開発した脂質ナノ粒子製剤技術STLNP (STLNP)とメッセンジャーリボ核酸キャッピング技術SmartCap (SmartCap)を前面に押し出し、オリゴヌクレオチド委託開發生産の世界第3位の企業という足場の上で、事業領域をメッセンジャーリボ核酸と遺伝子はさみ治療薬までを網羅する統合RNA (RNA) 委託開發生産へと広げると明らかにしました。

ST Pharmはこれに先立ち、Ewha Womans Universityの研究陣と手を組んで脂質ナノ粒子伝達プラットフォームと一緒に開発した経験も持っています。実験室での設計段階にとどまっていた国内のナノ医薬品技術が、大量生産と輸出まで見据える事業へと育った様子です。

しかし、人工知能がいくら精巧にナノ粒子を設計しても、まだ解けていない宿題が一つ残っています。カナダのトロント大学（University of Toronto）のWarren Chan教授の研究チームは2016年、それまでの10年間に発表されたナノ粒子の論文117編を集め、体に注射したナノ粒子のうち、実際に腫瘍に到達する量を計算しました。その値は中央値の基準で0.7パーセントにとどまりました。ナノ粒子を1000個注射すれば、わずか7個だけが腫瘍に届くという意味であり、研究チームはこの数値がその前の10年間でもほとんど変わっていないと明らかにしました。

コンピュータの中でいくら完璧に描かれた粒子であっても、血管に乗って流れながら肝臓や脾臓で濾過され、免疫細胞に捕まるという実際の体の中の旅路までは、人工知能が追いつけていないと思われます。設計の速度は速くなりましたが、伝達効率という古い壁はまだ大きく低くなっていないという意味であり、ナノ医学の次の勝負は、新しい粒子を描くことよりも、すでに描いた粒子を目的地まで無事に連れて行くことで決まる可能性が高いです。

大量生産の壁も甘くありません。実験室の微小流体装置で数ミリリットルずつきれいに抽出していたナノ粒子を工場の設備に移し、数百リットルずつ作り始めると、粒子の大きさと薬物の封入率が揺らぐことがよくあります。配管を流れる速度や温度がわずかに変わっただけでも脂質ナノ粒子の均一性が崩れるため、実験室のデータを工場の規格へとそのまま移すのは難しいというのが、業界の共通した悩みです。人工知能は設計を数日で終わらせてくれますが、その設計を人が安心して飲んだり打ったりできる安定した薬として固める仕事は、依然として指先の経験と繰り返される工程の検証に頼っています。

第4部

非臨床・臨床検証と精密医療の革新

4.1

前臨床研究のイノベーション：動物実験の代替および生物学的モデリング

びまん性胃がんの診断を受けたある患者の腕に、点滴の管が繋がれました。通常の臨床試験の初日と変わらないように見える光景でした。この瞬間が尋常でなかった理由は別にあります。その患者に初めて投与された新薬の候補物質は、誕生してからただの一度もマウスやサルの体を通ったことがありませんでした。設計から安全性の検証まで、すべての過程がコンピューターの中の人工知能と、培養皿の上のオルガノイドの中だけで行われました。

医薬専門メディアのバイオスペースは、この臨床試験を、人工知能とオルガノイドだけで前臨床検証を終えた世界初の事例として報じました。実験室の外に出た動物はただの1匹もいませんでした。それでも、薬は人の病床まで正確に届きました。

このような場面が生まれるまで、製薬産業は長らく違う方式で動いていました。新薬の候補物質が人に投与される前、研究者たちはマウスやウサギ、時にはサルの体を借りて、薬物の吸収や分解、毒性を確認してきました。この方式は数十年の間、新薬開発の唯一の関門でした。問題は、マウスの体と人の体が同じ設計図で作られていないという点にあります。

同じ鍵でも錠の形が違えば扉が開かないように、マウスの肝臓では無事に分解されていた物質が、人の肝臓では全く違った反応を示すことが繰り返されました。遺伝子とタンパク質構造の種間差異のために、動物において安全だと確認された薬物のうちの多くが、いざ人を対象とした臨床試験では副作用を起こしたり、効果を出せなかったりして挫折しました。

この隙間を示す事例として、2006年にイギリスであった出来事ほどふさわしいものは稀です。TGN1412という抗体治療の候補物質は、サルを対象とした試験において、人の予想投与量の500倍を注入しても特別な問題は現れませんでした。し

かし、健康なボランティア6人にいざそれよりはるかに少ない量を投与すると、数時間のうちに彼らの全身で免疫系が暴走するサイトカインストームが起き、複数の臓器が同時に壊れました。

サル免疫細胞と人の免疫細胞は、見た目は似ていても化学的には互いに異なっており、その小さな違いが実験台の上では明らかにならず、人の血管の中でようやく姿を現しました。この事件以降、学界と規制当局は、霊長類の試験結果をそのまま人に当てはめることはできないという事実を痛切に受け入れました。

動物実験の結果と人の実際の反応の間に生じるこの隙間を、業界では橋渡し研究のギャップと呼んでいます。そのギャップが積み重なった結果は、数字としてもはっきりと表れています。新薬開発プロジェクトのうち、90パーセントを超える数が臨床段階で失敗に終わります。この失敗の大部分は、候補物質が人の体に入った後、つまりすでに多くのお金と時間を使った後に初めて明らかになります。一つの新薬が承認を受けるまでにかかる時間は平均10年から15年、投入される費用は26億ドルに達します。

製薬業界は、この逆説をムーアの法則を逆さまにしてエールームの法則（Eroom's Law）と呼んでいます。コンピューターの性能は年々良くなるのに、一つの新薬を世に出す費用と時間は逆に増え続ける現象を指す言葉です。動物実験に対する長年の依存が、この逆説の根っこの一つとして指摘されてきた理由がここにあります。

動物実験を減らそうという声は、長らく生命倫理の言葉だけで扱われてきました。実験室で犠牲になる動物の数を減らそうという主張は、もちろん正当です。しかし、ここ数年の間に、この議論の重心は目に見えて移り変わっています。動物実験を代替することは、今や倫理の問題であると同時に、科学の正確度を高める問題として扱われています。人の細胞と組織で作ったモデルが、マウスの体よりも人の反応をはるかに正確に予測するならば、動物をあまり使わない選択は道徳的に正しいだけでなく、科学的にもより良い選択となります。倫理的な正当性と科学的な正確性という二つの理由が共に会うことで、前臨床研究の地図が描き直されています。

この変化は、規模の大きな製薬会社だけの話でもありません。動物の飼育施設を建てて維持するには莫大な初期費用と長い時間が必要ですが、オルガノイドの培養や仮想細胞のシミュレーションは、相対的に少ない資本でも始めることができます。小規模なバイオベンチャーや大学の研究室も、以前であれば大型の製薬会社だけが手を出せた前臨床検証の段階に直接飛び込める道が開かれています。資本の大きさではなく、仮説の斬新さとデータの質が勝負を分ける構図に変わってきているわけです。

この変化の最初の結び目は、法律において解けました。2022年12月、アメリカ議会は食品医薬品局近代化法2.0を通過させました。この法律は、新薬の承認を申請する際に必ず動物実験の資料を提出しなければならないという、数十年来の義務条項をなくしました。法律の条文に使われた表現も、動物試験から非臨床試験へと変わりましたが、この短い文言一つが、人体由来の細胞モデルやオルガノイド、臓器チップ、そして人工知能ベースのコンピューターシミュレーションを、法的に正当な代替手段として認める根拠となりました。製薬産業全体がざわめきました。動物実験はもはや唯一の正解ではなく、いくつかの選択肢のうちの一つとなりました。

法律が扉を開いた後にも、その扉を実際に広げることはアメリカ食品医薬品局が担わなければなりませんでした。食品医薬品局は、2024年に動物実験を減らすという方向性をまず明らかにし、2025年4月にはこれを具体的なロードマップへと発展させ、今後3年から5年以内に動物実験を新薬開発の標準ではなく例外にするという目標を打ち出しました。

食品医薬品局は、まずモノクローナル抗体、つまり一つの細胞から複製されて全く同じ形を持つ抗体治療薬と他の生物学的製剤から動物実験を減らしていき、その範囲を化学合成新薬へと広げていく計画です。2025年の初めには、医薬品と生物学的製剤の審査において人工知能をどのように扱うかを定めた指針の草案も併せて発表しました。

ロードマップを公開してから1年が経った後、食品医薬品局は初年度の目標を達成したと発表しました。抗体医薬品を開発する際にサルなどの霊長類試験を減らすか廃止する方向のガイドラインの草案を発表し、カブトガニの血液から抽出した内

毒素試験材料を合成物質に変える指針も更新しました。カブトガニの血液に頼ってきたこの試験1つを変えるだけでも、毎年100万頭以上の動物が実験台に上らなくて済みます。アメリカの下院では、ここで止まらず実施速度をさらに加速させるよう求める後続法案まで既に可決した状態です。

2025年初めに公開された人工知能ガイドラインの草案が扱う核心は、リスクレベルに応じて人工知能モデルの信頼性を異なる方法で審査するという考え方です。同じ人工知能でも、新薬候補物質を初めに選別するのに使われるモデルと、臨床試験なしに安全性を最終判断するのに使われるモデルは、全く異なる重さで扱われます。規制決定に与える影響が大きいほど、そのモデルがどんなデータで学習されたのか、予測が外れたときにどんな手順で除外されるのかを、はるかに丁寧に証明しなければなりません。このリスク基準の審査枠組みは、仮想細胞やシステム生物学モデルが臨床試験設計に実際に足を踏み入れるには、次にどんな関門を越えなければならないのかを事前に示しています。

この流れを主導する力は、1つの技術ではなく、互いに異なる2つの方向から一緒にやってきます。一方には、人間の細胞と組織を実際に培養する体外、すなわちインビトロ技術があります。もう一方には、人間の細胞をすべてコンピュータの中に移すインシリコ、すなわちコンピュータの中で行われる技術があります。オルガノイドと臓器チップが前で人間の体を物理的に模倣するなら、人工知能が作った仮想細胞とシステム生物学モデルは後ろで人間の体をデジタルに模倣します。

2つの技術は経系と緯系のように織り交ざりながら、互いに見落とし部分と部分を補い合います。コンピュータが出した予測は生きている人間の細胞で確認され、培養皿の中の実験結果は今度はコンピュータモデルを精密に調整するのに使われます。

このすべてのことが2020年代半ばに集中して起こったのには理由があります。人工知能モデルを訓練するための演算能力は、グラフィックス処理ユニットの性能が毎年大幅に向上するにつれて、過去には想像もできなかったレベルにまで増えました。同時に、単一細胞レベルで遺伝子活動を読み取る実験装置の価格は低下し、処理速度は高速化し、仮想細胞を訓練するのに十分な生物学的データがようやく十分に蓄積され始めました。

ここに、オルガノイドと臓器チップを作る細胞培養と微細加工技術まで成熟するにつれて、デジタルモデルが出した予測を物理的に確認する実験道具まで同時に備わりました。演算能力と生物学データ、そしてこれを検証する実験道具という3つの条件が似た時期に一緒に備わらなかったら、オルガノイドと臓器チップがどんなに精密になっても、それを支えるデジタルの相棒は出現しなかったでしょう。

仮想細胞は1つの細胞をすべて複製してコンピュータの中に移したデジタル双子です。研究者たちは、単一細胞トランスクリプトミクス、空間トランスクリプトミクス、プロテオミクスのような複数のレベルで得たデータを積み重ねて、その細胞の遺伝子発現ネットワークと代謝経路を数学的に再現します。単一細胞トランスクリプトミクスは、細胞1つ1つがどんな遺伝子をどのくらい多く使っているのかを読み取り、空間トランスクリプトミクスはその活動が組織の中のどこで起こっているのかを地図のように描き出し、プロテオミクスは遺伝子情報が最終的に作り出したタンパク質が細胞の中にどのくらいあるのかを測った結果です。

飛行機を実際に飛ばす前にコンピュータの中で何千回も模擬飛行を試みるのと同様です。新薬候補物質をこの仮想細胞に投与したり、特定の遺伝子を人為的に操作するインシリコ実験を通じて、研究者たちは実際の培養皿や動物なしでも、細胞がどのように反応するか、思わぬ標的に結合して副作用を引き起こさないかを事前に評価することができます。

仮想細胞が有用性を持つためには、分子1つ1つの動きから組織全体の変化まで一緒に計算できるシステム生物学モデリングに支えられなければなりません。このモデリングは、データを統計的に調べてパターンだけを見つけ出す機械学習とは異なる性質のものです。ある原因がある結果にどのようにつながるのかを微分方程式で解き明かした生物物理学的な機序モデルを人工知能と一緒に使います。細胞の中で信号が伝わり、物質代謝が回る過程はCOPASIやCellNOptのようなプラットフォームが再現し、腫瘍が成長したり免疫細胞が組織の中に入り込むように複数の細胞が絡み合って作り出す動きはPhysiCellやCompuCell3Dのような多細胞シミュレータが担います。

分子から細胞へ、細胞から組織へとつながるこの計算の鎖が備わってこそ、仮想細胞はようやく信頼して使える予測ツールとなります。

このようなアプローチを産業現場で実際に実装した事例としてRecursion Pharmaceuticalsを挙げることができます。この会社は、実験室用ロボット自動化設備とinVivoPRINTという機械学習プラットフォームを組み合わせ、前臨床研究の方法そのものを変えてしまいました。ロボットが24時間細胞と微細組織を撮影し、センサーで数値を測っている間、ディープラーニングモデルはこの膨大な映像とデータをリアルタイムで読み込んで、医薬品が投与された細胞に見られる毒性と有効性の表現型シグナルを抽出します。

このようにして得た化合物の指紋のようなシグナルを既知の毒性物質のデータベースと比較すれば、昔なら開発後期になってはじめて明らかになったであろう致命的な臓器障害の危険性を、研究の初期段階で事前に捉えることができます。危険な候補物質を早期に除外したおかげで、Recursion Pharmaceuticalsはプロジェクトを開始してからわずか18カ月で、新薬候補物質をヒト対象の臨床試験段階に進めることができました。

Recursion Pharmaceuticalsだけがこの道を歩んでいるわけではありません。人間の手を経ずにロボットが実験を設計し、実行し、次の実験まで自分で判断する、いわゆる自律の実験室が複数の製薬会社と研究所で一緒に増えています。人間の研究者が1日に処理できる実験皿の数は、どんなに頑張っても数十個を超えるのは難しいのですが、ロボットアームと人工知能が噛み合って動く自律の実験室は同じ時間に数千種類の組み合わせを試して、その結果をすぐに次の実験設計に反映させます。

夜通し休まずに動くこの実験室のおかげで、新薬の候補物質を選び出すスピードは、人が手で一つ一つ試験していた時代とは比較にならないほど速くなりました。

仮想細胞技術は、2025年と2026年の間に一気に数歩前進しました。アメリカのアーク研究所は2025年中頃、第1世代の人工知能仮想細胞モデルであるState3を公開しました。このモデルは、1億7千万個にのぼる細胞データと、70種類を超える細胞株から得た1億個以上の単一細胞攪乱反応データを学習して作られました。アーク研究所はここで立ち止まることなく、仮想細胞チャレンジという名前の公開コンテストを開き、さまざまな人工知能モデルが実際の生物学の実験結果とどれくらい近く一致するかを競わせました。

14か国から1200を超えるチームが集まり、300件を超える最終結果物を提出しました。学界はこのコンテストについて、仮想細胞のためのチューリングテストと呼び始めました。コンピューターと交わした会話が人と交わした会話なのか区別できないとき、そのコンピューターが試験に合格したとみなす元々のチューリングテストのように、仮想細胞が出した予測が、実際の生きている細胞の実験結果と区別できないほど正確かを競う試験という意味です。

巨大テクノロジー企業たちの動きも同じ時期に重なりました。チャン・ザッカーバーグ・イニシアチブは、生物学的データから直接学習した細胞と細胞状態の普遍的なデジタル表現、いわゆる人工知能仮想細胞を初めて公開しました。続いてエヌビディアと手を結んでバーチャルセルプラットフォームを開き、仮想細胞モデルを作り育てるために必要なデータを、規模感を持って扱える基盤を整えました。

この財団はrBioという推論モデルも一緒に発表しましたが、これは研究者が普通の文章で生物学の質問を投げかけると、仮想細胞のシミュレーションを元に答えを探してくれる道具です。グーグル・ディープマインドもまた、独自の仮想細胞プロジェクトを進めていると明らかにしました。仮想細胞はもはや一つの会社の賭けではなく、複数の巨大テクノロジー企業が同時に飛び込む競争の舞台になりました。

コンピューターの中の仮想細胞がいくら精巧になっても、生きている人の体の3次元的な生理現象をありのまま代わりにすることはできません。そのため、予測を実際に検証してくれる物理的なペアが必要です。その役割をオルガノイドが担います。マウスの細胞ではなく人の細胞で作られているので、通訳を通さずにネイティブスピーカーが直接話してくれるのと同じわけです。オルガノイドは、患者から得た成体幹細胞や、すでに成長しきった皮膚細胞のようなものを再び幼い頃の万能状態に戻した人工多能性幹細胞を、お皿の上で立体的に育て、実際の臓器に似た細胞の構成や構造、限られてはいますが生きている機能まで備えるように作った縮小版の臓器です。

平らな2次元の培養皿や動物の体では真似できない、人特有の細胞間の信号のやり取りや組織の微小環境を、オルガノイドはありのままに見せてくれます。新薬の

候補物質が人に実際どれくらいよく効くのかを、実験室の中で前もって確認できる道具として、オルガノイドは急速に使い道を広げています。

アルツハイマー病やパーキンソン病のように、神経が徐々に壊れていく疾患を前にして、オルガノイドはひととき大きな役に立ちます。マウスをはじめとする齧歯類の脳は、人の脳特有のしわのある皮質構造や複雑なグリア細胞のネットワーク、タンパク質が絡みつく病理現象をきちんと真似することができません。そのため、長い間、齧歯類の脳は中枢神経系の新薬開発の墓場という不名誉な名前で呼ばれてきました。人の細胞で作ったミニ脳や中脳のオルガノイドは違います。パーキンソン病の核心的な病理である、ドーパミンを作る神経細胞の退行や、アルファシヌクレインというタンパク質が異常に固まる現象を、実際の患者の脳とそっくり3次元で再現します。

最近では、光で細胞の活動を調節する光遺伝学の技術の中脳のオルガノイドに結びつけ、アルファシヌクレインが固まる過程を精密に統制する実験モデルまで出てきました。このモデルの上で研究者たちは、神経を保護する新しい薬の候補の効果を、人の生物学的な文脈の中ですぐに確認しています。

希少疾患の治療薬開発において、これらの技術の重みはさらに大きくなります。患者の数が少なく、動物モデルさえきちんと整っていない疾患も少なくありませんが、このような場合、患者一人の細胞から直接培養したオルガノイドが、事実上唯一の前臨床の検証手段になることもあります。動物モデルを新しく作って検証するのに数年を費やす代わりに、患者の組織検体一つから始めて、数週間のうちにオルガノイドを育て上げ、薬の反応を試験することができます。この数週間という時間差は、希少疾患の患者の家族にとっては決定的です。

オルガノイドにも隙はあります。血管のネットワークがなく、免疫系とやり取りする相互作用が抜けており、体の中の組織が実際に受ける物理的な力も取り込めていません。この隙を埋めようと、微小流体力学を組み合わせた臓器チップ、すなわち微小生理システム（MPS）が前臨床研究の最前線に出ました。この分野には、エミュレート以外にも、ミメタスやティッシュズ、ヘスペロスのような会社が、それぞれ異なる臓器を標的にして飛び込んでいます。臓器チップは、親指ほどの透明なポリマーチップの中に髪の毛より細い流体の通り道を刻み、その中で

人由来の組織細胞と血管内皮細胞を一緒に育てます。

この小さなチップ一つが、血流が流れながら作る圧力、腸がうごめく動き、肺が膨らんだり縮んだりするリズムまで、人の体の中で起きる力学的な環境をそのまま移し替えます。臓器チップが再現するのは細胞の形だけでなく、その細胞が実際に体の中で経験する物理的な条件、すなわち圧力と流れと動きの全体です。

臓器チップの一つ一つは、肝臓や心臓、腎臓、肺のように特定の臓器を一つ真似るだけにとどまりません。複数の臓器チップを細い管でつなぎ合わせれば、薬が体に吸収され、代謝されて体外に抜け出るまでの全過程、すなわち薬物動態学の全体を培養皿の外で続けて観察することができます。例えば、腸のチップと肝臓のチップを並べてつなげば、飲み薬が腸で吸収された後、すぐに肝臓へと移って分解される順序まで、実際の体の中の順序そのままに再現することができます。

