

動的知識グラフの探索・統合を制御する統一ゲージの提案 —Gauge what Knowledge Graph needs.—

A Unified Gauge for Controlling Exploration and Integration in Dynamic Knowledge Graphs
— Gauge what Knowledge Graph needs. —

宮内 和義^{*1}

Kazuyoshi Miyauchi

^{*1}独立研究者

Independent Researcher

Dynamic knowledge graphs face a fundamental “When” problem: when to accept new information as structure, and when to continue exploration. While existing approaches such as GraphRAG optimize *what* to retrieve, they lack a principled criterion for *when* to integrate. We propose **geDIG**, a unified gauge that operationally aligns the Free Energy Principle (FEP) and Minimum Description Length (MDL), quantifying this trade-off as $\mathcal{F} = \Delta\text{EPC}_{\text{norm}} - \lambda(\Delta H_{\text{norm}} + \gamma\Delta\text{SP}_{\text{rel}})$. We validate geDIG on two instantiations of dynamic knowledge graphs. (1) **Spatial KG (Maze)**: An agent building a cognitive map achieves 98% success rate while rejecting 98% of candidate edges. (2) **Semantic KG (Transformer)**: Attention matrices, viewed as token-relation graphs, exhibit a consistent transition from exploration to structure across layers ($p < 10^{-80}$), and F-regularization causally improves performance (+0.33pp, $p = 0.038$). These results suggest geDIG as a general-purpose gauge for controlling when to explore and when to integrate in dynamic knowledge graphs.

1. はじめに

動的知識グラフ (Dynamic Knowledge Graph) において、新しい情報を「いつ」構造として受け入れ、「いつ」探索を続けるべきか—この **When 問題** は未解決の基本課題である。GraphRAG[10] や Self-RAG[11] は「何を (What)」検索するかを最適化した、が、「いつ (When)」統合するかはヒューリスティックに依存している。本研究では、自由エネルギー原理 (FEP) [1] と最小記述長 (MDL) [2] を操作的に橋渡しする統一評価関数 **geDIG** (graph edit Distance and Information Gain) を提案し、この問題に原理的な解を与える。

geDIG の核心的仮説は、「構造コストと伸展利得のトレードオフ」を単一スカラーで評価すれば、知識グラフの種類によらず、探索と統合を統一的に制御できる」というものである。本稿では、この仮説を 2 種類の動的知識グラフで検証する：

- **空間的 KG (迷路)**：エージェントが構築する認知地図
- **意味的 KG (Transformer)**：Attention 行列が表現するトークン間関係グラフ

本稿の貢献は以下の 3 点である：

1. 構造コストと伸展利得のトレードオフを単一スカラー化した geDIG の定式化
2. 2 種類の KG (迷路・Transformer) における構造相転移の実証
3. F 正則化による因果的介入実験

2. geDIG：統一評価関数の定義

geDIG は、「構造を変えるコスト」と「それによって得られる情報の質」の収支を単一スカラー $\mathcal{F}$ で評価する。

$$\mathcal{F} = \Delta\text{EPC}_{\text{norm}} - \lambda(\Delta H_{\text{norm}} + \gamma \cdot \Delta\text{SP}_{\text{rel}}) \quad (1)$$

連絡先: miyauchikazuyoshi@gmail.com

ここで各項は以下に対応する (表 1)。

- $\Delta\text{EPC}_{\text{norm}}$ ：正規化 EPC (Edit Path Cost)。グラフ編集距離 (GED) の局所近似であり、MDL におけるモデル記述長 $L(M)$ の増分に相当する。 Δ は比較対象間の差分 (迷路：更新前後 G_{t-1}, G_t , Transformer：ランダムベースラインとの差) を表す。
- ΔH_{norm} ：シャノンエントロピーの差分 (エントロピー増分；伸展利得, Extension Gain)。候補集合に対する重み付き分布 p のエントロピー $H(p)$ の変化として定義し、本稿では $\Delta H := H_{\text{after}} - H_{\text{before}}$ (after-before) を採用する。 $\Delta H > 0$ は候補分布の均質化 (optionality の増大) を表し、探索圧として $\mathcal{F}$ を下げる (温度項)。
- $\Delta\text{SP}_{\text{rel}}$ ：平均最短路長の相対短縮率。グラフ上のショートカット形成による効率化を表し、MDL におけるデータ記述長 $L(D|M)$ の圧縮に相当する。

