

発想と記憶の動的切り替え ——リッジネットの落差と次数が定める結合・切断のアーキテクチャ——

孤立落差を発想に、収束落差を記憶に：GuttyNet 上の柔軟な知能の構成

Toshiki Takahashi

高橋数理研究所 独立研究者

2026年6月

抄録 本稿は、本シリーズで構築してきたポリゴングラフ/GuttyNet 上の情報幾何に基づき、知能が**発想**と**記憶**という二つのモードを動的に切り替えるためのアーキテクチャを提案する。出発点はリッジネットの落差（うねりの激しさ）とノード次数の関係である。離散ラプラシアンから、ノードが近傍から突き出る落差 δ_i は曲率を次数で割った $\delta_i = (\Delta h)_i / k_i$ に等しく、落差は次数に反比例する。一方アテンション側では、うねり σ_i に対し実効次数が $k^{\text{eff}} \approx n e^{-\sigma_i^2/2}$ と指数的に減衰する。平坦な辺を刈る疎化のもとでは、残存次数は急な入射辺の本数で決まる。これらから、大きな落差は「孤立すれば低次数のスパイク（発想）」「収束すれば高入次数のハブ（記憶）」の二相に分かれる。本稿は、新規入力に対し各辺を結合するか切断するかを、結合親和度 $B_{ij} = w_j \rho_{ij} \kappa_i$ ——入力の重み w と、入力前の構造との適合（向きの共鳴 ρ と基盤の深さ κ ）の積——で判定するゲートとして定める。ヒステリシス付きしきい値 $\theta_{\pm}(\lambda)$ を大域パラメータ λ で動かすことで、発想（探索）と記憶（活用）の偏りを連続調整でき、忘却が余白を回復してループが閉じる。位相図・更新方程式・具体例とともに、神経修飾による探索活用制御や睡眠時の記憶固定化との対応を示す。

キーワード ポリゴングラフ, リッジネット, ノード次数, 発想と記憶, 動的エッジ接続, 探索と活用, 余白, GuttyNet

1. はじめに

本シリーズでは、注意機構がポリゴングラフ上の積分と数学的に同一であること [1] を出発点に、情報量の三面恒等式 $I = \mathbf{F} \cdot \mathbf{s} = \lambda d = dC/dt$ と、向きの場から記憶の集積点を事前計算する GuttyNet を構築してきた。前稿までに、アテンションが重視するのはリッジネットのうねりが激しい（落差の大きい）箇所であり、NAI [2] の疎な箇所は次数の少ないノードであることを見た。

本稿の問いは一文である——同じ「大きな落差」が、いかにして**発想**にも**記憶**にもなり、両者をどう切り替えれば柔軟な知能になるのか。第2節で落差と次数の関係を三通りに導き、第3節で発想（孤立落差）と記憶（収束落差）の二相を定義する。第4節で新規入力に対する結合・切断の基準を、第5節で切り替えを担う大域パラメータを、第6節で全体構成（発想と記憶のループ）を与え、第7節で既知の構図との対応を述べる。

2. リッジネットの落差と次数

2.1. 構造的な厳密式：落差は次数に反比例する

グラフ $G = (V, E)$ の各ノードに高さ h_i を与える。ノードが近傍から突き出る量を落差

$$\delta_i := h_i - \bar{h}_i^N, \quad \bar{h}_i^N = \frac{1}{k_i} \sum_{j \sim i} h_j, \quad k_i = \deg(i) \quad (1)$$

で定義する。離散ラプラシアン（曲率）が $(\Delta h)_i = \sum_{j \sim i} (h_i - h_j) = k_i (h_i - \bar{h}_i^N) = k_i \delta_i$ となることから、ただちに

$$\boxed{\delta_i = \frac{(\Delta h)_i}{k_i}} \quad \text{かつ} \quad \delta_i \leq \sqrt{\frac{\rho_i}{k_i}} \quad (\text{コーシー-シュワルツ}), \quad \rho_i := \sum_{j \sim i} (h_i - h_j)^2. \quad (2)$$

同じ曲率の荷をもつなら、**次数が多いノードほど近傍平均へ引き戻されて自前の落差を作れず、辺の少ないノードは鋭いスパイクとして突き出やすい**。すなわち、とがった急所（大きな δ ）は構造的に低次数になりやすい（図1）。