臓器チップ同士のこのような接続順序は、個々の臓器チップだけでは決して得られない情報です。脳血管内皮細胞とアストロサイト、血管周囲細胞と一緒に培養した血液脳関門チップは、人工知能が設計した脳疾患治療薬が実際に脳を囲む障壁をどのくらいよく貫通して進むのかを、動物実験よりも正確に明らかにします。薬物が肝臓に損傷を与えたり、心臓の鼓動に異常を引き起こしたりするかどうかは、人細胞で作られた臓器チップの上に植え付けられた複数の微細電極がリアルタイムで捉えます。

これらの電極は、細胞が放出する微細な電気信号の変化を秒単位で読み取るのですが、心臓細胞ならば拍動リズムが乱れる瞬間を、肝細胞ならば細胞膜が損傷されて漏れ出る信号をすぐに察知します。コンピュータの中の仮想細胞が出した予測が不確実である場合、まさにこの物理的なチップ上のデータがその不確実性を払い除きます。インシリコが先に方向を示し、インビトロがその方向が正しいかどうかを体で確認してくれる循環がこのようにして完成します。

肝臓が特に新薬開発の足かせになってきた理由があります。薬物のほとんどが肝臓を通して分解されるのですが、このプロセスで予期しない副産物が作られて肝細胞を損傷させることが頻繁に起こります。薬物による肝損傷は、販売後の薬が市場から撤退する一般的な原因として指摘されてきました。また、動物実験では

無事だった物質が人間の肝臓でのみ問題を引き起こす事例も珍しくありませんでした。

この技術が実験室の垣根を超えて規制現場で力を得たのも最近のことです。米国食品医薬品局は、革新的な科学技術アプローチを審査する通路であるISTAND制度を2025年中盤に正式な制度として固めました。この通路を通じて、エミュレートのリバーチップは、薬物による肝損傷を事前に知る道具として適格性を認められる最終段階を踏んでいます。有毒物質に曝露された組織の形態変化とアルブミン、ALTタンパク質数値の変動だけで、実際の人間の肝臓で起こる損傷を予測することができるという点を検証される過程です。この適格性を得たならば、リバーチップは特定の会社一社ではなく、定められた用途の中で、すべての新薬開発プログラムが使用できる公認された道具となります。

ここで定められた用途、すなわちコンテキスト・オブ・ユーズは、このチップが正確にどのような質問に答えるのに使用できるかを決める規制言語です。肝毒性を予測する用途として認められたからといって、心毒性や腎毒性まで自動的に認められるわけではありません。臓器チップの一つが生み出したデータが、個別の実験の参考資料を超えて、規制審査の公式な根拠として認められる最初の事例が目前に迫っています。

オルガノイドも実験室を脱出して患者の元へ移動しています。2023年から2025年の間だけでも、患者由来オルガノイドを使用した臨床試験が139件登録されており、その流れは2026年にも続いています。切除不可能な肝細胞がん患者を対象に、オルガノイドで薬物反応を事前に試験して治療方針を決める2相臨床試験が2026年初頭に開始されました。同じ肝細胞がん診断を受けた患者でも、腫瘍の遺伝的特性は患者ごとに異なります。この臨床試験は、患者個々の腫瘍から取り出したオルガノイドが複数の薬物にどのように反応するかを皿の上でまず観察してから、その人のオルガノイドがよく効く薬を選んで処方する方法です。

一度薬を使ってみて効かなければ次の薬に変えるという従来の試行錯誤的な処方とは、順序そのものが異なります。オルガノイドが今では実験室で新薬候補物質を選ぶ道具に留まらず、すでに認可された薬物の中からどれを特定の患者に使用するかを決めるところまで関わっているということです。

韓国の動きもここから外れていません。食品医薬品安全処と傘下の韓国動物代替試験法検証センターは、オルガノイド、臓器チップ、人体組織モデルのような非動物試験法を国際基準まで引き上げる作業を主導しており、2028年までに経済協力開発機構の国際公認試験法として登載するという目標を設定しています。国内でもオルガノイドサイエンスやネクストアンドバイオのようなスタートアップがオルガノイド基盤の新薬評価プラットフォームを事業化し、この流れと一緒に乗っています。

2026年韓国FDA規制科学会春季学術大会では、オルガノイド基盤の動物代替試験法を扱う特別セッションが別に設けられて、国内研究者と規制当局が共にこの基準を磨いています。

これらすべての進展にもかかわらず、まだ解けていない結び目がいくつか残っています。数千万個の変数で絡み合った人工知能モデルが特定の化合物を危険と判定しながらも、そのような判断に至った理由を人間が理解できる言葉で説明できない場合が多いです。答えは当てるけれど、解く過程は示さない学生と変わりません。規制当局が新薬許可を与える時に要求するものは、結果だけではなく、その結果に至る論理であるため、根拠を説明できない予測は、どんなに正確でも審査の閾値を越えるのが難しいです。

学習に使用されたデータと似ている化合物では高い精度を示していたモデルが、構造が完全に見知らぬ新規物質や異なる人種の生物学的データに出会うと、予測力が突然崩れる事態も繰り返されています。訓練過程で答えに当たる情報が誤って混入するデータリークのために、実際よりもはるかに正確に見える偽陽性結果が出るという指摘も、最近複数の研究から出ています。

技術の外の問題もあります。異なる研究機関や会社が作った仮想細胞モデルを受け渡し、また検証しようとするには、共通の言語が必要なのですが、SBMLやCellMLのような標準モデル技術言語が広く広がっていないため、データとモデルがあちこちに散らばったままになっています。仮想細胞を適切に学習させようとするには、数十万人の遺伝体情報と診療記録が必要なのですが、これはヨーロッパの個人情報保護法やアメリカの健康保険情報保護法と正面から衝突します。

原本データを機関の外へ持ち出さず、暗号化された結果だけをやり取りする連合学習のような代替案が議論されていますが、これを支える制度と法条項はまだ多くの国で整備されていません。

これまで蓄積されてきたオルガノイド細胞株と生物学的なデータは、特定の地域や人口集団に偏っているため、このデータで学習したモデルを遺伝的な背景が異なる患者にそのまま適用した場合、どの程度正確になるかは、まだ自信を持って言うことは困難です。人工知能と微小生理システムが共に下した判断が実際の患者に予想外の危害を加えた場合、その責任をソフトウェアを作った会社が負うべきなのか、データを提供した病院が負うべきなのか、最終的な決定を下した製薬会社が負うべきなのかを分ける法的な基準も、まだ設けられていません。

それにもかかわらず、その胃がん患者の腕に刺さった点滴の管は、すでに別の物語を始めてしまっています。マウスもサルも経ていない薬が、人の静脈にまで到達しました。以前であれば、動物実験の段階だけで数年が経過していたはずの道のりです。この薬が今後数年後にどれほど多くの患者を救うのかは、まだ誰にも確信できません。臨床試験はまだ最初の一步を踏み出したばかりです。ただ、実験室のネズミのケージの代わりに、コンピューターのサーバーと培養皿が新薬開発の最初の関門として浮上する流れだけは、すでに後戻りできないところにまで来ています。

韓国でも、この流れに合わせて法律を新しく作ろうとする試みが続いています。これまで Nam In-soon 議員と Han Jeong-ae 議員がそれぞれ動物代替試験法制定案を提出しましたが、農林畜産食品部と気候エネルギー環境部、食品医薬品安全処、農村振興庁の間の立場の違いにより、二つの法案はどちらも国会の敷居を越えられずに消えました。2025年12月、Song Ok-joo 議員がこの法案を再び代表発議しました。今回は関係省庁が事前に内容を調整した上で出した案であるため、Nam In-soon 議員や Han Jeong-ae 議員の時とは事情が異なるという評価が出ています。

法案には、政府省庁が共同で協議体を構成して基本計画を立て、分野別に動物代替試験法を検証するセンターを置き、5年ごとに実態を調査するという内容が盛り込まれました。法案が出たことで共に明らかになった数字もあります。2024年の

1年間に国内で使用された実験動物は459万匹に達し、このうち半分を超える実験が、麻酔薬や鎮痛剤なしで毒性物質に曝露させたり腫瘍を起こさせたりする、苦痛が極めて激しい最高等級に分類されました。

オルガノイドに残された隙間は、血管と免疫系だけではありません。中脳オルガノイドは、ドーパミンを作る神経細胞一つ一つの動きを人間の脳と似たように見せてくれますが、その神経細胞が他の脳領域にまで長い軸索を伸ばして信号をやり取りする回路には、まだ追いついていません。人間の脳は、皮質と海馬、大脳基底核が互いに信号をやり取りしながら一つの回路として動きますが、皿の上で単独で育ったオルガノイド一つには、このような長距離のつながりが欠けています。

最近では、複数のオルガノイドを物理的につなぎ合わせてアセンブロイドという構造体を作り、その間を行き来する神経信号を電極で測定する研究が出てきます。方向は合っていますが、まだ歩き始めの段階であり、記憶を保存したり感情を起こしたりするように、脳全体が一緒に動いてこそ現れる現象は、皿の上ではそもそも観察する道がありません。人間の脳レベルの複雑な回路をそっくりそのまま再現することは、依然として遠い目標として残っています。

動物実験を減らそうという要求と、人間の安全を確実に守ろうという要求は、今後も簡単には和解しないと思われます。TGN1412 がサルでは安全に見えたのに、人間の体で免疫系を暴走させた事件は、オルガノイドや臓器チップ一つだけではまだすべてを防ぎきれない危険があるという事実を同時に示しています。免疫系全体が互いに信号をやり取りしながら作り出す反応や、複数の臓器が同時に経験する慢性毒性は、今の非動物試験法が追いつくのが難しい領域として残っています。

このような危険まで自信を持って取り除ける証拠が積み重なるまでは、動物実験を法律で一度に禁止するよりも、危険が低い領域から一つずつ引き継いでいく方が、より現実的な道である可能性が高いです。

4.2

デジタルツイン(Digital Twins)と次世代臨床試験の現代化

アメリカの食品医薬品局の会議室のモニター上で、人間の心臓が赤く鼓動していました。本物の心臓ではなく、フランスに本社を置くソフトウェア企業ダッソー・システムズ(Dassault Systèmes)が作成したデジタル複製本、いわゆるリビング・ハート(Living Heart)でした。2014年に開始されたこのプロジェクトは、世界で初めて人間の心臓をそのまま模した3次元のデジタル心臓を完成させ、その後、食品医薬品局と5年間にわたる共同研究協約を締結して、心血管医療機器を検証する公式な仮想試験環境として採用されました。似た時期、書類の前にいた別の会議室では、人工知能スタートアップのUnlearn.aiが提出した提案書が審査官の手に握られていました。

新薬を服用しない患者群、つまり対照群を実際の人間の代わりにコンピューターが計算した仮想患者データで埋めるという内容でした。新薬開発企業が臨床試験設計に関して規制当局と事前に話し合う公式なC型会議で、この提案は受け入れられました。人間ではなくコンピューターが作成した患者が臨床試験の半分を埋めることができるという点を、規制当局が初めて認めた瞬間でした。

伝統的な臨床試験は今でも新薬承認の最後の関門ですが、その関門を通過する道はますます険しくなっています。世界中の臨床試験の80パーセント以上が患者を時期に充足できず、スケジュールが延びており、製薬業界全体が1年間に投じる研究開発費は2000億ドルを超えています。これほどの労力をかけても、最終承認に至る新薬は12パーセント未満です。新薬を受け取らず、プラセボ、つまり偽薬だけを受け取る対照群患者は、最良の治療機会を失います。数千人を数年間追跡しなければ統計的に意味のある結果が出ない無作為対照臨床試験の枠組みは、患者にも予算にも大きな負担です。

この負担を減らすという試みの中で、人工知能と膨大な実使用データを組み合わせることで臨床試験を再設計するデジタル・ツインと臨床試験シミュレーション技術が、

精密医療を完成させるエンジンとして浮上しています。患者1人を登録し追跡するのに要する時間と費用を減らすことは、もはや数社の選択ではなく、新薬開発企業全体の生存条件に近づいています。つまり、この変化の目的は新薬1つを市場により早く投入することに留まらず、その新薬が本当に必要な患者に本当に必要な方法で届くようにすることにあります。

新米のパイロットは、実際の空に上がる前に、まず飛行シミュレーター、つまりシミュレーターに座ります。シミュレーター画面の空は偽りですが、エンジン音と計器盤の揺れ、突風が押し付ける感覚は実際の飛行とほぼ同じに設計されているので、パイロットは本当の乗客を乗せる前に間違えても構わない練習を思いっきり行うことができます。デジタル・ツインは、1人1人に対してこのシミュレーターを作成してあげる技術です。患者の遺伝情報と病院の診療記録、手首に着けたスマートウォッチが計測する心拍数と睡眠、画像検査の結果、血液中のタンパク質数値までをコンピューター内に集めておくと、その人とほぼ同じように動作するデジタル複製本が作成されます。

医師と研究者は、検証されていない新薬を実際の患者に先に与える代わりに、このデジタル複製本に薬を先に与えてみます。用量を増やしたときにどのような反応が出るのか、副作用が肝臓で始まるのか心臓で始まるのかをコンピューター内で先に確認してからこそ、実際の処方に移行する順序です。シミュレーター内では数十回間違えても本当の乗客が傷つかないように、デジタル複製本の中で数百回試してからこそ、新薬は実際の患者に近づきます。

似た試みは以前にもありました。ただし、過去のモデルは特定の薬が人口集団全体でどのように平均的に吸収され排出されるかを計算する程度に留まり、患者個々の実時間データとは接続されていませんでした。静的な仮説を確認する計算機レベルに留まっていた。最近、ディープラーニングが入ってきて、この計算機は完全に別の次元に移行しました。コンピューター内で実行されるという意味でインシリコ(in silico)ツインと呼ばれる現在の技術は、ウェアラブルデバイスと電子健康記録、高解像度画像、遺伝子とタンパク質を一緒に見つめるマルチオミクスデータを24時間リアルタイムで受け取り、患者の体内で病がどのように進行するかを自ら学習し、継続的に更新されます。

センサー1つが1日に出すデータだけで数万件に達するため、かつては病院に行くときだけ一度計測していた指標を、今はほぼ毎瞬間見つめることができるようになりました。昨日のデジタル・ツインと今日のデジタル・ツインが異なる予測を出すことができるということは、それだけ実際の患者の身体状態をより密接に追いつけるということでもあります。

精密腫瘍学のように腫瘍ごとに性質が全く異なる分野では、特定の標的治療薬に反応する可能性が大きい患者だけを選び抜くことが臨床試験の成否を決めます。このような階層化がない時代には、薬が効く患者と効かない患者を一緒にまとめて試験するため、実際には効果がある新薬でも平均値に埋もれて無効判定を受ける場合が珍しくはありませんでした。人工知能はデジタル・ツインに積み重ねられた遺伝子と病理データを分析して、その薬に反応する可能性が大きいバイオマーカー、つまり体内の変化を示す特定の物質を持つ患者の部分集団を見つけ出します。

遺伝子変異1つ、タンパク質数値1つまで見分ける精密さのおかげで、10人中1人にしか効かない薬でも、その1人を正確に見つけ出して臨床試験を成功に導くことができます。その結果、臨床試験設計者は必要な患者数を大幅に減らしながらも、薬効を証明するのに必要な統計的確信をむしろより強固にすることができます。より少ない患者を集めながらも、より信頼できる結果を得るというパラドックスが、ここから生じるのです。

患者を選び出すこの作業の反対側には、対照群を作り出す作業が鏡のように向かい合っています。現実世界では人工知能が新薬に反応する患者を選別しますが、デジタル世界では同じ人工知能がその患者と統計的に双子である仮想の対照群を作り出します。材料は、過去に複数の病院で蓄積された臨床試験の記録や、世界中の疾患登録簿、電子健康記録など、すでに世の中に存在するリアルワールドデータ（RWD）とリアルワールドエビデンス（RWE）です。研究チームはこれらのデータをもとに、新薬をまったく使わなかった場合に病気がどのように進行するのかをコンピューター内で再現し、その結果を実際に新薬を投与された患者の経過と並べて比較します。

実際の患者データを人工知能がふるいにかけて治療群を立てる前半の作業と、リアルワールドデータを人工知能がふるいにかけて仮想の対照群を立てる後半の作業が、同じ原理を両側で繰り返す構造です。患者を選び出す目と対照群を作り出す目が結局は同じ人工知能であるという事実は、臨床試験を設計するたびに現実世界とデジタル世界を同時に計算しなければならない時代が開幕したことを意味しています。

デジタルツインが臨床試験に導入されたことで、製薬会社は実際の患者を募集するずっと前に、一つの臨床試験の成功確率をあらかじめ見極めることができるようになりました。開発中の新薬の作用機序と目標患者群のデジタルツインデータをコンピューター内で先に突き合わせてみて、失敗しそうな設計を臨床試験を始める前に除外する手続きは、一種の仮想リハーサルに似ています。第3相臨床試験が一つ失敗に終われば数年の時間と莫大な費用がそのまま消えてしまうため、このリハーサルが除外する失敗の一つ一つは、業界全体から見れば相当な資源を節約する効果につながります。

コンピューター内での成功が実際の患者での成功をそのまま保証するわけではありませんが、最初から可能性の低い設計を除外するだけでも、新薬開発の成功率を引き上げるのには十分に役立ちます。臨床試験を設計する段階から人工知能が関与するため、昔なら患者の募集が終わった後ようやく明らかになったような設計上の欠陥を、あらかじめ直すことができるという利点もあります。

Unlearn.AIが開発した予後共変量調整（PROCOVA、Prognostic Covariate Adjustment）手法は、デジタルツインが予測した個々の患者の予後スコアを統計モデルに乗せて、対照群の規模を減らす方法です。患者個人の予後を反映せずに平均だけで比較していた昔の統計方法とは異なり、この手法は患者ごとに異なる回復速度や危険要因を先に計算に入れた後、新薬の純粋な効果だけを見分けます。欧州医薬品庁（EMA）は第2相と第3相の臨床試験でこの手法を使用できるという資格認定の意見を出しており、人工知能基盤の方法論に対して欧州の規制当局がこのような形で支持したのは初めてのことでした。

米国食品医薬品局も数回のタイプCミーティングで同じ手法に肯定的な意見を出しました。Unlearn.AIが発表した数値では、治療効果の推定値の分散を10パーセ

ントから20パーセントほど減らすことができ、これは対照群に必要な患者数を20パーセントから30パーセント、多くてそれ以上減らす結果につながります。実際の臨床試験データとデジタルトゥインが作った仮想データを一緒に使うハイブリッド臨床試験は、もはや例外ではなく標準として定着しつつあります。

Dassault SystemesとUnlearn.AIの他にも、Nova In Silico、Deep Intelligent Pharmaのような会社が同じ市場に飛び込み、競争は一段と広がりました。

このように実際の臨床試験の半分と、デジタルトゥインが作った残りの半分を一つに合わせた設計を、業界ではツイン臨床試験（TwinRCT）と呼び始めました。名前の通り、半分は人から、半分はその人のデジタル双子から答えを求める方法です。Unlearn.AIはこの名前を自社技術の正式名称として規制当局に提出しており、このような名付けは、まだ馴染みのない概念を臨床現場の言葉に移そうとする試みでもあります。

ツイン臨床試験は現在、いくつかの希少疾患の研究を超えて、ありふれた慢性疾患の領域にも少しずつ広がっています。一つの臨床試験を設計する際に人とデジタル双子の比重をどれくらい分けるのかが、今後臨床試験の設計者たちが新たに直面する質問になると思われます。

希少疾患の分野では、この技術の使い道が一段と際立ちます。患者数自体が全国を合わせても数百人にとどまる疾患であれば、そもそもプラセボ群と治療群に半分ずつ分ける余裕がなく、そのように分けること自体が統計的にも無理な要求になります。このような疾患では、デジタルトゥインが作った合成対照群が、事実上唯一の現実的な代案として挙げられます。親が子供をプラセボ群に割り当てることに快く同意するのが難しい小児希少疾患の領域ではなおさらです。

様々な国の希少疾患の新薬制度は、このような事情を考慮して対照群の要件を柔軟に適用するケースが多くなっています。臨床試験の参加者全員に新薬だけを投与する単一群臨床試験に、デジタルトゥインで作った比較群を後から付ける方式も徐々に増える傾向にあります。業界の専門誌も、希少疾患の新薬開発がデジタルトゥインと仮想臨床試験を通じて新しい転機を迎えていると明らかにしています。

このような患者にプラセボ群を別に設けることは、倫理的に正当化されにくい場合が多いです。病気が急速に進行する患者に偽薬だけを与えて数年を待たせることは、新薬が効く可能性があるならなおさら受け入れがたい選択です。デジタルツインが作った合成対照群は、このような患者にプラセボの代わりに新薬を受ける機会を開きながらも、臨床試験が要求する比較データはそのまま確保できるようにしてくれます。このような問題意識自体は新しいものではありません。人工知能が登場する前からでも、過去の記録だけで作った歴史的対照群は信頼性の議論にしばしば巻き込まれており、デジタルツインはその古い論争をはるかに精巧な統計で補完しようとする試みに似ています。