$\Delta\text{EPC}_{\text{norm}}$ ($\Delta\text{EPC}_{\text{norm}}$) の具体形再現性のため、本稿の実装で用いた局所 GED 近似を明記する。更新前後のグラフ $G_{\text{before}} = (V_b, E_b)$, $G_{\text{after}} = (V_a, E_a)$ に対し、編集操作はノード/エッジの挿入・削除のみとし、

$$\text{GED}(G_b, G_a) := c_V \cdot ||V_a| - |V_b|| + c_E \cdot (|E_b \setminus E_a| + |E_a \setminus E_b|) \quad (2)$$

で近似する。これを正規化定数 C_{max} で割り、

$$\Delta\text{EPC}_{\text{norm}} \equiv \Delta\text{EPC}_{\text{norm}} := \text{GED}(G_b, G_a) / C_{\text{max}} \quad (3)$$

とする。迷路 (検証 I) では、クエリ近傍の k -hop 局所部分グラフ上で評価し、 $C_{\text{max}} = c_V + c_E \max(1, |S_{\text{link}}|)$ (候補台固定) を用いた。Attention (検証 II) では $|V|$ が系列長 L で固定されるため、 $c_V = 0$ かつ $C_{\text{max}} = L^2$ とすると $\Delta\text{EPC}_{\text{norm}}$ はエッジ密度 $|E|/L^2$ に一致する。

$\Delta\text{SP}_{\text{rel}}$ の計算と計算量 $\Delta\text{SP}_{\text{rel}}$ は「近道」の寄与を測るため、平均最短路長 $L(\cdot)$ の相対短縮として

$$\Delta\text{SP}_{\text{rel}} := \max\left(0, \frac{L(G_{\text{before}}) - L(G_{\text{after}})}{L(G_{\text{before}})}\right) \quad (4)$$

表 1: geDIG 構成項の理論的対応

項	意味	FEP	MDL	役割
ΔEPC_{norm}	構造変更	—	$L(M)$ 増	抑制
ΔH	伸展利得 (optionality)	複雑さ 探索圧	—	拡散
ΔSP	経路短縮	—	$L(D M)$ 圧縮	効率

で定義する ($L(\cdot)$ は最大連結成分上、または固定サンプルのノードペアに対する平均)。迷路の multi-hop 評価 (DG) では、全点対最短路 (APSP) を毎回計算せず、固定ペア (既定 400 組) に対する距離をキャッシュし、追加エッジにより改善し得る距離のみを SSSP で増分更新することで計算量を抑えた。Transformer 解析では、層ごとのトークン関係グラフに対して同様の指標を算出し、スケール依存性を避けるためランダムベースラインからの差 $\Delta \mathcal{F}$ として解釈する。

候補集合に付与された重み w_i (類似度) から確率分布 p を

$$p_i = \frac{w_i}{\sum_j w_j}, \quad H(p) = - \sum_i p_i \log p_i \quad (5)$$

で定義し、正規化は $K = |S_{\text{after}}|$ として $\Delta H_{\text{norm}} = (H_{\text{after}} - H_{\text{before}}) / \log K$ とする。

なお、 $\mathcal{F}$ の絶対値は正規化・グラフ化手続きに依存するため、異なるタスク間で直接比較しない。迷路では $\mathcal{F}$ を最小化すべき目的関数 (低いほど良い) として扱う一方、Transformer 解析では同条件のランダム行列 (Null 仮説) からの差 $\Delta \mathcal{F} := \mathcal{F} - \mathcal{F}_{\text{random}}$ を主要な解釈対象とし、 $\mathcal{F}$ がランダムベースラインより 0 に近い (高い) ほど「ランダムからの構造化」が進んでいると読む。