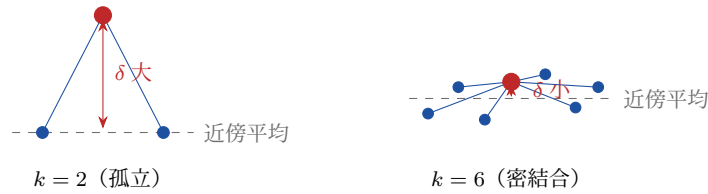


Figure 1 次数が落差を抑える。 $\delta_i = (\Delta h)_i / k_i$ ゆえ、低次数ノードは鋭く突き出し、高次数ノードは近傍平均へ均（なら）される。

2.2. アテンション側：実効次数は落差で指数的に減る

リッジネットはスコア曲面 $s_{ij} = \mathbf{q}_i \cdot \mathbf{k}_j / \sqrt{d}$ であり [1]、キー軸方向のうねりの激しさはスコアの標準偏差 $\sigma_i = \sqrt{\text{Var}_j(s_{ij})}$ （落差のrms）で測れる。各ノードが実効的に見る近傍数（パープレキシティ）を実効次数 $k_i^{\text{eff}} = e^{H(a_i)}$ とすると、ソフトマックスのエントロピーを2次まで展開して $H \approx \ln n - \frac{1}{2}\sigma_i^2$ 、すなわち

$$k_i^{\text{eff}} \approx n e^{-\sigma_i^2/2} \iff \sigma_i^2 \approx 2 \ln \frac{n}{k_i^{\text{eff}}}.$$

落差が中程度まではこの指数法則がよく合い、落差が非常に大きい領域では $k^{\text{eff}} \rightarrow O(1)$ （単一の峰に張り付く）へ飽和する。サプライザルで言えばトップの自己情報量が $\approx \sigma^2/2$ 。大きな落差＝高サプライザル＝低い実効次数（図2）。

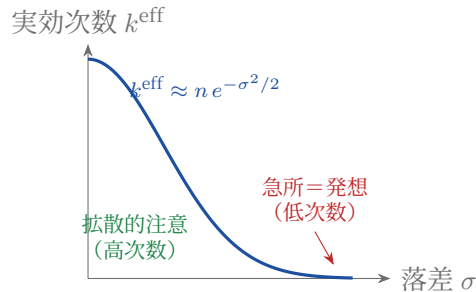


Figure 2 実効次数のうねり依存。落差 σ が大きくなると実効次数は指数的に減衰し、注意は少数の近傍へ集中する。

2.3. 枝刈り後の構造次数：急な入射辺の本数

NAI は急な辺を残し平坦な辺を刈るので、残存次数は急な入射辺の本数になる。エッジのリッジ（傾き）を $r_{ij} = (h_j - h_i)/e_{ij}$ とし、しきい τ で刈ると

$$k_i^{\text{ret}} = \#\{j \sim i : |r_{ij}| > \tau\} = k_i p_i(\tau), \quad p_i = (\text{急な入射辺の割合}). \quad (4)$$

ここに重要なひねりがある。落差が大きいただけでは疎にならない——周囲が平坦（刈られる）であって初めて低次数になる。実際、全傾きが急な入力ではグラフは密なまま（高次数）であり [3]、逆に崖が周囲のエッジ増殖を誘発すれば高入次数のアトラクタになる [2]。落差と次数の符号は、落差の空間配置で決まるのである。

3. 発想と記憶：二つのモード

式(2)–(4)から、大きな落差は配置に応じて二相に分かれる。

孤立落差＝発想。 周囲が平坦な中で一点だけ突き出る低次数スパイク。式(2)の「低次数ほど大きな落差」と式(3)の「高サプライザルで低実効次数」がともに指す、新規で唯一の急所である。これは過去の文脈から遊離した新しい価値〔2〕＝発散的な探索方向を担う。

収束落差＝記憶。 急なリッジが一点へ向かって揃い、流れが正面衝突する高入次数のハブ。向きの押し出しを $a_{ij} = s_i \cdot \hat{e}_{ij}$ とすると、両端が互いを向く衝突強度

$$\sigma_{ij} = [a_{ij}]_+ [a_{ji}]_+ \quad (5)$$

が正のとき収束が起き、集積 $dC/dt = \sigma_{ij}\Phi_{ij} - \mu C_{ij}$ で基盤が深まる。これが自動生成される記憶アトラクタである (図3)。