だからといって、合成対照群をどこにでもそのまま使えるわけではありません。過去のデータから作った仮想の対照群は、現在の診療環境とすでに異なっている危険性があり、標準治療がここ数年変わっていたら、昔のデータを今日の比較基準にすることは難しいのです。そのため米国食品医薬品局は2023年に発表した指針で、外部対照群を案件ごとに別々に審査するという原則を明らかにしながら、ほかの治療法が適切でない重症疾患に客観的に測定される評価指標があるときに限ってその使用を限定しておきました。

適応型臨床試験は、臨床試験が進行している最中にも入ってくる患者の反応データを見て、投与用量や対象患者群をそのつど調整する方法です。強化学習アルゴリズムが中間解析結果を読み、特定の患者には用量を上げ、反応がないか副作用のリスクが大きい患者は臨床試験から早く除外するという具合です。固形がん患者を対象にしたCURATE.AIの予備臨床試験が、この方法を実際に示した事例です。消化管がん患者に使用される標準抗がん化学療法であるXELOXやXELIRI処方、その中に含まれる経口抗がん剤Capecitabineの用量を患者個人の反応データに合わせて継続的に調整しました。

患者1人をあたかも独立した臨床試験1つのように扱うこのような方法をN-of-1と呼ぶのですが、標準用量表に患者を合わせるのではなく患者に合った用量をリアルタイムで探していくという点で、既存の処方とは順序自体が異なります。この予備臨床試験の結果は、2025年に国際学術誌npj Precision Oncologyに掲載され、患者1人1人に合わせた用量調整が実際に可能であることを確認してくれました。

米国食品医薬品局が2025年1月に発表した指針草案は、人工知能モデル1つ1つを無条件に信頼したり無条件に疑ったりせず、そのモデルをどこにどのように使うかによって審査の強度を異なるようにする危険度ベースのアプローチを採用しました。デジタルツインが初期用量を決める参考資料としてのみ使用されるのであれば、比較的軽い検証でも十分ですが、プラセボを完全に代替する合成対照群として使用されるのであれば、はるかに厳格な検証を通過する必要があります。新薬開発企業は、モデルを提出する前に、このモデルがどのようなデータから、どのような構造で作られたのか、そして学習に使用されなかった新しい患者からも同じ性能を発揮するのかを事前に文書で整理し、規制当局と協議しなければなりません。

煩雑な手続きですが、逆に見れば、人工知能が作った証拠も人間が作った証拠と同じ目線で審査を受けられる道が開かれたということでもあります。このようなアプローチは、自動車や航空機部品を認証するときに部品ごとに安全等級を異なるようにする方式と大きく異ならないので、規制当局の立場からしても非常に慣れた枠組みではあります。

このような流れが数社の実験にとどまらないことは、あちこちで進行中の他の臨床試験によっても確認されます。米国クリーブランド・クリニックでは、2型糖尿病患者を対象に全身デジタルツインベースの精密治療の効果を検証する無作為対照臨床試験が進行中であり、結果は2026年中に出される予定です。2025年4月には、デジタルツインベース臨床研究システム自体の妥当性を検証する別の研究も新たに開始されました。市場調査機関は、デジタルツイン関連の世界市場規模が2024年の177億ドル水準から2032年には2,593億ドルまで、毎年40パーセント近い速度で成長すると予想しています。

フランスの製薬企業Sanofiをはじめとする大型製薬企業も、自社メディアを通じて、人工知能が作ったデジタル双子で臨床試験を再設計する作業を公開的に紹介し始めました。患者が自分の治療経過を直接確認するのを助けるOutcomes4Meのようなプラットフォームまで登場しながら、デジタルツインは研究者の道具を超えて患者自身が使う道具にも広がっています。

合成対照群が対照群に必要な人の足を減らしたのであれば、ウェアラブル機器は残された参加者の病院への足をさらに減らします。スマートウォッチとスマートパッチのような機器を身につければ、患者が病院に来なくても心拍と睡眠、活動量といったデジタルバイオマーカーが24時間にわたって蓄積されます。病院という物理的場所に縛られないこのような分散型臨床試験（DCT）は、2026年基準で世界市場規模だけで100億ドルを超えると推定されています。

米国政府の生物医学先端研究開発局（BARDA）が大型薬局チェーンWalgreensと手を組み、地元の薬局を臨床試験の拠点とする事業を行うのも、病院の敷居を低くしようとする同じ流れです。食品医薬品局も2024年4月に臨床試験革新センターを新たに開き、同年9月に分散型臨床試験ガイドラインを確定してこの流れに規制の支援を加えました。

こうして流れ込むデータは、大規模言語モデルが読み取ります。人工知能は構造化されていない診療記録と患者が残した文章をリアルタイムで読み、薬物有害反応や薬物間相互作用を90パーセントを超える感度で捉えます。インフルエンザワクチンの有害反応を報告書から自動的に抽出した研究チームもありました。Tree of Thoughtsの推論方法を適用した大規模言語モデルで、心血管系有害反応を診療記録から自動的に識別した研究チームもありました。

病院に行き来するのが難しい地方の患者や少数民族患者もこのやり方では臨床試験に参加でき、分散型臨床試験は効率だけでなく臨床データの公平性の問題も同時に解決しています。患者が自分の症状を直接入力する方法より、このような自動検出が優れている点は、患者自身も気に留めずに報告しない初期信号までも見逃さないということにあります。

2024年下半期、ダッソー・システムズは食品医薬品局とともに米国機械学会の検証標準である四十（V&V 40）を基準とした世界初のインシリコ臨床試験ガイドブックを発行しました。医師とエンジニア、政府関係者が共に力を合わせて始めたこのプロジェクトは、2025年2月に患者一人ひとり、または特定の患者群に合わせて細かく調整できる次世代リビング・ハートモデルのベータテストに入ったとダッソー・システムズが明らかにしました。リビング・ハートに続きリビング・ブレインまで登場した現在は、肝臓と肺を含む人体のほぼすべての臓器へと、同じ手

法が広がりつつあります。

リビング・ハートプロジェクトは、2025年に米国の経済誌ファスト・カンパニーが選んだ世界を変えるアイデアのリストにも名を連ねました。国際学術誌ネイチャー・メディシンは、2026年にデジタルツインとインシリコ臨床試験が新薬開発の転換点に入ったという論文を掲載し、臨床試験をデジタルツインで補強する研究は、同時期に学術誌npjシステムズ・バイオロジーにも発表されました。

華やかな成果の後ろには、重い障害物も残っています。最初に直面する問題は、人工知能が出す判断の中を覗き見ることができない、いわゆるブラックボックス問題です。数千万個の媒介変数で絡み合ったディープラーニングモデルが、特定の治療薬が効果的であると結論づけたり肝毒性が発生すると警告したりしても、その判断がどのような生物学的経路を経て出たのかを医師と規制当局に明確に説明できなければ、新薬許可の審査で核心的な根拠として採択されるのは困難です。

患者の生命がかかった決定を、コンピュータがなぜそのように判断したのか誰も説明できない状態で下させるわけにはいかないからです。医師たちの間では、コンピュータが出した結果にそのまま従うよりも、その理由を納得した後に処方箋に署名するという態度が依然として強いです。そのため最近では、最終の予測値とともに、その予測に大きく影響を与えた要因のいくつかを並べて見せる説明可能な人工知能技法が、現実的な代案として浮上しています。

学習データの限界も足を引っ張ります。特定の人種や特定の疾病パターンに偏ったデータで学習したモデルは、見慣れない新薬候補や異なる人口集団に出会うと予測の正確度が大きく落ち、学界にはその正確度が0.3の水準まで落ちた事例も報告されています。学習の過程で、後に出るべき臨床結果の情報がミスで入力データに混じって入り、モデルの性能指標を実際より膨らませるデータ漏れの問題も繰り返し指摘されています。このようなエラーは臨床試験の段階ではよく現れませんが、実際の処方現場で、存在しない危険があると誤って知らせる偽陽性の判定として現れる危険があります。

このような理由から、最近、規制当局はモデル学習に全く使われなかった新しい患者群で別途に性能を確認する、いわゆる標本外検証の結果まで新薬の提出資料

に含めるよう要求し始めました。そのため学界と製薬業界は、ソフトウェアの単位テストから生物学的反応との一致度の検証までを包括する厳格な検証手続きを、新薬開発の必須関門として要求しています。

精巧なデジタルツインを作るには、ゲノム情報や電子医務記録のように極度に敏感な個人情報を大量に集めなければなりません。米国の健康保険通商責任法（HIPAA）や欧州の一般個人情報保護法（GDPR）の下では、病院と製薬会社が各自保有したデータを外部サーバーに移して一箇所に集めること自体が事実上塞がれています。その代案として浮上したのが、患者の原本データは病院内部のサーバーにそのまま置き、暗号化したモデルの重みだけをやり取りする連合学習と差分プライバシー技術です。国境を越える多国籍の臨床試験では、ある国で合法であるデータ処理方式が、他の国ではそのまま違法になる場合も珍しくなく、国際共同研究を設計する法務担当者たちの負担はますます大きくなっています。

異なる病院のデータ形式を一つに合わせ、鉄壁のようなサイバーセキュリティ水準を維持するのにかかる費用も依然として侮れません。そのため大型病院は、互いの原本データを直接やり取りする代わりに、第3の信頼機関を経て匿名化の検証を終えたデータだけを研究に使う方式も一緒に試みています。

データの無欠性を要求するアルコアプラスプラス（ALCOA++）の原則、つまり、誰がいつそのデータを作ったのかを追跡する帰属性や正確性、完全性、一貫性といった基準を、人工知能アルゴリズムの中でどのように文書として残すかについての悩みが、製薬業界の中で深まっています。これに加え、人工知能が設計した仮想患者と合成対照群を根拠に承認された新薬が、市販の後に予想できない副作用を起こした時、その責任を、ソフトウェアを作った企業とデータを提供した病院、処方を下した医師、承認のハンコを押した規制当局の中で誰に問うべきかを整理した法的な基準は、まだどの国にもきちんと設けられていません。

今のところは、新薬開発会社が契約書の中に責任の所在をあらかじめ分けて明記する方式でこの法的な空白を埋めているだけであり、裁判所の判例や別途の立法で明確な基準が立てられるまでは、このような不確実性が当分の間続く可能性が大きいです。多くの国の政府が関連する立法を検討していますが、法案の論議段階を越えて実際の施行に入った国はまだ見つけるのが難しいです。

米国議会は2022年と2023年にかけて食品医薬品局近代化法2.0（FDA Modernization Act 2.0）を通過させ、数十年間、新薬許可の前提条件であった動物実験の義務条項をなくし、人工知能ベースのインシリコ・シミュレーションの資料を安全性と有効性を立証する代替手段として公式に認めました。2025年1月には、食品医薬品局が医薬品と生物学的製剤の規制判断に人工知能を使う時に守るべき指針の草案を出し、モデルの用途をあらかじめ定義し、データ管理とモデル構造、学習に使わなかったデータでの性能まで検証するように要求しました。

同年12月には、実使用根拠に基づく医療機器審査ガイドラインを最終確定し、個別患者を特定できる原データなくても、匿名化された大規模医療データベースを根拠資料として受け入れることを明らかにしました。欧州医薬品庁は2024年9月、医薬品のライフサイクル全般にわたる人工知能の使用に関する考察文書を採用し、モデルの構造を透明に説明すること、および結果について人間が最後まで責任を持って監督することを製薬業界に勧告しました。2025年7月には、外部対照群に関する別途のガイドライン初案を公開意見募集に付しました。英国医薬品・医療製品規制庁(MHRA)も2025年5月、実使用データに基づく外部対照群ガイドライン初案を発表しました。

韓国食品医薬品安全庁もやはり2025年12月、デジタル機器で収集した臨床試験資料の認定基準を含むガイドラインを制定しました。2026年1月には、食品医薬品局と欧州医薬品庁が人工知能を新薬開発に用いる10の共通原則を共に発表する一方で、大西洋両側の基準を合わせ始めたのですが、これは製薬業界にすでに馴染みのある優良臨床試験管理基準、いわゆるGCP伝統に人工知能バージョンをもう一つ加えたものに近いです。

韓国内ではまだ行く道が遠いという評価が多いです。医療人工知能専門家は2026年にも国内の医療現場でデジタルツインを使用する事例が不足しており、信頼を支援する規制体系から整える必要があると指摘しています。米国と欧州では、規制当局と企業が5年、10年ずつデータを積み重ねてきた末に、認定と指針初案を発表しました。このギャップは予算や意志の問題というより、時間がかかる信頼蓄積の問題に近いです。ただし、その時間を単に流し去ることはできません。

データを十分に積み重ね、検証手順を備えた病院と企業が先に規制当局の信頼を得て、その経験が次の事例の審査基準となる方法で、ギャップが縮まる可能性が大きいです。デジタルツインと合成対照群が臨床試験の半分を代替するハイブリッドモデルは、今、規制当局の公式文書に入り込んでいます。ただし、すべての新薬開発のデフォルト値になるまでは、説明可能性とデータ公平性、責任所在を一つひとつ埋める時間がさらに必要に見えます。

春川のスタートアップ Optrimed を率いるCEO Hong Su-ji は、外国のデジタルツイン企業の仮想患者モデルがほとんど白人データで作られたという点を自社技術の差別化要因として挙げています。Optrimed が作成した人工知能ベースの人間デジタルツイン、OPTIVIS は、韓国人を含むアジア人患者集団のデータを別途学習させて、世界で稀にアジア人仮想患者を作成します。Optrimed は2026年7月、英国ロンドンで開催されたアルツハイマー協会国際学術大会(AAIC)でこの臨床試験デジタルツインプラットフォームを初めて公開しました。

新薬1つを開発するのに平均20億ドルと10年から15年がかかる今、Optrimed はOPTIVISを使用すればその時間と費用を約30パーセント削減できると明らかにしました。Optrimed は春川企業革新パークに研究拠点を新たに建設することにし、業務協定を締結しました。西洋人中心に積み重ねられたデータがアジア人患者によく適合しない可能性があるという懸念は、デジタルツインが世界に広がるにつれて、韓国を含む複数の国で等しく続く問題に見えます。

実際の臨床試験結果と実使用データで模倣した結果はいつも同じではありません。ハーバード医科大学附属ブリガム・アンド・ウィメンズ・ホスピタルのJessica Franklin と同僚は、すでに終わったランダム化比較試験10件を実使用データのみで再度再現してみました。規制判断と同じ結論に至った場合は10件のうち6件に止まり、ハザード比推定値が元の臨床試験の信頼区間内に入った場合は8件でした。

この結果は2021年のアメリカの学術誌、Circulation に掲載されました。臨床試験設計を元の研究とどの程度近く再現したかによって、2つの結果の間の正確度は大きく異なりました。デジタルツインと合成対照群をランダム化比較試験の代わりに広く使用する前に、どのような条件下で結果がずれるかから一つひとつ取り除く方が安全に見えます。

デジタルツインが切実に必要な患者であるほど、その患者のデータはむしろ少ないです。希少疾患患者と少数民族患者はこれまで臨床試験に少数参加してきており、そのため今、デジタルツインを学習させるデータの中にも少なく入っています。学界では少数の患者を統計的に補正するか、人工知能で新たにデータを作成してギャップを埋める方法を試みていますが、元のデータ自体が不足している限界を完全には埋めることはできません。昔の臨床試験から外された患者が新しい技術の中で再び外されるリスクは、新しいデータを積み重ねない限り自動的に消えないように見えます。

4.3

精密腫瘍学(Precision Oncology)と患者に合わせた治療

イスラエルのシェバ・メディカルセンターやアメリカのMDアンダーソンがんセンターをはじめとする国際WINコンソーシアムが進めたWINTHER臨床試験は、進行がん患者を300人以上登録し、腫瘍と正常組織を共に採取して遺伝子配列と遺伝子発現情報を同時に読み取った、大規模な多国籍研究でした。多発性骨髄腫と乳がんの患者にチロシンキナーゼ阻害剤という標的抗がん剤を投与しながら、研究陣は患者のがん細胞から採取した単一細胞一つ一つの遺伝子活動をコンピュータで分析しました。患者たちは二つの方式に分けられ、一方は236個の遺伝子を読み取るDNA塩基配列分析で、もう一方は腫瘍と正常組織の遺伝子発現を比較するRNA分析で治療薬が割り当てられました。

特定の遺伝子一つだけを狙って薬を選んでいた以前の方式は、すでに限界にぶつかった後でした。患者ごとに異なる遺伝情報と病理所見を一つにまとめ、その人だけに合う治療法を見つけ出す個人テラーメイド医療が、がん治療の中心へと移ってきていました。WINTHER試験では、遺伝子とタンパク質と代謝情報を一度に計算して最適の治療薬を見つけてくれるマルチオミクスマッチングアルゴリズムが、患者と薬を結びつける役割を担いました。腫瘍はじっと止まってはいません。人の体の中で時々刻々と姿を変える腫瘍を追いかけること、それがこの臨床試験の投げかけた問いでした。

がん治療薬の開発は、長らく一つの信念の上に立っていました。特定の遺伝子一つが故障してがんができたのなら、その故障した部位一つだけを正確に塞ぐ薬を作れば病気は治るという信念でした。この信念は、特定の遺伝子変異を標的にしたいいくつかの薬が実際に劇的な効果を上げた後、しばらく力を得ました。しかし時間が経つにつれて、同じ遺伝子変異を持つ患者でも反応が分かれ、最初はよく効いていた薬も数ヶ月、数年が過ぎると力を失う事例が積み重なりました。

問題は薬ではなく、がんを見つめる枠組み自体にありました。腫瘍は一枚の写真ではなく、動き続ける映像に近く、患者の遺伝情報や生活習慣、免疫状態までを網羅するはるかに大きな絵を一緒に読み取らなければならないという認識が広まり、個人テーラーメイド医療ががん治療戦略の中心へと移ってきました。

同じ乳がんでも人によって異なる理由は、あるクラスの生徒たちを思い浮かべれば理解しやすいです。同じ学年の同じクラスに座っていても、身長や性格、得意な科目がそれぞれ違うように、乳がんという名札一つの中にも、互いに異なる遺伝子変異を持った患者たちが混ざっています。腫瘍も違いはありません。見かけは一つのこぶのように見えても、その中を覗き込むと、それぞれ少しずつ異なる遺伝子変異を持ったがん細胞の群れが入り混じっており、これを腫瘍内不均一性と呼びます。同じ教室の中でも窓側の生徒と後ろの生徒が互いに異なる習慣を持つようになるように、腫瘍の中でも位置や周辺環境によって、細胞ごとに育つ速度や薬に反応する方式が変わってきます。

この違いを無視して腫瘍全体を一つの塊としてだけ扱うと、最初はよく効くように見えた薬が時間が経つにつれて力を失う理由を説明できません。そのため、同じ乳がんの診断を受けた二人であっても、一方にはよく効く治療薬がもう一方には全く効かないということが、実際の診察室ではよく起こります。

がん細胞が薬を避けて逃げる過程は、自然選択の縮図です。抗がん剤を投与すると大部分のがん細胞は死にますが、偶然その薬に耐えられる遺伝子変異を持った少数の細胞は死なずに生き残ります。生き残った細胞は再び分裂して増えながら、いつの間にか腫瘍全体をその変異で満たします。耐性はこのようにして生まれます。最初はよく効いていた薬が、もはや効かなくなる瞬間です。特定のがん誘発突然変異一つが登場すると、その余波で細胞内のシグナル伝達経路が再構成され、遺伝子を読み込む転写環境が変わり、エネルギーを作り出す代謝の流れまでが組み直されます。

このような連鎖反応は、薬一つで分子一つだけを防いだのでは鎮めることができません。最近、研究陣はリンチ症候群の患者を対象とした単一細胞空間免疫プロファイリング技術を動員し、腫瘍の中の細胞一つ一つの位置や遺伝子活動を細胞レベルで覗き込んでいます。間質細胞と免疫細胞とがん細胞の間の物理的な距離、

そしてこれらがやり取りする非線形的な相互作用を深層学習モデルで分析すれば、腫瘍がどの方向へ進化していくのか、その軌跡をあらかじめ描いてみるができます。

研究者たちの間では、人工知能を腫瘍の進化自体を推論する一つの統合エンジンと見る見方も力を得ています。異なる時点で得られた資料を一行に並べ、がん細胞の群れの系譜を再構成し、今後耐性が生じる方向まで見極める役割を、人工知能に任せるアプローチです。

このすべての予測は、原材料となる資料の解像度に頼っています。以前は腫瘍組織の一塊をすりつぶして、その中にある数多くの細胞の遺伝子シグナルを平均値一つにまとめて読み取っていましたが、単一細胞シーケンシングと空間トランスクリプトーム技術が普及するにつれて、今では細胞一つ一つの遺伝子活動と、その細胞が組織の中で正確にどの位置にあったのかまで一緒に記録できるようになりました。クラスの平均点だけを見ていたのが、生徒一人一人の名前と成績を確認できるようになったわけです。このように緻密になった資料が積み重なってこそ、人工知能も腫瘍の中の細胞の群れの微細な違いを見逃さずに捉えることができます。

