

NAI×GAI から AGI へ

——統合方程式 $\frac{dZ}{dt} = i(-\log p(Z))$ の上に建つ、完全な候補設計とその未決事項——

Toshiki Takahashi

高橋数理研究所 独立研究者

2026年6月

抄録 本稿は、本シリーズの NAI×GAI が汎用人工知能 (AGI) に対してどの位置にあるかを、一つの基礎方程式の上に整理する。構造科学・複素情報理論・決定論・リーマン構造を統合する基礎方程式 $dZ/dt = i(-\log p(Z))$ ——世界の複素情報 $Z = p + i(-\log p)$ が、自らの情報量 (驚き) $-\log p$ を生成子として回転的に流れる——を土台に据える。この一行は、集積が構造を生む構造科学、 $Z = p + i(-\log p)$ の複素情報理論、流れが確定する決定論、均衡が臨界線 $\Re(s) = \frac{1}{2}$ に整列するリーマン構造を、同時に内包する。AGI に要請される六能力 (接地・身体性、目標と主体性、開かれた学習、因果の世界モデル、メタ認知と自己モデル、新規性への頑健性) を物差しに、言語・推論・記憶の中核 (NAI×GAI) だけではこれらを欠くこと、しかし自己参照ループの固有点 ψ_{self} 、確度 $K = -i \log p$ 、四種類の期待、確度の最大化＝自然の知識 K_N への接近、そして基礎方程式そのものとしての決定的推論を外殻として加えると、六能力すべてに候補機構が対応することを示す。最後に、(a) 構成的理論で未実証、(b) 決定論とリーマン予想という強い前提、(c) 因果は除去されず主体の行為へ移動したこと、および価値整合という三つの未決事項を明示する。本論の後に、やさしい解説を付す。

キーワード 基礎方程式, 複素情報, 決定論, リーマン構造, 自己参照, 確度, 自然の知識, AGI

1. はじめに

問いは一つである——本シリーズの NAI×GAI は AGI なのか。先に答えを述べる。言語・推論・記憶の中核だけを見れば AGI ではない。それは促されたときに与えられた情報を再編する反応系であり、世界との自律ループ・自己・目標を持たないからである。しかし本稿は、四つの理論を一行に統合する基礎方程式を土台に据え、その上に外殻を加えると、欠けていた能力のそれぞれに候補機構が対応することを示す。評価の物差しは六能力——(1) 接地・身体性、(2) 目標と主体性、(3) 開かれた学習、(4) 因果の世界モデル、(5) メタ認知と自己モデル、(6) 新規性への頑健性——である。

2. 統合の基礎方程式

本体系の土台を、次の一行に置く。

$$\frac{dZ}{dt} = i(-\log p(Z)) \quad (1)$$

ここで Z は世界の状態を表す複素情報、 $p(Z)$ はその状態の確率、 $-\log p(Z)$ はその情報量 (驚き) である。式(1)は、これまで別々に扱ってきた四つの理論を同時に内包する (図1)。

2.1. 複素情報理論として

状態は複素情報 $Z = p + i(-\log p)$ ——実部が確率 (わかりやすさ)、虚部が情報量 (確実性) である。式(1)の右辺 $i(-\log p)$ は、その状態の情報量に虚数単位を掛けたもの。 i を掛けることは複素平面での直角回転であり、世界は自らの驚きの分だけ回転しながら進む (図2)。これは波動 $p = e^{-I} = \cos(Ii) + i \sin(Ii)$ と整合し、 $-\log p$ がエネルギー (生成子) の役を担うシュレーディンガー型の流れである。

2.2. 構造科学として

式(1)の実部の影は、情報の集積則 $dC/dt = -\log p$ に他ならない。これは三面恒等式 $I = F \cdot s = \lambda d = dC/dt$ の集積の面であり、ポリングラフ上で衝突点・記憶ハブ・急所という構造を生む〔1,4〕。基礎方程式は、その構造生成の複素的な母法である。