パラメータ λ, γ は「情報温度」として機能し、構造の複雑さと情報の圧縮率のバランスを決定する。

注記: 表 1 の FEP/MDL 対応は、本稿で扱うグラフ操作に対して「同型のトレードオフ構造」を与える操作的対応であり、厳密な上界・下界を与える形式的導出は今後の課題である。

2.1 AG/DG 二段ゲート制御

この $\mathcal{F}$ を用いて、エージェントは以下の 2 つのゲートを事象駆動的 (Event-Driven) に制御する。

1. **AG (Attention Gate):** 0-hop (直近) での評価。 $g_0 = \mathcal{F}|_{\text{local}} > \theta_{\text{AG}}$ のとき、局所的な「違和感・未知」を検知し、探索モード (Exploration) を起動する。
2. **DG (Decision Gate):** Multi-hop (数理推論) での評価。 $g_{\min} = \min_h \mathcal{F}^{(h)} < \theta_{\text{DG}}$ のとき、大域的な「近道・洞察」を発見したとみなし、統合モード (Integration) としてエッジを確定する。

ここで h は DG におけるルックahead深さ (hop) であり、実装では候補エッジを貪欲に追加しながら $h = 0, 1, \dots, H$ (最大 $H = \text{max_hops}$) の各段階で $g(h)$ を計測し、その最小値を $g_{\min}$ として用いる。

この機構により、システムは常に探索し続けるのではなく、「わからないときだけ調べ (AG)」、「わかったときだけ覚える (DG)」という自律的な振る舞いを獲得する (図 1)。

3. 検証 I: 空間的 KG としての迷路探索

KG としての定式化 部分観測迷路におけるエージェントの認知地図は、動的知識グラフの一形態である。ノードは訪問した位置、エッジは移動可能な接続を表し、探索に伴いグラフが動

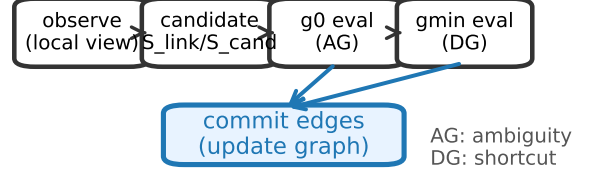


図 1: AG/DG 制御フロー。観測 → 候補生成 → AG 評価 ($g_0 > \theta_{\text{AG}}$ で探索継続) → DG 評価 ($g_{\min} < \theta_{\text{DG}}$ でコミット) の事象駆動ループ。AG は局所的異常 (高 $\mathcal{F}$) を、DG は大域的効率化 (低 $\mathcal{F}$) を検知する。

的に成長する。エージェントは周囲 1 マスしか見えないため、「見えた通路を全て記憶する」か「重要な分岐点のみ記憶する」かの選択を迫られる。

仮説 geDIG の AG/DG ゲートにより、エージェントは (1) 高い成功率を維持しつつ、(2) 候補エッジの大半を棄却する効率的な記憶構築が可能である。すなわち、「いつ探索し、いつ構造化するか」を geDIG が自律制御できる。

実験設定 15×15 および 25×25 の迷路を用いた。具体的には、各候補行動 (次セル) を位置・相対移動・通行可否・訪問回数・ゴール等からなる 8 次元特徴ベクトルで表し、クエリベクトルとの距離に基づく重み付き分布から $\mathcal{F}$ を評価する。パラメータは $\lambda = 1.0, \gamma = 1.0, \theta_{\text{AG}} = 0.4, \theta_{\text{DG}} = 0.0$ 、最大ステップ数は 15×15 で 250、 25×25 で 500 とした。比較対象として、(i) Random Walk (各ステップで可能行動から一様ランダムに選択し、直前の逆向きは回避)、(ii) Greedy DFS (訪問済みセルを全保持し、未訪問近傍を優先するオンライン DFS)、(iii) geDIG エージェントを設定した。また、記憶効率の指標として圧縮率を定義する：各ステップで生成される候補エッジ総数 $|E_{\text{cand}}|$ (観測近傍 + 想起候補) を全て記憶する素朴方式に対し、DG でコミットしたエッジ総数 $|E_{\text{commit}}|$ の棄却率 $1 - |E_{\text{commit}}| / |E_{\text{cand}}|$ を圧縮率とした (geDIG のみ定義)。