このように把握された腫瘍の複雑さは、直ちに患者別の薬物反応予測へとつながります。ここでマルチオミクスとは、遺伝子配列であるゲノム、遺伝子がどれほど活発に読み取られているかを示すトランスクリプトーム、そして実際に作られたタンパク質まで、細胞内の様々な層を一度に覗き込むという意味です。ディープラーニング基盤のモデルであるCDRscanは、複数のがん細胞株で確認された抗がん剤反応資料と、その細胞株のゲノム変異情報をペアにして学習したニューラルネットワークを土台とし、細胞実験の段階で積み上げられた知識を、患者個人の予測へと移してきます。

がん遺伝体シグナルだけを見て、特定の抗がん剤がその患者に実際にどの程度効果があるかを高い精度であらかじめ計算して出す方式です。PERCEPTIONという精密腫瘍学パイプラインは、大規模細胞株薬物スクリーニングから得られた大量の発現データと単一細胞発現データを共に学習して、多発性骨髄腫と乳がんと肺がん患者を対象とした実際の臨床試験において、標的治療薬に対する反応と耐性

発生の有無を予測することに成功しました。2つのモデル共に、患者の組織を検査室に入れる前に、コンピュータの中で先に治療結果を推し量ってみるというわけです。

最近では、複数の細胞株で先に学習してから、その知識を実際の患者データに移して適用する転移学習方式も薬物耐性の分子的原因を明らかにするのに使われています。複数オミクスデータを統合して免疫チェックポイント阻害剤によく反応する患者群をあらかじめ選び出す作業も同じ原理で行われており、合わない薬を投与して数ヶ月を無駄にした後になって効果がないという事実を確認するかつての試行錯誤を大きく減らしてくれます。

薬物耐性に打ち勝とうとする試みは、抵抗の経路自体を見つけ出すことにまで進みました。人工知能が実際の患者の複数オミクス臨床データを分析した結果、CDK4/6阻害剤に対する耐性がエストロゲン受容体に依存する経路とそうでない独立的な経路に分かれるという事実が明らかになりました。CDK4/6阻害剤はホルモン受容体陽性乳がん患者たちに広く使われている標的治療薬系列なので、この耐性経路を明らかにした成果は実際の診療現場に及ぼす影響が小さくありません。

VISTAは免疫細胞ががん細胞を敵として認識できないようにする蛋白質であり、STINGは体内に侵入した異物を感知して免疫反応をオンにする警報装置のような信号経路です。この発見に基づいて、研究チームはVISTAを阻害したりSTINGを刺激する併用療法を設計して、1つの薬だけでは塞ぐことができなかった耐性の隙間を埋めています。複数の標的を同時に狙うこのような併用療法は、料理に例えると理解しやすいです。材料1つだけで味を出そうとするのではなく、複数の材料の配合比率までを計算して一緒に入れる料理人のように、人工知能は薬の組み合わせと用量比率を一緒に計算して出します。

化学化合物同士のシナジー効果を予測するディープラーニング技法も同じ目的で使われており、がん細胞が片方の経路を遮断されると別の経路に迂回して生き残ろうとするため、最初から複数の標的を同時に狙ってその迂回経路そのものを塞ぐ戦略を立てるようになります。

標的を1つではなく複数に増やしたときに成績がどのように変わるのかを示した研究もあります。アメリカのカリフォルニア大学サンディエゴキャンパス研究チームが主導したI-PREDICT臨床試験は、患者の腫瘍で見つかった複数の遺伝子変異のうちどの程度の比率を薬と一緒に狙ったのかを基準として、その一致度が高いほど無増悪生存期間と全生存期間が共に長くなったという結果を出しました。1つの変異だけを狙っては腫瘍はすぐに別の経路に迂回して逃げ出しますが、複数の標的を同時に押さえれば逃げ場がそれだけ減るというわけです。

病理学ファウンデーションモデルが組織画像から抽出した特徴と遺伝体分析結果を臨床意思決定支援システムが一度に集めて調べながら、医療者が複数標的の組み合わせを設計してその根拠を確認する速度と正確度が顕著に上がっています。研究チームはこの結果に基づいて、がん遺伝体検査を最初に実施する段階から複数の標的を共に考慮する設計を標準診療手順に含めるべきだと提案しました。

顕微鏡の前に座って組織スライドを1枚ずつ目で読み下していた病理専門医の作業方式も変わっています。ファウンデーションモデルという言葉は建物の基礎工事に例えられます。数百万枚の組織病理画像と複数の形態の臨床データで広い基礎知識をまず固めてから、その上に胃がん診断や乳がん予後予測のような具体的な任務を1つずつ積み重ねて訓練させる方式だからです。このようにして作られた大型病理学ファウンデーションモデルが、がん診断と予後予測に本格的に使われ始めました。

これらのモデルは細胞の微細な形状と細胞が並んでいる空間的配列をパトホミクスという数値データに変えて、これを遺伝体情報や診療記録と合わせて進行性胃がん患者が手術前の抗がん化学療法や免疫治療にどのように反応するかをコンピュータの中で先にシミュレーションします。フィンランドのAiforiaは胃がんと乳がんと前立腺がんを診断する人工知能ソフトウェアでヨーロッパ体外診断医療機器認証(CE-IVD)を相次いで受けました。オランダのラドバウド大学病院から分社したAiosynも乳がん細胞の有糸分裂を数えるソフトウェアで同じ認証を受けて、スライド1枚を読む時間を最大60パーセント減らしてくれています。

国内ではルニットが開発した病理分析プラットフォーム「ルニット スコープ」が世界中の上位製薬会社20社のうち15社と手を握って臨床試験とバイオマーカー開

発に参加しており、2025年には米国診断企業ラボコープとデジタル病理分野のパートナーシップを結んで国際学会で共同研究成果を発表しました。同じ年ヨーロッパ腫瘍学会ではルニット スコープを活用して免疫抗がん薬治療反応を予測した研究3件が一緒に公開されて、病理画像分析結果が患者の実際の生存期間と繋がるという根拠を補強しました。

神経腫瘍学分野では、脳腫瘍患者の治療反応を標準化された方式で評価するためのAI-RANO指針が国際研究者協議体を通じて準備されて、人工知能ベースの画像判読が臨床試験で信頼できる根拠を示しています。

アメリカの精密医療企業Tempusはこのような流れの規模をよく示す事例です。ナスダックに上場されたこの会社は分子情報と画像と臨床記録と一緒に含まれた500ペタバイトを超えるデータを蓄積し、その中には4500万人を超える患者の診療履歴と150万人の遺伝子配列情報、そして遺伝体と転写体と画像と臨床情報をすべて備えた40万件を超えるがん患者記録が入っています。Tempusが出したNexusというプラットフォームは過去のデータを後で振り返る分析ではなく、診療が行われているまさにその瞬間に必要な情報を医師に渡すリアルタイム臨床知能を標榜しています。

2026年の米国臨床腫瘍学会で発表されたある研究は、6つの米国の医療機関、患者662名を対象にこのプラットフォームを導入した後、初期の非小細胞肺癌患者のバイオマーカー検査の割合が明らかに上がったという結果を出しました。同年7月には、微小残存病変の検査を専門とするPersonalisの買収を発表し、組織に基づく精密腫瘍学と血液に基づく残存病変の追跡を1つの会社の中でまとめる作業に乗り出しました。

Tempusは、自社が構築したファウンデーションモデルの研究結果の一部を外部の研究者に開放するPRECISION Challengeという国家単位のプログラムも開始し、1つの会社が独占していたデータとモデルを学界と共有する方向へと歩みを進めています。

膨大なマルチオミクスデータと予測アルゴリズムが出した結果は、最終的に担当医が処方箋を出す瞬間に1つに集まります。メモリアル・スローン・ケタリングが

んセンター（MSKCC）の知識と自然言語処理技術を組み合わせて米国包括的がんネットワーク（NCCN）のガイドラインを学習したIBMワトソン・フォー・オンコロジーは、電子カルテから患者の情報を抽出した後、強く推奨、考慮可能、推奨しないという3段階で治療の選択肢を医師に提示する代表的なコグニティブ・コンピューティング・システムです。KM2PSというフレームワークは、抗がん剤を処方する前に、肝機能と腎機能はもちろん、薬と遺伝子、薬と薬、薬と健康食品、薬と生活習慣の間で起こり得る相互作用をそれぞれ点検することで副作用を減らします。

OncoSOLVERのように臨床試験の適合性を判別する人工知能は、患者に合う臨床試験をリアルタイムで探し出し、複数の診療科の医師が集まる分子腫瘍委員会の議論を支援する仮想プラットフォームへも発展しました。米国食品医薬品局（FDA）は、人間と人工知能と一緒に事例を検討するxDECIDEという意思決定支援システムを通じて、仮想腫瘍委員会を運営しています。複数の医師が集まっていた従来の腫瘍委員会に、人工知能という出席者がもう1人増えたようなものです。

これに加え、FDAは次世代シーケンシング（NGS）で見つけ出した突然変異と、その突然変異を狙って臨床試験中である薬を、電子カルテやがんの知識ベース、臨床試験の登録情報から自動で繋げてくれるツールも開発しています。2026年には、複数の人工知能モデルを1つの指揮体系の下にまとめ、腫瘍学全般の意思決定を支援しようとするオーケストレーション・フレームワークの研究も登場し、異なるモデルが1人の患者について出す判断を調整する段階へと進んでいます。

ある地域のがん診療ネットワークでは、人工知能が提案した治療の決定が、大規模ながんセンターの腫瘍委員会の勧告と一致する割合が、導入前の60パーセントから導入後には94パーセントに上がったという結果も出ました。

治療の強さを決めることも、もはや標準の用量表1枚には依存しません。

CURATE.AIというプラットフォームは、固形がんの患者一人ひとりに必要な薬の用量を時間に応じて再計算し、患者の反応を見守りながら用量を少しずつ調整する適応型投与方式の妥当性を、実際の臨床で立証しました。昨日と今日の体重が同じでも、昨日飲んだ薬の効果と副作用が違って現れたなら、今日飲む用量も変わらなければならないという考えです。このような方式は、効果と毒性の間の隙

間が狭く、1つの誤差が大きな結果に繋がる小児がんの治療や、移植後の免疫抑制剤の調節において使い道が大きいです。

ナノ医学の分野では、複数の階層の生物学的な仕組みを網羅するモデルに人工知能を組み合わせた仮想腫瘍モデルが、ナノ粒子が血管の壁を通過して腫瘍の組織に溜まる透過性滞留効果（EPR）を数値で計算し、薬をナノ粒子に乗せて送る時にその効果が患者ごとにどう変わるかをあらかじめシミュレーションします。この計算が精巧になるほど、同じ薬でも人によって必要なナノ粒子の設計と投与経路を別々に定めることができる根拠が積み重なります。

手術で目に見える腫瘍を無くした後にも、体の中のどこかに残っているかもしれないがん細胞を探し出すことは、長い間精密腫瘍学の宿題でした。がん細胞は死ぬ時に自分のDNAの欠片を血液の中に流し出しますが、これを循環腫瘍DNA（ctDNA）と呼びます。一滴の血液が手がかかります。画像検査ではまだ何も見えない時点でも、血液中のctDNAの濃度を追跡すれば再発するかどうかを見極めることができ、このように手術の後に体の中に残ったがんの痕跡を微小残存病変（MRD）と呼びます。

2026年に入って学界に報告されたMRDsteerという人工知能システムは、ctDNAの分析過程で変異を読み取る際の信頼度をリアルタイムで監視し、信頼度が一定の基準より下に落ちると、危険度の高いゲノムの区間だけを選んで再び読み取る自律ループを回します。極めて低い濃度の変異まで捉えるにあたり、この方式が従来の方法より安定的であるという点が、検証データを通じて確認されました。乳がん患者を対象にctDNAを反復して追跡するGEMINI Breastという臨床試験も、2025年10月から患者を募集しており、胃腸管がんと肺がんを網羅する複数の研究において、手術や治療の後にMRD陽性の判定を受けた患者ほど再発の危険が高く、生存期間が短いという結果が繰り返し確認されています。

医療陣はこの情報を根拠に、再発の危険が高い患者には抗がん治療の期間を延ばし、危険が低い患者には不必要な治療を減らしてあげる方向で決定を下すことができます。

がん細胞は育ちながら、元々の体にはなかった変形したタンパク質の欠片を表面に表しますが、これを新生抗原と呼びます。新生抗原は、患者ごとに腫瘍の突然変異が違っただけに人によって全く異なる形をしており、まるでそれぞれの顔のようにその人の腫瘍にだけある目印になります。人工知能は遺伝子配列のデータを分析し、この目印の中から体の中の免疫細胞が見分けやすいものを選び出し、これを根拠に患者一人のためだけのワクチンを設計します。ビオンテックとジェネンテックが一緒に開発したautogene cevumeranは、膵臓がんの手術を受けた患者にこのような個人に合わせたワクチンを投与した結果、ワクチンに反応してT細胞が作られた患者は再発なしで過ごした期間がかなり長く続いた半面、反応しなかった患者は中央値13.4ヶ月で再発しました。

3年後の追跡調査では、反応した患者のT細胞が平均7.7年という長い寿命を維持していることも確認されて、現在は260人規模の3相臨床試験で範囲を広げて化学療法薬と免疫治療薬と一緒に投与する方法を検証しています。ModernaとMerkが黒色腫患者を対象に開発中のmRNA-4157(V940)も、免疫チェックポイント阻害剤のペンブロリズマブと一緒に投与した時に再発または死亡リスクを約49パーセント低下させたという臨床結果を発表した後、3相試験に入り、早ければ2026年末から2027年の間に最初の規制承認の可否が決まる見込みです。

腫瘍遺伝子配列を確認した後にこの個人カスタマイズワクチンを実際に製造するまでは、通常3週から4週かかります。個人カスタマイズワクチン1人分の価格は10万ドルから30万ドルの間と推定されるが、再発リスクが高い手術後補助療法段階に投入すれば、その費用を治療効果で正当化できるという分析も出ています。免疫チェックポイント阻害剤と個人カスタマイズワクチンを一緒に使う臨床試験は黒色腫を中心に150件を超えました。BioNTechが開発中のもう一つの個人カスタマイズワクチンBNT122も大腸がんを含む複数の固形がんで免疫チェックポイント阻害剤と一緒に3相試験段階に進んでいて、個人カスタマイズワクチン競争は1社や1つの癌種にとどまっていません。

ゲノムと転写体情報、病理画像と放射線画像、診療記録と血液中のctDNA、生活習慣データまで、異なる形式のデータが一箇所に集まれば、人工知能はこれを統合して患者を詳細なタイプに分け、その結果は再び個人カスタマイズ医学と各患

者に合った投与経路決定につながります。患者1名の治療方向を決めるために病理科と腫瘍内科と遺伝子解析チームが一箇所に集まることも今は珍しくありません。国内でもこのトレンドはもはや研究室の中だけにとどまりません。国立がんセンターは国際共同臨床研究に参加し精密腫瘍学データを蓄積しており、ソウル大学校病院を始めとした主要大学病院ではキメラ抗原受容体T細胞(CAR-T)治療と液体生検を融合させた臨床試験が進行しています。

標的治療、mRNAベースのカスタマイズワクチン、細胞治療薬、液体生検という4つのトレンドが2026年現在1つの病院の中で同時に回転し始め、治療の重心は病名1つで患者を束ねていた方法から患者1人1人を基準にする方法に移行しています。患者を呼ぶ名前も少しずつ変わっています。肺がん患者または乳がん患者という病名の代わりにEGFR陽性肺がんのように特定の変異を持つ患者群として、さらにその患者1人の腫瘍指紋として呼ぶ言葉が診療室の中に入ってきています。

しかし画面内の数式がいくら精巧でも、臨床現場の検証基準ははるかに冷徹です。WINTHER臨床試験がその現実を示す事例です。全患者の35パーセントにオミクスデータを根拠にカスタマイズ薬物をマッチングさせることに成功したが、マッチングされた治療の無進行生存期間が以前の治療より1.5倍以上長いはずという事前に指定された目標を満たした患者は22.4パーセントにとどまりました。ただし全身状態が良好でマッチングスコアが高かった患者たちだけ別に見ると、全生存期間中央値が25.8カ月で、そうでない患者たちの4.5カ月と比較して明らかな差を示しました。

格差は自動的に消えません。学習に使われたデータの中では優れた性能を示していた深層ニューラルネットワークモデルが見知らぬ病院の新しい患者データに出会えば性能が崩れることが一般的であるが、これを幻想的な一般化問題またはドメイン移動と呼びます。モデルを訓練する過程で後で知るべき治療結果情報が誤って入力値に混ざって入り性能を膨らませるデータ漏洩、そして本当の因果関係の代わりに表に出た相関関係だけを学んでしまうショートカット学習も実際の医療事故につながる可能性がある盲点として挙げられます。

そのため最近の研究者たちは、1つの病院で作られたモデルを異なる病院、異なる人種、異なる機器環境のデータで再度検証する外部検証手順を論文査読の必須条

件として要求し始めました。

ブラックボックス問題に対応しようとする努力も増えています。病理画像を分析するモデルは最終判定を下しながら、組織画像の中のどの部位を根拠として使用したかをヒートマップで表示し、病理専門医がその部位を直接目で確認できるようにして、ゲノムデータを扱うモデルは多くの変数の中のどの遺伝子とどの相互作用が予測結果に大きく影響したかを数値でさかのぼって示します。このような説明技術が完全な因果証明ではないが、医者が人工知能の推奨をそのまま従わず、その根拠を自分で判断する余地を広げるという点で規制当局の信頼を得る橋的役割を果たしています。

制度と技術の間隔は法廷での争いに発展する可能性も抱えています。人工知能が設計した仮想患者群データが盲検ランダム化臨床試験の対照群に代わることができるかどうかについて、米国食品医薬品局と欧州医薬品庁はリスクレベルに応じて事前に変更計画を決めておき透明に公開するよう指示するガイドラインを相次いで発表しています。単一細胞多重オミクスデータと膨大な診療記録を集め学習する過程は欧州一般個人情報保護法(GDPR)や米国保健保険携行可能性・説明責任法(HIPAA)とあちこちで衝突し、この問題はまだ明確な解決策を見つけていません。特定の人種または先進国患者のゲノムにデータが偏っていれば、そのアルゴリズムは疎外された集団に誤った診断または不適切な薬を勧める危険を抱えており、医療格差をむしろ拡大する可能性があります。

人工知能がなぜその抗がん剤を勧めたのか、なぜ特定の副作用を警告したのかを、医師に分子レベルの根拠で説明できないブラックボックス問題は、規制承認を遅らせる障害として依然として残っています。人工知能の臨床意思決定支援システムのエラーで患者が怪我をしたり命を落としたとき、その責任がアルゴリズムを作った会社にあるのか、データを提供した病院にあるのか、処方を下した医師にあるのかを問う法的条項は、まだどの国にも適切に設けられていません。保険給付の仕組みもこの変化に追いついていません。標準治療一つに値段をつけていた方式では、患者ごとに異なる併用療法や個人最適化ワクチンの費用をどのように分けて認めるか、答えを出すのが難しいからです。

結局、2025年と2026年の精密腫瘍学は、計算速度や予測の正確さを高める競争を越えて、データの公平性と判断根拠の透明性、そして規制当局と歩調を合わせた法的な安全網を共に備えてこそ、次の段階に進むことができます。決定を最終的に下す主体が人間である医師だという事実は、このすべての技術的進歩の中でも変わっていません。

Hangukでも精密腫瘍学はもはや馴染みのない言葉ではありません。Goryeo Daehakgyoが率いるK-MASTER事業団は、2017年6月から全国55の病院と手を結び、がん患者の遺伝子パネル検査の結果を進行中の臨床試験につなぐマッチマスターシステムを運営してきました。担当医師はこのシステムを通じて、患者に合う臨床試験があるか、どのような分子標的薬を考慮できるかを確認します。

次世代シーケンシング（NGS）基盤の遺伝子パネル検査も、2017年3月に10種類のがんを対象として健康保険が初めて適用され、2019年5月にはすべてのがん種へとその範囲が広がりました。ただし、検査費の半分は依然として患者自身が支払わなければならない、給付が認められる検査も、初めて診断を受けたときと、がんが再発したり別の場所に転移したときの、この2回に制限されています。

遺伝子変異を見つけ出したからといって、すぐに治療につながるわけではありません。KRAS変異を持つ肺がんのように、原因となる遺伝子は突き止めたものの、その遺伝子を防ぐ薬はまだ出ていないがん種もあります。薬があってもハードルは残ります。尿路上皮がんを使うPadcevとKeytrudaを共に投与すれば、患者3人のうち1人の割合でがんが完全に消えるという結果が出ましたが、この併用療法は国内でまだ健康保険の給付を受けられず、患者が薬代をすべて背負わなければなりません。