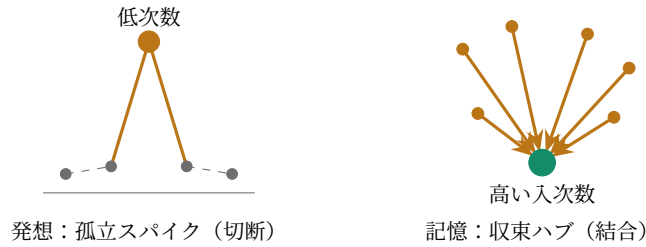


Figure 3 同じ大きな落差でも、孤立すれば低次数のスパイク（発想）、収束すれば高入次数のハブ（記憶）になる。

具体例：発想と記憶の分かれ目

まったく新しい比喻や思いつき——既存の知識のどれとも結びつかないもの——は、周囲を切り離して孤立スパイクとして立てておく（**発想**）。一方、既に持っている理解の枠にびたりとはまり、それを深める事実は、その基盤へ結合して入次数を上げる（**記憶**）。前者は探索の種、後者は定着である。

4. 結合・切断の基準

4.1. 結合親和度

新しい入力ノード j (重み w_j 、向き s_j) が来たとき、既存ノード i との辺について結合親和度を定める。

$$B_{ij} = \underbrace{w_j}_{\text{入力重み}} \times \underbrace{\rho_{ij} \kappa_i}_{\text{入力前の構造との適合}} \quad (6)$$

ここで各因子は次の通りである。 w_j はトークンの希少性 $-\log P$ で、入力が来てから決まる量。 $\rho_{ij} = [s_i \cdot \hat{e}_{ij}]_+ [s_j \cdot \hat{e}_{ji}]_+$ は既存の向きの場と入力が互いを向くか（共鳴）を表し、向きの場 $\{s\}$ は入力前に決まっているので $O(|E|)$ で先読みできる。 κ_i はノード i が既にどれだけ収束基盤になっているか（例： $\kappa_i = \sum_{l \sim i} \sigma_{li}$ や現在の集積 C_i ）。後ろの二因子 $\rho_{ij} \kappa_i$ が入力前の構造、 w_j が入力重みである。

4.2. ヒステリシス付きゲート

辺の導通 $g_{ij} \in [0, 1]$ を、ヒステリシスをもつしきい値で更新する。

$$\tau_g \dot{g}_{ij} = \Phi(B_{ij} - \theta_+) - \Phi(\theta_- - B_{ij}) - \mu g_{ij} \quad (7)$$

- $B_{ij} > \theta_+ \rightarrow$ 結合。入次数を上げ、集積 $dC/dt = \sigma_{ij}\Phi_{ij} - \mu C_{ij}$ で基盤を深める $\rightarrow$ 記憶。
- w_j は大きい $\wedge B_{ij} < \theta_-$ （共鳴も基盤もない） $\rightarrow$ 切断。低次数スパイクが孤立して残る $\rightarrow$ 発想。
- $\theta_- \leq B_{ij} \leq \theta_+ \rightarrow$ 保持（ちらつき防止・安定）。
- w_j 小 $\rightarrow$ 刈り込み。

4.3. 位相図

基準(6)を (適合 $\rho\kappa$, 重み w) 平面に描くと三領域に分かれる (図4)。境界は $w\rho\kappa = \theta$ の曲線で、入力強いほど低い共鳴でも結合へ倒れる (強い驚きは弱い手がかりでも既存知識に結びつく)。