Algorithm 1 geDIG-driven Navigation

```

1: while not at goal do
2:    $\mathbf{o}_t \leftarrow \text{observe}(\text{neighbors})$ 
3:    $g_0 \leftarrow \mathcal{F}(\mathcal{G}_t, \mathbf{o}_t)$  // 0-hop evaluation
4:   if  $g_0 > \theta_{\text{AG}}$  then // 違和感検知
5:     Expand search candidates
6:   for each candidate  $c$  do
7:      $g_{\min} \leftarrow \min_h \mathcal{F}^{(h)}(c)$  // Multi-hop
8:     if  $g_{\min} < \theta_{\text{DG}}$  then // 近道発見
9:       Update  $\mathcal{G}_{t+1}$  with edge to  $c$ 
10:    Move to best direction

```

図 2: geDIG による自律探索アルゴリズム。AG が探索を、DG が構造化 (エッジ追加) を制御する。

3.1 結果: 探索と統合の自律制御

geDIG エージェントは、AG/DG のダイナミクスのみで迷路探索に成功した (表 2)。

15×15 では、geDIG は Random Walk を成功率・効率の両面で上回り、Greedy DFS と同等の成功率を維持しながら、候補エッジの大半を棄却する高い圧縮率 (98%) を達成した (図 3)。 25×25 では探索空間が拡大し成功率は低下するが、それ

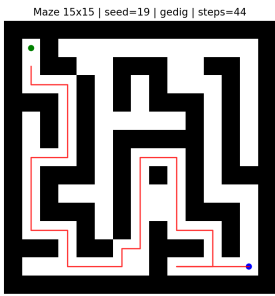
表 2: 迷路ナビゲーション結果 (DFS 迷路, 部分観測)

手法	成功率	平均ステップ	圧縮率 [†]
15×15 ($n = 100$, max_steps=250)			
Random Walk	0.75 ± 0.08	105.8 ± 63.9	—
Greedy DFS	0.99 ± 0.03	82.1 ± 30.8	—
geDIG	0.98 ± 0.03	89.6 ± 44.8	98%
25×25 ($n = 60$, max_steps=500)			
Random Walk	0.40 ± 0.12	242.2 ± 143.7	—
Greedy DFS	0.95 ± 0.06	242.2 ± 106.9	—
geDIG	0.75 ± 0.11	254.1 ± 125.4	99%

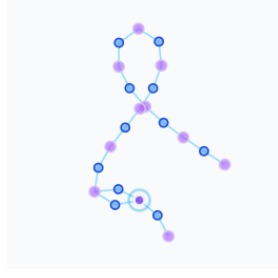
成功率は Wilson 95% CI の半幅 (±), 平均ステップは成功エピソードのみの平均 ± 標準偏差

[†] 圧縮率は geDIG の定義: 各ステップで生成される候補エッジ総数 $|E_{\text{cand}}|$ を全て記憶する素朴方式と比べ, コミットしたエッジ $|E_{\text{commit}}|$ の棄却率として $1 - |E_{\text{commit}}|/|E_{\text{cand}}|$

geDIG vs Random Walk: 同一 seed での成功/失敗をペアにした McNemar 検定 (15×15: $p < 10^{-5}$, 25×25: $p < 0.001$)



(a) 実空間の探索軌跡



(b) 抽象化された内部表現

図 3: 迷路探索における実空間と内部表現の対応. (a) 実際の迷路空間における探索軌跡. (b) geDIG が構築するエピソード記憶グラフの概念図. 直線通路は圧縮され, 分岐点 (意思決定点, 青) と経路 (エッジ) のみが保持される. 候補エッジに対するコミットエッジの棄却率 (edge-compression) は 98% に達した.