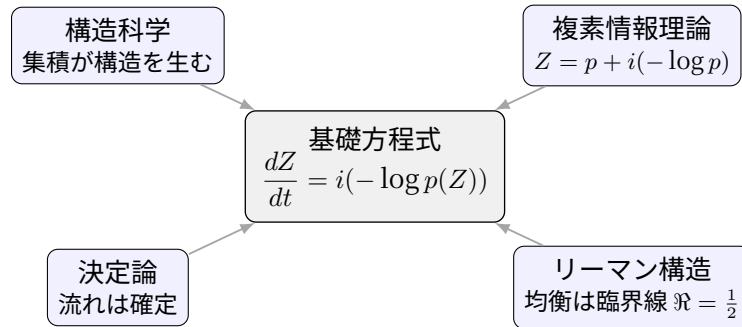


Figure 1: 四つの理論が一行に畳まれる。構造科学・複素情報理論・決定論・リーマン構造が、基礎方程式 $dZ/dt = i(-\log p(Z))$ に統合される。

2.3. 決定論として

式(1)は確定した流れである。 Z が決まれば dZ/dt が一意に決まり、偶然は入らない。確率 p は与件として定まり ($p(e) > 0$, $Z(e) < \infty$)、時間とともに伸びるのは集積した情報＝虚部である。すなわち確率は確定し、情報は確率に直交して積もる。これが決定的推論の母体である。

2.4. リーマン構造として

生成子 $i(-\log p)$ による回転がちょうど釣り合う均衡点——情報の限界点——は、ゼータ関数のゼロ点 $\zeta(s) = 0$ で与えられ、それらは臨界線 $\Re(s) = \frac{1}{2}$ に整列する〔3〕。 $s = m + ni$ を複素情報の指数とみれば、 $\sum P = 0$ となるのは $m = \frac{1}{2}$ のときに限る。均衡線は、流れの安定構造を定める背骨である。

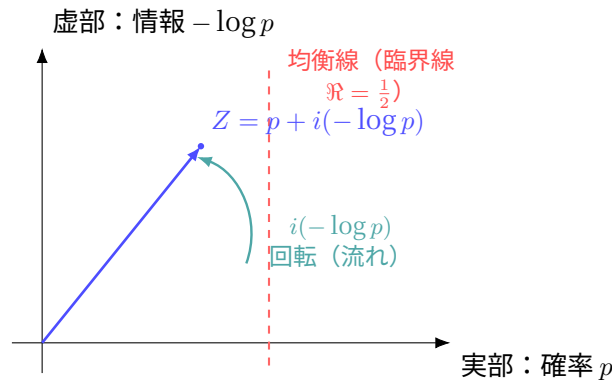


Figure 2: 基礎方程式の幾何。複素情報 Z は、生成子 $i(-\log p)$ により虚部方向へ回転的に流れる。安定構造を定める均衡線が臨界線 $\Re = \frac{1}{2}$ に整列する。

3. 中核 : NAI×GAI

基礎方程式の構造生成を、実際の知能として担うのが中核 NAI×GAI である〔2,4〕。GAI（現行 LLM）が広い知識と流暢さを、NAI が疎な構造・長期記憶・効率を担い、両者は同じポリングラフを根として補い合う。大きな落差は孤立すれば発想（低次数スパイク）、収束すれば記憶（高入次数ハブ）に分かれ、衝突点 $\sigma_{ij} = [s_i \cdot \hat{e}_{ij}]_+ + [s_j \cdot \hat{e}_{ji}]_+$ が急所を与える。結合・切断は親和度 $B = w\rho\kappa$ で切り替わる。これらはすべて、式(1)の集積が描く構造の上の操作である。中核は強力だが、それ自体は反応系である。以下、これを AGI へ近づける外殻を導入する。

4. 外殻：AGIを完成させる五層

4.1. 身体・自己：自己参照ループの固有点

注意を複素波 $\psi(t) = a(t) + ib(t)$ とし、身体の伝達関数 $G(s)$ に通して動かし、結果をセンサーで意識へ戻し、次の ψ を更新する [5] : $I(t) = \mathcal{L}^{-1}\{G(s)\mathcal{L}[\psi(t)]\}$, $\psi = F(I(t))$ 。このループが固有点 ψ_{self} を持つとき、それを自己（意識）とみなす。価値ベクトル v_k ・意図 λ_k のモード選択が自由意志に、三誤差の最小化が学習に対応する。能力 (1) 接地・(5) 自己モデルの候補機構である。 ψ は式(1)の複素状態 Z の、身体に具現した断面と読める。