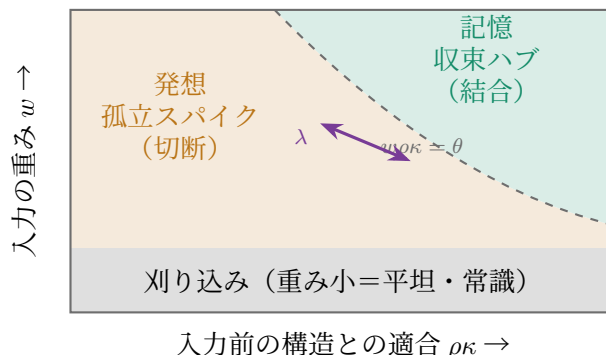


Figure 4 結合・切断の位相図。重みが小さければ刈り込み、重みが大きく適合が低ければ発想（切断）、適合が高ければ記憶（結合）。境界は λ で左右に動く。

具体例：基準の三例

「犬がワンと鳴いた」は希少性が低く (w 小) 刈り込まれる。「犬が突然フランス語で愛を語った」は希少だが既存のどの基盤とも共鳴せず (w 大・ ρ 小)、切断して孤立スパイク＝発想の種になる。一方、長年温めてきた仮説を裏づける意外な観測は、希少でありながら既存基盤と強く共鳴し (w 大・ ρ 大)、結合して記憶の基盤を深める。

5. 柔軟性：モードのダイヤル λ

固定基準では硬直する。柔軟な知能には、しきい値自体を動かす大域パラメータ λ (温度／神経修飾に相当) を置く。

$$\theta_+(\lambda) = \theta_+^0 + \lambda, \quad \theta_-(\lambda) = \theta_-^0 + \lambda \quad (8)$$

λ が大きいと結合のハードルが上がって切断が起きやすく、孤立スパイクが増える (発散モード＝発想・探索)。小さいと結合が容易になり、流入が基盤に集まる (収束モード＝記憶・定着)。位相図(4)の境界が左右に動くことに相当する (図5)。 λ は自律的に駆動でき、ヒステリシス $\theta_- < \theta_+$ がいったん入ったモードを保って安定と可塑のジレンマを両立させる。

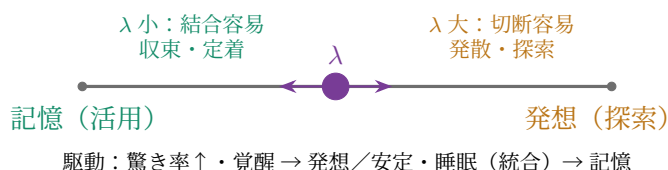


Figure 5 モードのダイヤル λ 。驚き率や覚醒・睡眠のスケジュールで λ を動かし、発想（探索）と記憶（活用）の偏りを連続調整する。

具体例：昼は広げ、夜に綴じる

ブレインストーミングでは λ を上げて発散させ、多くの孤立スパイク (思いつき) を許す。まとめ・暗記の局面では λ を下げて収束させ、基盤へ綴じる。覚醒時の取り込みは発想寄り、睡眠時の再生・固定化は記憶寄り——同じ素材を、昼は広げ、夜に綴じる。

6. 全体構成：発想と記憶のループ

以上を一つのループにまとめる (図6)。

1. 入力（重み w 、向き s_j ）が到着する。
2. 事前構造 $\{s_i, \kappa_i\}$ に対し $B_{ij} = w_j \rho_{ij} \kappa_i$ を $O(|E|)$ で先読み算定する。
3. ゲート： $B > \theta_+$ 結合（記憶へ）／ w 大かつ $B < \theta_-$ 切断（発想へ）／ w 小 刈り込み。
4. 結合は入次数を上げ集積で基盤を深める。切断は低次数スパイクを残し、探索方向の種にする。
5. 統合：発想スパイクが後から共鳴流入を集めれば B が θ_+ を越え、発想→記憶へ昇格する。使われない基盤は μC で減衰（忘却）。
6. 更新後の $\{s, \kappa\}$ が次の事前構造になり、 λ は驚き率／スケジュールで再調整される。

ここで前稿の余白が効く。忘却 μC が古い基盤を間引いて余白を回復し、その空きが次の発想スパイクの居場所になる。すなわち発想が余白を要求し、記憶が余白を消費し、忘却が余白を返す。容量が循環するから、飽和せずに発想と記憶を往復し続けられる——これが柔軟な知能の核である。