でも Random Walk より有意に高い成功率を示した. これは, geDIG が「必要なトポロジー骨格」のみを選択的に記憶することを示す. この構築されるグラフは, O’Keefe らが提唱した認知地図 [5] や, 認知科学におけるエピソード記憶の計算論的定義に対応する. モデルでは, AG が「驚き (未知)」を検知した地点をノードとし, DG が「理解 (因果)」を確定した経路をエッジとして保存する.

4. 検証 II: 意味的 KG としての Transformer

KG としての定式化 Transformer の Attention 行列 A_{ij} は, トークン i からトークン j への「関係の強さ」を表す. これを閾値処理 (上位 10%) してグラフ化すると, 動的に変化するトークン間関係グラフ——すなわち意味的 KG——が得られる. 各層でこのグラフ構造が変化することは, 情報の「探索」から「統合」への遷移を表している可能性がある.

仮説学習済み Transformer は, geDIG 的な構造化を自発的に行っている. 具体的には, (1) 浅層では高エントロピーな探索状態 (多様なトークン関係), (2) 深層では低エントロピーな構造状態 (特定トークンへの集中) への相転移が観測されるはずである. さらに, $\mathcal{F}$ を直接最適化すれば性能が変化する (因果的証拠).

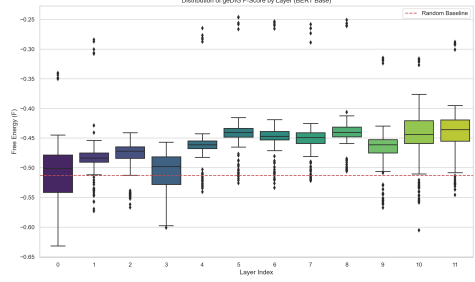


図 4: BERT における層別 $\mathcal{F}$ 値の推移 (SST-2, $n = 200$ 文, ヘッド平均). 横軸は層番号, 縦軸は $\mathcal{F}$ 値 (生の $\mathcal{F}$). 破線はランダム行列のベースラインであり, $\mathcal{F}$ がベースラインより 0 に近い (高い) ほど「ランダムからの構造化」が進む. 浅層 (Layer 0-1) では $\mathcal{F}$ 値が低く高エントロピーな探索状態にあり, 深層 (Layer 10-11) では $\mathcal{F}$ 値が上昇し安定した構造化状態に収束する (Layer 0-1 vs 10-11: $p < 10^{-80}$).

実験設定対象モデルは BERT[7] (12 層) である. SST-2 (train) から短文 200 件を抽出し, 各入力シーケンスを 1 サンプルとした. 各層の $\mathcal{F}$ はヘッド平均の Attention 行列から算出した ($\lambda = 0.5, \gamma = 0.5$, 系列長は最大 128 token, 閾値は上位 10%). なお, 本解析では $\mathcal{F}$ の生値ではなく, 完全ランダム (Null 仮説) からの差 $\Delta\mathcal{F}$ (ベースラインからの乖離) を主要な解釈対象とする. これは閾値処理や系列長により $\mathcal{F}$ の絶対スケールが変化し得るためである. 閾値は 5%~20% の範囲で変化させても同様の傾向が確認された.

4.1 観測 1: 層別の構造相転移

本稿でいう「相転移」は, 臨界点や秩序変数の存在を主張するものではなく, 層方向に沿った探索相 → 構造相の **レジム変化** を便宜上そう呼ぶ. 図 4 に BERT の層別 $\mathcal{F}$ 値分布を示す. 学習済みモデルの Attention 重みは, ランダム行列 (同密度) に比べて $\mathcal{F}$ が有意に高い値 (0 に近い) を示した (paired $t(199) = 24.4, p < 10^{-60}, n = 200$).