新しく出た抗がん新薬が、出てから1年以内にHangukに入ってくる割合は100件のうち5件にとどまり、経済協力開発機構（OECD）の平均である100件のうち18件にも及びません。肺がんや乳がんを使う分子標的薬は、1ヶ月分の値段が数百万ウォンを超えることも珍しくありません。標的を見つけ出す技術は速くなりましたが、その標的に合う薬を手に入れ、値段を負担することは、依然として別次元の戦いとして残っています。

この格差は病院の大きさとも噛み合います。2023年に新しくがんの診断を受けた患者のうち、首都圏に住む人は半分になりませんでした。実際に首都圏の病院で治療を受けたがん患者は10人中8人に近かったです。地方に住みながらSeoulのいわゆるBig 5大型病院まで訪ねていくがん患者の数も、2022年と比べて2024年に11.8パーセント増えました。最新のゲノム検査と人工知能基盤の分析は、たいていSeoulのいくつかの大型病院に先に備えられ、その設備と人材が地域の病院まで均等に広がるのには時間がかかります。

精密腫瘍学が出す恩恵は、その病院まで訪ねていく余裕がある患者に先に回る可能性が高いです。技術の正確度を高めることに劣らず、その技術が届く範囲を広げること、これからの残された宿題に見えます。

第5部

規制科学、データ倫理、そして未来の 融合エコシステム

5.1

AI新薬開発の規制フレームワークと品質管理ガバナンス

2026年1月14日、米国メリーランド州シルバースプリングにある食品医薬品局の本部と欧州医薬品庁のオフィスで、短い文書がほぼ同時に公開されました。名前は「医薬品開発の優れた人工知能実践に関する指導原則」であり、含まれている内容は10の原則がすべてでした。人間中心の設計、リスクに応じたアプローチ、データガバナンス、ライフサイクル管理のような言葉が、わずか数ページの文書にぎっしりと詰まっていました。太平洋を挟んだ二つの巨大規制機関が手を取り合い、声を一つにしたこの場面は、新薬開発の重心が試験管や動物のケージからコンピューターサーバーへと移りつつあるという事実を、遅ればせながら公式に認めた瞬間でした。

これより8ヶ月前の2025年5月、食品医薬品局は独自に開発した生成型人工知能エルサ(Elsa)を導入すると発表し、その年の6月末には審査官全員が使えるように全面的な配置を終えました。以前は3日かかっていた書類の比較作業をエルサが6分で終わらせたという内部報告が知られ、新薬審査の速度自体が変わってきているという噂は、すぐに事実として確認されました。エルサを実際に運用する方式には、それなりの安全装置があります。食品医薬品局はエルサを政府専用のクラウド環境内に閉じ込め、製薬会社が提出した機密の審査資料をエルサが学習してしまうことがないように設計しました。

審査官が内部文書にはアクセスしつつも、その内容がエルサの学習データとして流れ込み、他の会社の審査に影響を与えることは防いであるわけです。規制当局自らも、自分たちが使う人工知能に同じ基準を突きつけているという意味です。

この変化を規制機関は自ら規制科学という名前と呼んでいます。新薬が安全で効果があるかを判断する科学的根拠自体を再設計する作業という意味です。以前の規制科学は、臨床試験から出た数字を統計で検証することに重きが置かれていました。今の規制科学は、その数字を作り出した人工知能モデルがどれほど信頼で

きるかを検証することまで一緒に引き受けなければなりません。審査官には化学と医学の知識だけでは足りず、モデルがどんなデータで訓練されたのか、どんな前提の上で動作するのかを一緒に読み取れる目が新たに必要になりました。

このすべての動きを貫く全体図は、規制が人工知能を追いかける姿から、人工知能と並んで歩く姿へと変わっているということです。10年前だけでも、規制当局は新しい技術が市場に先に出てから、遅れて規則を作っていました。今は食品医薬品局と欧州医薬品庁がガイドラインの草案について産業界と事前に意見を交わし、開発の初期段階から相談窓口を開いておく方向に舵を切りました。決まった答えを当てるよう要求する試験官から、問題を一緒に解いていくパートナーへと、規制機関の役割が変わってきているわけです。

米国議会は2022年12月に食品医薬品局近代化法2.0を通過させながら、数十年にわたって新薬許可申請の前提条件として釘を刺されていた動物実験の義務条項をなくしました。その代わりに、コンピューターの中で化合物の動きを事前に計算するインシリコシミュレーション、人間の細胞で作った微細な臓器模型であるオルガノイド、半導体チップの上に人間の臓器の機能を移した臓器チップの実験結果が、安全性と有効性の公式な証拠として入ってきました。マウスやサルを経なくても、新薬の候補物質が人間の体内でどのように反応するかをコンピューターで事前に描いてみる道が、法律で開かれたわけです。

この条項は動物実験を直ちに禁止したというよりは、代替案を認めたことに近いです。製薬会社は依然として動物実験を選ぶこともできますが、インシリコデータだけでも審査台に立てる扉が初めて開かれたという事実が持つ重みは異なります。この対応戦略をどれほど早く消化するかが、その後の数年間、各製薬会社の臨床開発のスピードを分けるものと見られます。

食品医薬品局はすでに2021年、ソフトウェア型医療機器(SaMD)に分類される人工知能・機械学習ソフトウェアを扱う実行計画を出しながら、二つの装置を用意しておきました。一つは優れた機械学習慣行(GMLP)と呼ぶ10の原則です。データをどのように集めて分けて使うか、モデルをどのように設計して試験するか、人間と人工知能がどのように役割を分けるか、性能を市場に出した後もどのように見守り続けるかを定めた一種のレシピです。訓練に使ったデータと試験に使ったデー

タを必ず分けておかなければならないという原則、実際に薬を使うことになる患者集団をデータが均等に代表しなければならないという原則、人工知能単独ではなく人間と人工知能と一緒に働く性能を見なければならないという原則が、10の中でもよく引用されます。

国際医療機器規制当局フォーラム(IMDRF)は2025年1月にこの原則を国際標準文書として確定して発表し、米国とカナダと英国の規制の文言が事実上一つに集まりました。もう一つは事前特定変更管理計画(PCCP)です。人工知能は発売された後にも新しいデータに出会いながら学び続け、変わっていくソフトウェアであるため、変わるたびに最初から再び許可を受けさせると産業全体が止まってしまいます。食品医薬品局は製造会社が今後あり得る変更の範囲と検証方法をあらかじめ定め、一度に承認を受けられるようにする道を開き、2024年末から2025年の間に人工知能が入った医療機器ソフトウェアを対象とした最終指針を相次いで確定しました。

携帯電話を買う時に、今後出るソフトウェアのアップデートをどの範囲まで自動でインストールするか、あらかじめ規約に合意しておく方式と似ています。

2025年1月6日、食品医薬品局は「医薬品および生物学的製剤の規制上の意思決定を支援する人工知能の使用に関する考慮事項」という草案指針を出しました。この指針の骨組みは、7段階からなるリスクベースの信頼性評価体系です。製薬会社はまず、人工知能モデルにどんな質問を投げかけるのかを文章で明らかにしなければならず、そのモデルが新薬開発のどの段階でどんな役割を担うのか範囲を定めなければならず、モデルの判断が間違っていた時に患者が被るリスクの大きさを等級で評価しなければなりません。

リスクが大きいと判定されれば、それだけ分厚い証拠を集めて検証計画を立てて実行した後、その結果が当初定めた目的に合っているか最後に確認する手続きを経ます。小学生に説明するなら、このような図式になります。クラスの友達がテストの答案用紙に正解を書いたからといって、先生がすぐに100点を与えるわけではありません。その答えがどのような解答過程を経て出たのか、偶然勘で当てたのではないか、次のテストでも同じ方法で正解を出せるのかを一緒に見ます。

新薬を開発する人工知能も違いはありません。特定の化合物が肝臓に有害であると人工知能が予測した時、その予測が当たったという結果だけでは不十分です。なぜそう判断したのか解答過程を人が読めなければならず、他の化合物にも同じ論理が一貫して適用されるか確認できなければなりません。法廷で証人の陳述が事実だとしても、その陳述が出た経緯と信憑性を問わなければ証拠として採用されない論理と同じです。

規制機関が恐れる状況は、人工知能が偶然正解を当てて、いざその理由を誰も説明できない場合です。説明できない判断は繰り返すことができず、繰り返すことのできない判断を次の患者にそのまま適用した時、どのような結果を生むかは誰も保証できません。

信頼性評価体系が要求する書類には、いくつかの共通項目があります。モデルを訓練し検証するために使ったデータがどこから来たのかを明らかにする出处記録、モデルの予測がどの範囲までは信頼でき、どの範囲を超えれば揺らぐのかを数値で表した不確実性管理計画、そしてモデルを市場に出した後も性能が落ちていないかを定期的に点検するライフサイクルモニタリング計画がそれです。いくつかの化合物にだけよく当てはまるモデルを、まるで全ての化合物に広く通じるかのように誇張して申請する事例をふるい落とすための装置です。

危険度によって要求水準を変える態度は、アメリカとヨーロッパと韓国の指針の全てに共通して現れます。数千の化合物候補の中から有望な候補を絞り込む初期探索用人工知能は、最終判断を人が再び検討するため、相対的に軽い検証でも通過することができます。一方、患者に投与する用量を直接提案したり、臨床試験の参加者を危険群に分類したりする人工知能は、その判断が直ちに人の体に影響を及ぼすため、はるかに分厚い証拠とはるかに頻繁な再検証を要求されます。

同じ人工知能という名前をつけていても、実際に背負う規制の重さは、その判断が及ぶ範囲によって大きく変わるわけです。

同じ指針は、モデルマスターファイル(MMF)という制度の導入も一緒に扱っています。局所的に作用する医薬品のように、複数の新薬申請に似た人工知能モデルが繰り返し使われる場合、申請書のたびにモデル検証資料を最初から新しく提出

させると、製薬会社と審査官の双方に同じ書類を重ねて作成し検討する負担が積み重なります。モデルマスターファイルは、一度きちんと検証を受けた人工知能モデルの履歴書を別途登録しておき、その後、別の申請案件ではその履歴書を参照文書として持ってきて使えるようにする装置です。

大学の卒業証書を一度発行してもらっておけば、複数の会社に志願するたびに再び試験を受けずに卒業証書のコピーを提出する方式と似ています。産業界と規制当局はまだ詳細な運営方を議論している段階ですが、モデル審査が重なる非効率を減らす中核装置として挙げています。

新薬開発の初期段階に目を向けると、人工知能は臨床試験の設計そのものを変えています。以前は研究者が経験と直感に頼って患者を何人集めるか、投与量をいくらに分けるかを決めていたとすれば、今は計算モデルが過去の臨床データを学習し、効率の高い試験設計案をあらかじめ提示します。このようなアプローチをモデル情報基盤新薬開発と呼びますが、患者数をむやみに増やさなくても統計的に信頼できる結論を得られるという長所があるため、希少疾患のように患者の募集自体が難しい領域で重宝されています。

候補物質の最適用量を見つけるのにかかる時間と費用も一緒に減らしてくれます。ただし、規制当局はこのモデルが提示した設計案をすぐに受け入れはしません。モデルが仮定した統計的前提が、実際の患者集団の特性とどれほど一致するのか、標本を減らした分だけ結果の信頼度が落ちてはいないかを、前の信頼性評価体系の手続きの中で再び探ってみます。

2024年に制定され、2025年8月から順次効力を発揮し始めた欧州連合人工知能法(EU AI Act)は、臨床試験参加者の危険を予測したり、医薬品製造工程の品質を統制したりする人工知能システムを高リスク人工知能に分類しました。高リスクに分類されれば、ヨーロッパ市場に製品を出す前に別途の適合性評価を受けなければならない、発売後も偏向性と透明性とサイバーセキュリティを含む膨大な規制遵守義務を引き続き背負わなければなりません。適合性評価は、会社自ら行う自己評価で終わらない場合が多いです。

危険度の高い人工知能医療製品は、欧州連合が指定した独立審査機関、いわゆるノーティファイドボディの別途審査を経なければならず、この機関が発行する認証書がなければヨーロッパ市場に製品を出すことができません。ところが2026年5月、欧州連合がいわゆるデジタルオムニバスという政治的合意でこの日程に手を加え、医療機器に該当する高リスク人工知能システムの適用時点を、本来の予定であった2027年8月から2028年8月へ1年遅らせました。法が要求する目標はそのままにしつつ、企業が準備する時間を稼いでくれた措置です。

その間、人工知能が入った医療機器は、既存の医療機器規則と体外診断医療機器規則の枠組みの中で先に認証を受け、人工知能法の高リスク義務は猶予期間が終わった後に重ねて適用を受ける構造で運営されています。1つの法ではなく、2つの法が同じ製品を順次、時には同時に規律する二重構造であるわけなので、製薬会社の法務チームと許認可チームは2つの日程表を並べて対応戦略を練らなければならない立場に置かれています。

ヨーロッパ医薬品庁は2023年7月に反思文書の草案を先に発表し、意見を集約した後、2024年9月に医薬品ライフサイクル全体におけるAI使用を扱う反思文書を最終採択しました。この文書は、大規模言語モデルが規制文書を作成または臨床データを要約する際に、一見するともっともらしいが、実際には存在しない事実を作り出す幻覚現象を具体的に警告しました。

人間の目には滑らかな文章であるほど、検証なしにそのまま信じてしまう危険が大きいという指摘でした。そのため、この文書は、AIがどのような判断を下そうとも、最終責任は必ず人間が負うと決めました。AIがどれほど精巧な草案を提出しても、その草案を検討して署名する人間がいなければ、規制文書として成立しないという意味です。

規制フレームワークがいかに密にできていても、その基盤となるデータ自体が一方に偏っていれば役に立ちません。AIモデルを訓練するゲノム情報と臨床記録は、西方世界の数か国、その中でも特定の人種と特定の所得層に集中していることが少なくありません。このようなデータで学習したモデルを、異なる人種や異なる遺伝的背景を持つ患者にそのまま適用すると、予期しない副作用を見落としたり、用量を誤って推奨したりする誤りへとつながる危険が増します。

EUの一般データ保護規則（GDPR）やアメリカの医療保険の携行性と説明責任に関する法律（HIPAA）のような患者情報を厚く保護する法律と衝突しないようにしながら、複数の国のデータを一緒に使う方法として、元のデータは病院のサーバーにそのまま置き、学習の結果得られた値だけをやり取りする連合学習技術が注目されています。ただし国ごとにコンピュータシステムが異なり、計算コストも相当であるため、この技術が実際の規制文書に標準として定着するまでには、さらに時間が必要と思われます。

従来のソフトウェアを検証する方式とAIを検証する方式は根本的に異なります。計算機プログラムは同じ数字を入れると常に同じ答えが出るため、一度正しいと確認すれば、その後は再度試験する必要がありません。しかしAIは新しいデータに遭遇して自ら重みを少しずつ変えるソフトウェアであるため、昨日正しかった判断が今日も同じように正しいという保証はありません。

そのため規制当局は、AIを一度通過させて終わりにする試験としてではなく、製品が市場に留まっている限り繰り返して確認する対象として扱います。自動車を初めて買う時に一度検査を受けて終わりにするのではなく毎年定期検査を受けなければならないのと似た構造です。

数千万個のパラメータからなるAIモデルが特定の化合物を有毒物質と指定したとき、その判断の因果関係を人間が解きほぐして説明できなければ、予測がどれほどよく当たっても規制承認の敷居を越えることはできません。なぜ正確度だけでは不十分なのかを技術的に説き明かす作業が、検証と確認、すなわちV&Vと呼ぶ分野です。業界はShapley Additive exPlanations（SHAP）と呼ぶ技法とLIME（Local Interpretable Model-agnostic Explanations）と呼ぶ技法をしばしば一緒に用います。

両技法ともにAIが下した一つの判断について、その判断に入力値ひとつひとつがどの程度の影響を及ぼしたのかを逆算して計算し、人間が読むことのできる数字と図として示します。ブラックボックスの内部を覗き込むことはできなくても、ブラックボックスがどの入力にどの程度の重みを置いたのかは把握することができるようにするわけです。これに加えてアメリカ機械工学会が医療分野に合わせて作ったV&V 40原則と、医療情報学の分野で昔から用いられてきた評価報告指針

であるSTARE-HIチェックリストと一緒に動員されます。

モデルの予測が実際の人体の中で起きている現象とどの程度一致しているのかを交差検証し、その検証プロセス自体を別の研究者が同じように再現することができるほど詳しく記録しておくことが核となります。試験答案を採点する際に答案表だけを見るのではなく採点基準表も一緒に公開し、誰もが同じ基準で再度採点できるようにするのと同じ理屈です。

医薬品製造プロセスにAIとデジタルツインが入ってきたことで、品質を設計段階から予め組み込むという原則であるQuality by Design (QbD) と、プロセスをリアルタイムで監視するプロセス分析技術Process Analytical Technology (PAT) がAIと噛み合い始めました。温度がわずかに変動したり不純物が混じり込む瞬間をセンサーとAIがリアルタイムで捉え、自ら制御値を調整する方式です。この時規制当局が最初に確認することは、そのデータが改ざんされていないという証拠です。国際医薬品工学会が2022年に改定したGAMP 5第2版は、データの正確性と一貫性と追跡可能性を意味するALCOA++原則をソフトウェアバリデーション実務に密に組み込みました。

誰がいつそのデータを作成したのか確認できなければならず、読むことができなければならず、実際の仕事起きたそのまさかの瞬間に記録されなければならず、元のまま保存されなければならず、正確に漏れなく残っていなければならないという要求がその中に含まれています。訓練に使ったデータと性能を確認するのに使ったデータを物理的に分離しておいたのか、モデルが試験問題を予め暗記して解くように答えを合わせる過適合はなかったのかを監査記録として残しておかなければなりません。調理人がレシピ通りに調理したという主張だけでは不十分で、食材を買った領収書と調理過程を撮った写真まで一緒に保管しておいて初めて認められるのと変わりません。

この監査跡を密に残す習慣が優れた機械学習の実践を書類ではなく工場の現場で実行する方法であり、同時に医薬品品質管理部門の実質的な競争力へと繋がります。

AIが自ら知識グラフを遡ってこれまでなかった分子構造を描き出したとき、その結果物の特許権を誰に与えるのかもまだ国際的に整理されていない問いです。特許制度は本来人間の創作行為を前提に設計されたため、人間の関与がほぼなくアルゴリズム単独で見つけ出した発明をどの国の特許庁も積極的に同じ基準で判断することができていません。発明者欄に人間の名前を書くことができない発明を認めるのか、認めるとすればその権利をプログラムを作った会社に与えるのか、それともデータを提供した機関に分配するのかをめぐって国ごとに異なる判断が出ています。

事故が実際に起きたとき、誰が責任を負うのかという問いも残っています。人工知能が分析した合成対照群データや安全性予測結果を根拠に規制当局が新薬を承認したにもかかわらず、市販後に予想できなかった患者の死亡事故が起きたと仮定してみましょう。この場合、責任を人工知能を作ったソフトウェア会社が負うべきなのか、データを提供した病院が負うべきなのか、最終的な処方を下した医師が負うべきなのか、あるいは書類を検討して承認の印を押した規制当局が負うべきなのかを分ける明確な基準が、まだどの国にもありません。

自動車の自動運転事故で、運転手と製造元とソフトウェア業者の間で責任を分ける議論がすでに起こっているように、新薬開発の人工知能も近いうちに法廷で同じ問いに直面することになると思われます。

この分野で働く人々の構成も変わってきています。以前は化学者と医師と統計学者が新薬承認の書類を埋めていたとすれば、今はデータサイエンティストとソフトウェア検証の専門家と規制専門の弁護士が同じ会議室を囲んで座ります。人工知能モデルの危険等級を格付けし、責任の所在を分ける仕事は、結局のところ法律の言語で整理されてこそ、規制当局に提出する書類が完成するからです。技術を知る人と法律を知る人がお互いの言語を学ばなければ、書類一件もまともに完成できない時代になりました。

新薬が市場に出た後は、人工知能の役割が市販後の安全性監視、すなわちファーマコビジランスの領域へと移ります。以前は医師と患者が副作用を自発的に報告してこそ問題を知ることができたため、まれに現れる危険な薬物相互作用はずっと後になってから発見されることがよくありました。今は自然言語処理技術が電

子カルテやソーシャルメディアの投稿や膨大な臨床文献をリアルタイムで読み解き、人の目では見逃しやすい微細な兆候を早期に捉えます。病院の診療や日常生活でそのまま蓄積されるこのような資料を、実使用データと呼びます。