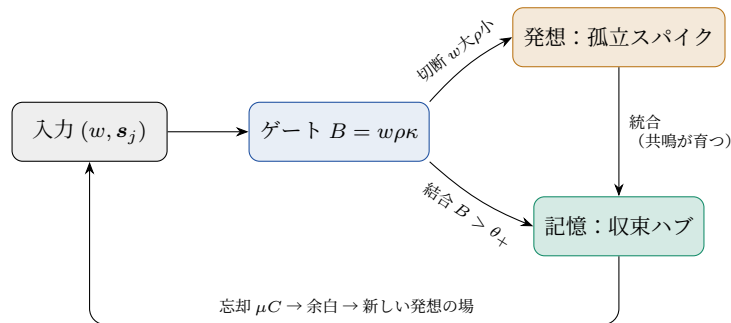


Figure 6 発想と記憶のループ。ゲート $B = w\rho\kappa$ が入力を切断（発想）か結合（記憶）に振り分け、忘却が余白を返してループが閉じる。

7. 対応する構図

この設計は、いくつかの確立した枠組みの幾何学版である。 λ による探索⇄活用の切り替えは神経修飾（青斑核–ノルアドレナリンのトニック／フェイジック制御）〔6〕に、覚醒時の取り込み／オフライン統合は睡眠時の記憶固定化に、孤立スパイク／収束ハブは発散的思考／収束的思考に、ヒステリシスは安定–可塑性ジレンマ〔7〕に対応する。NAIの言葉では、発想は「疎な読みの急所」、記憶は「密な集積のハブ」であり、結合親和度 B_{ij} がその振り分けゲートである。

8. 結言

本稿は、リッジネットの落差と次数の関係——構造的な $\delta_i = (\Delta h)_i / k_i$ 、アテンションの $k_i^{\text{eff}} \approx n e^{-\sigma_i^2/2}$ 、疎化後の $k_i^{\text{ret}} = k_i p_i(\tau)$ ——から出発し、大きな落差が孤立すれば発想、収束すれば記憶になることを示した。新規入力に対する結合・切断は、入力の重みと入力前の構造との適合の積 $B_{ij} = w_j \rho_{ij} \kappa_i$ で判定され、しきい値を λ で動かすことで発想（探索）と記憶（活用）を連続に切り替えられる。忘却が余白を回復してループが閉じることで、飽和せずに両者を往復する柔軟な知能の骨格が得られる。

構成的モデルゆえの選択を明示する。共鳴 ρ と基盤 κ の具体形、 λ コントローラ（驚き率からのフィードバック則）、しきい値とヒステリシス幅は校正・学習すべき項であり、本稿の関係式が実際のアテンションや人間の発想・記憶の挙動とどれだけ一致するかは実証を要する。今後は、入力列を流して発想スパイクと記憶ハブが切り替わる挙動の数値シミュレーション、および λ 制御の学習を進める。

参考文献

- [1] Takahashi, T. (2026). "On the Mathematical Identity between Transformers and Polygon Graphs: An Integration Structure Determined by Nodes and Edges," Journal of Takahashi Mathematics Laboratory (TML).
- [2] Takahashi, T. (2026). "NAI: Nuance AI——ポリゴングラフの幾何学的構造に基づく動的ニュアンス演算アーキテクチャの提案," TML.
- [3] Takahashi, T. (2026). "The Polygon Graph and the Conservation of Exponential Cost," TML.
- [4] Shannon, C. E. (1948). "A Mathematical Theory of Communication," Bell System Technical Journal, 27, pp. 379–423, 623–656.
- [5] Hopfield, J. J. (1982). "Neural networks and physical systems with emergent collective computational abilities," PNAS, 79(8), pp. 2554–2558.
- [6] Aston-Jones, G. & Cohen, J. D. (2005). "An integrative theory of locus coeruleus-norepinephrine function: adaptive gain and optimal performance," Annual Review of Neuroscience, 28, pp. 403–450.
- [7] Grossberg, S. (1987). "Competitive learning: From interactive activation to adaptive resonance," Cognitive Science, 11(1), pp. 23–63.