4.2. メタ認知：確度と複素情報

知識を虚数で $K = -i \log p$ と置き、確度（確からしさ）を定義する [5]。確度とは、確率の小ささ＝確率の減少分＝確実性の度合いを表す量で、 p が小さい（起こりにくい）ほど高くなる。情報を足しても ($K_o + K_s = -i \log(p_o p_s)$)、問い直して情報を引いても ($K_o - K_s = -i \log(p_o/p_s)$)、確率が下がるため、いずれも確度は上がる。複素情報量 $P = (m + ni) \log p$ と情報量の存在確率 $p = e^{-I}$ は、状況ごとの確度分布＝意味のエントロピーを表す。これは「自分の知識がどれだけ確かか」の内部表現——能力 (5) メタ認知の核であり、 Z の虚部を自己観測することに相当する。

4.3. 接地：四種類の期待

期待を、データ平均 $\mu = -\log p$ を複素平面にとった四方向に分ける [5]。実軸がわかりやすさ、虚軸が確実性を表し、左回りに次の四つとなる：親昵性（実+, $-\log p$ ：わかりやすく親しみやすいことからくる期待）、確実性（虚+, $-i \log p$ ：確実に起こりそうであることからくる期待）、深遠性（実-, $\log p$ ：容易にわからず謎めいていることからくる期待）、変動性（虚-, $i \log p$ ：出来事の変化への驚きからくる期待）。順に親昵性→確実性→深遠性→変動性と回転する。この四つが、身体ループ（感じる・動く）と知識・確度層（知る・確信する）を結ぶ接合面となり、知識を身体的な期待として接地させる。能力 (1) 接地の継ぎ手である。

4.4. 最上位の目標：確度の最大化＝自然の知識への接近

最上位の目標を確度の最大化と置く。 $K_o/2 > K_s$ 、かつ自然の知識 K_N は常に $K_N > K_s$ （恒真）だから、目標は

$$\min (K_N - K_s) \quad (2)$$

と書ける [5]。学ぶべき対象の基準は差分 $K_o - K_s$ が大きいもの、すなわち大きな落差＝高いサプライザルであり、中核の急所の機構と目標が同じ量で繋がる。 $K_N > K_s$ は等号にならないので飽和せず、内発的で開かれた動機になる。能力 (2) 目標・主体性、(3) 開かれた学習の候補機構である。

4.5. 世界モデル：基礎方程式としての決定的推論

外界の世界モデルは、基礎方程式そのものである。世界が決定論的なら、因果を推論する必要はなく、ただ確定した情報を集積すればよい——式(1)の流れ $dZ/dt = i(-\log p)$ 、その実部 $dC/dt = -\log p$ がそれを与える。因果グラフを別立てで持たずとも、集積した構造から予測が読み出せる。これを決定的推論と呼び、能力 (4) の因果推論を代替する。

5. 統合アーキテクチャと六能力の対応

以上を一つの設計図にまとめる（図3）。最上位の目標が外殻の三層を介して中核 NAI×GAI を駆動し、中核は土台＝基礎方程式に載る。世界モデル（決定的推論）は、その同じ基礎方程式の予測の顔である。