食品医薬品局が構築した情報視覚化プラットフォームInfoViPとSentinelシステムは、この実使用データを審査官が一目で確認できるように図へと変換するインフラです。Sentinelシステムは2024年1月に自然言語処理機能を加えたバージョン3.0へと改編され、これまで250件を超える安全性分析を実行しました。国際医学団体協議会(CIOMS)は人工知能と市販後監視を扱う第14次作業部会の報告書を2025年12月4日に最終確定して発表しました。この報告書は、リスクベースのアプローチ、人間の監督、モデルの堅牢性、透明性、データプライバシー、公平性、ガバナンスという7つの原則を骨組みとして、人工知能がまれな副作用を精密に選り出し、重複する報告書を自動的にフィルタリングし、患者の個人情報を自動的に隠す実務において、どこまで信頼できるのかを具体的に整理しました。

市販後に観察された安全性情報が、再び初期の開発段階の人工知能の訓練データに戻って反映される循環構造が、このように完成しつつあります。

2025年1月24日、韓国ではデジタル医療製品法が本格的に施行され、人工知能やロボットや仮想現実技術が入った先端デジタル医療製品を別のカテゴリーにまとめ、承認から市販後の管理までを一つの枠組みの中で扱う法的根拠が初めて設けられました。同日、食品医薬品安全処は世界で初めて生成AI医療機器の承認・審査ガイドラインを制定して発刊しました。学界や医療現場や産業界の専門家からなる協議体を構成し、生成AI医療機器が開発から承認と市販後の管理に至るまで、どのような危険要因を抱えているのかをあらかじめ探ってみた結果でした。

この流れは2026年にも続きました。食品医薬品安全処は2026年6月30日、大規模言語モデル(LLM)を基盤としたデジタル医療機器の承認と審査基準を盛り込んだ申請者向けの案内書を新たに出し、人工知能技術が適用されたデジタル医療機器の特性を反映した変更管理計画書の承認・審査ガイドラインも一緒に整備しました。承認以後に予想される変更事項をあらかじめ承認してもらっておけば、その都度再び承認と審査を経ることなく製品にすぐ反映できるように道を開いた措置であり、米国の事前特定変更管理計画と骨組みが似ています。

食品医薬品安全処は、2026年から2030年までの5年間で347億ウォンを投じ、人工知能による新薬開発と先端バイオとデジタル医療機器の分野を網羅する規制科学の専門人材1,100人を育成するという計画も明らかにしました。この予算は、学位課程と非学位課程を網羅する10の遂行機関を経て分けて執行されます。審査官が人工知能モデルの動作原理を理解できなければ、いくら精巧な指針を作っても実際の審査台では使い物にならなくなるという判断が、この投資の根本にあると思われます。米国や欧州が法律や指針で基準を立てている間に、韓国はその基準を実際に扱う人を育てることに予算を注ぎ込んでいるわけです。

このような流れは、韓国の製薬会社やバイオベンチャーにも実務的な課題を投げかけます。米国の食品医薬品局に新薬を申請するには米国の信頼性評価体系に従わなければならない、欧州に進出するには人工知能法の高リスク分類と考察文書の原則と一緒に満たさなければならない、国内ではデジタル医療製品法と食品医薬品安全処のガイドラインに従わなければならない。3つの地域の規定が要求する文書の形式と用語は少しずつ異なりますが、根本にある問いは結局一つに通じます。国内企業が開発の初期から国際基準に合わせて検証資料を蓄積しておけば、後で複数の国に同時に申請書を出す際に、同じ資料を繰り返して作成する手間を大幅に減らすことができます。

医薬品規制調和国際会議(ICH)は2026年1月29日、モデル基盤の新薬開発の一般原則を扱う新しいガイドラインM15を最終採択しました。人工知能を含む計算モデルが臨床試験の設計や用量の決定に使われる際、各国の規制当局が共通して参考にできる言語を作った作業です。まだ人工知能だけを切り離して扱う独立した医薬品規制調和国際会議のガイドラインはありませんが、2025年7月23日から発効された臨床試験設計の原則であるICH E6(R3)と一緒に置いて見れば、さまざまなガイドラインが一つの方向へと噛み合っていることがわかります。

米国の現代化法と信頼性評価体系、欧州の人工知能法と省察文書、国際機構の優秀機械学習慣行と販売後監視原則、そして韓国食薬処の法と人力養成計画は、それぞれ異なる文書名を帯びているが、要求する核心は一つに集まっています。人工知能が出した答えが正しいかどうかを超えて、その答えに至る過程を人が検証し、責任を持つことができなければならないという原則です。法律家の目で見ると

と、これは見慣れない要求ではありません。証拠の信憑性を問い、判断の根拠を記録に残すよう要求してきた法廷の長い慣行が、いま実験室と工場と病院のデータパイプラインの中へ移行しているだけです。

ただし、法廷の慣行と異なる点もあります。法廷は既に起きた事件について事後に是非を判断しますが、新薬規制は事故が起きる前に検証手続きを密に設けておかねばならない事前予防の領域です。人工知能が出した結果を人がどれほど理解し、責任を持つことができるかという問いは、今後も新薬開発の各段階ごとに繰り返して投げかけられるものと思われます。

責任の所在を決める基準がまだないという空白を埋めようとする動きは、欧州で先に現れました。欧州連合は2024年11月製造物責任指令を改正して公布し、ソフトウェアと人工知能システムも有形の物と同じく製造物として見なし、厳格責任を負うことに決めました。加盟国は2026年12月9日までにこの指令を国内法に転換しなければならず、被害を受けた患者は開発社の過失を別途証明しなくても損害賠償を請求する道が広がりました。ただし、その時点の技術水準では欠陥を知ることができなかったという開発危険抗弁はソフトウェアにも同様に適用されるため、人工知能開発社が責任を免れる道も新たに開かれました。

韓国の製造物責任法はまだソフトウェアを製造物の範囲に明示的に含めていないため、人工知能が設計した新薬で事故が起きても、ソフトウェア自体を相手にした責任追及は法廷で争いの余地が大きい状態のままです。国内立法が欧州のペースについていけないと、同じ事故に遭った韓国の患者がより重い立証負担を負う立場に置かれる可能性が大きいです。

審査官が人工知能を検証する能力を備えているかという問いには、人力統計がより冷たい答えを示しています。米国保健福祉部は2025年4月1日、食品医薬局職員3,500名を一度に出した後、残された審査官が同僚が担当していた業務までも引き受けたという報道が続きました。リスク基盤信頼性評価体系を発表してから3か月も経たない時点でした。食品医薬局はその格差を埋めようとするかのように、2025年12月1日、全職員が使えるエージェント型人工知能を新たに配置し、このツールは事前審査と審査結果検証と販売後監視業務まで分け持つことができるようになりました。

不足する審査人力を人工知能で補う方式は、人工知能の判断を人が検証しなければならないという原則と食い違って見えます。食薬処の事情も余裕がありません。2025年11月基準で食薬処の医療製品審査人力は369名で、米国食品医薬局審査人力9,000名、欧州医薬品庁5,500名、日本医薬品医療機器総合機構1,300名と比べると、桁から異なります。食薬処は行政安全部と協議して審査官295名を追加採用する方案を議論していますが、採用情報を知らせるホームページが2021年から止まっているという指摘が出るほど、準備はまだ遅い状態です。

人力養成予算をいくら増やしても、審査台の前に座る人自体がこのように少なければ、人工知能モデル一つ一つを密に見つめる余裕は国ごとに大きく異なるものと思われ、検証が不十分な国から先に申請書を出す戦略を取る企業も出てくるものと思われま

5.2

データ倫理、アルゴリズムの公平性、および連合学習エコシステム

2024年7月、ソウルの韓国製薬バイオ協会事務室に、10を超える製薬会社と大学研究チームの担当者たちが集まりました。事業の名前はK-MELLODDY、2028年12月まで続く5年間の国家課題で、保健福祉部と科学技術情報通信部が348億ウォンを投入し、連合学習ベースの新薬開発を加速させようという計画でした。会議室の空気は微妙でした。テーブルの上には各機関の名札が置かれていましたが、ノートパソコンの画面には互いに見せるべき共有画面は一つも表示されていませんでした。

競争関係にある製薬社の研究者たちが一か所に集まりましたが、誰も自分の会社の化合物データを競争相手に渡す気はありませんでした。それでもこの人たちは共に人工知能モデルを訓練することにしました。データを一行も授受しないままで、です。この矛盾のように聞こえる約束が実際に機能する理由を理解するには、新薬開発の人工知能が今直面している根本的なジレンマからまず見てみる番です。

精密医療を実現するには、数百万人の遺伝体と診療記録を集める必要があります。一つの病院、一つの製薬社が持つ患者数だけでは、希少疾患のパターンを読み取ることはできません。しかし患者の遺伝情報と診療記録を一か所に集める瞬間、そのデータベースはハッキングと流出の標的になります。米国はHIPAA（健康保険の移転と責任に関する法）で、欧州連合はGDPR（一般個人情報保護規則）で患者情報の国境間の移動と外部への流出を厳格に防いでいます。

製薬社の立場からは、良い人工知能モデルを作るデータが病院の塀の向こうにあるのに、法はその塀を越えてはいけないと言っています。この膠着状態を解くことができないと、人工知能新薬開発は少数の大型製薬社が自社で保有したデータだけで規模を広げる閉ざされたゲームに留まってしまいます。

この膠着状態を突き抜けた技術が連合学習（Federated Learning）です。小学生でも分かるように説明すると、こんな具合です。同じクラスの友達10人が試験勉強をしているのに、各自自分のノートは絶対に他人には見せないと約束しました。代わりに毎晩自分が解いた問題と正解を照らし合わせた後、どんな解き方のコツが点数をより上げるかだけを互いに教え合います。10日が経つと、10人全員が相手のノートの内容は一文字も見えていないのに、各自の成績は一人で勉強していた時より上がっています。

連合学習もこれと同じです。病院と製薬社は患者データを入れたノートを自分のサーバーの外に絶対出しません。代わりに人工知能がそのノートを見て少しずつ賢くなった結果、つまり数字でできたモデルの重みだけ暗号化して中央サーバーに送ります。中央サーバーは複数の機関から来た重みを一つに合わせてより賢くなったモデルを各機関に返します。このように複数の機関の重みを平均して一つに合わせる手続きを連合平均（Federated Averaging）と呼びます。

この重み交換方式は理論だけでは安全を保証しません。重み自体から元のデータを逆に推定しようとする攻撃が可能だからです。そのため最近の研究者たちは、重みを中央に送る前にわざと微細なノイズを混ぜ込む差分プライバシー（Differential Privacy）技術と、複数の機関が計算結果だけ照らし合わせて互いの入力値は見えないようにする安全多者計算技術を一緒に付けています。

2026年に発表されたPrivMol研究は、分子特性を予測するグラフ人工知能モデルにこの二つの技術を結合して、予測精度を大きく落とさないまま元のデータが漏れ出す危険を低くする結果を見せました。一つの錠前だけではなく二重、三重に閉じる方式と同じです。

連合学習が新薬開発の現場で実際に通じるという事実をはじめて証明したプロジェクトはヨーロッパのMELLODDYでした。ヨーロッパ革新医薬イニシアティブ（IMI2）とヨーロッパ製薬産業協会（EFPIA）が1800万ユーロを共同で出しながら2019年から3年間導いたこのプロジェクトには、グローバル製薬社10社と技術・学術機関7社が参加しました。これらが共に学習に使ったデータは26億個を超える実験活性値、2100万個を超える化合物、4万個を超える生物学的検査結果でした。どの会社も自分の化合物リストを競争相手に見せませんでした。

それでも連合学習で作った共同モデルは、各社が自分のデータだけで作ったモデルより予測性能が平均5.5パーセント、多いときは9.7パーセントまで良くなりました。互いに異なる冷蔵庫に各自異なる材料を入れておきながら、レシピ一つを共に調整してみんなの料理の味を上げた形です。この改善幅は実験室で化合物を直接合成して試さなくても、除外できる失敗作がそれだけ増えたということで、新薬開発初期段階の時間と費用を一緒に減らしてくれます。

同じ発想が韓国ではK-MELLODDYという名前で生まれ変わりました。韓国製薬バイオ協会が2024年7月に錨を上げたこの事業は、連合学習プラットフォームを作る課題、病院と製薬社のデータを集めて供給する課題、そのデータで人工知能モデルを訓練する課題、このように三つの軸で構成されています。事業団が出そうとしている結果物はFAM（Federated ADMET Model）という予測モデルです。ADMETとはある化合物が体に入った時どれだけよく吸収されて、どのように広がって、どのように分解されて、どれだけ毒性を起こすかを一つにまとめて呼ぶ言葉です。

2025年には高麗大、崇実大、韓国科学技術院、LG化学など複数の機関が新しい細部課題5件を新たに引き受けて着手しました。国内中堅製薬社とバイオベンチャー立場からはデータを買う資本が不足している場合が多くて、所有権は各自守りながら共に学習できるこの方式が魅力的だという評価を受けています。

製薬社間の協力を超えて、病院と病院を結ぶ連合学習も速く増えています。希少疾患は一つの国、一つの病院に集まる患者数がそもそも少なく、人工知能がパターンを学ぶ標本そのものが不足しています。例えば腫瘍画像や病理写真を扱う多国籍病院コンソーシアムでは、複数の国の病院がそれぞれ自分の患者の画像を自分のサーバー内に置いたまま、腫瘍を区別するモデルの重みだけをやり取りします。

どの病院も患者の身元がわかる元の画像を他の病院に送ることなく、モデルは一つの病院のデータだけでは到達できなかったであろう精度に達します。このような協力が可能なのは、元のデータが物理的に国境を越えないからです。

連合学習が守り抜く価値と、説明可能な人工知能が守り抜く価値は、天秤の両方の皿に置かれた二つの分銅のようなものです。片方の皿にはセキュリティという分銅があります。元のデータを元あった場所から一步も動かさず、知識だけを行き来させる連合学習が、この皿を押さえつけます。もう片方の皿には信頼性という分銅があります。人工知能がなぜそのような判断を下したのか、人間が理解できるように説明してくれるXAI(eXplainable AI, 説明可能な人工知能)が、この皿を押さえつけます。

どちらか片方だけが重くなると、天秤は傾いてしまいます。データを安全に守ったとしても、モデルがなぜその結論に至ったのかを誰も説明できなければ、医師も患者もその結果を信じて使うことはできません。反対に、説明は完璧でも、その説明を作り出したデータ自体が流出の危険にさらされているならば、その人工知能は最初から病院の敷居をまたぐことはできません。

しかし、データを安全に集めたからといって、そのデータ自体が公正であるという保証はありません。世界中に散らばっている臨床試験やゲノムデータベースの多くは、WEIRD、つまり西洋の教育を受けた、工業化された裕福な民主主義国の人口に偏っています。2025年半ばの時点で、ヒトゲノム全域関連解析(GWAS)に参加した人のうち、ヨーロッパ系の祖先を持つ人口の割合は90.53パーセントに達します。アフリカ系の祖先を持つ人口の割合は、過去10年間でほとんど増えていません。

2009年以降、15年以上にわたって、ヨーロッパ系中心の研究の割合は5パーセントポイント余り減ったにとどまりました。北米とヨーロッパの白人患者のデータだけで訓練したナノ粒子薬物送達モデルや腫瘍標的予測モデルを、遺伝的、生活習慣的な背景が全く異なるアジアやアフリカ、南アメリカの患者にそのまま当てはめると、予測性能が大きく落ち、予想外の毒性による副作用が現れる危険性が高まります。

皮膚がんを見つけ出す人工知能の診断ツールの中には、明るい肌の色の患者の写真を中心に学習されたせいで、暗い肌の色の患者では診断の精度が大きく落ちるものがいくつもありました。人工知能ソフトウェアだけの問題でもありません。コロナ19大流行の時期に広く使われたパルスオキシメーターも、肌の色が暗い患

者の血中酸素飽和度を実際よりも高く測ってしまう誤差を示し、米国食品医薬品局が別途安全性の警告を出したことがあります。細菌性膣炎を診断するとある人工知能ツールは、白人女性では精度が高かった半面、アジア人女性では精度が大きく落ち、ヒスパニック系女性には病気がないのに病気があると誤って判定する偽陽性が特に多く出ました。

米国のある大手保険会社が使っていたリスク予測アルゴリズムは、黒人患者を健康管理プログラムから体系的に締め出しました。医療費の支出額を健康リスクの代理指標としていたのですが、黒人患者は同じ病気にかかっているにもかかわらず医療へのアクセスが悪く、支出額そのものが少なかったからです。アルゴリズムは、人種を直接見ることなく、人種による差別をそのまま写し取りました。技術的には、偏ったパターンにわざと立ち向かうようにモデルを訓練する敵対的学習手法も、偏りを減らす方法として一緒に使われています。

データの公平性を数値で表した研究もあります。機械学習モデルの健康の公平性を測るHEALフレームワークで、ある皮膚疾患の診断モデルを分析した結果、性別による下位集団では公平性の点数が92.1パーセントと優秀でしたが、がんではない皮膚疾患を患う高齢患者の集団では、点数が0.0パーセントにまで落ちました。同じモデルの中で、ある患者は精密に診断され、ある患者は事実上放置されていたという意味です。訓練データのバランスが崩れた時に起こるもう一つの罫は、近道学習です。

健康な人の標本が80パーセント、疾患のある人の標本が20パーセントというデータでモデルを訓練すると、モデルはすべての入力値を無条件に健康だと当てずっぽうに選ぶだけでも、80パーセントというもっともらしい精度を得てしまいます。実際には病気を一つも見つけ出せないモデルなのに、数字だけ見ると優秀に見えるという錯覚が生まれます。この罫に陥ると、希少疾患の患者や有色人種のようにデータが少ない患者ほど、誤診や不適切な処方の危険をそのまま抱え込むことになります。

米国食品医薬品局(FDA)と欧州医薬品庁(EMA)は、臨床試験を設計する段階から少数人種や様々な年齢層、社会経済的に疎外された階層を意図的に含めるように求める多様性実行計画を、臨床試験のプロトコルに明記しました。米国食品医薬品

局は2025年8月、人工知能が組み込まれた医療機器ソフトウェアを、市販前から市販後までどのように管理するかを定めた指針を出しました。この指針は、訓練データの代表性と偏りの管理、サイバーセキュリティ、市販後の性能監視を求めつつ、あらかじめ定めておいた範囲の中では、新たに許可を受けなくてもアルゴリズムを修正していけるようにする、事前決定変更計画(PCCP)という制度を根付かせました。

2026年には、臨床意思決定支援ソフトウェアに関する個別の指針も出ました。データの量を増やすことにとどまらず、人口構成を実際に代表する高品質なデータを選んで盛り込まなければならないという原則は、もはや数枚の文書による勧告を超えて、新薬承認審査の実務基準になってきているようです。

データの偏りを根本から直そうとする試みも一緒に進められています。これまでヒトゲノム研究は、特定の人口集団の1人のゲノムを標準参照配列とし、他のすべての人のゲノムをその配列と比べる方式に依存してきました。この単一の参照配列は、最初から人類の多様性を反映するように設計されてはいなかったため、標準から外れた遺伝的変異が多い人口集団であるほど、診断や治療の予測において不利になりました。

最近では、複数の人口集団のゲノムを重ねて描くパンゲノム(pangenome)参照マップが、この限界を補う代案として浮上しました。これまでの研究で疎外されていた人口集団の標本をパンゲノムマップに十分に盛り込めば、人工知能モデルが最初からより均等なデータを学習の材料にすることができます。

データの偏りを直したとしても、ディープラーニングモデルの中には越えなければならない壁がもう一つあります。ブラックボックス問題です。数千万個から数十億個に及ぶパラメータで絡み合ったニューラルネットワークは、ある化合物が肝臓に有害であるとか、抗がん剤の標的として有望であるという結論を瞬時に導き出します。しかし、実際にどのような経路でその結論に至ったのかは、設計した研究者でさえ説明できないことが多いのです。蓋を開けられない箱に化合物の情報を入れると、答えだけが飛び出してくるのと同じです。

箱の中でどのような計算が交わされたのかは、誰にも見ることはできません。患者の命がかかった薬物承認審査において、このような不透明性がそのまま受け入れられることは困難です。医師が処方し、規制当局が承認の印を押すためには、予測確率という数字一つだけでは不十分であり、その背後にある判断過程を覗き見ることはできなければならないからです。

このブラックボックスを開けてみようとする試みが、説明可能な人工知能、すなわちXAIです。SHAPやLIMEのように人工知能の結論を逆戻りして辿る手法は、化合物のどの化学構造やタンパク質のどの配列が、結合力や毒性の予測にどれだけ寄与したかを、色や数字で目に見えるように描き出します。完成した家具だけをぽつんと出す代わりに、どのネジをどこにいくつ打ったのかが書かれた組み立て説明書を一緒に渡すのと同じです。医師や規制当局者は、この説明書を見て初めて、その結論が偶然の統計的な相関関係ではなく、実際の生物学的な根拠から出たものかどうかを検証できるのです。

ただし、2026年に出された抗生物質探索用のXAI評価研究は、このような説明手法の性能を測る基準そのものがまだ統一されていないという限界を指摘しました。説明を提示するツールは増えていますが、その説明がどれくらい信頼できるかを測る物差しはまだすべて揃っていないわけです。