ぶつかる場所が、効く場所

——発想と記憶のAI（NAI）を、創薬の例でやさしく解説——

Toshiki Takahashi / 高橋数理研究所 独立研究者 / 2026年6月

はじめに

これまでの研究報告では、情報を「風景（ランドスケープ）」として捉え、その起伏——**落差**——が急なところに知能の急所がある、と述べてきました。本稿は、最新報告『発想と記憶の動的切り替え』を、数式を使わずにやさしく解説します。

中心に置くのは「**衝突点**」という考え方です。向きが正面からぶつかる場所に、ものが集まる——これが記憶の正体であり、同時に**創薬**（薬を見つける研究）で大きく効きます。そして大事な点として、ここでの「**反応ネットワーク**」には二つの**レベル**（概念レベルと構造レベル）があり、NAIの理論はその両方を同じように扱えます。最後に、NAIといまのLLM（GAI）を組み合わせると、なぜ互いの弱点をちょうど補い合えるのかを示します。

二つのモード：発想と記憶

情報の風景には、急に高くなったり低くなったりする「**落差の大きい場所**」があります。面白いことに、**同じ大きな落差でも二通りの意味**を持ちます。

ひとつは、まわりが平坦な中にぽつんと突き出たスパイク。つながりが少ない（孤立した）急所で、これは**前例のない新しいもの＝発想（ひらめき）**です。もうひとつは、たくさんの流れが一点に集まってくるハブ。つながりが多く流れが溜まる場所で、これが**定着した記憶**です。



Figure 1 同じ大きな落差でも、孤立すれば発想（ひらめき）、集まれば記憶（定着）。知能はこの二つを切り替えて働く。

衝突点：向きがぶつかると、ものが集まる

記憶のハブはどうやってできるのでしょうか。鍵は「**向き**」です。各点には流れの向きがあり、向きがそろっていれば流れは**素通り**しますが、二つの向きが**正面から向かい合う**と、そこで流れがぶつかって止まり、ものが集まります。この**衝突点**が収束＝記憶のハブの正体です。

いちばん大事な点は、この衝突点が「**どこにできるか**」は、**流す前に、向きの配置だけから先に計算**できることです。さらにこの考え方は、**ノード・エッジ・向きさえあれば成り立つ**ので、次に見るように**複数のレベル**に同じように使えます。

反応ネットワークの二つのレベル

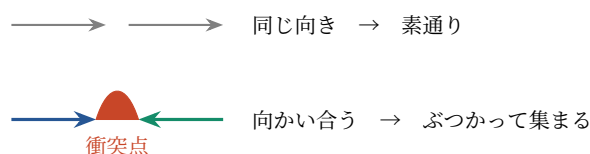


Figure 2 向きがそろえば素通り（上）、向かい合えば衝突して集まる（下）。集まる場所＝衝突点がハブになる。

ここで大事な区別があります。創薬で読む「反応ネットワーク」には、実は二つのレベルがあります。NAI（ポリゴングラフ）の理論は「ノード・エッジ・高さ・向き」だけで決まるので、どちらのレベルにも同じように使えます。これがこの理論の強みです。

概念（オントロジー）レベル。 ノードは名前——遺伝子名やタンパク質名——で、エッジは「活性化する／抑制する」という関係の矢印です。いわゆるパスウェイ図や知識グラフにあたります。ここでの衝突点は、**制御が正面からぶつかるハブ**——活性化と抑制の影響が一点に集まる「要（かなめ）」、いわゆるマスター制御因子やボトルネック——です。これは「どの分子を狙うか（標的特定）」を教えてください。

構造レベル。 ノードは原子や残基、立体構造の状態で、エッジは結合や反応経路、高さはエネルギー、向きは反応（電子）の流れです。ここでの衝突点は、**物理的に反応が集中する急所**——触媒部位や遷移状態、薬がはまる結合ポケット——です。これは「分子のどこを狙うか」を教えてください。

二つは、一つの理論の二つの解像度です。概念レベルで「どの分子か」を絞り、構造レベルで「そのどこか」を特定する。NAIは同じ衝突点の計算を両方のレベルで回し、概念から構造へとズームしていきます（図3）。

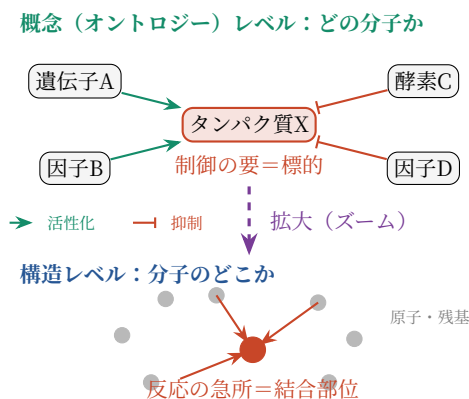


Figure 3 一つの理論を二つの解像度で。上＝概念レベル（名前と活性化・抑制の関係。制御がぶつかるハブ＝標的タンパク質）、下＝構造レベル（原子・残基。反応がぶつかる急所＝結合部位）。NAIは同じ衝突点の計算を両方で回す。