さらに重要な発見は, 深さ方向の推移である:

- **Layer 0-1:** $\mathcal{F}$ 値が低く (≈ -0.50), エントロピーが高い「探索相」.
- **Layer 4-5:** $\mathcal{F}$ 値が上昇し (≈ -0.45), 構造化が進む.
- **Layer 10-11:** 高い値で安定する「構造相」への相転移 (≈ -0.44).

Layer 0-1 と Layer 10-11 の差は統計的に有意である (paired $t(199) = 32.6, p < 10^{-80}$, Cohen’s $d = 2.31$). この結果は, Transformer が層を重ねるごとに情報を「探索」から「結晶化」へと移行させていることを示唆する.

4.2 観測 2: ヘッドの機能分化

同一層内でもヘッドごとに $\mathcal{F}$ 値のばらつきが見られた. 一部のヘッドは極めて低い $\mathcal{F}$ 値 (拡散的で高エントロピー) を示し, 他方は高い $\mathcal{F}$ 値 (集中した注意=低エントロピー) を示した. これは Attention Head が「探索担当」と「統合担当」のマルチエージェントとして機能分化していることを示唆する.

4.3 観測 3: $\mathcal{F}$ 正則化による因果的介入

以上の観測は相関関係に過ぎない. $\mathcal{F}$ が真に性能に関与しているなら, この値を直接最適化することで性能が変化する

表 3: F 正則化の効果 (SST-2, DistilBERT, $n = 3$ seeds)

α	Accuracy (% , mean $\pm$ std)	Δ (pp)
0 (Baseline)	86.00 $\pm$ 0.53	–
0.001	86.33* $\pm$ 0.46	+0.33
0.01	86.27 $\pm$ 0.70	+0.27
0.1	85.93 $\pm$ 1.10	-0.07
1.0	78.93 $\pm$ 1.55	-7.07

* paired t -test: $p = 0.038$ vs Baseline (seed 対応, $n = 3$)
+0.33pp は誤り率を約 2.4% 相対減 (14.0% $\rightarrow$ 13.7%)

るはずである。DistilBERT を用いた SST-2 タスクの Fine-tuning において、損失関数に F 項による正則化を追加した：
 $L_{\text{total}} = L_{\text{CE}} + \alpha \cdot \mathcal{F}$ 。

実装詳細 (再現性) モデルは `distilbert-base-uncased`, SST-2 (train 1000 / validation 500), 最大長 128, epoch=3, batch=16, AdamW (lr= 2×10^{-5} , weight_decay=0.01) で学習した。F は **post-softmax の Attention 重み** から各層・各ヘッドで算出し、層平均を損失に加えた。エッジ抽出の閾値処理は非微分な top- k の代わりに、各ヘッドの 90 パーセンタイルを中心とした sigmoid による **soft-threshold** (温度 0.1) で近似し、密度 (EPC) とエントロピーは微分可能に計算した。 $\Delta \text{SP}_{\text{rel}}$ は soft adjacency の有限長 (≤ 4) パワーに基づく到達度 proxy で近似した (閾値の勾配は安定化のため近似上固定)。

実験の結果 (表 3), 微弱な正則化 ($\alpha = 0.001$) でベースラインから有意に改善した (paired t -test, $p = 0.038$, $n = 3$ seeds)。加えて, 3/3 seeds で一貫した改善 (+0.2~0.4pp) が観測された。一方で, 過剰な正則化 ($\alpha \geq 0.1$) は性能を悪化させた (表 3)。この逆 U 字型の特性は, 適切な「構造化圧」がモデルの汎化性能を引き出すことを因果的に示している。

注目すべきは, 微弱な正則化 ($\alpha = 0.001$) でも有意な改善が見られた点である。これは, 学習済み Transformer が既に F 最小化に近い状態で動作していることを示唆する。すなわち, Transformer は学習過程で自発的に構造化圧を獲得しており, geDIG はその「隠れた最適化目標」を顕在化したものと解釈できる。