人工知能の予測モデルが実際の臨床現場に根付くためには、このような説明手法一つだけでは不十分です。米国機械工学会が定めたV&V 40の原則のように、計算モデルの信頼性を検証する手続きと、人が最終判断を再確認する監督手続きが、設計の段階から一緒に組み込まれていなければなりません。最近では、CONSORT-AIやDECIDE-AIのように、人工知能モデルの性能や訓練環境、臨床応用の範囲を文書として透明に残すように求める報告ガイドラインも、学界を中心に増えています。

その手続きを飛ばしたまま市場に出た人工知能の医療機器は、後日の監査や訴訟の過程で根拠資料を提出できずに足かせとなるケースが増えています。ただし、このようなガイドラインはまだ学術誌の編集方針に近いだけであり、法律で強制される規範ではありません。勧告と義務の間に置かれたこの隙間は、実際の医療事故が起きたとき、法的責任を誰に問うのかという問いにそのままつながります。

人工知能が診断や予測を超え、自ら新しい化合物を設計する水準に至ったことで、まったく異なる法的な争点が生まれました。生成型人工知能は、研究者が一つ一つ指示しなくても、構造と活性の間の関係を収めたデータやタンパク質の折り畳み情報を自ら学習し、最初から新しい分子構造を設計し出します。いわゆるデノボ（de novo）設計です。インシリコ・メディスンが独自の人工知能プラットフォームで設計し、特発性肺線維症の治療候補物質として臨床第2相まで推し進めた事例や、エクセンシアが設計して2020年に臨床第1相に入った強迫性障害の治療候補物質DSP-1181のように、人工知能が最初から最後まで設計に関与した化合物は、すでに数件が臨床段階に入っています。

このようにして生まれた分子に特許を出すとき、発明者の欄に誰の名前を書くべきかが問題となります。

この問いに正面からぶつかった事件があります。米国の発明家スティーブン・テーラーは、DABUSという人工知能が自ら発明を成し遂げたとして、人ではなくDABUSを発明者とした特許を、米国や欧州、英国、韓国を含む複数の国に出願しました。2022年、米国連邦巡回区控訴裁判所はテーラー対ビダル（Thaler v. Vidal）事件において、米国特許法がいう発明者は自然人でなければならないと釘を刺し、DABUSの発明者としての地位を認めませんでした。2025年11月28日には、米国特許商標庁が人工知能が関与した発明の発明者資格を扱う新しい審査ガイドラインを出し、人工知能は実験室の器具のように精巧な道具にすぎず、発明の核心である着想（conception）は人だけができる行為であるという原則を再確認しました。

2026年3月4日には、日本の最高裁判所でさえテーラーの上告を棄却し、7年に及んだDABUSの全世界での特許攻防が事実上幕を下ろしました。時間を遡ってみると、欧州特許庁は2021年に、英国最高裁判所は2023年12月の判決で、すでに同じ結論に至っていました。これまでに米国、欧州、英国、日本をはじめとするほぼすべての国の特許庁と裁判所は、DABUSを発明者として認めませんでした。

弁護士目から見れば、これらの判例が語る原則は明確です。人工知能が候補物質を実際に設計し出したとしても、特許書類の発明者欄には、その人工知能を設計して扱い、化合物の価値を実際に見出した研究者の名前を書かなければ特許を

守れないということです。問題は、数千億ウォンをかけて人工知能の新薬開発プラットフォームを構築した企業が、この原則をどのように実務に移すかです。人工知能が一晩の間に数百万個の候補分子を吐き出すとき、研究者がその中のどれにどれだけ創造的に介入したかを、特許出願の書類ごとにいちいち立証しなければならないからです。

この立証に失敗すると、苦労して発掘した化合物が特許で保護されないまま競合他社にそのまま露出する危険があります。反対に、研究者の寄与を過度に膨らませて書けば、後になって特許自体が無効に覆される口実を自ら作ることになります。このため、一部の企業は最初から特許の代わりに営業秘密として化合物を保護する戦略を選ぶこともあります。営業秘密は、競合他社が独自に同じ化合物を見つけ出すと保護膜がそのまま消え去るという限界を抱えています。このバランスを取る実務は、今後数年間、特許訴訟の核心的な争点として残る可能性が高いです。

知的財産権をめぐる紛争は、発明者地位の問題で終わりません。人工知能が大型言語モデルで伝統医学古書と先住民共同体の薬用植物知識を解読して新しい天然物新薬候補を見つけ出す研究が増えながら、その知識を数百年間守ってきた共同体にどのような権利を認めるかという問いも一緒に浮かび上がりました。2010年に採択されたナゴヤ議定書は、遺伝子資源とそれに付随する伝統知識を利用しようとする側があらかじめその知識を持つ共同体の同意を得て、利用して生じた利益を公正に分け合うよう定めた国際条約です。

これまでに142カ国がこの条約を批准し、130を超える国がそれぞれ利益共有関連法を持っています。しかし、人工知能が膨大な伝統知識文献を瞬く間に学習して新しい化合物を提案する今、どの資料がどの共同体から出たかを一つ一つ追跡して利益を分け合う手続はまだ適切に整備されていません。南アフリカ共和国のサン族とコイ族がルイボスティー産業界と結んだ利益共有協定は、伝統知識保有共同体に実際にロイヤルティが支払われた事例としてしばしば引用されます。

人工知能が伝統知識文献を調べて新しい化合物候補を見つけ出すときも、このレベルのロイヤルティ支払と出所表示が行われるべきだという声が学界で大きくなっています。伝統知識を私有化された特許に変えて莫大な利益を得ようとする誘因

と、その知識を守ってきた共同体に正当な対価を返すべきだという要求が正面から衝突しています。

特許問題を超えると、より先鋭な質問が待っています。人工知能の判断で患者が実際に傷ついたとき、責任は誰にあるのかという質問です。説明可能な人工知能で検証され、連合学習で訓練された臨床意思決定支援システムであっても、訓練データに隠れていた微細なバイアスのために特定の患者群に安全でない医薬品を安全だと誤って推奨することができます。例えば特定の遺伝型を持つ患者集団にのみ現れる副作用を訓練データが十分に含まなかったなら、人工知能はその危険を予測できないまま安全だという判定を下すことができます。

このとき責任が神経網を設計した情報技術企業にあるのか、データを提供した病院にあるのか、その推奨を信じて処方した医師にあるのか、それとも認可の判を押した規制当局にあるのかを見分ける確固とした法体系はまだ世界のどこにも完備されていません。

欧州連合はこの空白を埋めようと別の人工知能責任指針を準備しましたが、2025年2月この指針初案を取り下げると明らかにした後、その年10月公報掲載で撤回を確定し、人工知能関連事故は改正された製造物責任指針の中で扱うよう方向を変えました。欧州連合人工知能法そのものは2024年8月1日発効し、2025年8月からは汎用人工知能モデルを作る側に過料さえ科すことができるようになりました。医療機器を含む高リスク人工知能システムの義務は2026年オムニバス改正で日程が後ろにずれ、附属書3に該当する高リスクシステムは2027年12月2日から、製品に組み込まれる附属書1系列は2028年8月2日から適用されます。

高リスク人工知能を扱う企業は、リスク管理計画と良質な訓練データ、透明な文書、人間の監督装置、配置後の継続的な点検体系を備えなければならず、このうち一つでも違っていると製造物責任指針が言う欠陥として扱われる余地が生じます。義務を定めておき、その義務を破ると事故が起きたとき責任を逃れにくくする構造です。

欧州連合人工知能法は医療人工知能を二重に縛って規制します。患者診断や治療決定に関わる人工知能システムは医療機器規定に従った認証を受けなければなら

ないのと同時に、人工知能法が定めた高リスクシステム義務も一緒に守らなければなりません。二つの規制のうちどちらか一つでも守られなければ市場進入そのものが妨げられます。アルゴリズムの性能だけでは不足で、その性能を生み出したデータと設計過程を最初から最後まで文書に残さなければ市場に入ることができるという意味です。

連合学習コンソーシアムに参加する機関の間でも参加契約書に損害賠償と保険負担をあらかじめ定めておく実務が増えています。事故が起きたときに法廷で責任の所在を争う前に、契約条項で危険を先に分け合おうとする動きです。

韓国の法体系も似た方向に動いています。仮名情報とは名前や住民登録番号のように特定人をすぐに知ることができる項目を削除したり他の値に変えて、追加情報なしには誰かを知ることが難しくした個人情報と言います。韓国の個人情報保護法は異なる機関が持つ仮名情報を担当者同士直接対比しないよう、結合専門機関という第3の機関を経てのみデータを結合できるよう定めています。政府は2022年から2026年まで続く研究開発ロードマップに従って複数の機関が各自持つ仮名情報を互いに結合する制度と技術を整備してきており、2026年には仮名情報を結合した事例を競う競技大会さえ開きます。

保健福祉部は2026年研究開発事業で医療データと人工知能技術を新薬開発と医療現場のあちこちに繋ぐという計画を明らかにし、政府は2030年まで4,000億ウォンを人工知能ベース新薬開発と先端脳科学研究に投じることにしました。

欧州ではこれと対をなす制度として欧州保健データ空間（EHDS）規定が2025年3月26日発効しました。この規定に従い規定全体は2027年3月26日から適用され、個別診療のための第1次利用と研究や政策、人工知能開発のように元々の目的と異なる第2次利用はすべて2029年と2031年二度にわたって段階的に適用されます。患者サマリーと電子処方箋、そしてほとんどの研究用データが2029年にまず開かれ、画像と検査結果、退院記録、ゲノム情報のように敏感または重い資料は2031年に延期されています。

各加盟国は保健データアクセス機関を設け、データの利用許可を出し、責任を管理することになっています。欧州データ保護会議と欧州データ保護監督官は、こ

の規定が一般データ保護規則の定めた個人の権利を弱める危険があると懸念し、その結果、最終条文は二次利用の法的根拠を一般データ保護規則の上に再び乗せる形で妥協されました。プライバシーを守る規定と、データを集めて使おうとする規定が、一つの体の中で互いに引っ張り合っている様子です。

K-MELLODDYの会議室に集まった研究者たちが、互いのノートを最後まで見せ合わなくても一緒に成績を上げることができた理由は、技術が信頼とセキュリティという二つの価値を同時に守ることができるという事実を証明したからです。しかし、技術が開いてくれた扉の後ろには、まだ解かれていない問いが残っています。どの人口集団のデータで訓練したモデルなのか、人工知能が設計した化合物の特許は誰の名前で残るのか、事故が起きたときの責任はどの機関が負うのかを決める作業です。

患者と市民の目から見れば、このバランスが崩れたとき、先に危険にさらされる側はいつも、データが少なくて声を出しにくい人々であるという点を覚えておくべきです。この作業は実験室の中では終わらず、法廷や国会、国際機関の交渉テーブルへと引き継がれています。契約書に偏りに対する監査条項と、発明への貢献度を記録する義務を今の段階で前もって入れておく方が、後日紛争が起きたときに責任を明らかにする費用を減らしてくれると思われます。

アメリカと日本の裁判所が人工知能を発明者として認めないことでまとめたように、これからの数年間に、データの公平性と責任の所在を扱う規範も、国ごとに判例や指針として一つずつ積み重なっていく可能性が高いです。その規範が細くなるまで、連合学習と説明可能な人工知能が今作り出している信頼とセキュリティのバランスは、引き続き試験台に上がると思われます。

BogeonbokjibuとGaeinjeongbobohowiweonhoeは2025年12月31日、Bogeonuiryodeiteo hwaryong gaideurainを改正して施行に入りました。複数の病院が一緒に進める研究や、独自のGigangogeonuiryojeongbosimuiwiweonhoe、すなわちDRBを組むのが難しい中小機関のために、共用DRB制度が新しくできました。GigansaengmyeongyulliwiweonhoeとDRBの審議手続きを一つに合わせる標準手続きも設けられました。

世を去った患者の医療データを研究に使う時は、身元を示す識別子と遺族に関連するコードを消して使うように基準も新しく決めました。人工知能で新薬候補を探す国内のスタートアップや大学の研究チームの立場では、複数の病院のデータを集めて予測モデルを訓練させるまでにかかっていた審議の待機時間が、今回の改正で減ると思われます。

ところで、このすべての手続きは、データを出した患者本人ではなく、審議委員会の目だけで回っています。国内の個人情報の同意体系は、病院が情報を最初集める時に一度だけ同意をもらう方式にとどまっており、その後、仮名処理を経た情報がどの研究、どの企業の人工知能の学習に使われたのか、患者が後で振り返ってみる通路は設けられていません。その上、仮名情報は統計作成や科学研究のように本来の目的とは違って使われる時も、情報主体の追加の同意を再びもらわなくてもよいように、法律がすでに許可しています。

英国の国民保健サービスのように、情報主体が後からでも拒否の意思を明らかにすれば情報の使用を止めるオプトアウト装置を設けた国もありますが、韓国の法律はまだその段階まで進んでいません。同意が必要ないということが、監視まで必要ないという意味ではないはずですが、その監視を誰が担うのかははっきりしていません。データを元々あった場所から移さない連合学習がいくら安全でも、患者が自分の情報の行き先を自ら確認する道がなければ、その安全は半分のものにとどまる可能性が高いです。

連合学習が原本データを守ってくれるという信頼にも隙間はあります。2025年3月、国際学術誌Journal of Cheminformaticsには、研究者ファビアン・クリュガーが率いた研究チームが、新薬候補物質の性質を予測するニューラルネットワークモデルに対して、メンバーシップ推論攻撃(membership inference attack)を試みた結果が載りました。モデル内部の重み付けを覗き見ずに、外から質問を投げて答えを受け取るだけで、特定の化合物が訓練データに入っていたかどうかを言い当てる方式でした。

どのデータセット、どのモデル構造でも情報は漏れてしまい、標本数が少ない少数集団に属する化合物であるほど、より簡単に見破られました。新薬開発で標本が少ない化合物は、競合他社がまだ手をつけていない価値ある候補である場合が

多く、本当に守るべき情報から漏れてしまう危険を抱えているわけです。ファビアン・クリュガーの研究チームは、この攻撃手法をコードと一緒に公開し、他の研究者たちも自分のモデルの危険を自ら点検できるようにしました。

5.3

医薬品業界の未来：自律化研究所(Self-driving Labs)

2025年10月、Seoulのある実験室が静かに扉を開きました。Hanguk Jeyak Bio Hyeophoe傘下のAI Sinyak Yeonguwonが作ったこの空間には、白衣を着た研究員の代わりにロボットアームが一つ、秤の前に立っています。ロボットアームはガラス瓶の蓋を回して開け、0.005ミリグラム、塩一粒の重さの数百分の一に当たる試薬を正確に取り出します。液体を混ぜて再び重さを量る動作が、数秒間隔で繰り返されます。夜10時を過ぎると、研究員たちは一人二人と退勤し、実験室の明かりは消えます。ロボットはその暗闇の中で止まりません。

人々はこの空間を「Dark Lab」と呼びます。協会の関係者は、この実験室の目標が、人が眠っている間にも実験の処理量を最大に引き上げることにあると説明しました。ロボットが扱う反応容器ごとに実験の条件と結果が自動で記録され、翌日出勤した研究員は、昨夜積み重なった数十件の実験結果をコーヒ一杯と共に確認することから一日を始めます。出入り口の前のモニターには、その夜処理された実験の件数と成功率がリアルタイムの数字で上がり、温度や濃度があらかじめ決めておいた範囲から外れると自動で警告音が鳴るように設計されています。

2026年6月に初めてマスコミに公開された時、研究院側は今年中に150人の人材を新しく育て上げるという計画も共に明らかにし、国内の製薬会社数社はすでにこの実験室の運営方式を自社に取り入れる案を検討しています。わずか数年前までは映画の中の場面のように思われていた光景が、今ではSeoulのある実験室で毎晩繰り返される日常になっています。

2024年、国際学術誌Natureには、巨大言語モデルを基盤とした人工知能が分光器と液体取り扱いロボットを自ら制御し、人の手を通さずに有機ケイ素化合物を29種合成したという研究が載りました。このうち8種はそれまで世に知られていなかった化合物でした。化学者が試薬棚の前で何日も格闘しなければならなかった仕事を、人工知能は反応経路を自ら探し、装備の日程を調整して、一晩で終わらせました。過去の新薬発掘は、病気の原因となる標的を探し、その標的に合う化合物

を設計し、合成し、試験管と細胞で効能を確認するという4つの段階を、数年にわたって退屈に繰り返す仕事でした。

よく設計、合成、試験、分析の頭文字を取ってDMTA周期と呼ばれるこの過程は、人が一周回るだけで数週間かかる遅い循環でした。多国籍製薬会社Janssenがエイズ治療薬Prezistaの生産工程を、リアルタイムセンサーと機械学習基盤の自動調節体系に変え、アメリカ食品医薬品局とヨーロッパ医薬品庁の厳格な基準を満たした事も、同じ流れの上に置かれています。Seoulのロボットアーム1台から、地球の反対側の学術誌の1ページに至るまで、実験という行為自体の定義が変わっています。

この変化を導く力は、自ら思考し行動するAgent AIです。目標だけ与えられれば、計画を立てて実行まで責任を持つ人工知能を言います。過去の医療用人工知能は、一つの入力値に一つの出力値を出す、受動的な道具に近いものでした。Agent AIは違います。データを絶えず受け入れ、以前の判断を記憶し、自ら立てた計画が間違っていると判断されれば、次の計画を書き直します。ここに視覚と、言語、化学信号を一緒に読み取るマルチモダリティ能力まで備え、この人工知能はクラウド上のソフトウェアに留まらず、分光器やロボットアームのような実験室のハードウェアを直接指揮する水準にまで上がりました。

ChemCrowという名前が付けられた初期の構造は、この原理をよく見せてくれます。人が複数の部署に電話をかけて業務を分けて任せるように、巨大言語モデルが分光器やクロマトグラフィー装備ごとに決まった通信規格、いわゆるアプリケーション・プログラミング・インターフェースを通じて直接命令を下す方式です。研究者が目標物質と守るべき安全条件だけ入力すれば、人工知能は膨大な化学反応データベースを巡って合成経路を組み、試薬を注文し、装備の日程を調整し、実験を実行した後、その結果を分析して次の条件を自ら直します。

文献を読むことから仮説を立てること、ロボットで合成して分析することまで、人の手を通さずに一つにつながるこの流れを、閉ループシステムと呼びます。決められた順序だけ繰り返す既存の自動化設備とは違い、このシステムは状況が変われば、次に実行する手続き自体を自ら組み直すという点で根本的に異なります。

閉ループの最初のボタンは、仮説を立てることです。Google DeepMindが出した「Coscientist」は、Gemini巨大言語モデルを土台として、4つのエージェントが役割を分けて協業するシステムです。生成エージェントが論文とデータをざっと読んで仮説をいくつか出すと、省察エージェントが同僚の審査者のようにその隙を探り出します。順位エージェントは、生き残った仮説をトーナメント方式で競わせ、進化エージェントはその中で優秀な仮説を再び整えて結合します。

4つのエージェントはお互いの結果物をやり取りし、同じ仮説を何度も再検討しますが、この反復は人が学会発表を前に同僚と数ヶ月やり取りする論争を真似た構造に近いです。イギリスのImperial College LondonのJose Penades教授の研究チームは、細菌の間で抗生物質の耐性遺伝子がどのように広がるかを10年近く研究してきましたが、まだ発表していないこの研究テーマをCoscientistにそのまま任せてみました。人工知能は研究チームの未公開データに接近できない状態でも、48時間で同じ結論を指す仮説を出しました。

人の研究者が数週間から数ヶ月にわたって立てて反論していた仮説検証の過程を、この4つのエージェントははるかに短い時間の中で自動で回します。問題は、仮説が早く出れば出るほど、その仮説を実際に確認する手が新しいボトルネックとして浮かび上がるという点です。自律駆動実験室は、まさにこのボトルネックを埋めるために登場した次の段階です。

仮説がロボットの前に到着すると、次は物理的実行です。中国のクリスタルパイ(XtalPi)は、世界中で約300台のロボットワークステーションを24時間休まず稼働させ、独自に開発した約500個の人工知能モデルと量子力学計算を組み合わせ、化合物を合成・評価しています。同社はイーライリリーとの最大2億5000万ドル規模の共同研究を進めており、国内ではJW中外製薬と協力して、炎症疾患の標的タンパク質であるSTAT6を狙ったリード化合物と一緒に精製しています。人工知能アルゴリズムが描き出した分子構造やタンパク質配列は、直ちにロボット組立ラインと細胞培養システムに転送され、物理的な物質に変わります。

この過程で作られたすべての産出物は、コンピュータ内の設計図から最終生産ロットまで、一本でつながるユニークな番号が与えられ、最後まで追跡されます。これは、後でどのような実験条件がどのような結果につながったのかを、人間が逆