■ 概念レベルでは「どの分子か」、構造レベルでは「分子のどこか」。NAIは同じ衝突点の計算を、両方のレベルで回す。

発想の側も両レベルで効きます。概念レベルでは**思いがけない制御関係**（予想外の標的候補）が、構造レベルでは**予想外の結合様式**が、孤立スパイクとして立ちます。どちらも新しい薬の種（リード）——発想＝偶然の発見を、構造から拾えるのです。

NAIとGAIの組み合わせ：互いの弱点を補う

NAIだけでは創薬はできません。NAIは「どこが急所か（構造）」を両レベルで見つけられても、「その名前の分子が何者で、どんな関係をもち、どんな化合物が実際に効くか（知識）」は持っていないからです。

ここで、GAI（LLM）の知識の多くは概念レベルにあります。膨大な文献やデータベースから、遺伝子・タンパク質の名前と既知の関係（誰が誰を活性化・抑制するか）を知っている。一方、構造レベルの物理・化学は計算と知識を要します。

だから二段階で組み合わせます。概念レベルでは、GAIが名前と既知の関係を与え、NAIが制御ハブ（衝突点）を絞って標的分子を決める。構造レベルへズームしたら、NAIがその分子の反応の急所（結合部位）を絞り、GAIが化学知識で効く化合物を挙げて評価する（図4）。

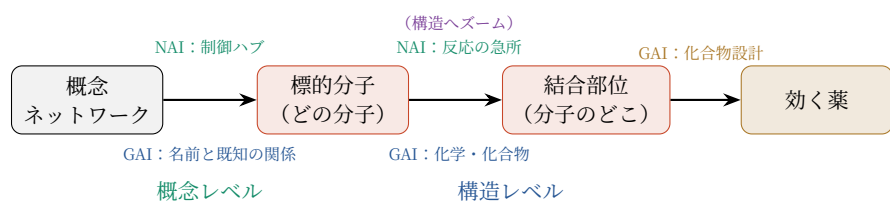


Figure 4 NAI×GAI の二段階パイプライン。概念レベルで標的分子（どれ）を、構造レベルで結合部位（どこ）を絞り、化合物を設計する。各段で NAI が「どこ」を、GAI が「何者・何を作るか」を担う。

両レベルを通じて、NAIは「どこを狙うか」（構造・急所・記憶）を、GAIは「何者で何を作るか」（名前と化学の知識）を受け持ちます。NAIの弱点（知識がない）をGAIが、GAIの弱点（力任せ・ハルシネーション）をNAIが、ちょうど補い合う。NAIだけでは中身が空っぽ、GAIだけでは的が定まらない——二つでひとつです。

■ NAIは「どこを狙うか」を両方のレベルで、GAIは「何者で何を作るか」を知識で。二つでひとつ。

前の解説の言葉でいえば、NAIが「余白（構造・骨格）」を、GAIが「充填（知識・中身）」を、概念と構造の両レベルで受け持つ組み合わせにほかなりません。

まとめ

ぶつかる場所が、効く場所です。向きが正面から向かい合う衝突点に、反応（や制御）が集まる。そしてその急所は、やみくもに探す前に構造だけから先に見つけられる。さらにこの理論は、概念（どの分子か）と構造（分子のどこか）の両レベルに同じように使えます。創薬でいえば、概念レベルで標的分子を絞り、構造レベルで結合部位を特定し、そこに効く化合物を設計する——という流れです。

そしてNAIとGAIを組み合わせると、「どこを狙うか（NAI・両レベル）」と「何者で何を作るか（GAI・知識）」が噛み合い、互いの弱点を補って、軽く・忘れず・筋の通った知能になります。発想で新しい急所を見つけ、記憶で定着させ、必要なところにだけ重い計算（GAI）を呼ぶ——それが、この一連の研究が描く知能の姿です。

（注）本稿は技術報告『発想と記憶の動的切り替え』のやさしい姉妹編であり、考え方を平易に述べたものです。概念・構造の両レベルへの創薬応用は構想であり、実際の有効性の検証は今後の課題です。