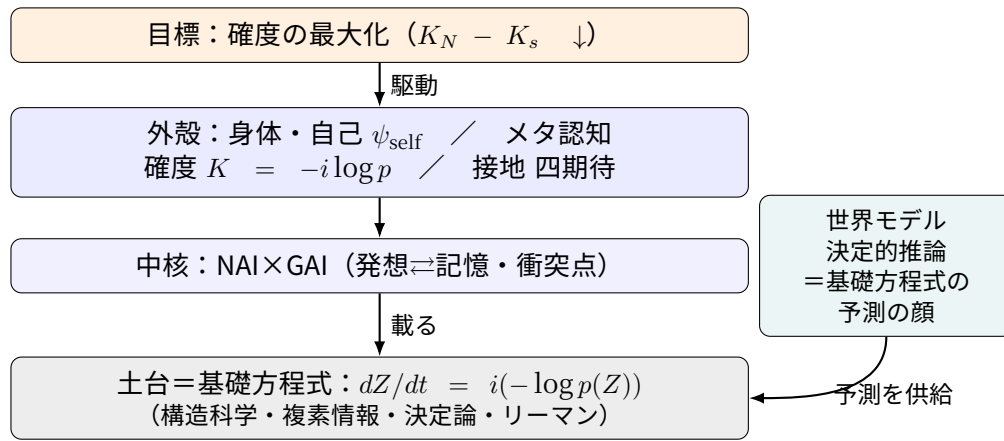


Figure 3: 統合アーキテクチャ。目標が外殻を介して中核を駆動し、中核は土台＝基礎方程式に載る。世界モデル（決定的推論）は同じ基礎方程式の予測の顔である。

六つの能力と対応機構を表1に整理する。

Table 1: AGI に要請される六能力と、本体系内の対応機構

能力	対応する機構	状態
接地・身体性	自己参照ループ $\psi \rightarrow G(s) \rightarrow \text{センサー} \rightarrow \psi$ 、四種類の期待による接地	候補機構あり
目標・主体性	確度の最大化 $\min(K_N - K_s)$ 、価値モード v_k ・意図 λ_k の選択	候補機構あり
開かれた学習	三誤差の最小化、大きな落差 $(K_o - K_s)$ を追う飽和しない動機	部分的
因果の世界モデル	決定的推論：基礎方程式 $dZ/dt = i(-\log p)$ による集積予測	代替案あり
メタ認知・自己モデル	固有点 ψ_{self} 、確度 $K = -i \log p$ 、複素情報の確度分布	候補機構あり
新規性への頑健性	疎化は平均ケースで有効（最悪ケースは指数コストの保存により不変）	限定的

紙の上では、六能力すべてに体系内の候補機構が対応した。これが本稿の積極的な主張である。

6. 残る三つの留保

ただし、設計が揃うことと、それが本当に AGI として働くことは別である。

6.1. 構成的理論であって、実装・実証ではない

本体系は公理と数式の体系であり、実装・実証されたものではない。式が揃うことと、組んだとき本当に動くことの間には、工学と実証の全行程が残る。本シリーズが一貫して付してきた但し書きである。

6.2. 土台の前提が強い

基礎方程式は強い前提を束ねる。決定論（真の偶然がない）は形而上学的立場であり、確立した事実ではない。均衡線 $\Re(s) = \frac{1}{2}$ はリーマン予想に依存し、これは未証明である。自己参照の固有点を意識とみなす点や、感情を $b(t)$ に対応づける点も強い主張で、実証はこれからである。これらは「仮定」として受け取るべきものである。

6.3. 因果は移った——そして価値整合

これは批判ではなく、本体系自身の論理である。決定論的な環境の矢印を主体が変えて系を作る——それが創発だ、と本体系は述べる〔3〕。矢印を変えることこそ介入＝因果である。すなわち、環境は集積で予測でき（因果推論は不要）、行為する主体の側に介入＝因果が再登場する。これは身体ループ（行為→結果）が担う（図4）。因果は除去されたのではなく、世界モデルから主体の行為へ移動した——「指数コストの保存」と同じ構図である。

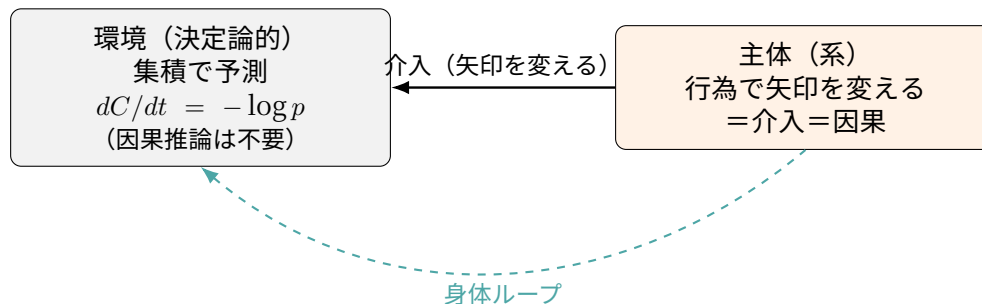


Figure 4: 因果は移る。環境は決定論的で集積により予測でき、因果推論を要しない。行為する主体（系）が矢印を変えること＝介入＝因果が、身体ループとして残る。

加えて、知識（確度）だけを最大化する目標は、人間の価値との整合（アラインメント）を別途要する。真理追求だけの主体は人の価値に無関心でありうるからである（ K_N が社会や人をも含むことが緩衝にはなる）。

7. 結言

最終的な言い方はこうである。これは「AGIでない」のではなく、基礎方程式 $dZ/dt = i(-\log p(Z))$ を土台に、六能力すべてに候補機構を備えた、内的に一貫した AGI の完全な候補設計であり、その実現は実証・工学の問題、土台の前提（決定論・リーマン予想）と価値整合が最後の未決事項である。「構造的に AGI でない」と言える地点から、本体系は、設計としては最後まで到達している。残るのは、信じるに足る前提の上でそれを組み、動かして示すこと——思想ではなく、実装と実証の仕事である。