向きに追ってみることができるようにする、最小限のセーフガードでもあります。ソウルのダークラボでのロボットアームが0.005ミリグラム単位で試薬を扱う動きも、地球の反対側の実験室の300台のロボット軍団も、結局同じ原理で動いています。人間なら間違いやすい反復作業こそ、ロボットの手は疲れることなく、同じ条件を何百回も何千回も全く同じように繰り返すことができるという点が、この物理的実行段階の価値を高めています。

実験が終わると、循環は止まらず次の段階に進みます。生物反応器の中に埋め込まれた複数のセンサーは、細胞の遺伝子発現と電気的信号、代謝産物の痕跡をリアルタイムで読み取り、人工知能制御装置に送ります。強化学習コントローラーは、この信号を受け取り、培養液の濃度やシグナル物質の量を自ら調整し、品質が基準から外れる前に先手を打ちます。結果分析から得られたデータは、再び最初の仮説設計段階に戻り、次の実験条件をより精密に整える材料となります。このように、仮説設計、実験室自動化、結果分析、データフィードバックという4つの段階は、円を描いて限りなく繰り返されます。

昨日の失敗した実験が今日の成功した実験を作る材料となり、今日の成功は再び明日の疑問を生みます。実験が1回回るたびにシステム自体が少しずつ賢くなる構造であるため、この実験室は人間が操作してこそ動く静的な空間ではなく、自ら学習しながら体を大きくしていく一つの生命体に近いです。人間が1日8時間働いて数件の実験を処理していた以前と異なり、今は同じ時間に数百件の条件を同時に試し、その結果を次の実験にすぐに反映させる規模の違いが生まれています。

自律実験室の可能性は、大企業のプレスリリース以外でも確認されています。複数の人工知能エージェントがチームを組んで、新型コロナウイルスに付着するナノボディ、つまり、ラクダやアルパカの体から見つけた抗体を真似して作った極めて小さな結合タンパク質を設計し、実際の実験でその効果まで検証した『バーチャルラボ』研究が、ピアレビューを経て学術誌に発表されました。米国のフューチャーハウスが作った『ロビン』というシステムは、仮説を立て、実験を設計し、結果を解釈するプロセスを一つに統合して、ドライ老年性黄斑変性、つまり、年を取りながら視力を徐々に失わせる眼疾患の治療候補として、緑内障治療薬のリパスジルを新たに見つけ出しました。

この研究は2026年5月の国際学術誌ネイチャーに掲載されました。いずれの事例も、大学と研究所の同僚研究者に検証され、論文として残されたという点で、自律実験室が投資誘致のためのプロモーション文句ではなく、実際に機能する科学的方法論であるという事実を支持しています。巨大資本を背景とする企業だけでなく、比較的予算に余裕のない大学研究チームも、クラウド上の人工知能と共有実験室機器を借りる方式でこの流れに参加する道が開かれています。

バイオベンチャーキャピタルのフラッグシップ・パイオニアリングは2025年3月、Lila Sciencesという名の会社を世に初めて知らせ、シード投資2億ドルを受けたと明らかにしました。その年の10月には、エヌビディアの投資部門が主導する1億1500万ドルを追加で調達し、累積投資額を5億5000万ドル、企業価値を13億ドル程度まで引き上げ、2026年6月には企業価値85億ドルを認められる後続投資の議論が進行中という報道も伝わりました。OpenAIでリサーチ部門副社長を務めたリアム・フェデロスと、Google DeepMindで研究員として働いていたエキン・ドゥス・チュブクが共に設立したPeriodic Labsも2025年3億ドルのシード投資を受け、超伝導体を含む新素材を狙った自律実験室を建設しています。

英国グラスゴー大学のリー・クローニン教授が創設したChemifyは、化学反応をコードに移し、ロボットがそのまま再現するようにする『ケムピューター』技術で、2025年5000万ドルを超えるシリーズB投資を受けました。三社とも新薬一つだけを狙っていません。彼らが狙うのは、化学と素材、エネルギーを包括する研究開発そのものの自動化であり、投資家リストにはエヌビディアとアマゾン創業者ジェフ・ベゾス、前Google CEOエリック・シュミットのような名前も上がっています。三社がこれまで集めた投資金を合わせても9億ドルに達するのに、数年前まで大学実験室の研究課題程度と見なされていた分野にこれほどの資金が集中するという事実自体が、この技術が通過した検証の重さを示しています。

自律実験室が生み出す成果は、既に複数の国で具体的な数字で確認されています。オーストラリアの連邦科学産業研究機構(CSIRO)は、新しい触媒を開発するのに必要だった期間を5年から6ヶ月に短縮しました。カナダのアクセラレーション・コンソーシアムは、自律実験室を通して21個を超える新しい候補物質を見つけ出しました。韓国でも、韓国科学技術研究院(KIST)が『オクトパス』という名のオペ

レーティングシステムを独自開発し、複数のロボット機器を1つの人工知能で指揮する実験を進めており、韓国製薬バイオ協会と延世大学校はK-NIBRT事業団を通して関連インフラを別途構築しています。

どちらの事業もまだ少数の研究者だけが接したことのある歩み始めの段階ですが、政府と大学がそれぞれ散らばって試みていた実験の自動化を一つの標準の下に集める作業であるという点で、次の段階へと移る橋渡しの役割を果たしています。国籍と研究対象はそれぞれ異なりますが、人間が数年がかりで取り組んでいた仕事を、人工知能とロボットが数ヶ月、時には数週間で終わらせているという点は共通しています。

Hanguk Gwahak Gisul Jeongbo Tongsinbuは、2026年から2028年までの3年間で495億ウォンを投じて自律実験室を6ヶ所建設するという計画を打ち出しました。汎用実験室の1ヶ所はCatholic Daehakgyoが担当して「K-Cell」プラットフォームを構築し、特化実験室の5ヶ所は、Daegu Gyeongbuk Gwahak Gisulwonが液体生検を、KAISTが感染症を、Goryeo Daehakgyoが遺伝子伝達体を、POSTECHが酵素工学を、Ulsan Gwahak Gisulwonがオルガノイド基盤の薬物評価をそれぞれ担当します。この事業を率いるOh Dae-hyeon未来戦略技術政策官は、人工知能とロボティクスが結合した自律実験室が、バイオ研究開発のパラダイムを根本的に変える契機になるだろうと明らかにしました。

政府はこれとは別に、物理的人工知能、つまりロボットの頭脳に該当する技術を国内技術として確保することにも、2年間で340億ウォンを追加で投入することにしました。SeoulのDark Labが先に見せてくれた徹夜の実験の風景が、2028年には全国の6ヶ所で同時に繰り広げられる絵へと拡大する準備をしているわけです。

論文を1編読む瞬間から特許を出願する瞬間までにかかる時間を90日前後に圧縮するという構想が、業界の一部から出ています。まだこれを実際に成し遂げたと公認された事例はありませんが、その方向が単なる掛け声だけではないことを推測させる数字はすでに出ています。Hong Kongに本社を置くInsilico Medicineは、標的発掘プログラムPandaOmicsと新薬候補物質設計プログラムChemistry42を前面に押し出し、前臨床候補物質、つまり動物試験に入る準備が終わった物質を見つけ出すのに、平均して12ヶ月から18ヶ月かかることが明らかにしました。

伝統的な新薬発掘にかかっていた2年半から4年という期間の半分にも満たない数値です。この会社が進めたあるプログラムは、候補物質を80個も合成することなく、新しい標的に対する新しい分子を見つけ出すのに260万ドルだけを使いましたが、これは既存の方式が要求していた費用の10分の1の水準です。この会社が発表した特発性肺線維症治療の候補物質は、人工知能が設計した新薬の中で初めて臨床第2相で肺機能改善効果が確認され、2025年には痛みの治療を狙った新しい候補物質ISM9528を前臨床段階へと引き上げ、年間の市場規模が1000億ドルに達する分野に挑戦状を突きつけました。

この会社が共に運営しているinClinicoというプログラムは、臨床第2相試験が第3相まで続いて成功する確率を、あらかじめデータで測ってみるツールですが、候補物質を見つける段階から臨床試験を設計する段階まで、同じ会社の人工知能が最初から最後まで関与するという点で、閉ループの縮小版と言えます。

量子力学の「もつれ」と「重ね合わせ」という現象を機械学習と混ぜ合わせた量子機械学習（QML）は、分子間の微細な相互作用を既存のコンピューターでは手に負えないほどの精度で計算し出します。フランスのQubit Pharmaceuticalsは、Pasqalが作った中性原子方式の量子コンピューターOrionを利用して、タンパク質の結合ポケットの中で水分子がどの位置に置かれるのかを量子アルゴリズムで計算し出しました。2026年3月には、アメリカのCleveland ClinicとIBMが手を組み、303個の原子からなるミニタンパク質Trp-cageの電子構造を、量子コンピューターと古典コンピューターを行き来する混合方式で計算し出すことに成功したと発表しました。

薬物が体内でコラーゲン組織に思いがけない毒性を引き起こすかどうかの問題のように、複数の原子が同時にもつれて作り出す微妙な相互作用をあらかじめ読み取る作業においても、量子アルゴリズムは既存の方式よりも精巧な答えを出しています。市場調査機関は、この分野の市場規模が2025年の1億2600万ドルから2035年には6億3800万ドルへと、年平均17.9パーセントずつ拡大すると見込んでいます。ここに、原料の投入から完成品の生産まで工程を止めずに流し続ける連続製造技術が噛み合います。乾燥や粉末の混合、錠剤を打ち出す打錠の過程で生じる微細な温度変化や不純物を、工程分析技術とデジタルツインがリアルタイム

で感知しシミュレーションしますが、Janssenが生産ラインにこのような体系を取り付けて規制基準を満たした事例が代表的です。

発見の速度と生産の速度が、今や同じリズムで動き始めたわけです。実験室で数週間で見つけ出した候補物質が、工場でも数週間で大量生産の条件を整える状況になり、これまで発見と生産の間に置かれていた時間の隙間自体が狭まっています。

産業界はこのような統合について、自動化を超えたインダストリー5.0という名前を付けています。機械が人間の反復労働の代わりにしていたインダストリー4.0の次の段階であり、人工知能とロボットが人間の認知的、身体的限界を広げてくれる同伴者となる時期であるという意味です。コンピュータのサーバー内にある設計図が実験室のロボットを経て工場の生産ラインまで途切れることなく続くのも、その結果の一つです。人の手が届かなかった夜や週末まで研究開発に組み込まれるようになり、同じ予算と同じ人員で生み出すことができる候補物質の数自体が何倍にも増えています。しかし、同伴者にも守るべき規則は必要です。

この閉ざされたループの中に入ってみると、一つの新薬が歩んでいく道が見えます。病気に関わるタンパク質を見つけて出す標的発掘の段階で、人工知能は膨大なゲノムや論文のデータをざっと調べて、有力な標的を絞り込みます。そのタンパク質の3次元構造が予測されると、人工知能はその構造にぴったり合う分子を最初から新しく描き出し、同時にその分子が肝臓でどのように分解されて体の外へどう抜け出すのか、毒性はないのかを前もって計算します。ある候補は、小さな化合物一つでは足りない判断され、複数の標的を同時に狙う多重標的抗体や細胞治療薬の形に再設計されたりもします。

このような候補物質は、脂肪の粒や高分子で作られたナノ粒子に乗せられ、肝臓や腫瘍のように希望する組織まで正確に配達されるように、もう一度整えられます。患者の細胞からお皿の上で育てられたミニ臓器であるオルガノイドと、患者の体全体をコンピュータの中にそのまま複製したデジタルツインは、動物実験の前の段階で候補物質の効果や副作用をあらかじめ試してみる舞台となります。実験室自体も例外ではありません。ロボットに実際の試薬を使わせる前に、実験条件をコンピュータの中の仮想実験室で先に回してみても、成功する可能性が高

い条件だけを選び出して物理的なロボットに引き渡す方法も、だんだん広く使われるようになっていきます。

ここで生き残った候補は、デジタルツインであらかじめ設計された臨床試験に入り、腫瘍の遺伝子マップを読んで患者さん一人一人に合った治療を見つける精密腫瘍学の道具たちと出会います。このすべての段階で出た記録は、規制当局に提出する書類の草案としても一緒に積み重なります。標的を探すことから、患者さんに合った処方を見つけることまで、別々の部署で別々に回っていた作業が、今では一つの循環の中でデータをやり取りしながら同時に回っています。

この循環が、滑らかにだけ見えるわけではありません。エージェントAIが偏ったデータやエラーのせいで危険な物質を自ら設計したり、連続製造ラインで間違った条件を適用して不良な医薬品を大量に作り出したりする可能性は、実際に存在する脅威です。そのため、実験室の現場では、人工知能の活動範囲をあらかじめ決められた安全の限界内に縛っておき、これを監視する「ガードエージェント」を別に置く場所が増えています。決められた限界を超えようとする試みが感知されると、ガードエージェントは実験自体をすぐに止めて人に判断を委ねるように設計されており、自律性と安全性が互いに牽制し合う二重の装置の役割を果たします。

センサーとソフトウェアの間の同期エラーによってデータがずれる技術的な欠陥もたびたび発生し、人工知能がなぜそのような判断を下したのかを人が完全に説明できないブラックボックスの問題も残っています。このような危険を統制するために、米国食品医薬品局と欧州医薬品庁は2026年1月、新薬の開発に人工知能を使う時に守るべき共同指針を初めて発表しました。核心は、品質に影響を与える判断であるほど、人工知能の提案を人が最終確認しなければならないという原則です。製薬会社は、国際製薬工学会が改訂したコンピューターシステム検証の標準指針と、医療機器に分類されるソフトウェアに人工知能を適用する時に従うべき実行計画、データを誰がいつ作ったのか最後まで追跡できなければならないというデータ完全性の原則に合わせて、人工知能の学習過程と偏向性の検証手続きを文書として残さなければなりません。

人工知能が自ら学習して少しずつ変わっていく現象、いわゆるモデルドリフトを統制するために、米国食品医薬品局は、市販された後でも人工知能が安全に自身をアップデートできる範囲をあらかじめ決めて承認を受ける事前変更管理計画制度も導入しました。

法の領域に移ると、問いはさらに鋭くなります。人間の発明者の実質的な介入なしに、人工知能が数千万件の論文を読んで白紙から描き出した分子に特許権を与えられるのかについての答えは国ごとに異なり、まだ世界的に整理されていません。自律実験室が最適化した用法が患者に思いがけない副作用を起こした時、その責任がソフトウェアを組んだ技術企業とこれを導入した製薬会社、承認を出した規制当局のうち誰にあるのかを明確に定めた法体系も、まだどこにもありません。米国と欧州がそれぞれ異なる速度と方式で指針を作っている間、Hangukを含む多くの国々は、まだこの問いに対する独自の答えを用意できていない状態です。

特許庁と裁判所の判断が国ごとに食い違えば、同じ人工知能が設計した同じ分子を巡っても、ある国では特許が認められ、ある国では拒絶されるということが実際に起こり得ます。技術は国境を越えて一晩の間に新しい化合物を29個ずつ作り出しているのに、その結果に責任を持つ法の言葉は、依然として国ごとに異なる速度で動いています。

ロボットが手の代わりをする実験室で、人がすべきことが何なのかについて、私の判断はこうです。研究者は消えません。ただ、担う役割が根本的に変わっています。試薬を運んで秤にかける手の労働はロボットに引き渡されますが、そもそもどんな質問を投げかけるのか、どのデータを信じてどのデータを疑うのかを決めることは、依然として人がやらなければなりません。人工知能が一晩中吐き出した数百個の結果の中で、何が偶然で何が本当の発見なのかを見分ける眼識、システムが越えてはいけない安全の境界線をどこに引くかを決める判断、事故が起きた時にその責任を誰が負うのか答えられる法的、倫理的な監督の役割は、ロボットには代わりができません。

研究者は、実験台の前でピペットを握る技術者から、実験室全体を設計してその実験室が犯すかもしれない失敗まで責任を負う監督者へと移り変わっているよう

です。大学で化学や生物学を勉強していた学生たちにも、今では統計やコーディング、ロボット工学の基礎が一緒に求められており、Dark Labを運営する協会が人材の育成を事業の最初の行に置いているのも、同じ理由からでしょう。実験の技術を指先で身につけていた徒弟式の教育の代わりに、人工知能が出した判断を検証してその限界を掘り下げる訓練が、次の世代の研究者の基本技になりつつあるわけです。

Dark Labの明かりが消えた夜にも、誰かはそのロボットに今日何を実験しろと命令を下しましたし、次の日の朝、その結果に署名する人もやはり決まっていなければなりません。

2028年、49,500百万ウォンが投入される韓国の自律実験室6ヶ所が計画通りに開設されれば、ソウルのダークラボとともに、全国各地の実験室が每晚同時に明かりを消し、ロボットを動かす光景が繰り広げられることになります。地球の反対側では、ロボット300台を操るクリスタルファイと、企業価値85億ドルを認められるライラ・サイエンスが同じ時間、同じ方式で実験を続けているでしょう。超電導体を狙うペリオディック・ラブスのロボットも、化学反応をコードに移すケミファイのケムピューターも、それぞれ異なる物質を狙いながら同じ原理で夜を明かすでしょう。

文献を読み、仮説を立て、物質を合成し、結果をフィードバックするこの循環は、人間が眠っている間も止まりません。標的を見つける人工知能も、分子を描く人工知能も、ロボットアームを動かす制御装置も、今では異なる企業と異なる国に散らばっていても、一つの循環のように噛み合って動作しています。ただし、ある早い夜明けに、自律実験室が一晩中最適化した配置から予想外の数値が飛び出してきたとき、その前に立つ人が誰であるかは、まだどの国の法律にも明確に記されていません。

ロボットはその問いに答えを出しません。その問いは翌日朝出勤した研究員の机の上に、一晩中積もった数十件のデータとともに静かに置かれているだけです。その答えを準備する仕事は、実験室の明かりが再び点く朝ごとに人間に戻ってきます。

アメリカのバークレー国立研究所が作った自律実験室A-Labは2023年、ロボットが人間の手を介さずに17日間で新しい無機化合物数十種を発見したという論文をネイチャーに掲載しました。数ヶ月後、ロンドン大学の化学者ロバート・パルグレイブをはじめとする外部研究者がその結果を改めて検討しました。「新たに発見した」とされた物質のうち多くは、すでに結晶構造データベース、つまり既知の物質の原子配列を集めたデータに登録されていた化合物でした。ネイチャーは2026年1月に訂正文を出しました。

この実験室を率いたバークレー研究チームは、その化合物をA-Labが実際に合成したのは正しいと反論し、論争はまだ終わっていません。仮説を立て、実験を実行し、その結果を解釈する仕事まで人工知能ひとつが担う閉じたループであればあるほど、その判断を改めて検討する外部の視点がむしろ減るという事実を、この事件が示したと見られます。試薬を直接触った経験なく、画面の中の数字だけで化学を学んだ次世代の研究者がこのような誤りをどれだけ捉えることができるかは、まだ誰も確実には言えません。

危険な物質を人工知能が設計しないようにするメカニズムも、すでに一度試験台に乗ったことがあります。2022年、アメリカの製薬会社コラボレーションズ・ファーマシューティカルズの研究者たちは、毒性を避けるように訓練した自社の新薬設計人工知能MegaSynの目標を逆さまにひっくり返して、毒性を最大に引き上げる分子を見つけるよう命令を出してみました。普通のノートパソコン一台で6時間走らせたただけなのに、人工知能は化学兵器禁止条約が禁止した神経作用剤VXを含め、4万個に達する毒性物質リストを吐き出しました。

研究者たちはこの結果を詳しく公開する代わりに、学術誌ネイチャー・マシン・インテリジェンスに危険性だけを知らせ、具体的な化学構造は閉じておきました。新薬を作るよう教えた人工知能と毒を作るよう教えた人工知能の間には、人間が目標ひとつを変えること以外には何の違いもなかったという事実が、この実験が残した不快な教訓です。

MegaSynの目標を逆転させるのに必要だったのは、コード数行を直すジェスチャーひとつだけでした。ロボットアームが一晩中試薬をはかりにかけ、人工知能が自ら次の実験を決める時代には、そのジェスチャーひとつに気づき、止める人間が

実験室の中に必ずいなければなりません。翌日の朝、誰かが人工知能が一晩中設計した薬ひとつをあなたの前に提示するなら、その薬を作ったのが人間か機械かより先に尋ねるべき質問があります。その薬を誰が、どのような方法で確認したのかという質問です。

AI創薬のパラダイム転換と未来の融合エコシステム

著者 | キム・ギョンジン

発行人 | キム・ギョンジン

発行所 | キム・ギョンジン弁護士出版社

価格 : USD 15

© キム・ギョンジン 2026

本書は著作者の知的財産であり、無断転載および複製を禁じます。

注) 本書の一部の内容と編集過程には人工知能ツールの助けを借りました。

kimkj008@gmail.com

この AI 書斎の本を応援してください

本書がお役に立ちましたら、今後の翻訳や無料公開版の制作、
そして開かれた人工知能の知識プロジェクトを PayPal で応援いただけます。



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。