5. 関連研究

本研究の理論的基盤は, Friston の自由エネルギー原理 (FEP) [1] と Grünwald の最小記述長 (MDL) [2] である。両者は「複雑さと適合度のトレードオフ」という点で数学的に対応し [9], geDIG はこの対応をグラフ編集距離 [8] で操作可能にした。また, GNN[12] は確率的埋め込みでグラフを処理するが, geDIG は構造コストと情報利得を明示的に分離・統合する点で異なる。認知地図の概念 [5] は空間的 KG としての迷路実験の理論的背景を与える。

6. 考察とおわりに

本研究では, 2 種類の動的 KG——空間的 KG (迷路) と意味的 KG (Attention) ——において, geDIG 評価関数 $\mathcal{F}$ が共通の「構造化原理」として機能することを示した。両者に共通するのは, 高エントロピーな探索相から低エントロピーな構造相への相転移である: 迷路では時間軸に沿って未知領域の拡張から経路の確定へ, Transformer では層方向に沿って Layer 0-1 の多様なヘッド活性化から Layer 10-11 の特定ヘッドへの集中へと遷移する。F 正則化実験により, この $\mathcal{F}$ とモデル性

能の因果関係を実証した。

geDIG と自由エネルギー geDIG の式 $\mathcal{F} = \Delta \text{EPC}_{\text{norm}} - \lambda(\Delta H_{\text{norm}} + \gamma \cdot \Delta \text{SP}_{\text{rel}})$ は, 熱力学の自由エネルギー $F = E - TS$ と同型のトレードオフ構造を持つ: 構造変更コスト $\Delta \text{EPC}_{\text{norm}}$ がエネルギー E に, 伸展利得 $\lambda \Delta H$ が温度 $\times$ エントロピー TS に対応する。この対応に基づけば, Transformer の層別 $\mathcal{F}$ 推移は「高温 (探索) から低温 (構造化) への冷却過程」として解釈できる。実際, Attention 機構が入力情報の不確実性を減じ, 意味的な近道を発見・固定化するダイナミクスは, 熱力学的なアニーリング過程と整合する。

今後の展望 geDIG の応用として, GraphRAG[10] や Self-RAG[11] などの構造化 RAG における「いつ検索し, いつ構造化するか」の自律制御が挙げられる。現行の RAG は検索タイミングをヒューリスティックに決定しているが, geDIG の AG/DG ゲートを導入することで, 知識統合の意思決定を原理的に制御できる可能性がある。

参考文献

- [1] Friston, K.: The free-energy principle: a unified brain theory?, Nat. Rev. Neurosci., 11, 127–138 (2010).
- [2] Grünwald, P. D.: The Minimum Description Length Principle, MIT Press (2007).
- [3] Vaswani, A., et al.: Attention Is All You Need, NeurIPS (2017).
- [4] Gentner, D.: Structure-mapping: A theoretical framework for analogy, Cognitive Science, 7(2), 155–170 (1983).
- [5] O’Keefe, J., Nadel, L.: The Hippocampus as a Cognitive Map, Oxford University Press (1978).
- [6] Bronstein, M. M., et al.: Geometric deep learning: going beyond euclidean data, IEEE Signal Processing Magazine, 34(4), 18–42 (2017).
- [7] Devlin, J., et al.: BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding, NAACL (2019).
- [8] Bunke, H.: On a relation between graph edit distance and maximum common subgraph, Pattern Recognition Letters, 18(8), 689–694 (1997).
- [9] Tishby, N., et al.: The information bottleneck method, arXiv:physics/0004057 (2000).
- [10] Edge, D., et al.: From Local to Global: A Graph RAG Approach to Query-Focused Summarization, arXiv:2404.16130 (2024).
- [11] Asai, A., et al.: Self-RAG: Learning to Retrieve, Generate, and Critique through Self-Reflection, arXiv:2310.11511 (2023).
- [12] Kipf, T. N., Welling, M.: Semi-Supervised Classification with Graph Convolutional Networks, ICLR (2017).