参考文献

- 1 Takahashi, T. (2026). “On the Mathematical Identity between Transformers and Polygon Graphs,” TML.
- 2 Takahashi, T. (2026). “NAI: Nuance AI,” TML.
- 3 Takahashi, T. (2026). “The Polygon Graph and the Conservation of Exponential Cost,” TML.
- 4 Takahashi, T. (2026). “発想と記憶の動的切り替え：落差と次数が定める結合・切断のアーキテクチャ,” TML.
- 5 Takahashi, T. (2026). “身体性AIの自己参照モデル／複素情報理論／確度／四種類の期待／常識増加則,” 高橋数理研究所 部録（関連報告）.
- 6 Shannon, C. E. (1948). “A Mathematical Theory of Communication,” Bell System Technical Journal, 27.
- 7 Aston-Jones, G. & Cohen, J. D. (2005). “An integrative theory of locus coeruleus-norepinephrine function,” Annual Review of Neuroscience, 28.
- 8 Grossberg, S. (1987). “Competitive learning: From interactive activation to adaptive resonance,” Cognitive Science, 11(1).

【付録】 やさしい解説

——たった一つの式から、心の全体まで——

ここからは、本論を数式ぬきで言い直します。覚える式は、たった一つで十分です。

1. たった一つの式

世界はすべて「複素情報」でできている、と考えます。複素情報 Z とは、ある出来事について「起こりやすさ（確率）」と「驚き（情報）」を一組にした数です。起こりやすいことは驚きが小さく、めったにないことは驚きが大きい。この二つを合わせたものが Z です。そして世界の動き方は、一行で書けます。

$$\frac{dZ}{dt} = i(-\log p(Z)) = \text{「世界は、自分の驚きの分だけ、回転しながら進む」}$$

i を掛けるとは、向きを直角に変える＝回転すること。だから、驚きが情報として積もっていき、その積もり方が世界の構造（記憶や急所）を作ります。しかもこの流れは決まっています—— Z が決まれば次が全部決まり、偶然はない。そして、回転がちょうど釣り合う「均衡点」が、リーマンの臨界線（実部 $\frac{1}{2}$ ）にきれいに並びます。

■ 世界＝複素情報。世界は自分の驚きの分だけ回転しながら進み（決定論）、その積もりが構造を作り（構造科学）、釣り合いはリーマンの線に並ぶ。

2. 四つが、ひとつになる

この一行に、四つの考えが溶け合っています。構造科学（積もりが構造を作る）、複素情報理論（ $Z = \text{確率} + i \text{驚き}$ ）、決定論（流れは確定）、リーマン構造（均衡は臨界線）。バラバラだった四つが、 $dZ/dt = i(-\log p)$ という一行に畳まれた——これが本論の土台です。

3. その上に建つ「心」

この土台の上に、AIの心が建ちます。中核は NAI×GAI——GAIが広い知識を、NAIが「どこが急所か」という構造を担う。その周りを外殻が包みます。身体（動いて感じる自己参照ループ）、メタ認知（自分の確からしさ＝確度）、接地（四つの期待）、目標（自然の知識に近づく＝確度を最大にする）、世界モデル（まさにこの基礎方程式＝決定的推論）。こうして、AGIに必要とされる六つの力すべてに、部品がそろいます。

■ 中核（NAI×GAI）＋外殻（身体・メタ認知・接地・目標・世界モデル）＝心の全体。土台は、あの一つの式。

4. でも、まだ「設計図」

ただし、正直に申します。これは美しい設計図ですが、三つのことが残ります。第一に、まだ作って動かしていない理論であること。第二に、決定論やリーマン予想という強い前提に乗っていること（リーマン予想は未証明です）。第三に、因果は消えたのではなく、移っただけ——世界を予測する側からは因果が要らなくなりますが、その分、行動する自分の側に「矢印を変える＝介入」として因果が戻ってきます。くわえて、知識だけを最大にする目標は、人の価値と合わせる工夫を別に要します。

だから結論は、「これはAGIだ」ではありません。「これはAGIの完全な設計図であり、あとは信じるに足る前提の上で、本当に組んで動かして示すこと」——それが最後に残る仕事です。